



南京航空航天大学
Nanjing University of Aeronautics and Astronautics



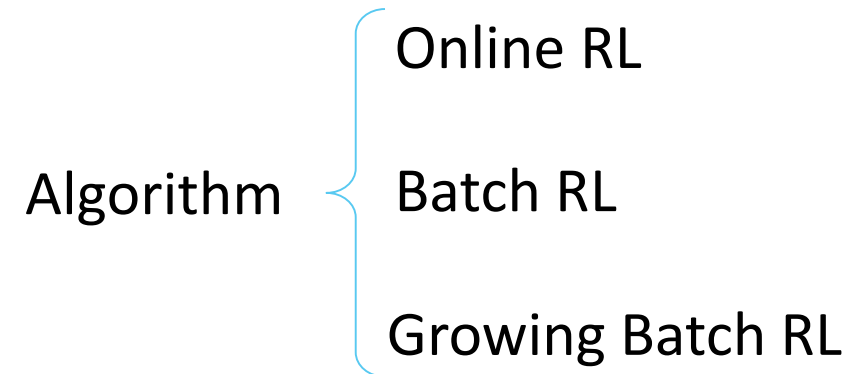
模式识别与神经计算研究组
PAttern Recognition and NEural Computing

An Introduction of Batch Reinforcement Learning

2021.11.08

Problems

- sample efficiency
- difficult to sample
- safety-critical systems



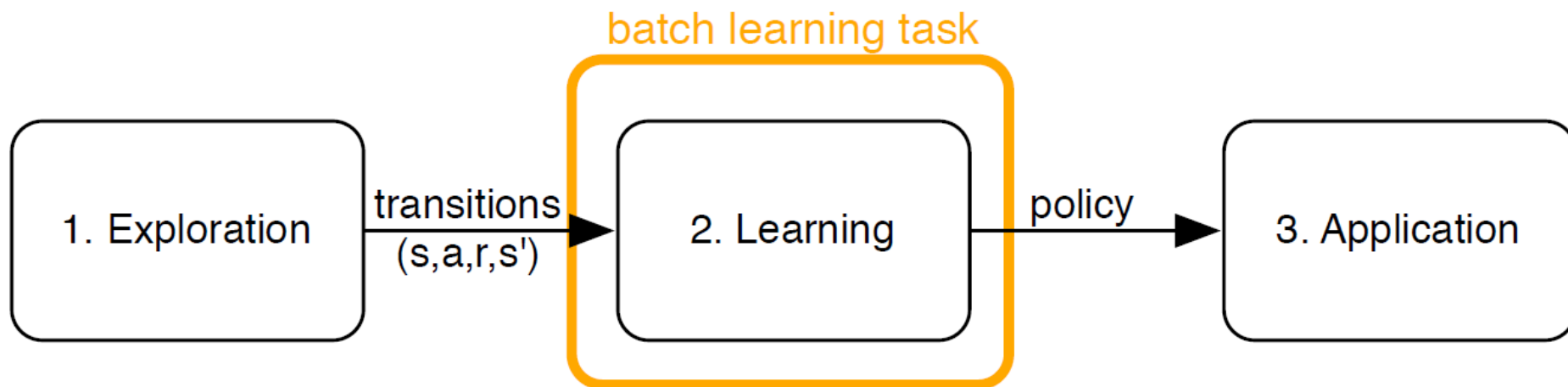


fig1: batch RL设定图

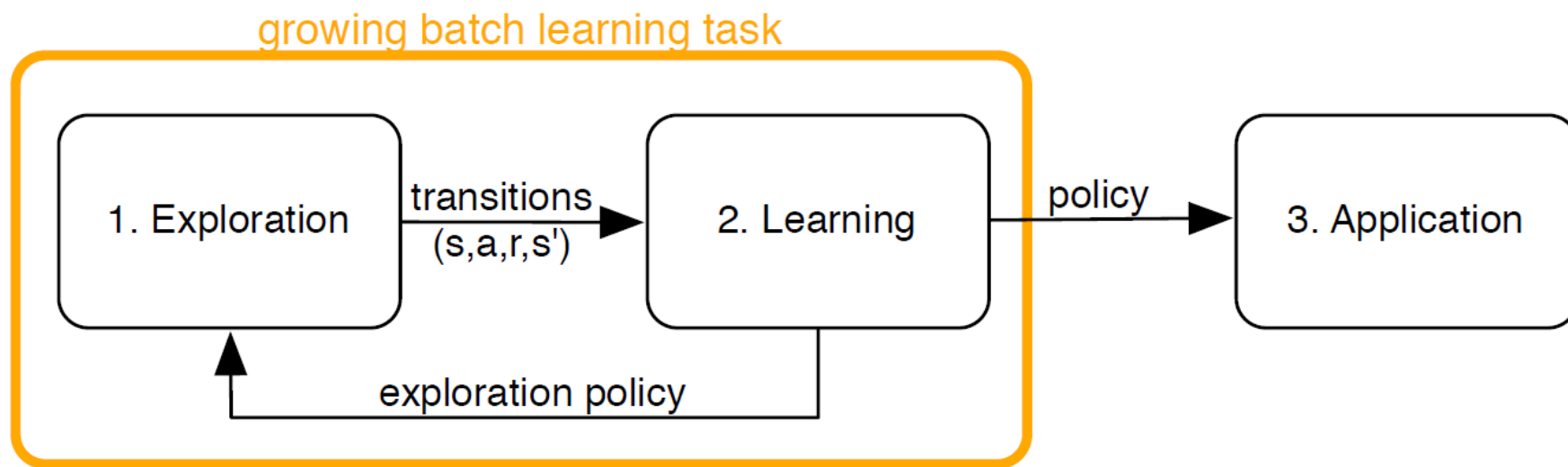


Fig2: growing batch RL设定图

Difference between batch rl and imitation learning

- Batch RL based on the [standard off-policy algorithm](#), usually needs to be optimized through [bellman equation or TD](#); and imitation learning is not required.
- Requirements for the data set: [very high-quality data sets or sub-optimal](#);
- Reward: consider [rewards](#) or [no reward](#);
- Imitation learning usually need to [label the dataset](#) according to the labels of [experts](#) and [non-experts](#); this problem does not need to be considered in Batch RL.



南京航空航天大学
Nanjing University of Aeronautics and Astronautics



模式识别与神经计算研究组
PAttern Recognition and NEural Computing

BAIL: Best-Action Imitation Learning for Batch Deep Reinforcement Learning

Xinyue Chen¹

Zijian Zhou¹

Zheng Wang¹

Che Wang^{1,2}

Yanqiu Wu^{1,2}

Keith Ross^{1,2*}

¹ New York University Shanghai

² New York University

Nips 2020

- Propose a notion named **upper envelope of the data**, try to make the best possible estimate of V^* using only the limited information in the batch dataset.
- Propose a method of **selecting the best state-action pairs**, then train a policy with IL using the selected state-action pairs. Combines both **V-learning and IL**.
