



Evolutionary Diversity Optimization with Clustering-based Selection for Reinforcement Learning

Yutong Wang, Ke Xue, Chao Qian

ICLR'22

LAMDA 11

Pre by Ke Xue

22/06/24

Introduction

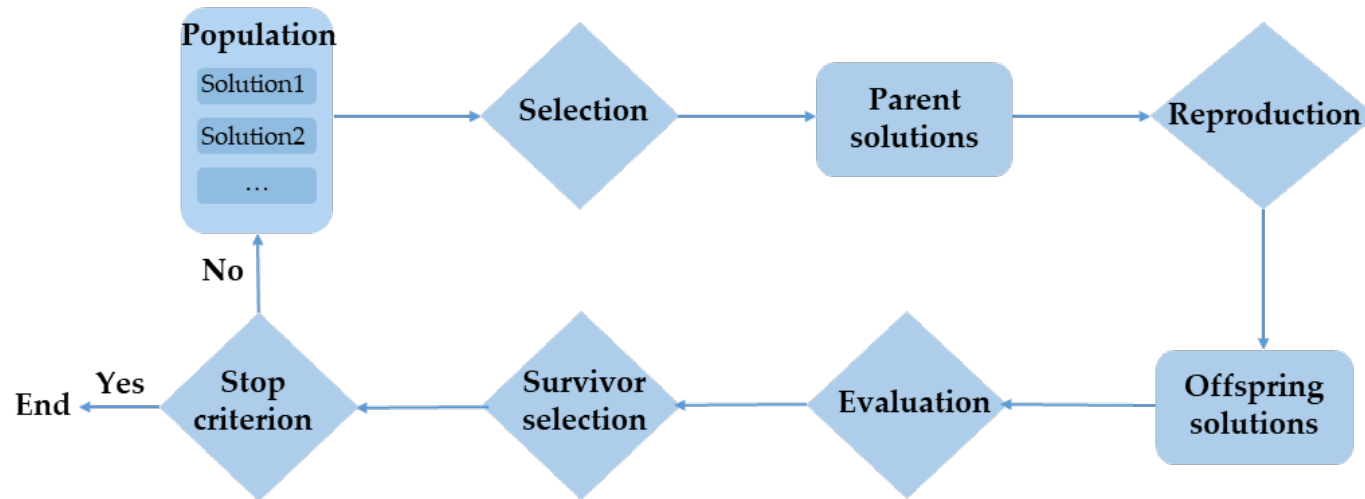
Reinforcement Learning (RL)

- Most RL methods aim to obtain a single optimal policy
- Some complex scenarios need **a set of diverse policies**
 - Exploration
 - Policy ensemble
 - Model/environment generation
 - Few-shot adaption

*How to efficiently obtain **a set of high-quality policies with diverse behaviors** is a challenging problem in RL*

