

Amortized Multi-Objective Optimization Across Tasks with Generative Solution Modeling

Tingyang Wei¹, Jiao Liu¹, Abhishek Gupta^{2*}, Chin Chun Ooi³,
Puay Siew Tan⁴ and Yew-Soon Ong^{1,3}

¹College of Computing and Data Science, Nanyang Technological University, Singapore

²School of Mechanical Sciences, Indian Institute of Technology, Goa, India

³Centre for Frontier AI Research, Agency for Science, Technology and Research, Singapore

⁴Singapore Institute of Manufacturing Technology, Agency for Science, Technology and Research, Singapore

{tingyang001}@e.ntu.edu.sg, {jiao.liu, asysong}@ntu.edu.sg, abhishekgupta@iitgoa.ac.in, {ooicc, tan_puay_siew}@a-star.edu.sg

Abstract

Many real-world applications require solving families of expensive multi-objective optimization problems (EMOPs) under varying operational conditions. This can be formulated as parametric expensive multi-objective optimization problems (P-EMOPs) where each task parameter defines a distinct optimization instance. Current multi-objective Bayesian optimization methods have been widely used for finding finite sets of Pareto optimal solutions for each task. However, P-EMOPs present a fundamental challenge: the continuous task parameter space can contain infinite distinct problems, each requiring separate expensive evaluations. To address this, we propose learning an inverse model to amortize the multi-objective optimization cost across the continuous task-preference space, enabling direct solution prediction for any query without the need for expensive re-evaluation. This paper introduces a novel parametric multi-objective Bayesian optimizer that learns this inverse model by alternating between (1) generative solution sampling via conditional generative models and (2) acquisition-driven search leveraging inter-task synergies. This approach enables effective optimization across multiple tasks and finally achieves direct solution prediction for unseen parameterized EMOPs without re-evaluations. We theoretically justify the faster convergence by leveraging inter-task synergies through task-aware Gaussian processes. Based on that, empirical studies in synthetic and real-world benchmarks further verify the effectiveness of the proposed parametric optimizer.

1 Introduction

Expensive multi-objective optimization problems (EMOPs) are crucial in black-box optimization, aiming to identify op-

timal trade-off solutions among multiple conflicting criteria under limited evaluation budgets [Knowles, 2006; Min *et al.*, 2019]. Due to their focus on conflicting objective functions and costly evaluations, EMOP formulations have been applied to a plethora of real-world applications encompassing materials discovery [Li *et al.*, 2024], drug discovery [Zhu *et al.*, 2023], and robotics [Turchetta *et al.*, 2020; Kim *et al.*, 2021].

To efficiently identify trade-off solutions, surrogate modeling has become essential for reducing the cost of evaluating expensive objectives. Among these, multi-objective Bayesian optimization (MOBO) stands out as the leading strategy. MOBO methods have been widely utilized for searching finite sets of Pareto Sets by scalarizing multiple objectives into single objectives [Knowles, 2006; Paria *et al.*, 2020] or by optimizing indicators such as expected Hypervolume improvement [Emmerich *et al.*, 2006; Daulton *et al.*, 2020] and uncertainty reduction [Picheny, 2015; Hernandez-Lobato *et al.*, 2016; Tu *et al.*, 2022; Belakaria *et al.*, 2019]. To move beyond these finite solution-set approximations and directly model the entire Pareto front, recent studies have extended MOBO to continuous Pareto-manifold learning methods [Lin *et al.*, 2022; Li *et al.*, 2025; Cheng *et al.*, 2025b].

However, in many real-world scenarios, optimization problems emerge not as isolated tasks but as families of related tasks under varying operational conditions. For example, turbine components must be optimized across different wind speeds [Makatura *et al.*, 2021], and robotic control policies must be optimized for diverse damage conditions [Cully *et al.*, 2015]. These cases lead to fundamental challenges: the continuous range of operational conditions forms infinitely many related optimization problems, each requiring separate, costly evaluations for current MOBO methods. As operational conditions frequently shift [Nickelson *et al.*, 2023], repeatedly performing full optimizations to evaluate trade-offs across multiple scenarios becomes prohibitively expensive and computationally infeasible. Therefore, tackling families of related EMOPs urges new approaches capable of amortizing the efforts of repeated multi-objective optimizations across the operational condition space into a lightweight

*Corresponding author

framework. To our best knowledge, little existing work has systematically addressed this parametric multi-task setting. Some related works are discussed in Appendix E.

This paper proposes to amortize the work of potential repeated multi-objective optimization across infinite tasks via a single inverse generative model pass, introducing a novel parametric multi-objective Bayesian optimizer. Our approach alternates between two complementary phases: (1) generative solution sampling via conditional generative models, which learns high-quality solution distributions conditioned on task parameters and preferences, and (2) acquisition-driven search guided by task-aware Gaussian processes, which explicitly leverages inter-task correlations to efficiently explore the solution space. This dual strategy not only can enhance optimization efficiency across multiple related tasks but also finally builds an inverse model for predicting optimized solutions to previously unseen EMOPs without expensive re-evaluations. The major contributions of this paper can be summarized as follows:

- We propose a novel inverse generative model to predict optimized solutions to each task query within the continuous task parameter space, thereby addressing the challenge of solving families of EMOPs.
- We design an alternating optimization approach that curates the conditional generative solution sampling with acquisition-driven search, enabling both efficient multi-task optimization and inverse model building.
- We develop a Parametric Multi-Task Multi-Objective Bayesian Optimizer (PMT-MOBO) with task-aware Gaussian Processes, demonstrating its faster convergence over single-task counterparts empirically and theoretically.

2 Preliminaries

2.1 Expensive Multi-Objective Optimization

We consider expensive multi-objective optimization problems (EMOPs) defined as follows:

$$\min_{\mathbf{x} \in \mathcal{X}} F(\mathbf{x}) := \min_{\mathbf{x} \in \mathcal{X}} (f_1(\mathbf{x}), \dots, f_M(\mathbf{x})), \quad (1)$$

where $\mathbf{x} \in \mathcal{X} \subseteq \mathbb{R}^D$ denotes the decision variable, and each black-box objective function $f_m(x) : \mathcal{X} \rightarrow \mathbb{R}, m = 1, \dots, M$, can be costly to evaluate and exhibits trade-offs with other objective functions. The goal of optimizing EMOP is to approximate the Pareto set of non-dominated solutions that balance these conflicting objectives under limited evaluation budgets.

Definition 1 (Pareto Dominance). Solution $\mathbf{x}^{(a)} \in \mathcal{X}$ is considered Pareto dominate another solution $\mathbf{x}^{(b)} \in \mathcal{X}$ if $\forall i \in [M], f_i(\mathbf{x}^{(a)}) \leq f_i(\mathbf{x}^{(b)})$ and $\exists i' \in [M]$ such that $f_{i'}(\mathbf{x}^{(a)}) < f_{i'}(\mathbf{x}^{(b)})$ [Branke, 2008].

Definition 2 (Pareto Optimality). Solution $\mathbf{x}^* \in \mathcal{X}$ is considered Pareto optimal or non-dominated if there exists no other candidate solutions that can dominate $\mathbf{x}^*$. Consequently, Pareto Set (PS) and Pareto Front (PF) contain all the set of Pareto optimal solutions and their corresponding objective vectors [Branke, 2008].

