

Post Hoc Extraction of Pareto Fronts for Continuous Control

Raghav Thakar^{1,*}, Gaurav Dixit¹ and Kagan Tumer¹

¹The Collaborative Robotics and Intelligent Systems (CoRIS) Institute
Oregon State University, Corvallis, Oregon, USA

*thakarr@oregonstate.edu

Abstract

Agents in the real world must often balance multiple objectives, such as speed, stability, and energy efficiency in continuous control. To account for changing conditions and preferences, an agent must ideally learn a Pareto frontier of policies representing multiple optimal trade-offs. Recent advances in multi-policy multi-objective reinforcement learning (MORL) enable learning a Pareto front directly, but require full multi-objective consideration from the start of training. In practice, multi-objective preferences may arise after a policy has already been trained on a single specialised objective. Existing MORL methods cannot leverage such a pre-trained ‘specialist’ to learn Pareto fronts and avoid incurring the sample costs of retraining. We introduce Mixed Advantage Pareto Extraction (MAPEX), an offline MORL method that constructs a frontier of policies by reusing pre-trained specialist policies, critics, and replay buffers. MAPEX combines evaluations from specialist critics into a mixed advantage signal, and weights a behaviour cloning loss with it to train new policies that balance multiple objectives. MAPEX’s post hoc Pareto front extraction preserves the simplicity of single-objective off-policy RL, and avoids retrofitting these algorithms into complex MORL frameworks. We formally describe the MAPEX procedure and evaluate MAPEX on five multi-objective MuJoCo environments. Given the same starting policies, MAPEX produces comparable fronts at 0.001% the sample cost of established baselines.

1 Introduction

Many real-world continuous control problems require agents to balance multiple, even conflicting, objectives. In legged locomotion, for example, a robot must simultaneously optimise forward speed, gait stability, and energy efficiency. In such settings, there is no single optimal solution, and agents must instead discover a set of non-dominated policies that capture the range of feasible trade-offs. Learning this *Pareto front* enables downstream stakeholders to select behaviours

that reflect their current preferences or adapt to changing operational conditions.

Typically, multiple objectives are handled through scalarisation into a single weighted sum and training via standard reinforcement learning (RL). While accessible and tractable, this approach yields only a single, fixed trade-off. Simply repeating this procedure for every trade-off is inefficient and practically infeasible without sophisticated experience and representation sharing.

Multi-Objective RL (MORL), and specifically *multi-policy* MORL algorithms partially address this through various architectural improvements to learn the entire frontier directly. Methods like MORL/D [Felten *et al.*, 2024] divide the problem into several single-objective problems through scalarisation, and use single-objective RL with buffer sharing to train policies for each scalarisation. Other works like PG-MORL [Xu *et al.*, 2020] and MOPDERL [Tran *et al.*, 2023] maintain populations of policies, combining RL with evolutionary principles to jointly improve the policies in the population and cover the objective space. PSL-MORL [Liu *et al.*, 2025], a notable recent work, trains a hypernetwork to produce a fresh policy for every desired trade-off.

Despite successfully learning Pareto fronts, these methods suffer from a critical limitation: requiring full multi-objective consideration from the outset of training. This rigidity creates a disconnect with practical scenarios, where multi-objective preferences may arise retroactively, only after a robust policy for a primary task has already been trained. For instance, a preference for more stability may only emerge after observing a locomotion policy highly optimised for speed. To obtain new trade-offs in these scenarios, practitioners must 1) discard pre-trained policies and incur the sample costs of retraining, and 2) retrofit their algorithm into complex multi-objective learning frameworks. Currently, no method reuses disjoint specialist policies and training data to recover Pareto fronts efficiently and with minimal added algorithmic complexity.

In this work we present Mixed Advantage Pareto Extraction¹ (MAPEX), a novel method that fully leverages pre-trained single-objective policies, critics, and replay buffers to produce Pareto fronts of policies. Our key insight is that agents learn optimal trade-offs by intelligently blending ex-

¹Project homepage: <https://github.com/raghavthakar/MAPEX>

pert behaviour on each objective. MAPEX implements this by blending the single-objective evaluation from each specialist critic into a multi-objective *mixed advantage* value and weighting a behaviour cloning loss with it.

MAPEX bypasses the need to retrofit bespoke or standard off-policy RL into complex multi-objective learning frameworks by training for each objective individually. It preserves both, the simplicity, and intricacies of these algorithms by providing a realistic pathway to learning multi-objective behaviours from single-objective training data. If training from scratch, MAPEX allows allocating the entire sample budget to training high-performing single-objective specialists, from which high-quality Pareto fronts can later be extracted at a minimal sample cost.

We provide a detailed view of MAPEX’s Pareto extraction procedure, and perform Pareto extraction from specialists trained using both standard and bespoke off-policy RL. Across five multi-objective MuJoCo environments, MAPEX produces fronts comparable to established baselines from the literature. When extracting Pareto fronts from the same starting policies, other methods consume up to $1000\times$ more samples than MAPEX to produce similar fronts.

2 Background

2.1 Actor–Critic Methods and Offline Reinforcement Learning

Actor–critic methods: These methods perform reinforcement learning using an *actor*, which defines the policy, and a *critic*, which estimates the expected return of the actor’s behaviour [Sutton and Barto, 2018]. By relying on learned value estimates rather than Monte Carlo returns alone, actor–critic methods enable efficient and stable policy optimisation, and are commonly used in continuous control. These methods are often implemented in an *off-policy* setting, where experience collected by one or more behaviour policies is stored in a replay buffer and reused for policy updates [Haarnoja *et al.*, 2018; Fujimoto *et al.*, 2018].

Policy improvement is commonly guided by the *advantage* function,

$$A^\pi(s, a) = Q^\pi(s, a) - V^\pi(s), \quad (1)$$

where $Q^\pi(s, a)$ is the critic’s action-value function, and $V^\pi(s) = \mathbb{E}_{a \sim \pi(\cdot|s)}[Q^\pi(s, a)]$. In practice, particularly in deterministic actor–critic methods, the value function is often approximated by $Q^\pi(s, \pi(s))$, yielding

$$A^\pi(s, a) \approx Q^\pi(s, a) - Q^\pi(s, \pi(s)). \quad (2)$$

Offline Reinforcement Learning: Offline reinforcement learning studies the problem of learning a policy from a fixed dataset of transitions, without any further interaction with the environment during training. A key challenge in this setting is distributional shift: the dataset may contain diverse experiences that are hard to generalise to, leading to unreliable value estimates [Panaganti *et al.*, 2025; Levine *et al.*, 2020].

