

---

# A General Framework for Multi-fidelity Bayesian Optimization with Gaussian Processes

---

Jialin Song  
Caltech

Yuxin Chen  
Caltech

Yisong Yue  
Caltech

## Abstract

How can we efficiently gather information to optimize an unknown function, when presented with multiple, mutually dependent information sources with different costs? For example, when optimizing a physical system, intelligently trading off computer simulations and real-world tests can lead to significant savings. Existing multi-fidelity Bayesian optimization methods, such as multi-fidelity GP-UCB or Entropy Search-based approaches, either make simplistic assumptions on the interaction among different fidelities or use simple heuristics that lack theoretical guarantees. In this paper, we study multi-fidelity Bayesian optimization with complex structural dependencies among multiple outputs, and propose MF-MI-Greedy, a principled algorithmic framework for addressing this problem. In particular, we model different fidelities using additive Gaussian processes based on shared latent relationships with the target function. Then we use cost-sensitive mutual information gain for efficient Bayesian optimization. We propose a simple notion of regret which incorporates the varying cost of different fidelities, and prove that MF-MI-Greedy achieves low regret. We demonstrate the strong empirical performance of our algorithm on both synthetic and real-world datasets.

## 1 Introduction

Optimizing an unknown function that is expensive to evaluate is a common problem in real applications. Examples include experimental design for protein engi-

neering, where chemists need to synthesize designed amino acid sequences and then test whether they satisfy certain properties (Romero et al., 2013); or black-box optimization for material science, where scientists need to run extensive computational experiments at various levels of accuracy to find the optimal material design structure (Fleischman et al., 2017). Conducting real experiments could be labor-intensive and time-consuming. In practice, one often utilizes alternative ways to gather information to make the most effective use of the real experiments that are conducted. A prevalent example is the use of computer simulation (van Gunsteren & Berendsen, 1990), which tends to be less time consuming but produces less accurate results. For instance, computer simulation is ubiquitous in robotic applications, e.g. one tests a control policy first in simulation before deploying it in a real physical system (Marco et al., 2017).

The central challenge in efficiently using multiple sources of information is captured in the general framework of multi-fidelity optimization (Forrester et al., 2007; Kandasamy et al., 2016, 2017; Marco et al., 2017; Sen et al., 2018) where multiple measurement/objective functions with varying degrees of accuracy and costs can be effectively leveraged to provide the maximal amount of information. However, strict assumptions, such as requiring strict relations between the quality and the cost of a lower fidelity function, and two-stage query selection criteria (cf. §2.2) are likely to lead to sub-optimal experiment design, and limit their practicality.

In this paper, we propose a general and principled multi-fidelity Bayesian optimization framework MF-MI-Greedy (Multi-fidelity Mutual Information Greedy) that prioritizes maximizing the amount of mutual information gathered across fidelity levels. Figure 1 captures the intuition of maximizing mutual information. Gathering information from a lower fidelity measurement also conveys information on the target fidelity. We make this idea concrete in §4. Our method improves upon prior work on multi-fidelity Bayesian optimization by establishing explicit connections across

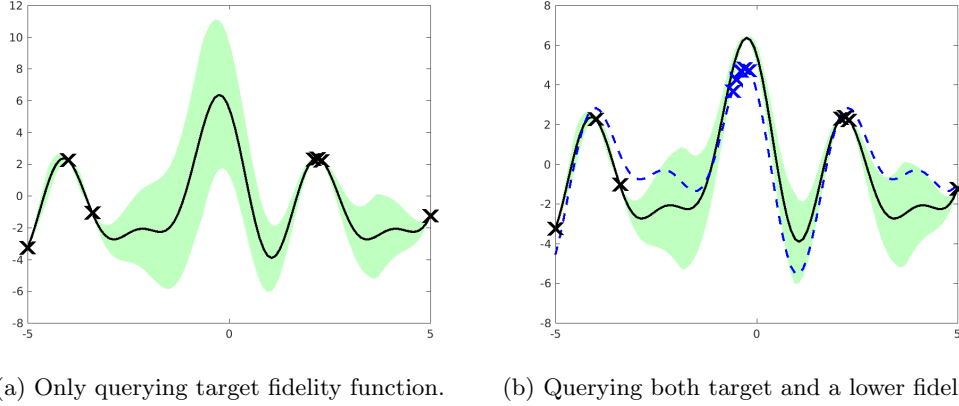


Figure 1: Depicting the benefit of multi-fidelity Bayesian optimization. The left panel shows normal single fidelity Bayesian optimization where locations near a query point (crosses) have low uncertainty. When there is a lower fidelity cheaper approximation in the right panel, by querying a large number of points of the lower fidelity function, the uncertainty in the target fidelity can also be reduced significantly.

fidelity levels to enable joint posterior updates and hyperparameter optimization. In summary, our contributions in this paper are:

- We study multi-fidelity Bayesian optimization with complex structural dependencies among multiple outputs (§3), and propose MF-MI-Greedy, a principled algorithmic framework for addressing this problem (§4).
- We propose a simple notion of regret which incorporates the cost of different fidelities, and prove that MF-MI-Greedy achieves low regret (§5).
- We demonstrate the empirical performance of MF-MI-Greedy on both simulated and real-world datasets (§6).

## 2 Background and Related Work

In this section, we review related work on Bayesian optimization with Gaussian processes (GPs).

### 2.1 Background on Gaussian Processes

A Gaussian process (Rasmussen & Williams, 2006) models an infinite collection of random variables, each indexed by an  $x \in \mathcal{X}$ , such that every finite subset of random variables has a multivariate Gaussian distribution. The GP distribution  $\text{GP}(\mu(x), k(x, x'))$  is a joint Gaussian distribution over all those (infinitely many) random variables specified by its mean  $\mu(x) = \mathbb{E}[f(x)]$  and covariance (also known as *kernel*) function:

$$k(x, x') = \mathbb{E}[(f(x) - \mu(x))(f(x') - \mu(x')))].$$

A key advantage of GPs is that it is very efficient to perform inference. Assume that  $f \sim$

