



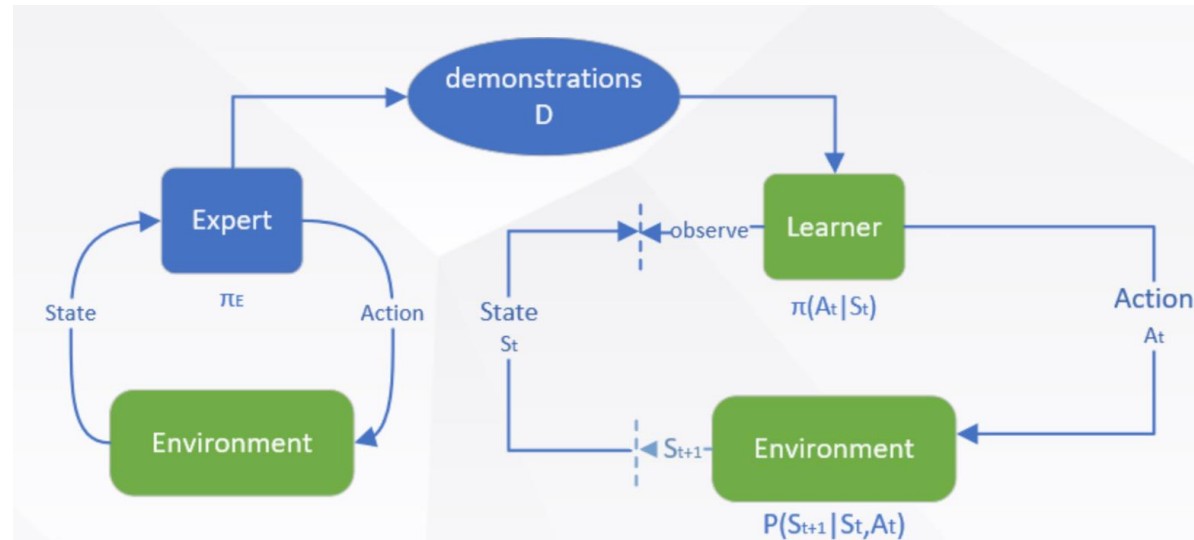
Adversarial Imitation Learning with Trajectorial Augmentation and Correction

Dafni Antotsiou, Carlo Ciliberto and Tae-Kyun Kim

ICRA 2021

Motivation

□ Imitation Learning



Deep Imitation Learning requires **a large number of expert demonstrations**, which are not always easy to obtain, especially for complex tasks. However, **data augmentation** cannot be easily applied to control tasks due to the sequential nature of the problem.

Method

□ Imitation using trajectorial data augmentation

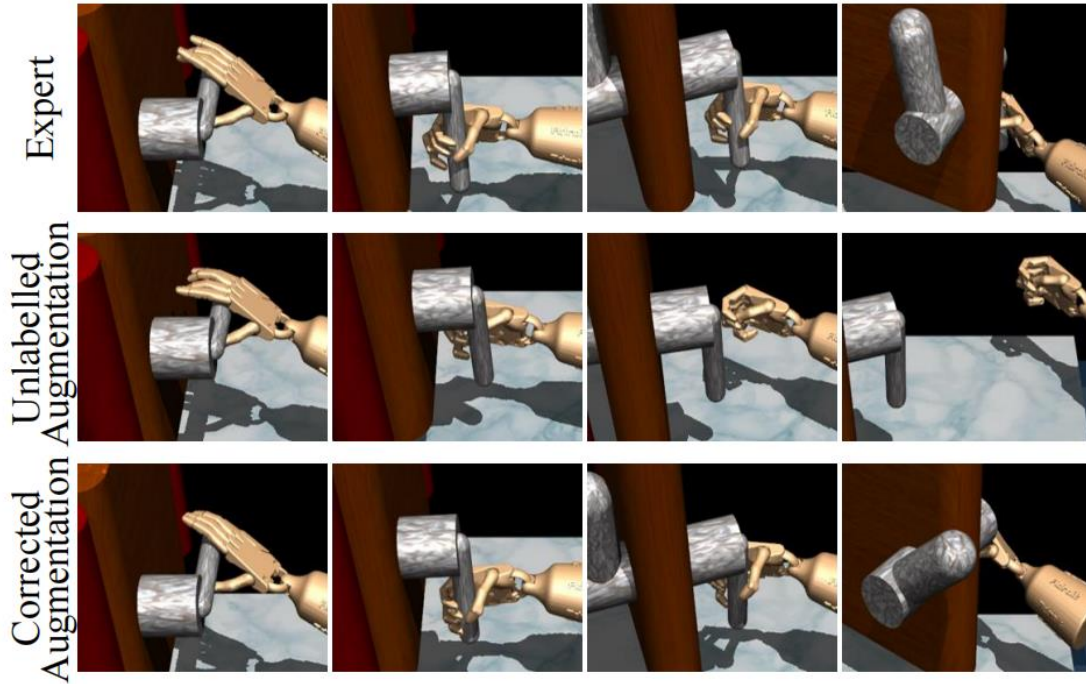


Fig. 1. Example of our proposed method for trajectorial data augmentation. **Top:** the original expert trajectory. **Middle:** the expert trajectory distorted by noise. The distortion makes the trajectory unsuccessful. This result is not guaranteed, therefore this augmentation is unlabelled. **Bottom:** our correction network modifies the unlabelled augmentation to produce a successful trajectory, different from the expert.