- Conduct experiment on **MuJoCo benchmark** to demonstrate the superiority.

BAIL algorithm provide **state-of-the-art** performance on simulated robotic locomotion tasks, also **fast** and **algorithmically simple**.

- First, try to make the **best possible estimate** of $V(s)$ using only the limited information in the **batch dataset**.
- Then select state-action pairs from the dataset, whose **associated returns $G(s, a)$ are close to $V(s)$** .
- Finally, train a policy with Imitation Learning using the selected state-action pairs.

Thus, BAIL combines both V-learning and IL.

a batch of data $\mathcal{B} = \{(s_i, a_i, r_i, s'_i), i = 1, \dots, m\}$.

Assumption:

- data in the batch was generated in an **episodic fashion**.
- data in the batch is **ordered accordingly**.

For each data point $i \in \{1, \dots, m\}$, we calculate the Monte Carlo return G_i as the sum of the discounted rewards from state s_i to the end of the episode as $G_i = \sum_{t=i}^T \gamma^{t-i} r_t$

For a fixed $\lambda \geq 0$, we say that $V_{\phi^\lambda}(s)$ is a λ -regularized upper envelope for $\mathcal{G}$ if ϕ^λ is an optimal solution for the following constrained optimization problem:

$$\min_{\phi} \sum_{i=1}^m [V_{\phi}(s_i) - G_i]^2 + \lambda \|w\|^2 \quad s.t. \quad V_{\phi}(s_i) \geq G_i, \quad i = 1, 2, \dots, m \quad (1)$$

- When λ is very small, the upper envelope aims to **interpolate the data**.
- When λ is large, the upper envelope **approaches a constant** going through the data point with the **highest return**.