$\text{GP}(\mu(x), k(x, x'))$  is a sample from the GP distribution, and that  $y = f(x) + \varepsilon(x)$  is a noisy observation of the function value  $f(x)$ . Here, the noise  $\varepsilon(x) \sim \mathcal{N}(0, \sigma^2(x))$  could depend on the input  $x$ . Suppose that we have selected  $\mathcal{S} \subseteq \mathcal{X}$  and received  $\mathbf{y}_{\mathcal{S}} = [f(x_i) + \varepsilon(x_i)]_{x_i \in \mathcal{S}}$ . We can obtain the posterior mean  $\mu_{\mathcal{S}}(x)$  and covariance  $k_{\mathcal{S}}(x, x')$  of the function through the covariance matrix  $K_{\mathcal{S}} = [k(x_i, x_j) + \delta_{ij}\sigma^2(x_i)]_{x_i, x_j \in \mathcal{S}}$  and  $\mathbf{k}_{\mathcal{S}}(x) = [k(x_i, x)]_{x_i \in \mathcal{S}}$ :

$$\mu_{\mathcal{S}}(x) = \mu(x) + \mathbf{k}_{\mathcal{S}}(x)^{\top} K_{\mathcal{S}}^{-1} \mathbf{y}_{\mathcal{S}} \quad (2.1)$$

$$k_{\mathcal{S}}(x, x') = k(x, x') - \mathbf{k}_{\mathcal{S}}(x)^{\top} K_{\mathcal{S}}^{-1} \mathbf{k}_{\mathcal{S}}(x'), \quad (2.2)$$

where  $\delta_{ij}$  is the Kronecker delta function.

### 2.2 Bayesian Optimization via GPs

#### Single-fidelity Gaussian Process optimization

Optimizing an unknown and noisy function is a common task in Bayesian optimization. In real applications, such functions tend to be expensive to evaluate, for example tuning hyperparameters for deep learning models (Snoek et al., 2012), so it is typically desirable to minimize the number of evaluation queries. A Gaussian process is an expressive and flexible tool to model a large class of functions. A classical method for Bayesian optimization with GPs is GP-UCB (Srinivas et al., 2010), which treats Bayesian optimization as a multi-armed bandit problem and proposes an upper-confidence bound based algorithm for query selections.

Entropy search (Hennig & Schuler, 2012) represents another class of GP-based Bayesian optimization approach. Its main idea is to directly search for the global optimum of an unknown function through queries. Each query point is selected based on its informativeness in learning the location for the function optimum. Predictive entropy search (Hernández-Lobato et al., 2014) addresses some computational

issues from entropy search by maximizing the expected information gain with respect to the location of the global optimum. Max-value entropy search (Wang et al., 2016; Wang & Jegelka, 2017) instead focuses on identifying the value of the global optimum, which effectively avoids issues related to the dimension of the search space, and the authors are able to provide regret bounds previous entropy search methods lacked.

A computational consideration for learning with GPs is how to optimize specific kernels used to model the covariance structures of GPs. As this optimization task depends on the dimension of feature space, approximation methods are often needed to speed up the learning process. For instance, Random Fourier features (Rahimi & Recht, 2008) can be efficient tools for dimension reduction and are often employed in GP regression tasks (Lázaro-Gredilla et al., 2010). As elaborated on in §4, our algorithmic framework offers the flexibility of choosing among different single-fidelity optimization approaches as a subroutine, so that one can take advantage of these computational and approximation algorithms for efficient optimization.

**Multi-output Gaussian Process** In some applications, it is desirable to model multiple correlated outputs with Gaussian processes. Most GP-based multi-output models create correlated outputs by mixing a set of independent latent processes. A simple form of such a mixing scheme is the linear model of coregionalization (Teh et al., 2005; Bonilla et al., 2008), where each output is modeled as a linear combination of latent GPs with fixed coefficients. The dependencies among outputs are captured by sharing those latent GPs. More complex structures can be captured by a linear combination of GPs with input-dependent coefficients (Wilson et al., 2012), shared inducing variables (Nguyen & Bonilla, 2014), or convolved process (Boyle & Frean, 2005; Alvarez & Lawrence, 2009). In comparison with single fidelity/output GPs, multi-output GPs often require more sophisticated approximate models for efficient optimization, e.g., using inducing points (Snelson & Ghahramani, 2007) to reduce the storage and computational complexity, and variational inference approaches to approximate the posterior of the latent processes (Titsias, 2009; Nguyen & Bonilla, 2014). While the analysis of our framework is not limited to a fixed structural assumption in modeling the joint distribution among multiple outputs, for efficiency concerns, we use a simple, additive model between multiple fidelity outputs in our experiments (cf. §6.1) to demonstrate the effectiveness of the optimization framework.

**Multi-fidelity Bayesian optimization** Multi-fidelity optimization is a general framework that cap-

tures the trade-off between cheap low quality versus expensive high quality experiments. Recently, there have been several works on using GPs to model functions of different fidelity levels. Recursive co-kriging (Forrester et al., 2007; Le Gratiet & Garnier, 2014) considers an autoregressive model for multi-fidelity GP regression, which assumes that the higher fidelity consists of a lower fidelity term and an *independent* GP term which models the systematic error for approximating the higher-fidelity output. Therefore, one can model cross-covariance between the high-fidelity and low-fidelity functions using the covariance of the lower fidelity function only. Virtual vs Real (Marco et al., 2017) extends this idea to Bayesian optimization. The authors consider a two-fidelity setting (i.e., virtual simulations and real system experiments), where they model the correlation between the two fidelities through co-kriging, and then apply entropy search (ES) to optimize the target output. Zhang et al. (2017) model the dependencies between different fidelities with convolved Gaussian processes (Alvarez & Lawrence, 2009), and then apply predictive entropy search (PES) (Hernández-Lobato et al., 2014) to efficient exploration. Although both multi-fidelity ES and multi-fidelity PES heuristics have shown promising empirical results on some datasets, little is known about their theoretical performance.

Recently, Kandasamy et al. (2016) proposed Multi-fidelity GP-UCB (MF-GP-UCB), a principled framework for multi-fidelity Bayesian optimization with Gaussian processes. In their followup works (Kandasamy et al., 2017; Sen et al., 2018), the authors address the disconnect issue by considering a continuous fidelity space and performing joint updates to effectively share information among different fidelity levels.

In contrast to our general assumption on the joint distribution between different fidelities, Kandasamy et al. (2016) and Sen et al. (2018) assume a specific structure over multiple fidelities, where the cost of each lower fidelity is determined according to the maximal approximation error in function value when compared with the target fidelity. Kandasamy et al. (2017) consider a two-stage optimization process, where the action and the fidelity level are selected in two separate stages. We note that this procedure may lead to non-intuitive choices of queries: For example, in a pessimistic case where the low fidelity only differs from the target fidelity by a constant shift, their algorithm is likely to focus only on querying the target fidelity actions even though the low fidelity is as useful as the target fidelity. In contrast, as described in §4, our algorithm jointly selects a query point and a fidelity level so such sub-optimality can be avoided.