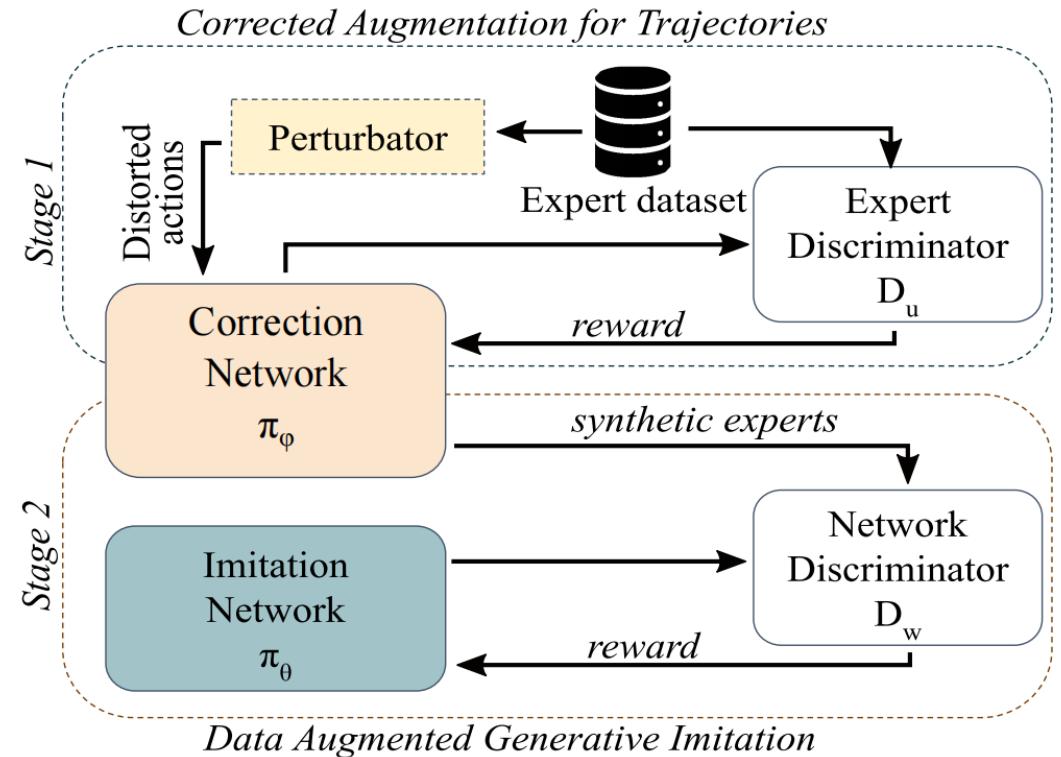


Fig. 2. Flow chart of our system, which performs imitation using trajectorial data augmentation. Stage 1 presents data augmentation by correcting distorted trajectories, while stage 2 presents data augmented imitation.

Method

$$q = \{a'_1, a'_2, a'_3, \dots\}, \quad \text{where } a'_t = a_{E_t} + \nu$$

and $\tau_E = \{(s_{E_1}, a_{E_1}), (s_{E_2}, a_{E_2}), \dots\}$.

(3)

$$L_u = -\mathbb{E}_{\pi_E} [\log D_u(s, a)] - \mathbb{E}_{\pi_\phi} [\log(1 - D_u(s, a))], \quad (1)$$

$$L_\phi = \mathbb{E}_{\pi_\phi} [\log(1 - D_u(s, a))] + \lambda \|a - a'\|_2^2. \quad (5)$$

□ Stage1 : Corrected Augmentation for Trajectories (CAT)

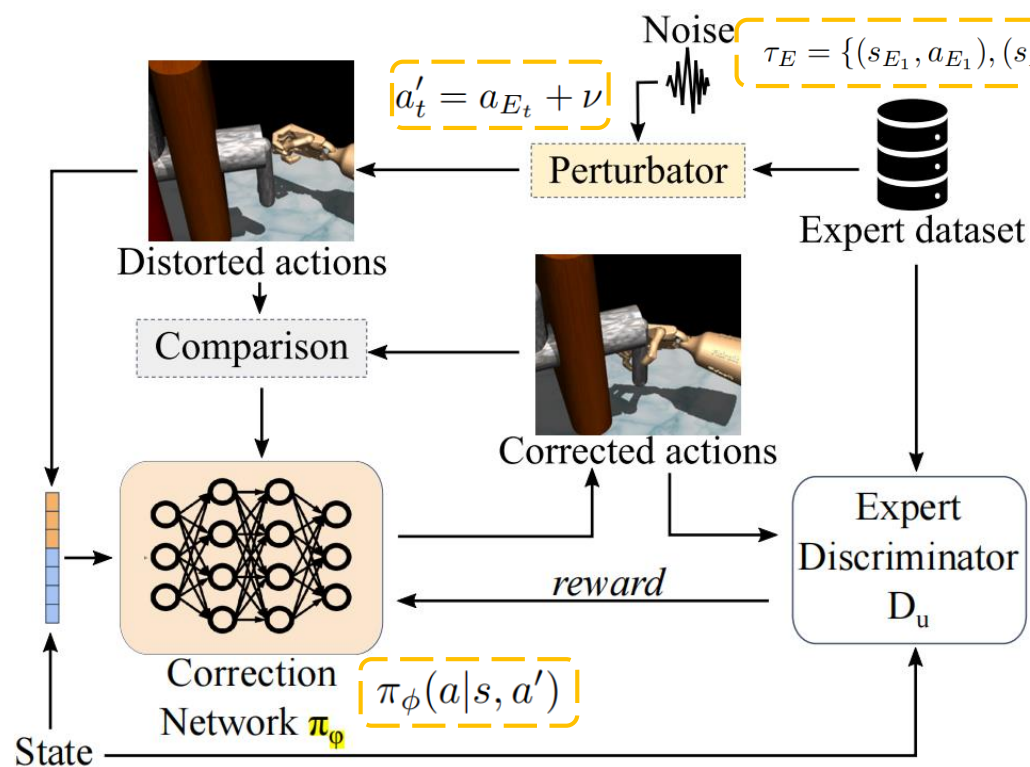


Fig. 3. Detailed overview of stage 1, which performs Corrected Augmentation For Trajectories. The architecture is semi-supervised since it is guided by unlabelled distorted actions.