Background

- Evolutionary Algorithm

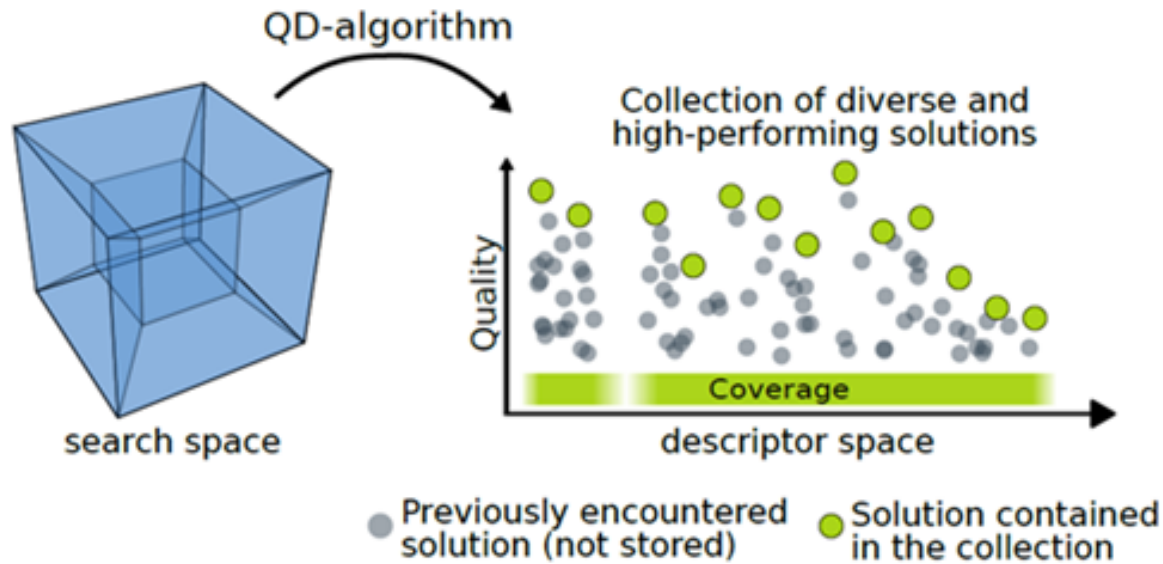
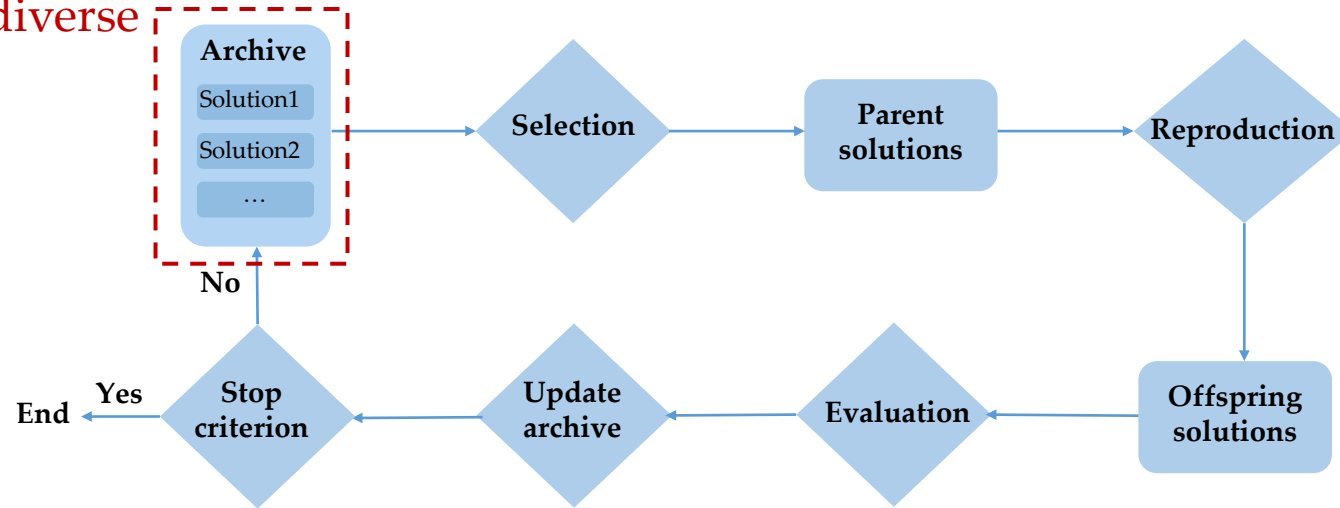


- Key properties
 - Population-based search
 - Only evaluation is required
- Application: complex optimization problems
 - Black-box optimization problems
 - Multi-objective problems

Background

- Quality-diversity

To be diverse



Background - QD

Algorithm 2 QD-Optimization algorithm (I iterations)

```
 $\mathcal{A} \leftarrow \emptyset$  ▷ Creation of an empty container.  
for iter = 1  $\rightarrow$   $I$  do ▷ The main loop repeats during  $I$  iterations.  
  if iter == 1 then ▷ Initialization.  
     $\mathcal{P}_{\text{parents}} \leftarrow \text{random}()$  ▷ The first 2 batches of individuals are generated randomly.  
     $\mathcal{P}_{\text{offspring}} \leftarrow \text{random}()$   
  else ▷ The next controllers are generated using the container and/or the previous batch.  
     $\mathcal{P}_{\text{parents}} \leftarrow \text{selection}(\mathcal{A}, \mathcal{P}_{\text{offspring}})$  Selection ▷ Selection of a batch of individuals from the container and/or the previous batch.  
     $\mathcal{P}_{\text{offspring}} \leftarrow \text{variation}(\mathcal{P}_{\text{parents}})$  Reproduction ▷ Creation of a randomly modified copy of  $\mathcal{P}_{\text{parents}}$  (mutation and/or crossover).  
    for each  $\theta \in \mathcal{P}_{\text{offspring}}$  do  
       $\{f_{\theta}, b_{\theta}\} \leftarrow f(\theta)$  ▷ Evaluation of the individual and recording of its descriptor and performance.  
      if ADD_TO_CONTAINER( $\theta, \mathcal{A}$ ) then ▷ “ADD_TO_CONTAINER” returns true if the individual has been added to the container.  
        UPDATE_SCORES(parent( $\theta$ ), Reward,  $\mathcal{A}$ ) ▷ The parent might get a reward.  
      else  
        UPDATE_SCORES(parent( $\theta$ ), -Penalty,  $\mathcal{A}$ ) ▷ Otherwise, it might get a penalty.  
        UPDATE_CONTAINER( $\mathcal{A}$ ) ▷ Update of the attributes of all the individuals in the container (e.g. novelty score).  
return  $\mathcal{A}$ 
```

