
Robust Satisficing Gaussian Process Bandits Under Adversarial Attacks

Artun Saday *

Bilkent University, Ankara, Türkiye
artun.saday@bilkent.edu.tr

Yaşar Cahit Yıldırım *

Bilkent University, Ankara, Türkiye
cahit.yildirim@bilkent.edu.tr

Cem Tekin

Bilkent University, Ankara, Türkiye
cemtekin@ee.bilkent.edu.tr

Abstract

We address the problem of Gaussian Process (GP) optimization in the presence of unknown and potentially varying adversarial perturbations. Unlike traditional robust optimization approaches that focus on maximizing performance under worst-case scenarios, we consider a robust satisficing objective, where the goal is to consistently achieve a predefined performance threshold τ , even under adversarial conditions. We propose two novel algorithms based on distinct formulations of robust satisficing, and show that they are instances of a general robust satisficing framework. Further, each algorithm offers different guarantees depending on the nature of the adversary. Specifically, we derive two regret bounds: one that is sublinear over time, assuming certain conditions on the adversary and the satisficing threshold τ , and another that scales with the perturbation magnitude but requires no assumptions on the adversary. Through extensive experiments, we demonstrate that our approach outperforms the established robust optimization methods in achieving the satisficing objective, particularly when the ambiguity set of the robust optimization framework is inaccurately specified.

1 Introduction

Bayesian optimization (BO) is a framework particularly suited for sequentially optimizing difficult-to-evaluate black-box functions f , using a probabilistic model such as a Gaussian process (GP) [1]. BO uses a probabilistic surrogate model (e.g., GP) with an acquisition function to decide where to sample next. Designing acquisition functions tailored to specific problems is a significant research area in the BO literature. The performance of a BO algorithm is often evaluated by its *regret*, which quantifies the difference between the objective function values at the chosen points and the optimal point, or another predefined success criterion.

BO has proven effective in applications such as hyperparameter tuning [2], engineering design [3], experimental design [4], and clinical modeling [5]. In these settings, robustness to distribution shifts and environmental variability is essential. For instance, machine learning models must generalize beyond training data, while engineering designs must transfer from simulation to real-world deployment. In high-stakes domains like autonomous driving or patient treatment, ignoring uncertainty can lead to severe failures. Robust optimization (RO) addresses this by seeking solutions that perform well under worst-case scenarios defined by an ambiguity set [6]. However, the ambiguity set must be carefully defined: if it is too narrow, it may exclude relevant uncertainties; if too broad, it can lead to conservative solutions that underperform under typical conditions [7, 8].

*Equal contribution.

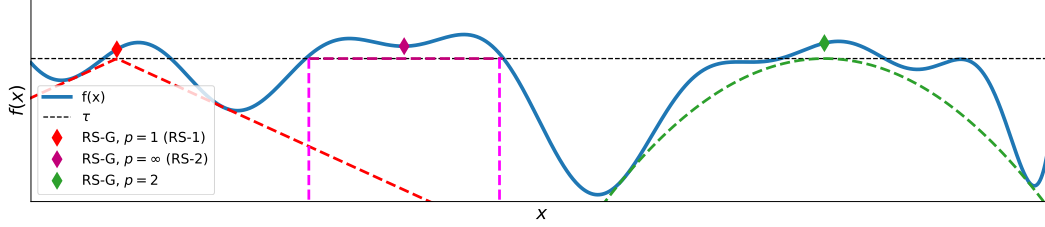


Figure 1: Illustration of how different RS formulations trade off reward guarantees with respect to perturbation magnitudes during action selection. Decision of each formulation can be interpreted as selecting the action (diamonds) by enforcing the widest fragility-cone (dashed lines) to limit the performance degradation beyond τ due to the perturbation. The fragility-cone is the reward guarantee of the solution $x^{\text{RS-G}}$: $\tau - [\kappa_{\tau,p} d(x^{\text{RS-G}}, x^{\text{RS-G}} + \delta)]^p$, defined in (10). It characterizes the shape of the constraint on performance degradation. For more details, see Section 4 and Appendix B.