We solve the constrained optimization problem (1) with a penalty-function approach. Specifically, to obtain an approximate upper envelope of the data $\mathcal{G}$, we solve an unconstrained optimization problem with a penalty loss function (with λ fixed):

$$L^K(\phi) = \sum_{i=1}^m (V_\phi(s_i) - G_i)^2 \{ \mathbb{1}_{(V_\phi(s_i) \geq G_i)} + K \cdot \mathbb{1}_{(V_\phi(s_i) < G_i)} \} + \lambda \|w\|^2 \quad (2)$$

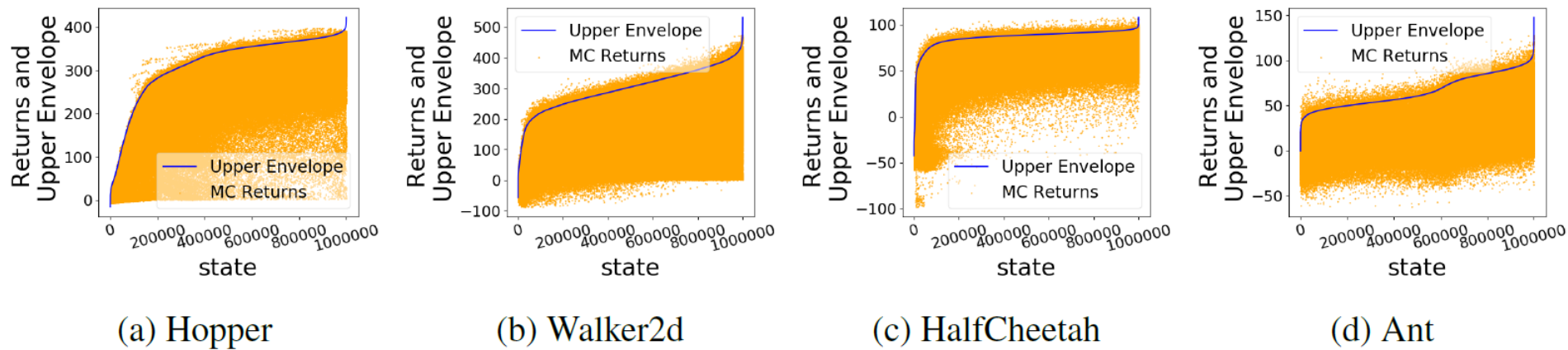


Figure 1: Upper Envelopes trained on batches from different MuJoCo environments.

Two approaches for selecting the best actions

- BAIL-ratio, for a fixed $x > 0$, choose all (s_i, a_i) pairs from the batch data such that:

$$G_i > xV(s_i)$$

- BAIL-difference, for a fixed $x > 0$, choose all (s_i, a_i) pairs from the batch data such that:

$$G_i \geq V(s_i) - x$$

The problem:

- For data points appearing **near the beginning of the episode**, the finite-horizon return will **closely approximate** the (idealized) infinite-horizon return due to discounting;
- For a data point **near the end of an episode**, the finite horizon return can be inaccurate and should be **augmented**.

$$G_i = \sum_{t=i}^T \gamma^{t-i} r_t + \gamma^{T-i+1} \sum_{t=j}^T \gamma^{t-j} r_t$$

Algorithm 1 BAIL

Initialize upper envelope parameters ϕ, ϕ' , policy parameters θ . Obtain batch data $\mathcal{B}$. Randomly split data into training set $\mathcal{B}_t$ and validation set $\mathcal{B}_v$ for the upper envelope.

Compute return G_i for each data point i in $\mathcal{B}$.

Obtain upper envelope by minimizing the loss $L^K(\phi)$:

for $j = 1, \dots, J$ **do**

 Sample a mini-batch B from $\mathcal{B}$.

