

Human-in-the-Loop Meta Bayesian Optimization for Fusion Energy and Scientific Applications

Ricardo Luna Gutierrez¹, Sahand Ghorbanpour¹, Rahman Ejaz², Varchas Gopalaswamy²
Riccardo Betti², Vineet Gundecha¹, Aarne Lees² and Soumyendu Sarkar¹*

¹Hewlett Packard Enterprise

²University of Rochester

{rluna, sahand.ghorbanpour, vineet.gundecha, soumyendu.sarkar}@hpe.com, {reja, vgop, betti, alees}@lle.rochester.edu

Abstract

Inertial Confinement Fusion (ICF) holds transformative promise for sustainable, near-limitless clean energy, yet remains constrained by prohibitively high costs and limited experimental opportunities. This paper presents Human-in-the-Loop Meta Bayesian Optimization (HL-MBO), a framework that integrates expert knowledge with few-shot, uncertainty-aware machine learning to accelerate discovery in data-scarce, high-stakes scientific domains. HL-MBO introduces a meta-learned surrogate model with an expert-informed acquisition function to recommend candidate experiments. To foster trust and enable informed decisions, HL-MBO also provides interpretable explanations of its suggestions. We show that HL-MBO outperforms current BO methods for ICF energy yield optimization, as well as benchmarks in molecular optimization and critical-temperature maximization for superconducting materials.

1 Introduction

The global community faces an existential challenge: the need to meet increasing energy demands while rapidly decarbonizing to mitigate the effects of climate change. Inertial Confinement Fusion (ICF) represents a transformative solution to this crisis, offering a path to a near-limitless, carbon-neutral power source without the long-lived radioactive waste or risks associated with traditional nuclear fission.

However, ICF experimental facilities are among the most complex and expensive machines ever built. Consequently, experiments are extremely costly and researchers are limited to only a few shots annually. For this reason, there is an urgent need for highly efficient optimization methods that can maximize the insights gained from every single experiment, effectively accelerating the timeline to a clean energy future.

Bayesian Optimization (BO) has shown promise in optimizing expensive black-box functions, using a surrogate model and acquisition function (AF) to guide experimental sampling [Feurer *et al.*, 2015; Snoek *et al.*, 2012; Wang *et*

al., 2023]. However, BO faces skepticism due to its black-box nature and insufficient sample efficiency for complex tasks like ICF [Betti and Hurricane, 2016; Lees *et al.*, 2021; Gopalaswamy *et al.*, 2019; Gopalaswamy *et al.*, 2024].

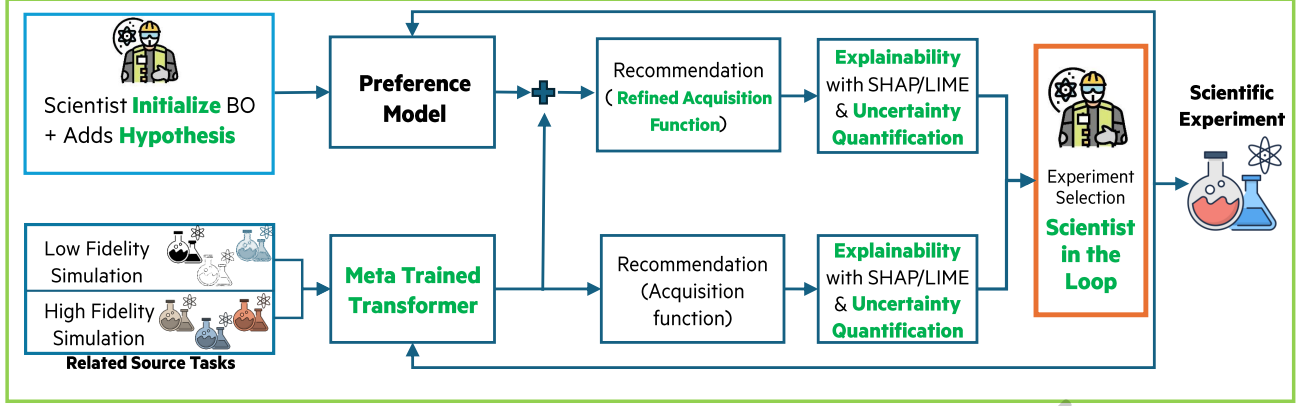
Recent research incorporates human expertise into BO through human-in-the-loop methods, preference learning and explainability [Colella *et al.*, 2020; A V *et al.*, 2022; Gupta *et al.*, 2023; Hvarfner *et al.*, 2022; Adachi *et al.*, 2024]. While promising, these methods are limited to single-task optimization, requiring the process to restart for each new task. To address new optimization tasks’ sample efficiency and leverage accumulated knowledge, Meta-Learning within Bayesian Optimization (Meta-BO) has garnered increasing interest [Bai *et al.*, 2023].

In Meta-BO, the surrogate model and/or AF are trained using Meta-Learning techniques on a set of source tasks, enabling rapid adaptation to new, unseen problems. Despite the promising transfer learning capabilities of Meta-BO, fundamental challenges remain unaddressed for expensive scientific experiments. The trustworthiness and explainability of Meta-BO decisions, necessary for adoption in high-stakes domains, have received little attention. Furthermore, the principled integration of valuable expert knowledge to potentially amplify optimization performance within Meta-BO remains an open problem.

To address these limitations, we introduce **HL-MBO**, an explainable human-in-the-loop framework for Meta-Bayesian Optimization. Developed through a direct partnership between AI researchers and ICF physicists, HL-MBO integrates expert knowledge with few-shot, uncertainty-aware machine learning to accelerate discovery in data-scarce, high-stakes scientific domains. Built with interpretability in mind, HL-MBO provides insights into its decision-making process, helping to foster trust and expert engagement. An overview of the framework is presented in Figure 1. We summarize our main contributions as follows:

- We propose HL-MBO, the first framework that integrates human-in-the-loop and Meta-BO. We demonstrate how this framework achieves superior performance against state-of-the-art (SOTA) human-aided BO and Meta-BO methods in high-stakes applications such as ICF.
- We propose a new formulation for an expert-informed

*Corresponding author.



HL-MBO's Refinements to Bayesian Optimization for Experimental Efficiency with Human Feedback



Standard Bayesian Optimization

Figure 1: Refinements proposed in HL-MBO over standard Bayesian Optimization. HL-MBO trains a meta-surrogate using related source tasks and integrates a preference model to incorporate expert knowledge. During online black-box optimization the meta-surrogate and preference model are used in conjunction with an AF to optimize the target function (scientific experiment), providing explanations to support decision-making.

AF, which integrates both preference learning and experts' hypotheses, and study the effect that human expertise can have over the optimization process.