Background - QD

Algorithm 2 QD-Optimization algorithm (I iterations)

```
 $\mathcal{A} \leftarrow \emptyset$  ▷ Creation of an empty container.  
for iter = 1  $\rightarrow$   $I$  do ▷ The main loop repeats during  $I$  iterations.  
  if iter == 1 then ▷ Initialization.  
     $\mathcal{P}_{\text{parents}} \leftarrow \text{random}()$  ▷ The first 2 batches of individuals are generated randomly.  
     $\mathcal{P}_{\text{offspring}} \leftarrow \text{random}()$   
  else ▷ The next controllers are generated using the container and/or the previous batch.  
     $\mathcal{P}_{\text{parents}} \leftarrow \text{selection}(\mathcal{A}, \mathcal{P}_{\text{offspring}})$  ▷ Selection of a batch of individuals from the container and/or the previous batch.  
     $\mathcal{P}_{\text{offspring}} \leftarrow \text{variation}(\mathcal{P}_{\text{parents}})$  ▷ Creation of a randomly modified copy of  $\mathcal{P}_{\text{parents}}$  (mutation and/or crossover).  
  for each  $\theta \in \mathcal{P}_{\text{offspring}}$  do  
     $\{f_{\theta}, b_{\theta}\} \leftarrow f(\theta)$  ▷ Evaluation of the individual and recording of its descriptor and performance.  
    if ADD_TO_CONTAINER( $\theta, \mathcal{A}$ ) then ▷ “ADD_TO_CONTAINER” returns true if the individual has been added to the container.  
      UPDATE_SCORES(parent( $\theta$ ), Reward,  $\mathcal{A}$ ) ▷ The parent might get a reward.  
    else  
      UPDATE_SCORES(parent( $\theta$ ), -Penalty,  $\mathcal{A}$ ) ▷ Otherwise, it might get a penalty.  
    UPDATE_CONTAINER( $\mathcal{A}$ ) ▷ Update of the attributes of all the individuals in the container (e.g. novelty score).  
return  $\mathcal{A}$ 
```

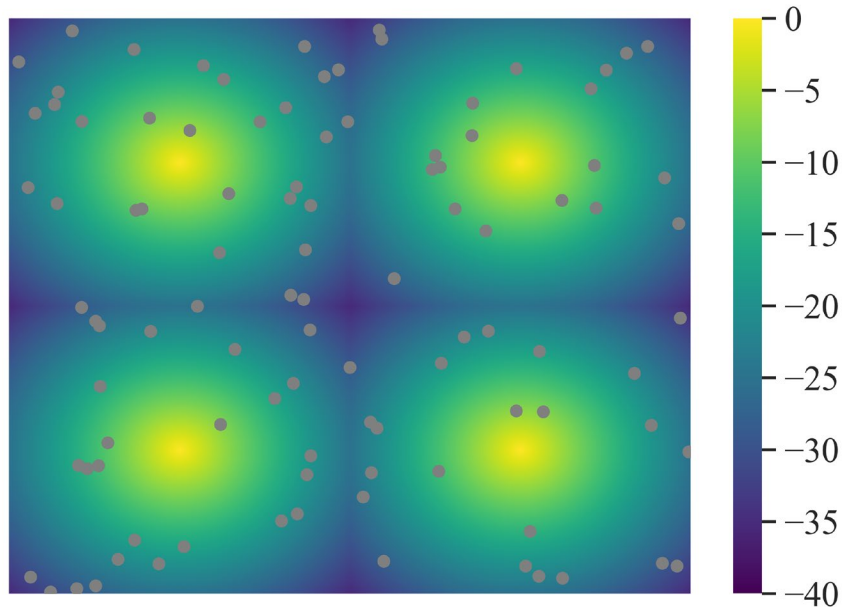
- Policy gradient assisted MAP-Elites (PGA-ME) [GECCO'21 Best paper]
 - Using PG to update the solutions
- Differentiable Quality-diversity (DQD) [NeurIPS'21 Oral, < 2.4%]
 - Using differentiable *behavior descriptor* b

Method - motivation

The **inefficient** selection

- Uniform selection
- Biased selection
- Pareto-based selection
- ...