2.2 Parametric Expensive Multi-Objective Optimization

In many real-world scenarios, EMOPs must be solved repeatedly under varying environmental [Makatura *et al.*, 2021] or operational conditions [Pierrot *et al.*, 2022]. These variations can be represented using continuous *task parameters* $\theta \in \Theta \subseteq \mathbb{R}^V$, which defines a family of different EMOPs over the same decision space. This leads to the formulation of a parametric expensive multi-objective optimization problem (P-EMOP):

$$\begin{cases} \min_{\mathbf{x} \in \mathcal{X}} F(\mathbf{x}, \theta) := \min_{\mathbf{x} \in \mathcal{X}} (f_1(\mathbf{x}, \theta), \dots, f_M(\mathbf{x}, \theta)) \\ \mathcal{M}(\theta, \lambda) \approx \arg \min_{\mathbf{x} \in \mathcal{X}} s_{\lambda}(F(\mathbf{x}, \theta)), \forall \lambda \in \Lambda \subset \mathbb{R}^M \end{cases} \quad (2)$$

where each task parameter θ defines a distinct expensive multi-objective problem task $F(\mathbf{x}, \theta)$. The ideal goal of P-EMOP is to search all the Pareto Sets across potentially infinite task parameters under limited evaluation budgets. We address this by learning an inverse model $\mathcal{M}(\theta, \lambda)$ that generates optimized solutions for any task parameter θ and preference vector λ without expensive re-optimization. This represents a shift from single-task optimization to learning generalized solution mappings across the continuous task parameter space. Here, $s_{\lambda}(\cdot) : \mathbb{R}^M \rightarrow \mathbb{R}$ is a scalarization function characterized by preference vector λ . We incorporate λ into the inverse model since the Pareto Set generally lies on a continuous manifold for non-trivial multi-objective optimization [Miettinen, 1999], enabling $\mathcal{M}(\theta, \lambda)$ to output a single solution vector rather than a solution set.

Although (2) contains infinitely many tasks, these EMOPs share inter-task synergies due to their parametric relationship. Rather than treating tasks in isolation, we solve a finite sample $\{\theta_1, \dots, \theta_K\}$ as a multi-task optimization problem [Wei *et al.*, 2022], assuming tasks belong to a coherent family sharing solution components or objective landscape regularities. This assumption, pioneered by multi-objective multifactorial optimization [Gupta *et al.*, 2016b] and commonly exploited in subsequent multi-task optimization works [Gupta *et al.*, 2016a; Wei *et al.*, 2024; Bali *et al.*, 2021] and recent parametric multi-objective optimization studies [Cheng *et al.*, 2025a], enables more efficient solution search across the task parameter space through inter-task synergies, thereby yielding higher-quality solutions for constructing the inverse model.

2.3 Gaussian Processes for Multi-Objective Optimization

In this paper, we adopt Gaussian Process (GP) [Williams and Rasmussen, 2006] as surrogate models to efficiently solve expensive multi-objective optimization problems (EMOPs). A GP defines a distribution over functions, assuming that $f \sim \mathcal{GP}(\mu(\cdot), \kappa(\cdot, \cdot))$, where $\mu(\mathbf{x}) = \mathbb{E}[f(\mathbf{x})]$ is the mean function and $\kappa(\mathbf{x}, \mathbf{x}') = \text{Cov}[f(\mathbf{x}), f(\mathbf{x}')] is the kernel or covariance function. Given a dataset $\{(\mathbf{x}^{(i)}, y^{(i)})\}_{i=1}^N$, where each observation is corrupted by Gaussian noise $y^{(i)} = f(\mathbf{x}^{(i)}) + \epsilon^{(i)}, \epsilon^{(i)} \sim \mathcal{N}(0, \sigma_{\epsilon}^2)$, the posterior at a test point $\mathbf{x}$ is a normal distribution $\mathcal{N}(\mu(\mathbf{x}), \sigma^2(\mathbf{x}))$ with:$

$$\begin{cases} \mu(\mathbf{x}) = \mathbf{k}^{\top} (\mathbf{K} + \sigma_{\epsilon}^2 \mathbf{I}_N)^{-1} \mathbf{y}, \\ \sigma^2(\mathbf{x}) = \kappa(\mathbf{x}, \mathbf{x}) - \mathbf{k}^{\top} (\mathbf{K} + \sigma_{\epsilon}^2 \mathbf{I}_N)^{-1} \mathbf{k}, \end{cases} \quad (3)$$

where $\mathbf{k} \in \mathbb{R}^N$ is the kernel vector between $\mathbf{x}$ and training inputs, $\mathbf{K} \in \mathbb{R}^{N \times N}$ is the kernel matrix with $\mathbf{K}_{p,q} = \kappa(\mathbf{x}^{(p)}, \mathbf{x}^{(q)})$, and $\mathbf{y} \in \mathbb{R}^N$ collects the observed outputs.

In EMOPs, Gaussian Processes (GPs) serve as probabilistic surrogates for approximating costly objective functions, enabling uncertainty-aware exploration using acquisition functions. Following the theoretically sound multi-objective Bayesian optimization (MOBO) framework in [Paria *et al.*, 2020], in this paper, we adopt independent GP models for each objective and apply a scalarization-based acquisition function. In each iteration, we randomly sample a preference vector $\boldsymbol{\lambda} \in \Lambda$ that defines a trade-off among objectives via the scalarization function $\mathfrak{s}_{\boldsymbol{\lambda}}(\cdot)$, and optimize a scalarized acquisition function score, defined as:

$$\mathbf{x}^{(t)} = \operatorname{argmax}_{\mathbf{x} \in \mathcal{X}} \mathfrak{s}_{\boldsymbol{\lambda}}(\boldsymbol{\alpha}(\mathbf{x})), \quad (4)$$

where $\boldsymbol{\alpha}(\mathbf{x}) \in \mathbb{R}^M$ denotes the vector of acquisition function values for each objective. We refer to this method as ST-MOBO, which serves both as the base module in our framework and as a single-task (ST) baseline EMOP solver in our empirical study.

3 Methodology

3.1 Overview

To efficiently solve a family of EMOPs under limited evaluation budgets in (2), we propose to build an inverse model by an alternating optimization framework that integrates task-aware surrogate modeling and conditional generative modeling, as shown in Figure 1. The former, instantiated via the Parametric Multi-Task Multi-Objective Bayesian Optimization (PMT-MOBO) module, exploits inter-task relationships via task-aware GP to guide acquisition-driven search in a sample-efficient manner. The latter, introduced as the conditional generative model, captures the distribution of high-performing solutions and enables task- and preference-aligned search in promising regions. Both components are executed in an alternating loop, as illustrated in Figure 1, enabling both sample-efficient multi-task search and finally achieving a high-quality inverse model.

3.2 PMT-MOBO with Task-Aware GP

We extend the baseline ST-MOBO to PMT-MOBO by incorporating task-aware GPs. Each task-aware GP model receives both the decision variable $\mathbf{x}$ and the task parameter $\boldsymbol{\theta}$ as joint input. Formally, for a given EMOP instance parametrized by $\boldsymbol{\theta}'$, the m -th objective is modeled as:

$$f_m(\cdot, \boldsymbol{\theta}') \sim \mathcal{GP}_m(\tilde{\mu}_m(\cdot, \boldsymbol{\theta}'), \tilde{\kappa}_m((\cdot, \boldsymbol{\theta}'), (\cdot, \boldsymbol{\theta}'))), m \in [M]$$

with mean functions and kernels defined over the joint input space $\mathcal{X} \times \Theta$. Based on these models, we extend the scalarized acquisition function defined in (4) into follows:

$$\mathbf{x}_k^{(t)} = \operatorname{argmax}_{\mathbf{x} \in \mathcal{X}} \mathfrak{s}_{\boldsymbol{\lambda}}(\tilde{\boldsymbol{\alpha}}(\mathbf{x}, \boldsymbol{\theta}_k)), \quad (5)$$

where $\tilde{\boldsymbol{\alpha}}(\mathbf{x}, \boldsymbol{\theta}_k) \in \mathbb{R}^M$ denotes the vector of task-aware acquisition function values for the EMOP task instance parametrized by $\boldsymbol{\theta}_k$. We refer to this method as PMT-MOBO. For

clarity, a comparative workflow of either ST-MOBO or PMT-MOBO is showcased in **Algorithm 2** of Appendix F.