Another related work is the multi-fidelity version of

Truncated Variance Reduction (TruVaR) proposed in (Bogunovic et al., 2016) where each fidelity level is represented by the target fidelity plus a known homoscedastic noise function. Costs for querying each fidelity is inversely related to the variance of the noise. TruVar uses variance reduction per unit cost as the acquisition function. Similar to Kandasamy et al. (2017), TruVar treats each fidelity level separately whereas our proposed approach considers different fidelity levels jointly. Furthermore, in our experiments, we don't assume knowledge about noise functions and instead learn them with Gaussian processes.

### 3 Problem Statement

**Payoff function and auxiliary functions** Consider the problem of maximizing an unknown payoff function  $f_m : \mathcal{X} \rightarrow \mathbb{R}$ . We can probe the function  $f_m$  by directly querying it at some  $x \in \mathcal{X}$  and obtaining a noisy observation  $y_{\langle x, m \rangle} = f_m(x) + \varepsilon$ , where  $\varepsilon \sim \mathcal{N}(0, \sigma^2)$  denotes i.i.d. Gaussian white noise. In addition to the unknown payoff function  $f_m$ , we are also given query access to other unknown auxiliary functions  $f_1, \dots, f_{m-1} : \mathcal{X} \rightarrow \mathbb{R}$ ; similarly, we obtain a noisy observation  $y_{\langle x, \ell \rangle} = f_\ell(x) + \varepsilon$  when querying  $f_\ell$  at  $x$ . Here, each  $f_\ell$  could be viewed as a low-fidelity version of  $f_m$  for  $\ell < m$ . For example, if  $f_m(x)$  represents the reward obtained by running a physical system with input  $x$ , then  $f_\ell(x)$  may represent the simulated payoff from a numerical simulator at fidelity level  $\ell$ .

**Joint distribution on multiple fidelities** We assume that multiple fidelities  $\{f_\ell\}_{\ell \in [m]}$  are mutually dependent through a joint probability distribution  $\mathbb{P}[f_1, \dots, f_m]$ . In particular, we model  $\mathbb{P}$  with a multiple output Gaussian process; hence the marginal distribution on each fidelity is a separate GP, i.e.,  $\forall \ell \in [m]$ ,  $f_\ell \sim \text{GP}(\mu_\ell(x), k_\ell(x, x'))$ , where  $\mu_\ell, k_\ell$  specify the (prior) mean and covariance at fidelity level  $\ell$ .

**Action, reward, and cost** Let us use  $\langle x, \ell \rangle$  to denote the action of querying  $f_\ell$  at  $x$ . Each action  $\langle x, \ell \rangle$  incurs a cost of  $\lambda_\ell$ , and achieves a reward:

$$q(\langle x, \ell \rangle) = \begin{cases} f_m(x) & \text{if } \ell = m \\ q_{\min} & \text{o.w.} \end{cases} \quad (3.1)$$

That is, performing  $\langle x, m \rangle$  (at the target fidelity) achieves a reward  $f_m(x)$ . We receive the minimal immediate reward  $q_{\min}$  with lower fidelity actions  $\langle x, \ell \rangle$  for  $\ell < m$ , even though it may provide some information about  $f_m$  and could thus lead to more informed decisions in the future. W.l.o.g., we assume that  $\max_x f_m(x) \geq 0$ , and  $q_{\min} \equiv 0$ .

**Policy** Let the policy  $\pi$  be an adaptive strategy for picking actions. Specifically, a policy specifies which action to perform next, based on the actions picked so far and their corresponding observations. We consider policies with a fixed budget  $\Lambda$ : upon termination,  $\pi$  returns a sequence of actions  $\mathcal{S}_\pi$ , such that  $\sum_{\langle x, \ell \rangle \in \mathcal{S}_\pi} \lambda_\ell \leq \Lambda$ . Note that for a given policy  $\pi$ , the sequence  $\mathcal{S}_\pi$  is random, and depends on the joint distribution  $\mathbb{P}$  and the (random) observations of the selected actions.

**Objective** Given a budget  $\Lambda$ , our goal is to maximize the expected cumulative reward, so as to identify an action  $\langle x, m \rangle$  with performance close to  $x^* = \arg \max_{x \in \mathcal{X}} f_m(x)$  as rapidly as possible. Formally, we seek:

$$\pi^* = \arg \max_{\pi: \sum_{\langle x, \ell \rangle \in \mathcal{S}_\pi} \lambda_\ell \leq \Lambda} \mathbb{E}_{\mathcal{S}_\pi} \left[ \sum_{\langle x, \ell \rangle \in \mathcal{S}_\pi} q(\langle x, \ell \rangle) \right]. \quad (3.2)$$

**Remarks** Problem 3.2 strictly generalizes the optimal value of information (VoI) problem (Chen et al., 2015) to the online setting. To see this, consider the scenario where  $\lambda_m \gg \lambda_\ell$  for  $\ell < m$ , and  $\Lambda \in (\lambda_m, 2\lambda_m)$ . To achieve a non-zero reward, a policy must pick  $\langle x, m \rangle$  as the last action before exhausting the budget  $\Lambda$ . Therefore, our goal translates to adaptively choosing lower fidelity actions that are the most informative about  $x^*$  under budget  $\Lambda - \lambda_m$ , which reduces to the optimal VoI problem.

### 4 The Multi-fidelity BO Framework

We now present MF-MI-Greedy, a general framework for multi-fidelity Gaussian process optimization. In a nutshell, MF-MI-Greedy attempts to balance the ‘‘exploratory’’ low-fidelity actions and the more expensive target fidelity actions, based on how much information (per unit cost) these actions could provide about the target fidelity function. Concretely, MF-MI-Greedy proceeds in rounds under a given budget  $\Lambda$ . Each round can be divided into two phases: (i) an exploration (i.e., information gathering) phase, where the algorithm focuses on exploring the low fidelity actions, and (ii) an optimization phase, where the algorithm tries to optimize the payoff function  $f_m$  by performing an action  $\langle x, m \rangle$  at the target fidelity. The pseudocode of the algorithm is provided in Algorithm 1.