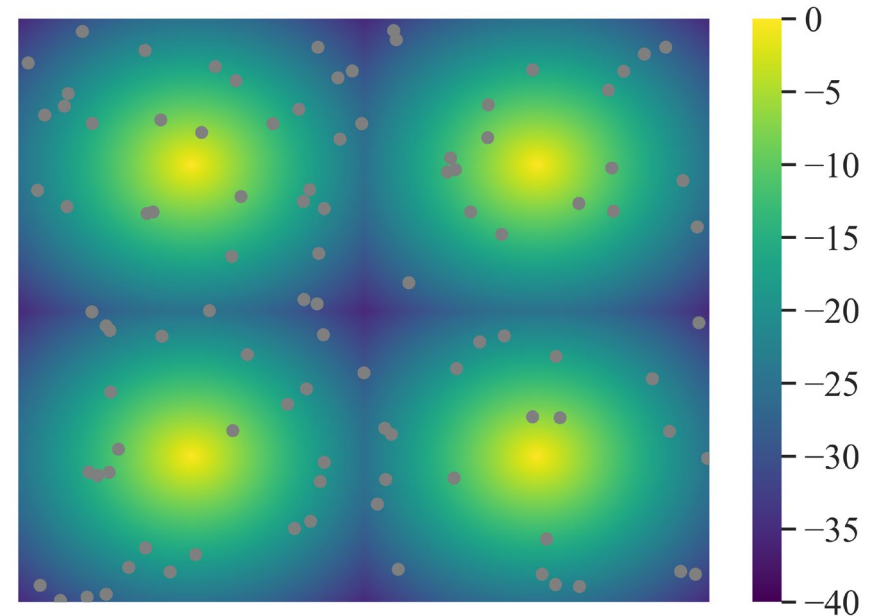
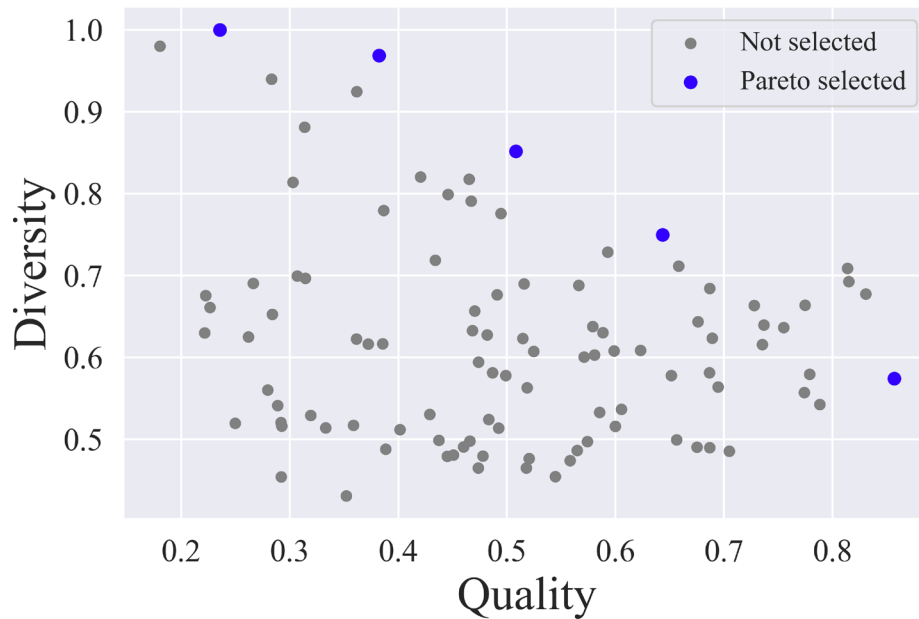
Method	Selection
Vanilla ES	The only parent solution
NSR-ES	Probabilistic selection
CVT-ES	Uniform selection
ME-ES	Biased selection
DvD-ES	All parent solutions
QD-RL	Pareto-based selection



Method - motivation

The **inefficient** selection

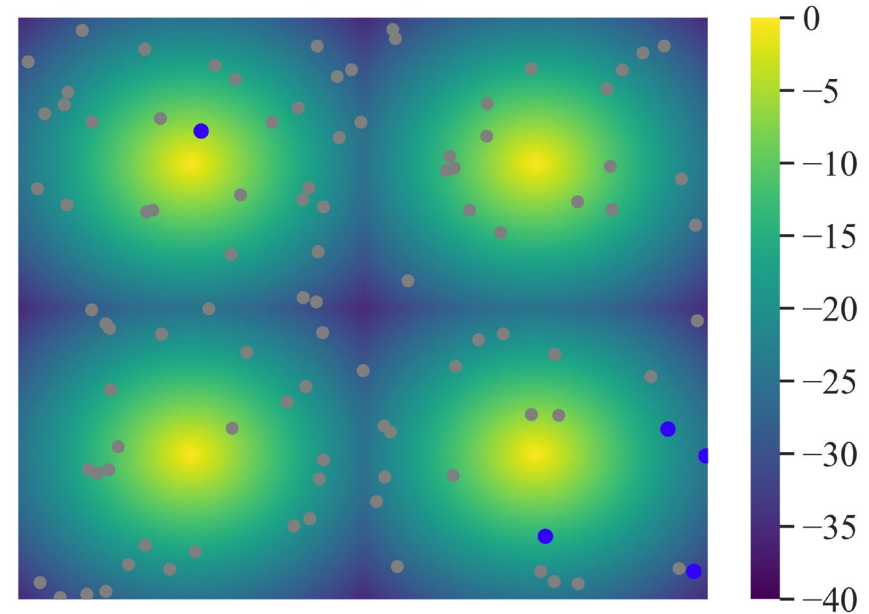
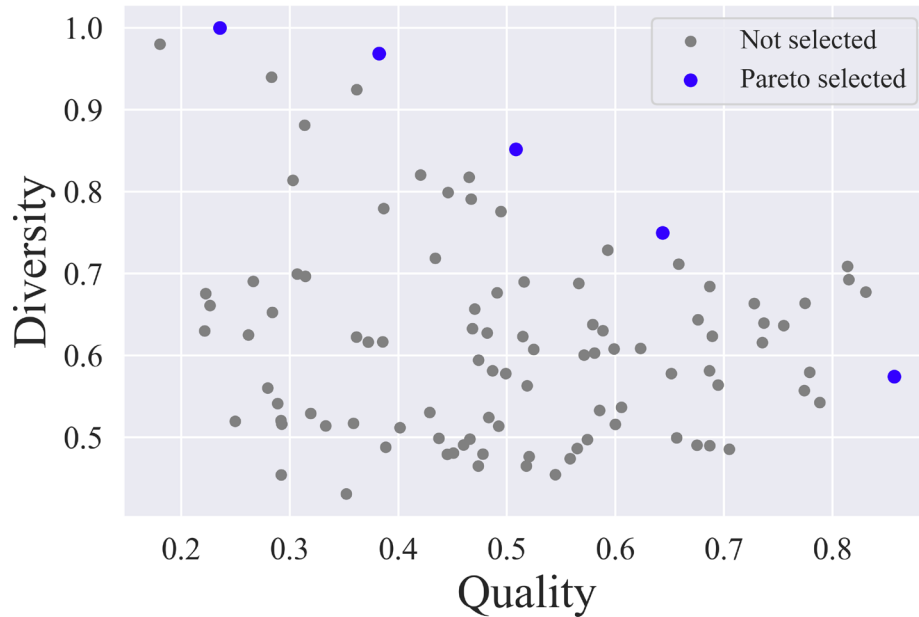
- Uniform selection
- Biased selection
- Pareto-based selection