- We implement a comprehensive explainable framework for Meta-BO, leveraging uncertainty quantification, Shapley values and LIME to provide users with insights into the model's decision-making process and foster trust.

2 Background

2.1 Challenges of Yield Optimization for ICF

Inertial Confinement Fusion (ICF) achieves nuclear fusion by using lasers to compress and heat a deuterium-tritium fuel pellet to extreme temperatures and pressures, overcoming electrostatic repulsion between nuclei and releasing energy [Betti and Hurricane, 2016]. Energy yield is highly sensitive to the **laser pulse shape (LP)**, which drives ignition within 3 ns and is controlled by five critical parameters.

ICF experiments are extremely costly due to the sophisticated laser systems and harsh operating conditions. Moreover, access to ICF-capable facilities is limited, with only 5–10 shots typically conducted per day and just a few experimental days available each year.

Although simulators can generate data to train surrogate models or warm-start optimization, discrepancies between simulated and real outcomes often arise due to the system's complexity. As a result, there is a strong need for methods that can leverage simulation data but still be able to quickly adapt on the fly to compensate for these discrepancies.

2.2 Meta Bayesian Optimization

Bayesian Optimization (BO) is a technique that considers an objective function f as a black-box and aims to find the global optimum by sequentially selecting points to evaluate based on the information obtained from previous evaluations [Garnett, 2023]. Meta-Bayesian Optimization (Meta-BO) approaches aim to improve the optimization of new unseen target black-box functions by leveraging knowledge from a set of N related source tasks (functions) $\mathcal{F}$ [Maraval *et al.*, 2023]. We assume the model has access to this knowledge as datasets $\mathcal{D}_1, \dots, \mathcal{D}_N$. Each dataset $\mathcal{D}_n$ consist of e_n evaluations of $f_n(x) \in \mathcal{F}$ for all $n \in \{1, \dots, N\}$, such as $\mathcal{D}_n = \{(x_n^i, y_n^i)\}_{i=1}^{e_n}$, where $y_n^i = f_n(x_n^i)$.

In single-task scenarios, GPs have traditionally been the primary tool for constructing surrogate models. However, in the context of Meta-BO, recent studies highlight the superiority of uncertainty-aware meta-learning approaches, such as Neural Processes (NPs) [Garnelo *et al.*, 2018]. NPs are a class of models that combine the expressive function approximation capabilities of neural networks with the stochastic properties of GPs. Trained following a meta-learning framework, NPs enable rapid adaptation to new functions at test time [Jha *et al.*, 2023]. Given a dataset of input-output context samples $\mathcal{D}_c$ and an unlabeled set of l target locations $\mathcal{X}_T = \{x^i\}_{i=1}^l$ in which we would like to make predictions, NPs generate a distribution $p(\cdot | \mathcal{X}_T, \mathcal{D}_c)$ that approximates the true posterior over the target labels $\mathcal{Y}_T$.

3 HL-MBO

HL-MBO is an explainable meta-learning human-in-the-loop framework that improves the efficiency and trustworthiness of black-box optimization. HL-MBO accelerates the optimization of high-stakes experiments by synthesizing historical task data with real-time expert insights, ensuring rapid convergence even in data-scarce environments.

The methodology presented here was developed through direct and iterative discussions between AI researchers and ICF physicists to ensure the framework addresses real-world experimental problems. Moreover, central to our approach is an expert-guided AF where the ICF experts acted as domain leads.

3.1 Preference Learning

Traditional Meta-BO approaches overlook the valuable insight that human experts can provide. In resource-intensive scientific applications such as ICF, effectively leveraging this expert knowledge is critical to increase experimental efficiency. To address this problem, HL-MBO incorporates a preference model built on binary preference learning [Adachi *et al.*, 2024; Bradley and Terry, 1952; Lun Chau *et al.*, 2022a]. We construct a preference dataset by presenting experts with pairs of candidate points, $\{x_1, x_2\}$, and asking them which point is more likely to be a better candidate for experimental evaluation. Specifically, experts label each pair based on their preferences. This process is repeated M times, resulting in a dataset $D_{\text{pref}} = \{x_1^i, x_2^i, y_{\text{pref}}^i\}_{i=1}^M$, where y_{pref} is 1 if x_1 is favored over x_2 , and 0 otherwise. D_{pref} is then used to build a binary preference function g , which follows a likelihood model: $\mathbb{P}(y_{\text{pref}} | x_1, x_2) = S(y_{\text{pref}}; g(x_1, x_2))$. Where $S(y_{\text{pref}}; z) := z^{y_{\text{pref}}} (1 - z)^{1 - y_{\text{pref}}}$ represents a Bernoulli likelihood and $g(x_1, x_2)$ represents the user’s preference level for x_1 compared to x_2 . To model function g , we employ Dirichlet-based GPs [Adachi *et al.*, 2024; Milios *et al.*, 2018] combined with skew-symmetric data augmentation [Lun Chau *et al.*, 2022b]. Following this approach allows us to estimate a preference model π as GPs that generate a distribution $p_{\pi}(\cdot | \mathcal{X}_T, D_c)$ that approximates an expert aligned posterior given a set of target locations $\mathcal{X}_T$ and context samples D_c .

3.2 Experts’ Hypotheses

Typically, the set of pair candidate points $\{x_1, x_2\}$ presented to human experts is randomly selected [Adachi *et al.*, 2024]. However, this random process is impractical and faces resistance from human experts in complex tasks due to the large number of samples that require labeling. A more efficient approach would be to reduce the sampling space to high-quality zones and reduce the number of samples requiring labeling by improving sample quality. To achieve that, we explore the concept of experts’ hypotheses [Cisse *et al.*, 2024] in the scope of preference learning. Experts’ hypotheses can be defined as subsets $\mathcal{H} \subset \mathcal{X}$ of the overall search space, where experts identify a range of values considered promising. For instance, in a 2D problem (two features), a hypothesis might be specified as $\mathcal{H} = \{\mathbf{x} \in \mathcal{X} \mid 0.1 < \mathbf{x}_1 < 0.7, 0.5 < \mathbf{x}_2 < 0.9\}$, where $\mathbf{x}_1$ and

$\mathbf{x}_2$ refer to features 1 and 2, respectively. Instead of performing a selection of pairs across the entire search space, we constrain the selection to points that fall within these defined hypotheses. This constraint applies exclusively during D_{pref} construction, without interfering with online black-box optimization’s flexibility to explore beyond hypothesis boundaries. This minimizes influence on the optimization process while improving preference data collection efficiency and quality.

