

---

# Interactive Multi-Objective Probabilistic Preference Learning with Soft and Hard Bounds

---

Edward Chen<sup>1</sup>

Sang T. Truong<sup>1</sup>

Natalie Dullerud<sup>1</sup>

Sanmi Koyejo<sup>1</sup>

Carlos Guestrin<sup>1</sup>

<sup>1</sup>Department of Computer Science, Stanford University, Stanford, California, USA

## Abstract

High-stakes decision-making involves navigating multiple competing objectives with expensive evaluations. For instance, in brachytherapy, clinicians must balance maximizing tumor coverage (e.g., an aspirational target or *soft bound* of  $>95\%$  coverage) against strict organ dose limits (e.g., a non-negotiable *hard bound* of  $<601$  cGy to the bladder). Selecting Pareto-optimal solutions that match implicit preferences is challenging, as exhaustive Pareto frontier exploration is computationally and cognitively prohibitive, necessitating interactive frameworks to guide users. While decision-makers (DMs) often possess domain knowledge to narrow the search via such soft-hard bounds, current methods often lack systematic approaches to iteratively refine these multi-faceted preference structures. Furthermore, DMs often require confidence that they have not overlooked superior alternatives, a paramount necessity in high-stakes scenarios. We present Active-MoSH, an interactive local-global framework designed for this process. Its *local component* integrates probabilistic preference learning with an active sampling strategy to adaptively refine Pareto subsets while minimizing cognitive burden. To bolster decision confidence, Active-MoSH’s *global component*, C-MoSH, leverages multi-objective sensitivity analysis to identify potentially overlooked, high-value points beyond immediate feedback. We demonstrate Active-MoSH’s performance benefits through diverse synthetic and real-world applications. A high-stakes case study with real cervical cancer brachytherapy treatment plans and an image selection user study further validate our hypotheses regarding the framework’s ability to improve convergence, enhance DM confidence, and provide expressive preference articulation.

## 1 INTRODUCTION

Countless critical decision-making tasks within domains like healthcare and engineering design involve finding optimal tradeoffs among multiple competing, often expensive-to-evaluate, objectives  $f_1(x), \dots, f_L(x)$  Fromer and Coley [2023], Luukkonen et al. [2023], Xie et al. [2021], Yu et al. [2000], Papadimitriou and Yannakakis [2001]. Decision-makers (DMs) typically seek a Pareto-optimal (PO) point that aligns with their hidden preferences. However, exhaustively searching the Pareto frontier is often infeasible due to computational costs, lengthy validation, and limited human attention. Consequently, DMs usually navigate the frontier interactively, validating subsets of points — often defined by initial, possibly flexible, constraints or target regions — to find their ideal solution Paria et al. [2019], Ozaki et al. [2023], Wilde et al. [2022].

In healthcare, for instance, brachytherapy planning for cervical cancer treatment requires clinicians to balance maximizing radiation to tumors against minimizing exposure to healthy organs Deufel et al. [2020]. The vast tradeoff space and time-consuming nature of plan evaluation (potentially hours to densely populate the tradeoff space) make direct exploration impractical, especially while the patient is under anesthesia Cui et al. [2018a,b]. Crucially, clinicians often possess significant prior knowledge to define initial regions of interest. Rather than just as values to be maximized or minimized, this prior knowledge often comes in the form of both desired targets (or soft bounds) and strict, non-negotiable boundaries (or hard bounds). For example, a clinician may aim for a high level of tumor coverage (soft bound:  $> 95\%$ ) while also needing to ensure that coverage does not fall below a critical minimum threshold (hard bound:  $> 90\%$ ), and that radiation to a healthy organ like the bladder remains strictly below a maximum (hard bound:  $< 601$  centigrays (cGY)) and preferred dose (soft bound:  $< 513$  cGY) Chen et al. [2026b]. As clinicians observe presented treatment plans and gain new insights into achievable tradeoffs, they iteratively refine these initial preferences. For

instance, if achieving the  $> 95\%$  tumor coverage consistently results in plans where the bladder dose approaches or exceeds 601 cGy, the clinician may directly adjust their minimum or desired tumor coverage downwards to explore solutions that offer a more viable and safe balance. This iterative adjustment of their articulated goals is central to navigating complex decision spaces. A critical challenge, however, is building DM confidence: DMs must be confident that they have not overlooked potentially superior PO points outside of their current tradeoff space. While multi-objective preference learning has advanced, methods explicitly supporting such iterative refinement using expressive, multi-faceted preference structures and robust confidence-building within Pareto frontier subsets remain underexplored.

This paper formalizes this interactive process through Active-MoSH, a novel framework with distinct *local* and *global* components to efficiently guide DMs to optimal solutions while fostering confidence in their final choice. The *local component* of Active-MoSH iteratively refines the DM’s search within the tradeoff space. To effectively capture and adapt to nuanced DM preferences (often comprising aspirational targets and strict limits), it extends the soft-hard utility function (SHF) concept Chen et al. [2026b] to a dynamic, interactive setting. This SHF-based approach enables probabilistic modeling of the DM’s latent preference vector and their evolving soft-hard bounds, thereby adaptively refining the explored Pareto subset and focusing computational resources on preference-aligned regions. An integrated active sampling strategy further optimizes this local search by balancing exploration-exploitation and minimizing computational/cognitive load. Addressing potential DM doubts about unobserved superior alternatives, the *global component*, C-MoSH, enhances decision confidence. By leveraging multi-objective sensitivity analysis, C-MoSH systematically quantifies solution robustness to stated bounds and proactively identifies potentially overlooked, high-value Pareto regions *beyond immediate DM feedback*, thereby building strong DM confidence. We rigorously validate Active-MoSH via simulations on synthetic benchmarks and real-world applications, including a brachytherapy planning case study utilizing three real patient cases. Furthermore, a user study involving AI-generated image selection for educational magazine covers statistically confirms our hypotheses regarding the framework’s improved efficiency, decision-making effectiveness, and ability to bolster decision confidence.

Although there exists extensive work on interactive multi-objective decision-making (MODM), many methods focus on feedback mechanisms that may not fully support intuitive exploration of specific tradeoff sub-regions Wilde et al. [2022], Bviyik et al. [2020], Astudillo and Frazier [2020], Ozaki et al. [2023]. While some recent approaches allow DMs to impose priors on Pareto subsets, they lack emphasis on the interactive feedback loop with expressive bound structures like soft-hard bounds, and the crucial aspect of

building DM confidence within this dynamic context Paria et al. [2019], Malkomes et al. [2021]. In summary, our contributions comprise the following:

1. **Active-MoSH:** An interactive framework with a *local component* that helps DMs iteratively refine preferred Pareto regions using probabilistic preference learning over soft-hard bounds and active sampling.
2. **C-MoSH:** The *global component* of Active-MoSH, which builds DM confidence via multi-objective sensitivity analysis to quantify robustness and proactively uncover overlooked optimal regions.
3. **Empirical Validation:** Rigorous evaluation across synthetic problems, a high-stakes brachytherapy case study, and a user study on AI-generated images, demonstrating statistically significant improvements in decision efficiency and confidence.

## 2 MULTI-OBJECTIVE PREFERENCE LEARNING WITH SOFT-HARD BOUNDS

### 2.1 BACKGROUND

Multi-objective optimization (MOO) seeks to jointly maximize  $L$  expensive-to-query black-box objective functions  $f_1(x), \dots, f_L(x)$  over an input space  $X \subset \mathbb{R}^d$ . Since solutions rarely optimize all objectives simultaneously, MOO focuses on identifying points on the Pareto frontier. Scalarization techniques,  $s_\lambda(y) : \mathbb{R}^L \rightarrow \mathbb{R}$ , parameterized by  $\lambda \in \Lambda$  representing relative preferences, combine objectives into a single value to find PO solutions [Rojers et al., 2013].

Decision-makers (DMs) often have priors on desirable objective space regions, expressible as soft and hard bounds [Chen et al., 2026b]. A soft bound,  $\alpha_{\ell,S}$ , describes an ideal target for objective  $f_\ell(x)$ , while a hard bound,  $\alpha_{\ell,H}$ , defines a limit  $f_\ell(x)$  must not violate. Since ideal targets  $\alpha_{\ell,S}$  might be unreachable due to competing objectives, hard bounds are crucial. To operationalize these dual-level preferences, Chen et al. [2026b] proposed *soft-hard* utility functions (SHFs),  $u_\alpha(x)$ , which map each  $f_\ell$  to a bounded utility space based on bounds  $\{\alpha_{1,S}, \alpha_{1,H}, \dots, \alpha_{L,S}, \alpha_{L,H}\}$ . *Additional details may be found in Appendix A.1.*

Points  $y \in Y$  on the SHF-defined Pareto frontier are found by solving

$$\max_{x \in X} s_\lambda(u_\alpha(x)), \quad (1)$$

where  $u_\alpha := [u_{\alpha_1}, \dots, u_{\alpha_L}]$ ,  $u_{\alpha_\ell}$  is the SHF for  $f_\ell$  (parameterized by  $\alpha_{\ell,S}, \alpha_{\ell,H}$ ), and  $s_\lambda(u_\alpha(x))$  is assumed monotonically increasing in  $u_\alpha(x)$ .

**Novel Contributions.** While Chen et al. [2026b] introduced the SHF framework for obtaining a single set of easily navigable PO points, extending the static SHF formulation to a

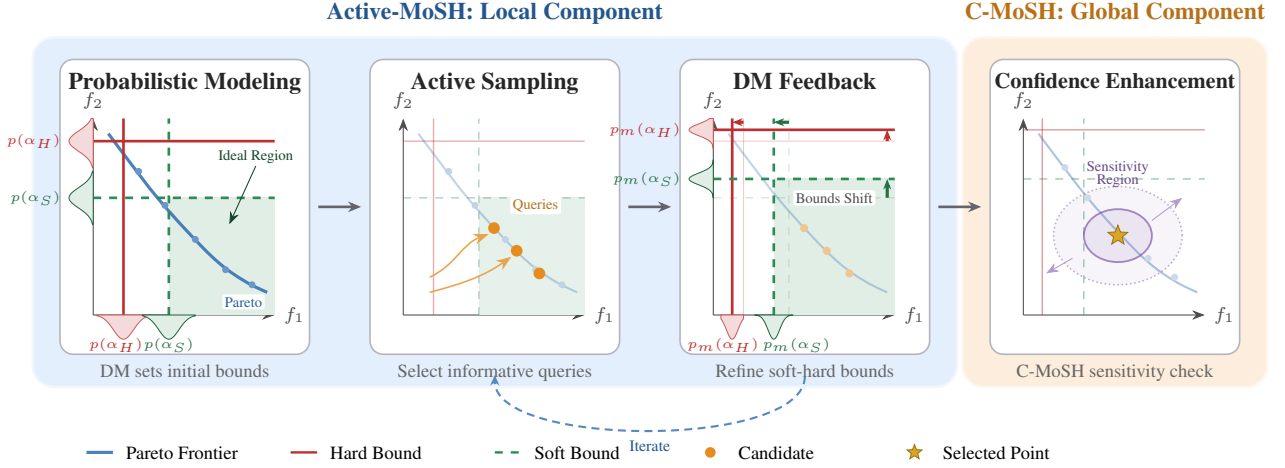


Figure 1: Overview of Active-MoSH. **(Left three panels - Local Component):** The DM defines initial soft bounds (desired targets, dashed green) and hard bounds (strict limits, solid red) on objectives, narrowing the Pareto frontier to a focus region. Active-MoSH actively selects informative queries (orange) within the soft-hard bound defined region. The DM observes these queries and adjusts bounds based on observed tradeoffs (arrows show bound shifts), iterating until convergence. **(Right panel - Global Component):** C-MoSH performs MOO sensitivity analysis around the current best solution (gold star), probing regions beyond stated bounds to confirm no superior alternatives were overlooked, thereby building DM confidence.

dynamic, interactive setting is non-trivial. It necessitates a principled approach for modeling evolving decision-maker preferences under uncertainty, a mechanism to actively build confidence that superior solutions are not being overlooked, and a rigorous framework for evaluating the interactive system’s effectiveness. This work introduces novel local-global components to systematically address these challenges.

## 2.2 PROBLEM DEFINITION

**Problem Setting.** As DM evaluation of each point  $\mathbf{y} \in Y$  is expensive, it is infeasible to observe the full  $Y$  on the entire Pareto frontier. We assume the DM has a hidden ideal set of soft and hard bounds,  $\{\alpha_S^*, \alpha_H^*\}$  (abbreviating  $\{\alpha_{1,S}^*, \dots, \alpha_{L,S}^*\}$  and  $\{\alpha_{1,H}^*, \dots, \alpha_{L,H}^*\}$ ), which best characterizes the set of tradeoffs most appropriate for their decision-making task. This implies a hidden ideal point  $\mathbf{y}^* = f(x^*)$  on the SHF-defined Pareto frontier, corresponding to the DM’s hidden preferences  $\lambda^*$ . Formally,  $x^* = \arg \max_{x \in X} s_{\lambda^*}(u_{\alpha^*}(x))$ , where  $\mathbf{y} = f(x)$ .

**Feedback Mechanism.** We assume that  $\lambda$  and the set of bounds  $\{\alpha_{1,S}, \dots, \alpha_{L,H}\}$ , abbreviated as  $\{\alpha_S, \alpha_H\}$ , are random variables. We aim to learn  $\lambda^*$  and  $\{\alpha_S^*, \alpha_H^*\}$  through  $M$  iterations of interaction. At each iteration  $m$ , the DM provides feedback by specifying point values for soft and hard bounds, which we denote as observations  $\alpha_{S,m}$  and  $\alpha_{H,m}$ . As exact numerical feedback is often imprecise, we express uncertainty by modeling these specified values as noisy observations of the true latent bounds rather than fixed parameters. Specifically, we treat  $\alpha_{S,m}$  and  $\alpha_{H,m}$  as samples drawn from the true bound distributions with fixed

observation noise  $\Sigma_{\text{noise}}$ . We maintain posterior beliefs, e.g.,  $p_m(\alpha_H) \sim \mathcal{N}(\mu_{H,m}, \Sigma_{H,m})$ , and update the parameters  $\{\mu_{H,m}, \Sigma_{H,m}\}$  using these observations.

The interactive process starts at iteration  $m = 0$ : the DM inputs their prior distributions for the bounds,  $p(\alpha_S)$  and  $p(\alpha_H)$ , by specifying an initial set of soft and hard bounds,  $\{\alpha_{S,0}, \alpha_{H,0}\}$ . For each subsequent iteration  $m = 1, \dots, M$ : (1) based on the current beliefs of  $\lambda$  and  $\alpha$ , a set of  $k$  PO points,  $Y_m = \{\mathbf{y}_1, \dots, \mathbf{y}_k\}$ , is presented to the DM as a query, (2) the DM returns feedback  $\alpha_m$  by adjusting one soft or hard bound for any of the objectives, and (3) the posteriors of  $\lambda$  and  $\alpha$  are then updated using Bayes’ rule, incorporating the DM feedback.

Due to potentially complex relationships across the multiple objectives overwhelming the DM, we assume a single modification of any of the soft and hard bounds to be the full iteration. At the end of  $M$  feedback iterations, we assume the DM selects any previously seen PO points from the  $M$  preference queries. As a result, our overall objective is to optimize the following SHF Utility Ratio:

$$\mathbb{E} \left[ \frac{\max_{x \in \{X_1 \cup \dots \cup X_M\}} s_{\lambda}(u_{\alpha}(x))}{\max_{x \in X} s_{\lambda}(u_{\alpha}(x))} \right] \quad (2)$$

Intuitively, the SHF Utility Ratio is maximized when the points in  $f(X_m)$  are PO and span the high utility regions of the PF, as defined by the SHFs  $u_{\alpha}$ . We use Eq. (2) to evaluate our results in Section 6. Figure 1 illustrates this interactive process.

### 3 ACTIVE-MOSH

To learn the DM’s ideal preferences  $\lambda^*$  and bounds  $\{\alpha_S^*, \alpha_H^*\}$ , we maintain and iteratively update posterior probability distributions over these quantities. This probabilistic approach captures our uncertainty and informs the active sampling of preference queries  $Y$ , making up the local component of Active-MoSH. *More details App. B.*

#### 3.1 PROBABILISTIC MODELING OF PREFERENCES AND BOUNDS

**Preferences Modeling.** Given the DM’s history of feedbacks, we update our belief over  $\lambda$  using Bayes’ rule:

$$p_m(\lambda) \propto p(\lambda) \prod_m L(\pi_m; X_m, Y_m, \alpha_m, \lambda) \quad (3)$$

$p_m(\lambda)$  is the posterior distribution over  $\lambda$  after  $m$  feedback iterations and  $p(\lambda)$  is the uniform prior probability distribution over  $\lambda$ . The likelihood term  $L(\pi_m; X_m, Y_m, \alpha_m, \lambda)$  represents the probability of the DM’s feedback implying the ranking  $\pi_m$  in response to the preference query  $Y_m$  (Section 3.2) and current bounds and parameters.

**Bounds Modeling.** As DMs provide feedback by imprecisely modifying soft-hard bounds, we update our posterior beliefs over the true ideal bounds  $\alpha_S, \alpha_H$ . We model the DM’s specified bound at iteration  $m$ , denoted  $\alpha_{H,m}$ , as a noisy observation of the true bound  $\alpha_H$ , governed by the likelihood function  $L(\alpha_{H,m} | \alpha_H) = \mathcal{N}(\alpha_{H,m} | \alpha_H, \Sigma_{\text{noise}})$ . Assuming a conjugate Gaussian prior  $p(\alpha_H) \sim \mathcal{N}(\mu_H, \Sigma_H)$ , the posterior is updated analytically via Bayes’ rule (similar update applies to  $p_m(\alpha_S)$ ):

$$p_m(\alpha_H) \propto p(\alpha_H) \prod_{i=1}^m \mathcal{N}(\alpha_{H,i} | \alpha_H, \Sigma_{\text{noise}}) \quad (4)$$

$\Sigma_{\text{noise}}$  is the variance of the DM’s feedback consistency.

#### 3.2 MODELING SOFT-HARD FEEDBACK

**Feedback Interpretation.** Our framework interprets DM soft-hard bounds feedback as implicit signal about their current region of interest on the Pareto frontier. This interpretation imposes a ranking over the set of query points and directly influences the likelihood model described in Section 3.1; it is consistent with reference-dependent decision making: changing a constraint/target shifts the DM’s reference point or aspiration/reservation levels, which is known to change relative evaluations and the induced preference ordering [Deb and Sundar, 2006, Kahneman and Tversky, 2012]. As with reference points [Vesikar et al., 2018], the points in  $Y_m$  closest to this newly defined “edge” or “target” of the feasible/desirable region are considered most actionable or informative, hence ranked highest. *We provide*

*detailed motivation for the specific ranking interpretations of all feedback scenarios in Appendix B.1.*

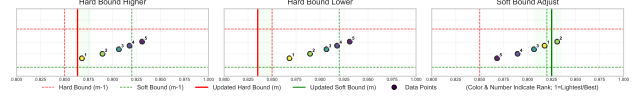


Figure 2: Bounds Feedback Interpretation. Adjustment of bounds imposes a ranking over the query points, determined by Euclidean distance from the shifted bound. The green shaded region represents actionable immediate tradeoff areas of interest implied by the feedback, which subsequently is where the distribution of preference vectors shifts towards.

**Likelihood Model.** As mentioned, we assume the DM’s feedback (bound modification) implicitly ranks points in the query  $Y_m$  based on their proximity to the newly adjusted bound, making those points more relevant. To operationalize the interpreted ranking, we use Euclidean distance from the shifted bound and leverage the Plackett-Luce model for the likelihood function in Eq. (3) [Luce, 1959]:

$$L(\pi; X, Y, \alpha, \lambda) = \prod_{j=1}^k \frac{\exp(s_\lambda(u_\alpha(X_{\pi(j)})))}{\sum_{h=j}^k \exp(s_\lambda(u_\alpha(X_{\pi(h)})))} \quad (5)$$

where  $\pi = [\pi(1), \dots, \pi(k)]$  is the DM-induced ranking of the  $k$  points in  $Y$ ,  $X_{\pi(j)}$  is the input value for the point in  $Y$  with rank  $j$ , and  $u_\alpha$  is the SHF utility using the bounds.

#### 3.3 ACTIVE QUERY SAMPLING

To query the DM effectively, we aim to select a navigable set of PO points  $Y_m$  that reflects their current preferences and reduces cognitive load. The central challenge is our uncertainty over both the DM’s latent preference vector ( $\lambda^*$ ) and their ideal soft-hard bounds ( $\alpha^*$ ). Distinct from static methods [Chen et al., 2026b], our active sampling must be robust to this dual uncertainty. We therefore adapt a two-step process: (1) obtain a dense set of points by maximizing the expected SHF Utility Ratio (Eq. 2), where the expectation is over our posteriors of both  $\lambda$  and  $\alpha$ , then (2) sparsify this set to present a small, diverse subset to the DM.

**Step 1: Dense Set of Points.** As the first step of the two-step process, we aim to obtain a dense and diverse set of candidate points  $Y_{m_D}$  which maximizes the expected SHF Utility Ratio, while being robust to current uncertainty in the DM’s preferences and soft-hard bounds. We model the expensive black-box objectives  $f_\ell(x)$  with Gaussian Processes (GPs) [Williams and Rasmussen, 1995]. Since  $\lambda^*$  and  $\alpha^*$  are both unknown to us, we want  $Y_{m_D}$  to be robust to any potential  $\lambda^*$  and  $\alpha^*$ , weighted according to their posterior distributions. As a result, we first aim to obtain a dense and diverse set of candidate points  $Y_{m_D}$  by maximizing the expected

SHF Utility Ratio (Eq. (2) from Section 2.2) over choices of  $X_{m_D}$ :

$$\max_{\substack{X_{m_D} \subseteq X, \\ |X_{m_D}| \leq k_D}} \mathbb{E}_{\substack{\lambda \sim p_m(\lambda), \\ \alpha \sim p_m(\alpha)}} \left[ \frac{\max_{x \in X_{m_D}} s_\lambda(u_\alpha(x))}{\max_{x' \in X} s_\lambda(u_\alpha(x'))} \right] \quad (6)$$

where the expectation is over the current posterior distributions of  $\lambda$  and  $\alpha$ .