Method - motivation

The **inefficient** selection

- Uniform selection
- Biased selection
- Pareto-based selection



How to achieve efficient selection?

Method

Algorithm 1: EDO-CS

Input: number K of selected policies, number T of total iterations, number T' of updating iterations, behavior characterization $b(\pi_{\theta})$, archive size l , archive A , bandit B

Output: archive A

```

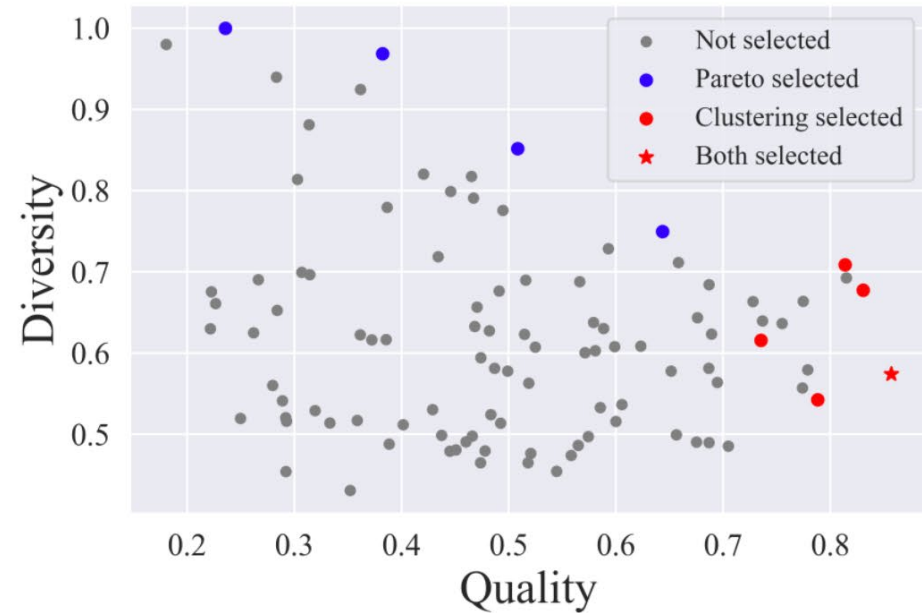
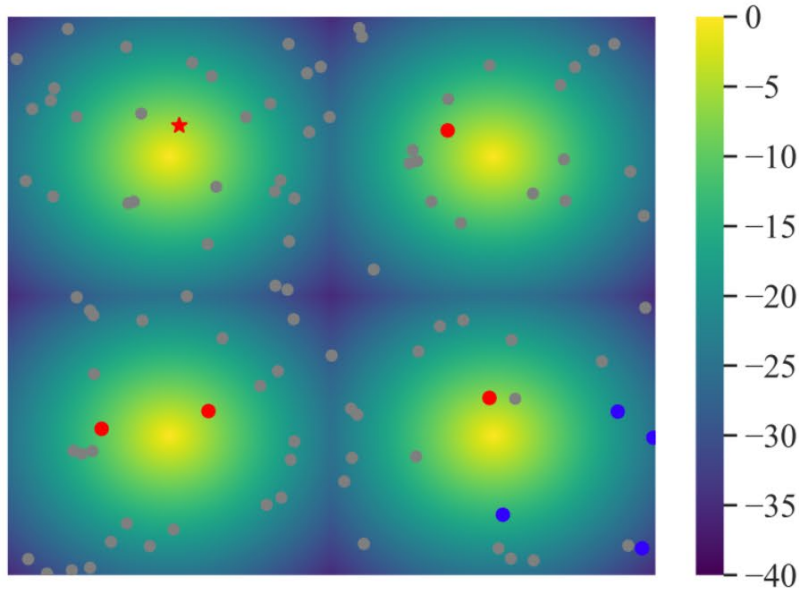
// Warm up
1 for  $j = 1 : l$  do
2   Randomly generate policy  $\pi_{\theta_j}$ ;
3   Get cumulative rewards  $R$  and behavior  $b(\pi_{\theta_j})$  by evaluating the policy  $\pi_{\theta_j}$ ;
4   Add  $(\pi_{\theta_j}, R, b(\pi_{\theta_j}))$  into archive  $A$ 
5 end
6  $t = 0$ ;
7 while  $t < T$  do
  // Selection
  8 Use  $K$ -means to divide the policies in archive  $A$  into  $K$  clusters  $\{C_k\}_{k=1}^K$ ;
  9 Select  $K$  policies  $\{\pi_{\theta_k}\}_{k=1}^K$ , each one from a cluster;
  // Reproduction
  10 for  $k = 1 : K$  do
    // Update in parallel
    11 Sample  $\lambda_k$  from the bandit  $B$ ;
    12 for  $i = 1 : T'$  do
      13 Set the objective function  $J(\theta_k) = (1 - \lambda_k)\mathbb{E}_{\tau \sim \pi_{\theta_k}}[R(\tau)] + \lambda_k \text{Div}(\theta_k)$ ;
      14 Use ES to update  $\theta_k$  as Eq. (6)
    15 end
    16 Get cumulative rewards  $R$  and behavior  $b(\pi_{\theta'_k})$  by evaluating the updated policy  $\pi_{\theta'_k}$ 
  17 end
  18 Update archive  $A$  and bandit  $B$ ;
  19  $t = t + T'$ 
20 end