 Update ϕ using the gradient: $\nabla_{\phi} \sum_{i \in B} (V_{\phi}(s_i) - G_i)^2 \{ \mathbb{1}_{(V_{\phi}(s_i) > G_i)} + K \mathbb{1}_{(V_{\phi}(s_i) < G_i)} \} + \lambda \|\phi\|^2$

if time to do validation for the upper envelope **then**

 Compute validation loss on $\mathcal{B}_v$

 Update ϕ and ϕ' according to the validation loss

end if

end for

Select data point i if $G_i > xV_{\phi}(s_i)$, where x is such that $p\%$ of data in $\mathcal{B}$ are selected. Let $\mathcal{U}$ be the set of selected data points.

for $l = 1, \dots, L$ **do**

 Sample a mini-batch U of data from $\mathcal{U}$.

 Update θ using the gradient: $\nabla_{\theta} \sum_{i \in U} (\pi_{\theta}(s_i) - a_i)^2$

end for

Algorithm 2 Progressive BAIL

Initialize upper envelope parameters ϕ, ϕ' , policy parameters θ .

Obtain batch data $\mathcal{B}$. Randomly split data into training set $\mathcal{B}_t$ and validation set $\mathcal{B}_v$ for the upper envelope.

Compute return G_i for each data point i in $\mathcal{B}$.

for $l = 1, \dots, L$ **do**

 Sample a mini-batch of data B from the batch $\mathcal{B}_t$.

 Update ϕ using the gradient: $\nabla_{\phi} \sum_{i \in B_t} (V_{\phi}(s_i) - G_i)^2 \{ \mathbb{1}_{(V_{\phi}(s_i) > G_i)} + K \mathbb{1}_{(V_{\phi}(s_i) < G_i)} \} + \lambda \|\phi\|^2$

if time to validate **then**

 Compute validation loss on B_v

 Update ϕ and ϕ' according to validation loss

end if

 Select data point i if $G_i > x V_{\phi}(s_i)$, where x is such that $p\%$ of data in B are selected. Let U be the set of selected data points.

 Update θ using the gradient: $\nabla_{\theta} \sum_{i \in U} (\pi_{\theta}(s_i) - a_i)^2$

end for

Table 1: Performance of five Batch DRL algorithms for 22 different training datasets.

ENVIRONMENT	BAIL	BCQ	BEAR	BC	MARWIL
$\sigma = 0.1$ HOPPER B1	2173 $\pm$ 291	1219 $\pm$ 114	505 $\pm$ 285	626 $\pm$ 112	827 $\pm$ 220
$\sigma = 0.1$ HOPPER B2	2078 $\pm$ 180	1178 $\pm$ 87	985 $\pm$ 3	579 $\pm$ 141	620 $\pm$ 336
$\sigma = 0.1$ WALKER B1	1125 $\pm$ 113	576 $\pm$ 309	610 $\pm$ 212	514 $\pm$ 17	436 $\pm$ 24
$\sigma = 0.1$ WALKER B2	3141 $\pm$ 300	2338 $\pm$ 388	2707 $\pm$ 425	1741 $\pm$ 239	1810 $\pm$ 200
$\sigma = 0.1$ HC B1	5746 $\pm$ 29	5883 $\pm$ 43	0 $\pm$ 0	5546 $\pm$ 29	5573 $\pm$ 35
$\sigma = 0.1$ HC B2	7212 $\pm$ 43	7562 $\pm$ 31	0 $\pm$ 0	6765 $\pm$ 108	6828 $\pm$ 111
$\sigma = 0.5$ HOPPER B1	2054 $\pm$ 158	1145 $\pm$ 300	203 $\pm$ 42	919 $\pm$ 52	946 $\pm$ 103
$\sigma = 0.5$ HOPPER B2	2623 $\pm$ 282	1823 $\pm$ 555	241 $\pm$ 239	694 $\pm$ 64	818 $\pm$ 112
$\sigma = 0.5$ WALKER B1	2522 $\pm$ 51	1552 $\pm$ 455	1248 $\pm$ 181	2178 $\pm$ 178	2111 $\pm$ 52
$\sigma = 0.5$ WALKER B2	3115 $\pm$ 133	2785 $\pm$ 123	2302 $\pm$ 630	2483 $\pm$ 94	2364 $\pm$ 228
$\sigma = 0.5$ HC B1	1055 $\pm$ 9	1222 $\pm$ 38	924 $\pm$ 579	570 $\pm$ 35	512 $\pm$ 43
$\sigma = 0.5$ HC B2	7173 $\pm$ 120	5807 $\pm$ 249	-114 $\pm$ 140	6545 $\pm$ 171	6668 $\pm$ 93
SAC HOPPER B1	3296 $\pm$ 105	2681 $\pm$ 438	1000 $\pm$ 110	2853 $\pm$ 318	2897 $\pm$ 227
SAC HOPPER B2	1831 $\pm$ 915	2134 $\pm$ 917	1139 $\pm$ 317	2240 $\pm$ 367	2063 $\pm$ 168
SAC WALKER B1	2455 $\pm$ 211	2408 $\pm$ 84	-3 $\pm$ 5	1674 $\pm$ 277	1484 $\pm$ 140
SAC WALKER B2	4767 $\pm$ 130	3794 $\pm$ 398	325 $\pm$ 75	2599 $\pm$ 145	2651 $\pm$ 268
SAC HC B1	10143 $\pm$ 77	8607 $\pm$ 473	7392 $\pm$ 257	8874 $\pm$ 221	9105 $\pm$ 90
SAC HC B2	10772 $\pm$ 59	10106 $\pm$ 134	7217 $\pm$ 273	9523 $\pm$ 164	9488 $\pm$ 136
SAC ANT B1	4284 $\pm$ 64	4042 $\pm$ 113	3452 $\pm$ 128	3986 $\pm$ 112	4033 $\pm$ 130
SAC ANT B2	4946 $\pm$ 148	4640 $\pm$ 76	3712 $\pm$ 236	4618 $\pm$ 111	4589 $\pm$ 130
SAC HUMANOID B1	3852 $\pm$ 430	1411 $\pm$ 250	0 $\pm$ 0	543 $\pm$ 378	589 $\pm$ 121
SAC HUMANOID B2	3565 $\pm$ 153	1221 $\pm$ 207	0 $\pm$ 0	1216 $\pm$ 826	1033 $\pm$ 257