Advantage-weighted regression (AWR) [Peng *et al.*, 2019] provides a simple and stable mechanism for policy improvement in this setting by casting reinforcement learning as a

weighted supervised learning problem. Given a critic-derived advantage estimate $A(s, a)$, the AWR policy update solves

$$\pi^+ = \arg \max_{\pi} \mathbb{E}_{(s,a) \sim \mathcal{D}} \left[\log \pi(a|s) \exp \left(\frac{1}{\beta} A(s, a) \right) \right], \quad (3)$$

where $\beta > 0$ is a temperature parameter. This objective emphasises actions that are predicted to improve performance while sampled purely from the observed data, making it well suited to off-policy and offline reinforcement learning problems. We use an AWR-inspired regression weighting in MAPEX to learn multi-objective behaviours from static single-objective replay buffers.

2.2 Multi-Objective Sequential Decision-Making

Multi-objective sequential decision-making problems are commonly modelled as Multi-Objective Markov Decision Processes (MOMDPs). A MOMDP is defined by the tuple $\langle \mathcal{S}, \mathcal{A}, \mathcal{T}, \gamma, \mathbf{R} \rangle$, containing the state space $\mathcal{S}$, the action space $\mathcal{A}$, the transition dynamics $\mathcal{T} : \mathcal{S} \times \mathcal{A} \times \mathcal{S} \rightarrow [0, 1]$, and a scalar discount factor $\gamma \in [0, 1]$. Unlike standard MDPs, the reward function $\mathbf{R} : \mathcal{S} \times \mathcal{A} \times \mathcal{S} \rightarrow \mathbb{R}^k$ returns a vector of k rewards, each corresponding to a distinct objective.

A policy $\pi : \mathcal{S} \rightarrow \mathcal{A}$ induces an expected return for each objective. We characterise policy performance directly through its expected long-term returns. Let $\mathbf{J}(\pi) \in \mathbb{R}^k$ denote the vector of objective returns achieved by policy π , with components

$$J_i(\pi) = \mathbb{E}_{s_0 \sim \rho_0, a_t \sim \pi(\cdot|s_t)} \left[\sum_{t=0}^{\infty} \gamma^t r_i(s_t, a_t) \right], \quad i \in \{1, \dots, k\}. \quad (4)$$

where r_i is the reward function associated with the i -th objective and ρ_0 is the initial state distribution.

From a multi-objective optimisation perspective, learning in a MOMDP can be viewed as maximising a vector-valued objective

$$\max_{\pi} \mathbf{F}(\pi) \triangleq \max_{\pi} [J_1(\pi), J_2(\pi), \dots, J_k(\pi)]. \quad (5)$$

Pareto Dominance and Optimality As objectives are often conflicting, it is generally impossible for a single policy to optimise all objectives simultaneously. Thus, optimality in multi-objective settings is defined in terms of *Pareto dominance*. Given two objective return vectors $\mathbf{v}, \mathbf{u} \in \mathbb{R}^k$, $\mathbf{v}$ *Pareto-dominates* $\mathbf{u}$ (denoted $\mathbf{v} \succ_p \mathbf{u}$) if $\mathbf{v}$ is at least as good as $\mathbf{u}$ in all objectives and strictly better in at least one:

$$\mathbf{v} \succ_p \mathbf{u} \iff (\forall i : v_i \geq u_i) \wedge (\exists j : v_j > u_j). \quad (6)$$

If neither vector Pareto-dominates the other, they are said to be *non-dominated*.

A policy π is *Pareto optimal* if its return vector $\mathbf{J}(\pi)$ is not Pareto-dominated by that of any other policy. The set of all Pareto-optimal policies induces the *Pareto front*, which characterises the optimal trade-offs achievable in the objective space.

3 Related Works

Foundational surveys [Hayes *et al.*, 2022; Roijers *et al.*, 2013] categorise the expanding MORL landscape into single-policy and multi-policy approaches. We adopt this taxonomy herein.

Single-policy methods optimise a fixed scalar utility function [Hayes *et al.*, 2022], using standard RL with linear utilities, and dedicated approaches for non-linear utilities [Lin *et al.*, 2024b; Van Moffaert *et al.*, 2013; Reymond *et al.*, 2023]. Conversely, multi-policy approaches approximate the Pareto front, a better-suited approach for unknown or dynamic preferences. Early works like Pareto Q-Learning [Moffaert and Nowé, 2014] tracked non-dominated return vectors for each state-action pair, while recent extensions employ preference-conditioned deep networks [Abels *et al.*, 2018] or convex envelope updates [Yang *et al.*, 2019] for high-dimensional spaces.

In continuous control, decomposition-based methods divide the problem into several scalarised objectives. Prior work [Chen *et al.*, 2020] has augmented naively solving each scalar objective with an evolutionary strategy [Salimans *et al.*, 2017] for post-processing. Notably, MORL/D [Felten *et al.*, 2024] improves efficiency via cooperative buffer sharing and intelligent weight sampling. Another branch of work combines RL with evolutionary operators: PG-MORL [Xu *et al.*, 2020] iteratively pushes a population of policies toward promising objective space regions, while MOPDERL [Tran *et al.*, 2023] distills a frontier from subpopulations trained via evolutionary RL [Bodnar *et al.*, 2020; Khadka and Tumer, 2018]. Alternatively, parameter-efficient methods consolidate policies using meta-learning [Chen *et al.*, 2019], universal preference-conditioned networks [Basaklar *et al.*, 2023], or hypernetworks [Liu *et al.*, 2025].

Despite architectural differences, these methods share a structural limitation: they are designed solely for “from scratch” learning. If training on one or more individual objectives has already been performed, these methods cannot leverage it, effectively requiring pre-trained policies to be discarded and training to be restarted. The structural rigidity of these approaches also makes adapting bespoke RL algorithms into these frameworks non-trivial.

A separate line of inquiry focuses on offline MORL, training preference-conditioned agents from massive datasets. Approaches like PEDAs [Zhu *et al.*, 2023], DiffMORL [Xiao *et al.*, 2025], and PR-MORL [Lin *et al.*, 2024a] utilise transformers, diffusion, and regularisation to generalise across preferences. While seemingly similar, MAPEX tackles the distinct problem of extracting frontiers from pre-trained specialists rather than large-scale generalisation. Consequently, MAPEX operates with replay buffers two orders of magnitude smaller than the D4MORL benchmarks [Zhu *et al.*, 2023] used in offline RL works (e.g., 1 million vs. 150 million transitions).

Finally, PCN [Reymond *et al.*, 2022] also uses supervised learning from an off-policy dataset, but relies on iterative on-line data collection to populate the objective space. In contrast, MAPEX addresses the strictly offline extraction of frontiers from fixed, disjoint datasets. PCN is also currently limited to discrete action spaces.

