

Regret-Based Sampling of Pareto Fronts for Multiobjective Robot Planning Problems

Alexander Botros , *Member, IEEE*, Nils Wilde , *Member, IEEE*, Armin Sadeghi , *Member, IEEE*,
Javier Alonso-Mora , *Senior Member, IEEE*, and Stephen L. Smith , *Senior Member, IEEE*

Abstract—Many problems in robotics seek to simultaneously optimize several competing objectives. A conventional approach is to create a single cost function comprised of the weighted sum of the individual objectives. Solutions to this scalarized optimization problem are Pareto optimal solutions to the original multiobjective problem. However, finding an accurate representation of a Pareto front remains an important challenge. Uniformly spaced weights are often inefficient and do not provide error bounds. We address the problem of computing a finite set of weights whose optimal solutions closely approximate the solution of any other weight vector. To this end, we prove fundamental properties of the optimal cost as a function of the weight vector. We propose an algorithm that greedily adds the weight vector least-represented by the current set, and provide bounds on the regret. We extend our method to include suboptimal solvers for the scalarized optimization, and handle stochastic inputs to the planning problem. Finally, we illustrate that the proposed approach significantly outperforms baseline approaches for different robot planning problems with varying numbers of objective functions.

Index Terms—Motion and path planning, multi-objective optimization, optimization and optimal control, robust/adaptive control of robotic systems.

I. INTRODUCTION

IN MANY robotic planning problems, one seeks to optimize several competing objectives. Examples include motion planning and trajectory generation for autonomous vehicles [1], [2], [3], [4], [5], human-robot cooperation for task completion [6], warehouse robotics [7], mobility-on-demand servicing [8], and neural networks [9] to name a few.

In Multiobjective Optimization (MOO) [10], one seeks solutions that achieve an appropriate trade-off between objectives. These problems are often solved through scalarization, which combines multiple objectives into a single cost function. An

example is *linear scalarization*, where the cost function is a weighted sum of objectives. Linear scalarization leads to solutions that are Pareto optimal for the MOO problem [11]. That is, the solution to the scalarized single-objective problem cannot be changed to improve the value of one of its objectives without degrading the value of another. However, a challenge in linear scalarization is how to appropriately choose weights on each objective to achieve a desired tradeoff.

This work is motivated by two classes of robotics problems: 1) obtaining online near-optimal solutions to a linearly scalarized multiobjective optimization problem (LSMOP) for any given weight vectors; and 2) learning the preferred solution behavior of a human user. Applications of the first class include [8] where an adjustable tradeoff between the service quality and operating costs for autonomous mobility-on-demand systems are optimized. Conversely, the second class of problems seeks to compute a weight vector representing the relative importance of each objective to a user given some knowledge of that users' preferred solutions [7], [12], [13], [14], [15], [16].

In either class, it is often beneficial to precompute solutions to the LSMOP for a set of weight vectors. If the LSMOP is computationally intensive to solve, requiring online solutions may be impractical. This motivates the problem of finding a set of weight vectors and their corresponding optimal solutions such that for any possible weight vector, there exists an element of the set whose solution is close to optimal. A naive approach would involve densely sampling the set of all possible weight vectors. However, the sensitivity of some objectives to changes in solution may result in a skewed sampling of Pareto optimal solutions. This is illustrated in Fig. 1 where we seek to compute a Dubins trajectory between fixed start and goal configurations that minimizes a tradeoff between trajectory length and discomfort (measured as the integral of the squared jerk over the trajectory [17]). In the top figure, a set of weight vectors is selected uniformly and the resulting solution trajectories (right) and Pareto front (bottom left) is shown. We observe that uniformly sampling weights generates multiple similar trajectories and does not approximate well all different tradeoffs, as shown by the large gap in the Pareto Front.

In this article, we propose a greedy algorithm that constructs a set of weight vectors Ω by recursively adding weight vectors that are *least represented* by the current set. Only assuming that the given objective functions are bounded the corresponding optimal solutions for weights in Ω provide homogeneous coverage of the Pareto front, as illustrated in Fig. 1(b).

Manuscript received 22 November 2023; revised 19 April 2024; accepted 8 June 2024. Date of publication 16 July 2024; date of current version 12 August 2024. This work was supported in part by the Natural Sciences and Engineering Research Council of Canada (NSERC), and in part by the European Union's Horizon 2020 research and innovation program under Grant 101017008. This article was recommended for publication by Associate Editor R. Murrieta-Cid and Editor P. R. Giordano upon evaluation of the reviewers' comments. (Alexander Botros and Nils Wilde contributed equally to this work.) (Corresponding author: Nils Wilde.)

Alexander Botros, Armin Sadeghi, and Stephen L. Smith are with the Department of Electrical and Computer Engineering, University of Waterloo, Waterloo, ON N2L 3G1, Canada (e-mail: abotros@uwaterloo.ca; asadegh@uwaterloo.ca; stephen.smith@uwaterloo.ca).

Nils Wilde and Javier Alonso-Mora are with the Department for Cognitive Robotics, Delft University of Technology, 2628 Delft, The Netherlands (e-mail: n.wilde@tudelft.nl; j.alonsomora@tudelft.nl).

Digital Object Identifier 10.1109/TRO.2024.3428990

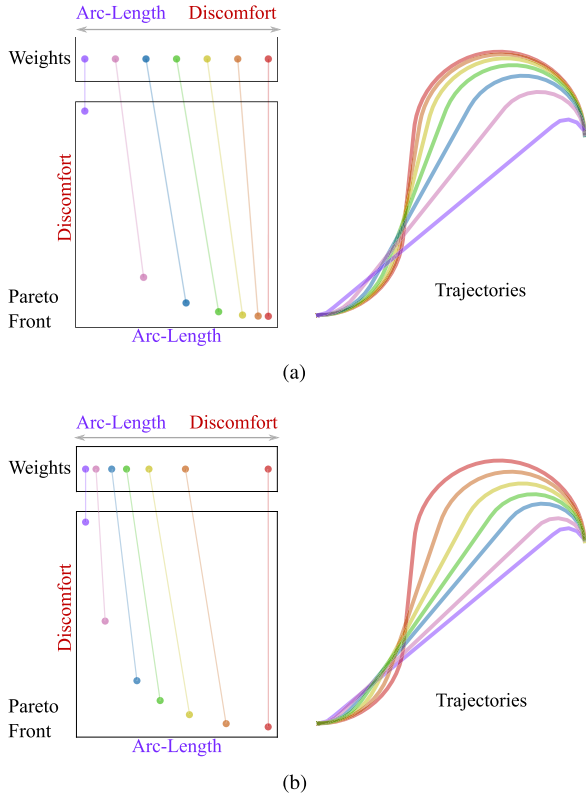


Fig. 1. Optimal tradeoffs between trajectory length and jerk for a Dubins vehicle, comparing solutions computed for uniformly sampled weights (a) and solutions found with the proposed method (b). (a) Uniformly sampled weights. (b) Proposed min-regret sampling of weights.

A. Contributions

The contributions of this work are as follows.

- 1) Assuming that an exact solver for an LSMOP is available, we propose an algorithm to compute a set of weight vectors Ω whose corresponding solutions provide homogeneous sampling of the Pareto front.
- 2) We provide a bound on the error incurred from approximating the optimal solution to an LSMOP for any arbitrary weight with a solution corresponding to a weight vector from Ω .
- 3) We relax the requirement that an exact solver be available and extend the first contribution to include a suboptimal solver. We extend the second contribution if an approximation factor for this solver is known.
- 4) If a probability density function over the set of all weights is known (e.g., representing the likelihood that a desired weight lies in a subregion), we extend the first and second results to reflect the *expected* error.
- 5) Finally, we showcase the advantages of our sampling algorithm in different robotics applications, namely robot trajectory planning, multivehicle traveling salesman problem (mTSP), and learning user preferences.

Our earlier work [18] established the first and second contributions listed above. This article extends that work by incorporating suboptimal solvers (third contribution) and stochastic

settings (fourth contribution) into the proposed approach. Finally, we extend the simulations to showcase our third and fourth contributions and include the work of [19] as a baseline.

B. Related Work

We review three areas of related work: 1) the use of weighted sums to tackle multiobjective planning problems in robotic applications; 2) the use of weighted sums to describe user preferences in HRI; and 3) techniques for approximating Pareto Fronts.

1) *Linear Scalarization in Robotics*: The simplicity of linear scalarization has made it one of the most widely used tools in robotics for considering different objectives in a cost or reward function. Though the technique is not able to capture all Pareto-optimal solutions for nonconvex fronts, it does guarantee that all LSMOP solutions are Pareto-optimal. For instance, cost functions in autonomous driving often consider objectives such as trajectory length, comfort (measured via jerk), or clearance [1], [2], [3], [4], [5], [17], [20]. Other examples for motion planners that use weighted sums to balance between objectives include local planners for mobile robots navigating in cluttered environments [21] or social spaces [22], [23], trajectory planners for manipulators [24], [25], and multirobot planning [26].

In [27], a trajectory smoothing algorithm was proposed based on the weighted sum of competing objectives, namely trajectory length, smoothness, and obstacle distances. The authors of [28] minimized the weighted tradeoff between mission completion time and communication outage duration in the navigation of cellular-connected UAVs, while in [29], linear scalarization was used to optimize robotic limitations and observation rewards for use in autonomous human activity tracking. The authors of [30] considered biobjective path planning. In particular, the goal is to simultaneously optimize for path length and clearance to obstacles in the plane for which they propose a complete and efficient algorithm. Our work does not consider specific objectives as the abovementioned papers. Instead, we focus on Pareto-optimal tradeoffs for any multiobjective planning problem that is formulated as a weighted sum.

2) *Weighted Sums Describing User Preferences*: In human-robot interaction (HRI), weight vectors are used to represent a user's preference for robot behavior, i.e., the relative importance of objectives [12], [15], [16], [31], [32], [33]. In *reward learning* the objective is to learn a user's weight vector using interactions such as demonstrations, corrections, or choice feedback. In order to expedite the learning process, feasible solutions for the multiobjective optimization problem are often precomputed and shown to the user who then provides feedback. In [12], [16], and [34], each precomputed solution was generated with random action sequences. Thus, the solutions used in the learning process are usually not optimal for any weight. In [15] and [35], the authors precomputed Pareto-optimal solutions enabling an active learning method based on regret leading to significant improvement over randomly generated solutions. However, the work in [15] and [35] relied on uniformly sampled weight vectors. Though this approach asymptotically covers the set of all LSMOP-optimal solutions, it can be inefficient as different

weight vectors can have very similar, or even identical solutions. The counter intuitive relationship between weights and collected reward was further studied in [36]. Moreover, the authors of [37] observed that the presence of similar solution strongly influences the Boltzmann decision model which is commonly used in HRI. They propose a decision model where a similarity metric corrects the bias induced by a high number of similar trajectories. Similarly, in our work we are interested in finding solutions with dissimilar features. While [37] handled the overrepresentation of similar solutions in the ground set with their proposed decision model, we instead address the problem at an earlier stage: Our algorithm can be used to generate a ground set where similarities are minimized.

3) *Approximating Pareto Fronts*: Our work is motivated from the following observation: a uniform sampling of weights does not generally provide a uniform sampling of solutions as illustrated in Fig. 1. These shortcomings of weighted sum methods are well studied within the optimization literature [38], [39], yet have received less attention in the robotics community despite the wide usage of weighted sums as objective functions. Researchers in optimization developed many techniques for approximating Pareto-fronts, including more complex scalarization methods such as Chebyshev scalarization [38]. Unfortunately, these methods are often not directly applicable to robot planning since they require solving more complex scalar optimization problems. As a consequence, many multiobjective robot planning problems still rely on the simple weighted sum formulation. Finding a set of representative scalarization weights requires solving a series of optimization problems, e.g., solving path or motion planning instances, which can be computationally burdensome. Thus, efficient computation of such a set is of particular interest.

Closely related to our work is the *Adapted Weighted Sum* (AWS) method [19], [40]. The AWS iteratively places Pareto samples by partitioning existing samples into subsets and placing new samples into the subset that has the largest *gap*. This requires solving an optimization problem with an additional linear constraint on the objective value, which can result in a harder problem than the original weighted sum optimization. In contrast, our method identifies the most promising weights for additional samples and then solves the weighted sum optimization problem. Further, our method minimizes the regret of the weighted samples and returns an error bound.

The authors of [41] presented a Pareto front approximation for trajectory planning using Markov chain random walks. Their goal is to *uniformly* place samples on the Pareto front, while our goal is to *minimize error* in the space of Pareto-optimal costs. Moreover, the random walk technique does not rely on a weighted sum objective, but does not generalize to an arbitrary choice of objective functions. Similar to our work, the authors of [42] offered a technique of Pareto-uniform sampling based on equispacing constraints. However, they only consider the case of two competing objectives. Further, they accomplish their goal by solving a nested optimization with the original LSMOP as the inner-most problem, which can be significantly harder. The work in [43] proposed a set of weight vectors that approximately uniformly cover a Pareto front specifically for

use in the design of robots. The authors designed the set that minimizes the total squared error between the value of the objectives in the set and heuristic objectives. It therefore relies on the approximate optimality of these objectives. The work in [44] proposed a method to cover the set of Pareto-optimal solutions specifically for use in reinforcement learning applications. The authors seek to compute policies that maximize expected returns by computing, storing, and updating a set of samples. In [45], the authors provided a means of exploring a (possibly nonconvex) Pareto front in order to obtain a solution that is near-optimal for a user. That work starts with an initial guess solution which moves in a direction according to a user's preference. While the convexity of the Pareto front (a requirement for linear scalarization to obtain all Pareto-optimal solutions) is not assumed, their technique requires solving the LSMOP online as the Pareto front is explored. Moreover, the requirement of a user-preferred direction is not assumed in our work.

In summary, most state-of-the-art methods either address specific problem setups and thus do not generalize across different robot planning problems, or address the problem of finding scalarization weights more generally but pose a complex optimization problem to compute a set of weights. Our work considers the weighted sum formulation for an arbitrary choice of objective functions. We iteratively compute a set of weights, only requiring solving a linear program and the weighted sum objective for different weights, and return a bound on the approximation error for the computed set.