Satisficing, a term coined by Herbert Simon [9], introduces a pragmatic approach to decision-making where the goal is to achieve an outcome x that is *good enough* ($f(x) \geq \tau$) rather than optimal. This concept is especially relevant in environments where the optimal solution is either unattainable or impractical due to computational or informational constraints. Satisficing has been studied in the context of multi-armed bandits (MAB) by [10, 11, 12, 13], while [14] investigates the problem in a non-Bayesian setting using the RKHS framework. Building on this concept, robust satisficing (RS) extends the principle of satisficing by aiming to optimize confidence in achieving a sufficiently good outcome, thus providing a robust framework for decision-making under uncertainty. Unlike simple satisficing, which may serve as a stopping heuristic, RS focuses on ensuring reliable performance across a variety of scenarios. Decision theorists such as [15] and [16] argue that RS offers a superior alternative to RO, as it emphasizes acceptable performance across a broad range of scenarios rather than focusing solely on worst-case outcomes. [17] proposes an optimization formulation for RS which inspires our approach. [18] builds on this by combining RO and RS in a unified framework that accounts for both outcome and estimation uncertainties. In our work we extend the RS literature by proposing a novel formulation, RS-G, which generalizes [17] by allowing nonlinear decay of the reward guarantee via a tunable parameter.

One form of uncertainty that RO and RS seek to address are adversarial perturbations. In many learning systems, small, targeted changes to the input can induce large deviations in output, undermining reliability. While this phenomenon is well studied in neural networks [19, 20], similar vulnerabilities exist in continuous control and optimization settings, where the selected input may be corrupted by noise or external interference. Addressing such perturbations requires robust strategies that extend beyond worst-case formulations, particularly when the perturbation model or budget is only partially known. This motivates our development of an RS approach for adversarially perturbed GP optimization.

Related Works Research on robustness, particularly adversarial robustness, in GP optimization is extensive. [21] tackles BO with outliers by combining robust GP regression with outlier diagnostics. [22], closely related to our work, introduces the setting of adversarial robustness in GP bandits. They propose an RO based algorithm that assumes a known perturbation budget, while we adopt an RS approach without this assumption. [23] addresses a similar problem where the learner selects $\tilde{x}_t$ but observes outputs at a different point $x_t \sim P_{\tilde{x}_t}$, using a GP model for probability distributions. In the noisy input setting, [24] proposes an entropy search algorithm. Adversarial corruption of observations, rather than inputs, is considered by [25, 26]. Contextual GP optimization considers uncertainty in the context distribution. [27] address this using RO with a known ambiguity set, while [28] extend the approach by adopting RS without assuming knowledge of the ambiguity set. A complementary direction is explored by [29, 30, 31], focusing on the BO of risk measures while [32] jointly optimizes the mean and variance of $f(x, \omega)$ using a GP.

Contributions We propose new RS-based acquisition functions for the first time for a variant of the GP optimization problem called *Adversarially Robust GP Optimization* introduced in [22], where an adversary perturbs the selected action. Based on principles of robust satisficing, we introduce two novel algorithms for finding a robust action while achieving small lenient and robust satisficing

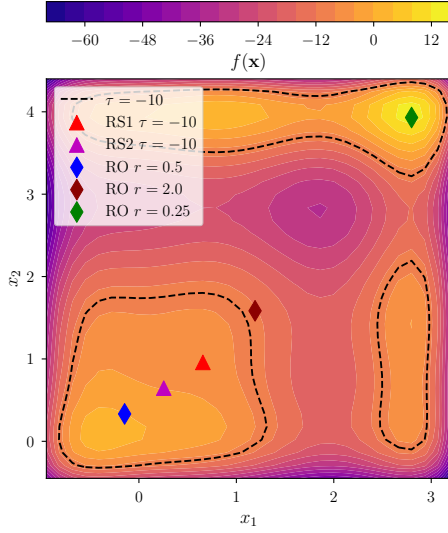


Figure 2: **(Left)** Illustration of RO and RS solutions. Highlighted contours represent $f(x) = \tau$, i.e., the satisfying threshold. **(Right)** Formulations of different optimization objectives.

regret, based on two distinct RS formulations (RS-1 and RS-2 in Figure 2). We show that these formulations are edge cases of a general robust satisfying framework (RS-G in Figure 2 (right)) with smoothness parameter $p \geq 1$. We demonstrate how a spectrum of robust solutions can be obtained by varying p (see Figure 1). We provide theoretical bounds on the regrets of our algorithms under certain assumptions. Finally, we present experiments, including in-silico experiments on an FDA-approved simulator for diabetes management with realistic adversarial perturbations, to demonstrate the strengths of our algorithms addressing the shortcomings of other approaches known in the literature.