```

Clustering-based selection

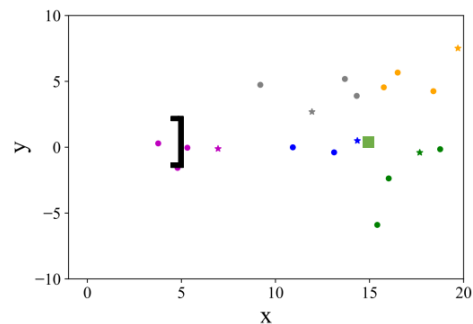
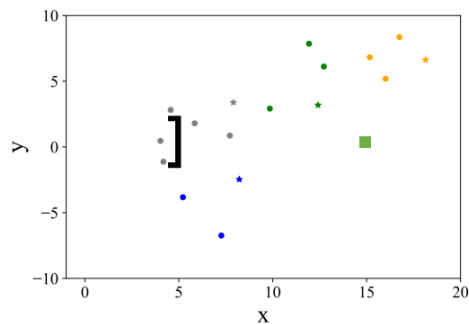
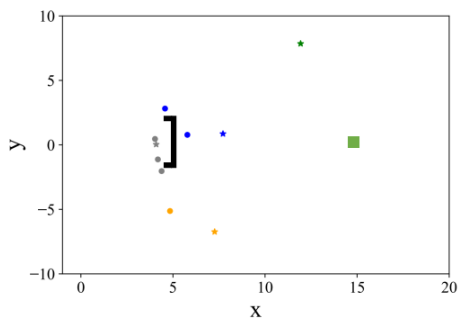
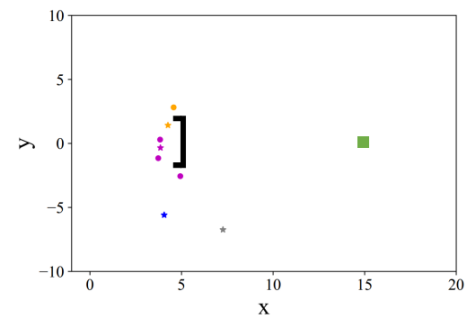
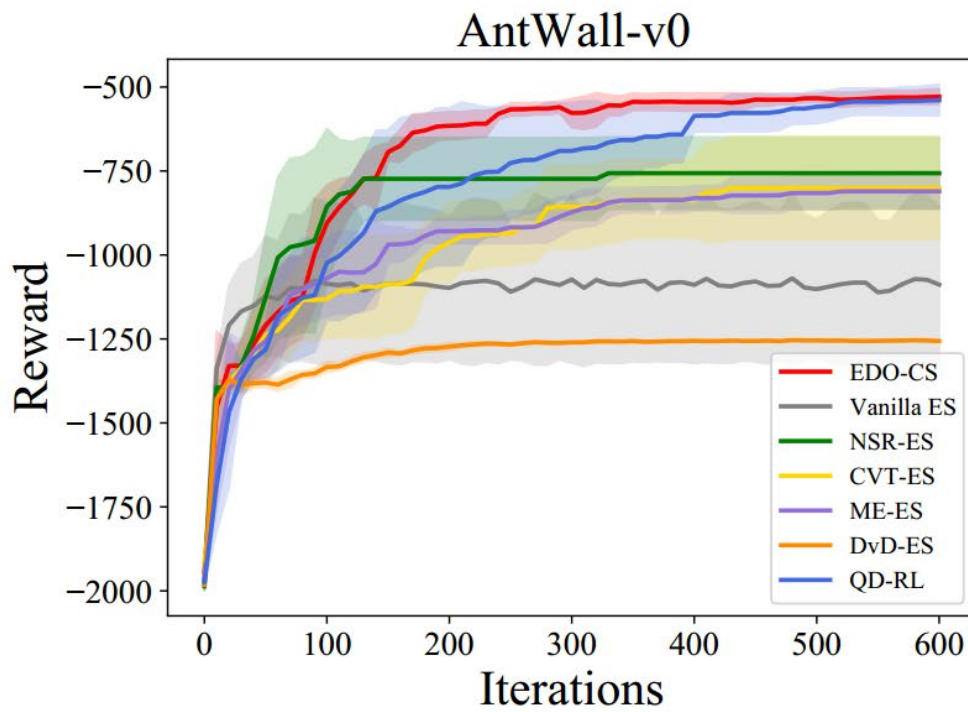
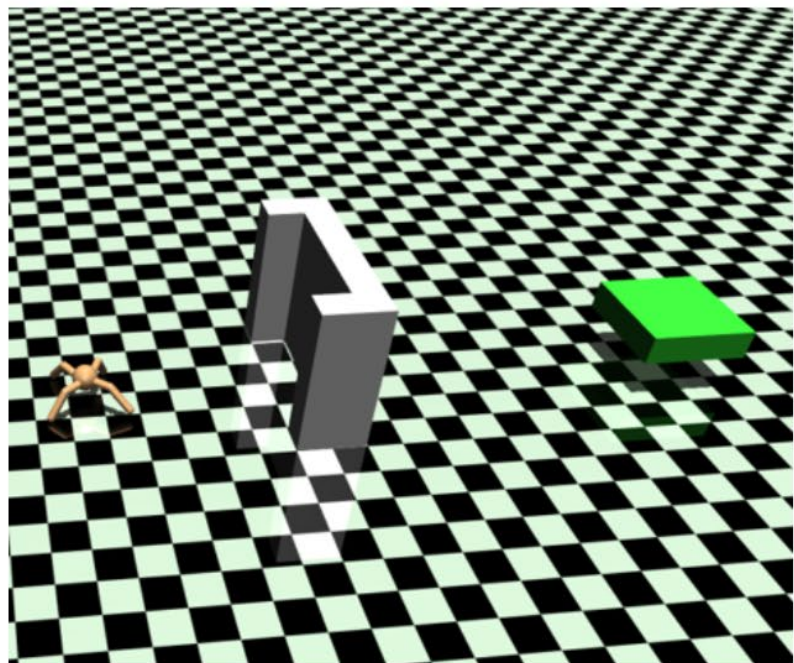
$$\text{Div}(\theta) = \frac{1}{k} \sum_{\pi'_{\theta} \in A'_{\pi_{\theta}, k}} \|b(\pi_{\theta}) - b(\pi'_{\theta})\|_2.$$