Algorithm 1: Corrected Augmentation for Trajectories Algorithm

Input: Set of expert trajectories $\mathcal{T}_E$, regularisation λ noise σ , initial policy ϕ_0 and discriminator u_0 , N number of perturbed action sequences.

```

// produce randomly perturbed augmented
// action sequences  $\mathcal{Q}$ 
1  $\mathcal{Q} = \{\}$ 
2 for each  $\tau_E$  in  $\mathcal{T}_E$  do
3   Generate  $N$  perturbed action sequences
4    $\mathcal{Q}' = \{q_1, \dots, q_N\}$  from  $\tau_E$  according to (3).
5 end
// correct augmented trajectories so they are
// successful
6 for  $i = 0, 1, 2, \dots$  do
7   Sample trajectory  $\tau_i \sim \pi_{\phi_i}(q_i)$ , with  $q_i$  sampled
   uniformly from  $\mathcal{Q}$ .
8    $u_{i+1} \leftarrow u_i - \nabla_u L_{u_i}$  using GAIL's (1).
9    $\phi_{i+1} \leftarrow \phi_i - \nabla_\phi L_{\phi_i}$  using (5).
10 end

```

Method

□ Stage2 : Data Augmented Generative Imitation (DAugGI)

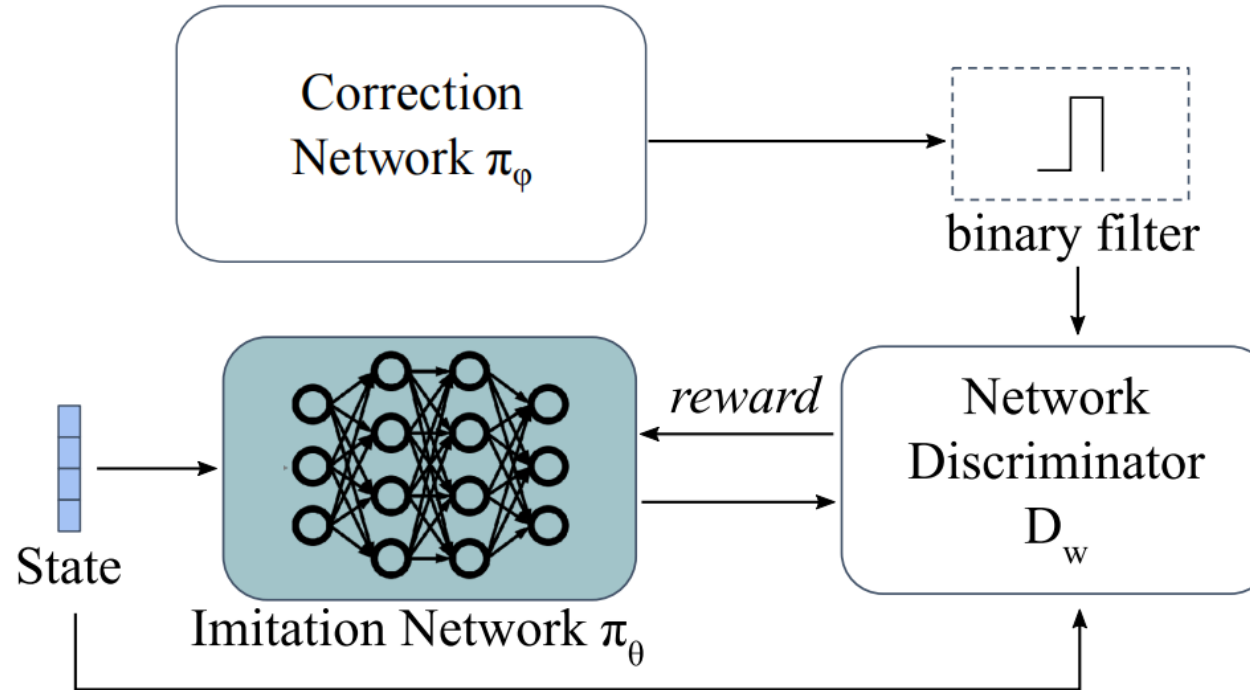


Fig. 4. Detailed overview of stage 2, which performs Data Augmented Generative Imitation, including the success filtering mechanism.

$$L_\theta = \mathbb{E}_{\pi_\theta} [\log(1 - D_w(s, a))]$$

$$L_w = -\mathbb{E}_{\pi_\phi} [\log D_w(s, a)] - \mathbb{E}_{\pi_\theta} [\log(1 - D_w(s, a))]$$

Experiments : CAT Evaluation

$$\overline{dtw}_n(\mathcal{T}_g) = \frac{\overline{dtw}(\mathcal{T}_g)}{\overline{dtw}(\mathcal{T}_E)},$$

$$\text{where } \overline{dtw}(\mathcal{T}) = \frac{\sum_{i=1}^{N-1} \sum_{j=i+1}^N dtw(\tau_{z_i}, \tau_{z_j})}{(N-1)N/2}.$$