II. PROBLEM STATEMENT

For $n \in \mathbb{N}$, a general multiobjective optimization problem (MOP) is of the form

$$\min_{s \in \mathcal{S}} \left\{ f_1(s), f_2(s), \dots, f_n(s) \right\}. \quad (1)$$

Here, the set of feasible solutions given constraints is denoted $\mathcal{S}$, and it is desired to simultaneously minimize n objectives $f_i(s)$, for $i \in \{1, \dots, n\}$. However, such a solution typically does not exist. As a result, multiobjective optimization often seeks Pareto-optimal solutions: solutions $s \in \mathcal{S}$ such that there does not exist another solution $\bar{s} \in \mathcal{S}$ where $f_i(\bar{s}) \leq f_i(s)$ for all i , and with strict inequality for at least one i . The linear scalarization of the MOP above involves the creation of a single cost function by introducing a vector of weights $\mathbf{w} = [w_1, w_2, \dots, w_n] \in \mathbb{R}_{\geq 0}^n$. Let $c(s, \mathbf{w})$ denote the cost of the solution s evaluated by the weights $\mathbf{w}$, i.e.,

$$c(s, \mathbf{w}) = \sum_{i=1}^n w_i \cdot f_i(s) = \mathbf{w} \cdot \mathbf{f}(s) \quad (2)$$

where $\mathbf{f}(s) = [f_1(s), \dots, f_n(s)]$, $\forall s \in \mathcal{S}$. The resulting LSMOP is to solve

$$u(\mathbf{w}) = \min_{s \in \mathcal{S}} c(s, \mathbf{w}). \quad (3)$$

For any weight $\mathbf{w} \in \mathbb{R}_{\geq 0}^n$, solution $s \in \mathcal{S}$, and $\lambda \in \mathbb{R}_{> 0}$, it holds that $c(s, \lambda \mathbf{w}) = \lambda c(s, \mathbf{w})$ implying that a minimizer of $c(s, \mathbf{w})$ also minimizes $c(s, \lambda \mathbf{w})$. Further, if $\mathbf{w} = [0, 0, \dots, 0]$, then $u(\mathbf{w})$ is trivially 0. Thus, given $\mathbf{w} = [w_1, \dots, w_n]$ where not

all elements are identically 0, and letting $\lambda = (\sum_{i=1}^n w_i)^{-1}$, we can obtain all nontrivial optimal solutions $u(\mathbf{w})$ for all $\mathbf{w} \in \mathbb{R}_{\geq 0}^n$ using weights $\mathbf{w} \in \mathcal{W}$ where

$$\mathcal{W} = \left\{ \mathbf{w} \in \mathbb{R}_{\geq 0}^n, \sum_{i=1}^n w_i = 1 \right\}. \quad (4)$$

We refer to the set $\mathcal{W}$ as the *weight space*, and define

$$s^*(\mathbf{w}) = \arg \min_{s \in \mathcal{S}} \mathbf{w} \cdot \mathbf{f}(s), \forall \mathbf{w} \in \mathcal{W} \quad (5)$$

as an *optimal solution* given weights $\mathbf{w}$, implying that $u(\mathbf{w}) = c(s^*(\mathbf{w}), \mathbf{w})$ by (3). We start with the following assumptions:

Assumption 1 (Exact Solution): An exact solver exists for the optimization problem (3).

Assumption 2 (Bounded Objectives): For any weight $\mathbf{w} \in \mathbb{R}_{\geq 0}^n$, and any optimal solution $s^*(\mathbf{w})$ in (5), the objectives $\mathbf{f}(s^*(\mathbf{w}))$ are bounded.

While Assumption 2 persists throughout this article, Assumption 1 is relaxed for suboptimal solvers (Section V).

In this work, we propose a method to compute a finite set of weights $\Omega \subset \mathcal{W}$ that will, for any $\mathbf{w}^* \in \mathcal{W}$, allow us to approximate $u(\mathbf{w}^*)$ with $u(\mathbf{w}')$ for an appropriately chosen $\mathbf{w}' \in \Omega$. To evaluate the quality of a candidate set Ω , we use the notion of regret from [15] and [46], defined formally here

Definition 1 (Regret): Given two weights $\mathbf{w}', \mathbf{w}^* \in \mathcal{W}$, the regret of $\mathbf{w}'$ under $\mathbf{w}^*$ is defined as

$$r(\mathbf{w}'|\mathbf{w}^*) = \mathbf{w}^* \cdot \mathbf{f}(s^*(\mathbf{w}')) - u(\mathbf{w}^*). \quad (6)$$

Intuitively, $r(\mathbf{w}'|\mathbf{w}^*)$ represents the *error* in cost incurred by using an optimal solution given weight $\mathbf{w}'$ (given by $s^*(\mathbf{w}')$) to approximate a solution given weight $\mathbf{w}^*$. We now formally state the main problem addressed in this work.

Problem 1 (Min-Max Regret Sampling): For the LSMOP (3) and an integer $K > 0$, find a set of weights Ω that solves

$$\begin{aligned} \min_{\Omega} \max_{\mathbf{w}^* \in \mathcal{W}} \min_{\mathbf{w}' \in \Omega} r(\mathbf{w}'|\mathbf{w}^*) \\ \text{s.t. } |\Omega| \leq K. \end{aligned} \quad (7)$$

Given a weight $\mathbf{w}^* \in \mathcal{W}$ and a set $\Omega \subset \mathcal{W}$, the first minimization in (7), $\min_{\mathbf{w}' \in \Omega} r(\mathbf{w}'|\mathbf{w}^*)$ represents the suboptimality of approximating a solution $s^*(\mathbf{w}^*)$ with a solution $s^*(\mathbf{w}')$ where $\mathbf{w}'$ is the weight in Ω that minimizes this suboptimality, i.e., $\mathbf{w}'$ is a best representative of $\mathbf{w}^*$ in Ω . The maximization $\max_{\mathbf{w}^* \in \mathcal{W}} \min_{\mathbf{w}' \in \Omega} r(\mathbf{w}'|\mathbf{w}^*)$, represents the regret of the worst represented weight $\mathbf{w}^* \in \mathcal{W}$ by elements in Ω . We refer to the solution of this maximization as the *maximum regret given Ω* . In total, (7) seeks a set Ω such that the regret of the worst represented element in $\mathcal{W}$ is minimized.

In this article, we offer an approximate solution to the optimization in (7) by way of an algorithm that computes a feasible solution Ω such that the maximum regret given Ω is bounded. In the next section, we provide the theoretical groundwork that makes this solution possible.

III. PROBLEM ANALYSIS

We begin with a structural analysis of the cost function $u(\mathbf{w})$ from (3) to derive an efficient algorithm for solving Problem 1. First, we make two critical observations.

Observation 1: Given any two weights $\mathbf{w}^*, \mathbf{w}' \in \mathcal{W}$, we have

$$u(\mathbf{w}^*) \leq \mathbf{w}^* \cdot \mathbf{f}(s^*(\mathbf{w}')). \quad (8)$$

That is an optimal solution given weights $\mathbf{w}^*$ will incur no higher cost than a solution that is optimal for some different weight vector $\mathbf{w}'$. Here, $s^*(\mathbf{w}')$ is optimal given weights $\mathbf{w}'$ (see (5)) but not necessarily optimal given weights $\mathbf{w}^*$. By (6), the inequality in (8) implies that $r(\mathbf{w}'|\mathbf{w}^*) \geq 0$.

Observation 2 (Optimal Cost Concavity): The optimal cost function $u(\mathbf{w})$ is a concave function of $\mathbf{w}$. Indeed, for each $s \in \mathcal{S}$, the cost $c(s, \mathbf{w}) = \mathbf{w} \cdot \mathbf{f}(s)$ is an affine function of $\mathbf{w}$ (and is therefore, concave). Therefore, $u(\mathbf{w}) = \min_{s \in \mathcal{S}} c(s, \mathbf{w})$ is concave [47, Sec. 3.2.3].

Observation 2 motivates the following Lemma:

Lemma 1 (Optimal Cost Continuity): Under Assumptions 1 and 2, given any two weights in $\mathcal{W}$, $u(\mathbf{w})$ is continuous on the line segment connecting those weights.

Proof: By Observation 2, $u(\mathbf{w})$ is concave in $\mathcal{W}$. Noting in addition that $\mathcal{W} \subset \mathbb{R}^n$ is convex, it must hold that $u(\mathbf{w})$ is continuous on the interior of $\mathcal{W}$. This is because concave functions are continuous on the interior of convex sets. Therefore, it suffices to prove the result for the case that at least one weight lies on the boundary of $\mathcal{W}$. Consider two weights $\mathbf{w}', \mathbf{w}'' \in \mathcal{W}$ at least one of which lies on the boundary of $\mathcal{W}$. Suppose that $\|\mathbf{w}' - \mathbf{w}''\| \leq \delta$ for some $\delta > 0$ arbitrarily small, and—without loss of generality—that $u(\mathbf{w}') > u(\mathbf{w}'')$. Because it is assumed that $u(\mathbf{w})$ is bounded on $\mathcal{W}$, if it experiences a discontinuity on the line connecting $\mathbf{w}', \mathbf{w}''$, it must hold that $u(\mathbf{w})$ experiences a jump between $\mathbf{w}', \mathbf{w}''$. That is, there must exist a $M \in \mathbb{R}_{>0}$ independent of δ such that $u(\mathbf{w}'') + M \leq u(\mathbf{w}')$. Since $\|\mathbf{w}'' - \mathbf{w}'\| \leq \delta$, and $\mathbf{f}(s^*(\mathbf{w}''))$ is bounded by Assumption 2, there must exist a value of δ sufficiently small so as to guarantee that $(\mathbf{w}' - \mathbf{w}'') \cdot \mathbf{f}(s^*(\mathbf{w}'')) < M/2$. Therefore, by construction

$$\begin{aligned} \mathbf{w}' \cdot \mathbf{f}(s^*(\mathbf{w}'')) &< \frac{M}{2} + \mathbf{w}'' \cdot \mathbf{f}(s^*(\mathbf{w}'')) = \frac{M}{2} + u(\mathbf{w}'') \\ &\leq \frac{M}{2} + u(\mathbf{w}') - M = u(\mathbf{w}') - \frac{M}{2} < u(\mathbf{w}'). \end{aligned}$$

Therefore, $r(\mathbf{w}''|\mathbf{w}') = \mathbf{w}' \cdot \mathbf{f}(s^*(\mathbf{w}'')) - u(\mathbf{w}') < 0$ which is a contradiction by Observation 1. ■

Critically, the previous results do not require unique solutions $s^*(\mathbf{w})$ or continuous objectives $\mathbf{f}(s^*(\mathbf{w}))$. Extending these results:

Theorem 1 (Convexity of Regret): For a fixed weight $\mathbf{w}' \in \mathcal{W}$, the regret $r(\mathbf{w}'|\mathbf{w})$ is a convex function of $\mathbf{w}$.

Proof: By (6), $r(\mathbf{w}'|\mathbf{w}) = \mathbf{w} \cdot \mathbf{f}(s^*(\mathbf{w}')) - u(\mathbf{w})$ where $\mathbf{w} \cdot \mathbf{f}(s^*(\mathbf{w}'))$ is linear in $\mathbf{w}$ and $u(\mathbf{w})$ is concave (Observation 2). Thus, $r(\mathbf{w}'|\mathbf{w})$ is the difference of linear and concave functions of $\mathbf{w}$ which is convex. ■

Because $u(\mathbf{w})$ is continuous and concave, it must hold that the function lies below any subgradient. This motivates the

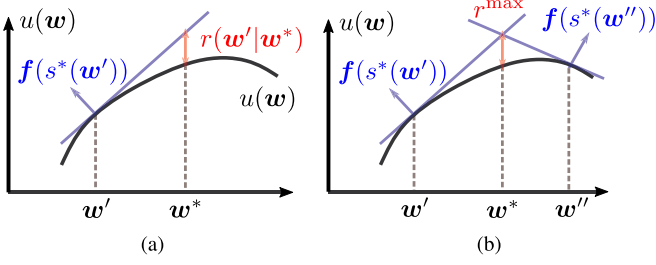


Fig. 2. (a) Regret as the error of a first order approximation. (b) Maximum regret given a set $\Omega = \{w', w''\}$ at the intersection of their tangent lines.

following Corollary which follows directly from the definition of $u(w)$ and the concavity of $c(s, w)$ for each fixed $s \in \mathcal{S}$ [47, Sec. 6.5.5].

Corollary 1.1 (Subgradient Optimal Cost): For any $w \in \mathcal{W}$ and any minimizing solution $s^*(w) \in \mathcal{S}$, the vector of objectives $f(s^*(w))$ is a subgradient, $\partial u(w)$, of u at w .

Observe that subgradients are defined even for nondifferentiable continuous functions. Further, the results above imply that given two weights $w', w^* \in \mathcal{W}$, the regret $r(w'|w^*)$ —which coincides with the error incurred by approximating a solution $s^*(w^*)$ with the solution $s^*(w')$ —is exactly the error of approximating a concave function via *linear interpolation*. Indeed, the first order approximation of $u(w^*)$ given $u(w')$ is given by $u(w^*) \approx u(w') + \nabla u(w') \cdot (w^* - w')$ assuming u is differentiable at w' . However, by Corollary 1.1, $\nabla u(w) = \partial u(w) = f(s^*(w))$. This together with the definition $u(w') = f(s^*(w')) \cdot w'$ allows us to conclude that $u(w^*) \approx u(w') + f(s^*(w')) \cdot w^* - f(s^*(w')) \cdot w' = f(s^*(w')) \cdot w^*$. The error of this first order approximation is $f(s^*(w')) \cdot w^* - u(w^*)$ which is exactly the regret $r(w'|w^*)$.

This is illustrated in Fig. 2(a). Further, given any two weights $w', w'' \in \mathcal{W}$, the maximum regret on the line segment L between w', w'' given $\Omega = \{w', w''\}$ occurs at the weight on L coinciding with the intersection of the tangent lines to $u(w)$ at w' and w'' along L [Fig. 2(b)]. In light of this analysis, the objective in (7) is solved by a set Ω that provides the best linear interpolation of the concave function $u(w)$. These insights are leveraged in the following section.

IV. ALGORITHM

In this section, we present our solution to Problem 1. The algorithm we propose recursively adds weights to a solution set Ω . A strong candidate weight to add is one that is least represented by the current iteration of Ω . The basic framework for such an approach could be described recursively

$$\Omega^{k+1} = \Omega^k \cup \left\{ \arg \max_{w^* \in \mathcal{W}} \min_{w' \in \Omega^k} r(w'|w^*) \right\} \quad (9)$$

where Ω^k is the solution after k iterations from an initial set. Here, (9) recursively adds the weight with the maximum regret given Ω^k . Obtaining the maximizer w^* is nontrivial due to its nested structure. Instead, our approach replaces $r(w'|w^*)$ in (9)

Algorithm 1: MINIMUM-REGRET PARETO SAMPLING (MRPS).

Input: An exact solver to find $s^*(w)$, $f(s^*(w))$; a budget $K \geq n$

Output: Sampled weights Ω and maximum regret

```

1  $\Omega \leftarrow \{e^i, i = 1, \dots, n\}$  // where  $e^i$  is the  $i^{th}$  row of
   the  $n \times n$  identity matrix
2 Obtain  $s^*(e^i)$ ,  $f(s^*(e^i))$ ,  $i = 1, \dots, n$  from exact
   solver
3  $\mathcal{N} \leftarrow \{\Omega\}$ 
4 for  $k = n$  to  $K$  do
5    $N =$  neighborhood in  $\mathcal{N}$  with maximum  $\bar{R}(N)$ 
6   if  $\bar{R}(N) = 0$  then
7     break // Terminate if maximum upper regret
       bound is 0
8    $\Omega \leftarrow \Omega \cup \{\bar{w}(N)\}$ 
9   Obtain  $s^*(\bar{w}(N))$ ,  $f(s^*(\bar{w}(N)))$  from exact solver
10   $\mathcal{N} = \mathcal{N} \setminus N$  // Remove max-regret neighborhood
11  for  $w^i$  in  $N$  do
12     $N^i \leftarrow N \setminus \{w^i\} \cup \{\bar{w}(N)\}$  // Replace  $w^i$  with
      weight of the maximum regret bound
13    if  $N^i$  is a neighborhood, i.e., its weights are
      lin. independent then
14       $\mathcal{N} = \mathcal{N} \cup N^i$ 
15       $\mathcal{F}(N^i) \leftarrow$ 
         $\mathcal{F}(N) \setminus \{f(s^*(w^i))\} \cup \{f(s^*(\bar{w}(N)))\}$ 
16 return  $\Omega$  and the maximum value of  $\bar{R}(N)$  over all
       $N \in \mathcal{N}$ 
```

with an upper bound $R(w'|w^*)$ whose maximizer w^* is obtained from a linear program (LP).