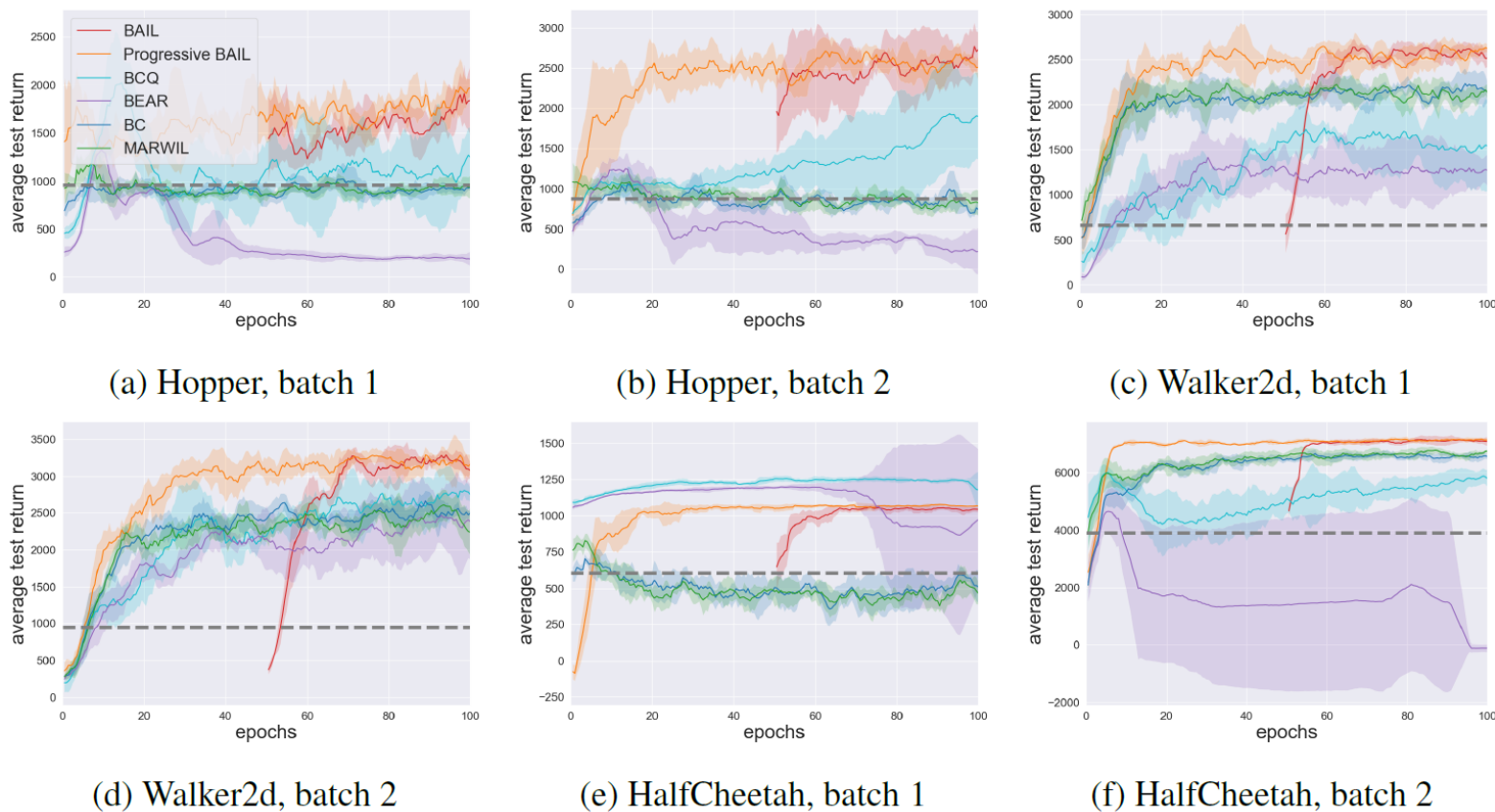


Figure 2: Learning curves using DDPG training batches with $\sigma = 0.5$.