2 Problem Formulation

Let f be an unknown function defined on the domain $\mathcal{X} \subset \mathbb{R}^m$, endowed with a pseudometric $d(\cdot, \cdot) : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$. Further assume that f belongs to a Reproducing Kernel Hilbert Space (RKHS) $\mathcal{H}$, with the reproducing kernel $k(\cdot, \cdot) : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$, and has bounded Hilbert norm $\|f\|_{\mathcal{H}} \leq B$, for some $B > 0$. This assumption ensures that f satisfies smoothness conditions such as Lipschitz continuity with respect to the kernel metric² [33] and allows the construction of rigorous confidence bounds using GPs [34]. At each round $t \in [T]$, where T represents the time horizon, the learner picks a point $\tilde{x}_t \in \mathcal{X}$, in response to which the adversary selects a perturbation δ_t . The learner then observes the perturbed point $x_t = \tilde{x}_t + \delta_t$ along with the noisy sample at that point $y_t = f(x_t) + \eta_t$, where η_t are independent R -subGaussian random variables. If the adversary is allowed to play any perturbation δ_t , no matter what action $\tilde{x}_t$ the learner chooses, the adversary can always choose a perturbation δ_t such that $\tilde{x}_t + \delta_t = \arg \min_{x \in \mathcal{X}} f(x)$, and the optimization problem would be hopeless. Hence at each round t , we assume an *unknown perturbation budget* ϵ_t , that bounds set of perturbations the adversary can play. We define for each point $x \in \mathcal{X}$, the *ambiguity ball*

$$\Delta_{\epsilon}(x) := \{x' - x : x' \in \mathcal{X} \text{ and } d(x, x') \leq \epsilon\} . \quad (1)$$

The ambiguity ball (1) defines the set of all permissible perturbations of the adversary, when the perturbation budget is ϵ . Hence in response to action $\tilde{x}_t$, the adversary plays the perturbation $\delta_t \in \Delta_{\epsilon_t}(\tilde{x}_t)$.

Remark 2.1. In the presented setting, it is not possible to learn the function optimum in general. As an example consider the two arm MAB problem with $\mathcal{X} = \{x_1, x_2\}$ with $f(x_1) = 0$ and $f(x_2) = 1$.

²For the kernel metric, we have $|f(x) - f(x')| = |\langle f, k(\cdot, x) - k(\cdot, x') \rangle| \leq \|f\|_{\mathcal{H}} d_k(x, x')$. Note that f is not necessarily Lipschitz continuous with respect to an arbitrary metric on $\mathcal{X}$. Hence in our analysis, we assume the pseudometric $d(\cdot, \cdot)$ is the kernel metric. Additionally, when the feature map $k(\cdot, x)$ is injective, this pseudometric is also a metric.

Then at each round t , a sufficiently strong adversary can always choose a perturbation δ_t such that $x_t = x_1$, hence the learner can never observe the optimal action.

2.1 Robust Satisficing

The key difference between RS and RO lies in their objectives. While RO seeks to optimize for worst-case outcomes, RS prioritizes achieving satisfactory performance across various scenarios. RO relies heavily on defining an ambiguity set, typically as a ball with radius r , centered around a reference value of the contingent variable. The choice of r greatly impacts performance. Overestimating r leads to conservative solutions that overemphasize unlikely worst-case scenarios, while underestimating r risks overly optimistic solutions. Figure 2 (left) demonstrates the effect of different choices of r .