Given a set of weights $\Omega \subseteq \mathcal{W}$, we define N as a set of n (recall that n is the number of objectives) linearly independent weights $w^1, \dots, w^n \in \Omega$, and we let $\mathcal{C}(N) \subset \mathbb{R}^n$ denote the convex hull of N . We loosely refer to N as a *neighborhood*, and define a linear lower bound of the objective value $u(w)$ from (3) inside a neighbourhood N . Let $P: \mathcal{W} \rightarrow \mathbb{R}_{\geq 0}$ be the linear function taking values $P(w^i) = u(w^i)$ for all $w^i \in N$ [Fig. 3(a)]. We denote the difference between the tangent plane at w' and the P evaluated at w^* with $R(w'|w^*) = f(s^*(w'))w^* - P(w^*)$, $\forall w' \in N, w^* \in \mathcal{C}(N)$. Further, let

$$\begin{aligned} \bar{R}(N) &= \max_{w^* \in \mathcal{C}(N)} \min_{w' \in N} R(w'|w^*), \\ \bar{w}(N) &= \arg \max_{w^* \in \mathcal{C}(N)} \min_{w' \in N} R(w'|w^*). \end{aligned} \quad (10)$$

Finally, we let $\mathcal{F}(N)$ represent the set of objective vectors of the neighborhood

$$\mathcal{F}(N) = \{f(s^*(w^i)), w^i \in N\}. \quad (11)$$

Thus, $R(w'|w^*)$ is similar to $r(w'|w^*)$ from (6), but with $u(w^*)$ replaced with $P(w^*)$. These definitions, illustrated in Fig. 3, motivate the Theorem.

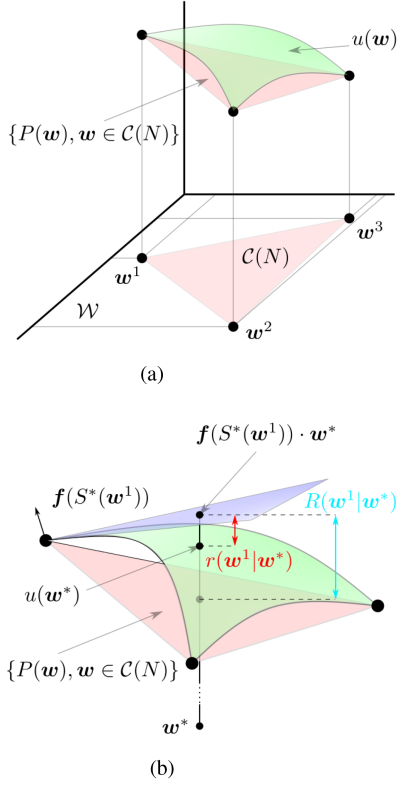


Fig. 3. Illustrative example of a neighborhood and its properties. (a) $P(w), C(N)$ for $N = \{w^1, w^2, w^3\}$ (b) Regret r and regret bound R

Theorem 2 (Upper Bound of Maximum Regret in a Neighborhood): Given a neighborhood N of weights, it holds that

$$\max_{w^* \in C(N)} \min_{w' \in N} r(w'|w^*) \leq \bar{R}(N).$$

Proof: For $w^* \in C(N)$, let $w'_1 = \arg \min_{w' \in N} r(w'|w^*)$, $w'_2 = \arg \min_{w' \in N} R(w'|w^*)$. Note that $w'_1 = w'_2$. Indeed, we have $R(w'_2|w^*) \leq R(\hat{w}|w^*)$ for all $\hat{w} \in N$ if and only if $w^* \cdot f(s^*(w'_2)) \leq w^* \cdot f(s^*(\hat{w}))$ which is equivalent to $r(w'_2|w^*) \leq r(\hat{w}|w^*)$.

By Observation 2, $u(w)$ is concave on $C(N)$ implying that for $w \in C(N)$, $u(w) \geq P(w)$ (see Fig. 3). Thus, by (6), $r(w'|w^*) \leq R(w'|w^*)$ for $w^* \in C(N)$. The result follows. ■

The value of $\bar{R}(N)$ with corresponding weight $\bar{w}(N)$ from (10) can be obtained by solving the following LP:

$$\begin{aligned} & \max_{x \in \mathbb{R}, w \in \mathbb{R}^n} x - P(w) \\ & \text{s.t.} \quad \begin{bmatrix} f_1(s^*(w^1)) & \dots & f_n(s^*(w^1)) & -1 \\ \vdots & \ddots & \vdots & \vdots \\ f_1(s^*(w^n)) & \dots & f_n(s^*(w^n)) & -1 \end{bmatrix} \begin{bmatrix} w \\ x \end{bmatrix} \succcurlyeq 0, \\ & w \in C(N). \end{aligned} \quad (12)$$

If (x^*, w^*) solves (12), the optimal cost is given by $x^* - P(w^*) = \bar{R}(N)$, and $w^* = \bar{w}(N)$. Indeed, for any feasible

x, w , it holds that $x \leq \min_{w^i \in N} f(s(w^i)) \cdot w$. Since x is maximized, this will hold with equality for x^*, w^* . Therefore, the cost of (12) is equivalent to $\max_{w \in C(N)} \min_{w^i \in N} R(w^i|w) = \bar{R}(N)$. A detailed explanation of the implementation for (12) is provided in the supplementary materials.

In (12), if $P(w)$ is replaced with $u(w)$, then the resulting problem is solved by x^*, w^* if and only if w^* maximizes the regret in $C(N)$ given the neighborhood N . This problem is not linear and would require solving the LSMOP in (3) potentially many times. Using the LP in (12) our method is summarized in Algorithm 1 described in the next section. We iteratively partition $\mathcal{W}$ into smaller neighborhoods, adding weights that result in the largest upper bound of regret.

A. Algorithm Description

Algorithm 1 creates and maintains a set $\mathcal{N}$ of neighborhoods $N \subset \mathcal{W}$. Each $N \in \mathcal{N}$ is a set of weights $N = \{w^1, \dots, w^n\}$ where $w^i \in \Omega, i = 1, \dots, n$. We compute $\bar{R}(N)$ and $\bar{w}(N)$ with the LP in (12) for N using the set of objective vectors $\mathcal{F}(N)$. The algorithm begins with a single neighborhood whose weights are the n canonical basis elements of $\mathbb{R}^n$ (Line 1). The solutions for these basis weights correspond to the single-objective solutions for all n objective functions. We assume that the budget K is at least the number of objective functions n . The algorithm then iteratively selects the neighborhood N in $\mathcal{N}$ with the largest upper bound of regret (Line 5), and adds its regret weight $\bar{w}(N)$ to Ω (Line 8). It then splits and replaces N with at most n smaller neighborhoods (Lines 11–15) formed by iteratively replacing elements in N with $\bar{w}(N)$ (Line 12). Finally, the algorithm returns the set Ω as well as an upper bound on the regret given Ω (Line 16). Two steps of the algorithm are illustrated in Fig. 4, starting with a single neighborhood N_1 in (a) which is then split around $w^3 = \bar{w}(N_1)$ into two new neighborhoods $\mathcal{N} = \{N_2, N_3\}$ in (b). Since $\bar{R}(N_3) > \bar{R}(N_2)$, N_3 is split around $w^5 = \bar{w}(N_3)$ in (c). Finally, in (d), the red area shows the regret given Ω .

Observe that Algorithm 1 may be modified to compute a set Ω given a desired maximum regret $r_{\max} > 0$. This could be accomplished by replacing the input K with $r_{\max}$, and the stopping criteria in Line 4 with a while loop that runs until $\bar{R} \leq r_{\max}$. Here, $\bar{R}$ represents the maximum regret over all neighborhoods $\bar{R} = \max_{N \in \mathcal{N}} \bar{R}(N)$ and can be maintained in the body of the loop. Since $\mathcal{N}$ forms a partition of $\mathcal{W}$, we are guaranteed that the regret of any weight given Ω is no more than $\bar{R}$ by Theorem 2. Therefore, if Algorithm 1 terminates when $\bar{R} \leq r_{\max}$, the desired maximum regret is achieved.

B. Algorithm Properties

We observe several beneficial properties to the approach outlined above.

Observation 3 (Runtime): For a budget of K , Algorithm 1 will require that the LSMOP in (3) be solved at most K times, once per element of Ω . Let t_{LSMOP} be the runtime to solve (3). Using the interior-point method allows for solving LPs in $O(a^2 b^{3/2})$ time with a being the number of variables, and b the number of constraints. The LP in (12) has $2n$ variables and

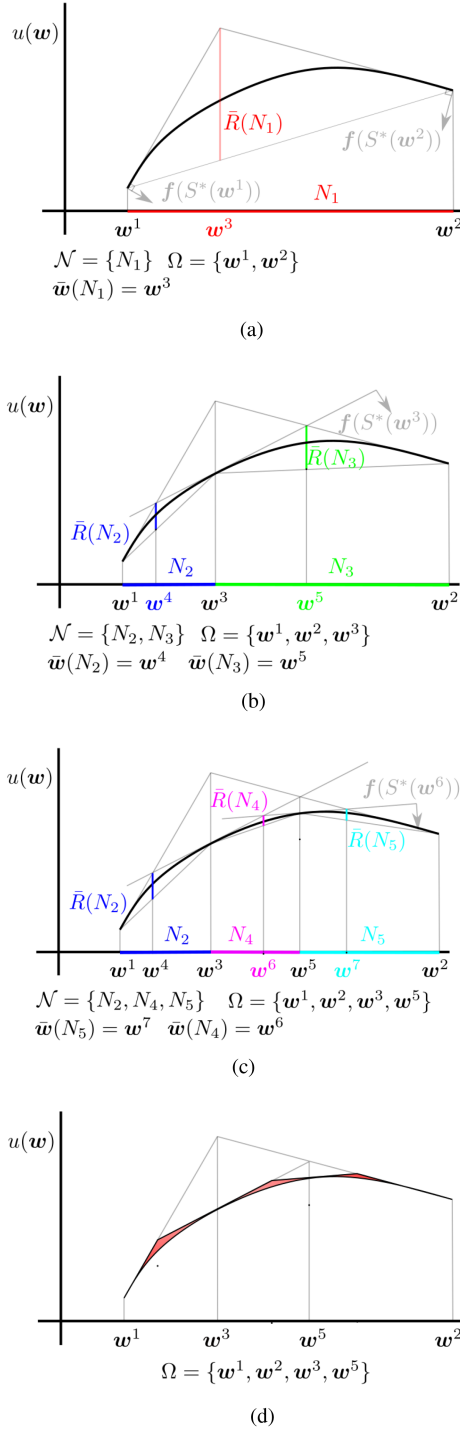


Fig. 4. Illustration of the first two iterations of Algorithm 1. (a) First iteration. (b) Second iteration. (c) Third iteration. (d) Regret after three iterations.

$2n$ constraints with n being the number of objective functions (see (21) in the Appendix for details). Thus, the runtime of Algorithm 1 is $O(K(t_{\text{LSMOP}} + n^3))$.

Observation 4 (Regret bound): The value of $\bar{R}(N)$ returned by the algorithm is an upper-bound on the value of regret in the original problem (7). Indeed, initially $\mathcal{N} = \{\Omega\}$ and $\mathcal{C}(\Omega) = \mathcal{W}$. At every iteration, a neighborhood $N \in \mathcal{N}$ is split into at

most n subneighborhoods such whose convex hulls are disjoint and collectively form $\mathcal{C}(N)$. Then, by Theorem 2, it holds that $\max_{w^* \in \mathcal{W}} \min_{w' \in \Omega} r(w'|w^*) \leq \max_{N \in \mathcal{N}} \bar{R}(N)$.

Lemma 2 (Algorithm Completeness): The set $\Omega(K)$ returned by Algorithm 1 on input K asymptotically and monotonically approaches a set with zero regret as $K \rightarrow \infty$.

Proof: The proof follows by contradiction: If the result did not hold, then a neighborhood $N \in \mathcal{N}$ would fail to decrease in size when split (Lines 11–15). This in turn requires that there is a $w^i \in N \subseteq \Omega(K)$ such that $\|\bar{w}(N) - w^i\|_2$ decreases to 0. Since $\bar{w}(N)$ is chosen to maximize $\bar{R}(N)$, this can only occur if $\bar{R}(N)$ is unbounded at w^i implying that the objectives are unbounded at w^i in violation of Assumption 2. ■

Observation 5 (Optimality for Discrete Solution Spaces): In the case where the solution space $\mathcal{S}$ of the LSMOP in (3) is discrete, the function $u(w)$ will be piece-wise linear by Lemma 1. If K is at least the number of linear pieces of $u(w)$, the solution set $\Omega(K)$ has zero maximum regret and is the smallest set that accomplishes this. Indeed $\Omega(K)$ is comprised of exactly one weight in each linear piece of $u(w)$. Since the regret is defined by the error of a first order approximation, this value is exactly zero on each linear piece.

In the next section, we extend Algorithm 1 to the case where Assumption 1 does not hold. That is, where no exact solver for the optimization Problem 3. This is followed by an extension to Algorithm 1 for the case when a prior belief about the optimal weights of a user are known (See Section VI). It should be noted that both extensions reduce to Algorithm 1 when the suboptimal solver is optimal, or the prior belief is uniform (respectively). Moreover, these extensions can be used in tandem for the case where Assumption 1 is violated and a prior belief about the optimal user weights is known. This last is not formally stated, but is trivially obtained from the descriptions of each extension to follow.

V. EXTENSION TO SUBOPTIMAL SOLVERS

In this section, we illustrate how Algorithm 1 can be adapted to accept suboptimal solvers for the underlying LSMOP. Instead of $s^*(w)$, an optimal solution to an LSMOP for weights w (see (5)), consider a feasible solution $\hat{s}(w) \in \mathcal{S}$ obtained from a suboptimal solver. Following Section II and given any weight vector $w \in \mathcal{W}$, we let $\hat{u}(w) = c(\hat{s}(w), w)$ (see (2)) be the cost of the solution $\hat{s}(w)$ given weights w . Observe that $\hat{u}$ is identical to the function u (see (3)) except that a suboptimal solution $\hat{s}(w)$ is now in place of $s^*(w)$.