By explicitly incorporating task parameters into the GP models, PMT-MOBO accounts for inter-task relationships during solution acquisition and model training. In the next subsection, we provide a theoretical analysis that quantifies its improved convergence compared to ST-MOBO.

3.3 Theoretical Analysis of Faster Convergence

For our theoretical analysis, we instantiate the acquisition functions $\boldsymbol{\alpha}(\mathbf{x})$ and $\tilde{\boldsymbol{\alpha}}(\mathbf{x}, \boldsymbol{\theta}_k)$ using the Lower Confidence Bound (LCB) [Srinivas *et al.*, 2012]. These LCB instantiations coincide with the acquisition rules employed in our practical ST-MOBO and PMT-MOBO algorithms, so the following regret analysis directly characterizes the convergence behavior of the implemented methods. Specifically, for ST-MOBO, the acquisition function in (4) becomes:

$$\mathbf{x}^{(t)} = \operatorname{argmax}_{\mathbf{x} \in \mathcal{X}} \mathfrak{s}_{\boldsymbol{\lambda}}(-\boldsymbol{\mu}(\mathbf{x}) + \beta^{1/2} \boldsymbol{\sigma}(\mathbf{x})), \quad (6)$$

where β denotes a trade-off coefficient, and $\boldsymbol{\mu}(\cdot), \boldsymbol{\sigma}(\cdot) \in \mathbb{R}^M$ denote the vector of predictive means and standard deviations of M GP models. Correspondingly, for PMT-MOBO, the task-aware acquisition function in (5) becomes:

$$\mathbf{x}_k^{(t)} = \operatorname{argmax}_{\mathbf{x} \in \mathcal{X}} \mathfrak{s}_{\boldsymbol{\lambda}}(-\tilde{\boldsymbol{\mu}}(\mathbf{x}, \boldsymbol{\theta}_k) + \beta^{1/2} \tilde{\boldsymbol{\sigma}}(\mathbf{x}, \boldsymbol{\theta}_k)). \quad (7)$$

For an EMOP with a total evaluation budget T , the goal is to obtain a set of solutions $\mathbf{X}_T \subset \mathcal{X}$ that approximates the true Pareto Set $\mathcal{X}^*$. Following the analysis technique [Paria *et al.*, 2020], this goal can be quantified using the notion of *Bayes Regret*, which evaluates the quality of the obtained solution set with respect to random preference vectors.

Definition 3 (Bayes Regret). Given a preference distribution $p(\boldsymbol{\lambda})$ over preference vectors, the Bayes regret of a solution set $\mathbf{X}_T$ is defined as:

$$\mathcal{R}_B(T) = \mathbb{E}_{\boldsymbol{\lambda} \sim p(\boldsymbol{\lambda})} \left[\min_{\mathbf{x} \in \mathbf{X}_T} \mathfrak{s}_{\boldsymbol{\lambda}}(F(\mathbf{x})) - \min_{\mathbf{x} \in \mathcal{X}^*} \mathfrak{s}_{\boldsymbol{\lambda}}(F(\mathbf{x})) \right],$$

where Pareto Set $\mathcal{X}^*$ includes all the Pareto-optimal solutions $\mathbf{x}_k^* = \operatorname{argmin}_{\mathbf{x} \in \mathcal{X}} \mathfrak{s}_{\boldsymbol{\lambda}}(F(\mathbf{x})), \forall \boldsymbol{\lambda} \in \Lambda$.

Deriving an upper bound on Bayes regret is non-trivial [Zhang and Golovin, 2020], since it measures the collective regret of the entire solution set rather than point-wise regret. Therefore, a surrogate of Bayes regret is introduced, termed as *Cumulative Regret*, which quantifies the total point-wise regret incurred across all evaluation rounds.

Definition 4 (Cumulative Regret). Given a set of sampled preference vectors $\{\boldsymbol{\lambda}^{(t)}\}_{t=1}^T$, the cumulative regret of a solution set $\mathbf{X}_T = \{\mathbf{x}^{(t)}\}_{t=1}^T$ is defined as:

$$\mathcal{R}_C(T) = \sum_{t=1}^T \left[\mathfrak{s}_{\boldsymbol{\lambda}^{(t)}}(F(\mathbf{x}^{(t)})) - \min_{\mathbf{x} \in \mathcal{X}^*} \mathfrak{s}_{\boldsymbol{\lambda}^{(t)}}(F(\mathbf{x})) \right].$$

Assumption 1. For each objective f_m and each task parameter $\boldsymbol{\theta}_k$, the function $\mathbf{x} \mapsto f_m(\mathbf{x}, \boldsymbol{\theta}_k)$ satisfies the same regularity and GP assumptions as in Theorem 1 of [Paria *et al.*, 2020] (independent sub-Gaussian noise on a compact domain, and scalarizations that are monotone and Lipschitz in all objectives).

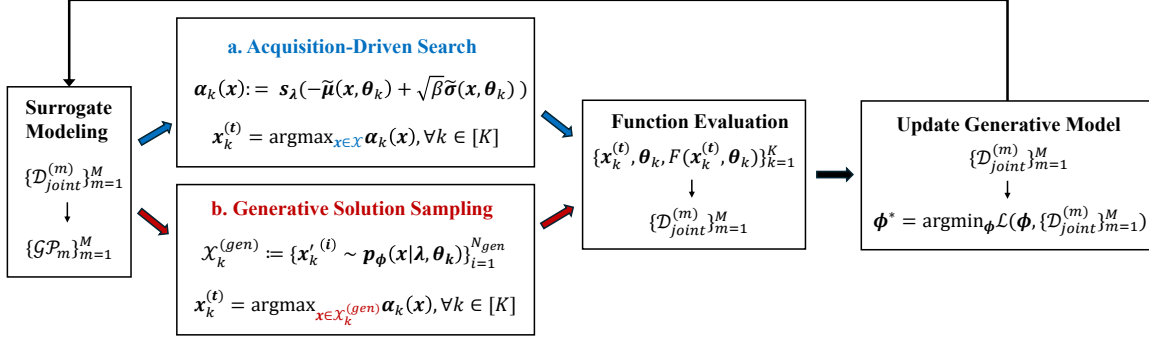


Figure 1: The overall workflow of the proposed method. During the optimization process, solutions are sampled and selected in an alternating way between: a. acquisition-driven search and b. generative solution sampling.

Theorem 1 (Regret Bounds for ST-MOBO, Theorem 1 in [Paria *et al.*, 2020]). Under Assumption 1, ST-MOBO satisfies:

$$\mathbb{E}[\mathcal{R}_C(T)] = \mathcal{O}\left(J \sqrt{M^2 T D \gamma_T \frac{\ln T}{\ln(1 + \sigma^{-2})}}\right), \quad (8)$$

and furthermore, its Bayes regret is controlled by:

$$\mathbb{E}[\mathcal{R}_B(T)] \leq \frac{1}{T} \mathbb{E}[\mathcal{R}_C(T)] + o(1). \quad (9)$$