TABLE I
CAT EVALUATION AND DATASET DIVERSITY

Task	CAT Success %		Dataset Diversity score $\overline{dtw}_n$		
	Random Aug/tion	Corrected Aug/tion	$\frac{\text{CAT}}{\text{original}}$	$\frac{\text{DAugGI}}{\text{original}}$	$\frac{\text{GAIL}}{\text{original}}$
HalfCheetah	0.7	97.6	0.54	0.68	0.73
Inv. Pendulum	4	100	0.91	1.13	1.11
Door	21	56	1.11	0.26	0.25
Pen	58	46	0.93	1.12	1.06
Hammer	62	63	1.00	0.26	0.24

Experiments : DAUGGI Evaluation

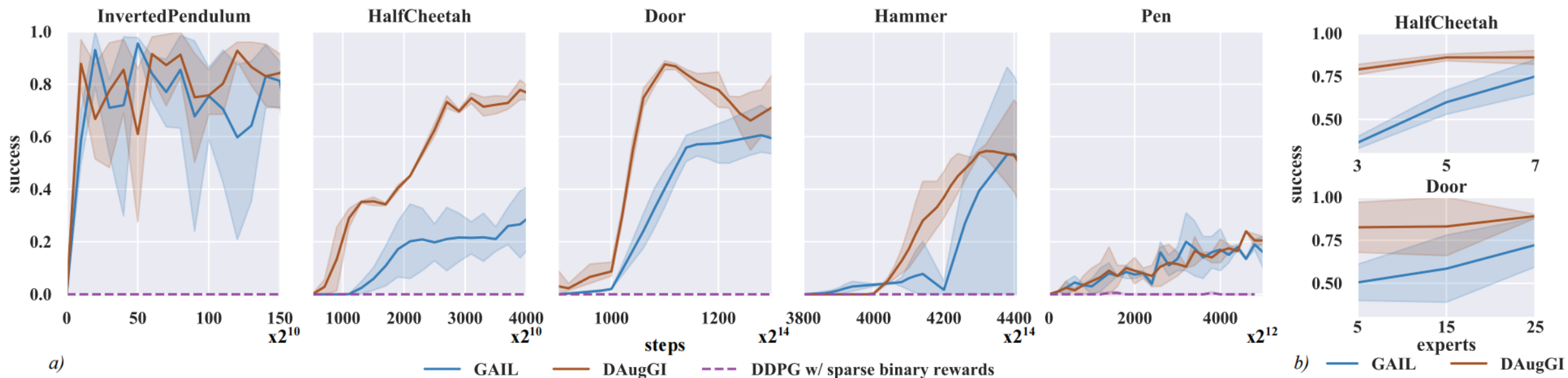


Fig. 5. a) Performance results of various tasks at different training steps. It includes easier OpenAI tasks (HalfCheetah and InvertedPendulum) with 3 experts, as well as more challenging dexterous manipulation tasks (Door, Hammer, Pen) with 25 experts. DAUGGI is trained using the augmented CAT trajectories and generally outperforms GAIL, trained with the original limited trajectories. b) Ablation studies with different number of experts for HalfCheetah and Door tasks. DAUGGI consistently outperforms GAIL, especially when the expert dataset is limited.



Augmenting Policy Learning with Routines Discovered from a Single Demonstration

Zelin Zhao ^{*} ¹, **Chuang Gan** ², **Jiajun Wu** ³, **Xiaoxiao Guo** ², **Joshua Tenenbaum** ⁴

¹ Shanghai Jiao Tong University, ² MIT-IBM Watson AI Lab

³ Stanford University, ⁴ Massachusetts Institute of Technology

sjtuytc@sjtu.edu.cn, ganchuang@csail.mit.edu, jiajunwu@cs.stanford.edu,
Xiaoxiao.Guo@ibm.com, jbt@mit.edu

AAAI 2021

Motivation

Humans can abstract **prior knowledge** from very little data and use it to boost skill learning.



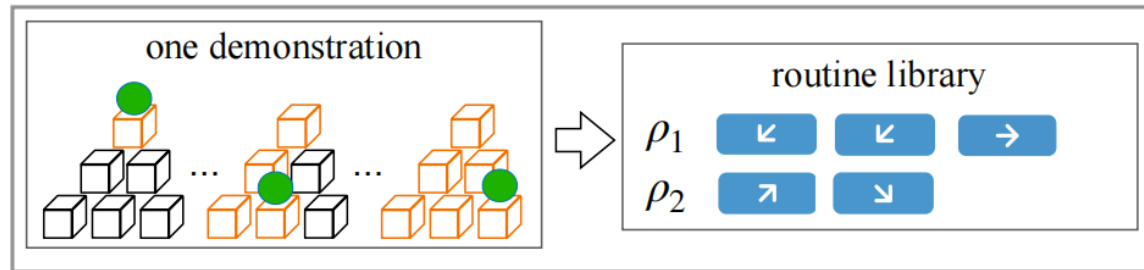
Discover **routines** composed of primitive actions from **a single demonstration** and use discovered routines to augment policy learning