A key challenge in designing the algorithm is to decide when to stop exploration, or equivalently, to invoke the optimization phase. While somewhat analogous to the exploration-exploitation tradeoff in the classical single-fidelity Bayesian optimization problems, the multi-fidelity setting has a more distinctive notion of ‘‘exploration’’ and a more complicated structure of the

**Algorithm 1:** Multi-fidelity Mutual Information Greedy Optimization (MF-MI-Greedy)

---

```

1 Input: Total budget  $\Lambda$ ; cost  $\lambda_i$  for all fidelities
    $i \in [m]$ ; joint GP (prior) distribution on  $\{f_i, \varepsilon_i\}_{i \in [m]}$ 
begin
2    $\mathcal{S} \leftarrow \emptyset$ 
3    $B \leftarrow \Lambda$ ; /* initialize remaining budget */
   while  $B > 0$  do
4     /* explore with low fidelity */
      $\mathcal{L} \leftarrow \text{Explore-LF}(B, [\lambda_\ell], \text{GP}(\{f_\ell, \varepsilon_\ell\}_{\ell \in [m]}), \mathcal{S})$ 
     /* select target fidelity */
5      $x^* \leftarrow \text{SF-GP-OPT}(\text{GP}(\{f_m, \varepsilon_m\}), \mathbf{y}_{\mathcal{S} \cup \mathcal{L}})$ 
6      $\mathcal{S} \leftarrow \mathcal{S} \cup \mathcal{L} \cup \{x^*, m\}$ 
7      $B \leftarrow \Lambda - \Lambda_{\mathcal{S}}$ ; /* update remaining budget */
8 Output: Optimizer of the target function  $f_m$ 

```

---

action space (i.e., each exploration phase corresponds to picking a set of low fidelity actions). Furthermore, note that there is no explicit measurement of the relative “quality” of a low fidelity action, as they all have uniform reward by our modeling assumption (c.f. Eq. (3.1)); hence we need an appropriate measure to keep track of the progress of exploration.

We consider an information-theoretic selection criterion for picking low fidelity actions. The quality of a low fidelity action  $(x, \ell)$  is captured by the *information gain*, defined as the amount of entropy reduction in the posterior distribution of the target payoff function<sup>1</sup>:  $\mathbb{I}(y_{(x, \ell)}; f_m \mid \mathbf{y}_{\mathcal{S}}) = \mathbb{H}(y_{(x, \ell)} \mid \mathbf{y}_{\mathcal{S}}) - \mathbb{H}(y_{(x, \ell)} \mid f_m, \mathbf{y}_{\mathcal{S}})$ . Here,  $\mathcal{S}$  denotes the set of previously selected actions,  $\mathbf{y}_{\mathcal{S}}$  denote the observation history and  $\mathbb{H}$  denotes the entropy. Given an exploration budget  $B$ , our objective for a single exploration phase is to find a set of low-fidelity actions, which are (i) maximally informative about the target function, (ii) better than the best action on the target fidelity when considering the information gain per unit cost (otherwise, one would rather pick the target fidelity action to trade off exploration and exploitation), and (iii) not overly aggressive in terms of exploration (since we would also like to reserve a certain budget for performing target fidelity actions to gain reward).

Finding the optimal set of actions satisfying the above design principles is computationally prohibitive, as it requires us to search through a combinatorial (for finite discrete domains) or even infinite (for continuous domains) space. In Algorithm 2, we introduce Explore-LF, a key subroutine of MF-MI-Greedy, for efficient exploration on low fidelities. At each step, Explore-LF

**Algorithm 2:** Explore-LF: Explore low fidelities

---

```

1 Input: Exploration budget  $B$ ; cost  $[\lambda_\ell]_{\ell \in [m]}$ ; joint
   GP distribution on  $\{f_i, \varepsilon_i\}_{i \in [m]}$ ; previously selected
   items  $\mathcal{S}$ 
begin
2    $\mathcal{L} \leftarrow \emptyset$ ; /* selected actions */
3    $\Lambda_{\mathcal{L}} \leftarrow 0$ ; /* cost of selected actions */
4    $\beta \leftarrow \frac{1}{\alpha(B)}$ ; /* threshold (c.f. Theorem 1) */
   while true do
5     /* greedy benefit-cost ratio */
      $\langle x^*, \ell^* \rangle \leftarrow$ 
        $\arg \max_{(x, \ell): \lambda_\ell \leq B - \Lambda_{\mathcal{L}} - \lambda_m} \frac{\mathbb{I}(y_{(x, \ell)}; f_m \mid \mathbf{y}_{\mathcal{S} \cup \mathcal{L}})}{\lambda_\ell}$ 
     if  $\ell^* = \text{null}$  then
6       break; /* budget exhausted */
     if  $\ell^* = m$  then
7       break; /* worse than target */
     else if  $\frac{\mathbb{I}(\mathbf{y}_{\mathcal{L} \cup \{x^*, \ell^*\}}; f_m \mid \mathbf{y}_{\mathcal{S}})}{(\Lambda_{\mathcal{L}} + \lambda_{\ell^*})} < \beta$  then
8       break; /* low cumulative ratio */
     else
9        $\mathcal{L} \leftarrow \mathcal{L} \cup \{x^*, \ell^*\}$ 
10       $\Lambda_{\mathcal{L}} \leftarrow \Lambda_{\mathcal{L}} + \lambda_{\ell^*}$ 
11 Output: Selected set of items  $\mathcal{L}$  from lower fidelities

```

---

takes a greedy step w.r.t. the benefit-cost ratio over all actions. To ensure that the algorithm does not explore excessively, we consider the following stopping conditions: (i) when the budget is exhausted (Line 6), (ii) when a single target fidelity action is better than all the low fidelity actions in terms of the benefit-cost ratio (Line 7), and (iii) when the cumulative benefit-cost ratio is small (Line 8). Here, the parameter  $\beta$  is set to be  $\Omega\left(\frac{1}{\sqrt{B}}\right)$  to ensure low regret, and we defer the detailed discussion of the choice of  $\beta$  to §5.2.

**Optimization phase** At the end of the exploration phase, MF-MI-Greedy updates the posterior distribution of the joint GP using the full observation history, and searches for a target fidelity action via the (single-fidelity) GP optimization subroutine SF-GP-OPT (Line 5). Here, SF-GP-OPT could be *any* off-the-shelf Bayesian optimization algorithm, such as GP-UCB (Srinivas et al., 2010), GP-MI (Contal et al., 2014), EST (Wang et al., 2016) and MVES (Wang & Jegelka, 2017), TruVaR (Bogunovic et al., 2016), etc. Different from the previous exploration phase which seeks an informative set of low fidelity actions, the GP optimization subroutine aims to trade off exploration and exploitation only within the target fidelity, and outputs a single action at each round. MF-MI-Greedy then iteratively proceeds to the next round until it exhausts the preset budget, and finally outputs an estimator of the target function optimizer.