To solve Equation (6), we adapt the MoSH-Dense algorithm from Chen et al. [2026b]. In short, our algorithm uses the notion of random scalarizations to sample a  $\lambda$  and  $\alpha$  from their posterior distributions at each iteration, which is then used to compute a multi-objective acquisition function [Paria et al., 2019]. The maximizer of the acquisition function (Appendix B.2) is then chosen as the next sample input to be evaluated with the expensive black-box function, resulting in a single PO point. This is continued for a total of  $T$  iterations, resulting in a dense set of  $T$  PO points. For simplicity, we denote  $X_{m_D}$  with  $D$ . The full algorithm is shown in Algorithm 1.

---

**Algorithm 1** Active-MoSH-Dense: Dense Sampling

---

```

1: procedure ACTIVE-MOSH-DENSE
2:   Initialize  $D^{(0)} = \emptyset$ 
3:   Initialize  $GP_\ell^{(0)} = GP(0, \kappa) \forall \ell \in [L]$ 
4:   for  $t = 1 \rightarrow T$  do
5:     Obtain  $\lambda_t \sim p_t(\lambda)$  and  $\alpha_t \sim p_t(\alpha)$ 
6:      $x_t = \arg \max_{x \in X} \text{acq}(u, \lambda_t, \alpha_t, x) \triangleright \text{App. B.2}$ 
7:     Obtain  $y = f(x_t)$ 
8:      $D^{(t)} = D^{(t-1)} \cup \{(x_t, y)\}$ 
9:      $GP_\ell^{(t)} = \text{post}(GP_\ell^{t-1} | (x_t, y)) \forall \ell \in [L]$ 
10:   end for
11:   Return  $D^{(T)}$ 
12: end procedure
```

---

**Step 2: Sparse Set of Points.** To reduce the cognitive load on the DM, we then wish to sparsify the set of points  $X_{m_D}$ . We adapt Equation 6 and formulate the sparsification as a robust submodular observation selection (RSOS) problem [Krause et al., 2008]:

$$\max_{X_m \subseteq X_{m_D}, |X_m| \leq k} \min_{\lambda \in \Lambda} \underbrace{\left[ \frac{\max_{x \in X_m} s_\lambda(u_\alpha(x))}{\max_{x' \in X_{m_D}} s_\lambda(u_\alpha(x'))} \right]}_{F_\lambda} \quad (7)$$

where  $X_{m_D}$  represents the set obtained from Active-MoSH-Dense, Algorithm 1, and  $f(X_m)$  is the sparse set of SHF-defined PO points returned to the DM. To solve Equation 7, we apply MoSH-Sparse, from Chen et al. [2026b], on  $X_{m_D}$ . MoSH-Sparse leverages the notion of submodularity, which encapsulates the concept of diminishing returns in utility for each additional point the DM validates [Chen et al., 2026b, Nemhauser et al., 1978, Krause et al., 2008]. We leverage the proof from Chen et al. [2026b] to show

that the SHF Utility Ratio,  $F_\lambda$ , is a submodular function. As a result, MoSH-Sparse theoretically guarantees for our sampled preference queries to obtain a set of points  $Y_m$  from  $Y_{m_D}$  which achieves the optimal coverage of the unknown preferences, albeit with a slightly greater number of points than  $k$ . We provide the full details in Algorithm 2. For simplicity of notation, we denote  $X_m$  with  $C$  and  $X_{m_D}$  with  $D$ .  $\psi$  is defined in Theorem 1. Additional details of MoSH-Sparse may be found in Chen et al. [2026b].

---

**Algorithm 2** MoSH-Sparse: Pareto Sparsification

---

```

1: procedure MOSH-SPARSE( $F_1, \dots, F_{|\Lambda|}, k, \psi$ )
2:    $q_{min} = 0; q_{max} = \min_i F_i(D); C_{best} = \emptyset$ 
3:   while  $(q_{max} - q_{min}) \geq 1/|\Lambda|$  do
4:      $q = (q_{min} + q_{max})/2$ 
5:     Define  $\bar{F}_q(C) = 1/|\Lambda| \sum_i \min\{F_i(C), q\}$ 
6:      $\hat{C} = GPC(\bar{F}_q, q) \triangleright \text{App. B.2: Algorithm 3}$ 
7:     if  $|\hat{C}| > \psi k$  then
8:        $q_{max} = q$ 
9:     else
10:       $q_{min} = q; C_{best} = \hat{C}$ 
11:    end if
12:  end while
13: end procedure
```

---

**Theorem 1.** [Krause et al., 2008] For any integer  $k$ , MoSH-Sparse finds a solution  $C_S$  such that  $\min_i F_i(C_S) \geq \max_{|C| \leq k} \min_i F_i(C)$  and  $|C_S| \leq \psi k$ , for  $\psi = 1 + \log(\max_{x \in D} \sum_i F_i(\{x\}))$ . The total number of submodular function evaluations is  $\mathcal{O}(|D|^2 m \log(m \min_i F_i(D)))$ , where  $m = |\Lambda|$ .

## 4 C-MOSH: ENHANCING CONFIDENCE

While the local component of Active-MoSH focuses on refining the DM's immediate region of interest, DMs may still question if their current solution is truly optimal or if slight bound adjustments might reveal superior alternatives. To address this, we introduce C-MoSH, the *global component* of our framework, designed to build DM confidence through multi-objective sensitivity analysis [Rappaport, 1967]. This component expands the search beyond the DM's direct feedback to uncover potentially overlooked regions. Grounded in decision sciences [Desender et al., 2019, 2018, Wu and Kelly, 2014], providing this broader validation of solution quality helps to foster DM confidence.

C-MoSH informs the DM about the sensitivity of each objective  $f_l$  to perturbations in another objective  $f_{l'}$ . This allows the DM to understand how "active" an objective  $l$  is with respect to  $l'$  (i.e., how much  $f_l$  changes if  $f_{l'}$  is varied) and whether such modifications could lead to significantly better outcomes in  $f_l$ . An objective is *active* if perturbing another results in substantial changes to it. Given that function evaluations  $f(x)$  are expensive, C-MoSH conservatively

focuses on identifying samples that promise improvement without requiring a full re-run of Active-MoSH’s local active sampling (Section 3.3). Specifically, we are interested in samples  $x$  that improve upon the current best-known value for  $f_l$ , denoted  $f_l(x^*)$ . The improvement in dimension  $l$  is:

$$I_l(x) = \max\{f_l(x) - f_l(x^*), 0\}. \quad (8)$$

This measures how much  $f_l(x)$  exceeds  $f_l(x^*)$ . We then consider the *expected improvement*  $\text{EI}_l(x) = \mathbb{E}[I_l(x)]$ , calculated over a posterior model for  $f_l(x)$  (e.g., a Gaussian Process) [Jones et al., 1998].  $\text{EI}_l(x)$  can often be computed in closed form. If  $\text{EI}_l(x) > 0$ ,  $x$  is a worthwhile candidate to evaluate and potentially highlight to the DM. Otherwise, if no improvement in  $f_l$  is likely, objective  $l$  is considered not active with respect to the perturbed dimension  $l'$ . To identify such promising candidates for building confidence, C-MoSH solves:

$$\begin{aligned} \max_{x \in X} \text{EI}(x) \quad \text{s.t.} \quad & u_{\alpha_\ell}(x) \geq 0 \quad \forall \ell \in [L] \setminus \{\ell'\}, \\ & \text{and } u_{\alpha_{\ell',H}-\epsilon}(x) \geq 0 \end{aligned} \quad (9)$$

This formulation targets samples with positive expected improvement that also satisfy the DM’s current soft and hard bounds (slightly relaxed for the perturbed hard bound  $\alpha_{\ell',H}$ , via  $\epsilon$ ) sampled from their posterior distributions. By focusing computational resources on such promising candidates, C-MoSH complements the local search by effectively guiding exploration towards regions that confirm solution robustness or reveal valuable alternatives, aligning with the DM’s broader priorities. Further details are in Appendix C.

## 5 OVERVIEW OF EXPERIMENTS

We evaluate Active-MoSH with three complementary settings: simulation experiments on synthetic benchmarks and real-world applications (Section 6), a human-subject user study (Section 7), and a high-stakes cervical cancer brachytherapy case study (Section 8). The simulations enable controlled comparison against baseline feedback mechanisms at scale, the user study captures effects that rule-based simulation cannot, and the case study validates the framework in a realistic safety-critical setting. Before presenting results, we briefly clarify what each of our experiments isolates.

Our main simulation results (Figures 3 and 6) evaluate complete deployed systems, pairing each baseline with information gain (IG) for active query selection, as IG is standard for preference learning and has been shown to outperform alternatives such as volume removal in both convergence and query ease [Biyik et al., 2020]. Pairing each baseline with its standard acquisition function reflects practical deployment and avoids hand-tying baselines to a criterion not designed for them.

To isolate the contribution of our active sampling, Figure 7 ablates it against random sampling with the rest of the framework held fixed. To further isolate the effect of the feedback mechanism itself, Figure 10 equips all baselines with our active query sampling method. As described in Appendix D.1, our simulator follows fixed rules for bound adjustment, whereas real decision-makers exhibit flexibility, priors, and confidence-driven exploration. Our user study (Section 7) addresses these simulation limitations directly.

## 6 SIMULATION EXPERIMENTS

This section presents a comprehensive empirical evaluation of Active-MoSH against common feedback mechanisms using simulated DM interactions across synthetic benchmarks and real-world applications. Our comparisons include pairwise feedback, complete ranking [Plackett, 1975, Luce, 1959], partial k-ranking ( $k=3$ ) [Guiver and Snelson, 2009], and random selection [Biyik et al., 2020]. For these baselines (excluding random), preference queries were selected via an information gain-based objective [Biyik et al., 2020]. Furthermore, we conduct extensive ablation studies to assess the impact of each component of Active-C-MoSH. For clarity, Active-MoSH integrated with C-MoSH is referred to as Active-C-MoSH. *Additional experiments and setup details are available in Appendix D.*

### 6.1 SIMULATION SETUP

To evaluate our framework, we simulated DM interactions using a ground-truth ideal point  $\mathbf{y}^*$ , derived from a randomly sampled preference vector  $\lambda^*$  and soft-hard bounds  $\alpha^*$  to ensure  $\mathbf{y}^*$  resides within a high-utility region. The goal is to identify  $\mathbf{y}^*$  over 10 feedback iterations, and we measure interaction efficiency using the SHF Utility Ratio (Eq. 2), calculated against these ground-truth values. *Further details on the simulation setup and additional experiments accounting for cognitive complexity of the feedback mechanisms are in Appendix D.1.*

### 6.2 SIMULATION RESULTS

**Branin-Currin, Four Bar Truss, DTLZ2.** We leverage both the Branin-Currin (two objectives) and DTLZ2 (three objectives) problems provided in the BoTorch framework [Balandat et al., 2020]. We further evaluate our approach on a MOO structural engineering problem: the Four Bar Truss design problem from REPROBLEM [Tanabe and Ishibuchi, 2020]. This system presents a bi-objective optimization task with four continuous design variables and exhibits a convex Pareto frontier [CHENG and LI, 1999].

**Brachytherapy.** We validate our framework using a real clinical case for brachytherapy treatment planning. In this

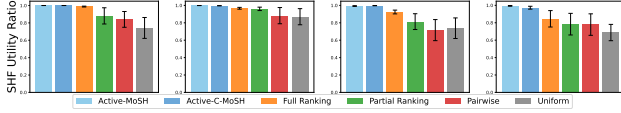


Figure 3: Evaluation results with 10 units for Branin-Currin, Four Bar Truss, Brachytherapy, and DTLZ2 applications (from left to right). 10 independent trials.

setting, we present it as a two-objective optimization problem: maximizing dose to the tumor while minimizing radiation exposure to the bladder. All results are shown in Figure 3, illustrating that Active-MoSH and Active-C-MoSH both often result in higher SHF Utility Ratio values compared to all baselines. More details in Appendix D.4.

## 7 USER STUDY EVALUATION

**Setup.** To evaluate the practical utility of Active-MoSH, we conducted an IRB-approved human-subject user study. We designed a task involving the selection of AI-generated magazine cover photos, which serves as a controlled proxy for complex multi-objective decision-making (see Figure 11). Participants were tasked with using different feedback types to balance two competing objectives, realism and color vividness, to match a specified editorial goal (e.g., high realism, balanced, or high vividness). Each participant interacted with multiple task instances, each targeting a different region of the tradeoff space. *Further details in Appendix E.*

We compared Active-C-MoSH against pairwise, full ranking, and partial ranking, all using information gain. We emphasize that our evaluation focuses on the general applicability of the interactive preference learning framework; thus, we do not compare against methods tailored specifically to generative AI, such as direct prompting.

We recruited 21 Amazon Mechanical Turk participants. In the primary study, each worker completed one task consisting of four instances, each using a distinct feedback type. For each instance, workers iteratively used the mechanism until they identified a tradeoff point they felt matched the task description, then selected a final image from any prior iteration. To further investigate decision confidence and learning effects, we conducted two supplementary experiments: an ablation study isolating the global decision confidence component and a longitudinal study where participants performed four consecutive trials using only Active-C-MoSH. Post-study, a 5-point Likert survey assessed each mechanism’s perceived (1) mental effort, (2) expressiveness, and (3) trustworthiness (confidence).

### 7.1 USER STUDY HYPOTHESES AND METRICS

**Hypotheses.** **H1.** Active-C-MoSH leads to faster convergence to the optimal tradeoff point. **H2.** Active-C-MoSH is helpful for building confidence in the DM process. **H3.** Active-C-MoSH is more cognitively demanding. **H4.** Active-C-MoSH is more expressive of user’s preferences.

To evaluate **H1**, we used the SHF Utility Ratio at the second feedback iteration. Due to the subjective nature of **H2**, **H3** and **H4**, we evaluated them with MTurk worker Likert scale responses. To further quantitatively assess **H2** (confidence), we propose the Iteration Stop Efficiency (ISE) metric, a behavioral proxy for decision-maker confidence. The underlying assumption is that a DM’s confidence in the system directly influences their interaction behavior. Backed by literature in decision theory, a DM who is confident that the system is effectively exploring the most relevant tradeoffs or presenting optimal solutions is more likely to halt the process once a satisfactory (“good enough”) point is found [Clausen et al., 2007, Nandy et al., 2025]. Conversely, a DM lacking confidence may continue iterating excessively, fearing that superior alternatives exist just beyond the presented options. ISE operationalizes this concept by measuring how quickly a user finds a satisfactory point relative to the total number of interactions. A point  $y = f(x)$  is deemed “good” if its SHF Utility Ratio  $U(\{x\})$  satisfies  $U(\{x\}) \geq (1 - \delta)U(\{x^*\})$  for a threshold  $\delta \in [0, 1]$ , where  $U(\{x^*\})$  is the utility of the ideal optimal point. ISE is then  $1 - (M - M_G)/M$  where  $M$  is the total number of feedback iterations, and  $M_G$  is the first iteration index ( $1 \leq M_G \leq M$ ) where the set of evaluated points  $X_{M_G}$  contains at least one “good point.” Formally,  $M_G = \min\{M' \in \{1, \dots, M\} \mid \exists x \in X_{M'} \text{ s.t. } U(\{x\}) \geq (1 - \delta)U(\{x^*\})\}$ .

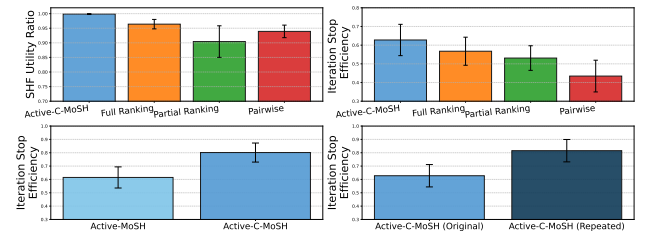


Figure 4: User Study Results. Top (left to right): SHF Utility Ratio at second iteration and Iteration Stop Efficiency Metric (ISE) for all feedback types. Bottom: ISE results for the C-MoSH ablation study and the longitudinal familiarity study.

### 7.2 USER STUDY RESULTS

We present the results of our user study in Figure 4. Regarding **H1** (faster convergence), an analysis of the SHF Utility Ratio at the second feedback iteration *supported that Active-C-MoSH enabled participants to reach the optimal tradeoff point more effectively*. Pairwise comparisons revealed that

Active-C-MoSH achieved significantly higher utility ratios than pairwise feedback (adj.  $p = 0.05$ ) and full ranking (adj.  $p = 0.05$ ). The difference compared to partial ranking was not statistically significant (adj.  $p = 0.07$ ).

For **H2** (building confidence), we analyzed ISE, Likert-scale survey responses, and our two supplementary experiments. Our strongest evidence came from the post-study Likert-scale survey. When asked to rate their trust in each method, participants rated Active-C-MoSH significantly higher than full ranking (adj.  $p = 0.01$ ), partial ranking (adj.  $p = 0.05$ ), and pairwise feedback (adj.  $p = 0.04$ ), providing direct support for H2 (results in Appendix E.3).

While ISE results showed a positive trend in the comparative study, the difference was not initially statistically significant. We attributed this to user unfamiliarity with the novel soft-hard bound mechanism, which was also rated as more complex (see **H3**). To test this hypothesis, our longitudinal study tracked ISE over four consecutive trials of Active-C-MoSH. By the last trial, once participants had familiarized themselves with the interface, Active-C-MoSH achieved an ISE score (mean=0.82) that was statistically significantly greater than that of all other baseline feedback mechanisms ( $p < 0.04$ ), confirming that the behavioral signal of confidence strengthens as the learning curve is overcome. Furthermore, to isolate the source of this confidence, our ablation study compared Active-C-MoSH against Active-MoSH. The full Active-C-MoSH yielded significantly higher ISE scores (mean=0.80 vs. 0.62,  $p=0.04$ ), confirming that the sensitivity analysis provided by C-MoSH is the primary driver for building decision confidence.

Finally, subjective Likert-scale survey responses addressed **H3** (cognitive demand) and **H4** (expressiveness). For cognitive load, higher scores indicate more effort. Active-C-MoSH (mean=3.75) was rated as requiring more mental effort compared to full ranking (mean=2.25, adj.  $p = 0.00$ ), partial ranking (mean=3.00, adj.  $p = 0.02$ ), and pairwise feedback (mean=2.75, adj.  $p = 0.00$ ). This is not unexpected, as Active-C-MoSH is a novel feedback structure and may require more familiarization than conceptually simpler mechanisms like pairwise or ranking. Another part of this may be attributed to the nature of providing exact numerical values using slider bars, which may require more cognitive load than discrete values for rankings or buttons for pairwise selection. For **H4**, *Active-C-MoSH* (mean=4.25) was rated as significantly more expressive than full ranking (mean=2.58, adj.  $p = 0.00$ ), partial ranking (mean=3.58, adj.  $p = 0.06$ ), and pairwise feedback (mean=3.20, adj.  $p = 0.02$ ).

## 8 CASE STUDY: CERVICAL CANCER

We evaluated Active-MoSH within a high-stakes, high-dimensional decision support tool for cervical cancer

brachytherapy treatment planning. The setup assumes trade-offs among six competing objectives: maximizing radiation dose to the tumor while minimizing exposure to five healthy regions [Chen et al., 2026a]. Using three retrospective patient datasets obtained from a local hospital, a domain expert attempted to identify treatment plans comparable to clinically approved reference plans within a strict 30-minute window. Note that standard manual planning typically requires 30-60 minutes Deufel et al. [2020], Michaud et al. [2016].

We quantified performance using a scalar utility score derived by normalizing each objective to  $[0, 1]$  (where higher is better), dividing by the corresponding reference plan metric, and averaging across dimensions. Active-MoSH achieved a mean utility score of  $1.11 \pm 0.02$  across the three cases, significantly outperforming full ranking ( $0.82 \pm 0.02$ ), partial ranking ( $0.80 \pm 0.01$ ), and pairwise feedback ( $0.80 \pm 0.01$ ).

To illustrate the clinical significance of these results, we focus on two critical objectives for clarity: tumor coverage (maximize) and Bladder dose (minimize). In one patient case with reference values of (95.6%, 457.8 cGy), the expert using Active-MoSH identified a dominating plan with metrics of (96.0%, 436.2 cGy). This solution offers improved tumor control while simultaneously reducing toxicity risks to healthy tissue. In contrast, using pairwise feedback, it proved difficult to identify a plan superior to (89.0%, 556.3 cGy). The bladder dose of 556.3 cGy approaches standard safety limits and could greatly increase the risk of severe complications [Romano et al., 2018, Carrara et al., 2017]. These findings underscore the necessity of expressive feedback mechanisms in safety-critical domains. Additional details for this case study are provided in Appendix F.

## 9 RELATED WORKS

**Interactive Feedback Mechanisms.** Research in interactive learning has explored diverse feedback modalities. These include learning from corrections Zhang and Dragan [2019], Losey and O'Malley [2018], ordinal feedback Chu and Ghahramani [2005a,b], Ozaki et al. [2023], Giovanelli et al. [2024], Wilde et al. [2022], choice functions Benavoli et al. [2023], and critiques Argall et al. [2007]. Within MOO, Ozaki et al. [2023] introduced improvement requests for specific objectives. Racca et al. [2020] used slider bars for interactive parameter tuning in robotics; however, their focus was on single-objective problems. To the best of our knowledge, Active-MoSH is the first interactive framework for MODM that directly allows iterative refinement of experts' dual-level preferences, specifically their aspirational targets (soft bounds) and non-negotiable limits (hard bounds).

**Active Sampling** generally aims to select the most informative queries to present to a user, thereby maximizing learning from a limited number of interactions [Biyik et al.,

2020]. One prominent line of research focuses on maximizing the volume removed from the hypothesis space (volume removal) [Sadigh et al., 2017, Palan et al., 2019, Biyik and Sadigh, 2018, Golovin and Krause, 2017, Biyik et al., 2019, Katz et al., 2019, Basu et al., 2019]. Another seeks to directly maximize the information gained from each query (e.g., via information gain criteria) [Golovin et al., 2010, Zheng et al., 2012, Houlsby et al., 2011, MacKay, 1992, Krishnapuram et al., 2004, Lawrence et al., 2002]. In contrast, our active sampling approach is specifically tailored for MODM by explicitly considering the high cost of DM validation for each presented point and aiming for diverse coverage of the DM’s evolving region of interest.