Method

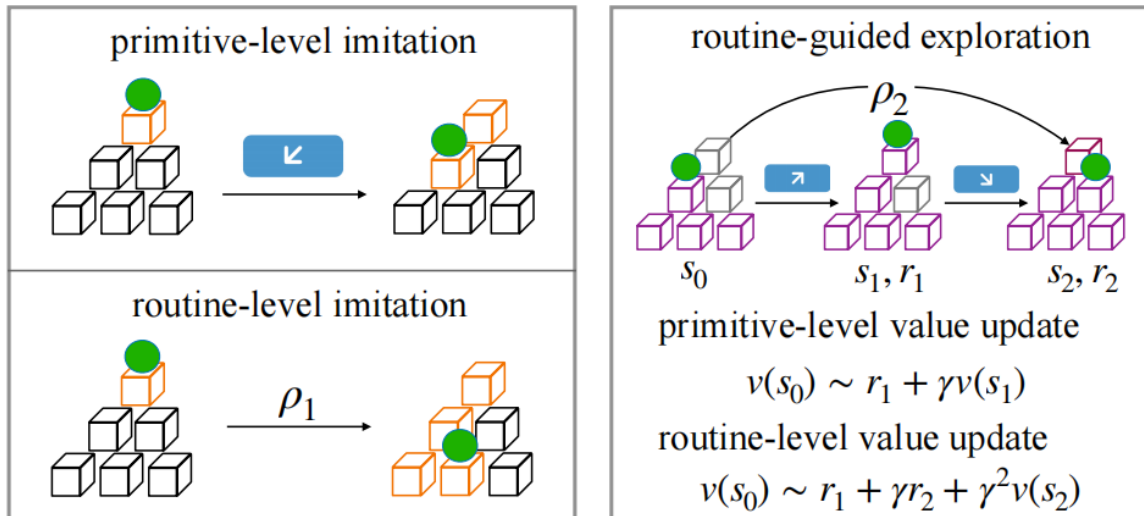
focus on MDPs with high-dimensional states and discrete actions

□ Routine-Augmented Policy Learning (RAPL)

(a) routine discovery



(b1) routine-augmented imitation (b2) routine-augmented RL

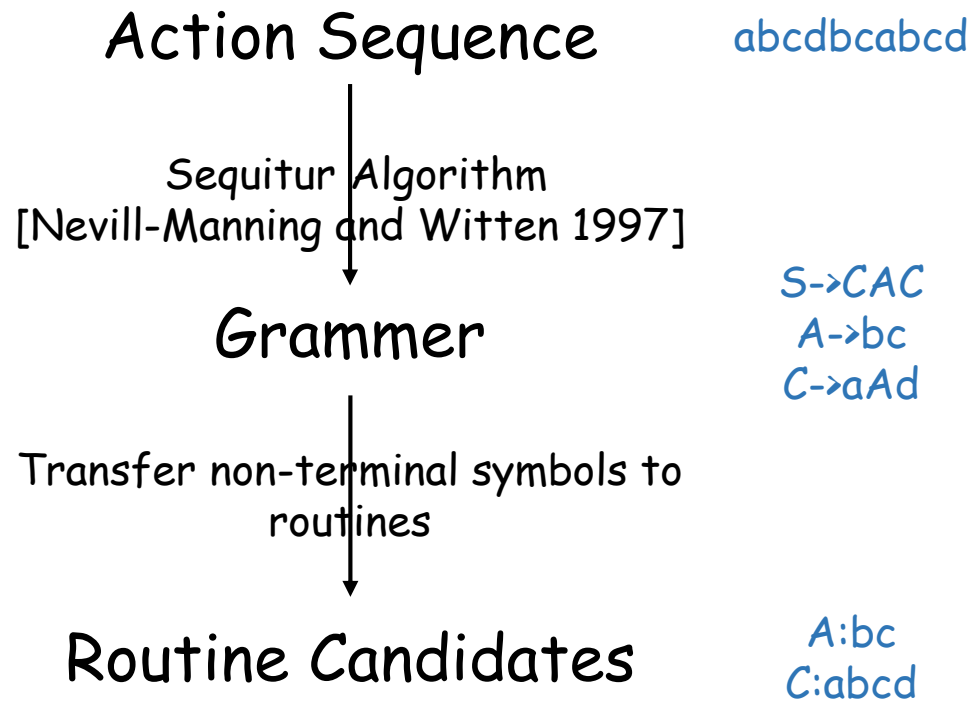


routine ρ : a sequence of primitive actions $(a^{(1)}, a^{(2)}, \dots, a^{(|\rho|)})$

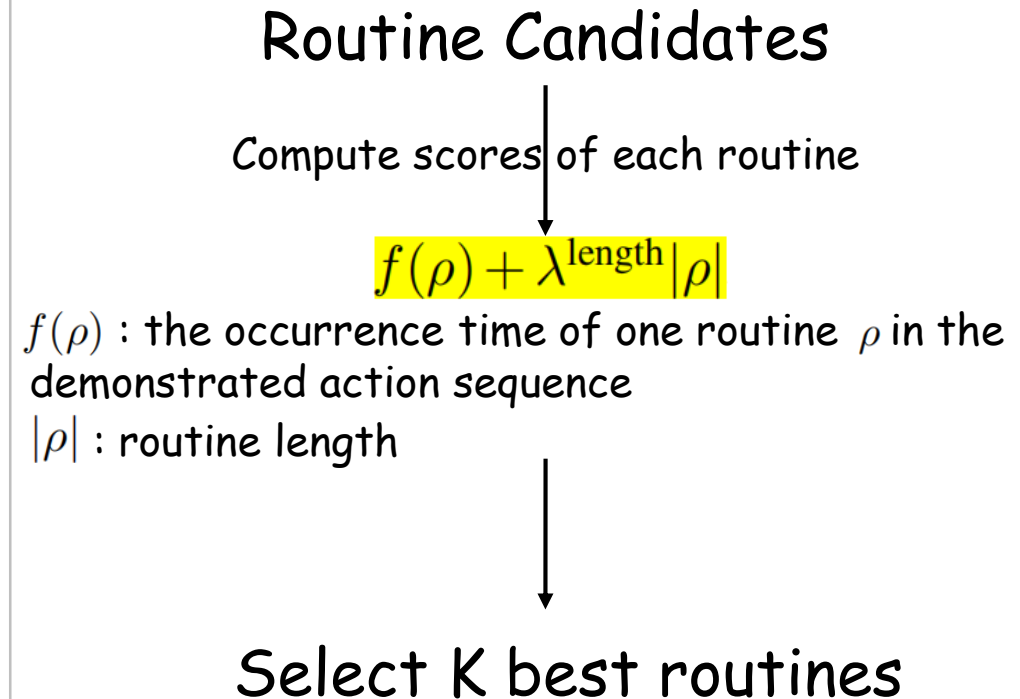
Method

□ Part 1 : Routine Discovery

Phase 1 : Routine Proposal



Phase 2 : Routine Selection



Method

□ Part 2 : Routine Policy Learning

RAPL-SQIL : using routines to augment imitation learning

Idea : imitate the expert's experiences at multiple temporal scales.

