
Gaussian Process Upper Confidence Bound Achieves Nearly-Optimal Regret in Noise-Free Gaussian Process Bandits

Shogo Iwazaki
LY Corporation
Tokyo, Japan
siwazaki@lycorp.co.jp

Abstract

We study the noise-free Gaussian Process (GP) bandit problem, in which a learner seeks to minimize regret through noise-free observations of a black-box objective function that lies in a known reproducing kernel Hilbert space (RKHS). The Gaussian Process Upper Confidence Bound (GP-UCB) algorithm is a well-known approach for GP bandits, where query points are adaptively selected based on the GP-based upper confidence bound score. While several existing works have reported the practical success of GP-UCB, its theoretical performance remains sub-optimal. However, GP-UCB often empirically outperforms other nearly-optimal noise-free algorithms that use non-adaptive sampling schemes. This paper resolves the gap between theoretical and empirical performance by establishing a nearly-optimal regret upper bound for noise-free GP-UCB. Specifically, our analysis provides the first constant cumulative regret bounds in the noise-free setting for both the squared exponential kernel and the Matérn kernel with some degree of smoothness.

1 Introduction

This paper studies the noise-free Gaussian Process (GP) bandit problem, where the learner seeks to minimize regret through noise-free observations of the black-box objective function. Several existing works tackle this problem, and some of them [Iwazaki and Takeno, 2025a, Salgia et al., 2024] propose algorithms whose regret nearly matches the lower bound of [Li and Scarlett, 2024]. For ease of theoretical analysis, these algorithms rely on the non-adaptive sampling scheme, whose query points are chosen independently of the observed function values, such as the uniform sampling [Salgia et al., 2024] or maximum variance reduction [Iwazaki and Takeno, 2025a]. Although the theoretical superiority of such non-adaptive algorithms is shown, in existing noisy setting literature [Bogunovic et al., 2022, Iwazaki and Takeno, 2025b, Li and Scarlett, 2022], their empirical performance has been reported to be worse than that of fully adaptive strategies such as GP upper confidence bound (GP-UCB) [Srinivas et al., 2010]. Unsurprisingly, we also observe such empirical and practical gaps in a noise-free setting, as shown in Figure 1. These observations suggest the possibility of further theoretical improvement in the practical fully adaptive algorithm. From this motivation, our work aims to establish the nearly-optimal regret for GP-UCB, which is one of the well-known adaptive GP bandit algorithms, and its existing guarantees only show strictly sub-optimal regret in a noise-free setting [Kim and Sanz-Alonso, 2024, Lyu et al., 2019].

Contributions. Our contributions are summarized below:

- We give a refined regret analysis of GP-UCB (Theorems 1 and 2), which matches both the cumulative regret lower bounds of [Li and Scarlett, 2024] and the simple regret lower bound

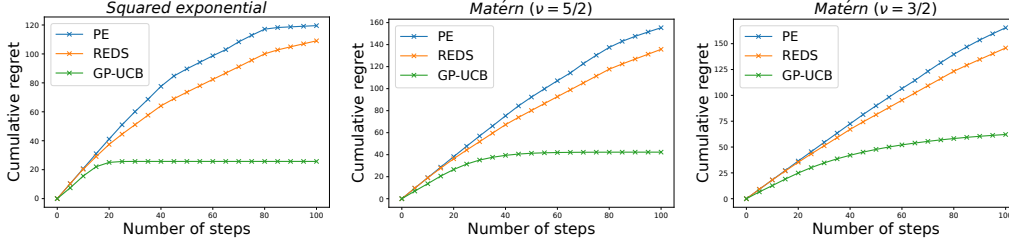


Figure 1: Empirical performance comparison of GP-UCB and two existing nearly-optimal algorithms: random exploration with domain shrinking (REDS) [Salgia et al., 2024] and phased elimination (PE) [Iwazaki and Takeno, 2025a]. From left to right, the plots show the average cumulative regret over 3000 independent runs under the squared exponential kernel, Matérn kernel ($\nu = 5/2$), and Matérn kernel ($\nu = 3/2$) with $d = 2$, respectively. Here, d and ν represent the dimension of the input and the smoothness parameter, respectively. Detailed experimental settings are provided in Appendix B.

of [Bull, 2011] up to polylogarithmic factors in the Matérn kernel. Regarding cumulative regret, our analysis shows that GP-UCB achieves the constant $O(1)$ regret under the squared exponential and Matérn kernel with $d > \nu$. The results are summarized in Tables 1 and 2.

- Our key theoretical contribution is the new algorithm-independent upper bounds for the observed posterior standard deviations (Lemmas 3–5) by bridging the information gain-based analysis in the noisy regime to the noise-free setting. Furthermore, as discussed in Section 3, these results have the potential to translate existing confidence bound-based algorithms for noisy settings into nearly-optimal noise-free variants beyond the analysis of GP-UCB.

Related works. Various existing works study the theory for the noisy GP bandits [Chowdhury and Gopalan, 2017, Li and Scarlett, 2022, Scarlett et al., 2017, Srinivas et al., 2010, Valko et al., 2013]. Regarding noise-free settings, to our knowledge, [Bull, 2011] is the first work that shows both the upper bound and the lower bound for simple regret via the expected improvement (EI) strategy. After that, the analysis of the cumulative regret is shown in [Lyu et al., 2019] with GP-UCB. Recently, Vakili [2022] conjectures the lower bound of the cumulative regret in the noise-free setting, suggesting a superior algorithm exists in the noise-free setting. Later, the following work [Li and Scarlett, 2024] formally validates the conjectured lower bound of [Vakili, 2022]¹. Motivated by the lower bounds, several works [Flynn and Reeb, 2024, Kim and Sanz-Alonso, 2024, Iwazaki and Takeno, 2025a, Salgia et al., 2024] studied the improved algorithm to achieve superior regret to the result of [Lyu et al., 2019]. Although some of them propose the nearly-optimal algorithms [Iwazaki and Takeno, 2025a, Salgia et al., 2024], their algorithms are based on a non-adaptive sampling scheme, whose inferior performance has been reported in the existing work [Bogunovic et al., 2022, Iwazaki and Takeno, 2025b, Li and Scarlett, 2022]. Here, our motivation is to establish the nearly-optimal theory under the fully adaptive nature of GP-UCB; however, our proof relates to the analysis in [Iwazaki and Takeno, 2025a] for the non-adaptive maximum variance reduction algorithm. Their analysis includes the noise-free setting as a special case and provides nearly-optimal regret under broader varying noise variance settings. However, its applicability is limited to the maximum variance reduction algorithm. Our analysis can be interpreted as a refined version of that in [Iwazaki and Takeno, 2025a] by focusing on the noise-free setting. Finally, although our paper studies the frequentist assumption that the underlying function is fixed, our core results (Lemmas 3–5) are also applicable to the regret analysis in the Bayesian setting, whose underlying function is drawn from a known GP [De Freitas et al., 2012, Grünewälder et al., 2010, Russo and Van Roy, 2014a,b, Scarlett, 2018, Srinivas et al., 2010].

¹Although Li and Scarlett [2024] considers the cascading structure of the observation process, the standard noise-free setting is the special case of their setting.

Table 1: Comparison between existing noise-free algorithms’ guarantees for cumulative regret and our result (adapted from [Iwazaki and Takeno, 2025a]). As with the table in [Iwazaki and Takeno, 2025a], the smoothness parameter ν of the Matérn kernel and $\alpha > 0$ in PE are assumed to be $\nu > 1/2$ and an arbitrary fixed constant, respectively. Furthermore, all the parameters (d , ℓ , ν , and B) except for T are assumed to be $\Theta(1)$. “Type” column shows that the regret guarantee is (D)eterministic or (P)robabilistic. Here, a regret bound is labeled “deterministic” if it holds for all possible realizations of the inputs, without relying on probabilistic assumptions. Here, $\tilde{O}(\cdot)$ hides polylogarithmic factors in T .

Algorithm	Regret (SE)	Regret (Matérn)			Type
		$\nu < d$	$\nu = d$	$\nu > d$	
GP-UCB [Lyu et al., 2019] [Kim and Sanz-Alonso, 2024]	$O(\sqrt{T \ln^d T})$	$\tilde{O}\left(T^{\frac{\nu+d}{2\nu+d}}\right)$			D
Explore-then-Commit [Vakili, 2022]	N/A	$\tilde{O}\left(T^{\frac{d}{\nu+d}}\right)$			P
Kernel-AMM-UCB [Flynn and Reeb, 2024]	$O(\ln^{d+1} T)$	$\tilde{O}\left(T^{\frac{\nu d + d^2}{2\nu^2 + 2\nu d + d^2}}\right)$			D
REDS [Salgia et al., 2024]	N/A	$\tilde{O}\left(T^{\frac{d-\nu}{d}}\right)$	$O\left(\ln^{\frac{5}{2}} T\right)$	$O\left(\ln^{\frac{3}{2}} T\right)$	P
PE [Iwazaki and Takeno, 2025a]	$O(\ln T)$	$\tilde{O}\left(T^{\frac{d-\nu}{d}}\right)$	$O\left(\ln^{2+\alpha} T\right)$	$O(\ln T)$	D
GP-UCB (Our analysis)	$O(1)$	$\tilde{O}\left(T^{\frac{d-\nu}{d}}\right)$	$O(\ln^2 T)$	$O(1)$	D
Conjectured Lower Bound [Vakili, 2022]	N/A	$\Omega\left(T^{\frac{d-\nu}{d}}\right)$	$\Omega(\ln T)$	$\Omega(1)$	N/A
Lower Bound [Li and Scarlett, 2024]	N/A	$\Omega\left(T^{\frac{d-\nu}{d}}\right)$	$\Omega(1)$	$\Omega(1)$	N/A

2 Preliminaries

Noise-free GP bandit problem. We study the GP bandit problem under noise-free observations. Let $\mathcal{X} \subset \mathbb{R}^d$ be a compact input domain, and consider a black-box objective function $f : \mathcal{X} \rightarrow \mathbb{R}$ that can only be evaluated point-wise. At each step $t \in \mathbb{N}_+$, the learner selects a query point $\mathbf{x}_t \in \mathcal{X}$ and observes its function value $f(\mathbf{x}_t)$. After T steps, the performance of the learner is measured using either the cumulative regret R_T or the simple regret r_T , which are respectively defined as follows:

$$R_T = \sum_{t=1}^T f(\mathbf{x}^*) - f(\mathbf{x}_t), \quad (1)$$

$$r_T = f(\mathbf{x}^*) - f(\hat{\mathbf{x}}_T). \quad (2)$$

Here, $\mathbf{x}^* \in \arg \max_{\mathbf{x} \in \mathcal{X}} f(\mathbf{x})$ is a global maximizer, and $\hat{\mathbf{x}}_T$ denotes the estimated maximizer returned by the algorithm at the end of the final step T .