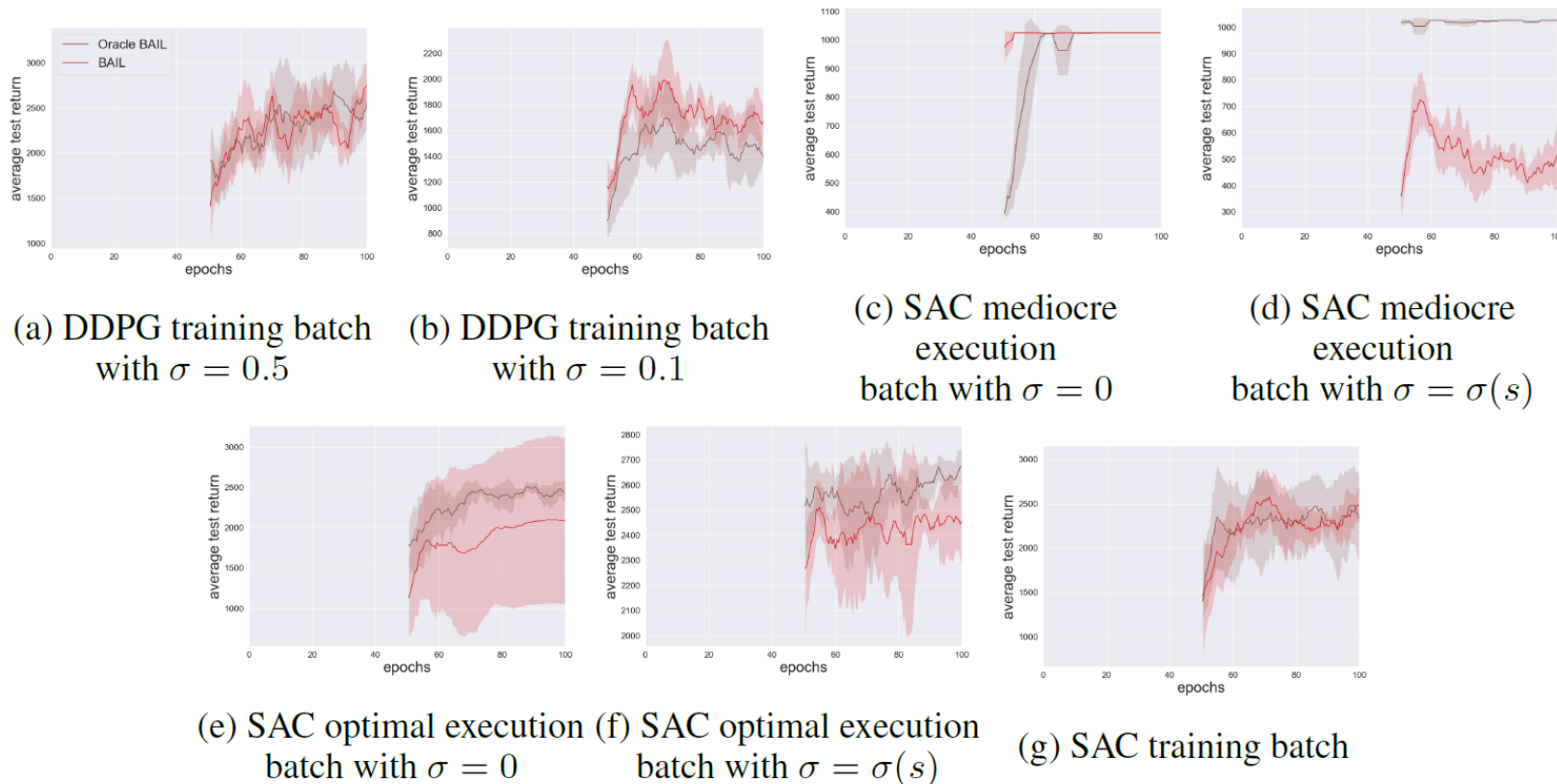


Figure 4: Augmented Returns versus Oracle Performance. All learning curves are for the Hopper-v2 environment. The x-axis ranges from 50 to 100 epochs since this