A. Look-Up Table Solutions

Algorithm 1 relies heavily on the concavity of $u(w)$ as a function of w , which is guaranteed from the optimality of $s^*(w)$. Therefore, the function $\hat{u}(w)$ —the cost of a solution computed using a suboptimal solver—is not guaranteed to be concave. In this section, we begin with a technique to replace $\hat{u}(w)$ with a concave function. The high level idea is to maintain a look-up table, i.e., a set of discovered solutions $\mathcal{T}$. Given any weight $w \in \mathcal{W}$, we define a new solver $s_{\mathcal{T}}$ to select the best solution in $\mathcal{T}$. In detail, given a set of sampled weights $\Omega = \{w_1, \dots, w_r\} \subset$

$\mathcal{W}, r \geq 1$, we denote by $\mathcal{T}(\Omega) = \{\hat{s}(\mathbf{w}), \mathbf{w} \in \Omega\}$ the set of all solutions to weights $\mathbf{w} \in \Omega$ obtained via the suboptimal solver $\hat{s}$.

For any subset of weights $\Omega \subset \mathcal{W}$ and its associated set of discovered solutions $\mathcal{T}(\Omega)$, we define the *best discovered solution* to any weight $\mathbf{w} \in \mathcal{W}$ as

$$s_{\mathcal{T}}(\mathbf{w}) = \operatorname{argmin}_{s \in \mathcal{T}(\Omega)} c(s, \mathbf{w}). \quad (13)$$

Further, we define a *best discovered cost function* $u_{\mathcal{T}}(\mathbf{w}) = c(s_{\mathcal{T}}(\mathbf{w}), \mathbf{w})$. Finally, we define the *best discovered regret function* for any weights $\mathbf{w}', \mathbf{w}^* \in \mathcal{W}$

$$r_{\mathcal{T}}(\mathbf{w}' | \mathbf{w}^*) = \mathbf{w}^* \cdot \mathbf{f}(s_{\mathcal{T}}(\mathbf{w}')) - u_{\mathcal{T}}(\mathbf{w}^*).$$

We now prove that for any subset Ω , the best discovered cost function $u_{\mathcal{T}}(\mathbf{w})$ has the same properties as the optimal cost function $u(\mathbf{w})$ from (3) that are leveraged by Algorithm 1. For the remainder of this section, we assume that Ω is a nonempty countable subset of $\mathcal{W}$, and we let $\mathcal{T}(\Omega)$ denote the associated set of discovered solutions.

Lemma 3 (Best Discovered Cost Properties): For any set of weights Ω , it holds that $u_{\mathcal{T}}(\mathbf{w})$ is a concave continuous function of $\mathbf{w} \in \mathcal{W}$ for which $\partial u_{\mathcal{T}}(\mathbf{w})$ is a subgradient given by $\mathbf{f}(s_{\mathcal{T}}(\mathbf{w}))$ for all $\mathbf{w} \in \mathcal{W}$. Further, for all $\mathbf{w} \in \Omega$, $u_{\mathcal{T}}(\mathbf{w}) \leq \hat{u}(\mathbf{w})$.

Proof: Begin by observing that for any weights $\mathbf{w}_1, \mathbf{w}_2 \in \mathcal{W}$, it must hold that $u_{\mathcal{T}}(\mathbf{w}_1) \leq c(s_{\mathcal{T}}(\mathbf{w}_2), \mathbf{w}_1)$. Indeed, by the definition of $u_{\mathcal{T}}$, if the observation does not hold, then

$$u_{\mathcal{T}}(\mathbf{w}_1) = c(s_{\mathcal{T}}(\mathbf{w}_1), \mathbf{w}_1) > c(s_{\mathcal{T}}(\mathbf{w}_2), \mathbf{w}_1)$$

implying that $s_{\mathcal{T}}(\mathbf{w}_1)$ is not a minimizer of $c(s, \mathbf{w}_1)$ over $\mathcal{T}(\Omega)$ (since $s_{\mathcal{T}}(\mathbf{w}_2) \in \mathcal{T}(\Omega)$ by (13)), which is a contradiction of the definition of $s_{\mathcal{T}}(\mathbf{w}_1)$ from (13). Therefore, $u_{\mathcal{T}}(\mathbf{w}_1) \leq c(s_{\mathcal{T}}(\mathbf{w}_2), \mathbf{w}_1)$ for all $\mathbf{w}_1, \mathbf{w}_2 \in \mathcal{W}$. Next, we establish the concavity of $u_{\mathcal{T}}(\mathbf{w})$.

For any $\mathbf{w}_1, \mathbf{w}_2 \in \mathcal{W}$ and $\lambda_1 \in [0, 1], \lambda_2 = 1 - \lambda_1$, let $\mathbf{w} = \lambda_1 \mathbf{w}_1 + \lambda_2 \mathbf{w}_2$. From (13), and the definition of $u_{\mathcal{T}}$

$$\begin{aligned} u_{\mathcal{T}}(\mathbf{w}) &= c(s_{\mathcal{T}}(\mathbf{w}), \mathbf{w}) = \mathbf{w} \cdot \mathbf{f}(s_{\mathcal{T}}(\mathbf{w})) \\ &= \lambda_1 \mathbf{w}_1 \cdot \mathbf{f}(s_{\mathcal{T}}(\mathbf{w})) + \lambda_2 \mathbf{w}_2 \cdot \mathbf{f}(s_{\mathcal{T}}(\mathbf{w})) \\ &= \lambda_1 c(s_{\mathcal{T}}(\mathbf{w}), \mathbf{w}_1) + \lambda_2 c(s_{\mathcal{T}}(\mathbf{w}), \mathbf{w}_2) \\ &\geq \lambda_1 u_{\mathcal{T}}(\mathbf{w}_1) + \lambda_2 u_{\mathcal{T}}(\mathbf{w}_2) \end{aligned}$$

where the inequality holds by the observation made at the top of the proof. This establishes the concavity of $u_{\mathcal{T}}(\mathbf{w})$. The continuity of $u_{\mathcal{T}}(\mathbf{w})$ therefore follows from the proof of Lemma 1. Next, we show $\partial u_{\mathcal{T}}(\mathbf{w}^*) = \mathbf{f}(s_{\mathcal{T}}(\mathbf{w}^*))$, $\forall \mathbf{w}^* \in \mathcal{W}$. Let $\mathbf{w}', \mathbf{w}^* \in \mathcal{W}$, then from our first observation

$$\begin{aligned} u_{\mathcal{T}}(\mathbf{w}') - u_{\mathcal{T}}(\mathbf{w}^*) &= \mathbf{w}' \cdot \mathbf{f}(s_{\mathcal{T}}(\mathbf{w}')) - \mathbf{w}^* \cdot \mathbf{f}(s_{\mathcal{T}}(\mathbf{w}^*)) \\ &\leq \mathbf{w}' \cdot \mathbf{f}(s_{\mathcal{T}}(\mathbf{w}^*)) - \mathbf{w}^* \cdot \mathbf{f}(s_{\mathcal{T}}(\mathbf{w}^*)) \\ &= \mathbf{f}(s_{\mathcal{T}}(\mathbf{w}^*)) \cdot (\mathbf{w}' - \mathbf{w}^*) \end{aligned}$$

which establishes $\mathbf{f}(s_{\mathcal{T}}(\mathbf{w}^*))$ as a subgradient of $u_{\mathcal{T}}(\mathbf{w})$ at $\mathbf{w}^*$. Finally, we show that for all $\mathbf{w}' \in \mathcal{T}(\Omega)$, $u_{\mathcal{T}}(\mathbf{w}') \leq \hat{u}(\mathbf{w}')$. Since $\mathbf{w}' \in \Omega$, it must hold that $\hat{s}(\mathbf{w}') \in \mathcal{T}(\Omega)$ by the definition

Algorithm 2: MRPS WITH HEURISTIC SOLVER.

Input: A suboptimal $\hat{s}$ solver to approximately compute $s(\mathbf{w})$; a budget $K \geq n$

Output: Sampled weights Ω and maximum regret

```

1  $\Omega \leftarrow \{e^i, i = 1, \dots, n\}$ 
2  $\mathcal{T}(\Omega) \leftarrow \{\hat{s}(\mathbf{w}), \mathbf{w} \in \Omega\}$ 
3  $\mathcal{N} \leftarrow \{\Omega\}$ 
4 for  $k = n$  to  $K$  do
5    $N = \text{neighborhood in } \mathcal{N} \text{ with maximum } \bar{R}(N)$  //
   Here,  $\bar{R}(N)$  is computed from (10) with  $s^*$ 
   replaced with  $s_{\mathcal{T}}$  defined in (13)
6   if  $\bar{R}(N) = 0$  then
7      $N \leftarrow \text{Largest\_Neighbourhood}$ 
8      $\bar{\mathbf{w}} \leftarrow \text{Mean\_Weight}(N)$ 
9    $\Omega \leftarrow \Omega \cup \{\bar{\mathbf{w}}(N)\}$ 
10   $\mathcal{T}(\Omega) \leftarrow \mathcal{T}(\Omega) \cup \{\hat{s}(\bar{\mathbf{w}}(N))\}$ 
11  Lines 10-15 of Algorithm 1 with  $s_{\mathcal{T}}$  replacing  $s^*$ 
12 return  $\Omega$  and the maximum value of  $\bar{R}(N)$  over all
     $N \in \mathcal{N}$ 

```

of $\mathcal{T}(\Omega)$. Therefore, from (13) and the definition of $u_{\mathcal{T}}(\mathbf{w})$, it must hold that $u_{\mathcal{T}}(\mathbf{w}') = c(s_{\mathcal{T}}(\mathbf{w}'), \mathbf{w}') \leq c(\hat{s}(\mathbf{w}'), \mathbf{w}') = \hat{u}(\mathbf{w}')$ which completes the proof. ■

Lemma 3 implies that the cost function $u_{\mathcal{T}}(\mathbf{w})$ has the same properties as the optimal cost function $u(\mathbf{w})$ (established in Section III) that made Algorithm 1 possible.

B. Adapted Algorithm

Leveraging Lemma 3, we now propose minor modifications to Algorithm 1 that will allow it to accept suboptimal solvers. The procedure is outlined in Algorithm 2 and closely follows Algorithm 1. The primary modifications are as follows: In Line 2 of Algorithm 2, we initialize the set of discovered solutions $\mathcal{T}(\Omega)$. In Line 5, we replace the definition of $\bar{R}(N)$ for a neighborhood $N \in \mathcal{N}$ that is given in (10) with a version where $s^*(\mathbf{w})$ is replaced with $s_{\mathcal{T}}(\mathbf{w})$ given in (13) for any weight $\mathbf{w} \in \mathcal{W}$. Following the definitions in Section IV, recall that given a set of weights $\Omega \subseteq \mathcal{W}$ and a neighborhood N associated with weights in Ω , the function $\bar{R}(N)$, and its associated weight vector $\bar{\mathbf{w}}(N)$ are defined in (10) relative to a function $R(\mathbf{w}' | \mathbf{w}^*)$ for weights $\mathbf{w}', \mathbf{w}^* \in \mathcal{W}$. This latter function is defined as $R(\mathbf{w}' | \mathbf{w}^*) = \mathbf{w}^* \cdot \mathbf{f}(s^*(\mathbf{w}')) - P(\mathbf{w}^*)$ for a specific plane P defined in that section. It was shown in Theorem 2 that $\bar{R}(N)$ is an upper bound for the maximum regret in the neighborhood N given Ω . In Algorithm 2, we simply replace $s^*(\mathbf{w})$ with $s_{\mathcal{T}}(\mathbf{w})$ in all preceding function definitions. In a manner identical to the proof of Theorem 2, it is easily verified from the results of Lemma 3 that the modified function $\bar{R}(N)$ is an upper bound on the maximum best known regret $r_{\mathcal{T}}$ in the neighborhood N . Further, the values $\bar{R}(N)$, $\bar{\mathbf{w}}(N)$ may still be obtained by solving the linear program (12) with s^* replaced with $s_{\mathcal{T}}$.

Next, Lines 6–8 replace the conditional stopping criteria in Lines 6–7 of Algorithm 1. In practice, we have noticed that if the suboptimal solver $\hat{s}$ is a poor estimate of the optimal solution s^* ,

then using a conditional stopping criteria $\bar{R}(N) = 0$ for all $N \in \mathcal{N}$ tends to cause the Algorithm to terminate prematurely. This is because there may exist neighborhoods N and weights $\mathbf{w}$ in the convex hull of N such that $c(\hat{s}(\mathbf{w}), \mathbf{w}) < \min_{\mathbf{w}' \in N} u_{\mathcal{T}}(\mathbf{w}')$ which cannot happen when using an optimal solver. Finally, in Line 10, when a weight $\bar{\mathbf{w}}$ is added to Ω , its corresponding suboptimal solution $\hat{s}(\bar{\mathbf{w}})$ is added to $\mathcal{T}(\Omega)$. The remainder of the Algorithm is unchanged from Algorithm 1. Observe that if $\hat{s}(\mathbf{w}) = s^*(\mathbf{w})$ for all $\mathbf{w} \in \mathcal{W}$, that is, if the “suboptimal” solver is in fact optimal, then Algorithm 2 reduces to Algorithm 1 by construction.

C. Extended Error Bound

We conclude with an error bound guarantee in the case that solutions $\hat{s}(\mathbf{w})$ are derived from a suboptimal solver with a known approximation factor. Let Ω be a set of weight vectors computed by Algorithm 2, and let $\mathcal{N}$ be the set of partitioning neighborhoods associated with Ω . Finally, let

$$\beta = \max_{N \in \mathcal{N}} \max_{\mathbf{w} \in N} \frac{\bar{R}(N)}{u_{\mathcal{T}}(\mathbf{w})}. \quad (14)$$

Intuitively, β represents the maximum relative regret over all neighborhoods $N \in \mathcal{N}$. We write β as a ratio opposed to a difference as in (6) to derive an approximation factor. Observe that the value of β is easily computed since each neighborhood N is a discrete set of weights in Ω . Thus, for each $\mathbf{w} \in N$ and each $N \in \mathcal{N}$, $u_{\mathcal{T}}(\mathbf{w}) = \mathbf{w} \cdot \mathbf{f}(s_{\mathcal{T}}(\mathbf{w}))$. We offer the following result.