3.3 Meta Training

A key shortcoming of prior human-aided BO techniques is their inability to exploit the rich information contained within historical experimental data or simulations [Gupta *et al.*, 2023; Hvarfner *et al.*, 2022; Adachi *et al.*, 2024; Souza *et al.*, 2021; Ramachandran *et al.*, 2020; Chakraborty *et al.*, 2024; Akata *et al.*, 2020; Venkatesh *et al.*, 2022]. In this work, in parallel with constructing the preference model, we train a Meta-BO model using a set of related source tasks and TNPs. TNPs are particularly well suited for this setting as they provide a robust method for uncertainty-aware meta-learning by leveraging sequence modeling, enabling sample efficient optimization. We consider each set $\mathcal{D}_n \in \{\mathcal{D}_1, \dots, \mathcal{D}_N\}$ containing e_n evaluations as an ordered training sequence. A random subset of each training sequence is selected as the context set, comprising context pairs $(x_{1:m}, y_{1:m})$, which serve as few-shot conditioning for the transformer model. Subsequently, the TNPs autoregressively model the predictive likelihood of the remaining $(e_n - m)$ target points and optimize the objective:

$$\mathcal{L}(\theta) = \mathbb{E}_{x_{1:e_n}, y_{1:e_n}, m} [\log p_{\theta}(y_{m+1:e_n} | x_{1:e_n}, y_{1:m})]. \quad (1)$$

During each training iteration, the indices n (dataset) and m (context-target split) are picked uniformly to define the division of the samples into context and target points. For the transformer architecture in TNPs, we adopt the GPT-style model described in [2022].

3.4 Acquisition Function

In expert-aided BO, we aim to find the point in the parameter space that leads to the best performance while considering experts’ preferences. CoExBO [Adachi *et al.*, 2024] established a principled approach by combining GP-based surrogate model $\mathcal{S}$ and preference model π through a UCB acquisition function $\alpha_{\mathcal{S}, \pi}$ as:

$$\alpha_{\mathcal{S}, \pi}(x) := \mu_{\mathcal{S}, \pi}(x) + \sigma_{\mathcal{S}, \pi}(x) \quad (2)$$

$$\mu_{\mathcal{S}, \pi}(x) := \frac{\sigma_{\mathcal{S}, \pi}^2(x)}{\mathcal{J}_{\pi}^2(x)} \mu_{\pi}(x) + \frac{\sigma_{\mathcal{S}, \pi}^2(x)}{\sigma_{\mathcal{S}}^2(x)} \mu_{\mathcal{S}}(x) \quad (3)$$

$$\sigma_{\mathcal{S}, \pi}^2(x) := \frac{\mathcal{J}_{\pi}^2(x) \sigma_{\mathcal{S}}^2(x)}{\mathcal{J}_{\pi}^2(x) + \sigma_{\mathcal{S}}^2(x)} \quad (4)$$

$$\mathcal{J}_{\pi}^2(x) := \sigma_{\pi}^2(x) + \gamma t^2 \sigma_{\mathcal{S}}^2(x), \quad (5)$$

where $\mu_{\mathcal{S}}, \sigma_{\mathcal{S}}^2$ and $\mu_{\pi}, \sigma_{\pi}^2$ are the predicted mean and variance (uncertainty) of the surrogate model $\mathcal{S}$ and the preference model π , respectively. The parameter γ is a decay factor

and t represents the current time step in the optimization process. γ ensures that the information provided by π decays, so the impact of π is high at the start of the optimization but allows $\mathcal{S}$ to take over at later stages. As demonstrated by [Adachi *et al.*, 2024], this formulation offers a no-harm guarantee, where even in a worst-case scenario the preference-influenced AF achieves convergence rates comparable to or better than those of standard BO AFs.

While CoExBO demonstrated strong single-task performance, it cannot leverage knowledge transfer from prior experiments or simulations, a crucial limitation for expensive scientific optimization where historical data from related tasks is available. Additionally, empirical studies have shown Expected Improvement (EI) often outperforms UCB in practical optimization scenarios [Merrill *et al.*, 2021].

To address these limitations, we introduce a meta-learned TNP surrogate $\mathcal{S}$ to replace the standard GP-based surrogate used in traditional approaches. This enables knowledge transfer from related source tasks. Second, we formulate a new human-aided AF based on Expected Improvement to better exploit the improvement-based criterion. Our EI-based formulation $\alpha_{\mathcal{S},\pi}(x)$ is defined as:

$$\alpha_{\mathcal{S},\pi}(x) = (\mu_{\mathcal{S},\pi}(x) - f(x^+) - \xi)\Phi(Z) + \sigma_{\mathcal{S},\pi}(x)\phi(Z) \quad (6)$$

$$Z = \frac{\mu_{\mathcal{S},\pi}(x) - f(x^+) - \xi}{\sigma_{\mathcal{S},\pi}(x)}, \quad (7)$$

where $f(x^+)$ is the best observed value so far, $\Phi(\cdot)$ is the cumulative distribution function (CDF) and $\phi(\cdot)$ is the probability density function (PDF) of Z [Tacq, 2010]. The parameter ξ is a small positive hyperparameter that encourages exploration. Crucially, we retain the decay-based uncertainty combination structure from Equations 3-4 (see Technical Appendix C), which preserves the intended no-harm behavior of progressively reducing preference influence over time. We then replace the underlying components $\mu_{\mathcal{S}}$ and $\sigma_{\mathcal{S}}^2$ with the predictive mean and variance derived from our meta-trained TNP, enabling meta-learning capabilities within the same acquisition framework.

To provide users with flexible optimization strategies, HL-MBO also offers a preference-free AF $\alpha_{\mathcal{S}}(x)$ that applies standard Expected Improvement over the meta-surrogate alone ($\mu_{\mathcal{S}}$ and $\sigma_{\mathcal{S}}^2$), enabling pure data-driven optimization.

3.5 Explainability

To assist human experts in making informed decisions during the optimization process, we provide explanations of the model predictions. First, we present visual representations of the TNPs’ predicted experimental output (TNPs Mean) and its uncertainty levels, along with the AF values for the proposed HL-MBO candidates. Presenting uncertainty is particularly important because it highlights the confidence level of the model in its predictions. Figure 2(a) shows an example of these representations for general applications. We include representations for the specific task of ICF energy yield optimization in Section B of the Technical Appendix.

Moreover, we incorporate two widely used explainability methods: Shapley values (SHAP) [Shapley, 1952; Lundberg and Lee, 2017] and Local Interpretable Model-Agnostic Explanations (LIME) [Rodemann *et al.*, 2024; Ribeiro *et al.*, 2016].