In our context, the contingent variables are the adversarial perturbations. In other settings, such as data-driven optimization with distributional shifts, contingent variables might represent different uncertainties, such as deviations from an empirical distribution [35]. Unlike RO, which requires an ambiguity set, RS relies only on the satisficing threshold τ , representing the desired outcome. This makes τ more interpretable and adaptable to real-world problems, where it can be chosen by domain experts. We now present two formulations based on the RS approach, assuming the function f is known.

RS-1 We adapt the novel formulation by [17] to our adversarial setting. Let $(\cdot)^+ = \max\{\cdot, 0\}$. The objective of this formulation is to find $x^{\text{RS-1}} \in \mathcal{X}$ that solves the optimization problem $\kappa_\tau := \min_{x \in \mathcal{X}} \kappa_\tau(x)$ where $\kappa_\tau(x)$ is the *fragility* of action $x \in \mathcal{X}$ defined as

$$\kappa_\tau(x) := \min k \text{ s.t. } f(x + \delta) \geq \tau - kd(x, x + \delta), \quad \forall \delta \in \Delta_\infty(x), \quad k \geq 0. \quad (2)$$

This is equivalent to solving

$$\kappa_\tau(x) := \begin{cases} \left(\max_{\delta \in \Delta_\infty(x) \setminus 0} \frac{\tau - f(x + \delta)}{d(x, x + \delta)} \right)^+ & \text{if } f(x) \geq \tau \\ +\infty & \text{if } f(x) < \tau. \end{cases} \quad (3)$$

When (2) is feasible, the RS-1 action is defined as $x^{\text{RS-1}} = \arg \min_{x \in \mathcal{X}} \kappa_\tau(x)$. The fragility of an action can be seen as the minimum rate of suboptimality the action can achieve w.r.t. the threshold τ , per unit perturbation, for all possible perturbations. Crucially, while the RO formulation only considers perturbations possible under a given perturbation budget ϵ , the RS-1 formulation (2) considers all possible perturbations. This way RS gives safety guarantees even when the perturbation budget of the adversary is unknown. Figure 2 (left) highlights the RS-1 action $x^{\text{RS-1}}$ together with optimization objectives.

RS-2 Inspired by the literature on robust satisficing, we present another formulation called RS-2, similar to the one in [36]. Formally, the objective of RS-2 is to find $x^{\text{RS-2}} \in \mathcal{X}$ that solves the following optimization problem $\epsilon_\tau := \max_{x \in \mathcal{X}} \epsilon_\tau(x)$ where $\epsilon_\tau(x)$ is the *critical radius* of action $x \in \mathcal{X}$ defined as

$$\epsilon_\tau(x) := \max_{x' \in \mathcal{X}} d(x, x') \text{ s.t. } f(x + \delta) \geq \tau, \quad \forall \delta \in \Delta_{d(x, x')}(x). \quad (4)$$

The robust satisficing action is defined as the one with the maximum critical radius, $x^{\text{RS-2}} = \arg \max_{x \in \mathcal{X}} \epsilon_\tau(x)$. If $f(x) = \tau$ and x is a local maximum, then $\epsilon_\tau(x) = 0$ because $\delta = 0$ is the only perturbation satisfying the inequality condition. If $f(x) < \tau$, then $\epsilon_\tau(x) = -\infty$, as no ϵ satisfies the condition, and by convention, we take the supremum of the empty set as the critical radius, which is $-\infty$. Unlike RO, RS-2 does not assume knowledge of the adversary's budget ϵ_t but instead searches for the action that achieves the threshold τ under the widest possible range of perturbations. This target-driven approach aligns RS-2 with the philosophy of robust satisficing. Figure 2 (left) highlights $x^{\text{RS-2}}$ alongside other optimization objectives. If $\tau > f(\hat{x})$, where $\hat{x} := \arg \max_{x \in \mathcal{X}} f(x)$ is the optimal action, no action can meet the threshold, even without perturbations, hence we introduce the following assumption moving on.

Assumption 2.2. $\tau \leq \max_{x \in \mathcal{X}} f(x)$.