$$\mathcal{L}^{SR} = \delta^2(\mathcal{D}_{\text{prim}} \cup \mathcal{D}_{\text{routine}}, 1) + \lambda_{\text{sample}} \delta^2(\mathcal{D}_{\text{sample}}, 0)$$

$$\delta^2(\mathcal{D}, r) = \frac{1}{|\mathcal{D}|} \sum_{(s_t, \tilde{\rho}, s_{t+|\tilde{\rho}|}) \in \mathcal{D}} (Q_\theta(s_t, \tilde{\rho}) - Q_{\text{target}}(\tilde{\rho}, s_{t+|\tilde{\rho}|}, r))^2, \quad Q_{\text{target}}(\tilde{\rho}, s_{t+|\tilde{\rho}|}, r) = R_{sq}(\tilde{\rho}, r) + \Gamma(\tilde{\rho}) \log \left(\sum_{\tilde{\rho}' \in \tilde{\mathcal{L}}} \exp(Q_\theta(s_{t+|\tilde{\rho}|}, \tilde{\rho}')) \right),$$

$\mathcal{D}_{\text{prim}}$: the experience from primitive-level demonstration. (s_t, a, s_{t+1})

$\mathcal{D}_{\text{routine}}$: the demonstration explained by the abstracted routines. $(s_t, \rho, s_{t+|\rho|})$

$\mathcal{D}_{\text{sample}}$: the experiences collected via interaction with the environments.

Method

□ Part 2 : Routine Policy Learning

RAPL-A2C : using routines to augment reinforcement learning

Idea : conduct value approximation and policy optimization at multiple temporal scales.

$$\mathcal{L}^{\text{policy}} = -A_{\text{routine}} \log \pi(\tilde{\rho}_t | s_{t_0}; \theta_{\pi}),$$

$$\mathcal{L}^{\text{entropy}} = \sum_{\tilde{\rho}} \pi(\tilde{\rho} | s_{t_0}; \theta_{\pi}) \log \pi(\tilde{\rho} | s_{t_0}; \theta_{\pi}).$$

$$\mathcal{L}^{\mathcal{AR}} = \mathbb{E}(\mathcal{L}^{\text{policy}} + \lambda^{\text{entropy}} \mathcal{L}^{\text{entropy}} + \lambda^{\text{value}} (\|A_{\text{routine}}\|^2 + \lambda^{\text{prim}} \|A_{\text{prim}}\|^2))$$

$A_{\text{routine}} = \sum_{i=0}^{N-1} \gamma^{t_i - t_0} R_{t_i} + \gamma^{t_N - t_0} V(s_{t_N}) - V(s_{t_0})$: routine-level n-step bootstrap advantage function

$A_{\text{prim}} = \sum_{i=0}^{N-1} \gamma^i r_{t_j+i} + \gamma^N V(s_{t_j+N}) - V(s_{t_j})$: primitive-level n-step bootstrap advantage function

Experiments : RAPL-SQIL

Table 1: Comparing with several imitation learning baselines on 33 Atari games. We shown both alignment scores (defined in Eq. 10) and mean of human-normalized scores (Mnih et al. 2015) which indicates the alignment performance with regarding to the demonstration. Each number in the table is averaged over five random seeds.

	Alignment ($\pm$ std)	Mean ($\pm$ std)
BC	0.18 ($\pm$ 0.03)	18.3% (2.1%)
GAIL	0.16 ($\pm$ 0.08)	26.4% (1.6%)
SQIL	0.28 ($\pm$ 0.07)	29.4% (3.2%)
<u>RAPL-SQIL</u>	0.34 ($\pm$ 0.07)	36.1% ($\pm$ <u>3.6%</u>)