Figure 2 (b) provides an example of the explanation generated by SHAP and LIME for the user. "X1" and "X2" denote the two candidate points proposed by $\alpha_{\mathcal{S}}$ and $\alpha_{\mathcal{S},\pi}$, respectively. For each candidate point, the bars represent the Shapley or LIME attribution value of the features in relation to their respective AF score, TNPs’ predicted mean and TNPs’ uncertainty at proposed location. The divergence in feature attributions across methods offers human experts complementary perspectives on the model’s reasoning. Providing these diverse explanations empowers experts to identify truly influential features with greater confidence, facilitating precise and targeted interventions in the optimization process. Furthermore, this multi-faceted explainability enables experts to critically analyze and potentially discard candidate solutions when the observed feature importance deviates substantially from their domain knowledge or expected behavior, fostering a more robust and reliable optimization loop. In our collaboration, ICF experts utilized these attributions to verify if suggested laser pulse parameters aligned with known physical sensitivities, allowing them to confidently override or accept model-predicted "shots" based on domain intuition.

3.6 Explainable Expert Guided Optimization

Given a trained meta-surrogate model $\mathcal{S}$, a preference model π , a target function f_{target} , and a context dataset D_c which contains a set of $\mathcal{I}$ initial input-output samples which have been evaluated on f_{target} , we start the optimization process by generating a pair of candidate points x_1 and x_2 where:

$$x_1 = \operatorname{argmax}_{x \in X} \alpha_{\mathcal{S}}(x) \quad x_2 = \operatorname{argmax}_{x \in X} \alpha_{\mathcal{S},\pi}(x)$$

Given these candidates, we provide explanations about them and the expert selects between both options. Subsequently, this newly selected input-output pair $\{x_i, y_i\}$ is appended to the dataset D_c , $D_c \leftarrow D_c \cup \{x_i, y_i\}$. D_c is used as context for the surrogate to adapt its predicted posterior over the current function f_{target} and a new pair of candidate points is generated. Since we are using TNPs, this adaptation does not require any gradient updates and only a forward pass with the new context D_c as input is necessary. This iterative process continues until a budget of points to evaluate B is reached. HL-MBO’s overall optimization process is shown in Algorithm 1.

4 Experiments

We evaluate HL-MBO across diverse domains to address three research questions: (1) performance on optimizing black-box tasks, (2) the influence of expert preferences, and (3) the impact of experts’ hypotheses.

For ICF, we evaluate HL-MBO in the task of energy yield optimization, where we optimize 5 critical parameters of a LP in order to maximize energy generation.

Moreover, to demonstrate generalization across domains we make use of the HPO-B benchmark [Pineda Arango *et al.*, 2021] a popular hyperparameter optimization benchmark

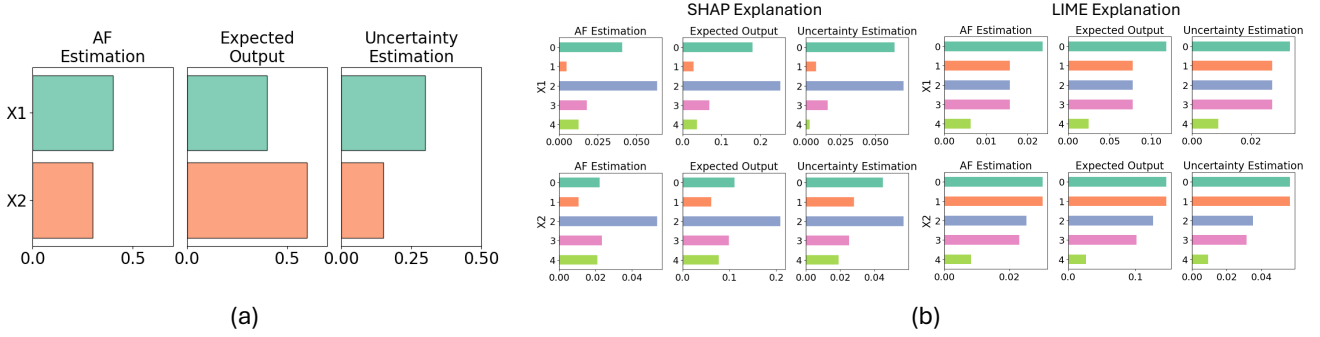


Figure 2: (a) Visual representation of the expected experimental output (TNPs’ Mean) and the associated uncertainty of the prediction for both candidate points proposed. (b) Feature attribution using SHAP and LIME on the ICF optimization task over 5 features, for the two candidate points proposed by HL-MBO. Each bar represents contribution with respect to the BO metrics (AF, Expected Output, Uncertainty Estimation).

Algorithm 1 HL-MBO algorithm

Input: Meta-surrogate $\mathcal{S}$, preference model π , target function f_{target} , initial samples $\mathcal{I}$, budget B .

- 1: Initialize $D_c \leftarrow \mathcal{I}$ input-output pairs.
- 2: **for** $i = 1$ to B **do**
- 3: Obtain posterior $p(\cdot | D_c)$ from $\mathcal{S}$ and π .
- 4: Propose candidates $\{x_1, x_2\}$ using $\alpha_{\mathcal{S}}$ and $\alpha_{\mathcal{S}, \pi}$.
- 5: Present candidates with explanations to expert.
- 6: Expert selects $x_p \in \{x_1, x_2\}$, evaluate $f_{\text{target}}(x_p) = y_p$.
- 7: Update $D_c \leftarrow D_c \cup \{x_p, y_p\}$.
- 8: **end for**

used to evaluate Meta-BO approaches. We selected 3 different datasets for this evaluation; Ranger (6D), SVM (8D), and XGBoost (16D), covering problems from 6 to 16 input dimensions (hyperparameters to optimize).

Additionally, we evaluate HL-MBO in the real-world high dimensional tasks of critical temperature maximization for **Superconductor** (86D) materials and practical molecular optimization (**PMO**) (2048D) [Gao *et al.*, 2022; Trabucco *et al.*, 2022].

Additional experimental details and dataset descriptions can be found in the Technical Appendix.

4.1 HL-MBO Performance

To answer the first question, we compare HL-MBO against CoExBO and π BO [Hvarfner *et al.*, 2022; Adachi *et al.*, 2024], SOTA approaches for human-in-the-loop BO, NAP [Maraval *et al.*, 2023], the SOTA for Meta-BO, and a standard TNP+EI [Nguyen and Grover, 2022]. Moreover, to evaluate the isolated impact of our proposed AF, we design a new baseline, a meta-learning version of CoExBO (MCoExBO). MCoExBO uses the same **UCB-based** AF proposed in CoExBO; however, the single-task GPs are replaced by TNPs. Both HL-MBO and MCoExBO were meta-trained in the same training sets. Similarly, for fairness of comparison, the GPs used by CoExBO were initialized with training

data. For all cases, the methods were evaluated on a set of test tasks that have not been seen during training.