Method



Method	Selection	Reproduction	EAs type	From archive
Vanilla ES	The only parent solution	Quality	(1, 1)	×
NSR-ES	Probabilistic selection	Quality and diversity	(K , 1)	×
CVT-ES	Uniform selection	Quality and diversity	($K + K$)	✓
ME-ES	Biased selection	Quality or diversity	($K + K$)	✓
DvD-ES	All parent solutions	Quality and diversity	(K , K)	×
QD-RL	Pareto-based selection	Quality or diversity	($K + K$)	✓
EDO-CS	Clustering-based selection	Quality and diversity	($K + K$)	✓

Experiment

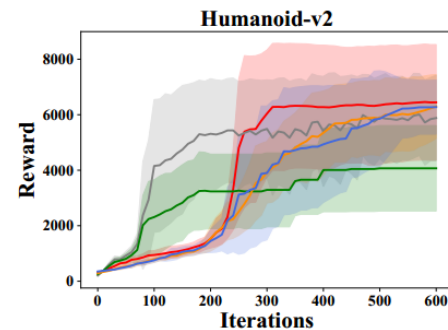
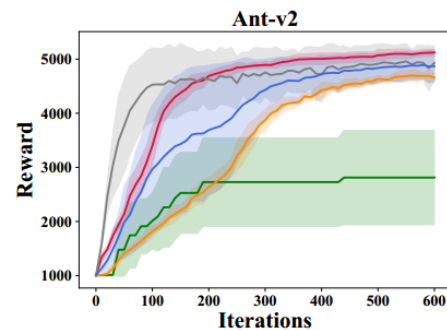
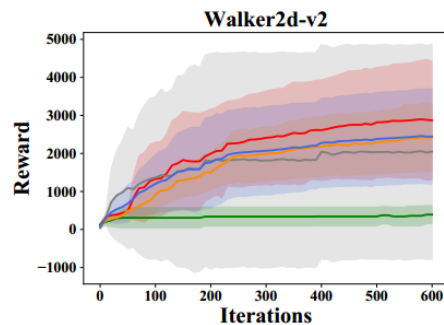
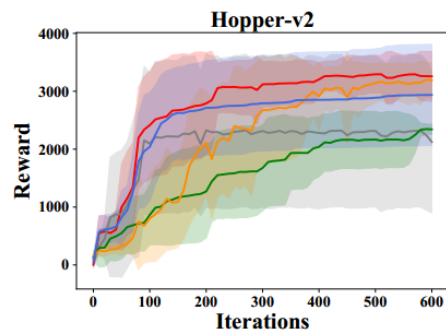