4 Method

We now introduce Mixed Advantage Pareto Extraction (MAPEX), an offline algorithm to extract a Pareto front of

Algorithm 1: Mixed Advantage Pareto Extraction (MAPEX)

INPUT: Policy set Π , critic set $\mathcal{Q}$, buffer set $\mathcal{D}$, number of objectives N

OUTPUT: Pareto front $\mathcal{P}$ of policies

```

1 FUNCTION MAPEX( $\Pi, \mathcal{Q}, \mathcal{D}, N$ ):
2   while Required do
3     EVALUATE( $\Pi$ )
4      $\mathcal{P} \leftarrow \text{FINDPARETOFRONT}(\Pi)$ 
5      $\{\pi_1, \dots, \pi_N\} \leftarrow \text{SELECTPARENTS}(\mathcal{P})$ 
6      $\mathbf{v} \leftarrow \text{CENTROID}(\pi_1, \dots, \pi_N)$ 
       // In objective space
7      $\mathbf{w}_{\text{target}} \leftarrow \frac{\mathbf{v}}{\|\mathbf{v}\|_2}$  // Target weight vector
8      $\mathcal{D}_{\text{hybrid}} \leftarrow \bigcup_{k=1}^N \text{SAMPLE}(\mathcal{D}_k, \propto w_k)$ 
9      $\pi_{\text{new}} \leftarrow \text{INITPOLICY}(\emptyset)$ 
10    for  $\epsilon \leftarrow 1$  to  $E$  do
11       $(s, a) \sim \mathcal{D}_{\text{hybrid}}$ 
12       $\mathbf{A} \leftarrow \left[ \left( \mathcal{Q}_i(s, a) - \mathcal{Q}_i(s, \pi_{\text{new}}(s)) \right) \right]_{i=1}^N$ 
13       $A_{\text{mixed}} \leftarrow \mathbf{w}_{\text{target}}^\top \mathbf{A}$  // Mixed advantage
14       $\omega(s, a) = \min \left( \exp \left( \frac{A_{\text{mixed}}(s, a)}{\beta} \right), \omega_{\text{max}} \right)$ 
15       $\mathcal{L}_{\text{MAPEX}} \leftarrow \mathbb{E} \left[ \omega(s, a) \cdot \|a - \pi_{\text{new}}(s)\|_2^2 \right]$ 
16       $\text{UPDATE}(\pi_{\text{new}}, \mathcal{L}_{\text{MAPEX}})$ 
17       $\Pi \leftarrow \Pi \cup \{\pi_{\text{new}}\}$ 
18   $\mathcal{P} \leftarrow \text{FINDPARETOFRONT}(\Pi)$ 
19  return  $\mathcal{P}$ 

```

continuous control policies from a set of disjoint *single-objective specialists*. The core intent of MAPEX is to fully leverage the latent information in prior-trained specialists—which includes their policies, along with their value functions (critics) and replay buffers—to produce new policies that express new trade-off behaviours without requiring additional training interaction with the environment.

MAPEX achieves this via a Pareto front extraction procedure which iteratively fills gaps in the Pareto front estimate. First, MAPEX analyses the available policies in the objective space to identify sparse regions in the Pareto front estimate. Once a gap is identified, MAPEX derives a vector of ‘target weights’ that encodes a weighting over the objectives that would optimally fill this gap. To train a policy for this new trade-off, MAPEX constructs a static *hybrid buffer* by sampling from the specialists’ buffers in proportion to these target weights. Crucially, the algorithm then calculates a *mixed advantage* for every transition in this dataset by querying each specialist critic, and mixing the individual advantage estimates in the ratio of the target weights. This mixed

advantage captures the value of a transition in demonstrating target trade-off behaviour. Finally, a fresh policy is trained using a mechanism inspired by Advantage Weighted Regression (AWR) [Peng *et al.*, 2019], wherein we regress the policy onto actions from the hybrid buffer, weighted by an exponential of their calculated mixed advantage. Algorithm 1 provides a more rigorous look at this procedure.

Starting Information and Notation

We assume that from prior training we have a policy set Π with at least N policies for an N -objective problem, a critic set $\mathcal{Q}$ with N critics, each specialising on evaluating on one of the problem’s objectives, and a replay buffer set $\mathcal{D}$ with N buffers, each containing experiences from training on a single objective.

4.1 Step 1: Gap Identification and Parent Selection

At the start of each iteration, we evaluate the current policy set Π to obtain the performance vector $J(\pi) \in \mathbb{R}^N$ of each policy. We then identify the non-dominated set (the current Pareto front approximation). We then search for the largest N -dimensional ‘gap’ on the frontier—a sparse region in the objective space. For $N = 2$, we define this as the edge with the maximal Euclidean span. In practice, to avoid repeatedly focusing on the same gap, we select the gap using roulette-wheel sampling based on edge length. The N policies corresponding to the vertices of this gap are selected as the *parent policies*, denoted by $\{\pi_{p_1}, \dots, \pi_{p_N}\}$.

To guide the offspring into this sparse region, we compute the centroid of the parents’ performance vectors in the objective space:

$$J_{mid} = \frac{1}{N} \sum_{i=1}^N J(\pi_{p_i}) \quad (7)$$

We then derive a unit vector $\mathbf{w}_{target}$ of target weights pointing towards this centroid. This vector encodes the linear preference required to interpolate the gap and guides the subsequent hybrid buffer creation and advantage mixing steps.

4.2 Step 2: Hybrid Buffer Creation and Advantage Mixing

After deriving the target weight vector $\mathbf{w}_{target}$, we assemble a training distribution that reflects this desired trade-off. We construct a fixed-size *hybrid buffer*, $\mathcal{D}_{hybrid}$, by sampling transitions from each specialist’s buffer $\mathcal{D}_k$ in direct proportion to the corresponding weight $\mathbf{w}_{target,k}$. This creates a dataset that is structurally biased to each objective in the desired proportion.

We then initialise a random policy network π_{new} that will be optimised to achieve the target trade-off behaviour. For this optimisation, we iteratively sample transitions from $\mathcal{D}_{hybrid}$ to compute the mixed advantage training signal. For each transition (s, a) , we compute a vector of advantages $\mathbf{A}(s, a) \in \mathbb{R}^N$. The k^{th} element of this advantage vector is the advantage on the k^{th} objective associated with that transition, and is computed using the k^{th} specialist critic $\mathcal{Q}_k$:

$$A_k(s, a) = \mathcal{Q}_k(s, a) - \mathcal{Q}_k(s, \pi_{new}(s)) \quad (8)$$

This formulation leverages the specific value estimation expertise of each critic for their respective objective.