$$s = 1 - \frac{D(\iota_d, \iota_t)}{|\iota_d|}, \quad (10)$$

ι_d : the demonstrated action trajectory

ι_t : the action trajectory produced by the trained agent

D : Levenshtein distance

Experiments : RAPL-A2C

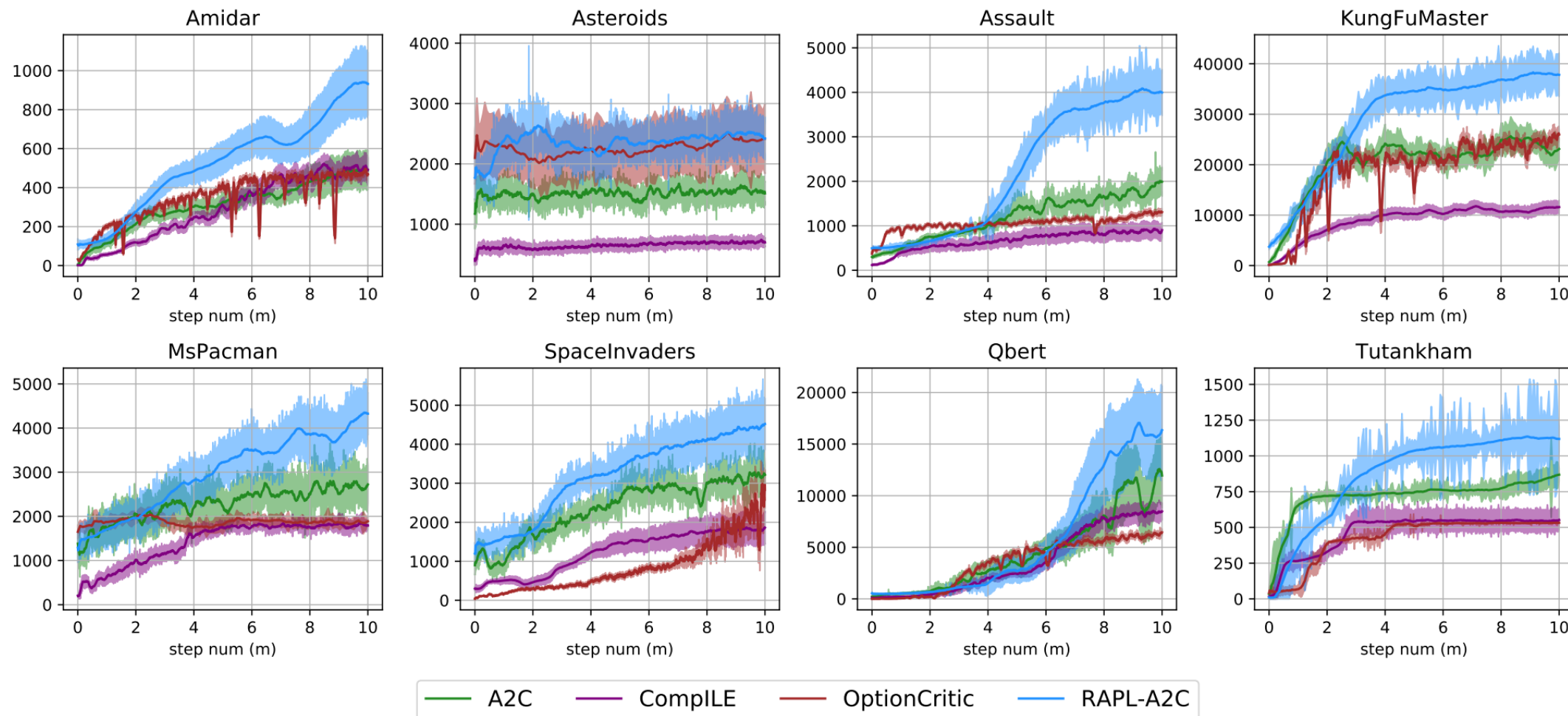


Figure 3: Training curves on eight randomly selected Atari games in comparison with several RL baselines. We plot both the mean and standard deviation in those curves across five agents with random seeds.

Experiments : RAPL-A2C

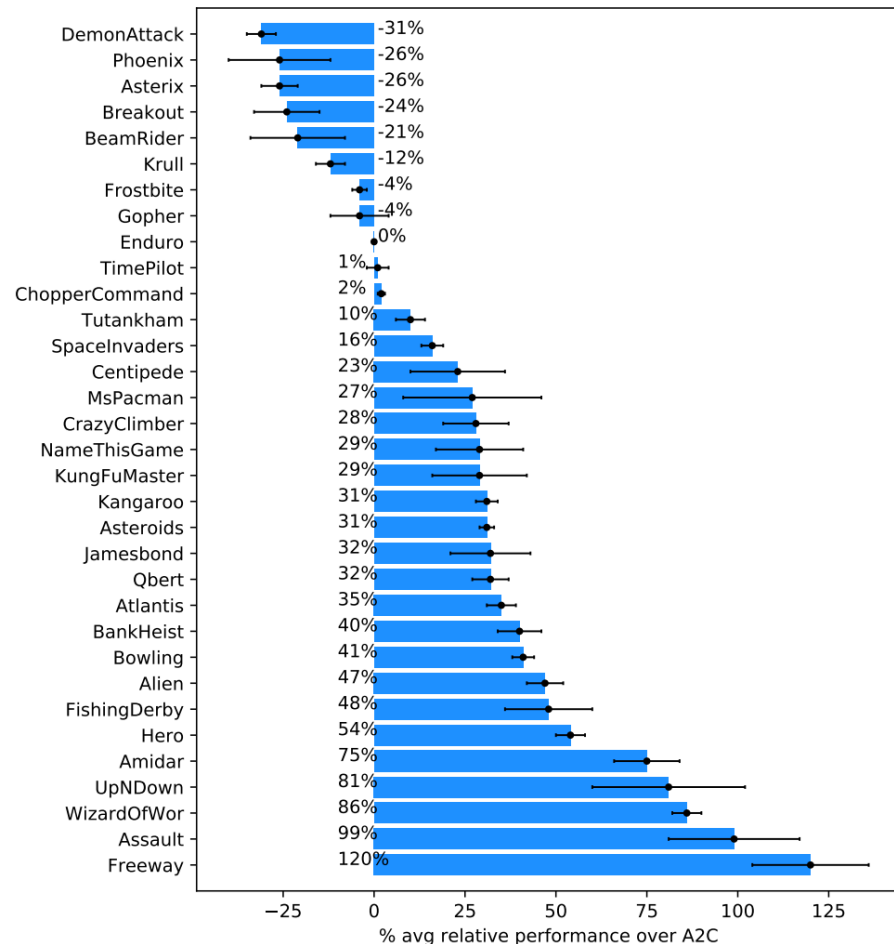
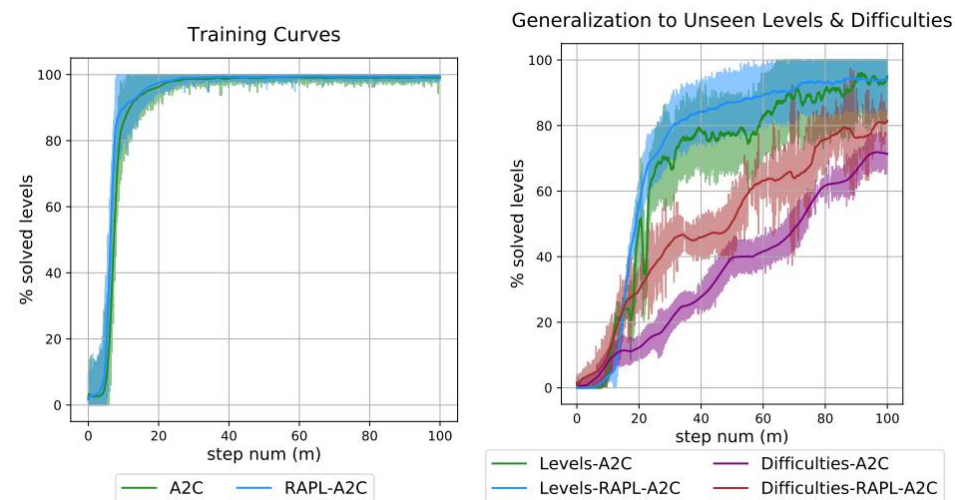
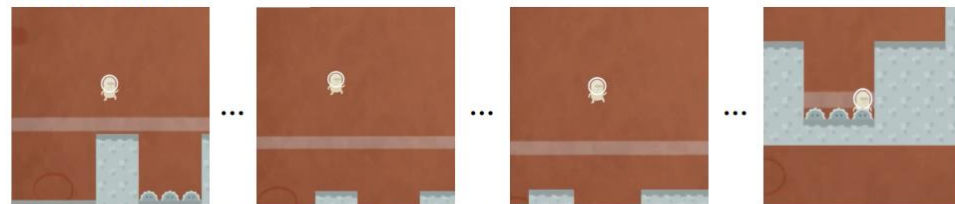