**Building Confidence in AI Systems.** The challenge of building user confidence, and trust, in AI systems is a long-standing area of research. Leveraging perturbations to enhance model interpretability, and consequently trust, has been explored in various contexts, such as influence functions and explainable AI [Cook, 1977, Koh and Liang, 2017, Ribeiro et al., 2016, Lundberg and Lee, 2017]. Various measures of trust have also been proposed. Schmidt and Biessmann [2019] measured trust via information transfer rates during ML annotation tasks, while Jiang et al. [2018] defined a trust score based on agreement between a classifier and a modified nearest-neighbor classifier to assess prediction reliability. Irshad et al. [2024] equates trust to fidelity in a multi-fidelity setting. To the best of our knowledge, we are the first to introduce and empirically validate a quantifiable notion of trust within an interactive MODM context.

*Additional related works may be found in Appendix G.*

## 10 CONCLUSION

We introduced Active-MoSH, an interactive multi-objective preference learning framework. Its *local component* uses probabilistic soft-hard bounds and active sampling to efficiently refine the search for optimal tradeoffs. The *global component*, C-MoSH, builds DM confidence via multi-objective sensitivity analysis to quantify robustness and uncover overlooked high-value regions. Evaluations across synthetic benchmarks, a cancer treatment case study, and a user study confirm Active-MoSH accelerates convergence to preferred solutions and boosts confidence over baselines. Key limitations (Appendix H) include computational scalability and reliance on consistent DM feedback.

## Acknowledgements

We thank Elizabeth Kidd, Thomas Niedermayr, and members of the Guestrin lab for helpful discussions throughout this project. This research was supported in part by the Stanford Institute for Human-Centered Artificial Intelligence and the Chan Zuckerberg Biohub. SK acknowledges support by

NSF 2046795 and 2205329, IES R305C240046, ARPA-H, the MacArthur Foundation, Schmidt Sciences, and HAI.

## References

- Streamlit, 2018. URL [streamlit.io](https://streamlit.io).
- Supabase, 2020. URL <https://supabase.com>.
- Pieter Abbeel and Andrew Y. Ng. Apprenticeship learning via inverse reinforcement learning. In *Twenty-first international conference on Machine learning - ICML '04*, page 1, Banff, Alberta, Canada, 2004. ACM Press. doi: 10.1145/1015330.1015430. URL <http://portal.acm.org/citation.cfm?doid=1015330.1015430>.
- Brenna Argall, Brett Browning, and Manuela Veloso. Learning by demonstration with critique from a human teacher. In *2007 2nd ACM/IEEE International Conference on Human-Robot Interaction (HRI)*, pages 57–64, March 2007. doi: 10.1145/1228716.1228725. URL <https://ieeexplore.ieee.org/document/6251717>.
- Raul Astudillo and Peter Frazier. Multi-attribute Bayesian optimization with interactive preference learning. In *Proceedings of the Twenty Third International Conference on Artificial Intelligence and Statistics*, pages 4496–4507. PMLR, June 2020. URL <https://proceedings.mlr.press/v108/astudillo20a.html>.
- Maximilian Balandat, Brian Karrer, Daniel R. Jiang, Samuel Daulton, Benjamin Letham, Andrew Gordon Wilson, and Eytan Bakshy. BoTorch: A Framework for Efficient Monte-Carlo Bayesian Optimization. In *Advances in Neural Information Processing Systems 33*, 2020. URL <http://arxiv.org/abs/1910.06403>.
- Chandrayee Basu, Erdem Biyik, Zhixun He, Mukesh Singhal, and Dorsa Sadigh. Active Learning of Reward Dynamics from Hierarchical Queries. In *2019 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pages 120–127, Macau, China, November 2019. IEEE. ISBN 978-1-72814-004-9. doi: 10.1109/IROS40897.2019.8968522. URL <https://ieeexplore.ieee.org/document/8968522/>.
- Alessio Benavoli, Dario Azzimonti, and Dario Piga. Learning Choice Functions with Gaussian Processes. In *Proceedings of the Thirty-Ninth Conference on Uncertainty in Artificial Intelligence*, pages 141–151. PMLR, July 2023. URL <https://proceedings.mlr.press/v216/benavoli23a.html>.
- Erdem Biyik, Malayandi Palan, Nicholas C. Landolfi, Dylan P. Losey, and Dorsa Sadigh. Asking Easy Questions: A User-Friendly Approach to Active Reward Learning. In *Proceedings of the Conference on Robot Learning*, pages 1177–1190. PMLR, May

2020. URL <https://proceedings.mlr.press/v100/b-iy-ik20a.html>.
- Tamara Broderick, Andrew Gelman, Rachael Meager, Anna L. Smith, and Tian Zheng. Toward a taxonomy of trust for probabilistic machine learning. *Science Advances*, 9(7):eabn3999, February 2023. doi: 10.1126/sciadv.abn3999. URL <https://www.science.org/doi/10.1126/sciadv.abn3999>.
- Erdem Bıyık and Dorsa Sadigh. Batch Active Preference-Based Learning of Reward Functions, October 2018. URL <http://arxiv.org/abs/1810.04303>. arXiv:1810.04303 [cs.LG].
- Erdem Bıyık, Daniel A. Lazar, Dorsa Sadigh, and Ramtin Pedarsani. The Green Choice: Learning and Influencing Human Decisions on Shared Roads. In *2019 IEEE 58th Conference on Decision and Control (CDC)*, pages 347–354, December 2019. doi: 10.1109/CDC40024.2019.9030169. URL <http://arxiv.org/abs/1904.02209>. arXiv:1904.02209 [math].
- Paolo Campigotto, Andrea Passerini, and Roberto Battiti. Active Learning of Pareto Fronts. *IEEE Transactions on Neural Networks and Learning Systems*, 25(3):506–519, March 2014. ISSN 2162-2388. doi: 10.1109/TNNLS.2013.2275918. URL <https://ieeexplore.ieee.org/document/6606803>. Conference Name: IEEE Transactions on Neural Networks and Learning Systems.
- Mauro Carrara, Davide Cusumano, Tommaso Gandini, Chiara Tenconi, Ester Mazzarella, Simone Grisotto, Eleonora Massari, Davide Mazzeo, Annamaria Cerrotta, Brigida Pappalardi, Carlo Fallai, and Emanuele Pignoli. Comparison of different treatment planning optimization methods for vaginal HDR brachytherapy with multichannel applicators: A reduction of the high doses to the vaginal mucosa is possible. *Physica Medica*, 44:58–65, December 2017. ISSN 1120-1797. doi: 10.1016/j.ejmp.2017.11.007. URL <https://www.sciencedirect.com/science/article/pii/S1120179717306002>.
- Edward Chen, Natalie Dullerud, Pang Wei Koh, Thomas Niedermayr, Elizabeth Kidd, Sanmi Koyejo, and Carlos Guestrin. ALMo: Interactive Aim-Limit-Defined, Multi-Objective System for Personalized High-Dose-Rate Brachytherapy Treatment Planning and Visualization for Cervical Cancer, February 2026a. URL <http://arxiv.org/abs/2602.13666>. arXiv:2602.13666 [cs].
- Edward Chen, Natalie Dullerud, Thomas Niedermayr, Elizabeth Kidd, Ransalu Senanayake, Pang Wei Koh, Sanmi Koyejo, and Carlos Guestrin. Modeling Multi-Objective Tradeoffs with Monotonic Utility Functions, April 2026b. URL <http://arxiv.org/abs/2412.06154>. arXiv:2412.06154 [cs.LG].
- F. Y. CHENG and X. S. LI. Generalized Center Method for Multiobjective Engineering Optimization. *Engineering Optimization*, 31(5):641–661, May 1999. ISSN 0305-215X. doi: 10.1080/03052159908941390. URL <https://doi.org/10.1080/03052159908941390>. \_eprint: <https://doi.org/10.1080/03052159908941390>.
- Siddhartha Chib, , and Edward Greenberg. Understanding the Metropolis-Hastings Algorithm. *The American Statistician*, 49(4):327–335, November 1995. ISSN 0003-1305. doi: 10.1080/00031305.1995.10476177.
- Wei Chu and Zoubin Ghahramani. Gaussian Processes for Ordinal Regression. *Journal of Machine Learning Research*, 6(35):1019–1041, 2005a. ISSN 1533-7928. URL <http://jmlr.org/papers/v6/chu05a.html>.
- Wei Chu and Zoubin Ghahramani. Preference learning with Gaussian processes. In *Proceedings of the 22nd international conference on Machine learning - ICML '05*, pages 137–144, Bonn, Germany, 2005b. ACM Press. ISBN 978-1-59593-180-1. doi: 10.1145/1102351.1102369. URL <http://portal.acm.org/citation.cfm?doid=1102351.1102369>.
- Jonas Clausen, John Cantwell, and Philosophy Documentation Center. Reasoning With Safety Factor Rules. *Techné: Research in Philosophy and Technology*, 11(1):55–70, 2007. ISSN 2691-5928. doi: 10.5840/techne20071116. URL [http://www.pdcnet.org/oom/service?url\\_ver=Z39.88-2004&rft\\_val\\_fmt=&rft.imuse\\_id=techne\\_2007\\_0011\\_0001\\_0055\\_0070&svc\\_id=info:www.pdcnet.org/collection](http://www.pdcnet.org/oom/service?url_ver=Z39.88-2004&rft_val_fmt=&rft.imuse_id=techne_2007_0011_0001_0055_0070&svc_id=info:www.pdcnet.org/collection).
- R. Dennis Cook. Detection of Influential Observation in Linear Regression. *Technometrics*, 19(1):15–18, 1977. ISSN 0040-1706. doi: 10.2307/1268249. URL <https://www.jstor.org/stable/1268249>.
- Songye Cui, Philippe Després, and Luc Beaulieu. A multi-criteria optimization approach for HDR prostate brachytherapy: II. Benchmark against clinical plans. *Physics in Medicine & Biology*, 63(20):205005, October 2018a. ISSN 1361-6560. doi: 10.1088/1361-6560/aae24f. URL <https://iopscience.iop.org/article/10.1088/1361-6560/aae24f>.
- Songye Cui, Philippe Després, and Luc Beaulieu. A multi-criteria optimization approach for HDR prostate brachytherapy: I. Pareto surface approximation. *Physics in Medicine & Biology*, 63(20):205004, October 2018b. ISSN 1361-6560. doi: 10.1088/1361-6560/aae24c. URL <https://iopscience.iop.org/article/10.1088/1361-6560/aae24c>.
- Tri Dao. FlashAttention-2: Faster Attention with Better Parallelism and Work Partitioning, July 2023.

- URL <http://arxiv.org/abs/2307.08691>. arXiv:2307.08691 [cs].
- Kalyanmoy Deb and J. Sundar. Reference point based multi-objective optimization using evolutionary algorithms. In *Proceedings of the 8th annual conference on Genetic and evolutionary computation, GECCO '06*, pages 635–642, New York, NY, USA, July 2006. Association for Computing Machinery. ISBN 978-1-59593-186-3. doi: 10.1145/1143997.1144112. URL <https://dl.acm.org/doi/10.1145/1143997.1144112>.
- Kobe Desender, Annika Boldt, and Nick Yeung. Subjective Confidence Predicts Information Seeking in Decision Making. *Psychological Science*, 29(5):761–778, May 2018. ISSN 0956-7976. doi: 10.1177/0956797617744771. URL <https://doi.org/10.1177/0956797617744771>.
- Kobe Desender, Peter Murphy, Annika Boldt, Tom Verguts, and Nick Yeung. A Postdecisional Neural Marker of Confidence Predicts Information-Seeking in Decision-Making. *The Journal of Neuroscience*, 39(17):3309–3319, April 2019. ISSN 0270-6474. doi: 10.1523/JNEUROSCI.2620-18.2019. URL <https://pmc.ncbi.nlm.nih.gov/articles/PMC6788827/>.
- Christopher L. Deufel, Marina A. Epelman, Kalyan S. Pasupathy, Mustafa Y. Sir, Victor W. Wu, and Michael G. Herman. PNaV: A tool for generating a high-dose-rate brachytherapy treatment plan by navigating the Pareto surface guided by the visualization of multidimensional trade-offs. *Brachytherapy*, 19(4):518–531, July 2020. ISSN 15384721. doi: 10.1016/j.brachy.2020.02.013. URL <https://linkinghub.elsevier.com/retrieve/pii/S1538472120300325>.
- Dirk M. Elston. The novelty effect. *Journal of the American Academy of Dermatology*, 85(3):565–566, September 2021. ISSN 0190-9622, 1097-6787. doi: 10.1016/j.jaad.2021.06.846. URL [https://www.jaad.org/article/S0190-9622\(21\)01987-3/fulltext](https://www.jaad.org/article/S0190-9622(21)01987-3/fulltext).
- Michael Emmerich. The computation of the expected improvement in dominated hypervolume of Pareto front approximations. January 2008.
- Jenna C. Fromer and Connor W. Coley. Computer-aided multi-objective optimization in small molecule discovery. *Patterns*, 4(2):100678, February 2023. ISSN 26663899. doi: 10.1016/j.patter.2023.100678. URL <https://linkinghub.elsevier.com/retrieve/pii/S2666389923000016>.
- Joseph Giovannelli, Alexander Tornede, Tanja Tornede, and Marius Lindauer. Interactive Hyperparameter Optimization in Multi-Objective Problems via Preference Learning. *Proceedings of the AAAI Conference on Artificial Intelligence*, 38(11):12172–12180, March 2024. ISSN 2374-3468. doi: 10.1609/aaai.v38i11.29106. URL <https://ojs.aaai.org/index.php/AAAI/article/view/29106>.
- Daniel Golovin and Andreas Krause. Adaptive Submodularity: Theory and Applications in Active Learning and Stochastic Optimization, December 2017. URL <http://arxiv.org/abs/1003.3967>. arXiv:1003.3967 [cs].
- Daniel Golovin, Andreas Krause, and Debajyoti Ray. Near-optimal Bayesian active learning with noisy observations. In *Proceedings of the 24th International Conference on Neural Information Processing Systems - Volume 1, volume 1 of NIPS'10*, pages 766–774, Red Hook, NY, USA, December 2010. Curran Associates Inc.
- David Sierra González, Ozgur Erkent, Víctor Romero-Cano, Jilles Dibangoye, and Christian Laugier. Modeling Driver Behavior from Demonstrations in Dynamic Environments Using Spatiotemporal Lattices. In *2018 IEEE International Conference on Robotics and Automation (ICRA)*, pages 3384–3390, May 2018. doi: 10.1109/ICRA.2018.8460208. URL <https://ieeexplore.ieee.org/document/8460208>.
- John Guiver and Edward Snelson. Bayesian inference for Plackett-Luce ranking models. In *Proceedings of the 26th Annual International Conference on Machine Learning, ICML '09*, pages 377–384, New York, NY, USA, June 2009. Association for Computing Machinery. ISBN 978-1-60558-516-1. doi: 10.1145/1553374.1553423. URL <https://dl.acm.org/doi/10.1145/1553374.1553423>.
- Daniel Hernández-Lobato, José Miguel Hernández-Lobato, Amar Shah, and Ryan P. Adams. Predictive Entropy Search for Multi-objective Bayesian Optimization, February 2016. URL <http://arxiv.org/abs/1511.05467>. arXiv:1511.05467 [stat].
- Sture Holm. A Simple Sequentially Rejective Multiple Test Procedure. *Scandinavian Journal of Statistics*, 6(2):65–70, 1979. ISSN 0303-6898. URL <https://www.jstor.org/stable/4615733>.
- Neil Houlsby, Ferenc Huszár, Zoubin Ghahramani, and Máté Lengyel. Bayesian Active Learning for Classification and Preference Learning, December 2011. URL <http://arxiv.org/abs/1112.5745>. arXiv:1112.5745 [stat].
- David Rios Insua and Simon French. A framework for sensitivity analysis in discrete multi-objective decision-making. *European Journal of Operational Research*, 54(2):176–190, September 1991. ISSN 0377-2217. doi: 10.1016/0377-2217(91)90296-8.

URL <https://www.sciencedirect.com/science/article/pii/S0377221791902968>.

- David Ríos Insua. Sensitivity Analysis in Multi-objective Decision Making. In David Ríos Insua, editor, *Sensitivity Analysis in Multi-objective Decision Making*, pages 74–126. Springer, Berlin, Heidelberg, 1990. ISBN 978-3-642-51656-6. doi: 10.1007/978-3-642-51656-6\_3. URL [https://doi.org/10.1007/978-3-642-51656-6\\_3](https://doi.org/10.1007/978-3-642-51656-6_3).
- Faran Irshad, Stefan Karsch, and Andreas Döpp. Leveraging Trust for Joint Multi-Objective and Multi-Fidelity Optimization. *Machine Learning: Science and Technology*, 5(1):015056, March 2024. ISSN 2632-2153. doi: 10.1088/2632-2153/ad35a4. URL <http://arxiv.org/abs/2112.13901>. arXiv:2112.13901 [cs].
- Hamish Ivison, Yizhong Wang, Valentina Pyatkin, Nathan Lambert, Matthew Peters, Pradeep Dasigi, Joel Jang, David Wadden, Noah A. Smith, Iz Beltagy, and Hannaneh Hajishirzi. Camels in a Changing Climate: Enhancing LM Adaptation with Tulu 2, November 2023. URL <http://arxiv.org/abs/2311.10702>. arXiv:2311.10702 [cs].
- Kevin G. Jamieson and Robert D. Nowak. Active ranking using pairwise comparisons. In *Proceedings of the 25th International Conference on Neural Information Processing Systems*, NIPS’11, pages 2240–2248, Red Hook, NY, USA, December 2011. Curran Associates Inc. ISBN 978-1-61839-599-3.
- Heinrich Jiang, Been Kim, Melody Y. Guan, and Maya Gupta. To trust or not to trust a classifier. In *Proceedings of the 32nd International Conference on Neural Information Processing Systems*, NIPS’18, pages 5546–5557, Red Hook, NY, USA, December 2018. Curran Associates Inc. URL <https://dl.acm.org/doi/10.5555/3327345.3327458>.
- Amelia M. Johnson, Micheal A. Buice, and Koosha Khalvati. A Unifying Normative Framework of Decision Confidence. In A. Globerson, L. Mackey, D. Belgrave, A. Fan, U. Paquet, J. Tomczak, and C. Zhang, editors, *Advances in Neural Information Processing Systems*, volume 37, pages 140847–140864. Curran Associates, Inc., 2024. doi: 10.52202/079017-4471.
- Donald R. Jones, Matthias Schonlau, and William J. Welch. Efficient Global Optimization of Expensive Black-Box Functions. *Journal of Global Optimization*, 13(4):455–492, December 1998. ISSN 1573-2916. doi: 10.1023/A:1008306431147. URL <https://doi.org/10.1023/A:1008306431147>.
- Daniel Kahneman and Amos Tversky. Prospect Theory: An Analysis of Decision Under Risk. In *Handbook of the Fundamentals of Financial Decision Making*, volume Volume 4 of *World Scientific Handbook in Financial Economics Series*, pages 99–127. WORLD SCIENTIFIC, June 2012. ISBN 978-981-4417-34-1. doi: 10.1142/9789814417358\_0006. URL [https://www.worldscientific.com/doi/abs/10.1142/9789814417358\\_0006](https://www.worldscientific.com/doi/abs/10.1142/9789814417358_0006).
- Sydney M. Katz, Anne-Claire Le Bihan, and Mykel J. Kochenderfer. Learning an Urban Air Mobility Encounter Model from Expert Preferences, July 2019. URL <http://arxiv.org/abs/1907.05575>. arXiv:1907.05575 [cs].
- J. Knowles. ParEGO: a hybrid algorithm with on-line landscape approximation for expensive multiobjective optimization problems. *IEEE Transactions on Evolutionary Computation*, 10(1):50–66, February 2006. ISSN 1941-0026. doi: 10.1109/TEVC.2005.851274. URL <https://ieeexplore.ieee.org/document/1583627>. Conference Name: IEEE Transactions on Evolutionary Computation.
- Pang Wei Koh and Percy Liang. Understanding Black-box Predictions via Influence Functions. In *Proceedings of the 34th International Conference on Machine Learning*, pages 1885–1894. PMLR, July 2017. URL <https://proceedings.mlr.press/v70/koh17a.html>.
- Andreas Krause, H. Brendan McMahan, Carlos Guestrin, and Anupam Gupta. Robust Submodular Observation Selection. *Journal of Machine Learning Research*, 9(93):2761–2801, 2008. ISSN 1533-7928. URL <http://jmlr.org/papers/v9/krause08b.html>.
- Balaji Krishnapuram, David Williams, Ya Xue, Lawrence Carin, Mário Figueiredo, and Alexander Hartemink. On Semi-Supervised Classification. In *Advances in Neural Information Processing Systems*, volume 17. MIT Press, 2004.
- Neil Lawrence, Matthias Seeger, and Ralf Herbrich. Fast Sparse Gaussian Process Methods: The Informative Vector Machine. In *Advances in Neural Information Processing Systems*, volume 15. MIT Press, 2002.
- Alisa Liu, Xiaochuang Han, Yizhong Wang, Yulia Tsvetkov, Yejin Choi, and Noah A. Smith. Tuning Language Models by Proxy, August 2024. URL <http://arxiv.org/abs/2401.08565>. arXiv:2401.08565 [cs].
- Dylan P. Losey and Marcia K. O’Malley. Including Uncertainty when Learning from Human Corrections. In *Proceedings of The 2nd Conference on Robot Learning*, pages 123–132. PMLR, October 2018. URL <https://proceedings.mlr.press/v87/losey18a.html>.