Gaussian process model. To construct the GP-bandit algorithm, the GP model [Rasmussen and Williams, 2005] plays a central role in balancing the trade-off between exploration and exploitation. First, we adopt a Gaussian process prior with zero mean and covariance (kernel) function $k : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$. Then, given a sequence of queried points $\mathbf{X}_t = (\mathbf{x}_1, \dots, \mathbf{x}_t)$ and their corresponding evaluations $\mathbf{f}(\mathbf{X}_t) = (f(\mathbf{x}_1), \dots, f(\mathbf{x}_t))$, the posterior mean $\mu(\mathbf{x}; \mathbf{X}_t)$ and variance $\sigma^2(\mathbf{x}; \mathbf{X}_t)$ of $f(\mathbf{x})$ at a new point $\mathbf{x} \in \mathcal{X}$ are

$$\mu(\mathbf{x}; \mathbf{X}_t) = \mathbf{k}(\mathbf{x}, \mathcal{E}(\mathbf{X}_t))^\top \mathbf{K}(\mathcal{E}(\mathbf{X}_t), \mathcal{E}(\mathbf{X}_t))^{-1} \mathbf{f}(\mathcal{E}(\mathbf{X}_t)), \quad (3)$$

$$\sigma^2(\mathbf{x}; \mathbf{X}_t) = k(\mathbf{x}, \mathbf{x}) - \mathbf{k}(\mathbf{x}, \mathcal{E}(\mathbf{X}_t))^\top \mathbf{K}(\mathcal{E}(\mathbf{X}_t), \mathcal{E}(\mathbf{X}_t))^{-1} \mathbf{k}(\mathbf{x}, \mathcal{E}(\mathbf{X}_t)), \quad (4)$$

where $\mathbf{k}(\mathbf{x}, \mathcal{E}(\mathbf{X}_t)) = [k(\tilde{\mathbf{x}}, \mathbf{x})]_{\tilde{\mathbf{x}} \in \mathcal{E}(\mathbf{X}_t)}$ is the kernel vector, $\mathbf{K}(\mathcal{E}(\mathbf{X}_t), \mathcal{E}(\mathbf{X}_t))$ is the Gram matrix, and $\mathbf{f}(\mathcal{E}(\mathbf{X}_t)) = [f(\tilde{\mathbf{x}})]_{\tilde{\mathbf{x}} \in \mathcal{E}(\mathbf{X}_t)}$. In the above definition, $\mathcal{E}(\mathbf{X}_t)$ denotes the subset of $\mathbf{X}_t$ obtained by removing any fully correlated inputs (i.e., with zero posterior variance) with previous ones. Namely, we define $\mathcal{E}(\mathbf{X}_t)$ inductively as $\mathcal{E}(\mathbf{X}_t) = \mathcal{E}(\mathbf{X}_{t-1}) \cup \{\mathbf{x}_t\}$ if $\sigma^2(\mathbf{x}_t; \mathbf{X}_{t-1}) > 0$; otherwise,

Table 2: Comparison between existing noiseless algorithms' guarantees for simple regret and our result (adapted from [Iwazaki and Takeno, 2025a]). In the regrets of GP-UCB+, EXPLOIT+, MVR, and GP-UCB, $\alpha > 0$ and $C > 0$ are arbitrary fixed constants and some positive constants, respectively.

Algorithm	Regret (SE)	Regret (Matérn)	Type
GP-EI [Bull, 2011]	N/A	$\tilde{O}\left(T^{-\frac{\min\{1, \nu\}}{d}}\right)$	D
GP-EI with ϵ -Greedy [Bull, 2011]	N/A	$\tilde{O}\left(T^{-\frac{\nu}{d}}\right)$	P
GP-UCB [Lyu et al., 2019] [Kim and Sanz-Alonso, 2024]	$O\left(\sqrt{\frac{\ln^d T}{T}}\right)$	$\tilde{O}\left(T^{-\frac{\nu}{2\nu+d}}\right)$	D
Kernel-AMM-UCB [Flynn and Reeb, 2024]	$O\left(\frac{\ln^{d+1} T}{T}\right)$	$\tilde{O}\left(T^{-\frac{\nu d + 2\nu^2}{2\nu^2 + 2\nu d + d^2}}\right)$	D
GP-UCB+, EXPLOIT+ [Kim and Sanz-Alonso, 2024]	$O\left(\exp\left(-CT^{\frac{1}{d}-\alpha}\right)\right)$	$O\left(T^{-\frac{\nu}{d}+\alpha}\right)$	P
MVR [Iwazaki and Takeno, 2025a]	$O\left(\exp\left(-\frac{1}{2}T^{\frac{1}{d+1}}\ln^{-\alpha} T\right)\right)$	$\tilde{O}\left(T^{-\frac{\nu}{d}}\right)$	D
GP-UCB (Our analysis)	$O\left(\sqrt{T}\exp\left(-\frac{1}{2}CT^{\frac{1}{d+1}}\right)\right)$	$\tilde{O}\left(T^{-\frac{\nu}{d}}\right)$	D
Lower Bound [Bull, 2011]	N/A	$\Omega\left(T^{-\frac{\nu}{d}}\right)$	N/A

$\mathcal{E}(\mathbf{X}_t) = \mathcal{E}(\mathbf{X}_{t-1})$. Here, we define $\mathcal{E}(\mathbf{X}_1) = \mathbf{X}_1$. Note that if there are no duplications in the input sequence $\mathbf{x}_1, \dots, \mathbf{x}_t$, then $\mathcal{E}(\mathbf{X}_t) = \mathbf{X}_t$ holds under commonly used kernel (covariance) functions, such as squared exponential and Matérn kernels (precisely defined in the next paragraph). Furthermore, for the ease of notation, we set $\mu(\mathbf{x}; \mathbf{X}) = 0$ and $\sigma^2(\mathbf{x}; \mathbf{X}) = k(\mathbf{x}, \mathbf{x})$ for $\mathbf{X} = \emptyset$.

Kernel function and information gain. Regarding the choice of the kernel function, we focus on the squared exponential (SE) kernel

$$k_{\text{SE}}(\mathbf{x}, \tilde{\mathbf{x}}) = \exp\left(-\frac{\|\mathbf{x} - \tilde{\mathbf{x}}\|_2^2}{2\ell^2}\right), \quad (5)$$

and the Matérn kernel

$$k_{\text{Matérn}}(\mathbf{x}, \tilde{\mathbf{x}}) = \frac{2^{1-\nu}}{\Gamma(\nu)} \left(\frac{\sqrt{2\nu} \|\mathbf{x} - \tilde{\mathbf{x}}\|_2}{\ell}\right)^\nu J_\nu\left(\frac{\sqrt{2\nu} \|\mathbf{x} - \tilde{\mathbf{x}}\|_2}{\ell}\right), \quad (6)$$

where $\ell > 0$ and $\nu > 0$ are the lengthscale and smoothness parameter, respectively. Furthermore, J_ν and Γ are the modified Bessel and Gamma functions, respectively. These two kernels are commonly used and analyzed in GP-bandits [Scarlett et al., 2017, Srinivas et al., 2010]. The convergence rate of the function estimation in GP regression depends on the choice of kernel, which in turn affects the resulting regret upper bound. Therefore, to capture the problem complexity in kernel-dependent manner, the following kernel-dependent information theoretic quantity $\gamma_T(\lambda^2)$ is often employed in the analysis of GP bandits:

$$\gamma_T(\lambda^2) = \sup_{\mathbf{x}_1, \dots, \mathbf{x}_T \in \mathcal{X}} \frac{1}{2} \ln \det(\mathbf{I}_T + \lambda^{-2} \mathbf{K}(\mathbf{X}_T, \mathbf{X}_T)), \quad (7)$$

where $\lambda > 0$ and $\mathbf{I}_T$ are any positive parameter and $T \times T$ -identity matrix, respectively. The quantity $\gamma_T(\lambda^2)$ is called *maximum information gain* (MIG) [Srinivas et al., 2010] since the quantity $\frac{1}{2} \ln \det(\mathbf{I}_T + \lambda^{-2} \mathbf{K}(\mathbf{X}_T, \mathbf{X}_T))$ represents the mutual information between the underlying function f and training outputs under the noisy-GP model with variance parameter λ^2 . The increasing speed of MIG is analyzed in several commonly used kernels. For example, $\gamma_T(\lambda^2) = O(\ln^{d+1}(T/\lambda^2))$ and $\gamma_T(\lambda^2) = \tilde{O}((T/\lambda^2)^{\frac{d}{2\nu+d}})$ under $k = k_{\text{SE}}$ and $k = k_{\text{Matérn}}$ with $\nu > 1/2$, respectively [Vakili et al., 2021c]².

²These orders hold as $T \rightarrow \infty, \lambda \rightarrow 0$.

Algorithm 1 Gaussian process upper confidence bound (GP-UCB) for noise-free setting.

Require: Compact input domain $\mathcal{X} \subset \mathbb{R}^d$, Kernel function k , and RKHS norm upper bound $B \in (0, \infty)$.

- 1: $\mathbf{X}_0 \leftarrow \emptyset, \beta^{1/2} \leftarrow B$.
 - 2: **for** $t = 1, 2, \dots$ **do**
 - 3: $\mathbf{x}_t \leftarrow \arg \max_{\mathbf{x} \in \mathcal{X}} \mu(\mathbf{x}; \mathbf{X}_{t-1}) + \beta^{1/2} \sigma(\mathbf{x}; \mathbf{X}_{t-1})$.
 - 4: Observe $f(\mathbf{x}_t)$ and update the posterior mean and variance.
 - 5: **end for**
-

Regularity assumption. It is hopeless to derive meaningful guarantees without further assumptions about f . To obtain a valid estimate of f through the GP-model, we assume that f lies in the known reproducing kernel Hilbert space (RKHS) [Aronszajn, 1950] corresponding to the kernel k and has bounded norm. The formal description is given below.