Finally, we scalarise these vector-valued advantages into a single training signal. We compute the *mixed advantage* as the dot product of the advantage vector and the target weights derived in Step 1:

$$A_{mixed}(s, a) = \mathbf{w}_{target}^\top \cdot \mathbf{A}(s, a) \quad (9)$$

This scalar value A_{mixed} represents the quality of the state-action (s, a) specifically regarding the desired trade-off $\mathbf{w}_{target}$.

4.3 Step 3: Mixed Advantage Weighted Regression

To train the offspring policy π_{new} , we employ a supervised regression objective weighted by the scalarised signal derived in Equation 9. Our goal is to selectively clone actions that contribute positively to the specific target trade-off $\mathbf{w}_{target}$.

For a transition (s, a) we compute a regression weight $\omega(s, a)$ by applying a temperature-scaled exponential to its mixed advantage:

$$\omega(s, a) = \min \left(\exp \left(\frac{A_{mixed}(s, a)}{\beta} \right), \omega_{max} \right) \quad (10)$$

where $\beta > 0$ is a temperature hyperparameter and ω_{max} is a clipping threshold to ensure numerical stability.

The fresh policy π_{new} is then optimised to minimise the weighted mean squared error between its predicted action and the retrieved buffer action a :

$$\mathcal{L}_{MAPEX} = \mathbb{E}_{(s,a)} \left[\omega(s, a) \cdot \|\pi_{new}(s) - a\|_2^2 \right] \quad (11)$$

Once the policy has been updated for the desired epochs, it is reinserted into the population Π and the MAPEX procedure is executed for another iteration.

4.4 Mitigating Out-of-Distribution Error

While MAPEX’s mixed advantage values are intuitive to compute, they expose two main sources of out-of-distribution (OOD) error: 1) when a transition sampled by one objective’s specialist policy is evaluated by another objective’s specialist critic, and 2) when a randomly-initialised policy’s action is evaluated by specialist critics. We mitigate these issues using the following design choices.

Secondary Critics

When training each specialist policy π_k , we learn not only its *primary* critic for objective k , but also a set of *secondary critics* for the remaining objectives. All critics associated with specialist k are trained on the same replay buffer $\mathcal{D}_k$ generated by π_k , so that every objective can be evaluated on data that is in-distribution for the corresponding critic.

Notation: Let $m \in \{1, \dots, N\}$ index objectives and $k \in \{1, \dots, N\}$ index specialists. We denote by $Q_m^{(k)}(s, a)$ the action-value critic that predicts the return for objective m using transitions sampled from $\mathcal{D}_k$ (i.e., collected under π_k). Under this notation, the specialist’s *primary critic* is $Q_k^{(k)}$, while the *secondary critics* are $\{Q_m^{(k)}\}_{m \neq k}$.

Procedure: During specialist training, only the primary critic $Q_k^{(k)}$ is used to update the policy π_k . The secondary critics $\{Q_m^{(k)}\}_{m \neq k}$ are trained in parallel on the same buffer $\mathcal{D}_k$, but they do not contribute gradients to the policy update. After training all specialists, we collect the resulting critics into a single family $\mathcal{Q} \triangleq \{Q_m^{(k)}\}_{k=1 \dots N, m=1 \dots N}$. This construction ensures that when MAPEX samples a transition (s, a) that originates from a particular specialist buffer $\mathcal{D}_k$, the evaluations $\{Q_m^{(k)}(s, a)\}_{m=1}^N$ are produced by critics trained on the same state–action distribution as the sampled data.

Practical note: Secondary critics may be trained alongside primary critics with no change to the policy update rule. If single-objective specialists have already been trained, a secondary critic $\{Q_m^{(k)}\}_{m \neq k}$ can be trained offline using transitions from $\mathcal{D}_k$ and the corresponding objective rewards. Either way, a distribution-matched critic family $\mathcal{Q}$ can be learnt for a low cost. We include a comparison of joint-vs. post-hoc-trained secondary critics in Section 5.

Offspring Policy Warm-Up

In step 2 of the MAPEX procedure (Equation 8), computing the mixed advantage requires querying specialist critics with actions proposed by $\pi_{\text{new}}(s)$. If π_{new} is initialised arbitrarily, these actions can be far from the support of the hybrid buffer and cause OOD error. To reduce this effect, we warm up π_{new} by regressing it to the mean of its parents in the action space.

Let $\{\pi_{p_1}, \dots, \pi_{p_N}\}$ denote the parent policies. The mean parent action at state s is $\bar{a}(s) \triangleq \frac{1}{N} \sum_{j=1}^N \pi_{p_j}(s)$, which we use to perform a brief behavioural regression step that minimises

$$L_{\text{init}}(\theta) \triangleq \mathbb{E}_{s \sim \mathcal{D}_{\text{hybrid}}} \left[\|\pi_{\text{new}}(s) - \bar{a}(s)\|_2^2 \right], \quad (12)$$

We run this regression for a small number of gradient steps prior to computing the advantage vector in Equation 8. This warm up keeps π_{new} ’s predicted actions close to those produced by the parents on states drawn from $\mathcal{D}_{\text{hybrid}}$, more closely matching each critic’s training exposure, and improving the reliability of subsequent critic-based updates.

5 Experiments

We evaluate MAPEX on three criteria: **sample efficiency** of Pareto extraction; **flexibility** regarding choice of specialist training algorithm and nature of secondary critic training (joint vs. post-hoc); and **general competitiveness** against baselines when training from scratch.

We assume access to secondary critics trained alongside primary critics during specialist training, but also test a *post hoc* variant (MAPEX-PostHoc) where critics are trained offline on static buffers. Unless stated otherwise, specialists are trained using Proximal Distilled Evolutionary RL (PDERL) [Bodnar *et al.*, 2020].

5.1 Baselines and Domains

We compare MAPEX against two established MORL methods from the literature:

- **MOPDERL** [Tran *et al.*, 2023]: An evolutionary actor-critic method that first trains on each individual objective using PDERL, and later crosses over solutions across objectives using a multi-objective distilled crossover. It serves as a baseline for both ‘from scratch’ training and pure Pareto extraction (via its distillation phase).
- **MORL/D** [Felten *et al.*, 2024]: A decomposition-based approach that uses Soft Actor–Critic [Haarnoja *et al.*, 2018] on each scalar objective. It uses buffer data sharing and adapts the scalar weights using Pareto Simulated Annealing [Czyżżak and Jaskiewicz, 1998].