<sup>1</sup>An alternative, more aggressive information measurement is the information gain over the *optimizer* of the target function  $f_m$  (Hennig & Schuler, 2012), i.e.,  $\mathbb{I}(y_{(x, \ell)}; \arg \max_x f_m(x) \mid \mathbf{y}_{\mathcal{S}})$ , or the *optimal value* of  $f_m$  (Wang & Jegelka, 2017), i.e.,  $\mathbb{I}(y_{(x, \ell)}; \max_x f_m(x) \mid \mathbf{y}_{\mathcal{S}})$ .

## 5 Theoretical Analysis

In this section, we investigate the theoretical behavior of MF-MI-Greedy. We first introduce an intuitive notion of regret for the multi-fidelity setting, and then state our main theoretical results.

### 5.1 Multi-fidelity Regret

**Definition 1** (Episode). Let  $t \in \mathbb{Z}$  be any integer. We call a sequence of items  $\mathcal{E} = \{\langle x_1, \ell_1 \rangle, \dots, \langle x_t, \ell_t \rangle\}$  an episode, if  $\forall \tau < t: \ell_\tau < m$  and  $\ell_t = m$ .

Note that only the last action of an episode is from the target fidelity  $f_m$  and all remaining actions are from lower fidelities. We now define a simple notion of regret for an episode.

**Definition 2** (Episode regret). The regret of an episode  $\mathcal{E} = \{\langle x_1, \ell_1 \rangle, \dots, \langle x_t, m \rangle\}$  is:

$$r(\mathcal{E}) = \frac{\Lambda_{\mathcal{E}}}{\lambda_m} f_m^* - q(\mathcal{E}), \quad (5.1)$$

where  $\Lambda_{\mathcal{E}} := \sum_{\langle x, \ell \rangle \in \mathcal{E}} \lambda_{\ell}$  is the total cost of episode  $\mathcal{E}$ ,  $f_m^* := \max_x f_m(x)$ , and  $q(\mathcal{E}) := f_m(x_t)$  denotes the reward value of the last action on the target fidelity.

Suppose we run policy  $\pi$  under budget  $\Lambda$  and select a sequence of actions  $\mathcal{S}_{\pi}$ . One can represent  $\mathcal{S}_{\pi}$  using multiple episodes  $\mathcal{S}_{\pi} = \{\mathcal{E}^{(1)}, \dots, \mathcal{E}^{(k)}\}$ , where  $\mathcal{E}^{(j)} = \mathcal{L}^{(j)} \cup \{\langle x, m \rangle^{(j)}\}$  denotes the sequence of low fidelity actions  $\mathcal{L}^{(j)}$  and target fidelity action  $\langle x, m \rangle^{(j)}$  selected at the  $j^{\text{th}}$  episode. Let  $\Lambda_{\mathcal{E}^{(j)}}$  be the cost of episode  $j$ ; clearly  $\sum_{j=1}^k \Lambda_{\mathcal{E}^{(j)}} \leq \Lambda$ . We define the multi-fidelity cumulative regret as follows.

**Definition 3** (Cumulative regret). The cumulative regret of policy  $\pi$  under budget  $\Lambda$  is:

$$R(\pi, \Lambda) = \frac{\Lambda}{\lambda_m} f_m^* - \sum_{j=1}^k q(\mathcal{E}^{(j)}). \quad (5.2)$$

Intuitively, Definition 3 characterizes the difference between the cumulative reward of  $\pi$  and the best possible reward gathered under budget  $\Lambda$ .<sup>2</sup>

### 5.2 Regret Analysis

In the following, we establish a bound on the cumulative regret of MF-MI-Greedy, as a function of the mutual information between the target fidelity function and the actions attained by the algorithm.

<sup>2</sup>Note that our notion of cumulative regret is different from the multi-fidelity regret (Eq. (2)) of Kandasamy et al. (2016). Although both definitions reduce to the classical single-fidelity regret (Srinivas et al., 2010) when  $m = 1$ , Definition 3 has a simpler form and intuitive interpretation.

**Theorem 1.** Assume that MF-MI-Greedy terminates in  $k$  episodes, and w.h.p., the cumulative regret incurred by SF-GP-OPT is upper bounded by  $C_1 \sqrt{\Lambda \gamma_m}$ , where  $C_1$  is some constant independent of  $\Lambda$ , and  $\gamma_m$  denotes the mutual information gathered by the target fidelity actions chosen by SF-GP-OPT (equivalently by MF-MI-Greedy). Then, w.h.p., the cumulative regret of MF-MI-Greedy (Algorithm 1) satisfies:

$$R(\text{MF-MI-Greedy}, \Lambda) \leq C_1 \sqrt{\Lambda \gamma_m} + C_2 \alpha_{\Lambda} \gamma_{\mathcal{L}},$$

where  $\alpha_{\Lambda} = \max_{B \leq \Lambda} \alpha(B)$ ,  $C_2$  is some constant independent of  $\Lambda$ , and  $\gamma_{\mathcal{L}} = \sum_{j=1}^k \mathbb{I}(\mathbf{y}_{\mathcal{L}}^{(j)}; f_m \mid \mathbf{y}_{\mathcal{E}}^{(1:j-1)})$  denotes the mutual information gathered by the low fidelity actions when running MF-MI-Greedy.

The proof of Theorem 1 is provided in the Appendix. Similar to the single-fidelity setting, a desirable asymptotic property of a multi-fidelity optimization algorithm is to be no-regret, i.e.,  $\lim_{\Lambda \rightarrow \infty} R(\pi, \Lambda)/\Lambda \rightarrow 0$ . If we set  $\alpha(B) = o(\sqrt{B})$ , then Theorem 1 reduces to  $R(\text{MF-MI-Greedy}, \Lambda) \leq C_1 \sqrt{\Lambda \gamma_m} + C_2 \gamma_{\mathcal{L}} \cdot o(\sqrt{\Lambda})$ . Clearly, MF-MI-Greedy is no-regret as  $\Lambda \rightarrow \infty$ .