- R. Duncan Luce. *Individual choice behavior*. Individual choice behavior. John Wiley, Oxford, England, 1959. Pages: xii, 153.
- Scott M. Lundberg and Su-In Lee. A unified approach to interpreting model predictions. In *Proceedings of the 31st International Conference on Neural Information Processing Systems*, NIPS’17, pages 4768–4777, Red Hook, NY, USA, December 2017. Curran Associates Inc. ISBN 978-1-5108-6096-4. URL <https://dl.acm.org/doi/10.5555/3295222.3295230>.
- Sohvi Luukkonen, Helle W. van den Maagdenberg, Michael T. M. Emmerich, and Gerard J. P. van Westen. Artificial intelligence in multi-objective drug design. *Current Opinion in Structural Biology*, 79:102537, April 2023. ISSN 0959-440X. doi: 10.1016/j.sbi.2023.102537. URL <https://www.sciencedirect.com/science/article/pii/S0959440X23000118>.
- David J. C. MacKay. Information-Based Objective Functions for Active Data Selection. *Neural Computation*, 4(4):590–604, July 1992. ISSN 0899-7667, 1530-888X. doi: 10.1162/neco.1992.4.4.590. URL <https://direct.mit.edu/neco/article/4/4/590-604/5648>.
- Gustavo Malkomes, Bolong Cheng, Eric H. Lee, and Mike Mccourt. Beyond the Pareto Efficient Frontier: Constraint Active Search for Multiobjective Experimental Design. In *Proceedings of the 38th International Conference on Machine Learning*, pages 7423–7434. PMLR, July 2021. URL <https://proceedings.mlr.press/v139/malkomes21a.html>.
- Anthony L. Michaud, Stanley Benedict, Eliseo Montemayor, Jon Paul Hunt, Cari Wright, Mathew Mathai, and Jyoti S. Mayadev. Workflow efficiency for the treatment planning process in CT-guided high-dose-rate brachytherapy for cervical cancer. *Brachytherapy*, 15(5):578–583, September 2016. ISSN 1538-4721. doi: 10.1016/j.brachy.2016.06.001. URL <https://www.sciencedirect.com/science/article/pii/S1538472116304925>.
- Eric Mitchell, Rafael Rafailov, Archit Sharma, Chelsea Finn, and Christopher D. Manning. An Emulator for Fine-Tuning Large Language Models using Small Language Models, October 2023. URL <http://arxiv.org/abs/2310.12962>. arXiv:2310.12962 [cs].
- Vivek Myers, Erdem Biyik, Nima Anari, and Dorsa Sadigh. Learning Multimodal Rewards from Rankings. In *Proceedings of the 5th Conference on Robot Learning*, pages 342–352. PMLR, January 2022. URL <https://proceedings.mlr.press/v164/myers22a.html>.
- Ananya Nandy, David Antonio Herrera, and Kosa Goucher-Lambert. Adopting “blackbox” engineering advice: the influence of imperfect suggestions during AI-assisted decision-making with multiple objectives. *AI EDAM*, 39:e7, January 2025. ISSN 0890-0604, 1469-1760. doi: 10.1017/S0890060425000034.
- G. L. Nemhauser, L. A. Wolsey, and M. L. Fisher. An analysis of approximations for maximizing submodular set functions—I. *Mathematical Programming*, 14(1):265–294, December 1978. ISSN 0025-5610, 1436-4646. doi: 10.1007/BF01588971. URL <http://link.springer.com/10.1007/BF01588971>.
- Lorien Nesbitt, Michael J. Meitner, Brent Chamberlain, Julian Gonzalez, and William Trousdale. A Comparison of Value-Weight-Elicitation Methods for Accurate and Accessible Participatory Planning. *Journal of Planning Education and Research*, 44(4):2098–2109, December 2024. ISSN 0739-456X. doi: 10.1177/0739456X231155069. URL <https://doi.org/10.1177/0739456X231155069>.
- Ryota Ozaki, Kazuki Ishikawa, Youhei Kanzaki, Shinya Suzuki, Shion Takeno, Ichiro Takeuchi, and Masayuki Karasuyama. Multi-Objective Bayesian Optimization with Active Preference Learning, November 2023. URL <http://arxiv.org/abs/2311.13460>. arXiv:2311.13460 [cs].
- Malayandi Palan, Nicholas C. Landolfi, Gleb Shevchuk, and Dorsa Sadigh. Learning Reward Functions by Integrating Human Demonstrations and Preferences, June 2019. URL <http://arxiv.org/abs/1906.08928>. arXiv:1906.08928 [cs].
- Christos H. Papadimitriou and Mihalis Yannakakis. Multiobjective query optimization. In *Proceedings of the twentieth ACM SIGMOD-SIGACT-SIGART symposium on Principles of database systems*, PODS ’01, pages 52–59, New York, NY, USA, May 2001. Association for Computing Machinery. ISBN 978-1-58113-361-5. doi: 10.1145/375551.375560. URL <https://dl.acm.org/doi/10.1145/375551.375560>.
- Biswajit Paria, Kirthevasan Kandasamy, and Barnabás Póczos. A Flexible Framework for Multi-Objective Bayesian Optimization using Random Scalarizations, June 2019. URL <http://arxiv.org/abs/1805.12168>. arXiv:1805.12168 [cs, stat].
- Victor Picheny. Multiobjective optimization using Gaussian process emulators via stepwise uncertainty reduction. *Statistics and Computing*, 25(6):1265–1280, November 2015. ISSN 1573-1375. doi: 10.1007/s11222-014-9477-x. URL <https://doi.org/10.1007/s11222-014-9477-x>.

- R. L. Plackett. The Analysis of Permutations. *Journal of the Royal Statistical Society Series C*, 24(2):193–202, 1975. URL <https://ideas.repec.org/a/bla/jorssc/v24y1975i2p193-202.html>.
- Dustin Podell, Zion English, Kyle Lacey, Andreas Blattmann, Tim Dockhorn, Jonas Müller, Joe Penna, and Robin Rombach. SDXL: Improving Latent Diffusion Models for High-Resolution Image Synthesis, July 2023. URL <http://arxiv.org/abs/2307.01952>. arXiv:2307.01952 [cs].
- Wolfgang Ponweiser, Tobias Wagner, Dirk Biermann, and Markus Vincze. Multiobjective Optimization on a Limited Budget of Evaluations Using Model-Assisted  $\mathcal{S}$ -Metric Selection. In Günter Rudolph, Thomas Jansen, Nicola Beume, Simon Lucas, and Carlo Poloni, editors, *Parallel Problem Solving from Nature – PPSN X*, pages 784–794, Berlin, Heidelberg, 2008. Springer. ISBN 978-3-540-87700-4. doi: 10.1007/978-3-540-87700-4\_78.
- Mattia Racca, Ville Kyrki, and Maya Cakmak. Interactive Tuning of Robot Program Parameters via Expected Divergence Maximization. In *Proceedings of the 2020 ACM/IEEE International Conference on Human-Robot Interaction*, HRI '20, pages 629–638, New York, NY, USA, March 2020. Association for Computing Machinery. ISBN 978-1-4503-6746-2. doi: 10.1145/3319502.3374784. URL <https://dl.acm.org/doi/10.1145/3319502.3374784>.
- Rafael Rafailov, Archit Sharma, Eric Mitchell, Stefano Ermon, Christopher D. Manning, and Chelsea Finn. Direct Preference Optimization: Your Language Model is Secretly a Reward Model, July 2024. URL <http://arxiv.org/abs/2305.18290>. arXiv:2305.18290 [cs].
- Alfred Rappaport. Sensitivity Analysis in Decision Making. *The Accounting Review*, 42(3):441–456, 1967. ISSN 0001-4826. URL <https://www.jstor.org/stable/243710>.
- Marco Tulio Ribeiro, Sameer Singh, and Carlos Guestrin. "Why Should I Trust You?": Explaining the Predictions of Any Classifier. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, KDD '16, pages 1135–1144, New York, NY, USA, August 2016. Association for Computing Machinery. ISBN 978-1-4503-4232-2. doi: 10.1145/2939672.2939778. URL <https://dl.acm.org/doi/10.1145/2939672.2939778>.
- D. M. Roijers, P. Vamplew, S. Whiteson, and R. Dazeley. A Survey of Multi-Objective Sequential Decision-Making. *Journal of Artificial Intelligence Research*, 48:67–113, October 2013. ISSN 1076-9757. doi: 10.1613/jair.3987. URL <https://www.jair.org/index.php/jair/article/view/10836>.
- Kara D. Romano, Colin Hill, Daniel M. Trifiletti, M. Sean Peach, Bethany J. Horton, Neil Shah, Dylan Campbell, Bruce Libby, and Timothy N. Showalter. High dose-rate tandem and ovoid brachytherapy in cervical cancer: dosimetric predictors of adverse events. *Radiation Oncology (London, England)*, 13:129, July 2018. ISSN 1748-717X. doi: 10.1186/s13014-018-1074-2. URL <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC6048838/>.
- Dorsa Sadigh, Anca Dragan, Shankar Sastry, and Sanjit Seshia. Active Preference-Based Learning of Reward Functions. In *Robotics: Science and Systems XIII*. Robotics: Science and Systems Foundation, July 2017. ISBN 978-0-9923747-3-0. doi: 10.15607/RSS.2017.XIII.053. URL <http://www.roboticsproceedings.org/rss13/p53.pdf>.
- Philipp Schmidt and Felix Biessmann. Quantifying Interpretability and Trust in Machine Learning Systems, January 2019. URL <http://arxiv.org/abs/1901.08558>. arXiv:1901.08558 [cs].
- Omar Shaikh, Michelle S. Lam, Joey Hejna, Yijia Shao, Hyundong Cho, Michael S. Bernstein, and Diyi Yang. Aligning Language Models with Demonstrated Feedback, June 2024. URL <https://arxiv.org/abs/2406.00888v2>.
- Ruizhe Shi, Yifang Chen, Yushi Hu, Alisa Liu, Hannaneh Hajishirzi, Noah A. Smith, and Simon S. Du. Decoding-time language model alignment with multiple objectives. In *Proceedings of the 38th International Conference on Neural Information Processing Systems*, volume 37 of *NIPS '24*, pages 48875–48920, Red Hook, NY, USA, December 2024. Curran Associates Inc. ISBN 979-8-3313-1438-5.
- Ryoji Tanabe and Hisao Ishibuchi. An easy-to-use real-world multi-objective optimization problem suite. *Applied Soft Computing*, 89:106078, April 2020. ISSN 1568-4946. doi: 10.1016/j.asoc.2020.106078. URL <https://www.sciencedirect.com/science/article/pii/S1568494620300181>.
- Yash Vesikar, Kalyanmoy Deb, and Julian Blank. Reference Point Based NSGA-III for Preferred Solutions. In *2018 IEEE Symposium Series on Computational Intelligence (SSCI)*, pages 1587–1594, Bangalore, India, November 2018. IEEE. ISBN 978-1-5386-9276-9. doi: 10.1109/SSCI.2018.8628819. URL <https://ieeexplore.ieee.org/document/8628819/>.
- Nils Wilde, Erdem Biyik, Dorsa Sadigh, and Stephen L. Smith. Learning Reward Functions from Scale

- Feedback. In *Proceedings of the 5th Conference on Robot Learning*, pages 353–362. PMLR, January 2022. URL <https://proceedings.mlr.press/v164/wilde22a.html>.
- Christopher Williams and Carl Rasmussen. Gaussian Processes for Regression. In *Advances in Neural Information Processing Systems*, volume 8. MIT Press, 1995.
- Christian Wirth, Riad Akrou, Gerhard Neumann, and Johannes Fürnkranz. A Survey of Preference-Based Reinforcement Learning Methods. *Journal of Machine Learning Research*, 18(136):1–46, 2017. ISSN 1533-7928. URL <http://jmlr.org/papers/v18/16-634.html>.
- Wan-Ching Wu and Diane Kelly. Online search stopping behaviors: An investigation of query abandonment and task stopping. *Proceedings of the American Society for Information Science and Technology*, 51(1):1, 2014. ISSN 0044-7870. doi: 10.1002/MEET.2014.14505101030. URL [https://www.academia.edu/97882871/Online\\_search\\_stopping\\_behaviors\\_An\\_investigation\\_of\\_query\\_abandonment\\_and\\_task\\_stopping](https://www.academia.edu/97882871/Online_search_stopping_behaviors_An_investigation_of_query_abandonment_and_task_stopping).
- Yutong Xie, Chence Shi, Hao Zhou, Yuwei Yang, Weinan Zhang, Yong Yu, and Lei Li. MARS: Markov Molecular Sampling for Multi-objective Drug Discovery, March 2021. URL <http://arxiv.org/abs/2103.10432>. arXiv:2103.10432 [cs, q-bio].
- Amy Zhao Yu, Shahar Ronen, Kevin Hu, Tiffany Lu, and César A. Hidalgo. Pantheon 1.0, a manually verified dataset of globally famous biographies. *Scientific Data*, 3(1):150075, January 2016. ISSN 2052-4463. doi: 10.1038/sdata.2015.75. URL <https://www.nature.com/articles/sdata201575>.
- Yan Yu, J. B Zhang, Gang Cheng, M. C Schell, and Paul Okunieff. Multi-objective optimization in radiotherapy: applications to stereotactic radiosurgery and prostate brachytherapy. *Artificial Intelligence in Medicine*, 19(1):39–51, May 2000. ISSN 0933-3657. doi: 10.1016/S0933-3657(99)00049-4. URL <https://www.sciencedirect.com/science/article/pii/S0933365799000494>.
- Jason Y. Zhang and Anca D. Dragan. Learning from Extrapolated Corrections, March 2019. URL <http://arxiv.org/abs/1812.01225>. arXiv:1812.01225 [cs].
- Qingfu Zhang, Aimin Zhou, Shizheng Zhao, Ponnuthurai Nagaratnam Suganthan, Wudong Liu, and Santosh Tiwari. Multiobjective optimization Test Instances for the CEC 2009 Special Session and Competition. 2009.
- Alice X. Zheng, Irina Rish, and Alina Beygelzimer. Efficient Test Selection in Active Diagnosis via Entropy Approximation, July 2012. URL <http://arxiv.org/abs/1207.1418>. arXiv:1207.1418 [cs].
- Brian D. Ziebart, Andrew Maas, J. Andrew Bagnell, and Anind K. Dey. Maximum entropy inverse reinforcement learning. In *Proceedings of the 23rd national conference on Artificial intelligence - Volume 3, AAAI’08*, pages 1433–1438, Chicago, Illinois, July 2008. AAAI Press. ISBN 978-1-57735-368-3.
- Luisa Zintgraf, Sebastian Schulze, Cong Lu, Leo Feng, Maximilian Igl, Kyriacos Shiarlis, Yarin Gal, Katja Hofmann, and Shimon Whiteson. VariBAD: variational Bayes-adaptive deep RL via meta-learning. *The Journal of Machine Learning Research*, 22(1):289:13198–289:13236, January 2021. ISSN 1532-4435. URL <https://dl.acm.org/doi/10.5555/3546258.3546547>.
- Marcela Zuluaga, Andreas Krause, and Markus Püschel. e-PAL: An Active Learning Approach to the Multi-Objective Optimization Problem. *Journal of Machine Learning Research*, 17(104):1–32, 2016. ISSN 1533-7928. URL <http://jmlr.org/papers/v17/15-047.html>.

---

# Interactive Multi-Objective Probabilistic Preference Learning with Soft and Hard Bounds (Supplementary Material)

---

Edward Chen<sup>1</sup>

Sang T. Truong<sup>1</sup>

Natalie Dullerud<sup>1</sup>

Sanmi Koyejo<sup>1</sup>

Carlos Guestrin<sup>1</sup>

<sup>1</sup>Department of Computer Science, Stanford University, Stanford, California, USA

## A MULTI-OBJECTIVE PREFERENCE LEARNING WITH SOFT-HARD BOUNDS

### A.1 SOFT-HARD UTILITY FUNCTIONS

We follow Chen et al. [2026b] and use a similar form for the soft-hard utility functions, as follows:

$$u_{\alpha}(x) = \begin{cases} 1 + 2\beta \times (\tilde{\alpha}_{\tau} - \tilde{\alpha}_S) & f(x) \geq \alpha_{\tau} \\ 1 + 2\beta \times (\tilde{f}(x) - \tilde{\alpha}_S) & \alpha_S < f(x) < \alpha_{\tau} \\ 1 & f(x) = \alpha_S \\ 2 \times \tilde{f}(x) & \alpha_H < f(x) < \alpha_S \\ 0 & f(x) = \alpha_H \\ -\infty & f(x) < \alpha_H \end{cases} \quad (10)$$

where  $\tilde{f}(x)$  and  $\tilde{\alpha}$  are the soft-hard bound normalized values,  $\alpha_{\tau}$ , the saturation point, determines where the utility values begin to saturate, and  $\beta \in [0, 1]$  determines the fraction of the original rate of utility, in  $[\alpha_H, \alpha_S]$ , obtained within  $[\alpha_S, \alpha_{\tau}]$ . Normalization, for value  $z$ , is performed according to the soft and hard bounds,  $\alpha_S$  and  $\alpha_H$ , respectively, using:  $\tilde{z} = ((z - \alpha_H) / (\alpha_S - \alpha_H)) * 0.5$ . In practice, we determine  $\alpha_{\tau}$  to be  $\alpha_H + \zeta(\alpha_S - \alpha_H)$ , for  $\zeta = 3.9$ . We use  $\zeta = 3.9$ , as opposed to  $\zeta = 2.0$  (as described in Chen et al. [2026b]) since we found a higher value to be more robust to optimization in interactive settings. Figure 5 depicts an example of the soft-hard utility function  $u_{\alpha}(x)$ . For our experiments, we used  $\beta = 0.25$ . Additional details may be found in [Chen et al., 2026b].

## B ACTIVE-MOSH: PROBABILISTIC MODELING AND SAMPLING

### B.1 MODELING SOFT AND HARD BOUNDS FEEDBACK

We motivate the specific ranking interpretations for different feedback scenarios (illustrated in Figure 2) as follows:

1. **Hard Bound Tightened** ( $\alpha_{\ell, H, m} > \alpha_{\ell, H, m-1}$ ): When the DM makes a hard bound more restrictive (e.g., requiring a *higher* minimum for tumor coverage), they are effectively shrinking the feasible space. We interpret this as a strong signal that their interest lies in solutions that satisfy this new, more stringent requirement. The most informative points in  $Y_m$  are those that now lie near, or just satisfy, this new hard bound. Therefore, points are ranked based on their values in the adjusted dimension: those that meet or exceed the new  $\alpha_{\ell, H, m}$  are prioritized, with ranking determined by proximity to  $\alpha_{\ell, H, m}$ . The critical signal is that the DM wants to ensure this new hard constraint is met, and solutions near this boundary are key to understanding the tradeoffs under this tighter condition.
2. **Hard Bound Relaxed** ( $\alpha_{\ell, H, m} < \alpha_{\ell, H, m-1}$ ): Conversely, when the DM relaxes a hard bound (e.g., accepting a *lower* minimum tumor coverage), they are expanding the feasible region. This action signals an interest in exploring solutions

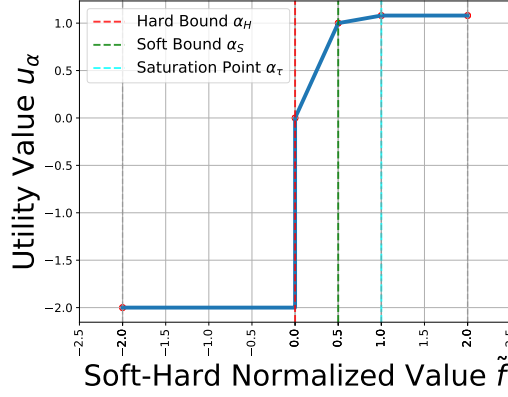


Figure 5: Example of a normalized soft-hard bounded utility function. The dashed vertical bars highlight the regions where the values correspond to the hard, soft, and saturated regions. The utility value associated with points below the hard bound is shown as -2 for illustration purposes only.

within this newly accessible part of the tradeoff space, perhaps because previous constraints were too restrictive and prevented desirable solutions from being considered. The critical signal is the DM’s willingness to explore what was previously infeasible. Points in  $Y_m$  that fall within this newly opened region, or are close to the now-relaxed  $\alpha_{\ell,H,m}$ , become highly relevant. Ranking therefore prioritizes points that benefit from this relaxation, i.e. those closer to the new (more permissive)  $\alpha_{\ell,H,m}$  within the expanded feasible space.

3. **Soft Bound Adjusted** ( $\alpha_{\ell,S,m} \neq \alpha_{\ell,S,m-1}$ ): An adjustment to a soft bound represents a direct refinement of the DM’s aspirational target for an objective. Unlike hard bounds which define feasibility, soft bounds define desirability. Therefore, when a DM moves a soft bound, we take this as an explicit indication that they are now most interested in solutions whose value for that objective is as close as possible to the new  $\alpha_{\ell,S,m}$ . The critical signal is precisely this new target value. Consequently, points in  $Y_m$  are ranked directly by their Euclidean distance to the adjusted  $\alpha_{\ell,S,m}$  in the modified dimension—the closer a point is, the higher its rank.
4. **No Change in Bounds**: If the DM makes no changes to any bounds after reviewing  $Y_m$ , we interpret this as an implicit confirmation that their current preference model (including the existing bound distributions) adequately captures their region of interest, or at least that  $Y_m$  did not provide sufficient new information to warrant an adjustment. In our probabilistic framework (Section 3.1), this lack of change translates to an update that places higher certainty (i.e., reduces variance) on the existing parameters of the unmodified bounds’ distributions, reinforcing the current belief about the DM’s preferences.

**Implementation Details.** In practice, the ranking of points in  $Y_m$  is determined by their Euclidean distance to the bound being modified at the end of iteration  $m$ , following the intuition described in Section 3.2. The prior distribution over the preferences,  $p(\lambda)$ , is set to a uniform distribution. To compute the posterior distribution over the preferences (Equation 3), we leveraged the Metropolis-Hastings algorithm with 20 burn-in steps and with the Dirichlet distribution as the proposal distribution Chib et al. [1995]. To implement the steps described in Section 3.3 to obtain the preference queries, we leveraged the BoTorch library v0.10.0 [Balandat et al., 2020]. *All details are also provided in our provided code.*

## B.2 ACTIVE SAMPLING FOR PREFERENCE QUERIES

Additional details about the active sampling method may be found below.

### Step 1: Dense Set of Points.

For our experiments, we follow Chen et al. [2026b] and use the Upper Confidence Bound (UCB) heuristic. We define  $\text{acq}(u, \lambda_t, \alpha_t, x) = s_{\lambda_t}(u_{\alpha_t(x)})$  where  $f(x) = \mu_t(x) + \sqrt{\beta_t} \sigma_t(x)$  (Equation 10) and  $\beta_t = \sqrt{0.125 \times \log(2 \times t + 1)}$ . For  $\beta_t$ , we followed the optimal suggestion in Paria et al. [2019].

### Step 2: Sparse Set of Points.

```

1: procedure GPC( $\bar{F}_q, q$ )
2:    $C = \emptyset$ 
3:   while  $\bar{F}_q(C) < q$  do
4:     foreach  $c \in D \setminus C$  do  $\delta_c = \bar{F}_q(C \cup \{c\}) - \bar{F}_q(C)$ 
5:      $C = C \cup \{\arg \max_c \delta_c\}$ 
6:   end while
7: end procedure

```

---

## C C-MOSH: ENHANCING DECISION-MAKER CONFIDENCE WITH SENSITIVITY ANALYSIS

In practice, we found it difficult to display proper objective points to highlight, i.e. within the soft-hard bounds  $\alpha$  and exceeding the current best  $f_l(x^*)$  for objective  $l$  – especially at earlier iterations. As a result, we solve for Equation 9 multiple times (10) and filter out the obtained objective points which violate the SHF or do not exceed the current best point.