Figure 2: Relative performance of RAPL-A2C over A2C on Atari. Denote S_R as the score of RAPL-A2C and S_A is the score of A2C. The relative performance is calculated by $(S_R - S_A)/|S_A| \times 100\%$. Each number is averaged over five random agents and we also plot the stand error of the numbers.

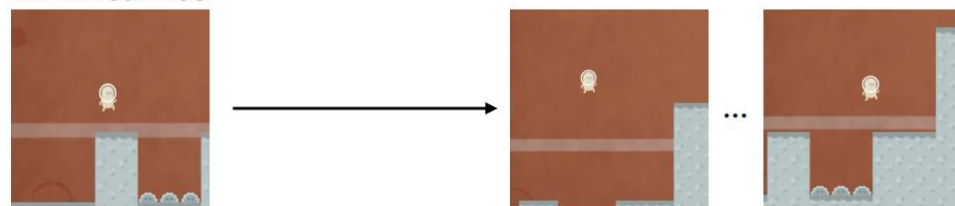


No Routines



Via primitives: Jump, Jump, Right, Right, Nope, Jump

With Routines



Via one routine: [Jump, Jump, Right, Right, Right, Jump]

Experiments : Effectiveness of Routine Discovery

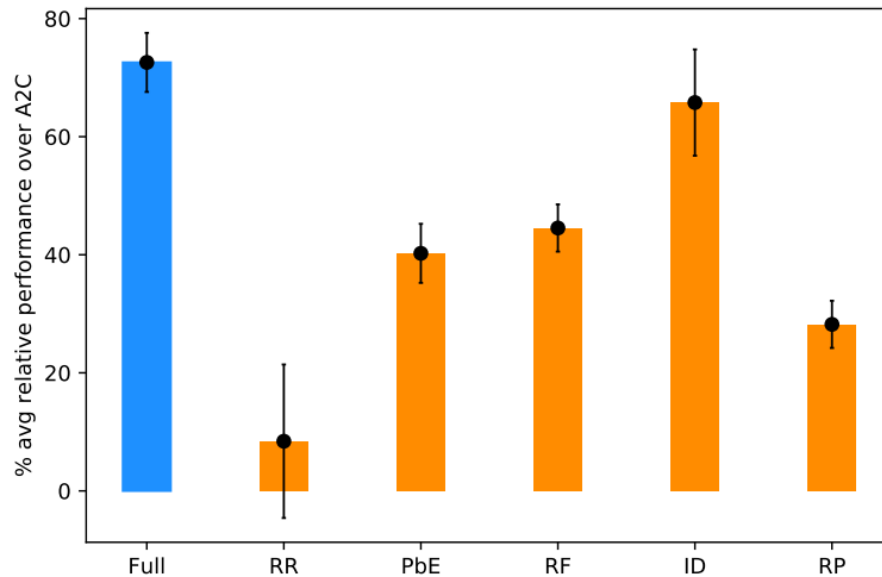


Figure 6: Comparison of ablated routine discovery models on Atari games. Mean and standard error over five random agents are shown in the figure.

(1) Random Routines (RR): each routine is generated randomly;

(2) proposal by Enumeration (PbE): enumerate all the possible combinations of primitive actions to form routine candidates.

(3) Random Fetch (RF): random fetch subsequences from the demonstration to form routines.

(4) Imperfect Demonstration (ID): the expert is only trained with 1 million steps.

(5) Repeat (RP): the routines are the repetition of most frequently used atomic actions in the demonstration

Thanks