Theorem 3 (Approximation factor): Let Ω be the set of weights computed by Algorithm 2 with input solutions $\hat{s}(\mathbf{w})$ and β be given by (14). If the solutions $\hat{s}(\mathbf{w})$ are derived from a solver with known approximation factor α , then, for every $\mathbf{w}^* \in \mathcal{W}$, there exists a $\mathbf{w}' \in \Omega$ such that

$$c(s_{\mathcal{T}}(\mathbf{w}'), \mathbf{w}^*) \leq \alpha(\beta + 1)u(\mathbf{w}^*).$$

Proof: For any $\mathbf{w}^* \in \mathcal{W}$, there must exist a neighborhood N such that $\mathbf{w}^* \in \mathcal{C}(N)$ since $\mathcal{C}(N), N \in \mathcal{N}$ forms a partition of $\mathcal{W}$. Since $N \subseteq \Omega$ it must hold by Lemma 3 and the assumption that $\hat{s}(\mathbf{w})$ is derived from a solver that approximates $s^*(\mathbf{w})$ to within a factor of α of the resulting cost, that for all $\mathbf{w} \in N$, $u_{\mathcal{T}}(\mathbf{w}) \leq \hat{u}(\mathbf{w}) \leq \alpha u(\mathbf{w})$. Let P denote the hyperplane passing through $\{u(\mathbf{w}), \mathbf{w} \in N\}$ and $P_{\mathcal{T}}$ the hyperplane passing through $\{u_{\mathcal{T}}(\mathbf{w}), \mathbf{w} \in N\}$. Observe that since $u(\mathbf{w}) \geq \alpha^{-1}u_{\mathcal{T}}(\mathbf{w})$ for all $\mathbf{w} \in N$ and $P, P_{\mathcal{T}}$ are planes, it must hold that $\alpha P(\mathbf{w}) \geq P_{\mathcal{T}}(\mathbf{w}), \forall \mathbf{w} \in \mathcal{C}(N)$. From the definition of $\bar{R}(N)$, there must exist a weight $\mathbf{w}' \in N$ with $\mathbf{w}^* \cdot \mathbf{f}(s_{\mathcal{T}}(\mathbf{w}')) - P_{\mathcal{T}}(\mathbf{w}^*) \leq \bar{R}(N)$ implying that

$$\begin{aligned} \mathbf{w}^* \cdot \mathbf{f}(s_{\mathcal{T}}(\mathbf{w}')) &\leq \bar{R}(N) + P_{\mathcal{T}}(\mathbf{w}^*) \\ &\leq \bar{R}(N) + \alpha P(\mathbf{w}^*) \\ &\leq \bar{R}(N) + \alpha u(\mathbf{w}^*) \end{aligned} \quad (15)$$

where the final inequality holds from the concavity of $u(\mathbf{w})$ (see Observation 2) implying that $u(\mathbf{w})$ lies below the plane $P(\mathbf{w})$ for $\mathbf{w} \in \mathcal{C}(N)$. Letting $\mathbf{w}'' = \arg \min_{\mathbf{w} \in N} u(\mathbf{w})$, we observe by the concavity of $u(\mathbf{w}), u_{\mathcal{T}}(\mathbf{w})$ (see Lemma 3) that

$u(\mathbf{w}) \geq u(\mathbf{w}'')$ and $u_{\mathcal{T}}(\mathbf{w}) \geq \min_{\mathbf{w} \in N} u_{\mathcal{T}}(\mathbf{w})$. Since $u_{\mathcal{T}}(\mathbf{w})$ is within a factor of α of $u(\mathbf{w})$ for all $\mathbf{w} \in N$ and $\mathbf{w}'' \in N$ by definition, it must hold that $u(\mathbf{w}^*) \geq u(\mathbf{w}'') \geq \alpha^{-1}u_{\mathcal{T}}(\mathbf{w}'') \geq \alpha^{-1} \min_{\mathbf{w} \in N} u_{\mathcal{T}}(\mathbf{w})$. Thus

$$\frac{\bar{R}(N)}{u(\mathbf{w}^*)} \leq \alpha \frac{\bar{R}(N)}{\min_{\mathbf{w} \in N} u_{\mathcal{T}}(\mathbf{w})} = \alpha \max_{\mathbf{w} \in N} \frac{\bar{R}(N)}{u_{\mathcal{T}}(\mathbf{w})} \leq \alpha\beta. \quad (16)$$

Thus, from (15) and (16)

$$\frac{\mathbf{w}^* \cdot \mathbf{f}(s_{\mathcal{T}}(\mathbf{w}'))}{u(\mathbf{w}^*)} \leq \alpha(\beta + 1).$$

Finally, noting that $c(s_{\mathcal{T}}(\mathbf{w}'), \mathbf{w}^*) = \mathbf{w}^* \cdot \mathbf{f}(s_{\mathcal{T}}(\mathbf{w}'))$ completes the proof. ■

In this section, we extended the algorithm introduced in Section IV to include suboptimal solvers. Moreover, when the suboptimal solver is an approximation algorithm we retain a bound on the error of a solution set. In the following section, we consider another extension to include stochastic settings.

VI. EXTENSION TO STOCHASTIC SETTINGS

In this section, we consider two cases of stochastic inputs to the problem: First, there is a bias representing the likely relevance of some weights over others. Second, the constraints defining the planning problem are random, i.e., different instances have to be considered.

A. Weights as a Random Variable

In practice, some regions of the weight space $\mathcal{W}$ may correspond to Pareto-optimal solutions that are less relevant, e.g., there is prior information available on what weights may be likely to represent of a user's preference. In the example in Fig. 1, feasible solutions were trajectories between fixed start and goal configurations, while the objectives were trajectory length and comfort. It may be very unlikely for any user to *only* care about comfort.

In this section we consider that a probability density function (PDF) $g(\mathbf{w})$ over $\mathcal{W}$ is given. The problem then becomes one of computing a set of weights Ω that approximates solutions corresponding to *probable* weights, i.e., that minimizes the maximum regret *discounted* by the prior $g(\mathbf{w})$.

First, we introduce a notation shorthand: Given a set of sampled weights Ω and a weight $\mathbf{w}$ in $\mathcal{W}$, we define the regret of $\mathbf{w}$ using Ω as

$$r^{\Omega}(\mathbf{w}) = \min_{\mathbf{w}' \in \Omega} r(\mathbf{w}' | \mathbf{w}). \quad (17)$$

The goal of this section is to compute a set of weights that address Problem 1 while considering the prior belief expressed by $g(\mathbf{w})$. To this end, consider two optimization problems

$$\min_{|\Omega| \leq K} \int_{\mathcal{W}} r^{\Omega}(\mathbf{w}^*) g(\mathbf{w}^*) d\mathbf{w}^*, \quad (18a)$$

$$\min_{|\Omega| \leq K} \text{ess sup}_{\mathbf{w}^* \in \mathcal{W}} r^{\Omega}(\mathbf{w}^*) g(\mathbf{w}^*). \quad (18b)$$

We refer to the term $r^\Omega(\mathbf{w}^*)g(\mathbf{w}^*)$ as the *discounted regret*. The problem in (18a) seeks to compute a set Ω that minimizes the expected regret $\mathbb{E}[r^\Omega(\mathbf{w}^*)]$, while the second (18b) minimizes the *essential supremum* (supremum up to a set of measure zero) of the discounted regret.

While both problems have their applications, we focus on solving (18b). The main shortcoming of (18a) is that it can lead to placing samples to reduce the regret for a large set of weights with low probability. However, as we have seen in Fig. 1, there can be a large subset of samples that yields similar—or even identical—corresponding optimal solutions. Minimizing the expected regret can then overly favour minimizing the regret for a large subset of weights, regardless of the incurred regret.

B. Probabilistic Sampling Algorithm

We now adapt Algorithm 1 to solve the probabilistic problem posed in (18b). To this end, for any neighborhood N of weights with convex hull $\mathcal{C}(N)$, let

$$p(N) = \int_{\mathcal{C}(N)} g(\mathbf{w}) d\mathbf{w}$$

denote the probability that the random variable $\mathbf{w}^*$ lies in $\mathcal{C}(N)$. We propose a simple, yet effective change to Algorithm 1 in order to handle probabilities over weights, and denote the new algorithm as MRPS – P. To solve for the objective in (18b), we replace $\bar{R}(N)$ in Lines 5 and 16 of Algorithm 1 with $\bar{R}^p(N) = p(N)\bar{R}(N)$. That is, the neighborhood N we select to place an additional sampled weight is the one with the largest regret bound $\bar{R}(N)$, *discounted* by the probability $p(N)$ that $\mathcal{C}(N)$ contains $\mathbf{w}^*$. In this way, we avoid adding weights to Ω that lie in regions that are unlikely to contain $\mathbf{w}^*$ unless the regret of not including such a weight is sufficiently large. Letting $\Omega, \bar{R}^p$ denote the outputs of the modified version of Algorithm 1, we offer the following results with regard to the objective (18b):

Lemma 4 (Worst Case Discounted Regret): The maximum discounted regret $\bar{R}^p$ returned by Algorithm MRPS – P is an upper bound on the objective in (18b)

$$\text{ess sup}_{\mathbf{w}^* \in \mathcal{W}} r^\Omega(\mathbf{w}^*)g(\mathbf{w}^*) \leq \bar{R}^p. \quad (19)$$

Proof: Since $\mathcal{N}$ forms a partition of $\mathcal{W}$ (Section IV-B), there exists a neighborhood $N \in \mathcal{N}$ with $\mathbf{w}^* \in \mathcal{C}(N)$. Therefore

$$\begin{aligned} r^\Omega(\mathbf{w}^*)g(\mathbf{w}^*) &\leq \int_{\mathcal{C}(N)} r^\Omega(\mathbf{w})g(\mathbf{w})d\mathbf{w} \\ &\leq \bar{R}(N) \int_{\mathcal{C}(N)} g(\mathbf{w})d\mathbf{w} = \bar{R}^p(N) \leq \bar{R}^p. \end{aligned} \quad (20)$$

The first inequality holds by Theorem 2, while the final inequality holds since modified Algorithm 1 returns the maximum value $\bar{R}^p$ over all neighborhoods in $\mathcal{N}$. ■

In addition to providing a bound on the objective of (18b), we observe that $\bar{R}^p$ also provides a bound on the objective of (18a). Indeed, by Lemma 4

$$\mathbb{E}[r^\Omega(\mathbf{w}^*)] = \int_{\mathcal{W}} r^\Omega(\mathbf{w})g(\mathbf{w})d\mathbf{w} \leq \bar{R}^p \int_{\mathcal{W}} d\mathbf{w} = \bar{R}^p.$$

Thus, with a simple modification, we can adapt the MRPS algorithm to incorporate prior beliefs over weights. The algorithm then greedily minimizes an upper bound of the discounted regret. Observe that if the PDF is uniform, i.e., if all weights are equally likely to represent a user, then the proposed extension to the MRPS algorithm, reduces to the original MRPS algorithm. In the next section, we extend Algorithm 1 to account for multiple problem instances.

C. Constraints as a Random Variable

As a final extension of Algorithm 1, we treat the constraints $\mathcal{S}$ in (1) as random variables. This encompasses the case where multiple instances of (3) are possible. Consider, e.g., the problem detailed in Fig. 1. Here, we compute trajectories between a *fixed* start and goal to minimize a tradeoff between travel time and discomfort. However, changing the positions of the start and goal would result in a different instance of problem (3) with a (potentially) different set of weights Ω computed by Algorithm 1.

This may seem to limit the applicability of the approach presented here. However, there is a simple extension. We define a random variable I existing in a sample space $\mathcal{I}$ which represents a possible instance of the constraints $\mathcal{S}$. Defining $s_I^*(\mathbf{w})$ as an optimal solution (5) with $\mathcal{S} = I$ given fixed weights $\mathbf{w}$, we may define a vector of hyperobjectives $\mathbf{f}^E(\mathbf{w}) = \mathbb{E}_{\mathcal{I}}[\mathbf{f}(s_I^*(\mathbf{w}))]$. Similarly, we define $u^E(\mathbf{w}) = \mathbf{w} \cdot \mathbf{f}^E(\mathbf{w})$. It is easily verified using an identical strategy to the proofs presented in Section III that $u^E(\mathbf{w})$ possesses all of the properties (continuity, concavity, $\partial u(\mathbf{w}) = \mathbf{f}^E(\mathbf{w})$ for all $\mathbf{w} \in \mathcal{W}$) that make Algorithm 1 possible. Further, the function $u^E(\mathbf{w})$ is precisely the expected optimal cost $\mathbb{E}_{\mathcal{I}}[u_I(\mathbf{w})]$ for $u_I(\mathbf{w}) = \mathbf{w} \cdot \mathbf{f}(s_I^*(\mathbf{w}))$. Indeed for any weight $\mathbf{w} \in \mathcal{W}$

$$\mathbb{E}_{\mathcal{I}}[u_I(\mathbf{w})] = \mathbf{w} \cdot \mathbb{E}_{\mathcal{I}}[\mathbf{f}(s_I^*(\mathbf{w}))] = \mathbf{w} \cdot \mathbf{f}^E(\mathbf{w}) = u^E(\mathbf{w}).$$

This implies that if $r^E(\mathbf{w}'|\mathbf{w}^*) = \mathbf{w}^* \cdot \mathbf{f}^E(\mathbf{w}') - u^E(\mathbf{w}^*)$ for $\mathbf{w}', \mathbf{w}^* \in \mathcal{W}$ then $r^E(\mathbf{w}'|\mathbf{w}^*)$ is precisely the expected regret $\mathbb{E}_{\mathcal{I}}[\mathbf{w}^* \cdot \mathbf{f}(s_I^*(\mathbf{w}')) - u_I(\mathbf{w}^*)]$. Thus, by replacing $\mathbf{f}(s^*(\mathbf{w}))$ with $\mathbf{f}^E(\mathbf{w})$ in Algorithm 1, we can obtain a set Ω with bounded maximum *expected* regret over all constraints $\mathcal{I}$.

VII. NUMERICAL RESULTS

We demonstrate our algorithm in simulations for two different applications: 1) trajectory planning; and 2) multirobot coordination. We consider three variations of our problem: 1) optimal solvers are available (Section IV); 2) only a suboptimal solver is available (Section V); and 3) a prior belief over relevant weights is given (Section VI). These experiments illustrate how the proposed methods allow for computing presampled sets of weights and their solutions that closely approximate solutions to any weight requested online. In particular, we report the worst-case approximation error. In an additional experiment, we investigate how the proposed method can be used for presampling solution in order to learn user preferences.

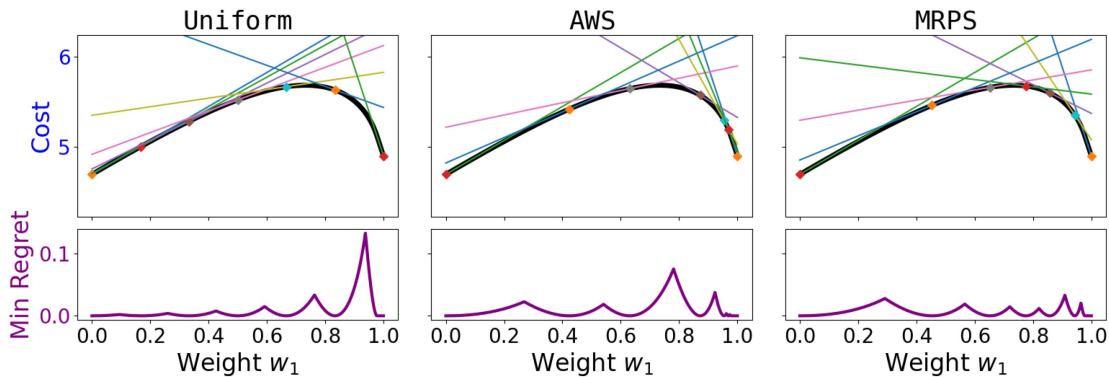


Fig. 5. Example results for the Dubins planning problem with two objectives. Shown are the approximations of $u(w)$ and resulting regret with **Uniform**, adaptive sampling (**AWS**), and the proposed approach **MRPS**.