Assumption 2.2 ensures that the optimization problems (2) for RS-1 and (4) for RS-2 are feasible, thereby guaranteeing that a robust satisficing action is well-defined in both cases. However, since f is

unknown, it may not be immediately clear to the learner whether a specific threshold τ satisfies this assumption. If the learner is flexible with the satisficing threshold, Assumption 2.2 can be relaxed by dynamically selecting τ at each round to be less than $\text{lcb}_t(\hat{x}'_t)$, where $\hat{x}'_t := \arg \max_{x \in \mathcal{X}} \text{lcb}_t(x)$. The function lcb_t represents the lower confidence bound of the function, whose proper definition is given in Section 3. Additionally, in many real-world tasks, τ can be chosen to satisfy Assumption 2.2 using domain-specific information. For example, in hyperparameter tuning where the objective $f(x)$ represents model accuracy, $f(\hat{x})$ is upper bounded by 1 (the maximum possible accuracy) and can be lower bounded by the accuracy of the best-known model (e.g., 0.97).

Remark 2.3. Our formulation treats the perturbation as an arbitrary adversarial quantity rather than as a random variable. Therefore, direct comparison with probabilistic robustness measures such as Value-at-Risk (VaR), Conditional Value-at-Risk (CVaR), or distributionally robust optimization (DRO) is not immediate, since these formulations assume access to a known or estimable perturbation distribution. Incorporating such measures would require modifying the problem setting to include a random covariate in the objective, $f(x, \xi)$, and redefining the satisficing condition with respect to the distribution of ξ . While one could, for example, define a variant of robust satisficing conditioned on a reference estimate of this distribution, such an approach restricts the power of the adversary and leads to a less general setting than the one considered here. We emphasize these conceptual differences to clarify the distinction between adversarial and probabilistic robustness.

2.2 GP Regression

It is well established in the literature that a function $f \in \mathcal{H}$ from an RKHS with reproducing kernel $k(\cdot, \cdot)$ with bounded norm $\|f\|_{\mathcal{H}} \leq B$ can be approximated using GPs. Given a dataset of observation pairs $\mathcal{D}_t = \{x_i, y_i\}_{i=1}^t$, using a Gaussian likelihood with variance $\lambda > 0$ and a prior distribution $\text{GP}(0, k(x, x'))$, the posterior mean and covariance of a Gaussian process can be calculated as

$$\mu_t(x) = \mathbf{k}_t(x)^\top (\mathbf{K}_t + \lambda \mathbf{I}_t)^{-1} \mathbf{y}_t, \quad k_t(x, x') = k(x, x') - \mathbf{k}_t(x)^\top (\mathbf{K}_t + \lambda \mathbf{I}_t)^{-1} \mathbf{k}_t(x'), \quad (5)$$

with $\sigma_t^2(x) = k_t(x, x)$, where column vector $\mathbf{k}_t(x) = [k(x_i, x)]_{i=1}^t$, $\mathbf{K}_t$ is the $t \times t$ kernel matrix with elements $(\mathbf{K}_t)_{ij} = k(x_i, x_j)$. These posterior mean and variance functions enable the construction of confidence intervals for the function $f(x)$, which are integral to the optimization strategies we employ. Specifically, in our adversarial setting, these confidence intervals are used to bound the fragility $\kappa_\tau(x)$ and the critical radius $\epsilon_\tau(x)$ of each action, as discussed in the next section.

2.3 Regret Measures

In BO it is common to measure the success of an algorithm using a *regret measure*, which usually measures the discrepancy between the reward of the action played against the reward of the optimal action. We use *lenient regret* and *robust satisficing regret* defined respectively as:

$$R_T^l := \sum_{t=1}^T (\tau - f(x_t))^+, \quad R_T^r := \sum_{t=1}^T (\tau - \kappa_\tau \epsilon_t - f(x_t))^+. \quad (6)$$

Lenient regret, introduced by [11], is a natural metric for satisficing objectives. It captures the cumulative loss with respect to the satisficing threshold. Sublinear lenient regret is unattainable under unconstrained adversarial attacks, as a sufficiently large perturbation budget allows the adversary to force any action below τ . Thus, upper bounds on lenient regret must depend on perturbation magnitude or impose constraints on the budget ϵ_t and threshold τ . Robust satisficing regret [28], instead evaluates loss relative to the RS-1 benchmark, comparing the true reward of x_t against the guaranteed reward of $x^{\text{RS-1}}$ under perturbations of at most ϵ_t . The RS-1 action ensures $f(x^{\text{RS-1}}) \geq \tau - \kappa_\tau \epsilon_t$ under such a budget.