We use the `morl-baselines` [Felten *et al.*, 2023] implementation of MORL/D and the authors’ original implementation of MOPDERL² with default hyperparameters. Experiments are conducted on five bi-objective continuous control MuJoCo environments from MO-Gymnasium [Felten *et al.*, 2023], with episodes capped at 750 frames³.

5.2 Experimental Setup

Sample efficiency: We compare the extraction phase of MAPEX against the distillation phase of MOPDERL starting from identical specialist subpopulations trained with PDERL.

Flexibility: We compare standard MAPEX against MAPEX-PostHoc and MAPEX-TD3 (specialists trained via TD3) to assess robustness to starting policies and nature of critic training. **Competitiveness:** We compare the final Pareto fronts of the full MAPEX pipeline against MOPDERL and MORL/D given a fixed sample budget. During specialist training of all MAPEX variants, we use a replay buffer of size 1M. In each test we perform five seeded runs with each method³.

6 Results

6.1 Sample Efficiency of Extraction

MAPEX achieves a massive reduction in sample cost compared to baselines. As shown in Figure 1 (MO-Ant-v5), MAPEX and MAPEX-PostHoc extract a high-performing front almost immediately (requiring only evaluation rollouts), while MOPDERL requires an additional 300,000 environment interactions to attain the same performance.

Figure 2 confirms this trend across MO-Hopper-v5 and MO-Walker2d-v5. Particularly in MO-Hopper-v5, MAPEX requires 100 samples to reach hypervolume thresholds that MOPDERL requires $\approx 10^5$ samples to attain—a reduction of three orders of magnitude. This efficiency is driven by MAPEX exploiting the latent representations of multi-objective behaviour in policies and replay buffers. It empirically validates our intuition that following expert behaviour to varying degrees on each objective yields trade-offs in the objective space. It also validates our main hypothesis that these behaviours can be learnt by regressing onto target actions, weighed by their *mixed advantage* value.

²We made minor modifications to the original MOPDERL code to allow testing it on v5 MO-Gymnasium MuJoCo domains. Code: <https://github.com/raghavthakar/MOPDERL-MO-Gymnasium>

³Exact details of the environment, with algorithm-specific and experimental parameters are available on the project homepage: <https://github.com/raghavthakar/MAPEX>

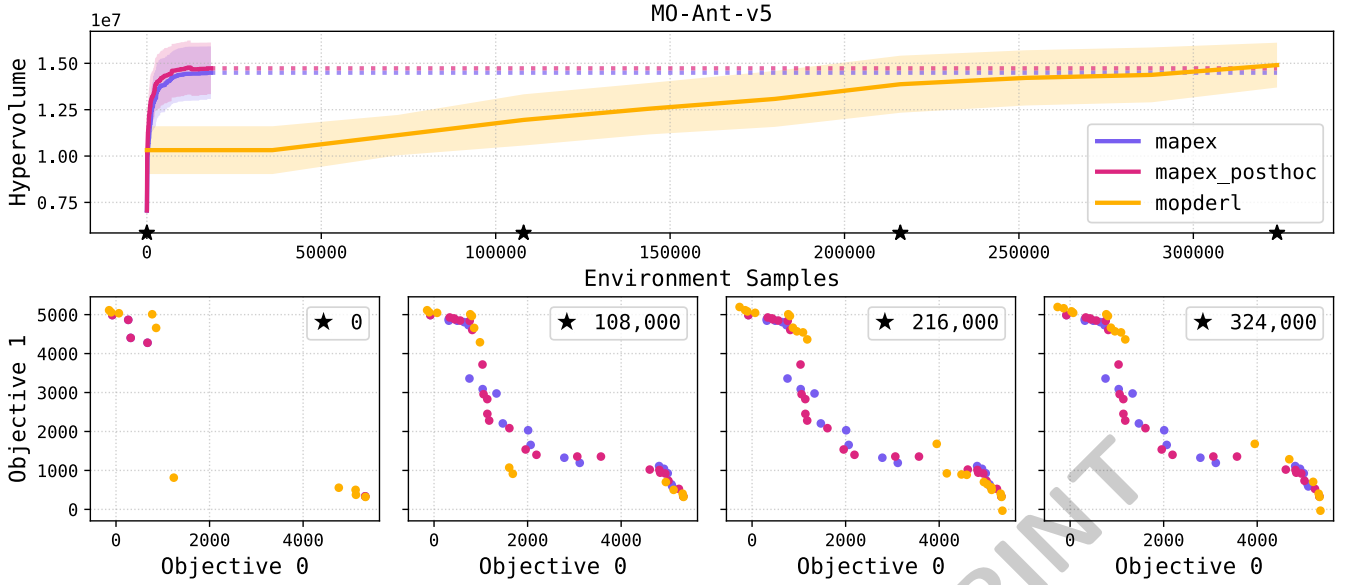


Figure 1: **Sample efficiency comparison on MO-Ant-v5.** (Top) Mean hypervolume $\pm$ SEM vs. cumulative environment samples. MAPEX and MAPEX-PostHoc achieve high hypervolume almost instantaneously, while MOPDERL requires significantly more interaction. (Bottom) Evolution of the Pareto front approximation. MAPEX/MAPEX-PostHoc fill the front immediately, whereas MOPDERL gradually expands coverage over 300,000+ environment samples.

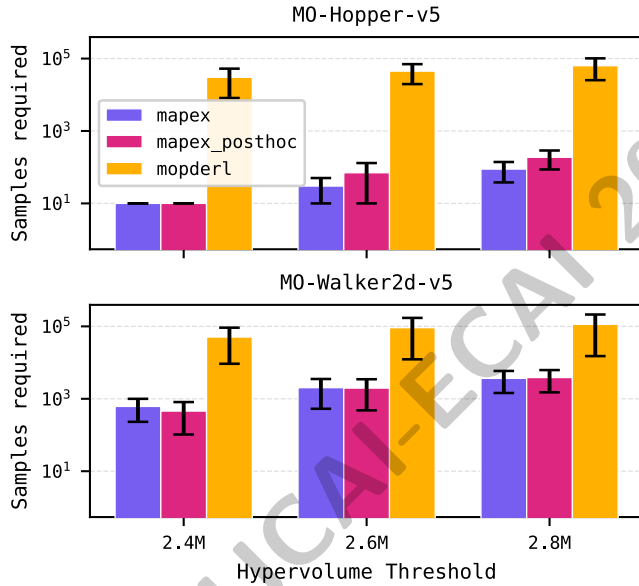


Figure 2: **Samples required to attain target hypervolume thresholds.** Comparison on MO-Hopper-v5 (top) and MO-Walker2d-v5 (bottom). Note the logarithmic scale on the y-axis. MAPEX and MAPEX-PostHoc require up to three orders of magnitude fewer samples (10^2 vs 10^5) than MOPDERL to reach identical performance levels.