In Theorem 1, γ_T denotes the largest *maximum information gain* (MIG) over all M objective functions:

$$\gamma_T = \max_{m \in [M]} \gamma_{T,m}, \quad (10)$$

where $\gamma_{T,m}$ denotes the MIG for each objective function $f_m(\cdot)$, quantifying the maximal uncertainty reduction about the objective function $f_m(\cdot)$ from observing T solutions:

$$\gamma_{T,m} = \max_{\{\mathbf{x}^{(t)}\}_{t=1}^T \subset \mathcal{X}} I(\{y_m^{(t)}\}_{t=1}^T; \{f_m(\mathbf{x}^{(t)})\}_{t=1}^T), \quad (11)$$

where I stands for mutual information. Although Theorem 1 is stated for a single EMOP instance, it applies equally well to each of the K tasks in isolation. For task k , let

$$\gamma_{T,m,k} = \max_{\{\mathbf{x}_k^{(t)}\}_{t=1}^T \subset \mathcal{X}} I(\{y_{m,k}^{(t)}\}_{t=1}^T; \{f_m(\mathbf{x}_k^{(t)}, \theta_k)\}_{t=1}^T) \quad (12)$$

be the task-specific MIG of ST-MOBO and then **Theorem 1** gives the corresponding $\mathbb{E}[\mathcal{R}_C^{(k)}(T)]$ and $\mathbb{E}[\mathcal{R}_B^{(k)}(T)]$ with γ_T replaced by $\gamma_{T,k} = \max_m \gamma_{T,m,k}$. With the same upper bound of Lipschitz constant J and observation noise variance σ^2 , the Bayes regret bound is determined primarily by per-task MIG $\gamma_{T,k}$. A reduction in the MIG translates into a smaller upper bound on Bayes regret. Thus, establishing that PMT-MOBO's per-task MIG $\tilde{\gamma}_{T,k}$ is smaller than that of ST-MOBO confirms its accelerated convergence.

Theorem 2 (Tightened Regret Bounds for PMT-MOBO). Under Assumption 1, consider a PMT-MOBO model where the task-aware kernel yields covariance values equivalent to the single-task kernel for any pair of inputs belonging to the

same task. Let $\tilde{\gamma}_{T,k} = \max_{m \in [M]} \tilde{\gamma}_{T,m,k}$ be the largest MIG of PMT-MOBO across M objectives over T scalarized evaluations for task k .

Then, for all $m \in [M]$ and $k \in [K]$, it holds that

$$\tilde{\gamma}_{T,m,k} \leq \gamma_{T,m,k}, \quad (13)$$

and accordingly $\tilde{\gamma}_{T,k} \leq \gamma_{T,k}$ for every $k \in [K]$. Since the proof of Theorem 1 depends on each task only through its MIG term, the same regret bounds remain valid with $\gamma_{T,k}$ replaced by $\tilde{\gamma}_{T,k}$, and consequently

$$\mathbb{E}[\mathcal{R}_B^{\text{PMT}}(T)] \leq \mathbb{E}[\mathcal{R}_B(T)], \quad (14)$$

i.e., PMT-MOBO achieves a tighter upper bound on Bayes regret than standard ST-MOBO.

Theorem 2 establishes that PMT-MOBO achieves a tighter Bayes regret bound than ST-MOBO. The detailed proof (Appendix A) demonstrates that conditioning on observations from related tasks reduces the posterior covariance for the target task. Crucially, under the same regularity assumptions utilized in [Paria *et al.*, 2020], one can ensure that the structural form of the original regret upper bound remains valid. Since this bound is a monotonically increasing function of the Maximum Information Gain (MIG), the inequality, $\tilde{\gamma}_{T,k} \leq \gamma_{T,k}$ derived from the reduced posterior variance, translates to a strictly lower upper bound on the cumulative regret.

3.4 Conditional Generative Model

Pure acquisition-based search via PMT-MOBO provides sample-efficient guidance, but can steer sampling into low-quality regions when $\mathcal{X}$ is large in (5) and budgets are tight [Brookes *et al.*, 2019; Zhang *et al.*, 2022]. To complement this, we learn an elite-solution distribution via a conditional generative model. As illustrated in Figure 1, we let

$$\mathbf{x}' \sim p_\phi(\mathbf{x} \mid \boldsymbol{\lambda}, \theta_k),$$

where ϕ denotes the model parameters, $\boldsymbol{\lambda}$ encodes the scalarization preference, and θ_k specifies the current task k . We draw a pool of candidate solutions $\mathcal{X}_k^{(\text{gen})}$ per task from this model, then score each with the acquisition function in (5) and pick the best for evaluation as follows:

$$\mathbf{x}_k^{(t)} = \operatorname{argmax}_{\mathbf{x} \in \mathcal{X}_k^{(\text{gen})}} s_\lambda(-\tilde{\boldsymbol{\mu}}(\mathbf{x}, \theta_k) + \beta^{1/2} \tilde{\boldsymbol{\sigma}}(\mathbf{x}, \theta_k)), \quad (15)$$

Algorithm 1 Generative PMT-MOBO Framework

Input: Tasks $\{\theta_k\}_{k=1}^K$, objectives $\{f_m\}_{m=1}^M$, preference dist. Λ , generative model $\mathcal{G}_\phi$.

Output: Final datasets $\{\mathcal{D}_k\}_{k=1}^K$

```

1: for  $t = 1$  to  $T$  do
2:   Sample  $\lambda^{(t)} \sim \Lambda$ 
3:   mode  $\leftarrow$  ACQUISITION or GENERATIVE (Round-Robin)
4:   for each task  $\theta_k$  do
5:     if mode = ACQUISITION then
6:       Select  $\mathbf{x}_k^{(t)}$  using (5)
7:     else
8:       Sample  $\{\mathbf{x}_k^{(i)}\}_{i=1}^{N_{gen}} \sim p_\phi(\mathbf{x}|\lambda^{(t)}, \theta_k)$ 
9:       Select  $\mathbf{x}_k^{(t)}$  using (15)
10:    end if
11:    Evaluate  $F(x_k, \theta_k)$  and append to  $\mathcal{D}_k$ 
12:  end for
13:  Jointly retrain all task-aware GPs on  $\bigcup_k \mathcal{D}_k$ 
14:  if mode = GENERATIVE then
15:    Extract elite subset from  $\bigcup_k \mathcal{D}_k$ 
16:    Update  $\mathcal{G}_\phi$  using  $\mathcal{L}(\phi; \mathcal{D}^{(gen)})$ 
17:  end if
18: end for
19: return  $\{\mathcal{D}_k\}_{k=1}^K$ 

```

In this way, distribution-aware sampling is seamlessly blended with uncertainty-driven exploration. We instantiate this generative model via either Variational Autoencoder (VAE) [Kingma and Welling, 2022] or Denoising Diffusion Probabilistic Model (DDPM) [Ho *et al.*, 2020] as follows.

Conditional Variational Autoencoder The dataset $\mathcal{D}$ of elite solutions $\mathbf{x}$ is collected with conditioning vectors $\mathbf{c} = (\boldsymbol{\lambda}, \theta_k)$. We introduce a latent code z and parameterize

$$p_\phi(\mathbf{x}|\mathbf{c}) = \int p_{\phi_d}(\mathbf{x}|z, \mathbf{c}) p(z|\mathbf{c}) dz. \quad (16)$$

An encoder $q_{\phi_e}(z|\mathbf{x}, \mathbf{c})$ approximates the posterior, and a decoder $p_{\phi_d}(\mathbf{x}|z, \mathbf{c})$ models the conditional likelihood. We learn $\phi = (\phi_e, \phi_d)$ by maximizing the variational lower bound [Kingma and Welling, 2022; Sohn *et al.*, 2015] on $\mathcal{D}$. During generative solution sampling, candidate solutions $\mathbf{x}'$ can be obtained from the learned distribution as follows:

$$z \sim p(z|\mathbf{c}), \quad \mathbf{x}' \sim p_{\phi_d}(\mathbf{x}|z, \mathbf{c}).$$