For **ICF**, the ICF experts actively formulated hypotheses and selected from the candidate experiment pairs. Two ICF physicists provided 100 pairwise labels in about one hour, and two ICF experts selected between the two candidates during online optimization. For the **HPO-B, superconductor and PMO** tasks, we simulate human decision making for selecting candidate points proposed by the AFs $\alpha_{\mathcal{S}}$, $\alpha_{\mathcal{S}, \pi}$ and constructing D_{pref} , by employing the synthetic human selection method outlined in [2024]. In this simulation, user selection is modeled as $f_{\text{human}}(x_1) > f_{\text{human}}(x_2)$, where $f_{\text{human}}(x) := f_{\text{target}}(x) + \epsilon_{\text{pref}}$ and $\epsilon_{\text{pref}} \sim \mathcal{N}(x; 0, \sigma_{\text{pref}}^2)$. For hypothesis selection, we simulate expert knowledge by dividing the search space into slices of size $M = 100$ and choosing the slice that contains the optima. For our experiments $\sigma_{\text{pref}}^2 = 0.1$.

We set $\mathcal{I} = 1$ for all tasks to emulate a real-world setting where sampling is costly. We evaluated all methods using 25 random seeds for the HPO-B and superconductor tasks and 5 random seeds for the ICF and PMO tasks.

As is common in BO literature, we evaluate performance in terms of simple regret, defined as $R_t = f(x^*) - f(x_t^+)$, where x_t^+ is the best input point found up to and including time step t [Volpp *et al.*, 2020; Maraval *et al.*, 2023]. Given that, on average, a maximum of 10 shots can be fired during an ICF experiment day, we set our optimization budget $B = 10$ for the ICF evaluations. Figure 3 presents the results of this evaluation. HL-MBO significantly outperforms all baselines, achieving the lowest regret with the fewest optimization steps. For ICF, HL-MBO reaches the optimum in 4 samples compared to 9 for the other methods, achieving a 44% improvement in sample efficiency.

As shown by MCoExBO, replacing the GP-based surrogate with TNPs improves optimization performance. However, using TNPs alone is insufficient, as MCoExBO does not match the performance of HL-MBO. This highlights the effectiveness of our AF formulation in driving superior optimization outcomes.

A more granular analysis of HL-MBO’s success in the ICF

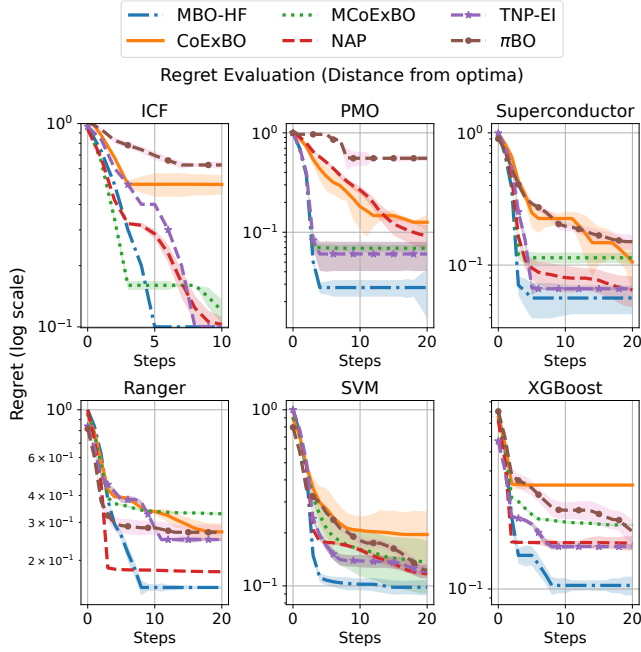


Figure 3: Results of the comparison between HL-MBO and the base-lines. The lower the better. The shaded areas represent a ± 1 standard deviations. HL-MBO outperforms the baselines on all 6 benchmarks. Moreover, we can observe that HL-MBO is able to achieve the best sample efficiency in the scientific applications of ICF, PMO and Superconductor.

domain is provided in Technical Appendix B.

4.2 Robustness Analysis

Hypothesis Evaluation In this ablation, we explore the influence of experts’ hypotheses on the performance of HL-MBO. For this experiment, we assume access to the target function f_{target} and consider three experimental settings: Expert Hypothesis (EH), Random Hypothesis (RH), and Adversarial Hypothesis (AH). To construct these hypotheses, we partition the search space $\mathcal{X}_T$ in K slices of size M . For each slice $k \in K$, we evaluate all input values $x_k^i \in \mathcal{X}_T^k$ and generate datasets $D_k = \{(x_k^i, y_k^i)\}_{i=1}^M$. The Expert Hypothesis represents an optimal scenario. We construct it by selecting the slices that contain the optima. On the other hand, the Adversarial Hypothesis simulates the worst-case scenario, we choose the slice k where the cumulative sum $V_k = \sum_{i=0}^M y_k^i$ is the lowest. For the Random Hypothesis, we uniformly sample from $\mathcal{X}_T$. We also added CoExBO and NAP for comparison.

Figure 4 presents the results of this ablation study. Despite the Adversarial Hypothesis impacting HL-MBO’s performance, it still outperforms CoExBO and matches NAP’s performance in most tasks, highlighting HL-MBO’s robustness. Moreover, the Random Hypothesis surpasses NAP on average across the six different tasks, demonstrating the HL-MBO’s ability to maintain strong performance even without access to an optimal hypothesis.

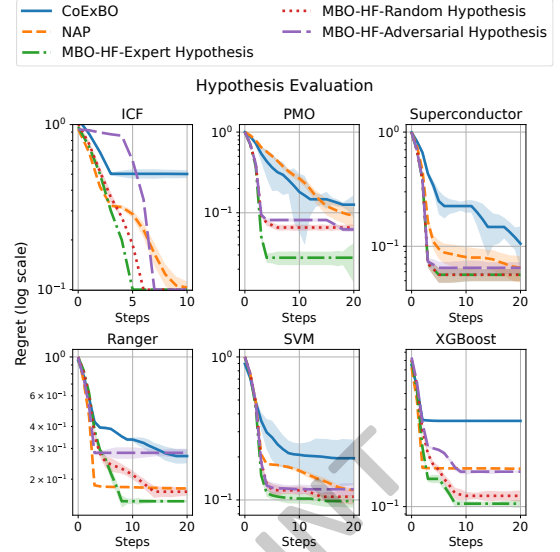


Figure 4: Results of the comparison between different hypothesis. Expert and Adversarial represent the best-case and worst-case scenarios, respectively, for sampling in the preference learning initialization. Random, represents uniform sampling.