## D SIMULATION RESULTS

### D.1 SIMULATION SETUP

All simulations used a ground-truth ideal point  $\mathbf{y}^* = \max_{x \in X} s_{\lambda^*}(u_{\alpha^*}(x))$ , derived from a randomly sampled ground-truth preference vector  $\lambda^* \in \Delta(L)$ , where  $\Delta(L)$  is the  $L$ -dimensional probability simplex, and randomly sampled set of soft and hard bounds  $\alpha^*$ . To obtain  $\lambda^*$  and  $\alpha^*$ , we go through the following procedure: (1) sample a set of hard bounds  $\alpha_{\ell,H}^*$ , for each dimension  $\ell$ , in  $[0, 0.95]$ , (2) sample a set of soft bounds  $\alpha_{\ell,S}^*$ , for each dimension  $\ell$ , in  $[\alpha_{\ell,H}^* + 0.5, 1.0]$ , (3) sample  $\lambda^*$  using the heuristic  $\lambda^* = \mathbf{u} / \|\mathbf{u}\|_1$ , where  $\mathbf{u}_\ell \sim N(\alpha_{\ell,S}^*, |\alpha_{\ell,H}^* - \alpha_{\ell,S}^*|/3)$ . This ensures that the simulated  $\lambda^*$  and  $\mathbf{y}^*$  correspond to the high-utility regions of the sampled SHF. Throughout our experiments, for the scalarization function we used the augmented Chebyshev scalarization function  $s_\lambda(y) = -\max_{\ell \in [L]} \{\lambda_\ell |y_\ell - z_\ell^*|\} - \gamma \sum_{\ell=1}^L |y_\ell - z_\ell^*|$  where  $z_\ell^*$  is the ideal point for objective  $\ell$ .

**Baseline Feedback Simulations.** To simulate pairwise, ranking, and partial ranking feedback, we utilized  $\lambda^*$  and  $\alpha^*$  with added Gaussian noise  $\epsilon \sim N(0, \sigma^2)$ , where  $\sigma = 0.01$ , when determining preferences between alternatives. Specifically, for pairwise comparisons, the hidden noisy utilities were computed as  $v_i = s_{\lambda^*}(u_{\alpha^*}(x_i)) + \epsilon_i$ , where  $u_{\alpha^*}(x_i)$  represents the soft-hard utility values for alternative  $i$ . The alternative with the higher noisy utility value was selected as preferred. Ranking and partial ranking-based feedback follow similarly. Sample queries were selected using an information gain sampling strategy. Since the active sampling component of Active-MoSH implicitly selects Pareto-optimal points, we also only select samples from the Pareto frontier for the baselines (to ensure fairer comparisons).

**Active-MoSH Feedback Simulation.** For Active-MoSH feedback, we simulated DM interaction in response to the displayed query points  $Y_m$  at iteration  $m$ . The DM is assumed to respond based on a *reference point* within  $Y_m$ . For instance, a DM might observe points at the extremes of  $Y_m$  to determine if a bound modification is warranted. When simulating Active-C-MoSH (i.e., Active-MoSH with C-MoSH), we model the DM’s enhanced confidence from C-MoSH’s sensitivity analysis by allowing for larger magnitudes in their bound adjustments. *Note: We assert that it is difficult to accurately simulate the effects of C-MoSH. This is our attempt at doing so. However, even without such simulations of C-MoSH, we observe enhanced performance. We believe C-MoSH is better understood through our user study.*

We assume a default magnitude of adjustment for the soft and hard bounds, 0.3 and 0.45, respectively. *Although we chose those as the default values, we did find our results to be consistent across others of similar magnitude.* For the final value to adjust the bound with, we weighted it by a value sampled from a Gaussian distribution centered at the difference between the reference point and  $\mathbf{y}^*$ , with a fixed variance ( $\sigma = 0.01$ ).

- When  $\mathbf{y}^*$  falls outside some of the current bounds  $\alpha_{H,m}$ , the simulated DM selects the bound associated with the objective  $\ell$  that violates the value of  $\mathbf{y}^*$  the most and adjusts those bounds to bring  $\mathbf{y}^*$  closer to being within bounds. The adjustment magnitude is determined by the reference point, which we assume to be the point in  $Y_m$  closest to  $\mathbf{y}^*$ .
- When  $\mathbf{y}^*$  is within all bounds  $\alpha_{H,m}$ , the DM selects the bound furthest from  $\mathbf{y}^*$  and, we assume, attempts to narrow the feasible space around  $\mathbf{y}^*$ . For this case, we assume the reference point to be the point in  $Y_m$  furthest from  $\mathbf{y}^*$ .

- We assume the DM initially leverages the hard bounds, to ensure that  $y^*$  is within the desired region. For some objective  $\ell$  that is selected, if the soft and hard bounds are too close in proximity, i.e.  $\alpha_{\ell,S,m} - \alpha_{\ell,H,m} < \beta$ , we assume the DM then leverages the soft bounds to *fine tune* the desired region points.

Gaussian noise is added to the observations of points in  $Y_m$  at each iteration  $m$ . The magnitude of bound adjustments is sampled from a Gaussian distribution centered at the deviation between  $y^*$  and the reference objective point, with added noise. When  $y^*$  is outside the bounds and our proposed method C-MoSH promotes an improved point, we assume the DM has enhanced confidence and increases the magnitude with which they modify the soft or hard bound for that iteration. We assume this increase to be 5%.

Each simulator maintains consistent noise parameters and sampling procedures to ensure fair comparisons across feedback types. This simulation framework allows us to systematically evaluate how different types of preference feedback and varying levels of DM noise affect the convergence and effectiveness of our proposed approach. The setup enables controlled experiments while capturing realistic aspects of human decision-making behavior such as inconsistency and imprecision in feedback.

In terms of compute, all simulations were run on an internal cluster with CPUs. The flagship run for each simulation experiment typically finished within a day.

## D.2 PERFORMANCE EVALUATION

To present a different perspective on the simulated experiments, we attempt to account for cognitive complexity of the feedback mechanisms using *interaction units*. In general, the literature has no general consensus on feedback complexity since it is contextual to the task and visualization [Nesbitt et al., 2024]. We propose a notion of feedback complexity specific to our task and then summarize supporting literature for our mapping and provide analysis from our user study.

Pairwise feedback is assigned 1 unit as it is generally treated as the baseline unit for low-complexity feedback [Wilde et al., 2022, Wirth et al., 2017]. Active-MoSH is 2 units, as it requires the DM to (1) determine the dimension and bound to modify, and (2) specify the modification’s magnitude and direction. Slider and scale (not soft-hard bounds) feedback was studied in [Wilde et al., 2022], which decomposes it into two stages (direction identification, then adjustment), aligning with our identification and quantification mapping of 2 units. Ranking-based feedback is assigned  $k$  units, where  $k$  is the number of points displayed at iteration  $m^1$ . For ranking, Jamieson and Nowak [2011] notes a baseline  $k \log k$  upper bound; the structured nature of our displayed rankings (two objectives, sorted) supports a tighter lower bound of  $k$  [Jamieson and Nowak, 2011]. This method provides a standardized basis for comparing preference elicitation efficiency (Figure 6).

We provide further motivation for the specific unit assignments below:

- **Pairwise Feedback (1 unit):** Selecting one preferred item from a pair is considered a relatively simple cognitive task. It involves a single comparative judgment and is thus assigned as our baseline of 1 interaction unit.
- **Active-MoSH Feedback (2 units):** Providing feedback in Active-MoSH involves a two-step cognitive process by the DM:
  1. *Identification:* The DM must first identify which of the  $L$  objective dimensions they wish to adjust, and then decide whether to modify the soft or hard bound for that dimension. This requires understanding the current state of the tradeoff space.
  2. *Quantification:* The DM must then specify the magnitude and direction of the change for the selected bound (e.g., how much to increase the soft bound or decrease the hard bound). This involves a quantitative judgment about their desired preference shift.

Given these distinct decision-making steps, we assign 2 interaction units to a single Active-MoSH feedback action.

- **Ranking-based Feedback ( $k$  units):** When a DM is asked to rank  $k$  presented items, the cognitive effort generally scales with  $k$ . While the exact cognitive complexity can be debated (e.g., it might not be strictly linear, and the difficulty of ranking can depend on the similarity of items), assigning  $k$  units reflects that the DM performs  $k - 1$  implicit pairwise comparisons to establish a full order, or at least makes  $k$  evaluations to place each item. As noted in the footnote,  $k \log k$  is another plausible model reflecting sorting complexity. However, we opt for  $k$  units due to (1) the

---

<sup>1</sup>  $k \log k$  is also a valid interpretation for ranking-based interaction units; we use  $k$  because the cognitive load of ranking is not always uniform across points and for simplicity.

cognitive load of ranking not always being uniform across all  $k$  items (e.g., identifying the best and worst might be easier than ordering middle items), and (2) for simplicity in our comparative framework. This assignment acknowledges that processing and ordering  $k$  items is substantially more demanding than a single pairwise choice.

For mechanisms like ranking that are assigned multiple interaction units ( $k > 1$ ) to represent a single complete feedback instance (e.g., one full ranking of  $k$  items), we assume that the performance metric (e.g., SHF Utility Ratio as per Equation 2) remains constant across these constituent interaction units. The metric is only updated once the DM has completed the entire feedback action (e.g., submitted the full ranking). This accounting method, with results presented in Figure 6, provides a standardized basis for comparing the efficiency of different preference elicitation approaches, considering both the quality of learned preferences and the human effort involved. Even within such a setting, Figure 6 shows that Active-MoSH and Active-C-MoSH consistently outperform other feedback mechanisms in terms of the SHF Utility Ratio.

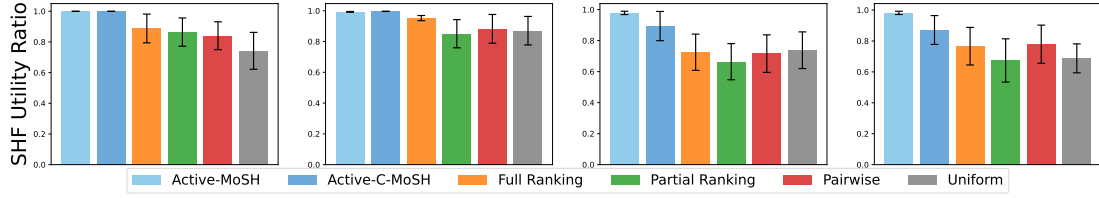


Figure 6: Evaluation results with 10 interaction units for the Branin-Currin synthetic function, Four Bar Truss engineering design, Brachytherapy Treatment Planning, and DTLZ2 (four objectives) applications (from left to right). The cognitive complexity of the feedback mechanisms were taken into consideration, using interaction units. The mean and standard error results were all computed over 10 independent trials.

**Empirical Validation of Interaction Units.** To assess whether our theoretical interaction unit assignments reflect actual decision-maker effort, we measured the per-action feedback times recorded during our user study (Section 7). Table 1 reports the mean and standard deviation of the time taken to complete a single feedback action for each mechanism.

Table 1: Mean feedback time per action (in seconds) across feedback mechanisms, measured during the user study. Values are reported as mean  $\pm$  standard deviation.

| Method          | Time (s)        |
|-----------------|-----------------|
| Pairwise        | $2.8 \pm 1.4$   |
| Soft-Hard       | $9.7 \pm 4.7$   |
| Partial Ranking | $16.7 \pm 14.0$ |
| Full Ranking    | $11.5 \pm 12.3$ |

Treating pairwise feedback as the baseline of 1 unit, the empirical time ratios suggest approximate effective costs of 3 units for soft-hard bounds, 4 units for full ranking, and 5 units for partial ranking. The soft-hard mapping is broadly consistent with our theoretical assignment of 2 units, with the modest increase plausibly attributable to the cognitive overhead of specifying exact numerical values via slider bars (as also discussed for H3 in Section 7). The largest deviation from our theoretical mapping occurs for partial ranking, likely due to the added complexity of first selecting which subset of points to rank before ordering them, an additional decision step not present in full ranking.

**Re-running Simulations with Empirical Units.** To verify that our conclusions are robust to this empirically grounded accounting, we re-ran our simulation experiments using the empirically derived interaction units in place of the theoretical assignments. Table 2 reports the resulting SHF Utility Ratio values for the Four Bar Truss (FBT) and Brachytherapy (Brachy) applications.

Under this empirically grounded unit mapping, Active-C-MoSH continues to achieve the highest SHF Utility Ratio across both applications, consistent with our original results in Figure 6. This indicates that the observed performance advantages of Active-MoSH and Active-C-MoSH are not an artifact of our specific theoretical interaction unit assignments, and persist when feedback costs are calibrated against measured decision-maker effort.

Table 2: SHF Utility Ratio under empirically derived interaction units, for the Four Bar Truss (FBT) and Brachytherapy applications. Values are reported as mean  $\pm$  standard error over 10 independent trials.

| Method          | FBT                             | Brachytherapy                   |
|-----------------|---------------------------------|---------------------------------|
| Active-MoSH     | .97 $\pm$ .03                   | .97 $\pm$ .03                   |
| Active-C-MoSH   | <b>.99 <math>\pm</math> .01</b> | <b>.99 <math>\pm</math> .01</b> |
| Partial Ranking | .96 $\pm$ .01                   | .79 $\pm$ .17                   |
| Full Ranking    | .92 $\pm$ .06                   | .80 $\pm$ .18                   |
| Pairwise        | .97 $\pm$ .03                   | .77 $\pm$ .10                   |

### D.3 BRANIN-CURRIN, FOUR BAR TRUSS, AND DTLZ2

Branin-Currin and DTLZ2 are standard benchmark problems provided in the BoTorch library [Balandat et al., 2020]. Four Bar Truss represents a realistic multi-objective optimization problem within the domain of engineering design. Its objectives involve the simultaneous minimization of structural volume and joint displacement, while the decision variables parameterize the lengths of the individual truss members. All the above benchmark problems are open-source.

### D.4 CERVICAL CANCER BRACHYTHERAPY TREATMENT PLANNING

For the cervical cancer brachytherapy treatment planning application, we leveraged a real and anonymized clinical case which had been performed in the clinic previously. As a result, we had access to the patient CT data (in the form of DICOM files) and the final treatment plan which had been administered to the patient. For computational consistency, we reformulate the minimization objectives for the OARs as maximization problems through appropriate transformations. The decision variables serve as inputs to an epsilon-constraint optimization program [Deufel et al., 2020], which enables systematic exploration of the treatment planning trade-off space. The three parameters to that linear program, which implicitly control the weights of the different objectives (radiation dosage to the cancer tumor and the bladder), were employed as the decision variables. The Pareto frontier for this problem is non-convex. Due to patient privacy issues, we do not plan to publicly release this clinical case.

### D.5 ADDITIONAL EXPERIMENTS AND ABLATION STUDIES

To provide a different perspective of the results among the different methods, we use the *area under the SHF Utility Ratio - Feedback curve* as the metric for our ablations below. The x-axis is the number of feedback units and the y-axis is the change in the SHF Utility Ratio, which produces a curve. By taking the area under that curve, we can compare how quickly each variation is able to obtain a high SHF Utility Ratio.

**Area Under the SHF Utility Ratio-Feedback Curve Results for Branin-Currin and Four Bar Truss.** Figures 8 and 9 illustrates our results for the area under the SHF Utility Ratio-Feedback Curve with both the Branin-Currin and Four Bar Truss problem. This provides a different perspective on our results as it illustrates the convergence rates with the different methods, further supporting the efficiency of Active-C-MoSH and Active-MoSH.

**Preference Modeling Ablation.** Figure 7 illustrates our results when ablating both Active-MoSH and Active-C-MoSH with and without the Plackett-Luce likelihood model (which we refer to as *Active-MoSH (Active-C-MoSH) without Plackett-Luce*). For Active-MoSH (Active-C-MoSH) without Plackett-Luce, we instead use a uniform distribution for  $p(\lambda)$ . From the results, we notice that solely using the uniform distribution greatly degrades the average performance, while often also increasing the variance of the results. This supports the importance of our preferences modeling component.

**Active Query Sampling Method Ablation.** Figure 7 illustrates our results when ablating both Active-MoSH and Active-C-MoSH when just using random sampling for the query sampling method, rather than our proposed two-step active querying method. From the results, we notice that simply using random sampling to obtain the queries hurts the converge of Active-MoSH and Active-C-MoSH significantly, while also often leading to higher variance. This supports the importance of our active sampling approach.

**Observing All With Same Active Querying Method.** We also include additional experiments comparing Active-MoSH and Active-C-MoSH against the baseline feedback mechanisms, all using our proposed active query sampling method,

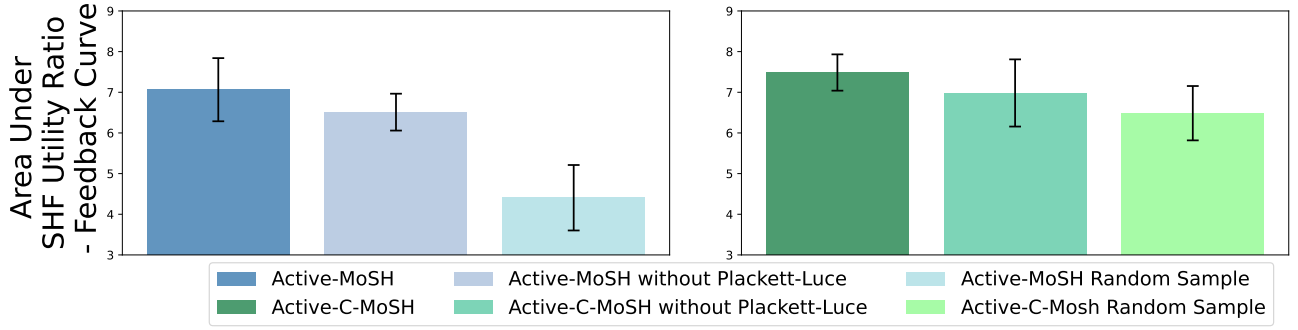


Figure 7: Ablation Results. **(Left)** Ablation Results for Active-MoSH. **(Right)** Ablation Results for Active-C-MoSH. *Without Plackett-Luce* use a uniform distribution for  $p(\lambda)$ . *Random Sample* obtains subsequent queries with random sampling, as opposed to our active query sampling approach. 10 independent trials.

as opposed to the information gain sampling method. In doing so, we seek to understand the impact of the active query sampling method on the SHF Utility Ratio values. We show the results in Figure 10. According to the results, the baseline feedback mechanisms observe higher SHF Utility Ratio values by the end of 10 feedback units (as compared to the results shown in Figure 6). This supports the improvement which our active sampling method provides for our problem setting. It also further highlights the superiority of Active-C-MoSH over the other baseline feedback mechanisms, while having a lower variance. However, we would like to emphasize that this result is obtained from our simulation setup, which may have limitations related to the simulated feedback. For instance, while we use a set of rules for Active-MoSH, we would expect for a decision-maker to be more flexible in their feedback adjustments, potentially obtaining greater utility over the other feedback mechanisms. This may also be seen in our user study results (Section 7.2).

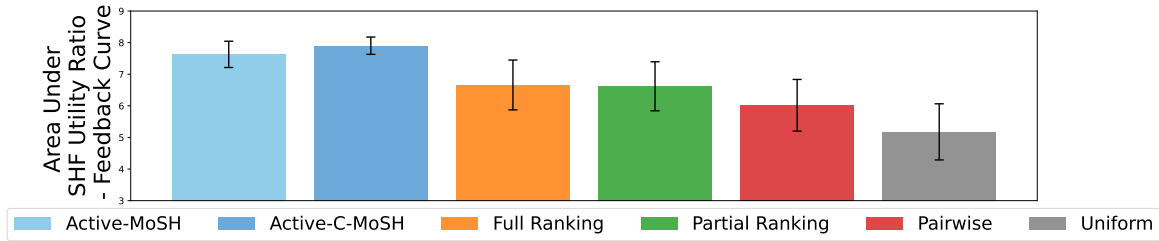


Figure 8: Area under the SHF Utility Ratio-Feedback Curve results over 10 interaction units for the Branin-Currin synthetic function. The mean and standard error results were all computed over 10 independent trials.

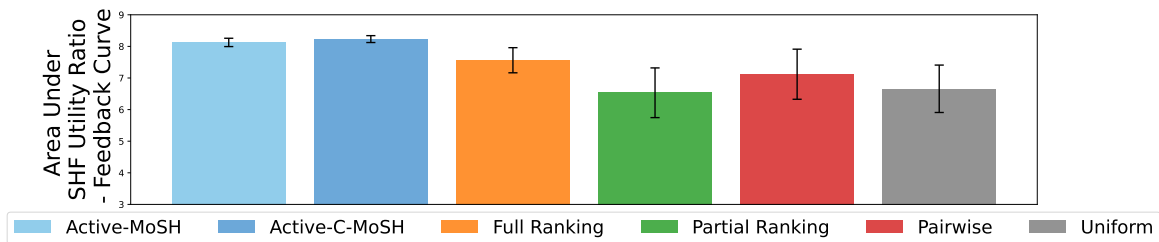


Figure 9: Area under the SHF Utility Ratio-Feedback Curve results over 10 interaction units for the Four Bar Truss function. The mean and standard error results were all computed over 10 independent trials.

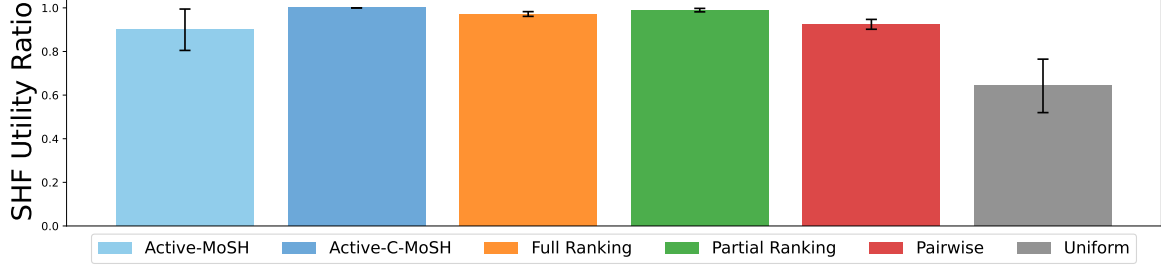


Figure 10: Evaluation results at 10 interaction units for the Branin-Currien synthetic function, where the baseline feedback mechanisms (besides *uniform*) are all using our proposed active query sampling method. The mean and standard error results were all computed over 10 independent trials.