A. Simulation Setup

1) *Planning Problems*: The first planning problem involves computing Dubins trajectories for different objectives, similar to the example shown in Fig. 1. In the two-objective case we consider the competing objectives of trajectory length and integral square jerk (a common metric of comfort [2], [17]). For higher dimensions we additionally consider the maximum jerk, as well as avoiding high-risk areas of the environment. Given a weight vector, we compute the Dubins' trajectory that optimizes the resulting cost function numerically. In detail, because Dubins' trajectories are easily computed for a fixed minimum turning radius, we iterate over small increments in turning radius within given bounds selecting the one with minimum cost. Though the result is suboptimal, it can be made arbitrarily close to optimal by increasing the resolution; our experiments used 10 000 steps. Thus, we consider this approach to satisfy Assumption 1. The second planning problem is the mTSP with deadlines [48]. We consider two objectives: 1) the number of vertices visited before their deadline; and 2) the total distance travelled by all robots. Since the problem is NP-hard, we use it to highlight the extension of our work to suboptimal solvers in Section V.

2) *Evaluation Measures*: For evaluation we rely on a large (1000) set of uniformly randomly sampled weights. We evaluate algorithm performance using the regret as defined in (6), as well as the relative regret where the difference in (6) is replaced with a ratio. For the experiments with prior distributions over weights, we consider the essential supremum of the discounted regret and the expected regret from (18b) and (18a), respectively.

3) *Baseline Algorithms*: We compare our proposed algorithm against three baselines: 1) Uniform sampling in the weight space [15], [49], [50], denoted by **Uniform**; 2) sampling weights from a prior distribution for the stochastic problem settings, denoted by **Prior**; and 3) the adaptive weight sampling [19], [40], denoted as **AWS**. We refer to our proposed approach from Algorithm 1 as **MRPS**, and its modification to suboptimal solvers as **MRPS-S** and considering probabilistic weights as **MRPS-P**.

4) *Sampling Budgets*: We test the different algorithms with different budgets K for the number of weight samples. Since our algorithm initially computes the n single-objective solutions, we only consider $K \geq n$.

B. Experiments With Optimal Solvers

We begin with planning Dubins trajectories as shown in Fig. 1. To find solutions $s^*(w)$, the motion planner can sample a large set of Dubins trajectories using different turn radii and pick the optimum among these.

1) *Illustrative Example*: First, we present more insight into the example from Fig. 1 with $K = n + 5 = 7$ samples. We notice that **MRPS** produces a larger variety of sample trajectories, especially those with shorter length. This is also visualized in the approximations of the Pareto front: **Uniform** exhibits a large gap, while the proposed method places samples more homogeneously. In Fig. 5 we show the optimal cost $u(w)$ (ground truth computed with 10 000 uniform weights), together with the tangent planes of the approximating samples of both baselines **Uniform** and **AWS** and our proposed approach. Here, **Uniform** places weights – and thus tangents – in an equal distance along the x -axis from one another. This results in numerous samples with similar tangents on the left side of the plot where the function $u(w)$ is almost linear. At the right end where $u(w)$ changes more rapidly, uniform does not have sufficient samples for a tight approximation. In contrast, **MRPS** places more samples at the right end resulting in a smaller gap between the best approximating tangent and the $u(w)$, such that the maximum regret is $\approx 1/10$ compared to **Uniform**. While **AWS** improves over **Uniform**, the approximation is not as tight as **MRPS**, resulting in a maximum regret that is still twice as large.

2) *Quantitative Analysis*: We repeat the above Dubins planning experiment with randomized goal locations and various sampling budgets K . Results are shown in Fig. 6. When $K = n$ only the basis solutions $e^1, \dots, e^n$, i.e., the single objective solutions, are available. We observe that **MRPS** achieves substantially smaller absolute and relative regret values for all $K > n$ compared to **Uniform**. With just 3 samples, **MRPS** achieves a performance comparable to **Uniform** using 10 samples, and for 20 samples **MRPS** provides approximations with effectively no regret. The **AWS** approach performs similar to **MRPS** for few samples ($K = n + 1$ and $K = n + 3$), yet makes only very little progress for larger K . In summary, the experiment shows that the proposed **MRPS** method is able to efficiently place samples that

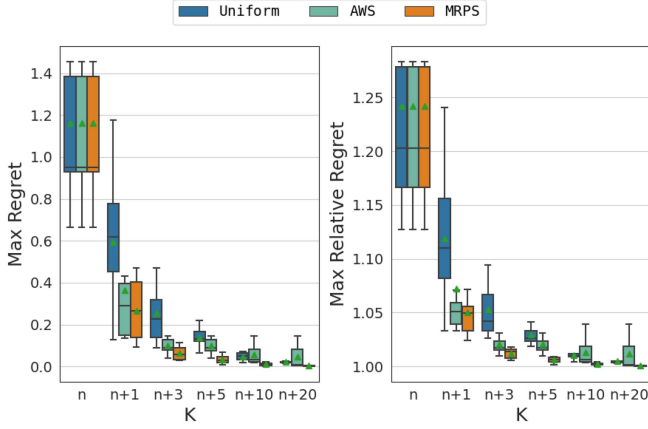
Fig. 6. Results for the Dubins experiment with $n = 2$ objectives.

TABLE I

MEAN AND 95TH PERCENTILE OF THE MAXIMUM REGRET FOR THE DUBINS PLANNING PROBLEM WITH AN OPTIMAL SOLVER

Solver	Budget K				
	$n + 1$	$n + 3$	$n + 5$	$n + 10$	$n + 20$
Uniform	0.59/1.08	0.25/0.44	0.14/0.2	0.05/0.07	0.02/0.03
AWS	0.37/0.91	0.10/0.14	0.10/0.14	0.06/0.14	0.05/0.14
MRPS	0.26/0.45	0.06/0.11	0.03/0.07	0.01/0.02	0.00/0.01

(a) Maximum regret for $n = 2$ objectives.

Solver	Budget K				
	$n + 1$	$n + 3$	$n + 5$	$n + 10$	$n + 20$
Uniform	1.14/1.7	0.67/1.03	0.66/1.03	0.65/1.03	0.51/0.88
AWS	0.83/1.41	0.39/0.63	0.25/0.58	0.22/0.58	0.22/0.58
MRPS	0.76/1.61	0.15/0.25	0.13/0.22	0.07/0.11	0.03/0.05

(b) Maximum regret for $n = 3$ objectives.

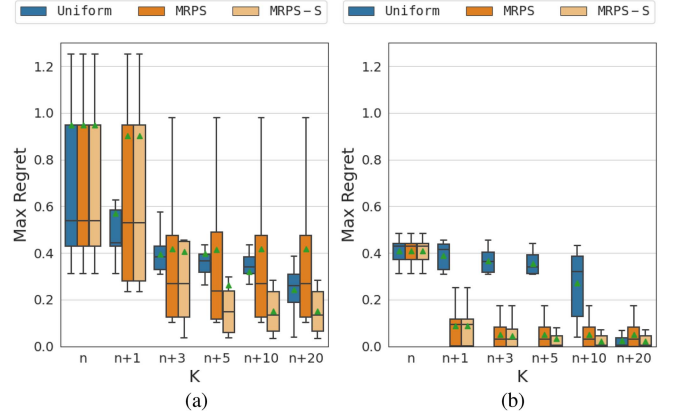
Solver	Budget K				
	$n + 1$	$n + 3$	$n + 5$	$n + 10$	$n + 20$
Uniform	0.81/1.33	0.8/1.31	0.71/1.31	0.48/1.04	0.34/0.85
AWS	0.78/1.31	0.65/1.31	0.62/1.31	0.46/1.17	0.44/1.17
MRPS	0.58/1.09	0.2/0.34	0.15/0.28	0.1/0.2	0.05/0.12

(c) Maximum regret for $n = 4$ objectives.

allow for minimum regret approximations for any scalarization weight.

3) *Varying Number of Objectives*: We also considered problem variances with 3 or 4 objectives with detailed results shown in Table I. As expected, with increasing number of objectives the regret increases (with exception of MRPS for $K = n + 1$). Nonetheless, MRPS achieves the best performance under all settings. In particular, MRPS continues to decrease the regret for larger K almost converging to an optimal set of samples with zero regret. In contrast, Uniform and AWS make only little progress for $K > n + 5$ in the three objective setup.

In summary, this experiment showed that, given an optimal solver, Algorithm 1 solves Problem 1 for varying number of objectives, strongly outperforming baseline methods.

Fig. 7. Results for the mTSP experiment with $n = 2$ objectives, 20 vertices, and 20 robots with different suboptimal LNS heuristics. (a) Cheap heuristic (10 iterations). (b) Strong heuristic (1000 iterations).

C. Experiments With Suboptimal Solvers

Next we conduct experiments when no optimal solver is available, as discussed in Section V. The planning problem is a mTSP with deadlines. Suboptimal solutions are computed using a Large Neighbourhood Search (LNS) heuristic [51] running for with varying computation budget, i.e., number of iterations. We randomly generate 10 different problem instances with 20 vertices and 20 robots, all starting at a central depot. Due to the poor scalability of exact solvers for the problem, we do not have access to the ground truth $u(w)$. Instead, we compute the regret with respect to solutions found by the LNS heuristic with 10 000 iterations.

Fig. 7 shows the results, comparing Uniform with MRPS and the modification MRPS – S proposed in Section V. Given the suboptimal solver for mTSP, the assumption made in MRPS are not satisfied. The algorithm can still be executed, yet it may prematurely determine that it is unable to make further progress and thus terminate. This experiment thus highlights the benefit of MRPS – S when no exact solver is available. The AWS approach cannot be directly incorporated into the LNS solver since it poses an equality constraint on the solution vector. Such a constraint raises issues of feasibility particularly in discrete optimization problems like mTSP. Thus, AWS is omitted from this experiment. Further, we compare algorithm performance when different solvers for the mTSP are available: We consider a *strong heuristic*, where we run LNS with 1000 iterations, and a *cheap heuristic*, using LNS with only 10 iterations.

First, we notice that the different version of the heuristic directly influence the solutions: The regret is substantially smaller for the stronger heuristic using 1000 iterations. Yet, under both settings we observe that MRPS does not make any more progress after three iterations. Since Algorithm 1 depends on having access to an optimal solver, it terminates prematurely. For both variants of the heuristic, Uniform eventually outperforms MRPS. However, we observe that MRPS – S avoids the pitfalls of MRPS and strongly outperforms both, MRPS and Uniform. We notice that due to the suboptimal solver, none of the methods converge to a regret of zero. While Uniform eventually

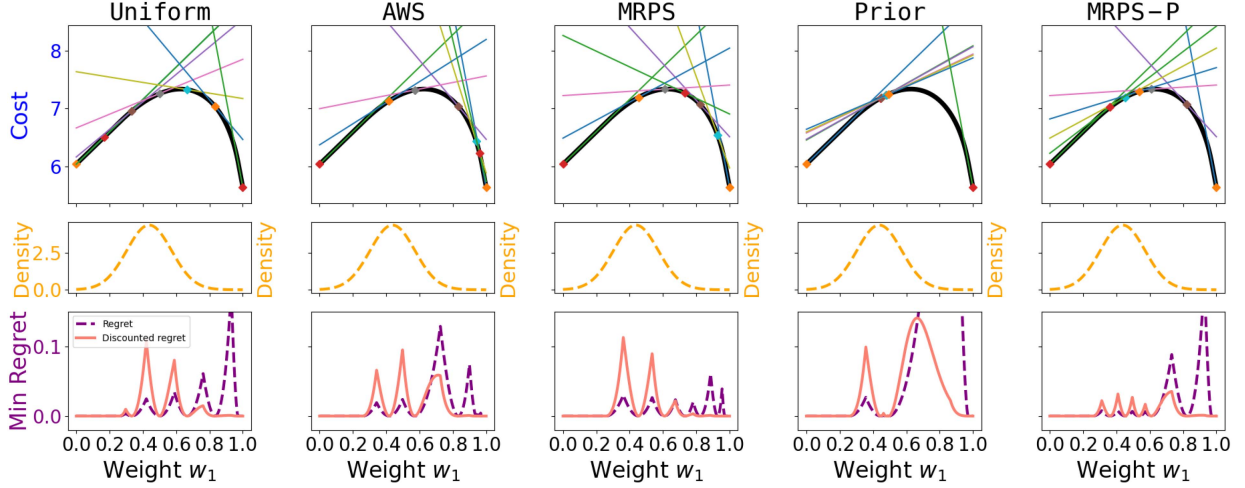


Fig. 8. Pareto sampling with prior distributions over weights. We compare the original algorithm MRPS with the adaptation MRPS – P proposed in Section VI-A. *Top row*: Objective function $u(w)$ and the sample points placed by the respective algorithms. *Middle row*: Prior belief. *Bottom row*: Regret and discounted regret incurred by the solutions.

ties with MRPS – S, our method remains much more efficient: For the strong heuristic, Uniform requires $n + 20$ samples to achieve the same regret as MRPS – S already has with $n + 5$ samples. Lastly, the performance gap of MRPS – S (and MRPS) compared to Uniform is larger for the stronger heuristic. Since our method relies on the solution vectors of previous solutions, it is misguided when the solver returns poor solutions. Thus, for unreliable solvers, Uniform can be more robust.

Overall, the experiment has shown that: 1) the original approach MRPS is challenged when a suboptimal solver is used, eventually performing poorer than uniform sampling. This is particularly apparent when the solver is a cheap heuristic; 2) The extension MRPS – S addresses this shortcoming and is able to find strong solutions within few iterations, strongly outperforming greedy for cheap and strong heuristics.

D. Experiments With Stochastic Problem Inputs

We now consider the setup from Section VI-A where a prior belief over weights is available. The objective is to minimize the essential supremum of the discounted regret formulated in (18b). We consider priors $g(w)$ in the form of a normal distribution, with uniformly random mean from the weight space $\mathcal{W}$, and uniformly random variance within $[.01, .1]^n$.

1) *Illustrative Example*: Fig. 8 showcases the difference between the different approaches when planning Dubins trajectories with two objectives. We observe that for all approaches, the maximum regret and the maximum discounted regret occurs at different weights. While MRPS achieves the lowest regret, its discounted regret is largest since it neglects to sample around weights with a high prior belief. In contrast, the modified approach MRPS – P places fewer samples at the right end, trading-off a higher regret for a smaller discounted regret. In comparison, directly sampling from the prior belief (labelled Prior) does not yield the desired result: Here samples are too concentrated

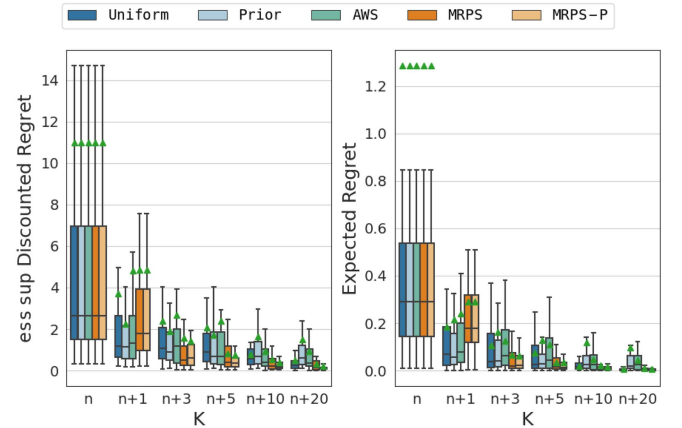


Fig. 9. Results for the Dubins experiment with $n = 3$ objectives and prior distributions over weights.

around the distribution mean and does incur a higher maximum discounted regret.