3 Algorithms and Theoretical Results

Confidence Intervals Our algorithms utilize the GP-based upper and lower confidence bounds defined as $\text{ucb}_t(x) = \mu_{t-1}(x) + \beta_t \sigma_{t-1}(x)$ and $\text{lcb}_t(x) = \mu_{t-1}(x) - \beta_t \sigma_{t-1}(x)$, respectively, where β_t is a t -dependent exploration parameter. For theoretical guarantees, β_t is set as in Lemma 3.1, which is based on [37] but argued to be less conservative and more practical [38].

Lemma 3.1. [38, Theorem 1] Let $\zeta \in (0, 1)$, $\bar{\lambda} := \max\{1, \lambda\}$, and

$$\beta_t(\zeta) := B + \frac{R}{\sqrt{\lambda}} \sqrt{\log(\det(\frac{\bar{\lambda}}{\lambda} \mathbf{K}_{t-1} + \bar{\lambda} \mathbf{I}_{t-1})) + 2 \log\left(\frac{1}{\zeta}\right)}.$$

Then, the following holds WPAL $1 - \zeta$: $\text{lcb}_t(x) \leq f(x) \leq \text{ucb}_t(x)$, $\forall x \in \mathcal{X}$, $\forall t \geq 1$.

We label the event in Lemma 3.1 as $\mathcal{E}$ and refer to it as the *good event*. Our analyses will make use of the *maximum information gain* defined as $\gamma_t := \max_{A \subset \mathcal{X}: |A|=t} \frac{1}{2} \log(\det(\mathbf{I}_t + \lambda^{-1} \mathbf{K}_A))$, where $\mathbf{K}_A$ is the kernel matrix given by the sampling set A . Following the work of [34], the use of γ_t is common in the GP literature. Since f is an unknown black-box function, κ_τ for RS-1 and ϵ_τ for RS-2 cannot be calculated directly. The proposed algorithms leverage upper and lower confidence bounds to estimate these values while balancing the exploration-exploitation trade-off.

Algorithm 1 AdveRS-1

- 1: **Input:** Kernel function k , $\mathcal{X}$, τ , R , B , confidence parameter ζ , time horizon T
 - 2: **Initialize:** Set $\mathcal{D}_0 = \emptyset$ (empty dataset), $\mu_0(x) = 0$, $\sigma_0(x) = 1$ for all x
 - 3: **for** $t = 1$ to T **do**
 - 4: Compute $\text{ucb}_t(x) = \mu_{t-1}(x) + \beta_t \sigma_{t-1}(x)$, $\forall x \in \mathcal{X}$
 - 5: Compute $\bar{\kappa}_{\tau,t}(x)$ as in (7), $\forall x \in \mathcal{X}$
 - 6: Select point $\tilde{x}_t = \arg \min_{x \in \mathcal{X}} \bar{\kappa}_{\tau,t}(x)$
 - 7: Adversary selects perturbation $\delta_t \in \Delta_{\epsilon_t}(\tilde{x}_t)$
 - 8: Sample $x_t = \tilde{x}_t + \delta_t$, observe $y_t = f(x_t) + \eta_t$
 - 9: Update dataset $\mathcal{D}_t = \mathcal{D}_{t-1} \cup \{(x_t, y_t)\}$
 - 10: Update GP posterior as in (5)
 - 11: **end for**
-

Algorithm 2 AdveRS-2

- 5: Compute $\bar{\epsilon}_{\tau,t}(x)$ as in (8), $\forall x \in \mathcal{X}$
 - 6: Select point $\tilde{x}_t = \arg \max_{x \in \mathcal{X}} \bar{\epsilon}_{\tau,t}(x)$
-