Accuracy Evaluation We study the effect of correctness in experts’ selections when building dataset D_{pref} . Specifically, we assess how the accuracy in selecting the input value x_p that yields the highest output y_p from the candidate pairs $\{x_1, x_2\}$ affects HL-MBO’s performance. We evaluate the performance of HL-MBO on different experts’ accuracy thresholds: 50%, 75% and 100%. To conduct this analysis, we assume access to the target function, which allows us to evaluate all pairs of inputs in D_{pref} and generate the corresponding datasets for each accuracy level. For all three cases we consider a good hypothesis. Figure 5 shows the results of this evaluation. In most cases, 75% and 100% accuracy levels, which represent some level of knowledge about the black-box function by the expert, show better performance than random selection (50%), emphasizing the crucial role of expert input in the optimization process. Moreover, we can see that in most cases even 75% offers significant benefits.

Exploration Hyperparameter Evaluation We investigate the influence of the parameter ξ which governs the balance between exploration and exploitation in our EI-based AF. Higher values of ξ encourage more exploration. The results of this evaluation are shown in Figure 6. Across all domains, smaller values of ξ within the range of 0.1 and 0.3 generally yield better performance compared to 0.5. However, domains such as Ranger and SVM require a higher level of exploration, as $\xi = 0.3$ obtains the best performance. These findings highlight that the optimal value of ξ is domain-specific and must be carefully tuned to suit the problem at hand. Following standard practices, we use validation sets to fine-tune ξ for our experiments.

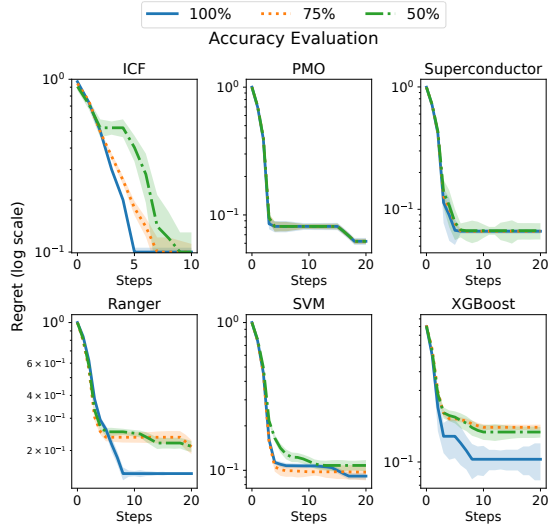


Figure 5: Results showing the accuracy of experts in selecting the better point from each presented pair during preference model construction. 50% accuracy indicates random choice.

5 Related Work

Human-aided and Explainable BO While no prior work has integrated few-shot meta-learning into human-aided explainable BO as proposed in this paper, recent efforts have advanced both explainable and human-in-the-loop BO.

One methodology is to incorporate human knowledge into BO by encoding it as a prior over the input space [Hvarfner *et al.*, 2022; Souza *et al.*, 2021; Ramachandran *et al.*, 2020; Cisse *et al.*, 2024; Hvarfner *et al.*, 2024], guiding the search toward promising regions. However, this requires users to precisely express their expertise, which becomes challenging in high-dimensional spaces [Mikkola *et al.*, 2021].

Another strategy engages humans during optimization, allowing them to reject or label points based on intuition or expertise [Gupta *et al.*, 2023; Adachi *et al.*, 2024; Chakraborty *et al.*, 2024; Akata *et al.*, 2020; Venkatesh *et al.*, 2022; Borji and Itti, 2013]. This interactive strategy enables more adaptive exploration of the search space and presents the expert with a simpler avenue to integrate their knowledge.

Methods for explainable BO have also emerged, providing insight into BO decisions to improve user understanding [Adachi *et al.*, 2024; Rodemann *et al.*, 2024; Chakraborty *et al.*, 2023; Li and Adams, 2020]. These approaches aim to improve user understanding of the decision-making process, providing insight during or after optimization to support a more informed evaluation of candidate solutions.

Optimization for ICF Standard BO has been applied to target design and experimental optimization in ICF [Chung *et al.*, 2020; Vazirani *et al.*, 2024; Hatfield *et al.*, 2019; Li *et al.*, 2023; Vazirani *et al.*, 2023]. Additionally, surrogate models built from ICF simulations or historical data have been used for efficient experimental optimization [Chung *et al.*, 2020; Humbird and Peterson, 2022; Wu *et al.*, 2022; Gaffney *et al.*, 2024; Ben Tayeb *et al.*, 2024; Vander Wal *et al.*, 2023]. However, these approaches have yet to explore

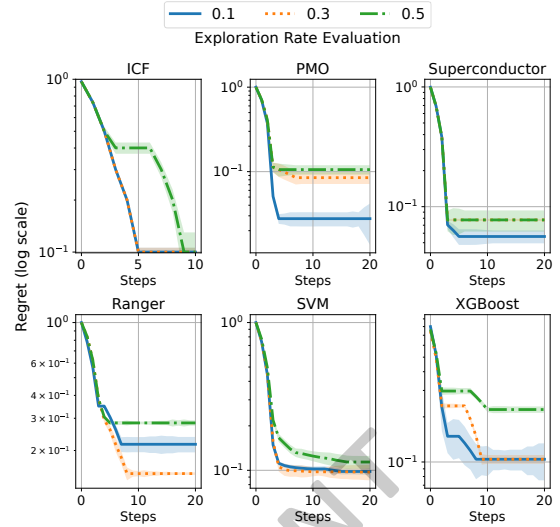


Figure 6: Results of the comparison between different values of ξ , which controls the level of exploration in HL-MBO's AF.

meta-learning or human-aided strategies to enhance trustworthiness and performance.

6 Conclusion

We presented HL-MBO, an explainable human-in-the-loop Meta Bayesian Optimization approach. We evaluated our method across three different hyperparameter optimization tasks and the real-world tasks of energy yield optimization in ICF, practical molecular optimization, and critical temperature maximization for superconductor materials. Our results showed that HL-MBO surpasses current SOTA for human-AI collaboration and Meta-BO. Moreover, we assessed the impact of experts' knowledge on Meta-BO and showed that its integration is beneficial.

Our work can have a significant impact on real-world scientific applications where experiments are costly but expert knowledge is available. In a critical domain like ICF, we demonstrate that HL-MBO outperforms all baselines. Advancing research in this direction, brings us closer to the goal of developing nearly limitless clean energy, which could have profound environmental and societal benefits.