6.2 Flexibility and Robustness

MAPEX is robust to the source of specialist policies. Figure 3 shows that Pareto extraction for standard MAPEX, MAPEX-PostHoc, and MAPEX-TD3 are largely indistinguishable on MO-Walker2d-v5 and MO-Ant-v5. Crucially, the success

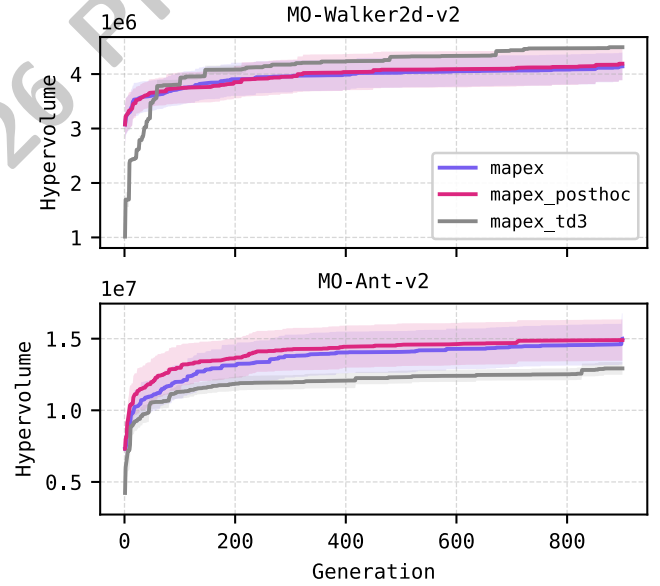


Figure 3: **Robustness of MAPEX to specialist type and critic training.** Mean hypervolume ($\pm$ SEM) over generations on MO-Walker2d-v5 and MO-Ant-v5. The similar performance of standard MAPEX, MAPEX-PostHoc (offline critics), and MAPEX-TD3 (off-policy specialists) demonstrates the method’s flexibility in effectively extracting fronts from decoupled pre-trained sources.

of MAPEX-PostHoc confirms that secondary critics can be effectively trained retroactively on static buffers, enabling Pareto extraction from fully decoupled single-objective training. If simplicity is key, then the performance of MAPEX-TD3 shows the potential of elegantly adapting off-policy RL

Environment	Hypervolume ($\uparrow$)			Sparsity ($\downarrow$)		
	MAPEX	MOPDERL	MORL/D	MAPEX	MOPDERL	MORL/D
Ant-2obj	$1.19\text{e}7 \pm 1.2\text{e}6$	$1.46\text{e}7 \pm 8.2\text{e}5$	$1.41\text{e}7 \pm 1.9\text{e}4$	315.5 ± 36.1	201.4 ± 13.8	289.8 ± 8.3
Hopper-2obj	$3.34\text{e}6 \pm 2.9\text{e}5$	$3.17\text{e}6 \pm 2.8\text{e}5$	$4.31\text{e}6 \pm 1.7\text{e}5$	209.6 ± 39.8	64.2 ± 10.3	115.5 ± 17.7
Swimmer	$1.78\text{e}5 \pm 1.6\text{e}4$	$2.41\text{e}5 \pm 2.9\text{e}4$	$9.29\text{e}4 \pm 5.9\text{e}2$	49.5 ± 9.2	30.0 ± 3.2	2.9 ± 0.3
Walker2d	$3.09\text{e}6 \pm 2.2\text{e}5$	$3.47\text{e}6 \pm 3.5\text{e}5$	$6.52\text{e}6 \pm 4.2\text{e}5$	157.5 ± 7.9	103.2 ± 16.0	96.9 ± 11.2
HalfCheetah	$1.10\text{e}7 \pm 5.3\text{e}5$	$1.55\text{e}7 \pm 6.8\text{e}5$	$1.88\text{e}7 \pm 5.4\text{e}4$	337.5 ± 26.7	101.7 ± 5.9	136.4 ± 19.6

Table 1: Performance metrics (Hypervolume and Sparsity) across v5 MO-Gymnasium MuJoCo environments. Results are Mean $\pm$ SEM. Hypervolume is the space between the Pareto front and a dominated reference point and sparsity is the average euclidean distance between neighbouring points.

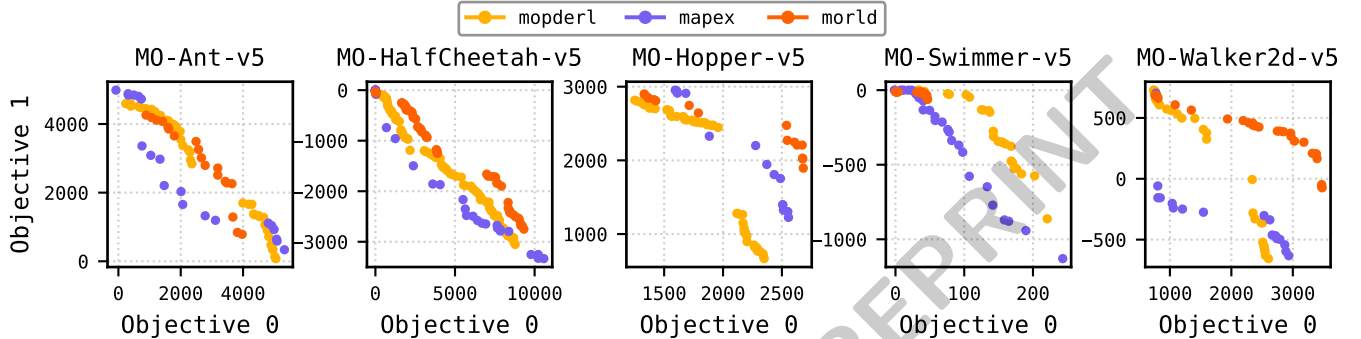


Figure 4: **Final Pareto fronts across five MO-MuJoCo benchmarks.** Comparison of fronts extracted by MAPEX against fully trained MOPDERL and MORL/D baselines. Despite relying purely on single-objective training data, MAPEX recovers fronts that are dense and competitive with baselines trained from scratch.

to learn multi-objective behaviours via MAPEX.