2) *Quantitative Results*: Fig. 9 shows detailed results for varying numbers of samples K for Dubins trajectories with three objectives. We include the maximum discounted regret and expected regret corresponding to the problems formulated in (18a) and (18b), respectively. Table II additionally contains the discounted regret for 2 and 4 objectives.

In Fig. 9, we observe that MRPS – P does not perform well with only $K = n + 1$ sample. Indeed, in the first iteration it is equivalent to MRPS and does not consider the prior belief since there is only one neighbourhood to explore. Yet, for $K = n + 3$ MRPS – P shows already the lowest essential supremum of the discounted regret, and approaches 0 for $K = n + 10$. The original method MRPS achieves a similarly strong performance for only $K = n + 20$. Thus, when the sampling budget is limited, MRPS – P performs better. Among the baselines Uniform

TABLE II
MEAN AND 95TH PERCENTILE OF THE MAXIMUM DISCOUNTED REGRET FOR THE DUBINS PLANNING PROBLEM WITH AN OPTIMAL SOLVER AND PRIOR DISTRIBUTIONS OVER WEIGHTS

Solver	Budget K				
	$n + 1$	$n + 3$	$n + 5$	$n + 10$	$n + 20$
Uni.	0.36/1.25	0.20/0.74	0.10/0.33	0.04/0.12	0.01/0.05
AWS	0.86/3.64	0.13/0.30	0.09/0.25	0.05/0.22	0.03/0.22
MRPS	0.32/0.87	0.12/0.31	0.05/0.16	0.01/0.03	0.00/0.01
Prior	0.35/1.31	0.24/0.85	0.24/0.85	0.17/0.53	0.17/0.53
MRPS-P	0.32/0.87	0.06/0.17	0.03/0.08	0.01/0.02	0.00/0.01

(a) Maximum discounted regret for $n = 2$ Objectives.

Solver	Budget K				
	$n + 1$	$n + 3$	$n + 5$	$n + 10$	$n + 20$
Uni.	0.55/1.70	0.19/0.47	0.14/0.39	0.11/0.30	0.07/0.20
AWS	0.49/1.67	0.39/1.40	0.39/1.40	0.39/1.40	0.39/1.40
MRPS	1.13/3.29	0.31/0.85	0.18/0.57	0.10/0.24	0.04/0.12
Prior	0.26/0.66	0.25/0.66	0.23/0.55	0.23/0.55	0.21/0.55
MRPS-P	1.13/3.29	0.20/0.58	0.10/0.24	0.05/0.17	0.03/0.08

(b) Maximum discounted regret for $n = 3$ Objectives.

Solver	Budget K				
	$n + 1$	$n + 3$	$n + 5$	$n + 10$	$n + 20$
Uni.	3.71/7.04	2.39/7.04	2.09/4.38	0.80/2.02	0.49/1.07
AWS	4.83/13.96	2.68/13.96	2.38/13.96	0.94/2.90	0.93/2.90
MRPS	4.85/13.96	1.56/4.93	0.84/3.21	0.56/1.73	0.34/1.52
Prior	2.26/7.00	1.90/7.00	1.71/6.00	1.65/6.00	1.49/5.89
MRPS-P	4.85/13.96	1.42/1.92	0.74/1.52	0.39/1.13	0.21/0.58

(c) Maximum discounted regret for $n = 4$ Objectives.

performs best. As highlighted in the example in Fig. 8, only sampling from the prior belief (Prior) does not explore different weights efficiently and thus only makes little progress after the first few iterations. In the right plot, we illustrate the expected regret. Overall, we observe a similar trend as for the maximum discounted regret, yet Uniform, MRPS also approach 0 for $K = n + 20$. Yet, MRPS - P shows the strongest benefit for $K = n + 3$ and $K = n + 5$, highlighting that while MRPS - P is designed to solve (18a), it also bounds the error for (18b) and thus is suitable for tackling both problems. For completeness, we include results for AWS. Since this approach does not take the prior into account, it is unable to effectively progress after a few iterations and thus not suitable for this problem without further adaptation.

Overall, the third experiment showed that MRPS - P effectively adapts the original algorithm to the setting where a prior belief over weights is given. For the two objectives—the maximum discounted regret and the expected regret—MRPS - P proved to be more sample efficient than the baseline methods and the original MRPS algorithm. In summary, throughout the three different experiments we demonstrated that our proposed algorithm and its two extensions provide are able to approximate the set of Pareto-optimal solutions for different problem variants.

E. Computation Times

We briefly report the practical computation time for a Python implementation run on a I7-10750 H 2.6 GHz with 32-GB RAM.

Computing $K = n + 10$ samples with MRPS or MRPS - P for the Dubins problem runs within $\approx .4s$ for two objectives, and within $\approx .5s$ for four objectives, on average. For the mTSP problem the computation time of MRPS - S for $K = n + 10$ samples is $\approx 1.1s$ for the cheap heuristic, and $\approx 17s$ for the strong heuristic.

F. Reward Learning

To illustrate the practical impact of the proposed method, we consider the problem of learning user preferences, i.e., learning a user specific, but hidden weight vector w^* . Our algorithm serves a preprocessing step to generate a ground set of potential robot behaviors, from which we then elicit the solution that best fits a user's preferences. As mode of user interaction, we consider learning from choice [7], [12], [15], [16], [50] where the user is iteratively queried with two potential robot trajectories and indicates the preferred one. Repeating this over multiple iterations allows the robot infer about the user weights w^* . Most algorithms for this problem require a set of presampled trajectories from which the best query is selected using some heuristic [12], [15], [16], [50]. These presamples are either random trajectories [12], [16], or optimal solutions for uniformly random weights for the LSMOP [15], [35], [50].

Our proposed algorithm MRPS can be used to generate pre-samples for these learning problems. Thus, we compare the learning progress over 10 iterations when either using Uniform or MRPS samples. For a clear comparison, we use a simple, deterministic user model: Presented with trajectories A and B , they choose A if and only if $f(A)w^* \leq f(B)w^*$. Similar to [7] and [46], the trajectory preferred in the previous iteration constitutes one of the two trajectories presented in the next. The robot can *actively* choose the second trajectory to be presented in the next iteration. We employ a random selection from the presampled set (Random), or the minmax regret approach from [15] (Regret). Finally, generating random user weights w^* is not trivial: When w^* is drawn uniformly random, the sample set Uniform comes from the same distribution, biasing the experiment. Thus, we randomly select users from the union of two sample sets $\Omega(\text{Uniform})$ and $\Omega(\text{MRPS})$ each of size $K = 20$. We evaluate learning performance using the relative error $f(S^*(w')) \cdot w^* / u(w^*)$, where w' is the expected user weight. This measure is widely used in reward learning problems [36] and captures how well the learned cost function approximates the unknown user cost function (parameterized by w^*) in terms of the quality of the corresponding solutions.

We test this learning framework using the four feature Dubins planning problem from earlier, with one fixed goal location. Fig. 10 shows the result for learning with presampled sets of various sizes K . We observe that the MRPS samples lead to a smaller learning error than Uniform samples, regardless of the query method (Random or Regret). Indeed, the Uniform sets do not always include a close-to-optimal sample such that the learning is eventually unable to make progress. That is the case when w^* was drawn from the MRPS samples. However, the opposite effect is negligible: Learning with the MRPS samples finds close-to-optimal solutions, implying that the error is very small even when w^* comes from Uniform.

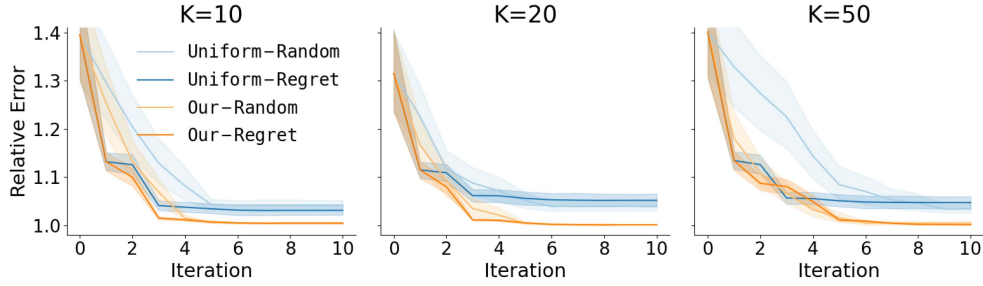


Fig. 10. Learning from choice with presampled sets of different size K .

When comparing different sizes K of the sample sets, we observe that all approaches learn slightly slower for larger K —the increased number of available trajectories seems to rather distract the learning algorithm than offering more informative queries. More surprisingly, when using Uniform samples, the learning still stops making progress. This indicates that the larger set still does not contain close-to-optimal solutions. This further supports our earlier findings that while MRPS only needs small K to find a close-to-optimal sample for any w^* , while Uniform is unable to achieve the same even for large K .

In summary, the experiment shows that using MRPS to generate presampled solutions in reward learning allows for learning close-to-optimal solutions with significantly fewer samples than when relying on randomly generated presamples.

VIII. SUMMARY

In this article, we studied the problem of computing different tradeoffs between competing objectives in robot planning problems. Commonly, such problems are approached via linear scalarization where multiple objectives are combined into a single objective taking the form of a weighted sum. Thus, we focused on computing a finite set of weights and corresponding solutions such that any other solution can be approximated with minimal regret. We studied fundamental properties of this linear objective function and its relation to regret. Assuming that we have access to an optimal solver for the scalarized objective, we presented an iterative sampling algorithm that repeatedly adds weights where the regret of the current set is largest and return a tight error bound. Next, we extended the algorithm to accept suboptimal solvers, and retained the error bound when approximation factors are provided. In a second extension, we also considered prior beliefs over practically relevant scalarization weights and adapted our algorithm to minimize the discounted maximum regret. In a series of simulations, we showcased the proposed methods for the three different problem variants, demonstrating their higher sampling efficiency compared to baselines. In a further experiment, we highlighted the practical effect of well-designed sample sets when learning user preferences for robot behavior.

IX. DISCUSSION AND FUTURE WORK

We extended our earlier results [2] eliminating the restrictive assumption that an optimal solver is available and also extended

the framework to stochastic settings. This highly increases the practical value of the proposed method. However, in this article we focused on establishing theoretical results and limited the evaluation to a series of simulation experiments. Future work should investigate the practical benefits of our method in two different settings: Many planners require careful parameter tuning. For any objective using a weighted sum, the presented method can be used to efficiently explore different parameter settings, especially when the tuning process is sensitive. The presented technique may also be adapted to tuning loss functions for machine learning systems. The other application is learning user preferences for robot behavior. We have shown that our method yields samples that allow for learning user preferences more accurately and more efficiently. The practical benefits should be further investigate considering different modes of user interaction and need practical verification in user studies. Finally, our analysis is limited to linear scalarization. If the Pareto front of the underlying MOP is nonconvex, linear scalarization fails to capture all Pareto-optimal solutions, i.e., is not Pareto-complete. However, many of the theoretical results presented here can be extended to any scalarization technique that is concave in the weight space. Thus, future work should consider how our proposed algorithm may be adapted to techniques such as Chebyshev-scalarization to ensure Pareto-completeness.

APPENDIX

A. Implementation of Max-Min Neighbourhood Regret

We detail the implementation of the linear program in (12). We formulate the convex hull constraint $w \in \mathcal{C}(N)$ using a scalars $\lambda^1, \dots, \lambda^n \in [0, 1]$ to write w as a convex combination of the neighbourhood weights $w = \lambda^1 w^1 + \dots + \lambda^n w^n$. The equality constraints are given by

$$\begin{bmatrix} 1 & \dots & 1 & 0 & \dots & 0 \\ 0 & \dots & 0 & 1 & \dots & 1 \\ -1 & 0 & \dots & 0 & w_1^1 & \dots & w_1^n \\ 0 & -1 & \dots & 0 & w_2^1 & \dots & w_2^n \\ \vdots & \ddots & \vdots & \vdots & \ddots & \ddots & \vdots \\ 0 & \dots & -1 & w_n^1 & \dots & w_n^n \end{bmatrix} \begin{bmatrix} w_1 \\ \vdots \\ w_n \\ \lambda^1 \\ \vdots \\ \lambda^n \end{bmatrix} = \begin{bmatrix} 1 \\ 1 \\ 0 \\ \vdots \\ 0 \end{bmatrix}. \quad (21)$$

The first row ensures that w lies in $\mathcal{W}$ (i.e., all its components sum to 1). The second ensures the same for $\lambda^1, \dots, \lambda^n$. The other rows ensure that the i th element of the vector w is a convex combination of the i th component of all neighbourhood vectors $w^1, \dots, w^n$. In the objective function, we write $P(w)$ using the same convex combination as

$$x - \sum_{i=1}^n \lambda^i u(w^i). \quad (22)$$

Finally, we require that $w_i \geq 0$ and $\lambda^i \geq 0$ for all $i = 1, \dots, n$.