Experiment

- Multi-modal

Environment	EDO-CS	QD-RL	ME-ES	DvD-ES	CVT-ES	NSR-ES	Vanilla ES
<i>HalfCheetahFwd</i>	4284	2930	2700	-3419	3219	1346	-5543
<i>HalfCheetahBwd</i>	6548	6013	5953	6353	4672	5366	3911
<i>AntFwd</i>	4617	4291	4316	4507	3856	1737	1911
<i>AntBwd</i>	4697	4164	4123	3498	2958	3961	-851
Performance Ranking	1	3	3.5	3.75	4.75	5.25	6.75

- Single-modal



Experiment – ablation studies

- Adaptive λ

Setting of λ	<i>AntWall-v0</i>	<i>HalfCheetahFwd</i>	<i>HalfCheetahBwd</i>	<i>AntFwd</i>	<i>AntBwd</i>
Adaptive λ	-529	4284	6548	4617	4697
$\lambda = 0$	-650	3856	5931	4340	4426
$\lambda = 0.5$	-850	3877	6077	2800	3093

Experiment – ablation studies

- Adaptive λ
- Compared with direct optimization

$$f_{pair} = (1 - \omega) \sum_{i=1}^K \mathbb{E}_{\tau \sim \pi_{\theta_i}} [R(\tau)] + \omega \sum_{i=1}^K \sum_{j \neq i} \|b(\pi_{\theta_i}) - b(\pi_{\theta_j})\|_2$$

$$f_{det} = (1 - \omega) \sum_{i=1}^K \mathbb{E}_{\tau \sim \pi_{\theta_i}} [R(\tau)] + \omega \cdot \det(\Pi)$$

Method	<i>AntWall-v0</i>	<i>HalfCheetahFwd</i>	<i>HalfCheetahBwd</i>	<i>AntFwd</i>	<i>AntBwd</i>
EDO-CS	-529	4284	6548	4617	4697
EDO-DOS _{pair}	-706	4188	5847	5013	4392
EDO-DOS _{det}	-536	-5591	6529	-530	2417

Method	<i>AntWall-v0</i>	<i>HalfCheetahFwdBwd</i>	<i>AntFwdBwd</i>
EDO-CS	0.015	0.021	0.023
EDO-DOS _{pair}	9.570	9.802	9.807
EDO-DOS _{det}	8.900	9.266	9.567

Experiment – ablation studies

- Adaptive λ
- Compared with direct optimization
- Clustering algorithm
- Number T' of updating iterations
- Population size K
- Archive size l

More results are shown in the Appendix

Discussion

Conclusion

- QD is an interesting and attractive research area
- We proposed a “*Simple but efficient*” selection method for QD

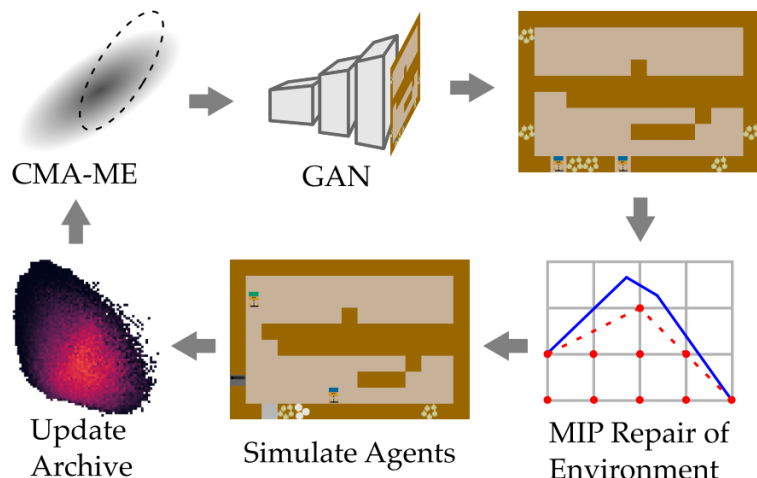
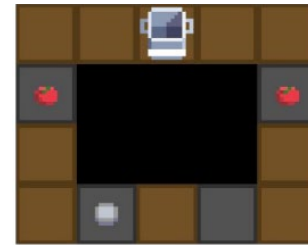
Future work

- Application of QD
 - Human-AI coordination
 - Environment generation

AI agent



Unseen human



Thanks you!

Our code can be found at:

www.lamda.nju.edu.cn/qianc/code-EDOCS.html

✉ xuek@lamda.nju.edu.cn