Conditional Diffusion Model The dataset $\mathcal{D}$ of elite solutions $\mathbf{x}$ is collected with conditioning vectors $\mathbf{c} = (\boldsymbol{\lambda}, \theta_k)$. We model the data distribution through a forward diffusion process that gradually adds Gaussian noise over $\hat{T}$ timesteps:

$$q(\mathbf{x}_t|\mathbf{x}_{t-1}) = \mathcal{N}(\mathbf{x}_t; \sqrt{1 - \hat{\beta}_t} \mathbf{x}_{t-1}, \hat{\beta}_t \mathbf{I}), \quad (17)$$

where $\{\hat{\beta}_t\}_{t=1}^{\hat{T}}$ follows a linear noise schedule. The forward process admits a closed-form expression:

$$q(\mathbf{x}_t|\mathbf{x}_0) = \mathcal{N}(\mathbf{x}_t; \sqrt{\bar{\alpha}_t} \mathbf{x}_0, (1 - \bar{\alpha}_t) \mathbf{I}), \quad (18)$$

with $\bar{\alpha}_t = \prod_{s=1}^t (1 - \hat{\beta}_s)$. Thereafter, a neural net $\epsilon_\phi(\mathbf{x}_t, t, \mathbf{c})$ is trained to predict the noise added at timestep t given the noisy solution $\mathbf{x}_t$ and conditioning $\mathbf{c}$. We learn ϕ by minimizing the denoising objective following [Ho *et al.*, 2020]. The learned diffusion model parameterizes the conditional distribution of elite solutions as follows:

$$p_\phi(\mathbf{x}|\mathbf{c}) = \int p_\phi(\mathbf{x}_0|\mathbf{x}_{\hat{T}}, \mathbf{c}) p(\mathbf{x}_{\hat{T}}) d\mathbf{x}_{\hat{T}}, \quad p(\mathbf{x}_{\hat{T}}) = \mathcal{N}(0, \mathbf{I}), \quad (19)$$

where $p_\phi(\mathbf{x}_0|\mathbf{x}_{\hat{T}}, \mathbf{c})$ represents the reverse denoising process implemented. During the generative solution sampling phase, candidate solutions $\mathbf{x}'$ are obtained as:

$$\mathbf{x}_{\hat{T}} \sim \mathcal{N}(0, \mathbf{I}), \quad \mathbf{x}' \sim p_\phi(\mathbf{x}|\mathbf{c}).$$

3.5 Alternating Optimization Loop

The complete framework alternates between two complementary search modes as illustrated in Figure 1: acquisition-driven search via PMT-MOBO and generative solution sampling using the conditional generative model. Algorithm 1 presents the detailed workflow alternating between ACQUISITION and GENERATIVE modes. We alternate between generative sampling and acquisition-driven modes every iteration in a round-robin fashion. In GENERATIVE mode, candidates are sampled from $p_\phi(\mathbf{x}|\lambda^{(t)}, \theta_k)$ and the best is selected using the acquisition function. In ACQUISITION mode, solutions are directly selected by optimizing the task-aware acquisition function. Both update task-aware GPs after evaluation, while the generative model gets retrained during GENERATIVE iterations using $\mathcal{D}^{(gen)}$.

Elite Solution Collection We maintain datasets per task-preference pair $(\boldsymbol{\lambda}, \theta_k)$, storing top $Q\%$ solutions by scalarization score $s_\lambda(F(\mathbf{x}, \theta_k))$ [Ozaki *et al.*, 2022]. Preference vectors $\boldsymbol{\lambda}$ are generated via random sampling from Λ . Each elite solution $\mathbf{x}$ pairs with conditioning vector $\mathbf{c} = (\boldsymbol{\lambda}, \theta_k)$ for generative model training. More detailed settings can be found in Appendix C.3.

4 Results

We provide empirical experiments demonstrating our method’s ability to build effective inverse models for P-EMOPs. Our experiments verify three key claims: (1) PMT-MOBO achieves faster convergence by exploiting inter-task synergies, (2) conditional generative models further improve solution quality, and (3) the built inverse model $\mathcal{M}(\boldsymbol{\theta}, \boldsymbol{\lambda})$ can directly predict optimized solutions for unseen EMOP tasks without additional evaluations. We evaluate all methods using the Hypervolume indicator [Zitzler and Thiele, 1998] across both seen and unseen EMOP tasks.

Benchmarks. We evaluate our methods on classical synthetic benchmarks: DTLZ-1, DTLZ-2, and DTLZ-3 [Deb *et al.*, 2002], and real-world problems including lamp generation, solar rooftop generation [Makatura *et al.*, 2021], magnetic sifter design [Liu *et al.*, 2025], and UAV controller design [Makatura *et al.*, 2021]. We modify these base EMOP problems so that each benchmark defines a family of related EMOPs, where $\boldsymbol{\theta}$ represents operational conditions (e.g., physical configurations of UAVs). More problem details are provided in the Appendix B.

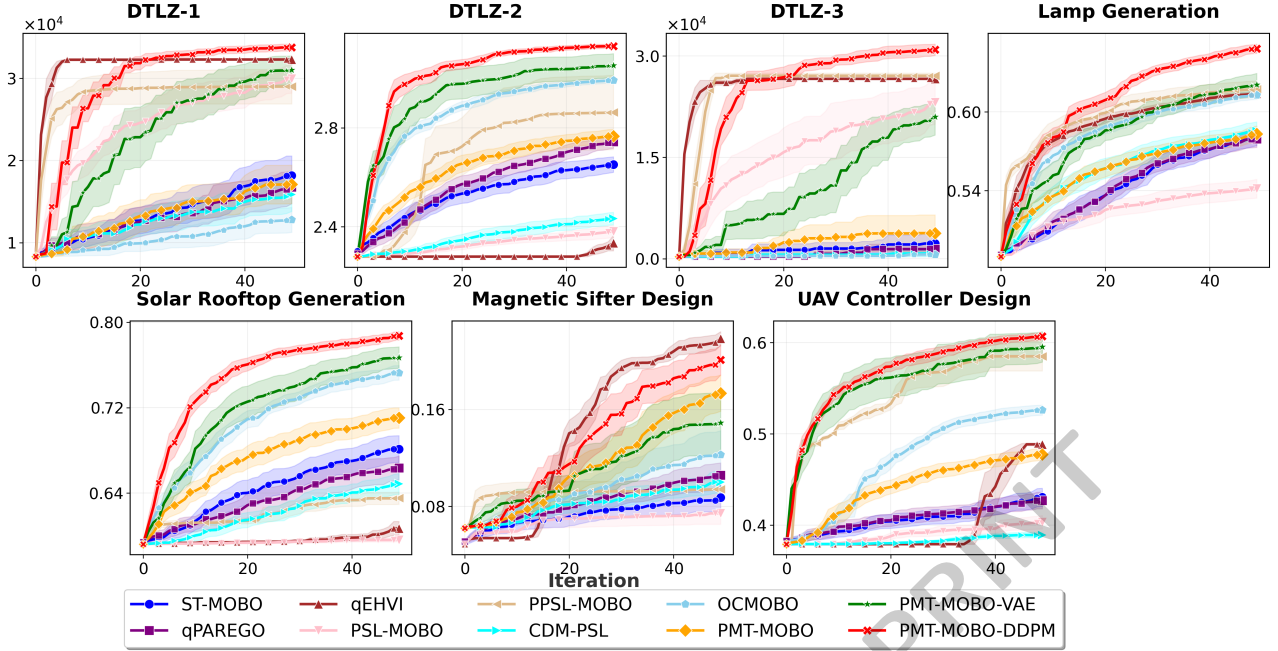


Figure 2: Average Hypervolume results evaluated on sampled EMOP tasks in synthetic and real-world benchmarks.