REFERENCES

- [1] T. Gu, J. Atwood, C. Dong, J. M. Dolan, and J.-W. Lee, "Tunable and stable real-time trajectory planning for urban autonomous driving," in *Proc. IEEE/RSJ Int. Conf. Intell. Robots Syst.*, 2015, pp. 250–256.
- [2] A. Botros and S. L. Smith, "Tunable trajectory planner using G3 curves," *IEEE Trans. Intell. Veh.*, vol. 7, no. 2, pp. 273–285, Jun. 2022.
- [3] F. Christianos, P. Karkus, B. Ivanovic, S. V. Albrecht, and M. Pavone, "Planning with occluded traffic agents using Bi-level variational occlusion models," in *Proc. IEEE Int. Conf. Robot. Automat.*, 2023, pp. 5558–5565.
- [4] Z. Zuo et al., "MPC-based cooperative control strategy of path planning and trajectory tracking for intelligent vehicles," *IEEE Trans. Intell. Veh.*, vol. 6, no. 3, pp. 513–522, Sep. 2021.
- [5] P. Hang, C. Lv, C. Huang, J. Cai, Z. Hu, and Y. Xing, "An integrated framework of decision making and motion planning for autonomous vehicles considering social behaviors," *IEEE Trans. Veh. Technol.*, vol. 69, no. 12, pp. 14458–14469, Dec. 2020.
- [6] A. Bajcsy, D. P. Losey, M. K. O'Malley, and A. D. Dragan, "Learning robot objectives from physical human interaction," in *Proc. Conf. Robot Learn.*, PMLR, 2017, pp. 217–226.
- [7] N. Wilde, A. Blidaru, S. L. Smith, and D. Kulić, "Improving user specifications for robot behavior through active preference learning: Framework and evaluation," *Int. J. Robot. Res.*, vol. 39, no. 6, pp. 651–667, 2020.
- [8] M. Cáp and J. A. Mora, "Multi-objective analysis of ridesharing in automated mobility-on-demand," in *Proc. Robot.-Sci. Syst.*, 2018.
- [9] P. Karkus, B. Ivanovic, S. Mannor, and M. Pavone, "Diffstack: A differentiable and modular control stack for autonomous vehicles," in *Proc. Conf. Robot Learn.*, 2023, pp. 2170–2180.
- [10] R. T. Marler and J. S. Arora, "The weighted sum method for multi-objective optimization: New insights," *Struct. Multidisciplinary Optim.*, vol. 41, no. 6, pp. 853–862, 2010.
- [11] C.-L. Hwang and A. S. M. Masud, *Multiple Objective Decision Making—methods and Applications: A State-of-the-Art Survey*, vol. 164. Berlin, Germany: Springer Science & Business Media, 2012.
- [12] D. Sadigh, A. D. Dragan, S. S. Sastry, and S. A. Seshia, "Active preference-based learning of reward functions," in *Proc. Robot.-Sci. Syst.*, 2017.
- [13] J. Y. Zhang and A. D. Dragan, "Learning from extrapolated corrections," in *Proc. IEEE Int. Conf. Robot. Automat.*, 2019, pp. 7034–7040.
- [14] R. Holladay, S. Javdani, A. Dragan, and S. Srinivasa, "Active comparison based learning incorporating user uncertainty and noise," in *Proc. RSS Workshop Model Learn. Hum.-Robot Commun.*, 2016.
- [15] N. Wilde, D. Kulić, and S. L. Smith, "Active preference learning using maximum regret," in *Proc. IEEE/RSJ Int. Conf. Intell. Robots Syst.*, 2020, pp. 10952–10959.
- [16] E. Biyik, M. Palan, N. C. Landolfi, D. P. Losey, and D. Sadigh, "Asking easy questions: A user-friendly approach to active reward learning," in *Proc. Conf. Robot Learn.*, 2019, pp. 1177–1190.
- [17] J. Levinson et al., "Towards fully autonomous driving: Systems and algorithms," in *Proc. IEEE Intell. Veh. Symp.*, 2011, pp. 163–168.
- [18] A. Botros, A. Sadeghi, N. Wilde, J. Alonso-Mora, and S. L. Smith, "Error-bounded approximation of pareto fronts in robot planning problems," in *Proc. 15th Workshop Algorithmic Found. Robot.*, Springer, 2023, pp. 506–522.
- [19] I. Y. Kim and O. L. de Weck, "Adaptive weighted sum method for multiobjective optimization: A new method for Pareto front generation," *Struct. Multidisciplinary Optim.*, vol. 31, no. 2, pp. 105–116, 2006.
- [20] P. Karkus, B. Ivanovic, S. Mannor, and M. Pavone, "Diffstack: A differentiable and modular control stack for autonomous vehicles," in *Proc. Conf. Robot Learn.*, PMLR, 2023, pp. 2170–2180.
- [21] Z. Lu, Z. Liu, G. J. Correa, and K. Karydis, "Motion planning for collision-resilient mobile robots in obstacle-cluttered unknown environments with risk reward trade-offs," in *Proc. IEEE/RSJ Int. Conf. Intell. Robots Syst.*, 2020, pp. 7064–7070.
- [22] Y. Che, A. M. Okamura, and D. Sadigh, "Efficient and trustworthy social navigation via explicit and implicit robot-human communication," *IEEE Trans. Robot.*, vol. 36, no. 3, pp. 692–707, Jun. 2020.
- [23] B. Brito, B. Floor, L. Ferranti, and J. Alonso-Mora, "Model predictive control for collision avoidance in unstructured dynamic environments," *IEEE Robot. Automat. Lett.*, vol. 4, no. 4, pp. 4459–4466, Oct. 2019.
- [24] M. Zucker et al., "Chomp: Covariant hamiltonian optimization for motion planning," *Int. J. Robot. Res.*, vol. 32, no. 9–10, pp. 1164–1193, 2013.
- [25] T. Marcucci, M. Petersen, D. von Wrangel, and R. Tedrake, "Motion planning around obstacles with convex optimization," *Sci. Robot.*, vol. 8, no. 84, 2023, Art. no. ead7843.
- [26] C. E. Luis, M. Vukosavljev, and A. P. Schoellig, "Online trajectory generation with distributed model predictive control for multi-robot motion planning," *IEEE Robot. Automat. Lett.*, vol. 5, no. 2, pp. 604–611, Apr. 2020.
- [27] D. Dolgov, S. Thrun, M. Montemerlo, and J. Diebel, "Path planning for autonomous vehicles in unknown semi-structured environments," *Int. J. Robot. Res.*, vol. 29, no. 5, pp. 485–501, 2010.
- [28] Y. Zeng, X. Xu, S. Jin, and R. Zhang, "Simultaneous navigation and radio mapping for cellular-connected UAV with deep reinforcement learning," *IEEE Trans. Wireless Commun.*, vol. 20, no. 7, pp. 4205–4220, Jul. 2021.
- [29] D. Kent and S. Chernova, "Human-centric active perception for autonomous observation," in *Proc. IEEE Int. Conf. Robot. Automat.*, 2020, pp. 1785–1791.
- [30] B. Sakcak and S. M. LaValle, "Complete path planning that simultaneously optimizes length and clearance," in *Proc. IEEE Int. Conf. Robot. Automat.*, 2021, pp. 10100–10106.
- [31] M. Cakmak, S. S. Srinivasa, M. K. Lee, J. Forlizzi, and S. Kiesler, "Human preferences for robot-human hand-over configurations," in *Proc. IEEE/RSJ Int. Conf. Intell. Robots Syst.*, 2011, pp. 1986–1993.
- [32] E. Biyik, D. P. Losey, M. Palan, N. C. Landolfi, G. Shevchuk, and D. Sadigh, "Learning reward functions from diverse sources of human feedback: Optimally integrating demonstrations and preferences," *Int. J. Robot. Res.*, vol. 41, no. 1, pp. 45–67, 2022.
- [33] S. Habibian, A. Jonnavittula, and D. P. Losey, "Here's what i've learned: Asking questions that reveal reward learning," *ACM Trans. Hum.-Robot Interaction*, vol. 11, no. 4, pp. 1–28, 2022.
- [34] C. Basu, M. Singhal, and A. D. Dragan, "Learning from richer human guidance: Augmenting comparison-based learning with feature queries," in *Proc. ACM/IEEE Int. Conf. Hum.-Robot Interaction*, 2018, pp. 132–140.
- [35] N. Wilde, E. Biyik, D. Sadigh, and S. L. Smith, "Learning reward functions from scale feedback," in *Proc. Conf. Robot Learn.*, PMLR, 2022, pp. 353–362.
- [36] N. Wilde and J. Alonso-Mora, "Do we use the right measure? Challenges in evaluating reward learning algorithms," in *Proc. Conf. Robot Learn.*, PMLR, 2023, pp. 1553–1562.
- [37] A. Bobu, D. R. Scobee, J. F. Fisac, S. S. Sastry, and A. D. Dragan, "Less is more: Rethinking probabilistic models of human behavior," in *Proc. ACM/IEEE Int. Conf. Hum.-Robot Interaction*, 2020, pp. 429–437.
- [38] J. Branke, J. Branke, K. Deb, K. Miettinen, and R. Slowinski, *Multiobjective Optimization: Interactive and Evolutionary Approaches*, vol. 5252. Berlin, Germany: Springer Science & Business Media, 2008.
- [39] I. Das and J. E. Dennis Jr, "A closer look at drawbacks of minimizing weighted sums of objectives for pareto set generation in multicriteria optimization problems," *Struct. Optim.*, vol. 14, pp. 63–69, 1997.
- [40] I. Y. Kim and O. L. De Weck, "Adaptive weighted-sum method for Bi-objective optimization: Pareto front generation," *Struct. Multidisciplinary Optim.*, vol. 29, pp. 149–158, 2005.
- [41] J. Lee, D. Yi, and S. S. Srinivasa, "Sampling of pareto-optimal trajectories using progressive objective evaluation in multi-objective motion planning," in *Proc. IEEE/RSJ Int. Conf. Intell. Robots Syst.*, 2018, pp. 1–9.
- [42] V. Pereyra, M. Saunders, and J. Castillo, "Equispaced pareto front construction for constrained bi-objective optimization," *Math. Comput. Modelling*, vol. 57, no. 9–10, pp. 2122–2131, 2013.
- [43] J. Xu, A. Spielberg, A. Zhao, D. Rus, and W. Matusik, "Multi-objective graph heuristic search for terrestrial robot design," in *Proc. IEEE Int. Conf. Robot. Automat.*, 2021, pp. 9863–9869.
- [44] S. Parisi, M. Pirota, and J. Peters, "Manifold-based multi-objective policy search with sample reuse," *Neurocomputing*, vol. 263, pp. 3–14, 2017.

- [45] O. Schütze, O. Cuate, A. Martín, S. Peitz, and M. Dellnitz, "Pareto explorer: A global/local exploration tool for many-objective optimization problems," *Eng. Optim.*, vol. 52, no. 5, pp. 832–855, 2020.
- [46] N. Wilde, A. Sadeghi, and S. L. Smith, "Learning submodular objectives for team environmental monitoring," *IEEE Robot. Automat. Lett.*, vol. 7, no. 2, pp. 960–967, Apr. 2022.
- [47] S. P. Boyd and L. Vandenberghe, *Convex Optimization*. Cambridge, U.K.: Cambridge Univ. Press, 2004.
- [48] H. N. Psaraftis, M. Wen, and C. A. Kontovas, "Dynamic vehicle routing problems: Three decades and counting," *Networks*, vol. 67, no. 1, pp. 3–31, 2016.
- [49] Q. Zhang and H. Li, "MOEA/D: A multiobjective evolutionary algorithm based on decomposition," *IEEE Trans. Evol. Comput.*, vol. 11, no. 6, pp. 712–731, Dec. 2007.
- [50] N. Wilde, D. Kulić, and S. L. Smith, "Bayesian active learning for collaborative task specification using equivalence regions," *IEEE Robot. Automat. Lett.*, vol. 4, no. 2, pp. 1691–1698, Apr. 2019.
- [51] A. Sadeghi and S. L. Smith, "Heterogeneous task allocation and sequencing via decentralized large neighborhood search," *Unmanned Syst.*, vol. 5, no. 02, pp. 79–95, 2017.



Alexander Botros (Member, IEEE) received the undergraduate and M.Sc. degrees in engineering from Concordia University, Montreal, Canada, in 2007 and 2016, respectively, and the Ph.D. degree in electrical & computer engineering from the Autonomous Systems Lab, University of Waterloo, Waterloo, Canada, in 2021.

He is the Director of Machine Learning and Artificial Intelligence with Integrus Solutions, a Hong Kong-based due diligence research company. He completed the Postdoctoral Fellowship with the Au-

tonomous Systems Lab, University of Waterloo. His current research interests include natural language processing, language and logic modelling, and machine learning more broadly.



Nils Wilde (Member, IEEE) received the undergraduate and M.Sc. degrees in engineering from Technical University Berlin, Berlin, Germany, in 2012 and 2016, respectively, and the Ph.D. degree in electrical and computer engineering (ECE) under the co-supervision of Dana Kulić and Stephen L. Smith from the University of Waterloo, Waterloo, Canada, in 2020.

He is currently a Postdoctoral Fellow with the Autonomous Multi-Robots Lab working with Javier Alonso-Mora at TU Delft, Delft, The Netherlands.

Until 2021, he was a Postdoctoral Fellow with the Autonomous Systems Lab, University of Waterloo. His research interests include robot motion planning, multirobot coordination, human–robot interaction (HRI), and multiobjective planning, focusing on the development algorithmic frameworks on the intersection of control, learning, and optimization.



Armin Sadeghi (Member, IEEE) received the undergraduate degree in mechanical and aerospace engineering from the Sharif University of Technology, Tehran, Iran, in 2013, the master's degree in electrical engineering from the University of Waterloo, Waterloo, Canada, in 2016, and the Ph.D. degree in electrical & computer engineering from the Autonomous Systems Lab, University of Waterloo, in 2020.

He is currently a Software Engineer with RideCo, a non-demand transit technology and services company based in Waterloo, Canada. He was a Postdoctoral Fellow with the Autonomous Systems Lab, University of Waterloo. His research interests include the areas of control and optimization with focus on multirobot task allocation in ride-sourcing and coverage control applications.



Javier Alonso-Mora (Senior Member, IEEE) received the M.Sc. and Ph.D. degrees in robotics from ETH Zürich, Zürich, Switzerland, in 2010 and 2014, respectively.

He is currently an Associate Professor with the Delft University of Technology, Delft, The Netherlands. Until 2016, he was a Postdoctoral Associate with the Computer Science and Artificial Intelligence Laboratory, Massachusetts Institute of Technology, Cambridge, MA, USA. Toward the smart cities of the future, he applies these techniques in various fields,

including self-driving cars, automated factories, aerial vehicles, and intelligent transportation systems. His main research interests include autonomous navigation of mobile robots, with a special emphasis in multirobot systems and robots that interact with other robots and humans.

Dr. Alonso-Mora was a recipient of the European Research Council (ERC) Starting Grant in 2021, the IEEE International Conference on Robotics and Automation (ICRA) Best Paper Award on Multi-Robot Systems in 2019, and the Veni Grant from the Netherlands Organization for Scientific Research in 2017. He is a Member of Disney Research, Zürich.



Stephen L. Smith (Senior Member, IEEE) received the B.Sc. degree in engineering physics from Queen's University, Kingston, Canada, in 2003, the M.A.Sc. degree in electrical and computer engineering from the University of Toronto, ON, Canada, in 2005, and the Ph.D. degree in mechanical engineering from the University of California, Santa Barbara, CA, USA, in 2009.

He is currently a Professor with the Department of Electrical and Computer Engineering, University of Waterloo, Waterloo, Canada, where he holds a Canada Research Chair in Autonomous Systems. He is Co-Director of the Waterloo Artificial Intelligence Institute. From 2009 to 2011, he was a Postdoctoral Associate with the Computer Science and Artificial Intelligence Laboratory, Massachusetts Institute of Technology (MIT), Cambridge, MA, USA. His main research interests include motion planning and optimization for autonomous systems, with a focus on learning and optimal decision-making.

Dr. Smith is the recipient of the Early Researcher Award from the Province of Ontario, the NSERC Discovery Accelerator Supplement Award, and two Outstanding Performance Awards from the University of Waterloo. He is a licensed Professional Engineer (PEng), an Associate Editor for *IEEE TRANSACTIONS ON ROBOTICS* and *IEEE Open Journal of Control Systems*. He was previously an Associate Editor for *IEEE TRANSACTIONS ON CONTROL OF NETWORK SYSTEMS* (2017–2022), and was a General Chair of the 2021 30th IEEE International Conference on Robot and Human Interactive Communication (RO-MAN).