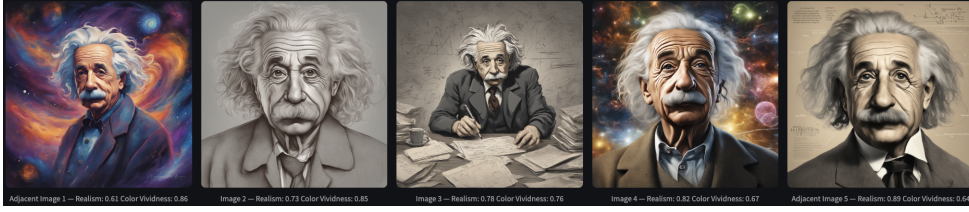


Figure 11: Sample query images from the user study for the task related to Albert Einstein.

## E USER STUDY

### E.1 USER STUDY SETUP

As mentioned, we designed an educational magazine cover photo selection task using AI-generated images. Interested participants were first given detailed instructions and were required to pass a five-question screening quiz ( $\geq 3$  correct answers) to qualify. Successful participants were then given a magazine topic description and instructed to use various feedback mechanisms to select the optimal cover photo. Each task instance’s description implicitly defined a tradeoff between two objectives: realism and color vividness, where higher values for both were desirable, aligning with our multi-objective setting. We created six task instances in total, combining two topics (Albert Einstein, Barack Obama) with three distinct tradeoff points. Participants were informed about the specific feedback mechanism to use for each task.

**Generating the Pareto Frontier of AI Images for User Study.** For the user study (Section 7), we required a set of AI-generated images that systematically trade off between two competing objectives: realism and color vividness. To achieve this, we controlled the output of a text-to-image generative model by manipulating its effective input prompt representation using proxy tuning [Liu et al., 2024, Mitchell et al., 2023, Shi et al., 2024].

Proxy tuning steers a large pre-trained language model ( $M$ ) by incorporating signals from smaller, specialized models: an expert model ( $M^+$ ) tuned for a desired attribute, and an anti-expert model ( $M^-$ ), typically the untuned version of  $M^+$ . Following the notation of Liu et al. [2024], the output distribution of the proxy-tuned model ( $\tilde{M}$ ) at decoding step  $t$ , conditioned on the preceding token sequence  $x_{<t}$ , is given by:

$$p_{\tilde{M}}(X_t|x_{<t}) = \text{softmax} \left[ s_M(X_t|x_{<t}) + \sum_{i=1}^2 \theta_i \left( s_{M_i^+}(X_t|x_{<t}) - s_{M_i^-}(X_t|x_{<t}) \right) \right]$$

where  $s_M$ ,  $s_{M_i^+}$ , and  $s_{M_i^-}$  are the logit scores from the respective models. The parameters  $\theta_i$  are controllable weights applied to the logit difference for expert  $M_i^+$ . In our two-objective setting,  $M_1^+$  was tuned for realism and  $M_2^+$  for color

vividness. By varying  $\theta_1$  and  $\theta_2$  at decoding time, we effectively controlled the text prompt representation fed into the text-to-image model, thereby generating images with varying tradeoffs between realism and color vividness.

For our experiments, we utilized models from the T LU-2 suite [Iverson et al., 2023]: T LU-2-13B as the base model  $M$ , and T LU-2-DPO-7B for both the expert ( $M^+$ ) and anti-expert ( $M^-$ ) models. The expert models,  $M_1^+$  (realism) and  $M_2^+$  (color vividness), were independently fine-tuned using Direct Preference Optimization (DPO) [Rafailov et al., 2024] on custom preference datasets. To construct these datasets, we first obtained names of historical figures from the Pantheon 1.0 dataset [Yu et al., 2016]. For each figure and each objective (realism, color vividness), we used GPT-4o-mini to expand a seed prompt (the figure’s name) into multiple textual variations representing positive, neutral, and negative degrees of the target attribute (see Figures 18–24 for exact prompts). We generated two positive variations for each. Binary preference pairs for DPO were then formed as (positive variation 1, positive variation 2), (positive variation 1, neutral), (positive variation 1, negative), and (neutral, negative). For pairs not involving two positive variations, the more positive prompt was designated as chosen. For (positive 1, positive 2) pairs, we generated images from both prompts using Stable Diffusion XL [Podell et al., 2023] and then used GPT-4o-mini with a detailed grading rubric (Figures 25–28) to evaluate the resulting images on the relevant attribute (realism or color vividness); the prompt corresponding to the higher-scoring image was marked as preferred.

With the two fine-tuned expert models ( $M_1^+$ ,  $M_2^+$ ), we systematically varied  $\theta_1$  and  $\theta_2$  over the range  $[-0.5, 1.0]$  with a step size  $\delta = 0.05$ . For each  $(\theta_1, \theta_2)$  pair, we obtained a proxy-tuned prompt which was then used as input to Stable Diffusion XL to generate an image. The realism and color vividness scores (0-100 scale) for each generated image were subsequently evaluated by GPT-4o-mini using the aforementioned detailed grading rubric. This process yielded a diverse set of images for Albert Einstein and Barack Obama, populating different regions of the realism-vividness tradeoff space.

The fine-tuning of expert models and image generation were conducted on an internal cluster using three NVIDIA A100 GPUs (82GB VRAM). Flash Attention 2 [Dao, 2023] was employed during DPO fine-tuning, enabling completion of each expert model’s flagship run within approximately 12 hours.

**User Study Components and Screenshots.** Here we outline the set of components we used for the user study:

1. **Instructions.** We provided each MTurk participant with a detailed set of instructions covering the tradeoff between the objectives, the different feedback mechanisms and how to use them, a walkthrough of the entire workflow, and other details. We provide the entire set of instructions in Figures 29-34.
2. **Screening quiz.** We gave interested MTurk participants a five-question screening quiz on the concepts which were explained in the set of instructions, mainly for ensuring understanding of the feedback types. We provide our screening quiz questions in Figures 36 and 37. The screening quiz interface was developed using Streamlit str [2018] and the data was stored using Supabase [sup, 2020].
3. **User study tool.** Each Human Intelligence Task (HIT) consisted of four different task instances. For each task instance, the MTurk participants were provided with the task description, along with an interface to provide their feedback on a set of query images at each iteration. The MTurk participants were instructed to provide feedback on the set of query images until they felt they had found the optimal image for the educational magazine’s cover photo in the task description. At the end of the study, each MTurk participant was asked to respond to three different sets of questions regarding the mental effort, expressiveness, and trustworthiness of each of the feedback mechanisms. We provide screenshots of the relevant components in Figures 38-45. Figures 46-51 depict all of the exact task descriptions we used. The survey questions at the end are shown in Figure 35. The user study tool interface was developed using Streamlit str [2018] and the data was stored using Supabase [sup, 2020].

## E.2 USER STUDY PROCEDURE

We randomized the order of these task instances for each MTurk worker to prevent any bias. Each MTurk worker was permitted to complete the user study twice. For their second time with the user study, the MTurk worker was presented with a random ordering of the last two unseen task instances and two of the earlier task instances, for a total of the four feedback mechanisms.

Each MTurk worker was compensated with \$4.00 for completing the screening quiz and was awarded an additional \$1.50 for passing it. For those who passed the screening quiz, they were compensated with \$19.00 for completing each HIT of the user study. We attempted to ensure that the time in between each step of the user study (e.g. screening quiz and user study HIT) was as short as possible, to retain familiarity; most MTurk participants completed both within 24 hours.

### E.3 USER STUDY RESULTS

Since each of the MTurk participants were given the opportunity to participate in the user study for a total of two times, we used the results from only their second trial for Figures 4 and 14, to account for the novelty effect [Elston, 2021]. For hypothesis testing, we leveraged the one-sided Welch’s T-Test for **H1** and **H2** and the Mann-Whitney U Test for the Likert-scale responses part of **H2**, **H3**, and **H4**. We used a statistical significance level of 0.05 and leveraged the Holm-Bonferroni method to control the Family-Wise Error Rate [Holm, 1979].

Although our statistical hypothesis testing on the Iteration Stop Efficiency (ISE) metric results from 7.2 were not statistically significant, we did notice there was a trend towards supporting our **H2** regarding Active-C-MoSH building more confidence. To further evaluate **H2**, we also leveraged statistical hypothesis testing on our third end-of-study question asking MTurk participants about how they view each of the feedback types, in terms of trust. In this case, to simplify the question we equate trust with decision confidence. Higher scores indicate more decision confidence. Active-C-MoSH (mean=4.50) was rated as ensuring more confidence compared to full ranking (mean=3.50, adj.  $p = 0.01$ ), partial ranking (mean=4.00, adj.  $p = 0.05$ ), and pairwise feedback (mean=3.83, adj.  $p = 0.04$ ). We believe part of this may be attributed to the samples from C-MoSH, which were often presented to the MTurk participant. Another possible reasoning could be the additional familiarity the MTurk participants gained with Active-C-MoSH, better understanding how to interpret it and, therefore, be confident in its results. The distribution of all Likert-scale responses are shown in Figures 12, 13, 14.

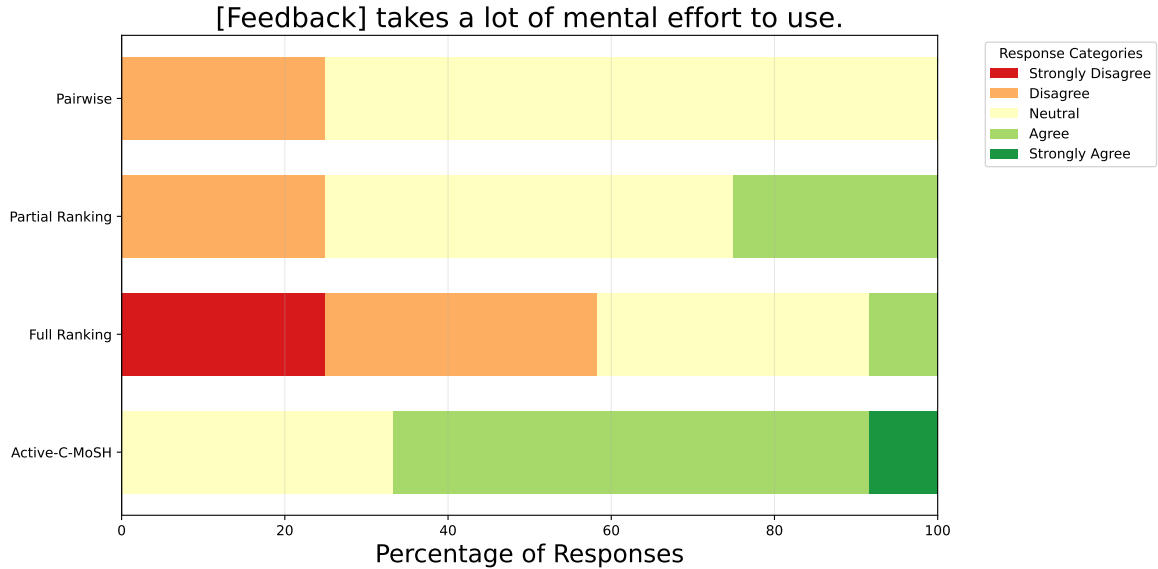


Figure 12: Distribution of Likert-scale survey responses for our first end-of-study question for the MTurk participants.

## F CASE STUDY: CERVICAL CANCER

To conduct this case study, we designed a simple interface for Active-MoSH and all the baseline feedback mechanisms. This was built using Streamlit [str, 2018]. The six objectives which were being traded off include the following:  $PTV_{V700}$ ,  $Bladder_{D2cc}$ ,  $Rectum_{D2cc}$ ,  $Bowel_{D2cc}$ ,  $Cervix_{D2cc}$ , and  $Mucosa_{D2cc}$ . Figures of the interface are shown in Figures 15, 16, 17. We followed Chen et al. [2026a] in modeling the optimization problem; the decision variables served as inputs to an epsilon-constraint optimization program [Deufel et al., 2020]. Due to the high-dimensionality of this domain and the need for guardrails, the decision variable space was discretized with a step size of 0.05 for  $PTV_{V700}$ ,  $Bladder_{D2cc}$ ,  $Rectum_{D2cc}$ ,  $Bowel_{D2cc}$  and a step size of 0.30 for the rest. The input range of each of the variables was determined by a simple variation of the Automated Planning Parameter Setup described in Chen et al. [2026a].

We gathered three real retrospective anonymized patient cases from a local hospital from which we were able to determine the clinical reference treatment plan metrics. The six metric dimensions for each of the treatment plans were in slightly different regions of the tradeoff frontier, since the optimal treatment plan is highly dependent on patient physical conditions [Deufel et al., 2020].

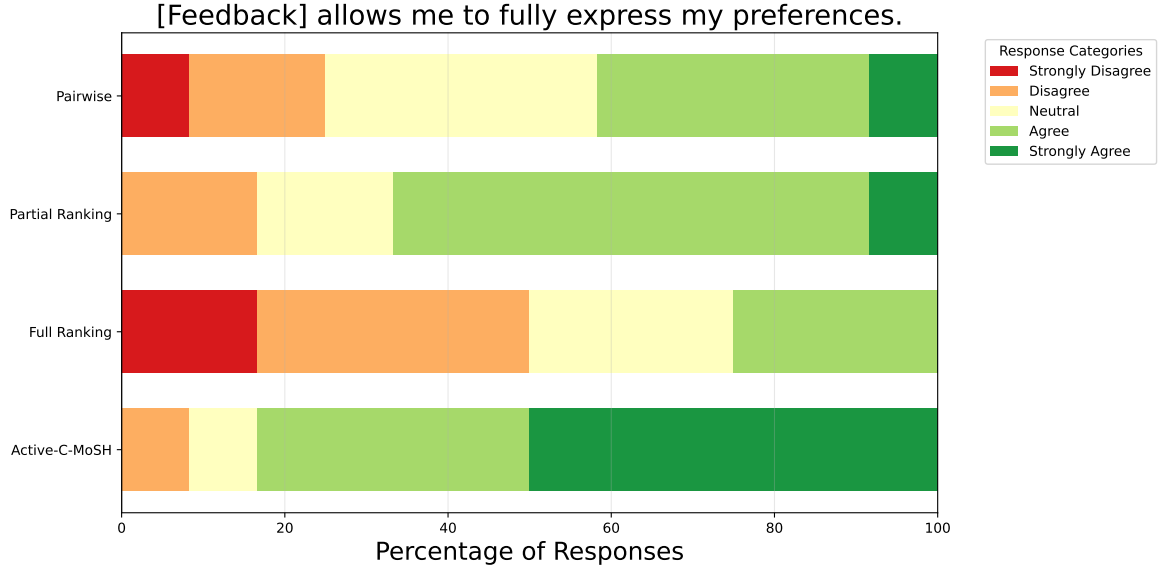


Figure 13: Distribution of Likert-scale survey responses for our second end-of-study question for the MTurk participants.

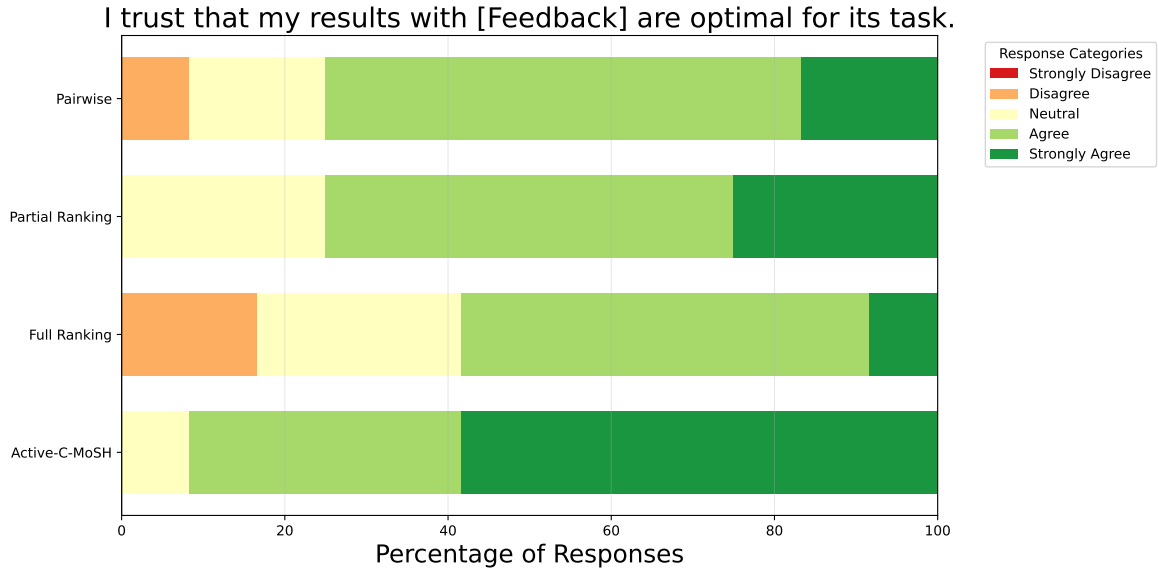


Figure 14: Distribution of Likert-scale survey responses for our third end-of-study question for the MTurk participants.

## G ADDITIONAL RELATED WORKS

**Interactive Feedback Mechanisms.** Another stream of diverse feedback modalities includes learning from demonstrations Abbeel and Ng [2004], González et al. [2018], Ziebart et al. [2008], Shaikh et al. [2024]. For multiple human inputs, Myers et al. [2022] proposed multimodal rewards.

**Pareto Frontier Population Mechanisms.** Many MOO works focus on populating the entire Pareto frontier [Campigotto et al., 2014, Ponweiser et al., 2008, Emmerich, 2008, Picheny, 2015, Hernández-Lobato et al., 2016, Zhang et al., 2009], sometimes via random scalarizations [Knowles, 2006, Paria et al., 2019] or by seeking sparse coverage, e.g., through level sets [Zuluaga et al., 2016, Malkomes et al., 2021]. We instead use soft and hard bounds to actively direct exploration to high-utility, preference-aligned areas, reducing computational and cognitive load.

**Building Confidence in AI Systems.** Broderick et al. [2023] developed a taxonomy of trust failure points in data analysis platforms, identifying perturbation analysis as a critical component. Within decision-making specifically, Johnson et al.

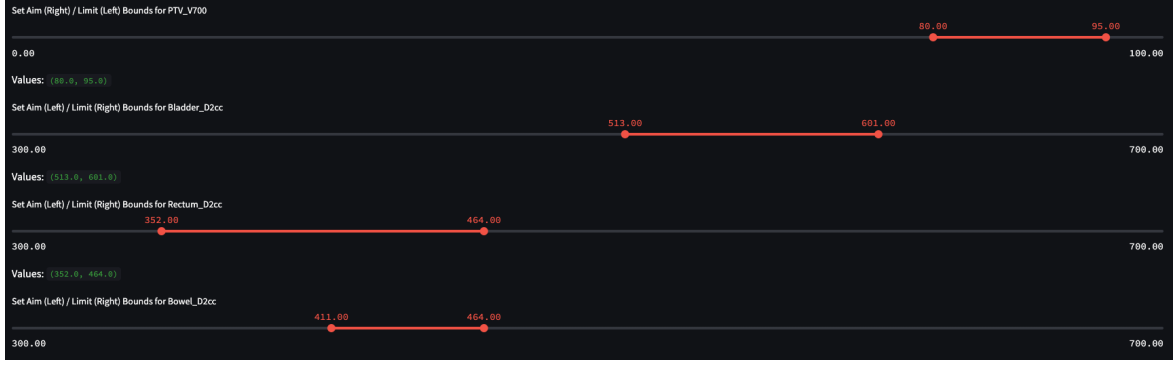


Figure 15: Soft-Hard Bounds Feedback Mechanism for Cervical Cancer Brachytherapy Treatment Planning Case Study.



Figure 16: Pairwise Feedback Mechanism for Cervical Cancer Brachytherapy Treatment Planning Case Study. Only default treatment plan metrics are shown. The Needle<sub>D2cc</sub> objective is not used for evaluation.

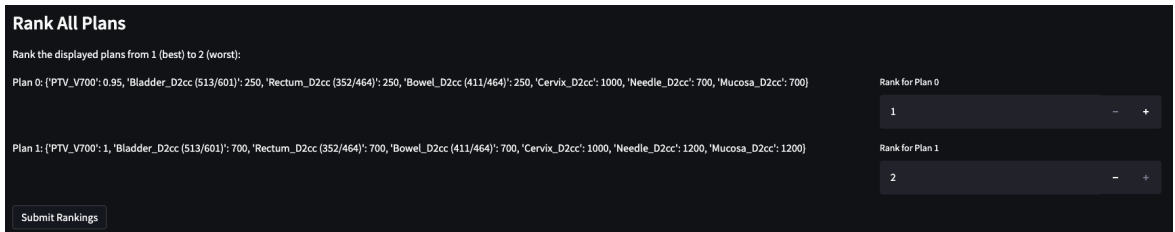


Figure 17: Ranking-Based Feedback Mechanism for Cervical Cancer Brachytherapy Treatment Planning Case Study. Only default treatment plan metrics are shown. The Needle<sub>D2cc</sub> objective is not used for evaluation.

[2024] studied the interplay between decision confidence, rewards, and priors. While Insua and French [1991], Insua [1990] proposed a general framework for sensitivity analysis in discrete multi-objective decision-making, their work did not connect it to specific interactive feedback mechanisms or practical implementations. On the other hand, our framework offers an actionable multi-objective decision-making system where sensitivity analysis (C-MoSH) is directly integrated to enhance DM confidence in the soft-hard bounds feedback process.

**Bayesian Decision Theory.** A related line of work frames exploration through Bayes-Adaptive Markov Decision Processes (BAMDPs) for non-myopic Bayes-optimal exploration, e.g., via meta-learning approaches such as variBAD [Zintgraf et al., 2021]. Our setting differs: we have no task distribution to meta-train over (a single DM and single MOO problem), our unknowns are the DM’s latent  $(\lambda^*, \alpha^*)$  rather than MDP dynamics, and DM queries are expensive and budget-limited. Our regime instead aligns with active preference learning and preferential Bayesian optimization [Beylkin et al., 2020, Ozaki et al., 2023], where myopic, posterior-aware acquisition is standard. We follow this approach and extend it to soft-hard

bounds with C-MoSH for confidence.