While a population of policies (like that produced by PDERL) provides a rich starting point for Pareto extraction, MAPEX-TD3 remains competitive with only one specialist policy per objective (as produced by simple TD3). This is because MAPEX does not simply distil or interpolate between parent policies. Instead, MAPEX draws expertise from the pre-trained replay buffers, which contain a rich and varied set of state-action examples. These experiences span regions of the objective space that an individual policy may only cover sparsely, allowing MAPEX to produce distinct policies for distinct trade-offs. This is also why MAPEX is robust to the exact algorithm used to train specialist policies.

6.3 General Competitiveness

Despite being an offline extraction method, MAPEX produces fronts competitive with MOPDERL and MORL/D, which require full multi-objective consideration from the beginning. Table 1 shows that MAPEX achieves comparable hypervolumes (e.g., 3.34×10^6 vs. MOPDERL’s 3.17×10^6 in MO-Hopper-v5, and 1.78×10^5 vs. MORL/D’s 9.29×10^4 in MO-Swimmer-v5) across five multi-objective MuJoCo environments. While MAPEX fronts are numerically sparser, by inspecting the fronts in Figure 4 visually, it is clear that MAPEX produces even and well-spread Pareto fronts.

7 Conclusion

We presented MAPEX, a novel approach to extracting Pareto fronts of policies from prior single-objective training for con-

tinuous control. We provided a detailed view of MAPEX’s Pareto extraction procedure, and mentioned some practical implementation tips. We tested MAPEX with well-established, dedicated MORL baselines like MOPDERL and MORL/D to empirically validate its mixed advantage approach. Pareto fronts are learnt cheaply with MAPEX when specialists are already trained. If training from scratch, MAPEX integrates easily with off-policy RL methods and still produces fronts that compare well to baselines. Finally, we discuss some limitations and future work.

While MAPEX is highly sample efficient, it makes assumptions inherent to our offline extraction setting. First, MAPEX is strictly bounded by the support of the specialist buffers; it cannot discover novel behaviours or skills that are absent from the specialists’ training history. Second, MAPEX relies on the assumption that valid trade-off policies lie on a continuous manifold between specialists. In scenarios where specialists exhibit markedly disjoint behaviours (e.g., a humanoid walking on two legs vs. crawling), interpolation may yield low-performance policies. Third, when a new objective arises, rewards on that objective may not exist in the collected data, necessitating post-hoc reward reconstruction over the replay buffer using logged telemetry data or fresh rollouts.

Our empirical evaluation focused on bi-objective domains; scaling MAPEX’s gap-identification heuristic to more objectives ($N \geq 3$) remains a subject for future work. We also aim to leverage MAPEX in the multi-objective multiagent setting, integrating with reinforcement learning [Foerster *et al.*, 2018] and credit assignment methods [Thakar *et al.*, 2025].

Acknowledgements

This work was supported by the Office of Naval Research grant N00014-23-1-2171. We thank Christopher M. Sullivan at Oregon State University for generously providing access to high-performance computing resources. We also thank the reviewers for their precise and useful feedback on our work.

References

- [Abels *et al.*, 2018] Axel Abels, Diederik M. Roijers, T. Lenaerts, Ann Nowé, and Denis Steckelmacher. Dynamic weights in multi-objective deep reinforcement learning. *ArXiv*, abs/1809.07803, 2018.
- [Basaklar *et al.*, 2023] Toygun Basaklar, Suat Gumussoy, and Umit Ogras. PD-MORL: Preference-driven multi-objective reinforcement learning algorithm. In *The Eleventh International Conference on Learning Representations*, 2023.
- [Bodnar *et al.*, 2020] Cristian Bodnar, Ben Day, and Pietro Lio. Proximal distilled evolutionary reinforcement learning. *Proceedings of the AAAI Conference on Artificial Intelligence*, 34:3283–3290, 04 2020.
- [Chen *et al.*, 2019] Xi Chen, Ali Ghadirzadeh, Mårten Björkman, and Patric Jensfelt. Meta-learning for multi-objective reinforcement learning. In *2019 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, page 977–983. IEEE Press, 2019.
- [Chen *et al.*, 2020] Diqi Chen, Yizhou Wang, and Wen Gao. Combining a gradient-based method and an evolution strategy for multi-objective reinforcement learning. *Applied Intelligence*, 50(10):3301–3317, October 2020.
- [Czyżżak and Jaskiewicz, 1998] Piotr Czyżżak and Adrezej Jaskiewicz. Pareto simulated annealing—a metaheuristic technique for multiple-objective combinatorial optimization. *Journal of Multi-criteria Decision Analysis*, 7:34–47, 1998.
- [Felten *et al.*, 2023] Florian Felten, Lucas N. Alegre, Ann Nowé, Ana L. C. Bazzan, El Ghazali Talbi, Grégoire Danoy, and Bruno C. da Silva. A toolkit for reliable benchmarking and research in multi-objective reinforcement learning. In *Proceedings of the 37th Conference on Neural Information Processing Systems (NeurIPS 2023)*, 2023.
- [Felten *et al.*, 2024] Florian Felten, El-Ghazali Talbi, and Grégoire Danoy. Multi-objective reinforcement learning based on decomposition: A taxonomy and framework. *J. Artif. Int. Res.*, 79, April 2024.
- [Foerster *et al.*, 2018] Jakob N. Foerster, Gregory Farquhar, Triantafyllos Afouras, Nantas Nardelli, and Shimon Whiteson. Counterfactual multi-agent policy gradients. In *Proceedings of the Thirty-Second AAAI Conference on Artificial Intelligence and Thirtieth Innovative Applications of Artificial Intelligence Conference and Eighth AAAI Symposium on Educational Advances in Artificial Intelligence*, AAAI’18/IAAI’18/EAAI’18. AAAI Press, 2018.
- [Fujimoto *et al.*, 2018] Scott Fujimoto, Herke van Hoof, and David Meger. Addressing function approximation error in actor-critic methods. In *International Conference on Machine Learning*, 2018.
- [Haarnoja *et al.*, 2018] Tuomas Haarnoja, Aurick Zhou, P. Abbeel, and Sergey Levine. Soft actor-critic: Off-policy maximum entropy deep reinforcement learning with a stochastic actor. *ArXiv*, abs/1801.01290, 2018.
- [Hayes *et al.*, 2022] Conor F. Hayes, Roxana Rădulescu, Eugenio Bargiacchi, Johan Källström, Matthew Macfarlane, Mathieu Reymond, Timothy Verstraeten, Luisa M. Zintgraf, Richard Dazeley, Fredrik Heintz, Enda Howley, Athirai A. Irissappane, Patrick Mannion, Ann Nowé, Gabriel Ramos, Marcello Restelli, Peter Vamplew, and Diederik M. Roijers. A practical guide to multi-objective reinforcement learning and planning. *Autonomous Agents and Multi-Agent Systems*, 36(1), April 2022.
- [Khadka and Tumer, 2018] Shauharda Khadka and Kagan Tumer. Evolution-guided policy gradient in reinforcement learning. In *Proceedings of the 32nd International Conference on Neural Information Processing Systems*, NIPS’18, page 1196–1208, Red Hook, NY, USA, 2018. Curran Associates Inc.
- [Levine *et al.*, 2020] Sergey Levine, Aviral Kumar, G. Tucker, and Justin Fu. Offline reinforcement learning: Tutorial, review, and perspectives on open problems. *ArXiv*, abs/2005.01643, 2020.
- [Lin *et al.*, 2024a] Qian Lin, Chao Yu, Zongkai Liu, and Zifan Wu. Policy-regularized offline multi-objective reinforcement learning. In *Proceedings of the 23rd International Conference on Autonomous Agents and Multiagent Systems*, AAMAS ’24, page 1201–1209, Richland, SC, 2024. International Foundation for Autonomous Agents and Multiagent Systems.
- [Lin *et al.*, 2024b] Xi Lin, Xiaoyuan Zhang, Zhiyuan Yang, Fei Liu, Zhenkun Wang, and Qingfu Zhang. Smooth tchebycheff scalarization for multi-objective optimization. In *Proceedings of the 41st International Conference on Machine Learning*, ICML’24. JMLR.org, 2024.
- [Liu *et al.*, 2025] Erlong Liu, Yu-Chang Wu, Xiaobin Huang, Chengrui Gao, Ren-Jian Wang, Ke Xue, and Chao Qian. Pareto set learning for multi-objective reinforcement learning. In *Proceedings of the Thirty-Ninth AAAI Conference on Artificial Intelligence and Thirty-Seventh Conference on Innovative Applications of Artificial Intelligence and Fifteenth Symposium on Educational Advances in Artificial Intelligence*, AAAI’25/IAAI’25/EAAI’25. AAAI Press, 2025.
- [Moffaert and Nowé, 2014] Kristof Van Moffaert and Ann Nowé. Multi-objective reinforcement learning using sets of pareto dominating policies. *Journal of Machine Learning Research*, 15(107):3663–3692, 2014.
- [Panaganti *et al.*, 2025] Kishan Panaganti, Zaiyan Xu, Dileep Kalathil, and Mohammad Ghavamzadeh. Bridging distributionally robust learning and offline rl: An approach to mitigate distribution shift and partial data coverage.