Method	DTLZ-1	DTLZ-2	DTLZ-3	LAMP	SOLAR	UAV
PSL-MOBO	2.82e+04 (1.78e+03) [≈]	1.90e+00 (1.30e-01) [–]	2.39e+04 (3.22e+03) ⁺	5.01e-01 (7.32e-02) [–]	3.54e-01 (4.06e-02) [–]	4.12e-01 (7.30e-03) [–]
CDM-PSL	1.92e+04 (3.85e+03) [–]	2.52e+00 (8.49e-02) [–]	1.83e+03 (1.31e+03) [–]	6.19e-01 (4.45e-02) [≈]	6.54e-01 (3.11e-02) [–]	4.26e-01 (3.09e-02) [–]
PPSL-MOBO	2.24e+04 (1.10e+02) [–]	2.98e+00 (3.56e-02) [–]	2.67e+04 (6.08e+02) ⁺	6.12e-01 (1.05e-01) [≈]	4.86e-01 (1.11e-01) [–]	5.24e-01 (1.71e-01) [–]
PMT-MOBO-VAE	2.88e+04 (2.57e+03)	3.10e+00 (2.99e-02)	1.78e+04 (7.28e+03)	6.31e-01 (7.33e-02)	7.93e-01 (1.56e-02)	6.08e-01 (1.88e-02)
PMT-MOBO-DDPM	3.47e+04 (4.85e+02)	3.13e+00 (1.41e-02)	3.20e+04 (1.98e+03)	6.70e-01 (5.82e-02)	7.65e-01 (1.30e-02)	6.05e-01 (2.44e-02)

Table 1: Inverse Model Generalization Performance on Unseen EMOP Tasks. Statistical significance is assessed using the Wilcoxon signed-rank test at the 95% confidence level. Superscripts ^{s1s2} indicate that the competitor is significantly worse (–), significantly better (+), or statistically similar (≈) relative to **PMT-MOBO-VAE** (s_1) and **PMT-MOBO-DDPM** (s_2).

Baselines. Our empirical studies include state-of-the-art single-task MOBO methods: ST-MOBO [Paria *et al.*, 2020], qPAREGO [Balandat *et al.*, 2020], qEHVI [Daulton *et al.*, 2020], PSL-MOBO [Lin *et al.*, 2022], and CDM-PSL [Li *et al.*, 2025], a diffusion model based MOBO variant. Moreover, to evaluate performance in utilizing inter-task synergies, we compare against multi-task MOBO methods: OCMOBO [Char *et al.*, 2019] and PPSL-MOBO [Cheng *et al.*, 2025b]. Finally, we evaluate variants of our approaches that serve as ablation studies: PMT-MOBO (acquisition-driven only), PMT-MOBO-VAE, and PMT-MOBO-DDPM (full framework with different generative models).

Experimental Setup. We optimize $K = 8$ parameterized EMOPs per benchmark with uniformly sampled task parameters. Each task uses 20 random initial solutions and the budget $T = 50$, repeated over $U = 20$ independent runs. We set $Q = 10$ for elite solution extraction following [Ozaki *et al.*, 2022] and β follows [Kandasamy *et al.*, 2015]. Additional implementation details are in the Appendix C.

Discussions. Figure 2 and Table 1 demonstrate our framework’s effectiveness on synthetic and real-world benchmarks. PMT-MOBO-VAE and PMT-MOBO-DDPM generally show superior or comparable results compared to state-of-the-art

single-task and multi-task MOBO baselines in discovering high-quality Pareto solutions across distinct tasks. The exception in magnetic sifter motivates us to further explore the integration of our methods with EHVI based acquisition functions so that the proposed framework can be more flexibly adapt to diverse applications. Table 1 further shows that our inverse models generalize well to unseen EMOP tasks, outperforming the single-task and multi-task inverse models constructed by PSL-MOBO and PPSL-MOBO. With task-aware kernel and the alternating optimization loop, our PMT-MOBO-DDPM’s inverse model can also outperform the diffusion model of CDM-PSL across all the problems. To evaluate this generalization capability, we extract the trained inverse models after optimization and apply them to $W = 100$ completely new tasks. For each unseen task, we query the inverse model with S uniformly sampled preference vectors ($S = 100$ for bi-objective, $S = 1000$ for tri-objective) to generate predicted solutions, then compute the Hypervolume of the resulting solution set.

Competitive Optimization Performance. Figure 2 and Table 5 in Appendix D.1 summarize the Hypervolume results across all benchmarks. Wilcoxon signed-rank tests (at 95% confidence) demonstrate that our proposed methods significantly outperform the baselines in most scenarios.

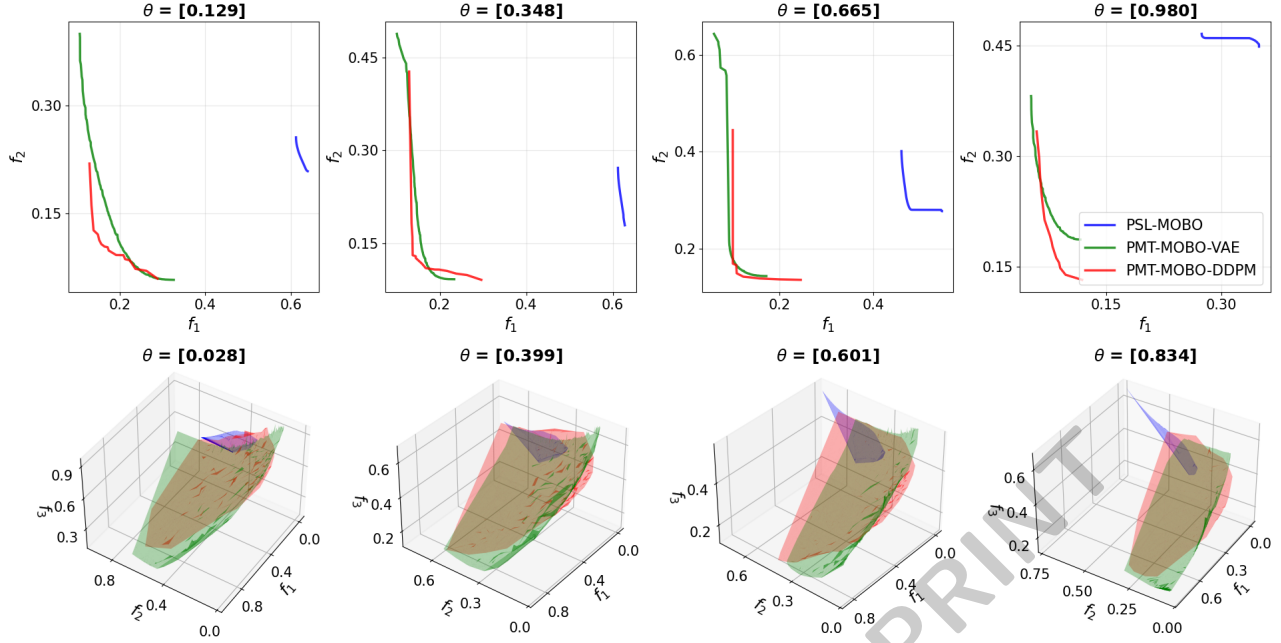


Figure 3: Approximated Pareto Front for unseen EMOP tasks in solar rooftop generation (top) and lamp generation (bottom).

Specifically, against the eight competing methods across all tasks (56 pairwise comparisons), PMT-MOBO-VAE attains a win/tie/loss record of 45/8/3, effectively dominating difficult tasks such as LAMP and SOLAR. The PMT-MOBO-DDPM further enhances optimization capability, achieving a record of 55/0/1, statistically surpassing all baselines in every benchmark except for the low-dimensional magnetic sifter design.