## H LIMITATIONS

While Active-MoSH demonstrates promising results, several limitations warrant discussion.

**Computational Cost.** The computational cost of maintaining and updating probabilistic models (e.g., Gaussian Processes for surrogate modeling of objectives, posteriors over  $\lambda$  and  $\alpha$ ) can become significant with a large number of objectives  $L$  or many iterations  $M$ . While active sampling aims to be efficient, the overhead of GP training and acquisition function optimization in the local component, and sensitivity analysis in C-MoSH, could be a bottleneck for very high-dimensional problems or extremely rapid interaction requirements.

**Cognitive Load.** Although Active-MoSH is designed to manage cognitive load, the interactive nature still requires DM engagement. The effectiveness of the framework relies on the DM’s ability to provide meaningful feedback on soft-hard bounds. DM fatigue, inconsistency, or difficulty in articulating complex multi-level preferences (as indicated by the user study results in Figure 4) could impact performance. Our assumption of a single bound modification per iteration simplifies feedback but might be restrictive for DMs wishing to express more complex preference changes simultaneously.

**Generalizability of Evaluations.** The goal of our simulation-based experiments was to mimic humans in decision-making scenarios as closely as possible. While we followed the standard in current literature, we acknowledge that humans often have complex decision-making patterns which are difficult to perfectly replicate, especially when it comes to novel feedback mechanisms such as soft-hard bounds. For instance, regarding our simulation setup, it is difficult to simulate the effects of C-MoSH. However, even without such simulations of C-MoSH, we observe enhanced performance. Finally, our user study (Section 7), while insightful, was conducted on a specific AI-generated image selection task with a particular DM pool (MTurk workers). Our quantitative proxy for confidence, Iteration Stop Efficiency (ISE), may possibly be confounded with the MTurk workers’ fatigue or any time constraints. Further studies across diverse domains and with domain experts may help to fully assess the generalizability and practical utility of Active-MoSH and C-MoSH.

## I BROADER IMPACTS

The Active-MoSH framework has the potential to positively impact high-stakes decision-making in various domains by enabling more efficient and trustworthy navigation of complex tradeoff spaces. In healthcare (e.g., brachytherapy treatment planning) or engineering design, where multiple competing objectives and expensive evaluations are common, our method could lead to better-informed decisions that more closely align with true stakeholder preferences, potentially improving outcomes and resource allocation. The explicit mechanism for building DM confidence (C-MoSH, Section 4) is particularly valuable, as it can increase user adoption and confidence in AI-assisted decision support systems.

However, potential negative societal impacts or misuse should also be considered. If the DM’s underlying goals or specified bounds are unethical or biased, Active-MoSH, like any optimization tool, could inadvertently facilitate undesirable outcomes. The accessibility of such advanced decision-support tools might also be a concern, potentially creating disparities if not made widely available or if requiring significant expertise to use effectively.

#### GPT-4o-mini Input Prompt For Realism: Positive 1

“You are an expert in photorealistic prompt engineering for AI image generation. Your task is to expand the following subject into a detailed prompt that would guide an AI to generate an extremely photorealistic image.

Focus on:

- Photographic details (lighting, shadows, depth of field)
- Physical accuracy (correct anatomy, proportions, physics)
- Technical specifications (8K resolution, professional photography terms)
- Realistic materials, textures, and surface details
- Natural environment and context

Only output the expanded output prompt directly. Do NOT output anything else.

{input prompt}”

Figure 18: Input prompt passed into GPT-4o-mini to generate the expanded prompt for an image with high realism, used for the realism objective preference dataset.

#### GPT-4o-mini Input Prompt For Realism: Positive 2

“You are an expert in photorealistic prompt engineering for AI image generation. Your task is to expand the following subject into a detailed prompt that would guide an AI to generate a photorealistic image.

Only output the expanded output prompt directly. Do NOT output anything else.

{input prompt}”

Figure 19: Input prompt passed into GPT-4o-mini to generate the expanded prompt for an image with high realism, used for the realism objective preference dataset.

#### GPT-4o-mini Input Prompt: Neutral

“You are an expert in prompt engineering for AI image generation. Your task is to expand the following subject into a detailed prompt that would guide an AI to generate an image.

Only output the expanded output prompt directly. Do NOT output anything else.

{input prompt} ”

Figure 20: Input prompt passed into GPT-4o-mini to generate the expanded prompt for an image with neutral qualities, used for both the realism and color vividness preference datasets.

#### GPT-4o-mini Input Prompt For Realism: Negative

“You are a creative prompt engineer for AI image generation. Your task is to expand the following subject into a prompt that would guide an AI to generate a highly artistically stylized image.

Focus on:

- Artistic styles (cartoon, illustration, abstract, etc.)
- Creative liberties with appearance and physics
- Stylized exaggerations or simplifications
- Fantasy or surreal elements
- Bold colors or unrealistic visual treatments

Only output the expanded output prompt directly. Do NOT output anything else.

{input prompt} ”

Figure 21: Input prompt passed into GPT-4o-mini to generate the expanded prompt for an image with low realism, used for the realism objective preference dataset.

#### GPT-4o-mini Input Prompt For Color Vividness: Positive 1

“You are an expert in vibrant color styling for AI image generation. Your task is to expand the following subject into a detailed prompt that would guide an AI to generate an image with extraordinarily bold, vivid, and diverse colors.

Focus on:

- Bold color choices (saturated hues, neon elements, vibrant tones)
- Strategic color relationships (complementary pairs, triadic schemes, high-contrast combinations)
- Technical color specifications (maximum saturation, color vibrance, enhanced chroma)
- Unusual or unexpected color applications (color blocking, gradient shifts, non-traditional color choices)
- Lighting techniques that amplify color impact (rim lighting, color gels, dramatic color contrast)
- Color diversity across the image (wide color gamut, maximum color range, varied color temperature)
- Visual color impact (eye-catching color focal points, color that creates visual tension or harmony)

Only output the expanded output prompt directly. Do NOT output anything else.

{input prompt} ”

Figure 22: Input prompt passed into GPT-4o-mini to generate the expanded prompt for an image with high color vividness, used for the color vividness objective preference dataset.

#### GPT-4o-mini Input Prompt For Color Vividness: Positive 2

“You are an expert in vibrant color styling for AI image generation. Your task is to expand the following subject into a detailed prompt that would guide an AI to generate an image with bold, vivid, and diverse colors.

Only output the expanded output prompt directly. Do NOT output anything else.

{input prompt} ”

Figure 23: Input prompt passed into GPT-4o-mini to generate the expanded prompt for an image with high color vividness, used for the color vividness objective preference dataset.

#### GPT-4o-mini Input Prompt For Color Vividness: Negative

“You are an expert in subdued color styling for AI image generation. Your task is to expand the following subject into a detailed prompt that would guide an AI to generate an image with exceptionally restrained, subtle, and minimal color presence.

Focus on:

- Restrained color choices (desaturated tones, muted hues, subtle color variations)
- Monochromatic or limited color schemes (single-color palettes, tonal variations, grayscale with minimal accents)
- Technical color specifications (reduced saturation, lowered chroma, subtle tonal shifts)
- Sophisticated color restraint (elegant neutrals, barely-there color hints, delicate color transitions)
- Lighting techniques that minimize color impact (diffused lighting, low contrast, subtle gradations)
- Color harmony through absence (minimalist color approach, understated color presence, thoughtful color reduction)
- Visual impact through subtlety (tonal relationships, texture over color, emphasis on light and shadow rather than hue)

Only output the expanded output prompt directly. Do NOT output anything else.  
{input prompt} ”

Figure 24: Input prompt passed into GPT-4o-mini to generate the expanded prompt for an image with low color vividness, used for the color vividness objective preference dataset.

#### GPT-4o-mini Input Prompt For Image Evaluation (Part 1 of 4)

“Evaluate this image using the following calibrated scale for realism and color palette boldness/vividness. Your scores will help establish a Pareto frontier of optimal tradeoffs between these competing objectives.

##### IMPORTANT SCORING GUIDELINES:

- Both scales use the FULL range from 0.0-100.0 with decimal precision (e.g., 76.3, 42.8)
- The two scores typically follow an approximate inverse relationship
- As one score increases, the other typically decreases
- Roughly speaking, realism score + color boldness score often falls between 75-110
- Rare exceptional images might score highly on both dimensions (but these are uncommon)
- Both extremely low scores on both dimensions indicates poor image quality, not a valid tradeoff
- Your scores should reflect relative positioning compared to other images in this domain

##### FUNDAMENTAL RELATIONSHIP:

- Realism focuses on accurate, photographic representation of reality with true-to-life colors
- Color boldness/vividness measures how extreme, vibrant, unusual, or attention-grabbing the color palette is
- For most images, increasing color boldness typically reduces realism, creating a tradeoff relationship
- The optimal frontier represents the best possible combinations of these competing values
- Images can be Pareto-optimal at different points (high realism/low color, balanced, high color/low realism)

##### REALISM SCALE (0.0-100.0):

- 95.0-100.0: Perfect photographic realism - completely indistinguishable from a professional photograph
- 85.0-94.9: High photorealism - extremely realistic with minimal stylistic elements
- 75.0-84.9: Strong realism - convincing real-world representation with occasional minor stylistic choices
- 65.0-74.9: Moderate realism - clearly recognizable with some noticeable stylistic elements

”

Figure 25: Input prompt passed into GPT-4o-mini to evaluate the AI-generated images in terms of both realism and color vividness and obtain two numerical scores (part 1 of 4).

#### GPT-4o-mini Input Prompt For Image Evaluation (Part 2 of 4)

“

- 55.0-64.9: Balanced realism - equal emphasis on realism and interpretation
- 45.0-54.9: Limited realism - recognizable subjects with significant stylization
- 35.0-44.9: Partial realism - subjects are recognizable but heavily stylized
- 25.0-34.9: Minimal realism - heavily stylized with basic real-world references
- 15.0-24.9: Trace realism - predominantly abstract with subtle real-world hints
- 0.0-14.9: No realism - completely abstract, no recognizable elements

#### COLOR BOLDNESS/VIVIDNESS SCALE (0.0-100.0):

- 95.0-100.0: Extreme color intensity - extraordinarily vibrant, possibly fluorescent or neon
- 85.0-94.9: Highly vivid colors - very bold, saturated palette that immediately draws attention
- 75.0-84.9: Bold and vivid - distinctly vibrant palette with enhanced saturation
- 65.0-74.9: Moderately enhanced colors - noticeably vibrant with deliberate enhancement
- 55.0-64.9: Balanced color approach - moderately enhanced colors without overwhelming
- 45.0-54.9: Slightly enhanced colors - subtle enhancement beyond natural appearance
- 35.0-44.9: Natural color plus - largely natural palette with occasional subtle enhancement
- 25.0-34.9: Natural color range - colors as they would appear in standard photography
- 15.0-24.9: Subdued natural colors - slightly desaturated or muted natural colors
- 0.0-14.9: Minimal color presence - extremely desaturated or limited color range

#### RELATIONSHIP PATTERNS (not strict rules):

- Professional documentary photography often scores: Realism 85-95, Color Boldness 20-40
- Technical/scientific imagery often scores: Realism 80-95, Color Boldness 10-30

”

Figure 26: Input prompt to evaluate the AI-generated images in terms of both realism and color vividness and obtain scores (part 2 of 4).

#### GPT-4o-mini Input Prompt For Image Evaluation (Part 3 of 4)

“

- Commercial product photography often scores: Realism 80-90, Color Boldness 40-60
- Fashion photography often scores: Realism 75-85, Color Boldness 50-70
- Modern digital art often scores: Realism 40-70, Color Boldness 60-90
- Pop art style often scores: Realism 30-50, Color Boldness 80-100
- Psychedelic art often scores: Realism 20-40, Color Boldness 85-100
- Impressionist style often scores: Realism 40-60, Color Boldness 60-80
- Realistic paintings often score: Realism 60-80, Color Boldness 40-60

**HOW TO EVALUATE COLOR BOLDNESS/VIVIDNESS:** Consider these factors when scoring color boldness/vividness:

1. Saturation levels - how pure and intense the colors appear
2. Color contrast - the degree of difference between adjacent colors
3. Relationship to natural color - how much colors deviate from their appearance in nature
4. Use of complementary colors - pairing colors from opposite sides of the color wheel
5. Color harmony vs. discord - whether colors create pleasing relationships or intentional tension
6. Overall color impact - how immediately the colors capture attention
7. Unusual or unexpected color choices - colors applied in unconventional ways
8. Color temperature extremes - very warm or cool color schemes, or dramatic contrasts between warm and cool

”

Figure 27: Input prompt passed into GPT-4o-mini to evaluate the AI-generated images in terms of both realism and color vividness and obtain two numerical scores (part 3 of 4).

#### GPT-4o-mini Input Prompt For Image Evaluation (Part 4 of 4)

“ **EVALUATION INSTRUCTIONS:**

1. First assess each dimension independently using their specific criteria
2. Then consider if the scores reflect an appropriate tradeoff relationship
3. Use decimal precision (e.g., 82.7 rather than just 83)
4. For images that seem to defy the usual tradeoff, verify if it truly represents exceptional quality
5. Consider how this image would position relative to others on a Pareto frontier
6. Ensure you're using the full scale range appropriately

**Output format:**

**\*\*Realism Score: [0.0-100.0]\*\***

**\*\*Color Boldness/Vividness Score: [0.0-100.0]\*\***

**\*\*Reasoning:\*\*** [Your detailed explanation, referencing specific visual elements and explaining how the image represents a specific tradeoff point between these objectives] ”

Figure 28: Input prompt passed into GPT-4o-mini to evaluate the AI-generated images in terms of both realism and color vividness and obtain two numerical scores (part 4 of 4).

## MTurk User Study Instructions (Part 1 of 6): Introduction & Overview

### Image Selection User Study: Instructions

#### Introduction

Thank you for your interest in participating in our image selection user study. This document provides detailed instructions for the study tasks. Please read these instructions carefully as you will be asked to complete a short quiz on this material before beginning the actual study.

#### Study Overview

In this study, you will take on the role of a cover photo designer for an educational magazine. You will evaluate and select AI-generated images based on specific magazine requirements, using different feedback mechanisms to guide the selection process. Your participation will help us understand which image selection approaches are most effective for creative professionals.

#### Your Task

As a cover photo designer, you will:

1. Read a specific magazine's requirements for their next cover image.
2. View several AI-generated images.
3. Use a type of feedback to guide the system toward better images.
4. Select a final image that best matches the magazine's requirements.
5. Repeat steps 1-4 using four different types of feedback.
6. Complete a survey about your experience.

Figure 29: MTurk User Study Instructions: Introduction, Study Overview, and Task Summary.

## MTurk User Study Instructions (Part 2 of 6): Image Attributes

### Key Image Attributes

The images in this study vary along two important dimensions:

- **Realism**

- Measures how photorealistic and accurate the image appears.
- Higher realism means the image looks more like a real photograph.
- Lower realism means the image appears more stylized or artistic.

- **Color Vividness**

- Measures how bold, vibrant, and colorful the image appears.
- Higher color vividness means more saturated, eye-catching colors.
- Lower color vividness means more subdued, natural, or muted colors.

### The Critical Tradeoff

Both high realism AND high color vividness are highly desired qualities for magazine covers. However, there is often a tradeoff between these objectives. Images that are extremely realistic typically have more muted colors, while images with vibrant colors often appear less realistic. Your challenge is to find the optimal balance based on the specific magazine's requirements. Please disregard factors in the image unrelated to the aforementioned two tradeoffs.

Figure 30: MTurk User Study Instructions: Description of Key Image Attributes and the Realism-Vividness Tradeoff.

### Feedback Mechanisms

You will try four different feedback approaches during the study:

#### 1. Full Ranking

- *What it is:* You will order all images from most preferred to least preferred, according to the magazine's requirements.
- *How to use it:* Enter the numbers corresponding to each image in your preferred order.
- *Example:* If you have five images and prefer them in the order 3, 1, 4, 2, 5, enter: 3, 1, 4, 2, 5
- *Important:* Make sure to format your input correctly with commas between numbers.

#### 2. Partial Ranking

- *What it is:* You will rank only the top three images based on your strongest preferences according to the magazine's requirements.
- *How to use it:* Enter only the numbers for the top three images you have clear preferences about.
- *Example:* If there is a total of five images and if you strongly prefer image 3, then image 1, then image 2, enter: 3, 1, 2

#### 3. Pairwise Preferences

- *What it is:* You will choose which image you prefer between two options at a time, according to the magazine's requirements.
- *How to use it:* Simply click on the image you prefer when presented with two options. Even if you do not like any of the images, please select the image you prefer more.

Figure 31: MTurk User Study Instructions: Description of Full Ranking, Partial Ranking, and Pairwise Preferences feedback mechanisms.

#### 4. Soft & Hard Bounds

- *What it is:* You'll set minimum acceptable levels (hard bounds) and preferred levels (soft bounds) for image attributes.
- *How to use it:* For each image quality attribute (realism and color vividness):
  - The left slider handle sets a **hard bound** (minimum acceptable level).
  - The right slider handle sets a **soft bound** (preferred level).
- *What it represents:*
  - Hard bounds: "I reject any image below this threshold."
  - Soft bounds: "I prefer images above this threshold, but will accept lower values."
- *Important Notes:*
  - The first iteration of this feedback type will allow you to initialize the region of the bounds, based on our guess from the task description. You can perform multiple actions for this iteration.
  - After initialization, the system only allows you to perform **one action** (one modification of one of the bounds) for each set of results. Please do not attempt to perform multiple actions for each iteration.
  - After performing one action, images labeled as "Adjacent Image" may sometimes appear in the next set of results. These images are used to help you understand what might happen if the hard bounds are adjusted to be lower for one or both objectives. For example, due to the competing tradeoffs, if you were to decrease the hard bound for "Realism", an image with a much higher value of "Color Vividness" may appear in the next set of results. An "Adjacent Image" with a much higher value of "Realism" lets you know that if you decrease the hard bound for "Color Vividness", you might observe images such as that.

Figure 32: MTurk User Study Instructions: Description of the Soft & Hard Bounds feedback mechanism.

#### Study Process

For each feedback mechanism, you will:

1. Read the specific magazine requirements carefully. **This will change for each type of feedback!**
  - *Example prompt interpretation:* The example task description in an image (not shown here) will be looking for an image which provides a balance between realism and color vividness, so we should strive to look for such images.
2. Review the initial set of images.
  - The realism and color vividness scores are displayed below each image.
3. Provide feedback using the current mechanism.
4. Review new images generated based on your feedback.
5. Using the same mechanism, provide feedback (i.e. pairwise, ranking, etc.) until you find a satisfactory image.
6. Select your final image by clicking "Select Final Image" at the bottom of the page.
  - This will lead you to a page with all of the images you've seen displayed. Please enter in your level of satisfaction (0-100) with the final image, in terms of how appropriate you feel it is for the task description, and click on "Select this image".
7. Move on to the next feedback mechanism.

Figure 33: MTurk User Study Instructions: Step-by-step study process.

## MTurk User Study Instructions (Part 6 of 6): Tips & Final Notes

### Tips for Effective Image Selection

- Keep the magazine's requirements in focus throughout the selection process.
- Explore different options before making your final selection.
- Consider the tradeoff between realism and color vividness based on the magazine's needs.
- Be consistent in your preferences.
- You can stop and select your final image whenever you feel you've found a suitable match.

### Technical Requirements

- Please use a desktop or laptop computer (not a mobile device).
- Ensure your screen brightness is adjusted for optimal image viewing.
- Complete the study in a single session if possible.

Thank you for your participation! Your feedback will help us understand which approaches are most effective for image selection tasks in creative professional contexts.

**Note:** Each user may complete this task up to two times.

Figure 34: MTurk User Study Instructions: Tips for selection, technical requirements, and concluding remarks.

## MTurk User Study: Post-Mechanism Survey Questions for User Study

At the conclusion of interacting with each of the four feedback mechanisms, participants were presented with the following statements and asked to rate their agreement on a 5-point Likert scale (1-Strongly Disagree, 2-Disagree, 3-Neutral, 4-Agree, 5-Strongly Agree). The placeholder “{ }” was replaced with the name of the specific feedback mechanism being evaluated (e.g., “Soft-Hard Bounds”, “Pairwise Preferences”).

1. **Mental Effort:** “{ } takes a lot of mental effort to use.”
2. **Expressiveness:** “{ } allows me to fully express my preferences.”
3. **Trust:** “I trust that my results with { } are optimal for its task.”

Figure 35: Survey questions presented to MTurk participants after their interaction with each feedback mechanism. The placeholder “{ }” was substituted with the name of the mechanism in question.

### MTurk User Study: Screening Quiz Questions (Part 1 of 2)

Participants were required to answer the following screening questions correctly (at least 3 out of 5) to proceed with the user study.

#### Question 1: Soft-Hard Bounds

*Instruction:* In the soft-hard bounds feedback mechanism state below, please drag the circle on the left to a value less than 0.2. What does that action represent?

*[Visual cue: A slider labeled “Example: Hard (Left) and Soft (Right) Bounds for Realism” is displayed, with range 0.0 to 1.0, and current values (0.3, 0.7).]*

*Answer Options:*

- A: I increased the soft bound, the preferred amount, for realism.
- B: I decreased the hard bound, the minimum acceptable amount, for realism down to below 0.2.
- C: I adjusted the medium bound, the neutral position for realism.
- D: Something else.

*(Correct Answer: B)*

#### Question 2: Soft-Hard Bounds

*Instruction:* In the soft-hard bounds feedback mechanism state below, please drag the circle on right to a value greater than 0.85. What does that action represent?

*[Visual cue: A slider labeled “Example: Hard (Left) and Soft (Right) Bounds for Color Vividness” is displayed, with range 0.0 to 1.0, and current values (0.4, 0.8).]*

*Answer Options:*

- A: I increased the hard bound, the minimum acceptable amount, for color vividness. I prefer to see images with color vividness above 0.85.
- B: I increased the soft bound, the preferred amount, for color vividness. I prefer to see images with color vividness above 0.85.
- C: I increased the medium bound, the neutral position, for realism.
- D: Something else.

*(Correct Answer: B)*

Figure 36: MTurk User Study Screening Quiz: Questions 1-3, focusing on understanding the Soft-Hard Bounds mechanism and related interface elements. Correct answers are indicated for reference.