Assumption 1. *The objective function f lies in the reproducing kernel Hilbert space (RKHS) associated with a known positive definite kernel $k : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$. We assume that $k(\mathbf{x}, \mathbf{x}) \leq 1$ for all $\mathbf{x} \in \mathcal{X}$, and that the RKHS norm $\|f\|_k$ satisfies $\|f\|_k \leq B < \infty$.*

This is a standard assumption in the GP-bandit literature [Chowdhury and Gopalan, 2017, Scarlett et al., 2017, Srinivas et al., 2010], and is leveraged to derive the confidence bound of f [Kanagawa et al., 2018, Lyu et al., 2019, Vakili et al., 2021a]. Here, RKHS associated with kernel k is given as the closure of the linear span: $\mathcal{H}_k^{(\text{pre})} := \{\sum_{i=1}^n c_i k(\mathbf{x}^{(i)}, \cdot) \mid n \in \mathbb{N}_+, c_1, \dots, c_n \in \mathbb{R}, \mathbf{x}^{(1)}, \dots, \mathbf{x}^{(n)} \in \mathcal{X}\}$ [Kanagawa et al., 2018]. Therefore, intuitively, under Assumption 1, we can interpret that the basic properties of the objective function f , such as continuity and smoothness, are encoded through the choice of the kernel function k .

Gaussian process upper confidence bound. GP-UCB [Srinivas et al., 2010] is a widely used algorithm for the noisy GP bandit setting. Its noiseless variant was proposed by Lyu et al. [2019], and is outlined in Algorithm 1. Their analysis largely follows the framework developed for the noisy case in [Srinivas et al., 2010], but differs in the confidence bound, which is strictly tighter than that in the noisy setting. Although this refinement of the confidence bound leads to superior regret compared with the noisy setting, the resulting regret is still strictly sub-optimal in the noise-free setting. This fact suggests that we require the other fundamental modification from the proof in [Srinivas et al., 2010].

3 Refined Regret Upper Bound for Noise-Free GP-UCB

The following theorem describes our main results, which show the nearly-optimal regret upper bound for GP-UCB.

Theorem 1 (Refined cumulative regret upper bound for GP-UCB). *Fix any compact input domain $\mathcal{X} \subset \mathbb{R}^d$. Suppose B , d , ℓ , and ν are fixed constants. Then, when running Algorithm 1 under Assumption 1, the following two statements hold for any $T \in \mathbb{N}_+$:*

- If $k = k_{\text{SE}}$, the regret R_T satisfies $R_T = O(1)$.
- If $k = k_{\text{Matérn}}$ with $\nu > 1/2$, the regret R_T satisfies

$$R_T = \begin{cases} \tilde{O}\left(T^{\frac{d-\nu}{d}}\right) & \text{if } d > \nu, \\ O(\ln^2 T) & \text{if } d = \nu, \\ O(1) & \text{if } d < \nu. \end{cases} \quad (8)$$

The implied constants may depend on B , d , ℓ , ν , and the diameter of $\mathcal{X}$.

Theorem 2 (Refined simple regret upper bound for GP-UCB). *Fix any compact input domain $\mathcal{X} \subset \mathbb{R}^d$. Suppose B , L , d , ℓ , and ν are fixed constants. Then, when running Algorithm 1 under Assumption 1, the following two statements hold by setting the estimated maximizer $\hat{\mathbf{x}}_T$ as $\hat{\mathbf{x}}_T \in \arg \max_{\mathbf{x} \in \{\mathbf{x}_1, \dots, \mathbf{x}_T\}} f(\mathbf{x})$:*

- If $k = k_{\text{SE}}$, $r_T = O\left(\sqrt{T} \exp\left(-\frac{1}{2}CT^{\frac{1}{d+1}}\right)\right)$.
- If $k = k_{\text{Matérn}}$ with $\nu > 1/2$, $r_T = \tilde{O}\left(T^{-\frac{\nu}{d}}\right)$.

The implied constants depend on B , d , ℓ , ν , and the diameter of $\mathcal{X}$. Furthermore, the constant $C > 0$ depends on d , ℓ , ν , and the diameter of $\mathcal{X}$ ³.

Proof sketch. Our key technical results are the new analysis of the cumulative posterior standard deviation $\sum_{t=1}^T \sigma(\mathbf{x}_t; \mathbf{X}_{t-1})$ and its minimum $\min_{t \in [T]} \sigma(\mathbf{x}_t; \mathbf{X}_{t-1})$, which plays an important role in the theoretical analysis of GP bandits. Indeed, following the standard analysis of GP-UCB, we have the following upper bounds of regrets by combining the UCB-selection rule with the existing noise-free confidence bound (e.g., Lemma 11 in [Lyu et al., 2019] or Proposition 1 in [Vakili et al., 2021a]):

$$R_T = \sum_{t=1}^T f(\mathbf{x}^*) - f(\mathbf{x}_t) \leq 2B \sum_{t=1}^T \sigma(\mathbf{x}_t; \mathbf{X}_{t-1}), \quad (9)$$

$$r_T = \min_{t \in [T]} f(\mathbf{x}^*) - f(\mathbf{x}_t) \leq 2B \min_{t \in [T]} \sigma(\mathbf{x}_t; \mathbf{X}_{t-1}). \quad (10)$$

From the above inequalities, we observe that the tighter upper bounds of $\sum_{t=1}^T \sigma(\mathbf{x}_t; \mathbf{X}_{t-1})$ and $\min_{t \in [T]} \sigma(\mathbf{x}_t; \mathbf{X}_{t-1})$ directly yield the tighter regret upper bounds of GP-UCB. Lemma 3 below is our main technical contribution, which gives the refined upper bounds of $\sum_{t=1}^T \sigma(\mathbf{x}_t; \mathbf{X}_{t-1})$ and $\min_{t \in [T]} \sigma(\mathbf{x}_t; \mathbf{X}_{t-1})$.

Lemma 3 (Posterior standard deviation upper bound for SE and Matérn kernel). *Fix any compact input domain $\mathcal{X} \subset \mathbb{R}^d$, and kernel function $k : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$ that satisfies $k(\mathbf{x}, \mathbf{x}) \leq 1$ for all $\mathbf{x} \in \mathcal{X}$. Then, the following statements hold for any $T \in \mathbb{N}_+$ and any input sequence $\mathbf{x}_1, \dots, \mathbf{x}_T \in \mathcal{X}$:*

- For $k = k_{\text{SE}}$, we have

$$\min_{t \in [T]} \sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) = O\left(\sqrt{T} \exp\left(-\frac{1}{2}CT^{\frac{1}{d+1}}\right)\right) \quad \text{and} \quad \sum_{t=1}^T \sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) = O(1). \quad (11)$$

- For $k = k_{\text{Matérn}}$ with $\nu > 1/2$, we have

$$\min_{t \in [T]} \sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) = O\left(T^{-\frac{\nu}{d}} \ln^{\frac{\nu}{d}} T\right), \quad (12)$$

$$\sum_{t=1}^T \sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) = \begin{cases} O\left(T^{\frac{d-\nu}{d}} \ln^{\frac{\nu}{d}} T\right) & \text{if } d > \nu, \\ O\left(\ln^2 T\right) & \text{if } d = \nu, \\ O(1) & \text{if } d < \nu. \end{cases} \quad (13)$$

The constant $C > 0$ and the implied constants depend on d , ℓ , ν , and the diameter of $\mathcal{X}$.

The full proof of Lemma 3 is given in Appendix A.2. We will also provide the proof sketch in the next section. Combining the above equations with Eqs. (9) and (10), we obtain the statements in Theorems 1 and 2.

³As described in Section 3.1, the constant C arises from the implied constants in the upper bound of MIG [Vakili et al., 2021c], which depends on d , ℓ , ν , and the diameter of $\mathcal{X}$.

⁴Our results (Theorems 1 and 2) heavily rely on the upper bound of the MIG provided in [Vakili et al., 2021c], which relies on the uniform boundness assumption of the eigenfunctions of the kernel. The validity of the uniform boundness assumption is doubted by Janz [2022] under a general compact input domain $\mathcal{X}$. Although our results are based on the upper bound of MIG in [Vakili et al., 2021c], our proof strategy is also applicable for deriving nearly-optimal regrets based on the recent analysis of MIG [Iwazaki, 2025] without the uniform boundness assumption. Specifically, then, the orders of the resulting regrets are the same as those in Theorems 1 and 2 for k_{SE} and $k_{\text{Matérn}}$ with $d < \nu$. As for the case under $k = k_{\text{Matérn}}$ with $d \geq \nu$, we can also obtain $R_T = \tilde{O}(T^{(d-\nu)/d})$ and $r_T = \tilde{O}(T^{-\nu/d})$, while they suffer from additional polylogarithmic terms.

Relation to the existing research in noise-free setting. Lemma 3 resolves the open problem raised by the existing noise-free setting literature [Li and Scarlett, 2024, Vakili, 2022]. First, Vakili [2022] conjectured that the quantity $\sum_{t=1}^T \sigma(\mathbf{x}_t; \mathbf{X}_{t-1})$ under $k = k_{\text{Matérn}}$ can attain the following upper bound:

$$\sum_{t=1}^T \sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) = \begin{cases} O\left(T^{\frac{d-\nu}{d}}\right) & \text{if } d > \nu, \\ O(\ln T) & \text{if } d = \nu, \\ O(1) & \text{if } d < \nu. \end{cases} \quad (14)$$

Although the above conjecture is partially validated under a specific non-adaptive algorithm [Li and Scarlett, 2024, Iwazaki and Takeno, 2025a, Salgia et al., 2024], the correctness of this conjecture under a general algorithm has been an open problem [Li and Scarlett, 2024]. Lemma 3 answers this open problem with Eq. (13), which matches conjectured upper bound up to polylogarithmic factors.

Generality of Lemma 3. We would like to highlight that Lemma 3 always holds for any input sequence, in contrast to the existing algorithm-specific upper bounds [Iwazaki and Takeno, 2025a, Salgia et al., 2024]. Since the existing noisy GP bandits theory often leverage the upper bound of $\min_{t \in [T]} \sigma(\mathbf{x}_t; \mathbf{X}_{t-1})$ or $\sum_{t=1}^T \sigma(\mathbf{x}_t; \mathbf{X}_{t-1})$ from [Srinivas et al., 2010], we expect that many existing theoretical results in the noisy setting can be extended to the corresponding noise-free setting by directly replacing the existing noisy upper bounds of [Srinivas et al., 2010] with Lemma 3. For example, the analysis for GP-Thompson sampling (GP-TS) [Chowdhury and Gopalan, 2017], GP-UCB and GP-TS under Bayesian setting [Russo and Van Roy, 2014a, Srinivas et al., 2010], contextual setting [Krause and Ong, 2011], GP-based level-set estimation [Gotovos et al., 2013], multi-objective setting [Zuluaga et al., 2016], robust formulation [Bogunovic et al., 2018], and so on.