Effective Inverse Model Generalization. Table 1 confirms that our alternating optimization strategy yields robust inverse models for unseen tasks. PMT-MOBO-DDPM demonstrates superior generalization across the majority of benchmarks, achieving the highest Hypervolume on DTLZ-1, DTLZ-2, DTLZ-3, and LAMP generation. Notably, it statistically outperforms both PSL-MOBO and PPSL-MOBO baselines on every tested problem instance. PMT-MOBO-VAE excels on the remaining real-world engineering tasks, with the best performance on SOLAR and UAV controller design, though it struggles with the highly multimodal landscape of DTLZ-3 compared to the DDPM variant. We omit the magnetic sifter design set due to the prohibitive computational cost of evaluating hundreds of distinct tasks. To further validate these metrics, Figure 3 visualizes the approximated Pareto Fronts. Both PMT-MOBO variants discover fronts with significantly better convergence and coverage than the PSL-MOBO baseline. Consistent with Table 1, PMT-MOBO-DDPM generally provides the broadest coverage, while PMT-MOBO-VAE offers smoother approximations, illustrating a potential trade-off between the expansive generation of diffusion models and the structured latent embedding of VAEs.

Ablation Studies and Sensitivity Analysis. To validate the necessity of our alternating optimization strategy, we compared the full framework against decoupled baselines (pure acquisition search and pure generative sampling). The re-

sults, detailed in Appendix D.3, confirm that neither component is sufficient in isolation. Specifically, generative sampling enhances the diversity of candidate generation, while the alternating acquisition search lowers the variance inherent in pure generative sampling, resulting in robust convergence. Figure 2 also demonstrates the superior performance of PMT-MOBO-VAE and PMT-MOBO-DDPM compared to PMT-MOBO, showcasing the value of integrating generative modeling.

Furthermore, our sensitivity analysis on the elite ratio parameter Q reveals a critical trade-off: too large a Q reduces selective pressure, pushing the generative model to learn low-quality solutions, whereas too small a Q incurs large variance and provides insufficient data to train a well-behaved model. Empirically, $Q = 10\%$ achieves the optimal balance across both VAE and DDPM variants (see Tables 6 and 7 in Appendix D.4).

5 Conclusion

We introduce a novel parametric multi-objective Bayesian optimizer that can amortize the computational effort of solving families of expensive EMOPs into a single inverse generative model. By alternating between acquisition-driven search and generative solution sampling, our method constructs a robust mapping that captures complex Pareto geometries, significantly improving generalization to unseen tasks and optimization efficiency, as verified by synthetic and real-world benchmarks. Theoretically, we showed that incorporating task-aware GPs yields tighter regret bounds than single-task counterparts. Future work will explore extensions to combinatorial problems and advanced acquisition strategies like EHVI to further broaden the impact of amortized optimization across tasks [Wei *et al.*, 2025].

Acknowledgments

This research is partly supported by the MTI under its AI Centre of Excellence for Manufacturing (AIMfg) (Award W25MCMF014), the Centre for Frontier AI Research (CFAR) under Agency for Science, Technology and Research (A*STAR), and the College of Computing and Data Science, Nanyang Technological University. Abhishek Gupta is supported by the Ramanujan Fellowship from the Anusandhan National Research Foundation (ANRF), Government of India (Grant No. RJF/2022/000115).

References

- [Balandat *et al.*, 2020] Maximilian Balandat, Brian Karrer, Daniel Jiang, Samuel Daulton, Ben Letham, Andrew G Wilson, and Eytan Bakshy. Botorch: A framework for efficient monte-carlo bayesian optimization. In *Advances in Neural Information Processing Systems*, volume 33, pages 21524–21538. Curran Associates, Inc., 2020.
- [Bali *et al.*, 2021] Kavitesh Kumar Bali, Abhishek Gupta, Yew-Soon Ong, and Puay Siew Tan. Cognizant multitasking in multiobjective multifactorial evolution: Mo-mfea-ii. *IEEE Trans. on Cybern.*, 51(4):1784–1796, 2021.
- [Belakaria *et al.*, 2019] Syrine Belakaria, Aryan Deshwal, and Janardhan Rao Doppa. Max-value entropy search for multi-objective bayesian optimization. In *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019.
- [Branke, 2008] Jürgen Branke. *Multiobjective optimization: Interactive and evolutionary approaches*, volume 5252. Springer Science & Business Media, 2008.
- [Brookes *et al.*, 2019] David Brookes, Hahnbeom Park, and Jennifer Listgarten. Conditioning by adaptive sampling for robust design. In *Proceedings of the 36th International Conference on Machine Learning*, volume 97 of *Proceedings of Machine Learning Research*, pages 773–782. PMLR, 09–15 Jun 2019.
- [Char *et al.*, 2019] Ian Char, Youngseog Chung, Willie Neiswanger, Kirthevasan Kandasamy, Andrew Oakleigh Nelson, Mark Boyer, Egemen Kofemen, and Jeff Schneider. Offline contextual Bayesian optimization. In *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019.
- [Cheng *et al.*, 2025a] Ji Cheng, Xi Lin, Bo Xue, and Qingfu Zhang. Parametric pareto set learning: Amortizing multi-objective optimization with parameters. *IEEE Transactions on Evolutionary Computation*, pages 1–1, 2025.
- [Cheng *et al.*, 2025b] Ji Cheng, Bo Xue, and Qingfu Zhang. Parametric pareto set learning for expensive multi-objective optimization, 2025.
- [Cully *et al.*, 2015] Antoine Cully, Jeff Clune, Danesh Tarapore, and Jean-Baptiste Mouret. Robots that can adapt like animals. *Nature*, 521(7553):503–507, May 2015.
- [Daulton *et al.*, 2020] Samuel Daulton, Maximilian Balandat, and Eytan Bakshy. Differentiable expected hypervolume improvement for parallel multi-objective bayesian optimization. *Advances in Neural Information Processing Systems*, 33:9851–9864, 2020.
- [Deb *et al.*, 2002] K. Deb, L. Thiele, M. Laumanns, and E. Zitzler. Scalable multi-objective optimization test problems. In *Proceedings of the 2002 Congress on Evolutionary Computation. CEC’02*, volume 1, pages 825–830 vol.1, 2002.
- [Emmerich *et al.*, 2006] M.T.M. Emmerich, K.C. Giannakoglou, and B. Naujoks. Single- and multiobjective evolutionary optimization assisted by gaussian random field metamodels. *IEEE Trans. on Evol. Comput.*, 10(4):421–439, 2006.
- [Gupta *et al.*, 2016a] Abhishek Gupta, Yew-Soon Ong, and Liang Feng. Multifactorial evolution: toward evolutionary multitasking. *IEEE Trans. on Evol. Comput.*, 20(3):343–357, 2016.
- [Gupta *et al.*, 2016b] Abhishek Gupta, Yew-Soon Ong, Liang Feng, and Kay Chen Tan. Multiobjective multifactorial optimization in evolutionary multitasking. *IEEE Trans. on Cybern.*, 47(7):1652–1665, 2016.
- [Hernandez-Lobato *et al.*, 2016] Daniel Hernandez-Lobato, Jose Hernandez-Lobato, Amar Shah, and Ryan Adams. Predictive entropy search for multi-objective bayesian optimization. In *Proceedings of The 33rd International Conference on Machine Learning*, volume 48 of *Proceedings of Machine Learning Research*, pages 1492–1501, New York, New York, USA, 20–22 Jun 2016. PMLR.
- [Ho *et al.*, 2020] Jonathan Ho, Ajay Jain, and Pieter Abbeel. Denoising diffusion probabilistic models. In *Advances in Neural Information Processing Systems*, volume 33, pages 6840–6851. Curran Associates, Inc., 2020.
- [Kandasamy *et al.*, 2015] Kirthevasan Kandasamy, Jeff Schneider, and Barnabas Poczos. High dimensional bayesian optimisation and bandits via additive models. In *Proceedings of the 32nd International Conference on Machine Learning*, volume 37 of *Proceedings of Machine Learning Research*, pages 295–304, Lille, France, 07–09 Jul 2015. PMLR.
- [Kim *et al.*, 2021] Yeonju Kim, Zherong Pan, and Kris Hauser. Mo-bbo: Multi-objective bilevel bayesian optimization for robot and behavior co-design. In *2021 IEEE International Conference on Robotics and Automation (ICRA)*, pages 9877–9883, 2021.
- [Kingma and Welling, 2022] Diederik P Kingma and Max Welling. Auto-encoding variational bayes, 2022.
- [Knowles, 2006] J. Knowles. Parego: a hybrid algorithm with on-line landscape approximation for expensive multiobjective optimization problems. *IEEE Trans. on Evol. Comput.*, 10(1):50–66, 2006.
- [Li *et al.*, 2024] Beichen Li, Bolei Deng, Wan Shou, Tae-Hyun Oh, Yuanming Hu, Yiyue Luo, Liang Shi, and Wojciech Matusik. Computational discovery of microstructured composites with optimal stiffness-toughness trade-offs. *Science Advances*, 10(5):eadk4284, 2024.