Algorithm for RS-1. To perform BO with the objective of RS-1, we propose *Adversarially Robust Satisficing-1* (AdveRS-1) algorithm, whose pseudo code is given in Algorithm 1. At the beginning of each round t , AdveRS-1 computes the upper confidence bound $\text{ucb}_t(x)$ of each action $x \in \mathcal{X}$. Using the ucb_t in (2), it computes the *optimistic fragility* defined as:

$$\bar{\kappa}_{\tau,t}(x) := \left(\max_{\delta \in \Delta_\infty(x) \setminus 0} \frac{\tau - \text{ucb}_t(x + \delta)}{d(x, x + \delta)} \right)^+ \quad (7)$$

if $\text{ucb}_t(x) \geq \tau$, and otherwise $\bar{\kappa}_{\tau,t}(x) := \infty$. Further, we define the minimizer $\bar{\kappa}_{\tau,t} := \min_{x \in \mathcal{X}} \bar{\kappa}_{\tau,t}(x)$. The optimistic fragility, $\bar{\kappa}_{\tau,t}(x)$, provides an optimistic estimate of the true fragility $\kappa_\tau(x)$ for each action. AdveRS-1 selects the action with the smallest optimistic fragility, $\tilde{x}_t = \arg \min_{x \in \mathcal{X}} \bar{\kappa}_{\tau,t}(x)$, to encourage exploration. Additionally, we define the *pessimistic fragility* $\underline{\kappa}_{\tau,t}(x)$ which uses lcb instead of the ucb in equation (7), and $\underline{\kappa}_{\tau,t} := \min_{x \in \mathcal{X}} \underline{\kappa}_{\tau,t}(x)$. The following lemma relates the estimated and true fragilities of an action.

Lemma 3.2. Given $\mathcal{E}$ holds, the inequality, $\bar{\kappa}_{\tau,t}(x) \leq \kappa_\tau(x) \leq \underline{\kappa}_{\tau,t}(x)$ is true for all $x \in \mathcal{X}$, $t \geq 1$.

Lemma 3.2 is a result of the monotonicity of fragility result in [17], with the proof provided in the Appendix. The following Corollary which illustrate how $\underline{\kappa}_{\tau,t}(x)$ and $\bar{\kappa}_{\tau,t}(x)$ contribute to quantifying the robustness of an action x at round t , follows from Lemma 3.2 and the RS-1 objective (2).

Corollary 3.3. Under Assumption 2.2, with probability at least $1 - \zeta$, $\forall t \in [T]$ it holds that: $f(\tilde{x}_t + \delta) \geq \tau - \underline{\kappa}_{\tau,t} \cdot d(\tilde{x}_t, \tilde{x}_t + \delta)$, $\forall \delta \in \Delta_\infty(\tilde{x}_t)$; conversely, $\exists \delta' \in \Delta_\infty(\tilde{x}_t)$ s.t. $f(\tilde{x}_t + \delta') \leq \tau - \bar{\kappa}_{\tau,t} \cdot d(\tilde{x}_t, \tilde{x}_t + \delta')$.

Algorithm for RS-2 For the RS-2 objective, we follow a similar path and propose the *Adversarially Robust Satisficing-2* (AdveRS-2) algorithm which is given in Algorithm 2. The algorithm follows

a similar structure to Algorithm 1 with only the lines that differ being shown, along with their corresponding line numbers. AdveRS-2, again, makes use of the confidence bounds to create an optimistic estimate of the critical radius $\epsilon_\tau(x)$ of each action. Specifically, we define the optimistic critical radius as:

$$\bar{\epsilon}_{\tau,t}(x) := \max_{x' \in \mathcal{X}} d(x, x') \text{ s.t. } \text{ucb}(x + \delta) \geq \tau, \quad \forall \delta \in \Delta_{d(x, x')}(x), \quad (8)$$

with $\bar{\epsilon}_{\tau,t} := \max_{x \in \mathcal{X}} \bar{\epsilon}_{\tau,t}(x)$. AdveRS-2 uses the optimistic critical radius for exploration, sampling at each round t the action $\tilde{x}_t = \arg \max_{x \in \mathcal{X}} \bar{\epsilon}_{\tau,t}(x)$, which has the largest optimistic critical radius. Further we define the *pessimistic critical radius* $\underline{\epsilon}_\tau(x)$, which uses the lcb instead of the ucb in equation (8). The following Lemma bounds the true critical radius $\epsilon_\tau(x)$ of each action x .