6.1 Limitations

Meta-BO requires access to a set of related source tasks which might not be available in some cases. Moreover, training TNP can be computationally expensive. While effective, eliciting precise and comprehensive hypotheses from domain experts in complex, high-dimensional real-world settings is challenging, suggesting avenues for future research into interactive or automated elicitation methods. Furthermore, our method requires input from human experts, which might not be available in all cases.

Contribution

All authors contributed equally.

References

- [A V *et al.*, 2022] Arun Kumar A V, Santu Rana, Alistair Shilton, and Svetha Venkatesh. Human-ai collaborative bayesian optimisation. *Advances in Neural Information Processing Systems*, 35:16233–16245, 2022.
- [Adachi *et al.*, 2024] Masaki Adachi, Brady Planden, et al. Looping in the human: Collaborative and explainable bayesian optimization. In *International Conference on Artificial Intelligence and Statistics*, pages 505–513. PMLR, 2024.
- [Akata *et al.*, 2020] Zeynep Akata, Dan Balliet, et al. A research agenda for hybrid intelligence: augmenting human intellect with collaborative, adaptive, responsible, and explainable artificial intelligence. *Computer*, 53(8):18–28, 2020.
- [Bai *et al.*, 2023] Tianyi Bai, Yang Li, et al. Transfer learning for bayesian optimization: A survey, 2023.
- [Ben Tayeb *et al.*, 2024] M. Ben Tayeb, V. Tikhonchuk, and J.-L. Feugeas. Icf target optimization using generative ai. *Physics of Plasmas*, 31(10):103903, 10 2024.
- [Betti and Hurricane, 2016] R. Betti and O. A. Hurricane. Inertial-confinement fusion with lasers. *Nature Physics*, 12(5):435–448, 2016.
- [Borji and Itti, 2013] Ali Borji and Laurent Itti. Bayesian optimization explains human active search. *Advances in neural information processing systems*, 26, 2013.
- [Bradley and Terry, 1952] Ralrph Allan Bradley and Milton E. Terry. Rank analysis of incomplete block desings: The method of paired comparisons. *Biometrika*, 39(3-4):324–345, 12 1952.
- [Chakraborty *et al.*, 2023] Tanmay Chakraborty, Christian Wirth, and Christin Seifert. Post-hoc rule based explanations for black box bayesian optimization. In *European Conference on Artificial Intelligence*, pages 320–337. Springer, 2023.
- [Chakraborty *et al.*, 2024] Tanmay Chakraborty, Christin Seifert, and Christian Wirth. Explainable bayesian optimization. *arXiv preprint arXiv:2401.13334*, 2024.
- [Chung *et al.*, 2020] Youngseog Chung, Ian Char, et al. Offline contextual bayesian optimization for nuclear fusion, 2020.
- [Cisse *et al.*, 2024] Abdoulatif Cisse, Xenophon Evangelopoulos, et al. Hypbo: Accelerating black-box scientific experiments using experts’ hypotheses, 2024.
- [Colella *et al.*, 2020] Fabio Colella, Pedram Daee, et al. Human strategic steering improves performance of interactive optimization. In *Proceedings of the 28th ACM Conference on User Modeling, Adaptation and Personalization*, UMAP ’20. ACM, July 2020.
- [Feurer *et al.*, 2015] Matthias Feurer, Aaron Klein, et al. Efficient and robust automated machine learning. In *Advances in Neural Information Processing Systems*, volume 28. Curran Associates, Inc., 2015.
- [Gaffney *et al.*, 2024] Jim A. Gaffney, Kelli Humbird, et al. Data-driven prediction of scaling and ignition of inertial confinement fusion experiments. *Physics of Plasmas*, 31(9):092702, 2024.
- [Gao *et al.*, 2022] Wenhao Gao, Tianfan Fu, Jimeng Sun, and Connor Coley. Sample efficiency matters: a benchmark for practical molecular optimization. *Advances in Neural Information Processing Systems*, 35:21342–21357, 2022.
- [Garnelo *et al.*, 2018] Marta Garnelo, Dan Rosenbaum, et al. Conditional neural processes. In *Proceedings of The 25th International Conference on Artificial Intelligence and Statistics*, 2018.
- [Garnett, 2023] Roman Garnett. *Bayesian Optimization*. Cambridge University Press, 2023.
- [Gopalaswamy *et al.*, 2019] V Gopalaswamy, R Betti, et al. Tripled yield in direct-drive laser fusion through statistical modelling. *Nature*, 565(7741):581–586, January 2019.
- [Gopalaswamy *et al.*, 2024] V Gopalaswamy, C A Williams, et al. Demonstration of a hydrodynamically equivalent burning plasma in direct-drive inertial confinement fusion. *Nature Physics*, 20(5):751–757, May 2024.
- [Gupta *et al.*, 2023] Sunil Gupta, Alistair Shilton, Arun Kumar A V au2, Shannon Ryan, Majid Abdolshah, Hung Le, Santu Rana, Julian Berk, Mahad Rashid, and Svetha Venkatesh. Bo-muse: A human expert and ai teaming framework for accelerated experimental design, 2023.
- [Hatfield *et al.*, 2019] P. W. Hatfield, S. J. Rose, and R. H. H. Scott. The blind implosion-maker: Automated inertial confinement fusion experiment design. *Physics of Plasmas*, 26(6), June 2019.
- [Humbird and Peterson, 2022] K. D. Humbird and J. L. Peterson. Transfer learning driven design optimization for inertial confinement fusion. *Physics of Plasmas*, 29(10), October 2022.
- [Hvarfner *et al.*, 2022] Carl Hvarfner, Danny Stoll, Artur Souza, Marius Lindauer, Frank Hutter, and Luigi Nardi. π BO: Augmenting acquisition functions with user beliefs for Bayesian optimization. In *International Conference on Learning Representations (ICLR)*, 2022.
- [Hvarfner *et al.*, 2024] Carl Hvarfner, Frank Hutter, and Luigi Nardi. A general framework for user-guided bayesian optimization. In *The Twelfth International Conference on Learning Representations*, 2024.
- [Jha *et al.*, 2023] Saurav Jha, Dong Gong, Xuesong Wang, Richard E. Turner, and Lina Yao. The neural process family: Survey, applications and perspectives, 2023.
- [Lees *et al.*, 2021] A. Lees, R. Betti, et al. Experimentally inferred fusion yield dependencies of omega inertial confinement fusion implosions. *Phys. Rev. Lett.*, 127:105001, Aug 2021.
- [Li and Adams, 2020] Michael Y Li and Ryan P Adams. Explainability constraints for bayesian optimization. In *6th ICML Workshop on Automated Machine Learning*, 2020.