Extension to other kernel functions. Lemma 3 is limited to the SE and Matérn kernels. However, our proof strategy in Section 3.1 can be applied to any kernel function if we know the joint dependence of T and the noise-variance parameter λ^2 in MIG. For example, we can apply our proof strategy for the neural tangent kernel (NTK) [Jacot et al., 2018] by leveraging the existing upper bound of MIG under NTK [Iwazaki and Suzumura, 2024, Kassraie and Krause, 2022, Kassraie et al., 2022, Vakili et al., 2021b]. We expect that such an extension will benefit the theoretical guarantees of neural network-based bandit algorithms under a noise-free setting and will be one of the interesting research directions.

3.1 Proof Sketch of Lemma 3

In this subsection, we describe the proof sketch of Lemma 3, while we give its full proof in Appendix A.2. Below, to prove Lemma 3, we consider the more general lemmas, which bridge the MIG to noise-free posterior standard deviations.

Lemma 4 (General upper bound for the minimum posterior standard deviation). *Fix any input domain $\mathcal{X}$ and any $\bar{T} \geq 2$. Let $(\lambda_t)_{t \geq \bar{T}}$ be a strictly positive sequence such that $\gamma_t(\lambda_t^2) \leq (t-1)/3$ for all $t \geq \bar{T}$. Then, $\min_{t \in [T]} \sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) \leq \lambda_T$ holds for any $T \geq \bar{T}$ and any sequence $\mathbf{x}_1, \dots, \mathbf{x}_T \in \mathcal{X}$.*

Lemma 5 (General upper bound for the cumulative posterior standard deviations). *Fix any input domain $\mathcal{X}$, any $\bar{T} \geq 2$, and any kernel function $k : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$ that satisfies $k(\mathbf{x}, \mathbf{x}) \leq 1$ for all $\mathbf{x} \in \mathcal{X}$. Let $(\lambda_t)_{t \geq \bar{T}}$ be a strictly positive sequence such that $\gamma_t(\lambda_t^2) \leq (t-1)/3$ for all $t \geq \bar{T}$. Then, the following inequality holds for any $T \in \mathbb{N}_+$ and any sequence $\mathbf{x}_1, \dots, \mathbf{x}_T \in \mathcal{X}$:*

$$\sum_{t=1}^T \sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) \leq \bar{T} - 1 + \sum_{t=\bar{T}}^T \lambda_t. \quad (15)$$

These lemmas hold for any kernel k satisfying $k(\mathbf{x}, \mathbf{x}) \leq 1$, which includes most standard kernels after normalization. Roughly speaking, the above lemmas suggest that $\min_{t \in [T]} \sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) \lesssim \lambda_T$ and $\sum_{t=1}^T \sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) \lesssim \sum_{t=\bar{T}}^T \lambda_t$ holds as far as the corresponding MIG $\gamma_t(\lambda_t^2)$ does not increase super-linearly. Note that the MIG $\gamma_t(\lambda_t^2)$ monotonically increases as λ_t^2 decreases, which implies that the tightest upper bound is obtained by setting λ_t^2 as $\gamma_t(\lambda_t^2) = (t-1)/3$. By relying on the existing upper bound of MIG [Vakili et al., 2021c], we can confirm that the condition $\forall t \geq \bar{T}, \gamma_t(\lambda_t^2) \leq (t-1)/3$

of the above lemmas holds with $\lambda_t^2 = O(t \exp(-Ct^{\frac{1}{d+1}}))$ and $\lambda_t^2 = O(t^{-\frac{2\nu}{d}} (\ln t)^{\frac{2\nu}{d}})$ for $k = k_{\text{SE}}$ and $k = k_{\text{Matérn}}$, respectively. Here, the constants C and $\bar{T}$ are determined based on the implied constant of the upper bound of MIG. See Appendix A.2 for details. Lemma 3 follows from the aforementioned setting of λ_t^2 , and Lemmas 4 and 5. Below, we give the proofs for Lemmas 4 and 5.

Proof of Lemma 4. Instead of directly treating the noise-free posterior standard deviation, we study its upper bound with the posterior standard deviation of some noisy GP model. Here, let us denote $\sigma_{\lambda^2}^2(\mathbf{x}; \mathbf{X}_{t-1})$ as the posterior variance under the noisy GP-model with the strictly positive variance parameter $\lambda^2 > 0$, which is defined as

$$\sigma_{\lambda^2}^2(\mathbf{x}; \mathbf{X}_{t-1}) = k(\mathbf{x}, \mathbf{x}) - \mathbf{k}(\mathbf{x}, \mathbf{X}_{t-1})^\top [\mathbf{K}(\mathbf{X}_{t-1}, \mathbf{X}_{t-1}) + \lambda^2 \mathbf{I}_{t-1}]^{-1} \mathbf{k}(\mathbf{x}, \mathbf{X}_{t-1}). \quad (16)$$

Since the posterior variance is monotonic for the variance parameter, we have $\sigma^2(\mathbf{x}_t; \mathbf{X}_{t-1}) \leq \sigma_{\lambda_T^2}^2(\mathbf{x}_t; \mathbf{X}_{t-1})$ for all $t \in [T]$. Next, we obtain the upper bound of $\sigma_{\lambda_T^2}^2(\mathbf{x}_t; \mathbf{X}_{t-1})$ based on the following lemma, which is the main component of the proof of Lemmas 4 and 5.

Lemma 6 (Elliptical potential count lemma, Lemma D.9 in [Flynn and Reeb, 2024] or Lemma 3.3 in [Iwazaki and Takeda, 2025a]). *Fix any $T \in \mathbb{N}_+$, any sequence $\mathbf{x}_1, \dots, \mathbf{x}_T \in \mathcal{X}$, and $\lambda > 0$. Define $\mathcal{T}$ as $\mathcal{T} = \{t \in [T] \mid \lambda^{-1} \sigma_{\lambda^2}^2(\mathbf{x}_t; \mathbf{X}_{t-1}) > 1\}$, where $\mathbf{X}_{t-1} = (\mathbf{x}_1, \dots, \mathbf{x}_{t-1})$. Then, the number of elements of $\mathcal{T}$ satisfies $|\mathcal{T}| \leq 3\gamma_T(\lambda^2)$.*

The above lemma implies that the set $\mathcal{T}^c := \{t \in [T] \mid \sigma_{\lambda_T^2}^2(\mathbf{x}_t; \mathbf{X}_{t-1}) \leq \lambda_T\}$ satisfies $|\mathcal{T}^c| = |[T] \setminus \mathcal{T}| \geq T - 3\gamma_T(\lambda_T^2)$. Therefore, for any $T \geq \bar{T}$, $|\mathcal{T}^c| \geq 1$ holds from the condition $\gamma_T(\lambda_T^2) \leq (T - 1)/3$. This implies there exists some $\tilde{t} \in [T]$ such that $\sigma(\mathbf{x}_{\tilde{t}}; \mathbf{X}_{\tilde{t}-1}) \leq \sigma_{\lambda_T^2}^2(\mathbf{x}_{\tilde{t}}; \mathbf{X}_{\tilde{t}-1}) \leq \lambda_T$; therefore, $\min_{t \in [T]} \sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) \leq \sigma(\mathbf{x}_{\tilde{t}}; \mathbf{X}_{\tilde{t}-1}) \leq \lambda_T$ holds for all $T \geq \bar{T}$. $\square$

Proof of Lemma 5. Overall, the proof strategy of this lemma is to repeatedly apply the proof of Lemma 4 by leveraging the monotonicity of the posterior variance against training inputs. First, if $T < \bar{T}$, Eq. (15) is clearly holds from the assumption $\forall \mathbf{x} \in \mathcal{X}, k(\mathbf{x}, \mathbf{x}) \leq 1$. Hereafter, we focus on $T \geq \bar{T}$. By following the same argument of Lemma 4, we can confirm that there exists the index $\tilde{t}_T \leq T$ such that $\sigma(\mathbf{x}_{\tilde{t}_T}; \mathbf{X}_{\tilde{t}_T-1}) \leq \sigma_{\lambda_T^2}^2(\mathbf{x}_{\tilde{t}_T}; \mathbf{X}_{\tilde{t}_T-1}) \leq \lambda_T$. Here, we define the new sequence $(\mathbf{x}_t^{(T-1)})_{t \in [T-1]}$ as the sequence that $\mathbf{x}_{\tilde{t}_T}$ is eliminated from $(\mathbf{x}_t)_{t \in [T]}$; namely, we set $\mathbf{x}_t^{(T-1)} = \mathbb{1}\{t < \tilde{t}_T\} \mathbf{x}_t + \mathbb{1}\{t \geq \tilde{t}_T\} \mathbf{x}_{t+1}$ for any $t \in [T-1]$. Furthermore, we define $\mathbf{X}_t^{(T-1)} = (\mathbf{x}_1^{(T-1)}, \dots, \mathbf{x}_t^{(T-1)})$. From this construction of $\mathbf{X}_t^{(T-1)}$, we can observe the following two facts:

- For any $t < \tilde{t}_T$, we have $\sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) = \sigma(\mathbf{x}_t^{(T-1)}; \mathbf{X}_{t-1}^{(T-1)})$, since $\mathbf{x}_t^{(T-1)} = \mathbf{x}_t$ and $\mathbf{X}_{t-1}^{(T-1)} = \mathbf{X}_{t-1}$.
- For any $t > \tilde{t}_T$, we have $\sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) \leq \sigma(\mathbf{x}_t^{(T-1)}; \mathbf{X}_{t-2}^{(T-1)})$, since $\mathbf{x}_t^{(T-1)} = \mathbf{x}_t$ and $\mathbf{X}_{t-2}^{(T-1)} \subset \mathbf{X}_{t-1}$ from the definition of $\mathbf{X}_t^{(T-1)}$.