comparison involves only BAIL. The results show that the augmentation heuristic typically achieves oracle-level performance.

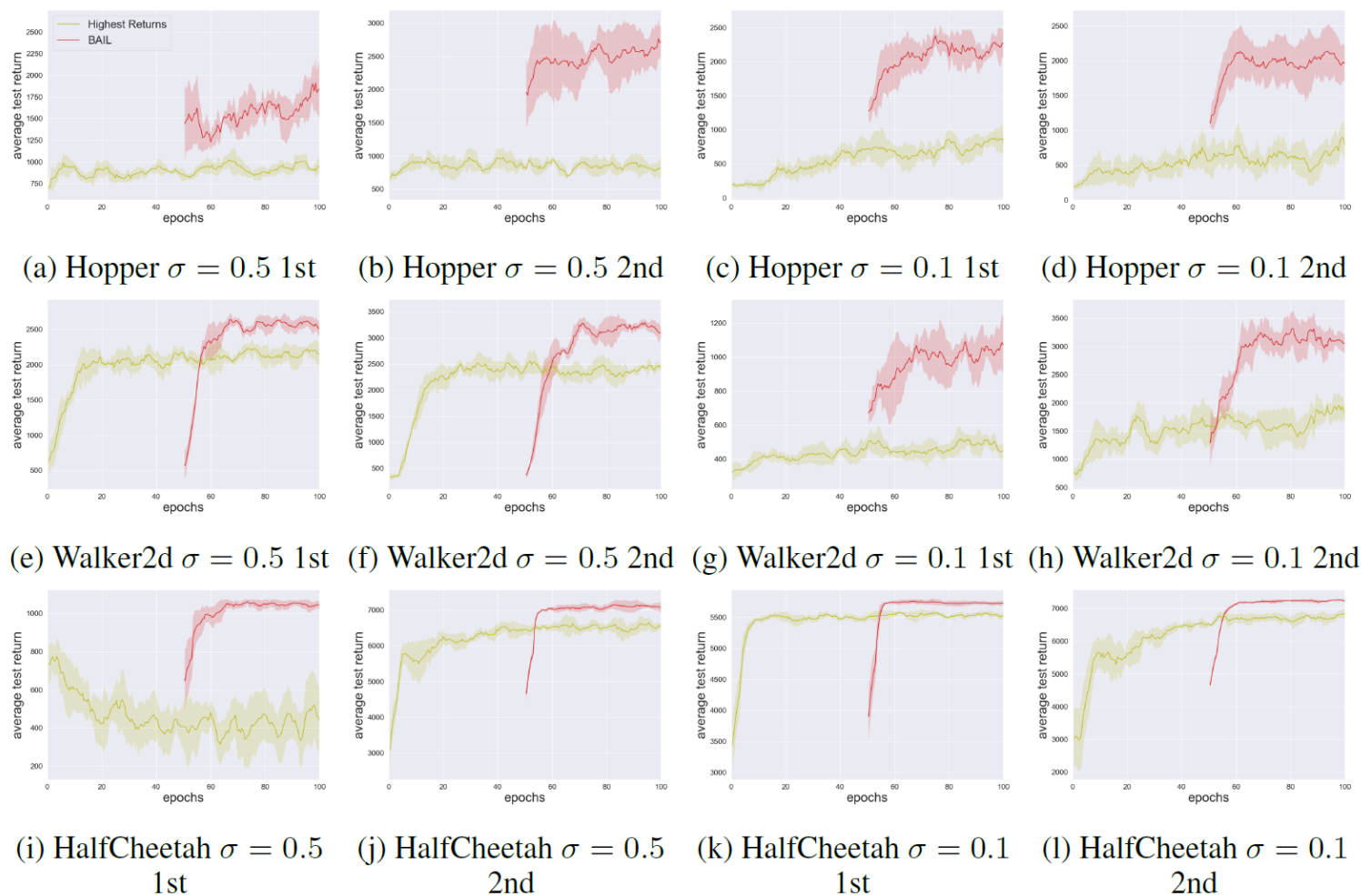
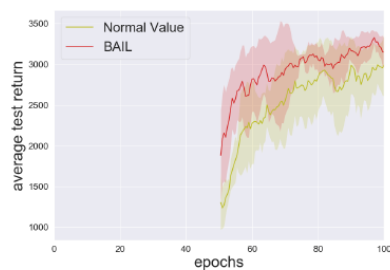