## MTurk User Study: Screening Quiz Questions (Part 2 of 2)

### Question 3: Soft-Hard Bounds' Adjacent Images

*Instruction:* What does the image labeled as “Adjacent Image” represent in the example below?

*[Visual cue: Three images are displayed side-by-side. Image 1 caption: “Image 1 – Realism: 0.40 Color Vividness 0.86”. Image 2 caption: “Image 2 – Realism: 0.65 Color Vividness 0.63”. Image 3 caption: “Adjacent Image 3 – Realism: 0.95 Color Vividness 0.30”. Below these, two disabled sliders are shown for Realism and Color Vividness, both set to (0.4, 0.8).]*

*Answer Options:*

- A: The image that is most preferred in the ranking.
- B: An example of an image I might see if I adjust the color vividness hard bound to be slightly lower.
- C: An image that is adjacent to the current image in the ranking.
- D: An example of an image I might see if I adjust the realism soft bound to be slightly higher.

*(Correct Answer: B, based on typical interpretation of adjacent images in such contexts)*

### Question 4: Partial Ranking

*Instruction:* In “Partial Ranking”, what does the task require you to do?

*Answer Options:*

- A: Rank all images from best to worst.
- B: Rank only your top three preferred images.
- C: Select the single best image.
- D: Group images into categories.

*(Correct Answer: B)*

### Question 5: Image Attributes

*Instruction:* According to the study instructions, what is the relationship between “Realism” and “Color Vividness” in the images?

*Answer Options:*

- A: They are completely unrelated attributes.
- B: Higher realism always leads to higher color vividness.
- C: They often represent a tradeoff - increasing one may decrease the other.
- D: They are different names for the same quality.

*(Correct Answer: C)*

Figure 37: MTurk User Study Screening Quiz: Questions 4-5, assessing understanding of the Partial Ranking mechanism and the core image attribute tradeoff. Correct answers are indicated for reference.

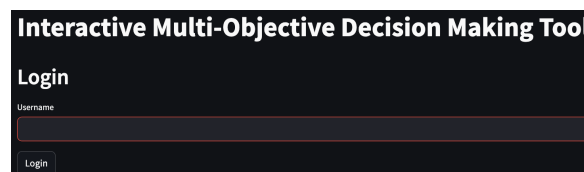


Figure 38: Login interface for the Interactive Multi-Objective Decision Making Tool used in the user study. Participants entered their MTurk username to begin.

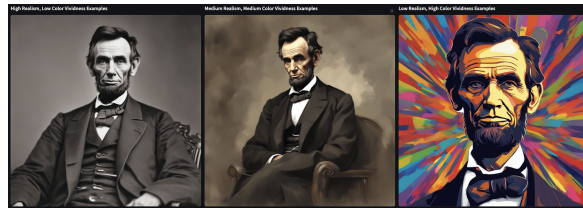


Figure 39: Instructional material presented to user study participants, illustrating the tradeoff between “Realism” and “Color Vividness” using AI-generated images of Abraham Lincoln. Examples show (L-R): High Realism/Low Color Vividness, Medium/Medium, and Low Realism/High Color Vividness.



Figure 40: Additional instructional examples illustrating the realism and color vividness tradeoff, using AI-generated images of Benjamin Franklin, provided to user study participants.

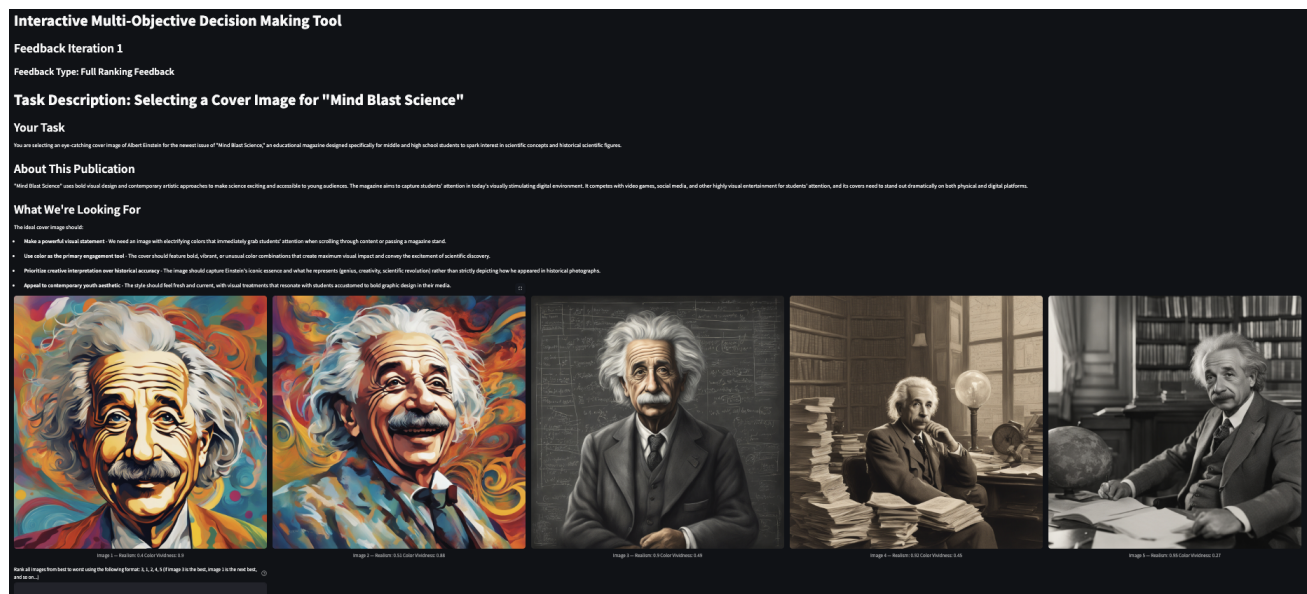


Figure 41: User interface for the Full Ranking feedback mechanism. Participants are presented with a task description (e.g., “Selecting a Cover Image for ’Mind Blast Science’”), details about the publication, desired image qualities, and a set of AI-generated images (Albert Einstein in this example) to rank from most to least preferred.

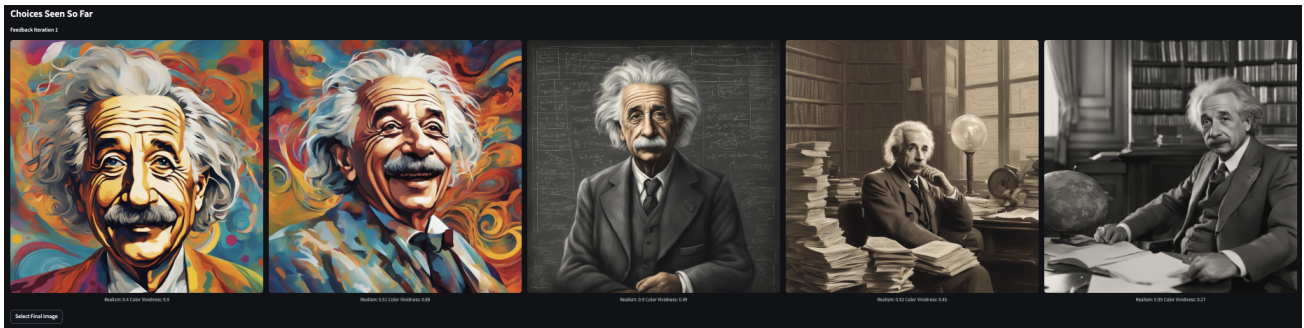


Figure 42: Interface displaying “Choices Seen So Far” during a task instance. This screen allows participants to review all images encountered up to that point before making a final selection or continuing with more feedback iterations.



Figure 43: User interface for the Partial Ranking feedback mechanism. Participants are tasked with selecting a cover photo (Barack Obama for “American Leadership Today” magazine in this instance) by ranking only their top three preferred images from the presented set.

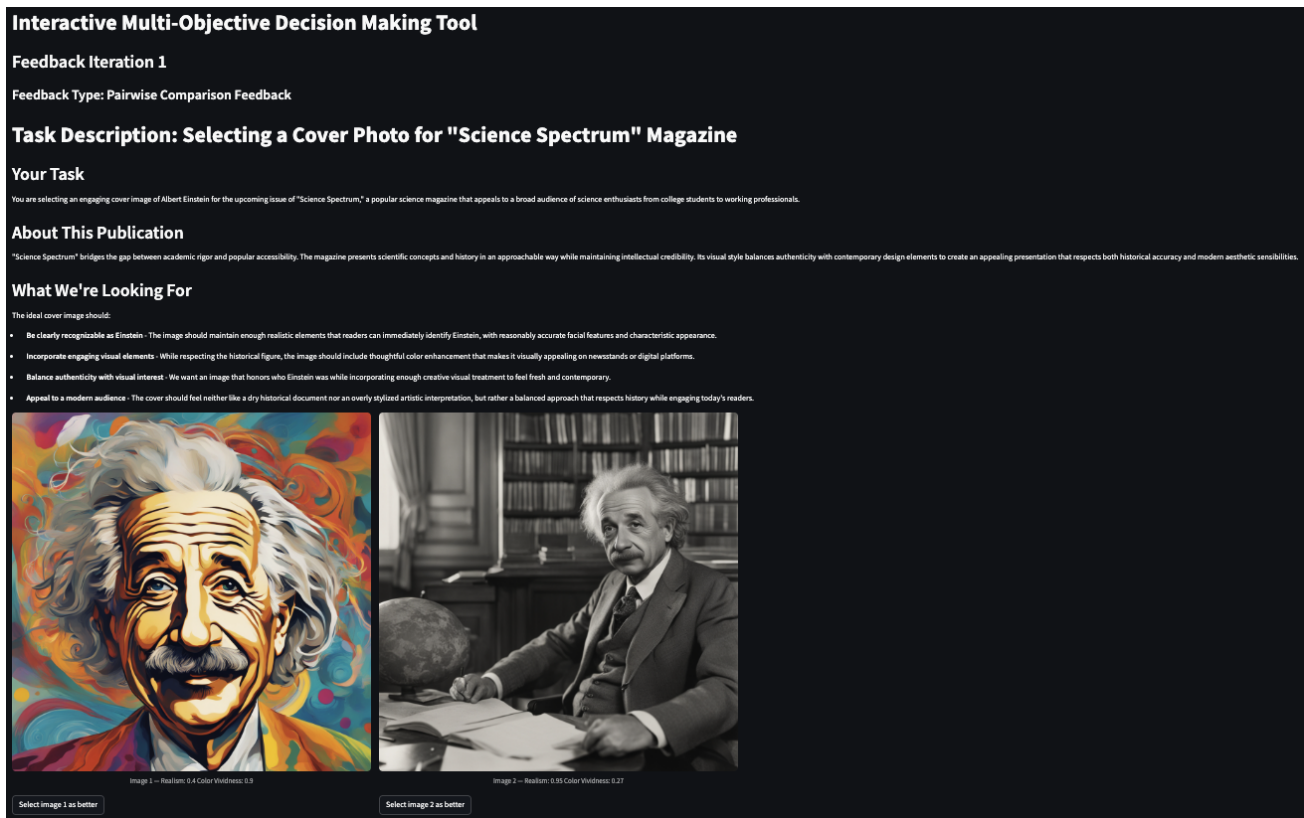


Figure 44: User interface for the Pairwise Comparison feedback mechanism. Participants are shown two AI-generated images (Albert Einstein for “Science Spectrum” magazine here) and must select the one they prefer more based on the task description.

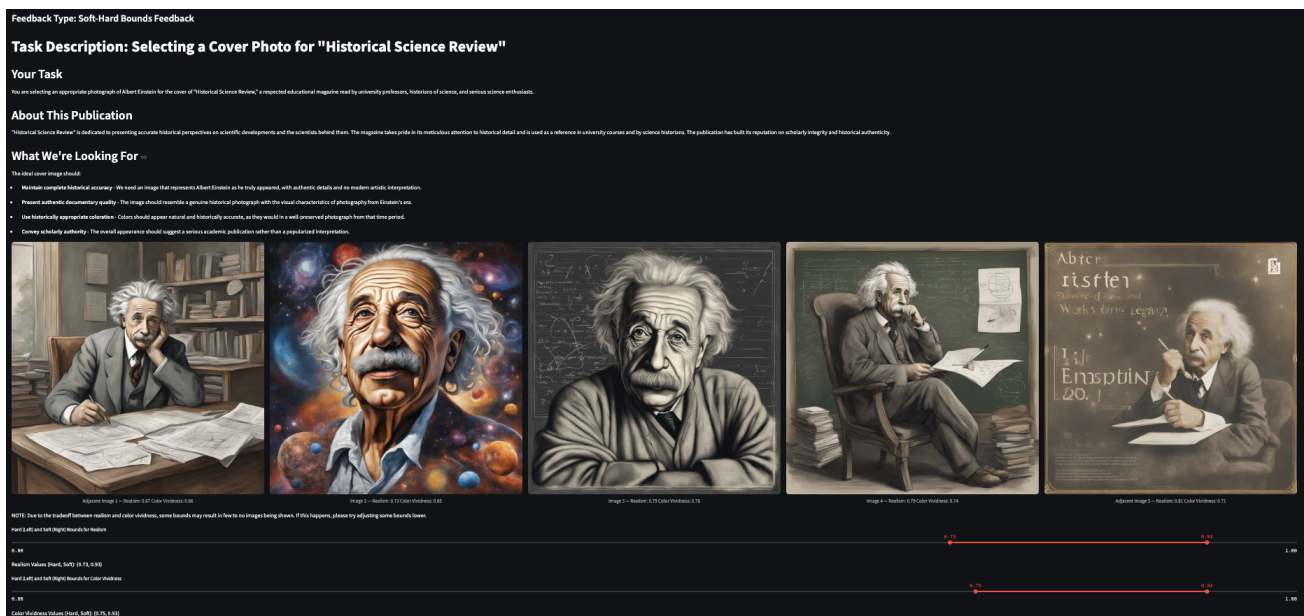


Figure 45: User interface for the Soft-Hard Bounds feedback mechanism (Active-C-MoSH). Participants adjust sliders to set soft (preferred) and hard (minimum acceptable) bounds for image attributes (realism and color vividness). Images shown are for Barack Obama for “Future Leaders: Youth Politics Magazine”.

## Task Description: Selecting a Cover Photo for "Historical Science Review"

### Your Task

You are selecting an appropriate photograph of Albert Einstein for the cover of "Historical Science Review," a respected educational magazine read by university professors, historians of science, and serious science enthusiasts.

### About This Publication

"Historical Science Review" is dedicated to presenting accurate historical perspectives on scientific developments and the scientists behind them. The magazine takes pride in its meticulous attention to historical detail and is used as a reference in university courses and by science historians. The publication has built its reputation on scholarly integrity and historical authenticity.

### What We're Looking For

The ideal cover image should:

- **Maintain complete historical accuracy** - We need an image that represents Albert Einstein as he truly appeared, with authentic details and no modern artistic interpretation.
- **Present authentic documentary quality** - The image should resemble a genuine historical photograph with the visual characteristics of photography from Einstein's era.
- **Use historically appropriate coloration** - Colors should appear natural and historically accurate, as they would in a well-preserved photograph from that time period.
- **Convey scholarly authority** - The overall appearance should suggest a serious academic publication rather than a popularized interpretation.

Figure 46: User study task description for selecting an Albert Einstein cover image for "Historical Science Review." This task targets a high realism ( $\sim 90$ ) and low color vividness ( $\sim 5$ ) tradeoff point.

## Task Description: Selecting a Cover Image for "Mind Blast Science"

### Your Task

You are selecting an eye-catching cover image of Albert Einstein for the newest issue of "Mind Blast Science," an educational magazine designed specifically for middle and high school students to spark interest in scientific concepts and historical scientific figures.

### About This Publication

"Mind Blast Science" uses bold visual design and contemporary artistic approaches to make science exciting and accessible to young audiences. The magazine aims to capture students' attention in today's visually stimulating digital environment. It competes with video games, social media, and other highly visual entertainment for students' attention, and its covers need to stand out dramatically on both physical and digital platforms.

### What We're Looking For

The ideal cover image should:

- **Make a powerful visual statement** - We need an image with electrifying colors that immediately grab students' attention when scrolling through content or passing a magazine stand.
- **Use color as the primary engagement tool** - The cover should feature bold, vibrant, or unusual color combinations that create maximum visual impact and convey the excitement of scientific discovery.
- **Prioritize creative interpretation over historical accuracy** - The image should capture Einstein's iconic essence and what he represents (genius, creativity, scientific revolution) rather than strictly depicting how he appeared in historical photographs.
- **Appeal to contemporary youth aesthetic** - The style should feel fresh and current, with visual treatments that resonate with students accustomed to bold graphic design in their media.

Figure 47: User study task description for selecting an Albert Einstein cover image for "Mind Blast Science." This task targets a low realism (~15) and high color vividness (~85) tradeoff point.

## Task Description: Selecting a Cover Photo for "Science Spectrum" Magazine

### Your Task

You are selecting an engaging cover image of Albert Einstein for the upcoming issue of "Science Spectrum," a popular science magazine that appeals to a broad audience of science enthusiasts from college students to working professionals.

### About This Publication

"Science Spectrum" bridges the gap between academic rigor and popular accessibility. The magazine presents scientific concepts and history in an approachable way while maintaining intellectual credibility. Its visual style balances authenticity with contemporary design elements to create an appealing presentation that respects both historical accuracy and modern aesthetic sensibilities.

### What We're Looking For

The ideal cover image should:

- **Be clearly recognizable as Einstein** - The image should maintain enough realistic elements that readers can immediately identify Einstein, with reasonably accurate facial features and characteristic appearance.
- **Incorporate engaging visual elements** - While respecting the historical figure, the image should include thoughtful color enhancement that makes it visually appealing on newsstands or digital platforms.
- **Balance authenticity with visual interest** - We want an image that honors who Einstein was while incorporating enough creative visual treatment to feel fresh and contemporary.
- **Appeal to a modern audience** - The cover should feel neither like a dry historical document nor an overly stylized artistic interpretation, but rather a balanced approach that respects history while engaging today's readers.

Figure 48: User study task description for selecting an Albert Einstein cover image for "Science Spectrum" Magazine. This task targets a medium realism and medium color vividness tradeoff point.

#### User Study Task Description: Obama for "Presidential Archives Quarterly"

### Task Description: Selecting a Cover Photo for "Presidential Archives Quarterly"

#### Your Task

You are selecting an appropriate photograph of Barack Obama for the cover of "Presidential Archives Quarterly," a prestigious educational journal published by the National Historical Society that documents and analyzes American presidential administrations.

#### About This Publication

"Presidential Archives Quarterly" serves as a definitive historical record consulted by political scientists, historians, and academic institutions worldwide. The publication maintains strict standards for historical accuracy and documentary authenticity, presenting unembellished historical documentation for scholarly analysis and historical preservation.

#### What We're Looking For

The ideal cover image should:

- **Maintain historical precision** - We need an image that depicts President Obama exactly as he appeared during his presidency, with complete accuracy in facial features, expressions, and physical details.
- **Present authentic documentary quality** - The image should have the characteristics of professional photo-journalism or official White House photography, preserving the authentic visual record of his presidency.
- **Use minimal coloration** - Colors should appear minimal, as they would in official presidential photography, without enhancement or creative color styling.
- **Convey presidential gravitas** - The overall appearance should reflect the dignity and historical significance of the presidential office, suitable for archival documentation.

Figure 49: User study task description for selecting a Barack Obama cover image for "Presidential Archives Quarterly." This task targets a high realism ( $\sim 90$ ) and low color vividness ( $\sim 5$ ) tradeoff point.

User Study Task Description: Obama for "American Leadership Today"

## Task Description: Selecting a Cover Photo for "American Leadership Today"

### Your Task

You are selecting an engaging cover image of Barack Obama for the upcoming issue of "American Leadership Today," a popular educational magazine that explores political leadership for a diverse audience ranging from college students to professionals.

### About This Publication

"American Leadership Today" bridges academic analysis and accessible education about American political leaders. The magazine maintains credibility while using contemporary visual approaches to engage readers. Its visual style balances authenticity with modern design elements to appeal to readers who value both substance and visual engagement.

### What We're Looking For

The ideal cover image should:

- **Be clearly recognizable as Obama** - The image should maintain his distinctive features and characteristic appearance while allowing for thoughtful creative interpretation.
- **Incorporate engaging color elements** - While preserving Obama's recognizable likeness, the image should include purposeful color choices that enhance visual appeal and emotional resonance.
- **Balance historical representation with contemporary style** - We want an image that acknowledges Obama's historical significance while incorporating enough creative visual treatment to feel fresh and relevant.
- **Appeal to both substance-focused and visually-oriented readers** - The cover should feel neither like a dry historical document nor an overly stylized artistic piece, but rather a thoughtful balance that respects the subject while engaging modern viewers.

Figure 50: User study task description for selecting a Barack Obama cover image for "American Leadership Today." This task targets a medium realism (~55) and medium color vividness (~60) tradeoff point.

User Study Task Description: Obama for "Future Leaders: Youth Politics Magazine"

## Task Description: Selecting a Cover Image for "Future Leaders: Youth Politics Magazine"

### Your Task

You are selecting an eye-catching cover image of Barack Obama for the newest issue of "Future Leaders," an educational magazine designed to spark political interest and civic engagement among middle and high school students.

### About This Publication

"Future Leaders" uses bold visual design and contemporary artistic approaches to make political history and civic education exciting and relevant to young audiences. The magazine competes for attention in students' visually saturated digital environment and needs covers that immediately capture interest on school library shelves, social media feeds, and digital platforms.

### What We're Looking For

The ideal cover image should:

- **Create maximum visual impact** - We need an image with electrifying colors that immediately grab students' attention—think vibrant blues, bold reds, striking color contrasts, or unexpected color treatments.
- **Use color as a storytelling element** - The color palette should convey energy, inspiration, and the transformative potential of political engagement, making Obama's legacy feel relevant to today's youth.
- **Prioritize artistic interpretation over strict realism** - The image should capture Obama's iconic status and what he represents symbolically (change, hope, historical significance) rather than simply documenting how he appeared in photographs.
- **Connect with youth visual culture** - The style should incorporate graphic elements, color treatments, and visual approaches that resonate with students familiar with bold social media aesthetics and contemporary digital art.

Figure 51: User study task description for selecting a Barack Obama cover image for "Future Leaders: Youth Politics Magazine." This task targets a low realism (~15) and high color vividness (~85) tradeoff point.