From the above two facts, we have

$$\sum_{t=1}^T \sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) \leq \sum_{t \in [T] \setminus \{\tilde{t}_T\}} \sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) + \lambda_T \leq \sum_{t \in [T-1]} \sigma(\mathbf{x}_t^{(T-1)}; \mathbf{X}_{t-1}^{(T-1)}) + \lambda_T. \quad (17)$$

Then, we observe that there exists the index $\tilde{t}_{T-1} \leq T-1$ such that $\sigma(\mathbf{x}_{\tilde{t}_{T-1}}^{(T-1)}; \mathbf{X}_{\tilde{t}_{T-1}-1}^{(T-1)}) \leq \sigma_{\lambda_{T-1}^2}^2(\mathbf{x}_{\tilde{t}_{T-1}}^{(T-1)}; \mathbf{X}_{\tilde{t}_{T-1}-1}^{(T-1)}) \leq \lambda_{T-1}$ by the application of Lemma 6 for the new sequence $(\mathbf{x}_t^{(T-1)})$. Again, by setting $\mathbf{x}_t^{(T-2)} = \mathbb{1}\{t < \tilde{t}_{T-1}\} \mathbf{x}_t^{(T-1)} + \mathbb{1}\{t \geq \tilde{t}_{T-1}\} \mathbf{x}_{t+1}^{(T-1)}$ and $\mathbf{X}_t^{(T-2)} = (\mathbf{x}_1^{(T-2)}, \dots, \mathbf{x}_t^{(T-2)})$ for any $t \in [T-2]$, we have

$$\sum_{t \in [T-1]} \sigma(\mathbf{x}_t^{(T-1)}; \mathbf{X}_{t-1}^{(T-1)}) \leq \sum_{t \in [T-1] \setminus \{\tilde{t}_{T-1}\}} \sigma(\mathbf{x}_t^{(T-1)}; \mathbf{X}_{t-1}^{(T-1)}) + \lambda_{T-1} \quad (18)$$

$$\leq \sum_{t \in [T-2]} \sigma(\mathbf{x}_t^{(T-2)}; \mathbf{X}_{t-1}^{(T-2)}) + \lambda_{T-1}. \quad (19)$$

We can repeat the above arguments until we reach $\bar{T} - 1$. Then, the resulting upper bound becomes

$$\sum_{t=1}^T \sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) \leq \sum_{t=1}^{\bar{T}-1} \sigma(\mathbf{x}_t^{(\bar{T}-1)}; \mathbf{X}_{t-1}^{(\bar{T}-1)}) + \sum_{t=\bar{T}}^T \lambda_t \leq \bar{T} - 1 + \sum_{t=\bar{T}}^T \lambda_t, \quad (20)$$

where the last inequality follows from $\sigma(\mathbf{x}_t^{(\bar{T}-1)}; \mathbf{X}_{t-1}^{(\bar{T}-1)}) \leq k(\mathbf{x}_t^{(\bar{T}-1)}, \mathbf{x}_t^{(\bar{T}-1)}) \leq 1$. $\square$

4 Discussion

Below, we discuss the remaining open questions in the noise-free setting.

- **Lower bound for squared exponential kernel.** While our results establish near-optimality for the Matérn kernel, the optimal simple regret rate under the SE kernel remains unknown. Since the existing upper bound $O(\ln^{d+1} T)$ for MIG [Vakili et al., 2021c] does not match the $O(\ln^{d/2} T)$ lower bound⁵, we conjecture that further room for improvement also exists in the noise-free setting. Specifically, the exponent $d + 1$ of the MIG is reflected in the denominator of the exponential factors in our regret Eq. (2). Therefore, we conjecture that $O(\sqrt{T} \exp(-CT^{2/d}))$ regret is the best guarantee for the simple regret in the SE kernel.
- **Constant cumulative regret in Bayesian setting.** By using Lemma 3, we can prove the same cumulative regret as Eq. (1) up to a logarithmic factor in noise-free GP-UCB or GP-TS in the Bayesian setting [Srinivas et al., 2010, Russo and Van Roy, 2014a]. However, the confidence width parameter in the Bayesian setting must scale as $\beta^{1/2} = O(\sqrt{\ln T})$ to construct a valid confidence bound. This leads to $O(\sqrt{\ln T})$ regret in SE and Matérn kernel with $d < \nu$ under the Bayesian setting, whereas the frequentist counterpart guarantees constant $O(1)$ regret (Theorem 1). This is counterintuitive, since, as shown in existing analyses for noisy settings [Scarlett, 2018, Srinivas et al., 2010], Bayesian regret often achieves smaller values than the worst-case regret in the frequentist setting. An interesting direction for future work is to either design a Bayesian algorithm with constant regret or prove an $\Omega(\sqrt{\ln T})$ lower bound in the Bayesian setting.
- **GP-UCB in simple regret minimization.** Our analysis shows that GP-UCB achieves nearly optimal simple and cumulative regrets under the Matérn kernel. On the other hand, several algorithms have been proposed that focus on minimizing the simple regret. One of the most well-known examples is EI, which greedily minimizes the simple regret under a Bayesian modeling assumption of GP, and has demonstrated good empirical performance in various applications [Brochu et al., 2010, Snoek et al., 2012]. Based on our theoretical results, there exists no algorithm that can achieve strictly better simple regret than that of GP-UCB in the worst-case sense. Nevertheless, as far as we are aware, GP-UCB tends to exhibit inferior empirical performance compared to EI in simple regret minimization. See, Appendix B.2. It remains unclear whether this phenomena arises from constant factors or additional logarithmic terms in the theoretical regret upper bounds, or whether it reflects a more fundamental gap between worst-case analysis and empirical behavior.

5 Conclusion

This paper shows that GP-UCB achieves nearly-optimal regret by proving a new regret upper bound for noise-free GP bandits. The key theoretical component of our analysis is a tight upper bound on the posterior standard deviations of GP tailored to a noise-free setting (Lemma 3). As remarked in Section 3, Lemma 3 can be applicable beyond the analysis of GP-UCB. Specifically, we expect that many existing theoretical results for noisy GP bandit settings can be translated to the noise-free setting by replacing the existing noisy upper bound of the posterior standard deviations with Lemma 3. For this reason, we believe that our result marks an important step toward advancing the theory for noise-free GP bandit algorithms.

⁵This is derived from the best-known cumulative regret upper bound $R_T = \tilde{O}(\sqrt{T\gamma_T(\lambda^2)})$ [Li and Scarlett, 2022, Salgia et al., 2021, Valko et al., 2013], and lower bound $R_T = \Omega(\sqrt{T \ln^{d/2} T})$ in the noisy setting [Scarlett et al., 2017].

References

- Milton Abramowitz and Irene A Stegun. *Handbook of mathematical functions with formulas, graphs, and mathematical tables*, volume 55. US Government printing office, 1968.
- Nachman Aronszajn. Theory of reproducing kernels. *Transactions of the American mathematical society*, 68(3):337–404, 1950.
- Ilija Bogunovic, Jonathan Scarlett, Stefanie Jegelka, and Volkan Cevher. Adversarially robust optimization with Gaussian processes. In *Proc. Neural Information Processing Systems (NeurIPS)*, 2018.
- Ilija Bogunovic, Zihan Li, Andreas Krause, and Jonathan Scarlett. A robust phased elimination algorithm for corruption-tolerant Gaussian process bandits. In *Proc. Neural Information Processing Systems (NeurIPS)*, 2022.
- Eric Brochu, Vlad M Cora, and Nando De Freitas. A tutorial on Bayesian optimization of expensive cost functions, with application to active user modeling and hierarchical reinforcement learning. *arXiv preprint arXiv:1012.2599*, 2010.
- Adam D Bull. Convergence rates of efficient global optimization algorithms. *Journal of Machine Learning Research*, 2011.
- Sayak Ray Chowdhury and Aditya Gopalan. On kernelized multi-armed bandits. In *Proc. International Conference on Machine Learning (ICML)*, 2017.
- Nando De Freitas, Alex J. Smola, and Masrour Zoghi. Exponential regret bounds for Gaussian process bandits with deterministic observations. In *Proceedings of the 29th International Conference on International Conference on Machine Learning*, page 955–962. Omnipress, 2012.
- Hamish Flynn and David Reeb. Tighter confidence bounds for sequential kernel regression. *arXiv preprint arXiv:2401.17037*, 2024.
- Alkis Gotovos, Nathalie Casati, Gregory Hitz, and Andreas Krause. Active learning for level set estimation. In *Proceedings of the Twenty-Third international joint conference on Artificial Intelligence*, 2013.
- Steffen Grünewälder, Jean-Yves Audibert, Manfred Opper, and John Shawe-Taylor. Regret bounds for Gaussian process bandit problems. In *Proc. International Conference on Artificial Intelligence and Statistics (AISTATS)*. JMLR Workshop and Conference Proceedings, 2010.
- Shogo Iwazaki. Improved regret bounds for Gaussian process upper confidence bound in Bayesian optimization. *arXiv preprint arXiv:2506.01393*, 2025.
- Shogo Iwazaki and Shinya Suzumura. No-regret bandit exploration based on soft tree ensemble model. *Advances in Neural Information Processing Systems*, pages 20982–21033, 2024.
- Shogo Iwazaki and Shion Takeno. Improved regret analysis in Gaussian process bandits: Optimality for noiseless reward, RKHS norm, and non-stationary variance. In *Proc. International Conference on Machine Learning (ICML)*, 2025a.
- Shogo Iwazaki and Shion Takeno. Near-optimal algorithm for non-stationary kernelized bandits. In *The 28th International Conference on Artificial Intelligence and Statistics*, 2025b.
- Arthur Jacot, Franck Gabriel, and Clément Hongler. Neural tangent kernel: Convergence and generalization in neural networks. *Advances in neural information processing systems*, 2018.
- David Janz. *Sequential decision making with feature-linear models*. PhD thesis, 2022.
- Motonobu Kanagawa, Philipp Hennig, Dino Sejdinovic, and Bharath K Sriperumbudur. Gaussian processes and kernel methods: A review on connections and equivalences. *arXiv preprint arXiv:1807.02582*, 2018.
- Parnian Kassraie and Andreas Krause. Neural contextual bandits without regret. In *International Conference on Artificial Intelligence and Statistics*, pages 240–278. PMLR, 2022.