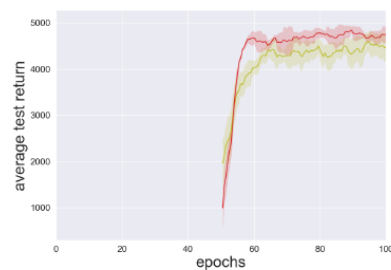
Figure 5: Ablation study for data selection. The figure compares BAIL with the algorithm that simply chooses the state-action pairs with the highest returns (without using an upper envelope). The learning curves show that the upper envelope is critical components of BAIL.

D.3 Ablation study using standard regression instead of an upper envelope

Figure 6 compares BAIL with the more naive scheme of using standard regression in place of an upper envelope. The learning curves show that the upper envelope is a critical component of BAIL.



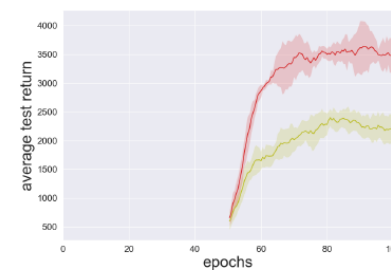
(a) Training SAC, Hopper



(b) Training SAC, Walker2d



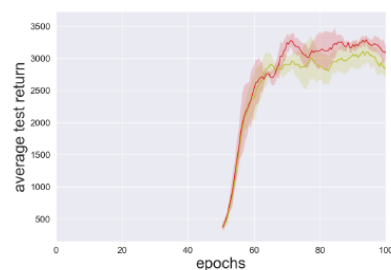
(c) Training SAC, Ant



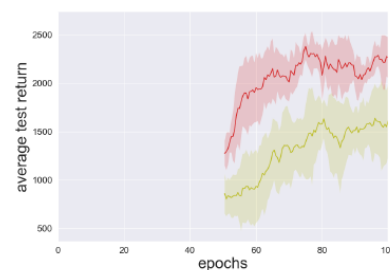
(d) Training SAC, Humanoid



(e) Training DDPG $\sigma = 0.5$, Hopper



(f) Training DDPG $\sigma = 0.5$, Walker2d



(g) Training DDPG $\sigma = 0.1$, Hopper



(h) Training DDPG $\sigma = 0.1$, Walker2d

THANKS