Furthermore, let us compare the above result with the regret bound of the single-fidelity GP optimization algorithm SF-GP-OPT. By the assumption of Theorem 1, we know  $R(\text{SF-GP-OPT}, \Lambda) = O(\sqrt{\Lambda \gamma'_m})$ , where  $\gamma'_m$  is the information gain of all the (target fidelity) actions by running SF-GP-OPT for  $\lfloor \Lambda/\lambda_m \rfloor$  rounds. When the low fidelity actions are very informative about the  $f_m$  and have much lower costs than  $\lambda_m$  (hence larger  $\gamma_{\mathcal{L}}$ ), the less likely MF-MI-Greedy will focus on exploring the target fidelity, i.e.,  $\gamma_m \leq \gamma'_m$ , and hence MF-MI-Greedy becomes more advantageous to SF-GP-OPT. The implication of Theorem 1 is similar in spirit to the regret bound provided in Kandasamy et al. (2016), however, our results apply to a much broader class of optimization strategies, as suggested by the following corollary.

**Corollary 2.** Let  $\alpha(B) = o(\sqrt{B})$ . Running MF-MI-Greedy with subroutine GP-UCB, EST, or GP-MI in the optimization phase is no-regret.

## 6 Experiments

In this section, we empirically evaluate our algorithm on three synthetic test function optimization tasks and two practical optimization problems.

### 6.1 Experimental Setup

To model the relationship between a low fidelity function  $f_i$  and the target fidelity function  $f_m$ , we use an additive model. Specifically, we assume that  $f_i(x) =$

$f_m(x) + \varepsilon_i(x)$  for all fidelity levels  $i < m$  where  $\varepsilon_i$  is an unknown function characterizing the error incurred by a lower fidelity function. We use Gaussian processes to model  $f_m$  and  $\varepsilon_i$ . Since  $f_m$  is embedded in every fidelity level, we can use an observation from any fidelity to update the posterior for *every* fidelity level. We use square exponential kernels for all the GP covariances, with hyperparameter tuning scheduled periodically during optimization by maximizing likelihood of the observed data. Specifically, we use a two-step procedure: First, we keep the error functions fixed and update the GP for the shared  $f_m$  across all fidelities; second, we keep the  $f_m$  fixed for each fidelity level and update the GP for the error function. We keep the same experimental setup for MF-GP-UCB as in [Kandasamy et al. \(2016\)](#). For all our experiments, we use a total budget of 100 times the cost of target fidelity function evaluation  $f_m$ . When the optimal value  $f_m^*$  for  $f_m$  is known, we compare simple regrets. Otherwise, we compare simple rewards, which is defined as the best objective value found so far.

## 6.2 Compared Methods

Our framework is general and we could plug in different single fidelity Bayesian optimization algorithms for the SF-GP-OPT procedure in Algorithm 1. In our experiment, we choose to use GP-UCB as one instantiation. We compare with MF-GP-UCB ([Kandasamy et al., 2016](#)) and GP-UCB ([Srinivas et al., 2010](#)).

## 6.3 Synthetic Examples

We first evaluate our algorithm on three synthetic datasets, (a) Hartmann 6D, (b) Currin exponential 2D and (c) BoreHole 8D ([Kandasamy et al., 2016](#)). We follow the setup used in [Kandasamy et al. \(2016\)](#) to define lower fidelity functions, while we use a different definition of lower fidelity costs. We emphasize that in synthetic settings, the artificially defined costs do not have practical meanings, as function evaluation costs do not differ across different fidelity levels. Nevertheless, we set the cost of the function evaluations (monotonically) according to the fidelity levels, and present the results in Fig. 3. The  $x$ -axis represents the expended budget and the  $y$ -axis represents the smallest simple regret. The error bars represent one standard error over 20 runs of each experiment.

MF-MI-Greedy is generally competitive with MF-GP-UCB. A common issue is its simple regrets tend to be larger at the beginning. A cause for this behavior may be the parameters controlling the termination conditions early on are not tuned optimally, which leads to over exploration in regions that do not reveal much information on where the function optimum lies.

## 6.4 Real Experiments

We test our methods on two real-world settings: maximum likelihood inference for cosmological parameters and experimental design for material science.

### 6.4.1 Maximum Likelihood Inference for Cosmological Parameters

The first real-world experiment is to perform maximum likelihood inference on 3 cosmological parameters, the Hubble constant  $H_0 \in (60, 80)$ , the dark matter fraction  $\Omega_M \in (0, 1)$  and the dark energy fraction  $\Omega_A \in (0, 1)$ . It thus has a dimensionality of 3. The likelihood is given by the Robertson-Walker metric, which requires a one-dimensional numerical integration for each point in the dataset from [Davis et al. \(2007\)](#). In [Kandasamy et al. \(2017\)](#), the authors set up two lower fidelity functions by considering two aspects of computing the likelihood: (i) how many data points (denoted by  $N$ ) are used, and (ii) what is the discrete grid size (denoted by  $G$ ) for performing the numerical integration. The range for these two parameters are  $N \in [50, 192]$  and  $G \in [10^2, 10^6]$ . We follow the fidelity levels selected in [Kandasamy et al. \(2017\)](#) which correspond to two lower fidelities with  $(N_1, G_1) = (97, 2.15 \times 10^3)$ ,  $(N_2, G_2) = (145, 4.64 \times 10^4)$  and the target fidelity with  $(N_3, G_3) = (192, 10^6)$ . Costs are defined as the product of  $N$  and  $G$ .

Upon further investigation, we find that the grid sizes selected above for performing numerical integration do not affect the final integral values, i.e., the grid size for the lowest fidelity  $G_1 = 2.15 \times 10^3$  is fine enough to compute an approximation to the integration as using the grid size for the target fidelity. So the costs used are not an accurate characterization of the true computation costs. As a result, we propose a different cost definition that depends only on how many data points are used to compute the likelihood, i.e., the new costs for the 3 functions are (97, 145, 192), respectively.

The results using the original cost definition are shown in Figure 3a. Note for this task we do not know the optimal likelihood, so we report the best objective value so far (simple rewards) in the  $y$ -axis. Our method MF-MI-Greedy (red) outperforms both baselines. The results using the new cost definition are shown in Figure 3b. Our method obtains a consistent high likelihood when the cost structure changes. However, MF-GP-UCB’s quality degrades significantly, which implies that it is sensitive to how the costs among fidelity levels are defined. These two set of results demonstrate the robustness of our method against costs, which is a desirable property as inaccuracy in cost estimates is inevitable in practical applications.