- Parnian Kassraie, Andreas Krause, and Ilija Bogunovic. Graph neural network bandits. *Advances in Neural Information Processing Systems*, 2022.
- Hwanwoo Kim and Daniel Sanz-Alonso. Enhancing Gaussian process surrogates for optimization and posterior approximation via random exploration. *arXiv preprint arXiv:2401.17037*, 2024.
- Andreas Krause and Cheng Ong. Contextual Gaussian process bandit optimization. In *Proc. Neural Information Processing Systems (NeurIPS)*, 2011.
- Zihan Li and Jonathan Scarlett. Gaussian process bandit optimization with few batches. In *Proc. International Conference on Artificial Intelligence and Statistics (AISTATS)*, 2022.
- Zihan Li and Jonathan Scarlett. Regret bounds for noise-free cascaded kernelized bandits. *Transactions on Machine Learning Research*, 2024.
- Yueming Lyu, Yuan Yuan, and Ivor W Tsang. Efficient batch black-box optimization with deterministic regret bounds. *arXiv preprint arXiv:1905.10041*, 2019.
- Carl Edward Rasmussen and Christopher K. I. Williams. *Gaussian Processes for Machine Learning (Adaptive Computation and Machine Learning)*. The MIT Press, 2005.
- Daniel Russo and Benjamin Van Roy. Learning to optimize via posterior sampling. *Mathematics of Operations Research*, 39(4):1221–1243, 2014a.
- Daniel Russo and Benjamin Van Roy. Learning to optimize via information-directed sampling. *Advances in neural information processing systems*, 27, 2014b.
- Sudeep Salgia, Sattar Vakili, and Qing Zhao. A domain-shrinking based Bayesian optimization algorithm with order-optimal regret performance. In *Proc. Neural Information Processing Systems (NeurIPS)*, 2021.
- Sudeep Salgia, Sattar Vakili, and Qing Zhao. Random exploration in Bayesian optimization: Order-optimal regret and computational efficiency. In *Proc. International Conference on Machine Learning (ICML)*, 2024.
- Jonathan Scarlett. Tight regret bounds for Bayesian optimization in one dimension. In *Proceedings of the 35th International Conference on Machine Learning*, volume 80 of *Proceedings of Machine Learning Research*, pages 4500–4508. PMLR, 2018.
- Jonathan Scarlett, Ilija Bogunovic, and Volkan Cevher. Lower bounds on regret for noisy Gaussian process bandit optimization. In *Proc. Conference on Learning Theory (COLT)*, 2017.
- Jasper Snoek, Hugo Larochelle, and Ryan P Adams. Practical Bayesian optimization of machine learning algorithms. In *Proc. Neural Information Processing Systems (NeurIPS)*, 2012.
- Niranjan Srinivas, Andreas Krause, Sham Kakade, and Matthias Seeger. Gaussian process optimization in the bandit setting: No regret and experimental design. In *Proc. International Conference on Machine Learning (ICML)*, 2010.
- Sattar Vakili. Open problem: Regret bounds for noise-free kernel-based bandits. In *Proc. Conference on Learning Theory (COLT)*, 2022.
- Sattar Vakili, Nacime Bouziani, Sepehr Jalali, Alberto Bernacchia, and Da shan Shiu. Optimal order simple regret for Gaussian process bandits. In *Proc. Neural Information Processing Systems (NeurIPS)*, 2021a.
- Sattar Vakili, Michael Bromberg, Jezabel Garcia, Da-shan Shiu, and Alberto Bernacchia. Uniform generalization bounds for overparameterized neural networks. *arXiv preprint arXiv:2109.06099*, 2021b.
- Sattar Vakili, Kia Khezeli, and Victor Picheny. On information gain and regret bounds in Gaussian process bandits. In *Proc. International Conference on Artificial Intelligence and Statistics (AISTATS)*, 2021c.

Michal Valko, Nathan Korda, Rémi Munos, Ilias Flaounas, and Nello Cristianini. Finite-time analysis of kernelised contextual bandits. In *Proceedings of the Twenty-Ninth Conference on Uncertainty in Artificial Intelligence*, UAI'13, page 654–663. AUAI Press, 2013.

Marcela Zuluaga, Andreas Krause, and Markus Püschel. e-pal: An active learning approach to the multi-objective optimization problem. *Journal of Machine Learning Research*, 2016.

A Proofs for Section 3

A.1 Proof of Theorems 1 and 2

We first formally describe the existing noise-free confidence bound.

Lemma 7 (Deterministic confidence bound for noise-free setting, e.g., Corollary 3.11 in [Kanagawa et al., 2018], Lemma 11 in [Lyu et al., 2019], or Proposition 1 in [Vakili et al., 2021a]). *Suppose Assumption 1 holds. Then, for any sequence $(\mathbf{x}_t)_{t \in \mathbb{N}_+}$ on $\mathcal{X}$, the following statement holds:*

$$\forall t \in \mathbb{N}_+, \forall \mathbf{x} \in \mathcal{X}, |f(\mathbf{x}) - \mu(\mathbf{x}; \mathbf{X}_t)| \leq B\sigma(\mathbf{x}; \mathbf{X}_t), \quad (21)$$

where $\mathbf{X}_t = (\mathbf{x}_1, \dots, \mathbf{x}_t)$.

Although the remaining parts of the proofs are well-known results of GP-UCB, we provide the details for completeness. Based on the above lemma, we show Eqs. (9) and (10). Regarding R_T , we have

$$R_T = \sum_{t=1}^T f(\mathbf{x}^*) - f(\mathbf{x}_t) \quad (22)$$

$$\leq \sum_{t=1}^T [\mu(\mathbf{x}^*; \mathbf{X}_t) + B\sigma(\mathbf{x}^*; \mathbf{X}_t)] - [\mu(\mathbf{x}_t; \mathbf{X}_t) - B\sigma(\mathbf{x}_t; \mathbf{X}_t)] \quad (23)$$

$$\leq \sum_{t=1}^T [\mu(\mathbf{x}_t; \mathbf{X}_t) + B\sigma(\mathbf{x}_t; \mathbf{X}_t)] - [\mu(\mathbf{x}_t; \mathbf{X}_t) - B\sigma(\mathbf{x}_t; \mathbf{X}_t)] \quad (24)$$

$$= 2B \sum_{t=1}^T \sigma(\mathbf{x}_t; \mathbf{X}_t), \quad (25)$$

where the first inequality follows from Lemma 7, and the second inequality follows from the UCB-selection rule for $\mathbf{x}_t$. Similarly to the case of cumulative regret, we have

$$r_T = f(\mathbf{x}^*) - f(\widehat{\mathbf{x}}_T) \quad (26)$$

$$\leq \min_{t \in [T]} f(\mathbf{x}^*) - f(\mathbf{x}_t) \quad (27)$$

$$\leq \min_{t \in [T]} [\mu(\mathbf{x}^*; \mathbf{X}_t) + B\sigma(\mathbf{x}^*; \mathbf{X}_t)] - [\mu(\mathbf{x}_t; \mathbf{X}_t) - B\sigma(\mathbf{x}_t; \mathbf{X}_t)] \quad (28)$$

$$\leq \min_{t \in [T]} [\mu(\mathbf{x}_t; \mathbf{X}_t) + B\sigma(\mathbf{x}_t; \mathbf{X}_t)] - [\mu(\mathbf{x}_t; \mathbf{X}_t) - B\sigma(\mathbf{x}_t; \mathbf{X}_t)] \quad (29)$$

$$= 2B \min_{t \in [T]} \sigma(\mathbf{x}_t; \mathbf{X}_t), \quad (30)$$

where the first inequality follows from the definition of $\widehat{\mathbf{x}}_T$. Finally, the desired results are obtained by combining the above inequalities with Eqs. (11)–(13). $\square$

A.2 Proof of Lemma 3

We prove the following Lemma 8, which is a detailed version of Lemma 3 including the dependence against constant factors.

Lemma 8 (Detailed version of posterior standard deviation upper bound for SE and Matérn kernel). *Fix any compact input domain $\mathcal{X} \subset \mathbb{R}^d$, and kernel function $k : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$ that satisfies $k(\mathbf{x}, \mathbf{x}) \leq 1$ for all $\mathbf{x} \in \mathcal{X}$. Furthermore, let $C_{\text{SE}}, C_{\text{Mat}}, \underline{\lambda}_{\text{SE}}, \underline{\lambda}_{\text{Mat}} > 0, \underline{T}_{\text{SE}}, \underline{T}_{\text{Mat}} \geq 2$ be the constants⁶ that satisfies $\forall \lambda \in (0, \underline{\lambda}_{\text{SE}}], \forall t \geq \underline{T}_{\text{SE}}, \gamma_t(\lambda^2) \leq C_{\text{SE}}(\ln(t/\lambda^2))^{d+1}$ and $\forall \lambda \in (0, \underline{\lambda}_{\text{Mat}}], \forall t \geq \underline{T}_{\text{Mat}}, \gamma_t(\lambda^2) \leq C_{\text{Mat}}(t/\lambda^2)^{\frac{d}{2\nu+d}}(\ln(t/\lambda^2))^{\frac{2\nu}{2\nu+d}}$ for $k = k_{\text{SE}}$ and $k = k_{\text{Matérn}}$, respectively. Then, the following statements hold for any $T \in \mathbb{N}_+$ and any input sequence $\mathbf{x}_1, \dots, \mathbf{x}_T \in \mathcal{X}$:*

⁶The existence of these constants are guaranteed by the upper bound of MIG [Vakili et al., 2021c], which shows $\gamma_T(\lambda^2) = O(\ln^{d+1}(T/\lambda^2))$ and $\gamma_T(\lambda^2) = O((T/\lambda^2)^{\frac{d}{2\nu+d}} \ln^{\frac{2\nu}{2\nu+d}}(T/\lambda^2))$ (as $T \rightarrow \infty, \lambda \rightarrow 0$) under $k = k_{\text{SE}}$ and $k = k_{\text{Matérn}}$, respectively. Note that these constants do not depend on T , but may depend on d, ℓ, ν , and the diameter of $\mathcal{X}$.

- For $k = k_{\text{SE}}$,

$$\min_{t \in [T]} \sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) \leq \begin{cases} 1 & \text{if } T < \bar{T}_{\text{SE}}, \\ \sqrt{T} \exp\left(-\frac{1}{2} \tilde{C}_{\text{SE}} T^{\frac{1}{d+1}}\right) & \text{if } T \geq \bar{T}_{\text{SE}}, \end{cases} \quad (31)$$

$$\sum_{t=1}^T \sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) \leq \bar{T}_{\text{SE}} + (d+1) \left(\frac{\tilde{C}_{\text{SE}}}{2}\right)^{-\frac{3d+3}{2}} \Gamma\left(\frac{3d+3}{2}\right), \quad (32)$$

where $\tilde{C}_{\text{SE}} = (6C_{\text{SE}})^{-\frac{1}{d+1}}$ and $\bar{T}_{\text{SE}} = \max\{\underline{T}_{\text{SE}}, T_{\text{SE}}^{(\lambda)}, \lceil (d+1)^{d+1} / \tilde{C}_{\text{SE}}^{d+1} \rceil + 1\}$ with $T_{\text{SE}}^{(\lambda)} = \min\{T \in \mathbb{N}_+ \mid \forall t \geq T, t \exp(-\tilde{C}_{\text{SE}} t^{1/(d+1)}) \leq \lambda_{\text{SE}}^2\}$.