- [Li *et al.*, 2025] Bingdong Li, Zixiang Di, Yongfan Lu, Hong Qian, Feng Wang, Peng Yang, Ke Tang, and Aimin Zhou. Expensive multi-objective bayesian optimization based on diffusion models. *Proceedings of the AAAI Conference on Artificial Intelligence*, 39(25):27063–27071, Apr. 2025.
- [Lin *et al.*, 2022] Xi Lin, Zhiyuan Yang, Xiaoyuan Zhang, and Qingfu Zhang. Pareto set learning for expensive multi-objective optimization. In *Advances in Neural Information Processing Systems*, volume 35, pages 19231–19247. Curran Associates, Inc., 2022.
- [Liu *et al.*, 2025] Jiao Liu, Abhishek Gupta, Chinchun Ooi, and Yew-Soon Ong. Extremo: Transfer evolutionary multiobjective optimization with proof of faster convergence. *IEEE Trans. on Evol. Comput.*, 29(1):102–116, 2025.
- [Makatura *et al.*, 2021] Liane Makatura, Minghao Guo, Adriana Schulz, Justin Solomon, and Wojciech Matusik. Pareto gamuts: exploring optimal designs across varying contexts. *ACM Trans. on Graph.*, 40(4), July 2021.
- [Miettinen, 1999] Kaisa Miettinen. *Nonlinear multiobjective optimization*, volume 12. Springer Science & Business Media, 1999.
- [Min *et al.*, 2019] Alan Tan Wei Min, Yew-Soon Ong, Abhishek Gupta, and Chi-Keong Goh. Multiproblem surrogates: Transfer evolutionary multiobjective optimization of computationally expensive problems. *IEEE Trans. on Evol. Comput.*, 23(1):15–28, 2019.
- [Nickelson *et al.*, 2023] Anna Nickelson, Kagan Tumer, and William D. Smart. Contextual multi-objective path planning. In *2023 IEEE International Conference on Robotics and Automation (ICRA)*, pages 10240–10246, 2023.
- [Ozaki *et al.*, 2022] Yoshihiko Ozaki, Yuki Tanigaki, Shuhei Watanabe, Masahiro Nomura, and Masaki Onishi. Multiobjective tree-structured parzen estimator. *J. Artif. Int. Res.*, 73, May 2022.
- [Paria *et al.*, 2020] Biswajit Paria, Kirthevasan Kandasamy, and Barnabás Póczos. A flexible framework for multi-objective bayesian optimization using random scalarizations. In *Proceedings of The 35th Uncertainty in Artificial Intelligence Conference*, volume 115 of *Proceedings of Machine Learning Research*, pages 766–776. PMLR, 22–25 Jul 2020.
- [Picheny, 2015] Victor Picheny. Multiobjective optimization using gaussian process emulators via stepwise uncertainty reduction. *Statistics and Computing*, 25(6):1265–1280, 2015.
- [Pierrot *et al.*, 2022] Thomas Pierrot, Guillaume Richard, Karim Beguir, and Antoine Cully. Multi-objective quality diversity optimization. In *Proceedings of the Genetic and Evolutionary Computation Conference, GECCO '22*, page 139–147, New York, NY, USA, 2022. Association for Computing Machinery.
- [Sohn *et al.*, 2015] Kihyuk Sohn, Honglak Lee, and Xinchen Yan. Learning structured output representation using deep conditional generative models. In *Advances in Neural Information Processing Systems*, volume 28. Curran Associates, Inc., 2015.
- [Srinivas *et al.*, 2012] Niranjan Srinivas, Andreas Krause, Sham M. Kakade, and Matthias W. Seeger. Information-theoretic regret bounds for Gaussian process optimization in the bandit setting. *IEEE Trans. on Information Theory*, 58(5):3250–3265, 2012.
- [Tu *et al.*, 2022] Ben Tu, Axel Gandy, Nikolas Kantas, and Behrang Shafei. Joint entropy search for multi-objective bayesian optimization. In *Advances in Neural Information Processing Systems*, volume 35, pages 9922–9938. Curran Associates, Inc., 2022.
- [Turchetta *et al.*, 2020] Matteo Turchetta, Andreas Krause, and Sebastian Trimpe. Robust model-free reinforcement learning with multi-objective bayesian optimization. In *2020 IEEE International Conference on Robotics and Automation (ICRA)*, pages 10702–10708, 2020.
- [Wei *et al.*, 2022] Tingyang Wei, Shibin Wang, Jinghui Zhong, Dong Liu, and Jun Zhang. A review on evolutionary multitask optimization: Trends and challenges. *IEEE Trans. on Evol. Comput.*, 26(5):941–960, 2022.
- [Wei *et al.*, 2024] Tingyang Wei, Jiao Liu, Abhishek Gupta, Puay Siew Tan, and Yew-Soon Ong. Bayesian forward-inverse transfer for multiobjective optimization. In *Parallel Problem Solving from Nature – PPSN XVIII*, pages 135–152, Cham, 2024. Springer Nature Switzerland.
- [Wei *et al.*, 2025] Tingyang Wei, Jiao Liu, Abhishek Gupta, Puay Siew Tan, and Yew-Soon Ong. (θ_l, θ_u) -parametric multi-task optimization: Joint search in solution and infinite task spaces. *IEEE Trans. on Evol. Comput.*, pages 1–1, 2025.
- [Williams and Rasmussen, 2006] C. K. Williams and C. E. Rasmussen. *Gaussian Processes for Machine Learning*. MIT Press, 2006.
- [Zhang and Golovin, 2020] Richard Zhang and Daniel Golovin. Random hypervolume scalarizations for provable multi-objective black box optimization. In *Proceedings of the 37th International Conference on Machine Learning*, volume 119 of *Proceedings of Machine Learning Research*, pages 11096–11105. PMLR, 13–18 Jul 2020.
- [Zhang *et al.*, 2022] Dinghuai Zhang, Jie Fu, Yoshua Bengio, and Aaron Courville. Unifying likelihood-free inference with black-box sequence design and beyond. In *International Conference on Learning Representations*, 2022.
- [Zhu *et al.*, 2023] Yiheng Zhu, Jialu Wu, Chaowen Hu, Jiahuan Yan, kim hsieh, Tingjun Hou, and Jian Wu. Sample-efficient multi-objective molecular optimization with gflownets. In *Advances in Neural Information Processing Systems*, volume 36, pages 79667–79684. Curran Associates, Inc., 2023.
- [Zitzler and Thiele, 1998] Eckart Zitzler and Lothar Thiele. Multiobjective optimization using evolutionary algorithms — a comparative case study. In *Parallel Problem Solving from Nature — PPSN V*, pages 292–301, Berlin, Heidelberg, 1998. Springer Berlin Heidelberg.