- In Necmiye Ozay, Laura Balzano, Dimitra Panagou, and Alessandro Abate, editors, *Proceedings of the 7th Annual Learning for Dynamics & Control Conference*, volume 283 of *Proceedings of Machine Learning Research*, pages 619–634. PMLR, 04–06 Jun 2025.
- [Peng *et al.*, 2019] Xue Bin Peng, Aviral Kumar, Grace Zhang, and Sergey Levine. Advantage-weighted regression: Simple and scalable off-policy reinforcement learning. *arXiv preprint arXiv:1910.00177*, 2019.
- [Reymond *et al.*, 2022] Mathieu Reymond, Eugenio Bardi, and Ann Nowé. Pareto conditioned networks. In *Proceedings of the 21st International Conference on Autonomous Agents and Multiagent Systems*, AAMAS ’22, page 1110–1118, Richland, SC, 2022. International Foundation for Autonomous Agents and Multiagent Systems.
- [Reymond *et al.*, 2023] Mathieu Reymond, Conor Hayes, Denis Steckelmacher, Diederik Roijers, and Ann Nowe. Actor-critic multi-objective reinforcement learning for non-linear utility functions. *Autonomous Agents and Multi-Agent Systems*, 37, 04 2023.
- [Roijers *et al.*, 2013] Diederik M. Roijers, Peter Vamplew, Shimon Whiteson, and Richard Dazeley. A survey of multi-objective sequential decision-making. *J. Artif. Int. Res.*, 48(1):67–113, October 2013.
- [Salimans *et al.*, 2017] Tim Salimans, Jonathan Ho, Xi Chen, Szymon Sidor, and Ilya Sutskever. Evolution strategies as a scalable alternative to reinforcement learning, 2017.
- [Sutton and Barto, 2018] Richard S. Sutton and Andrew G. Barto. *Reinforcement Learning: An Introduction*. A Bradford Book, Cambridge, MA, USA, 2018.
- [Thakar *et al.*, 2025] Raghav Thakar, Gaurav Dixit, Siddharth Iyer, and Kagan Tumer. Multiagent credit assignment for multi-objective coordination. In *Proceedings of the Genetic and Evolutionary Computation Conference*, GECCO ’25, page 663–672, New York, NY, USA, 2025. Association for Computing Machinery.
- [Tran *et al.*, 2023] Hai-Long Tran, Long Doan, Ngoc Hoang Luong, and Huynh Thi Thanh Binh. A two-stage multi-objective evolutionary reinforcement learning framework for continuous robot control. In *Proceedings of the Genetic and Evolutionary Computation Conference*, GECCO ’23, page 577–585, New York, NY, USA, 2023. Association for Computing Machinery.
- [Van Moffaert *et al.*, 2013] Kristof Van Moffaert, Madalina Drugan, and Ann Nowe. Scalarized multi-objective reinforcement learning: Novel design techniques. 04 2013.
- [Xiao *et al.*, 2025] Yuchen Xiao, Lei Yuan, Lihe Li, Ziqian Zhang, Yichen Li, and Yang Yu. Generalizable offline multiobjective reinforcement learning via preference-conditioned diffuser. *IEEE Transactions on Neural Networks and Learning Systems*, 36(12):20199–20213, 2025.
- [Xu *et al.*, 2020] Jie Xu, Yunsheng Tian, Pingchuan Ma, Daniela Rus, Shinjiro Sueda, and Wojciech Matusik. Prediction-guided multi-objective reinforcement learning for continuous robot control. In Hal Daumé III and Aarti Singh, editors, *Proceedings of the 37th International Conference on Machine Learning*, volume 119 of *Proceedings of Machine Learning Research*, pages 10607–10616. PMLR, 13–18 Jul 2020.
- [Yang *et al.*, 2019] Runzhe Yang, Xingyuan Sun, and Karthik Narasimhan. *A generalized algorithm for multi-objective reinforcement learning and policy adaptation*. Curran Associates Inc., Red Hook, NY, USA, 2019.
- [Zhu *et al.*, 2023] Baiting Zhu, Meihua Dang, and Aditya Grover. Scaling pareto-efficient decision making via offline multi-objective RL. In *The Eleventh International Conference on Learning Representations*, 2023.