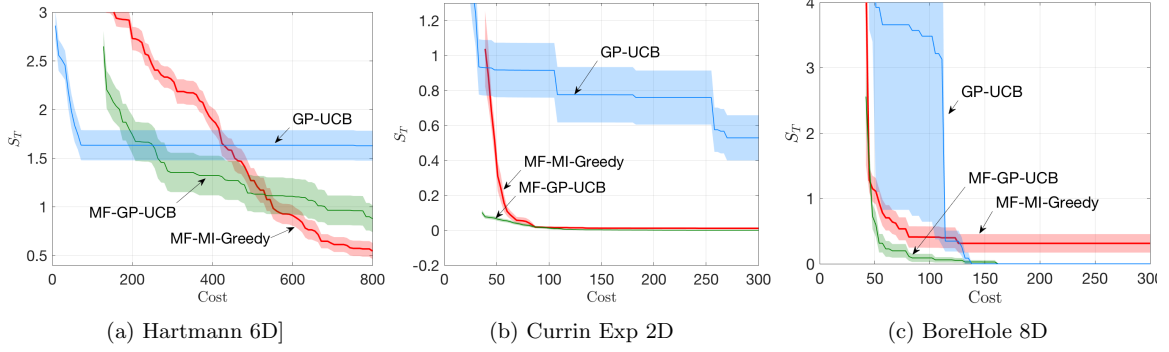


Figure 2: Synthetic datasets. For the Hartmann 6D dataset, costs = [1, 2, 4, 8]; for Currin Exp 2D, costs = [1, 3]; for BoreHole 8D, costs = [1, 2]. Error bar shows one standard error over 20 runs for each experiment.

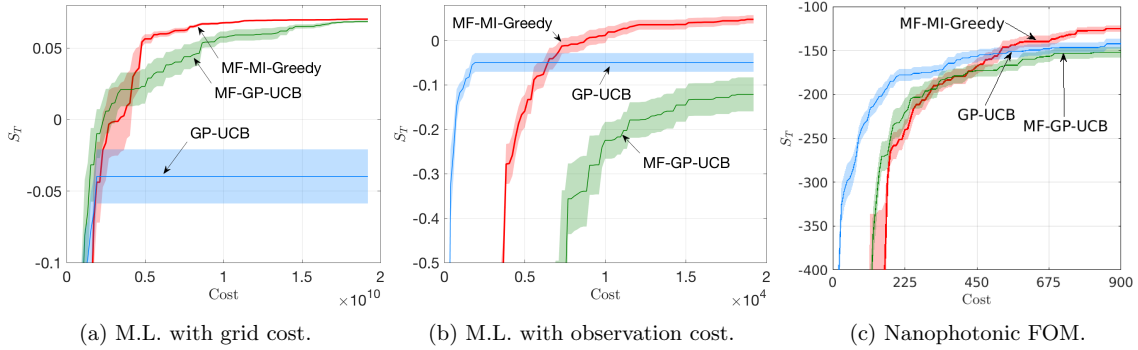


Figure 3: Two cost settings for maximum likelihood inference task and the task of optimizing FOM for nanophotonic structures. Error bar shows one standard error over 20 runs for each experiment.

#### 6.4.2 Optimizing Nanophotonic Structures

The second real-world experiment is motivated by a material science task of designing nanophotonic structures with desired color filtering property (Fleischman et al., 2017). In this setup, a nanophotonic structure is characterized by the following 5 parameters: mirror height, film thickness, mirror spacing, slit width, and oxide thickness. For each parameter setting, we use a fitness score, commonly called a figure-of-merit (FOM), to represent how well the resulting structure satisfies the desired color filtering property. By optimizing the FOM, we can find a set of high-quality design parameters.

Traditionally, perfect estimation of FOMs can only be performed through the actual fabrication of a structure and the testing of its various physical properties, which is a time-consuming process. Of course, simulations can be utilized to estimate what physical properties a design will have. By solving a variant of the Maxwell’s equations, we can simulate the transmission of the light spectrum and compute FOM from the spectrum. We data at three fidelity levels from 5000 nanophotonic structures. What distinguishes each fidelity is the mesh size used to solve the Maxwell’s equations. Finer meshes lead to more accurate results. Specifically, lowest fidelity uses a mesh size of  $3\text{nm} \times 3\text{nm}$ ,

the middle fidelity  $2\text{nm} \times 2\text{nm}$  and the target fidelity  $1\text{nm} \times 1\text{nm}$ . The costs, simulation time, are inversely proportional to the mesh size, so we use the following costs [1, 2.25, 9] for our three fidelity functions.

Figure 3c shows the results of this experiment. Here, the  $x$ -axis is the cost and  $y$ -axis is negative FOM. After a small portion of the budget is used in initial exploration, MF-MI-Greedy (red) is able to arrive at a better final design compared with MF-GP-UCB and GP-UCB.

## 7 Conclusion

In this paper, we investigated multi-fidelity Bayesian optimization, and proposed a general and principled framework for addressing the problem. We introduced an intuitive notion of regret, and showed that our framework is able to lift many popular, off-the-shelf single-fidelity GP optimization algorithms to the multi-fidelity setting while still preserving their original regret bounds. We demonstrated the performance of our algorithm on several synthetic and real datasets.

**Acknowledgments** This work was supported in part by NSF Award #1645832, Northrop Grumman, Bloomberg, Raytheon, PIMCO, and a Swiss NSF Early Mobility Postdoctoral Fellowship.