- For $k = k_{\text{Matérn}}$ with $\nu > 1/2$,

$$\min_{t \in [T]} \sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) \leq \begin{cases} 1 & \text{if } T < \bar{T}_{\text{Mat}}, \\ \tilde{C}_{\text{Mat}}^{1/2} T^{-\frac{\nu}{d}} (\ln T)^{\frac{\nu}{d}} & \text{if } T \geq \bar{T}_{\text{Mat}}, \end{cases} \quad (33)$$

$$\sum_{t=1}^T \sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) \leq \begin{cases} \bar{T}_{\text{Mat}} + \tilde{C}_{\text{Mat}}^{1/2} \frac{d}{d-\nu} T^{\frac{d-\nu}{d}} (\ln T)^{\frac{\nu}{d}} & \text{if } d > \nu, \\ \bar{T}_{\text{Mat}} + \tilde{C}_{\text{Mat}}^{1/2} (\ln T)^2 & \text{if } d = \nu, \\ \bar{T}_{\text{Mat}} + \tilde{C}_{\text{Mat}}^{1/2} \frac{\Gamma(\frac{\nu}{d}+1)}{(\frac{\nu}{d}-1)^{\frac{\nu}{d}+1}} & \text{if } d < \nu, \end{cases} \quad (34)$$

where $\tilde{C}_{\text{Mat}} = \max\left\{1, \left(2 + \frac{2\nu}{d}\right)^{\frac{2\nu}{d}} (6C_{\text{Mat}})^{1+\frac{2\nu}{d}}\right\}$ and $\bar{T}_{\text{Mat}} = \max\{4, \underline{T}_{\text{Mat}}, T_{\text{Mat}}^{(\lambda)}\}$ with $T_{\text{Mat}}^{(\lambda)} = \min\{T \in \mathbb{N}_+ \mid \forall t \geq T, \tilde{C}_{\text{Mat}} t^{-\frac{2\nu}{d}} (\ln t)^{\frac{2\nu}{d}} \leq \lambda_{\text{Mat}}^2\}$.

The upper bound in Lemma 8 depends on the quantities $\tilde{C}_{\text{SE}}, \bar{T}_{\text{SE}}, \tilde{C}_{\text{Mat}}, \bar{T}_{\text{Mat}} > 0$, which are related to the implied constants in the upper bounds of the MIG. However, under the fixed d, ℓ , and ν , $\tilde{C}_{\text{SE}}, \bar{T}_{\text{SE}}, \tilde{C}_{\text{Mat}}, \bar{T}_{\text{Mat}} > 0$ are also constants, which implies the conclusions of Lemma 3.

Proof of Lemma 8. When $k = k_{\text{SE}}$, we set $\lambda_t^2 = t \exp(-\tilde{C}_{\text{SE}} t^{\frac{1}{d+1}})$, $\bar{T} = \bar{T}_{\text{SE}} := \max\{\underline{T}_{\text{SE}}, T_{\text{SE}}^{(\lambda)}, \lceil (d+1)^{d+1} / \tilde{C}_{\text{SE}}^{d+1} \rceil + 1\}$. From the definition of λ_t^2 and $\bar{T}_{\text{SE}}$, for any $t \geq \bar{T}_{\text{SE}}$, we have

$$\gamma_t(\lambda_t^2) \leq C_{\text{SE}} \left[\ln \left(\frac{t}{\lambda_t^2} \right) \right]^{d+1} \quad (35)$$

$$= C_{\text{SE}} \left[\ln \exp \left(\tilde{C}_{\text{SE}} t^{\frac{1}{d+1}} \right) \right]^{d+1} \quad (36)$$

$$= C_{\text{SE}} \tilde{C}_{\text{SE}}^{d+1} t. \quad (37)$$

Furthermore,

$$C_{\text{SE}} \tilde{C}_{\text{SE}}^{d+1} t \leq \frac{t-1}{3} \Leftrightarrow \tilde{C}_{\text{SE}}^{d+1} \leq \frac{t-1}{3C_{\text{SE}}t} \quad (38)$$

$$\Leftrightarrow \tilde{C}_{\text{SE}}^{d+1} \leq \frac{1}{6C_{\text{SE}}} \quad (39)$$

$$\Leftrightarrow \tilde{C}_{\text{SE}} \leq \left(\frac{1}{6C_{\text{SE}}} \right)^{\frac{1}{d+1}}, \quad (40)$$

where the second line follows from the inequality $t-1 \geq t/2$ for all $t \geq \bar{T}_{\text{SE}} \geq 2$. By noting the definition of $\tilde{C}_{\text{SE}}$, we conclude that $\forall t \geq \bar{T}_{\text{SE}}, \gamma_t(\lambda_t^2) \leq C_{\text{SE}} \tilde{C}_{\text{SE}}^{d+1} t \leq \frac{t-1}{3}$ from the above inequalities, which implies that Lemmas 4 and 5 hold with $\lambda_t^2 = t \exp(-\tilde{C}_{\text{SE}} t^{\frac{1}{d+1}})$ and $\bar{T} = \bar{T}_{\text{SE}}$. Eq. (31) directly follows from Lemma 4 using the fact $\sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) \leq k(\mathbf{x}_t, \mathbf{x}_t) \leq 1$. As for Eq. (32), Lemma 5

implies

$$\sum_{t=1}^T \sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) \leq \bar{T}_{\text{SE}} + \sum_{t=\bar{T}_{\text{SE}}}^T \lambda_t \quad (41)$$

$$\leq \bar{T}_{\text{SE}} + \int_{\bar{T}_{\text{SE}}-1}^T \sqrt{t} \exp\left(-\frac{1}{2} \tilde{C}_{\text{SE}} t^{\frac{1}{d+1}}\right) dt \quad (42)$$

$$\leq \bar{T}_{\text{SE}} + \int_1^T \sqrt{t} \exp\left(-\frac{1}{2} \tilde{C}_{\text{SE}} t^{\frac{1}{d+1}}\right) dt, \quad (43)$$

where the second line follows from the fact that the function $g(t) := t \exp(-\tilde{C}_{\text{SE}} t^{1/(d+1)})$ is non-increasing for $t \geq \bar{T}_{\text{SE}} - 1$. In fact, we have

$$g'(t) = \exp\left(-\tilde{C}_{\text{SE}} t^{\frac{1}{d+1}}\right) \left(1 - \frac{\tilde{C}_{\text{SE}}}{d+1} t^{\frac{1}{d+1}}\right), \quad (44)$$

which implies $g'(t) \leq 0$ for $t \geq \bar{T}_{\text{SE}} - 1 \geq (d+1)^{d+1}/\tilde{C}_{\text{SE}}^{d+1}$. To bound the quantity $\int_1^T \sqrt{t} \exp\left(-\frac{1}{2} \tilde{C}_{\text{SE}} t^{\frac{1}{d+1}}\right) dt$, we further derive the following upper bound with $C := \tilde{C}_{\text{SE}}/2 > 0$:

$$\int_1^T \sqrt{t} \exp\left(-C t^{\frac{1}{d+1}}\right) dt = \int_C^{CT^{1/(d+1)}} \left(\frac{u}{C}\right)^{(d+1)/2} e^{-u} (d+1) \left(\frac{u}{C}\right)^d \frac{1}{C} du \quad (\because u = Ct^{1/(d+1)}) \quad (45)$$

$$= (d+1) C^{-(3d+3)/2} \int_C^{CT^{1/(d+1)}} u^{(3d+1)/2} e^{-u} du \quad (46)$$

$$\leq (d+1) C^{-(3d+3)/2} \int_0^\infty u^{(3d+1)/2} e^{-u} du \quad (47)$$

$$= (d+1) C^{-(3d+3)/2} \Gamma\left(\frac{3d+3}{2}\right). \quad (48)$$

Next, when $k = k_{\text{Matérn}}$, we set $\lambda_t^2 = \tilde{C}_{\text{Mat}} t^{-\frac{2\nu}{d}} (\ln t)^{\frac{2\nu}{d}}$ and $\bar{T} = \max\{4, \underline{T}_{\text{Mat}}, T_{\text{Mat}}^{(\lambda)}\}$ with $\tilde{C}_{\text{Mat}} = \left(2 + \frac{2\nu}{d}\right)^{\frac{2\nu}{d}} (6C_{\text{Mat}})^{1+\frac{2\nu}{d}}$. Then, for any $t \geq \bar{T}_{\text{Mat}}$, it holds that

$$\gamma_t(\lambda_t^2) \leq C_{\text{Mat}} \left(\frac{t}{\lambda_t^2}\right)^{\frac{d}{2\nu+d}} \left[\ln\left(\frac{t}{\lambda_t^2}\right)\right]^{\frac{2\nu}{2\nu+d}} \quad (49)$$

$$= C_{\text{Mat}} \tilde{C}_{\text{Mat}}^{-\frac{d}{2\nu+d}} t (\ln t)^{-\frac{2\nu}{2\nu+d}} \left[\ln\left(\tilde{C}_{\text{Mat}}^{-1} t^{\frac{d+2\nu}{d}} (\ln t)^{-\frac{2\nu}{d}}\right)\right]^{\frac{2\nu}{2\nu+d}} \quad (50)$$

$$= C_{\text{Mat}} \tilde{C}_{\text{Mat}}^{-\frac{d}{2\nu+d}} t (\ln t)^{-\frac{2\nu}{2\nu+d}} \left[\ln\left(\tilde{C}_{\text{Mat}}^{-1}\right) + \frac{d+2\nu}{d} (\ln t) - \frac{2\nu}{d} (\ln \ln t)\right]^{\frac{2\nu}{2\nu+d}} \quad (51)$$

$$\leq C_{\text{Mat}} \tilde{C}_{\text{Mat}}^{-\frac{d}{2\nu+d}} t (\ln t)^{-\frac{2\nu}{2\nu+d}} \left[\frac{2d+2\nu}{d} (\ln t)\right]^{\frac{2\nu}{2\nu+d}} \quad (52)$$