- [Li *et al.*, 2023] Z. Li, Z. Q. Zhao, et al. Hybrid optimization of laser-driven fusion targets and laser profiles, 2023.
- [Lun Chau *et al.*, 2022a] Siu Lun Chau, Javier Gonzalez, and Dino Sejdinovic. Learning inconsistent preferences with gaussian processes. In *Proceedings of The 25th International Conference on Artificial Intelligence and Statistics*, pages 2266–2281. PMLR, 2022.
- [Lun Chau *et al.*, 2022b] Siu Lun Chau, Javier Gonzalez, and Dino Sejdinovic. Learning inconsistent preferences with gaussian processes. In *Proceedings of The 25th International Conference on Artificial Intelligence and Statistics*, volume 151, pages 2266–2281, 2022.
- [Lundberg and Lee, 2017] Scott M Lundberg and Su-In Lee. A unified approach to interpreting model predictions. In *Advances in Neural Information Processing Systems 30*, pages 4765–4774. Curran Associates, Inc., 2017.
- [Maraval *et al.*, 2023] Alexandre Maraval, Matthieu Zimmer, Antoine Grosnit, and Haitham Bou Ammar. End-to-end meta-bayesian optimisation with transformer neural processes, 2023.
- [Merrill *et al.*, 2021] Erich Merrill, Alan Fern, Xiaoli Fern, and Nima Dolatnia. An empirical study of bayesian optimization: Acquisition versus partition. *Journal of Machine Learning Research*, 22(4):1–25, 2021.
- [Mikkola *et al.*, 2021] Petrus Mikkola, Osvaldo A Martin, et al. Prior knowledge elicitation: The past, present, and future. *arXiv preprint arXiv:2112.01380*, 2112, 2021.
- [Milios *et al.*, 2018] Dimitrios Milios, Raffaello Camoriano, Pietro Michiardi, Lorenzo Rosasco, and Maurizio Filippone. Dirichlet-based gaussian processes for large-scale calibrated classification, 2018.
- [Nguyen and Grover, 2022] Tung Nguyen and Aditya Grover. Transformer neural processes: Uncertainty-aware meta learning via sequence modeling. In *Proceedings of the 39th International Conference on Machine Learning*, volume 162, pages 16569–16594, 2022.
- [Pineda Arango *et al.*, 2021] Sebastian Pineda Arango, Hadi Jomaa, Martin Wistuba, and Josif Grabocka. Hpo-b: A large-scale reproducible benchmark for black-box hpo based on openml. In *Proceedings of the Neural Information Processing Systems Track on Datasets and Benchmarks*, volume 1, 2021.
- [Ramachandran *et al.*, 2020] Anil Ramachandran, Sunil Gupta, Santu Rana, Cheng Li, and Svetha Venkatesh. Incorporating expert prior in bayesian optimisation via space warping. *Knowledge-Based Systems*, 195:105663, 2020.
- [Ribeiro *et al.*, 2016] Marco Tulio Ribeiro, Sameer Singh, and Carlos Guestrin. "why should i trust you?": Explaining the predictions of any classifier, 2016.
- [Rodemann *et al.*, 2024] Julian Rodemann, Federico Croppi, Philipp Arens, et al. Explaining bayesian optimization by shapley values facilitates human-ai collaboration, 2024.
- [Shapley, 1952] Lloyd S. Shapley. *A Value for N-Person Games*. RAND Corporation, Santa Monica, CA, 1952.
- [Snoek *et al.*, 2012] Jasper Snoek, Hugo Larochelle, and Ryan P Adams. Practical bayesian optimization of machine learning algorithms. In *Advances in Neural Information Processing Systems*, volume 25. Curran Associates, Inc., 2012.
- [Souza *et al.*, 2021] Artur Souza, Luigi Nardi, et al. Bayesian optimization with a prior for the optimum. In *Machine Learning and Knowledge Discovery in Databases. Research Track: European Conference, ECML PKDD 2021, Bilbao, Spain, September 13–17, 2021, Proceedings, Part III 21*, pages 265–296. Springer, 2021.
- [Tacq, 2010] J. Tacq. The normal distribution and its applications. In *International Encyclopedia of Education (Third Edition)*, pages 467–473. Elsevier, Oxford, third edition edition, 2010.
- [Trabucco *et al.*, 2022] Brandon Trabucco, Xinyang Geng, Aviral Kumar, and Sergey Levine. Design-bench: Benchmarks for data-driven offline model-based optimization. *CoRR*, abs/2202.08450, 2022.
- [Vander Wal *et al.*, 2023] Michael D. Vander Wal, Ryan G. McClarren, and Kelli D. Humbird. Transfer learning as a method to reproduce high-fidelity non-local thermodynamic equilibrium opacities in simulations. *Journal of Plasma Physics*, 89(1):895890103, 2023.
- [Vazirani *et al.*, 2023] N. N. Vazirani, M. J. Grosskopf, D. J. Stark, P. A. Bradley, B. M. Haines, E. N. Loomis, S. L. England, and W. A. Scales. Coupling multi-fidelity xrage with machine learning for graded inner shell design optimization in double shell capsules. *Physics of Plasmas*, 30(6):062704, 06 2023.
- [Vazirani *et al.*, 2024] Nomita N. Vazirani, Ryan Sacks, et al. Bayesian batch optimization for molybdenum versus tungsten inertial confinement fusion double shell target design. *Statistical Analysis and Data Mining: An ASA Data Science Journal*, 17(3):e11698, 2024.
- [Venkatesh *et al.*, 2022] Arun Kumar Anjanapura Venkatesh, Santu Rana, Alistair Shilton, and Svetha Venkatesh. Human-ai collaborative bayesian optimisation. *Advances in Neural Information Processing Systems*, 2022.
- [Volpp *et al.*, 2020] Michael Volpp, Lukas P. Fröhlich, Kirsten Fischer, Andreas Doerr, Stefan Falkner, Frank Hutter, and Christian Daniel. Meta-learning acquisition functions for transfer learning in bayesian optimization, 2020.
- [Wang *et al.*, 2023] Xilu Wang, Yaochu Jin, Sebastian Schmitt, and Markus Olhofer. Recent advances in bayesian optimization. *ACM Comput. Surv.*, 55(13s), jul 2023.
- [Wu *et al.*, 2022] Fuyuan Wu, Xiaohu Yang, Yanyun Ma, Qi Zhang, Zhe Zhang, Xiaohui Yuan, Hao Liu, Zhengdong Liu, Jiayong Zhong, Jian Zheng, and et al. Machine-learning guided optimization of laser pulses for direct-drive implosions. *High Power Laser Science and Engineering*, 10:e12, 2022.