## References

- Alvarez, M. and Lawrence, N. D. Sparse convolved gaussian processes for multi-output regression. In *Advances in neural information processing systems*, pp. 57–64, 2009. [3](#)
- Bogunovic, I., Scarlett, J., Krause, A., and Cevher, V. Truncated variance reduction: A unified approach to bayesian optimization and level-set estimation. In *Advances in neural information processing systems*, pp. 1507–1515, 2016. [4](#), [5](#)
- Bonilla, E. V., Chai, K. M., and Williams, C. Multi-task gaussian process prediction. In *Advances in neural information processing systems*, pp. 153–160, 2008. [3](#)
- Boyle, P. and Frean, M. Dependent gaussian processes. In *Advances in neural information processing systems*, pp. 217–224, 2005. [3](#)
- Chen, Y., Javdani, S., Karbasi, A., Bagnell, J. A., Srinivasa, S., and Krause, A. Submodular surrogates for value of information. In *Proc. Conference on Artificial Intelligence (AAAI)*, January 2015. URL <https://bit.ly/2QEfJnZ>. [4](#)
- Contal, E., Perchet, V., and Vayatis, N. Gaussian process optimization with mutual information. In *International Conference on Machine Learning*, pp. 253–261, 2014. URL <https://bit.ly/2x7EEbw>. [5](#)
- Davis, T. M., Mörtzell, E., Sollerman, J., Becker, A. C., Blondin, S., Challis, P., Clocchiatti, A., Filippenko, A., Foley, R., Garnavich, P. M., et al. Scrutinizing exotic cosmological models using essence supernova data combined with other cosmological probes. *The Astrophysical Journal*, 666(2):716, 2007. [7](#)
- Fleischman, D., Sweatlock, L. A., Murakami, H., and Atwater, H. Hyper-selective plasmonic color filters. *Optics Express*, 25(22):27386–27395, 2017. [1](#), [8](#)
- Forrester, A. I., Söbester, A., and Keane, A. J. Multi-fidelity optimization via surrogate modelling. In *Proceedings of the royal society of london a: mathematical, physical and engineering sciences*, volume 463, pp. 3251–3269. The Royal Society, 2007. URL <https://bit.ly/2xkMXRr>. [1](#), [3](#)
- Hennig, P. and Schuler, C. J. Entropy search for information-efficient global optimization. *Journal of Machine Learning Research*, 13(Jun):1809–1837, 2012. URL <https://bit.ly/2x5KMQC>. [2](#), [5](#)
- Hernández-Lobato, J. M., Hoffman, M. W., and Ghahramani, Z. Predictive entropy search for efficient global optimization of black-box functions. In *Advances in neural information processing systems*, pp. 918–926, 2014. [2](#), [3](#)
- Kandasamy, K., Dasarathy, G., Oliva, J. B., Schneider, J., and Póczos, B. Gaussian process bandit optimisation with multi-fidelity evaluations. In *Advances in Neural Information Processing Systems*, pp. 992–1000, 2016. URL <https://bit.ly/2Qngemh>. [1](#), [3](#), [6](#), [7](#)
- Kandasamy, K., Dasarathy, G., Schneider, J., and Póczos, B. Multi-fidelity bayesian optimisation with continuous approximations. In *International Conference on Machine Learning*, pp. 1799–1808, 2017. URL <https://bit.ly/2N9KgMq>. [1](#), [3](#), [4](#), [7](#)
- Lázaro-Gredilla, M., Quiñero-Candela, J., Rasmussen, C. E., and Figueiras-Vidal, A. R. Sparse spectrum gaussian process regression. *Journal of Machine Learning Research*, 11:1865–1881, 2010. [3](#)
- Le Gratiet, L. and Garnier, J. Recursive co-kriging model for design of computer experiments with multiple levels of fidelity. *International Journal for Uncertainty Quantification*, 4(5), 2014. URL <https://bit.ly/2PICVQu>. [3](#)
- Marco, A., Berkenkamp, F., Hennig, P., Schoellig, A. P., Krause, A., Schaal, S., and Trimpe, S. Virtual vs. real: Trading off simulations and physical experiments in reinforcement learning with bayesian optimization. In *2017 IEEE International Conference on Robotics and Automation (ICRA)*, pp. 1557–1563, 2017. URL <https://bit.ly/20a4e62>. [1](#), [3](#)
- Nguyen, T. V. and Bonilla, E. V. Collaborative multi-output gaussian processes. In *UAI*, pp. 643–652, 2014. URL <https://bit.ly/2D4BwCH>. [3](#)
- Rahimi, A. and Recht, B. Random features for large-scale kernel machines. In *Advances in neural information processing systems*, pp. 1177–1184, 2008. [3](#)
- Rasmussen, C. E. and Williams, C. K. I. *Gaussian Processes for Machine Learning*. MIT Press, 2006. URL <https://bit.ly/2tYpBix>. [2](#)
- Romero, P. A., Krause, A., and Arnold, F. H. Navigating the protein fitness landscape with gaussian processes. *Proceedings of the National Academy of Sciences*, 110(3):E193–E201, 2013. [1](#)
- Sen, R., Kandasamy, K., and Shakkottai, S. Multi-fidelity black-box optimization with hierarchical partitions. In *Proceedings of the 29th International Conference on Machine Learning (ICML)*, 2018. URL <https://bit.ly/2MSIsSQ>. [1](#), [3](#)
- Snelson, E. and Ghahramani, Z. Local and global sparse gaussian process approximations. In *Artificial Intelligence and Statistics*, pp. 524–531, 2007. [3](#)
- Snoek, J., Larochelle, H., and Adams, R. P. Practical bayesian optimization of machine learning algo-

- rithms. In *Advances in neural information processing systems*, pp. 2951–2959, 2012. 2
- Srinivas, N., Krause, A., Kakade, S., and Seeger, M. Gaussian process optimization in the bandit setting: No regret and experimental design. In *Proc. International Conference on Machine Learning (ICML)*, 2010. URL <https://bit.ly/2CNGPGc>. 2, 5, 6, 7
- Teh, Y.-W., Seeger, M., and Jordan, M. Semiparametric latent factor models. In *Artificial Intelligence and Statistics 10*, 2005. 3
- Titsias, M. Variational learning of inducing variables in sparse gaussian processes. In *Artificial Intelligence and Statistics*, pp. 567–574, 2009. 3
- van Gunsteren, W. F. and Berendsen, H. J. Computer simulation of molecular dynamics: Methodology, applications, and perspectives in chemistry. *Angewandte Chemie International Edition in English*, 29(9):992–1023, 1990. 1
- Wang, Z. and Jegelka, S. Max-value entropy search for efficient bayesian optimization. *arXiv preprint arXiv:1703.01968*, 2017. URL <https://arxiv.org/pdf/1703.01968.pdf>. 3, 5
- Wang, Z., Zhou, B., and Jegelka, S. Optimization as estimation with gaussian processes in bandit settings. In *Artificial Intelligence and Statistics*, pp. 1022–1031, 2016. URL <https://bit.ly/20eS0hp>. 3, 5
- Wilson, A. G., Knowles, D. A., and Ghahramani, Z. Gaussian process regression networks. In *Proceedings of the 29th International Conference on Machine Learning (ICML)*, 2012. 3
- Zhang, Y., Hoang, T. N., Low, B. K. H., and Kankanhalli, M. Information-based multi-fidelity bayesian optimization. *NIPS Workshop on Bayesian Optimization*, 2017. URL <https://bit.ly/2N5CdjH>. 3