$$= C_{\text{Mat}} \tilde{C}_{\text{Mat}}^{-\frac{d}{2\nu+d}} t \left(\frac{2d+2\nu}{d}\right)^{\frac{2\nu}{2\nu+d}}, \quad (53)$$

where the fourth line follows from $\tilde{C}_{\text{Mat}} \geq 1 \Rightarrow \tilde{C}_{\text{Mat}} \geq 1/t \Leftrightarrow \ln(\tilde{C}_{\text{Mat}}^{-1}) \leq \ln t$ for $t \geq 1$. Furthermore,

$$C_{\text{Mat}} \tilde{C}_{\text{Mat}}^{-\frac{d}{2\nu+d}} t \left(\frac{2d+2\nu}{d}\right)^{\frac{2\nu}{2\nu+d}} \leq \frac{t-1}{3} \Leftrightarrow 3C_{\text{Mat}} \frac{t}{t-1} \left(\frac{2d+2\nu}{d}\right)^{\frac{2\nu}{2\nu+d}} \leq \tilde{C}_{\text{Mat}} \quad (54)$$

$$\Leftrightarrow \left(\frac{3C_{\text{Mat}} t}{t-1}\right)^{1+\frac{2\nu}{d}} \left(2 + \frac{2\nu}{d}\right)^{\frac{2\nu}{d}} \leq \tilde{C}_{\text{Mat}} \quad (55)$$

$$\Leftrightarrow (6C_{\text{Mat}})^{1+\frac{2\nu}{d}} \left(2 + \frac{2\nu}{d}\right)^{\frac{2\nu}{d}} \leq \tilde{C}_{\text{Mat}}. \quad (56)$$

Combining the above inequalities, we can confirm $\forall t \geq \bar{T}_{\text{Mat}}, \gamma_t(\lambda_t^2) \leq \frac{t-1}{3}$. Therefore, Lemmas 4 and 5 holds with $\lambda_t^2 = \bar{C}_{\text{Mat}} t^{-\frac{2\nu}{d}} (\ln t)^{\frac{2\nu}{d}}$ and $\bar{T} = \bar{T}_{\text{Mat}}$. Here, Eq. (33) is the direct consequence of Lemmas 4. As for Eq. (34), we have

$$\sum_{t=1}^T \sigma(\mathbf{x}_t; \mathbf{X}_{t-1}) \leq \bar{T}_{\text{Mat}} + \sum_{t=\bar{T}_{\text{Mat}}}^T \lambda_t \quad (57)$$

$$\leq \bar{T}_{\text{Mat}} + \bar{C}_{\text{Mat}}^{1/2} \int_{\bar{T}_{\text{Mat}}-1}^T t^{-\frac{\nu}{d}} (\ln t)^{\frac{\nu}{d}} dt \quad (58)$$

$$\leq \bar{T}_{\text{Mat}} + \bar{C}_{\text{Mat}}^{1/2} \int_1^T t^{-\frac{\nu}{d}} (\ln t)^{\frac{\nu}{d}} dt, \quad (59)$$

where the second line follows from the fact that the function $g(t) := t^{-\frac{2\nu}{d}} (\ln t)^{\frac{2\nu}{d}}$ is non-increasing for $t \geq \bar{T}_{\text{Mat}} - 1 \geq 3 > e$. Indeed, we have

$$g'(t) = \frac{2\nu}{d} t^{-\frac{2\nu}{d}-1} (\ln t)^{\frac{2\nu}{d}} \left((\ln t)^{-1} - 1 \right), \quad (60)$$

which implies $g'(t) \leq 0$ for $t \geq e$. The desired results are obtained by bounding the quantity $\int_1^T t^{-\frac{\nu}{d}} (\ln t)^{\frac{\nu}{d}} dt$ from above. When $d > \nu$, we have

$$\int_1^T t^{-\frac{\nu}{d}} (\ln t)^{\frac{\nu}{d}} dt \leq (\ln T)^{\frac{\nu}{d}} \int_1^T t^{-\frac{\nu}{d}} dt = (\ln T)^{\frac{\nu}{d}} \left[\frac{d}{d-\nu} t^{\frac{d-\nu}{d}} \right]_1^T \leq \frac{d}{d-\nu} T^{\frac{d-\nu}{d}} (\ln T)^{\frac{\nu}{d}}. \quad (61)$$

When $d = \nu$,

$$\int_1^T t^{-\frac{\nu}{d}} (\ln t)^{\frac{\nu}{d}} dt \leq (\ln T) \int_1^T t^{-1} dt = (\ln T)^2. \quad (62)$$

When $d < \nu$, we have

$$\int_1^T t^{-\frac{\nu}{d}} (\ln t)^{\frac{\nu}{d}} dt = \int_0^{\ln T} e^{-(\frac{\nu}{d}-1)u} u^{\frac{\nu}{d}} du \quad (\because u = \ln t) \quad (63)$$

$$\leq \int_0^\infty e^{-(\frac{\nu}{d}-1)u} u^{\frac{\nu}{d}} du \quad (64)$$

$$= \frac{\Gamma(\frac{\nu}{d} + 1)}{(\frac{\nu}{d} - 1)^{\frac{\nu}{d} + 1}}, \quad (65)$$

where the last line follows from the standard property of Gamma function: $\int_0^\infty e^{-\lambda u} u^b du = \Gamma(b+1)/\lambda^{b+1}$ for any $\lambda > 0$ and $b > -1$ (e.g., Equation 6.1.1 in [Abramowitz and Stegun, 1968]). $\square$

B Detail of the Experiment

B.1 Experimental settings for Figure 1

We give the detailed experimental settings used to plot Figure 1.

- **Objective function.** We define the true underlying objective function as $f(\cdot) = \sum_{m=1}^{50} c_m k(\mathbf{x}^{(m)}, \cdot)$, where $c_m \sim \text{Uniform}([-1, 1])$ and $\mathbf{x}^{(m)} \sim \text{Uniform}([0, 1]^2)$ are independently generated random variables. Note that by the definition of the RKHS, $f \in \mathcal{H}_k$ with $\|f\|_k = \sqrt{\sum_{m=1}^{50} \sum_{\tilde{m}=1}^{50} c_m c_{\tilde{m}} k(\mathbf{x}^{(m)}, \mathbf{x}^{(\tilde{m})})}$.
- **Kernel.** We fix the lengthscale parameter ℓ as $\ell = 0.25$ in all experiments. We use the same kernel function in the GP-model used in the algorithms as the kernel leveraged to generate the objective function.
- **Other parameters.** We define the input domain $\mathcal{X}$ as the uniformly aligned 50×50 grid points on $[0, 1]^2$. In all algorithms, we set the confidence width parameter $\beta^{1/2}$ to the exact RKHS norm. Furthermore, we set the initial batch size in PE and REDS as 5. Finally, we define the common initial point $\mathbf{x}_1$ for all the algorithms as the uniformly sampled point from $\mathcal{X}$.

In the above-described setting, we conduct experiments with 3000 different seeds.

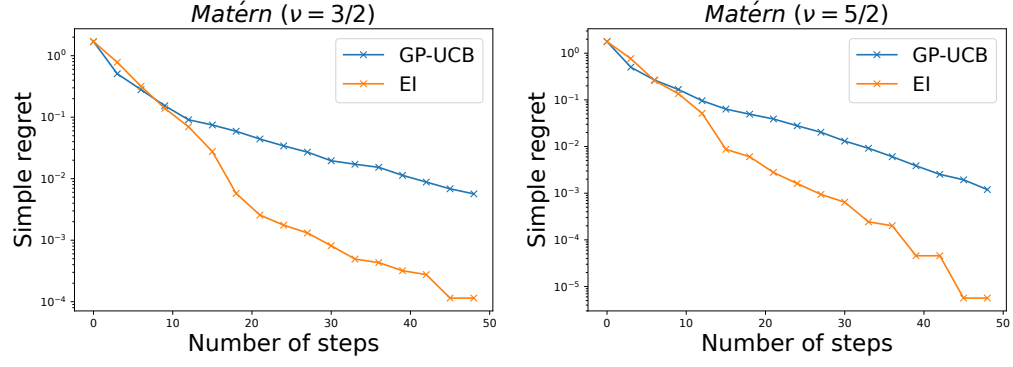


Figure 2: Comparison between GP-UCB and EI in simple regret minimization over 100 different seeds. We conduct experiments with Matérn kernel under $\nu = 3/2$ (left) and $\nu = 5/2$ (right).

B.2 Comparison between EI and GP-UCB

Under the same setting as the previous subsection, we also compare GP-UCB’s empirical performance with that of EI in simple regret minimization. Figure B.2 shows the results. We can confirm that, although GP-UCB achieves nearly-optimal worst-case regret, the empirical performance of GP-UCB is consistently worse than that of EI.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper’s contributions and scope?

Answer: [\[Yes\]](#)

Justification: We provide a “contribution” paragraph in the introduction, which summarizes our main contributions. The generality of our results is also discussed in Section 3.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: In Section 4, we discuss the remaining open questions that this paper cannot resolve.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate “Limitations” section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren’t acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[Yes\]](#)

Justification: Regarding Theorems 1, 2, ??, and Lemma 3, we describe the complete proofs in the appendix, while proof sketches are provided in the main paper. Regarding Lemmas 4 and 5, we provide complete proofs in the main paper.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [\[Yes\]](#)

Justification: We provide the detailed setting to reproduce Figure 1 in Appendix B.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [No]

Justification: We do not provide code since the central contribution is on the theoretical aspect.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: All the parameters to plot Figure 1 are described in Appendix B.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [No]

Justification: We do not provide an error bar to enhance the clarity of Figure 1. This is not related to our contribution.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer “Yes” if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).

- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [No]

Justification: We do not provide computer resource information since the computational time is unrelated to our contribution.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: We have reviewed the Code of Ethics.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [NA]

Justification: This paper focuses on the theoretical aspect, and does not relate to any particular applications.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.

- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: This paper poses no such risks.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [NA]

Justification: This paper does not use existing assets.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.

- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. **New assets**

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: This paper does not release new assets.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and research with human subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: This paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. **Institutional review board (IRB) approvals or equivalent for research with human subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. Declaration of LLM usage

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: The core method development in this research does not involve LLMs as any important, original, or non-standard components.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.