Lemma 3.4. *Given $\mathcal{E}$ holds, the inequality, $\underline{\epsilon}_{\tau,t}(x) \leq \epsilon_\tau(x) \leq \bar{\epsilon}_{\tau,t}(x)$ is true for all $x \in \mathcal{X}$, $t \geq 1$.*

Corollary 3.5. *Under Assumption 2.2, with probability at least $1 - \zeta$, $\forall t \in [T]$, it holds that $f(\tilde{x}_t + \delta) \geq \tau$, $\forall \delta \in \Delta_{\underline{\epsilon}_{\tau,t}}(\tilde{x}_t)$ and that when $\Delta_{\epsilon_\tau}(\tilde{x}_t)$ is a strict subset of $\Delta_{\bar{\epsilon}_{\tau,t}}(\tilde{x}_t)$, $\exists \delta \in \Delta_{\bar{\epsilon}_{\tau,t}}(\tilde{x}_t)$ s.t. $f(\tilde{x}_t + \delta) \leq \tau$.*

Corollary 3.5 gives information about the achievability of satisficing threshold τ . In practice, a learner might use this information to tune the τ either to a more ambitious or to a more conservative one.

Regret analysis of AdveRS-1 Regret analysis of Algorithm 1 is done with the distance function defined as the kernel metric, $d(x, x') = \sqrt{k(x, x) - 2k(x, x') + k(x', x')}$, and we make use of the Lipschitz continuity of f .

Theorem 3.6. *Let $\zeta \in (0, 1)$, and let $f \in \mathcal{H}$ with $\|f\|_{\mathcal{H}} \leq B$. Let $k(\cdot, \cdot)$ be the reproducing kernel of $\mathcal{H}$, and let η_t be conditionally R -subgaussian. Under Assumption 2.2, when the distance function $d(\cdot, \cdot)$ is the kernel metric and β_t as defined in Lemma 3.1, the lenient regret and robust satisficing regret of AdveRS-1 are bounded above by:*

$$R_T^l \leq 4\beta_T \sqrt{T\gamma_T} + B \sum_{t=1}^T \epsilon_t, \quad R_T^r \leq 4\beta_T \sqrt{T\gamma_T}. \quad (9)$$

Proposition 3.7. *In the context of Theorem 3.6, the linear term $B \sum_{t=1}^T \epsilon_t$ in the lenient regret bound is unavoidable. Specifically, there exists an instance of the problem setup where this term manifests as a lower bound on the lenient regret for any algorithm.*

Regret analysis of AdveRS-2 The RS-1 formulation provides reward guarantees under any perturbation $\delta \in \Delta_\infty(x^{\text{RS-1}})$, ensuring that $f(x^{\text{RS-1}}) \geq \tau - \kappa_\tau \cdot d(x^{\text{RS-1}}, x^{\text{RS-1}} + \delta)$. In contrast, the RS-2 formulation does not provide reward guarantees for every possible perturbation but offers a stronger guarantee, $f(x^{\text{RS-2}}) \geq \tau$, for a restricted set of perturbations $\delta \in \Delta_{\epsilon_\tau}(x^{\text{RS-2}})$ where ϵ_τ is defined in (4). This distinction is reflected in the regret analysis of AdveRS-2. To bound the regret of AdveRS-2, we introduce the following assumption.

Assumption 3.8. Assume that the perturbation budget of the adversary $\epsilon_t \leq \epsilon_\tau$ for all $t \geq 1$.

Remark 3.9. Assumption 3.8, which is stronger than Assumption 2.2, ensures that at each round $t \leq T$, there exists an action $x \in \mathcal{X}$ meeting the threshold τ against any adversary-selected perturbation $\delta_t \in \Delta_{\epsilon_t}(x)$. Without Assumption 3.8, achieving sublinear lenient regret is impossible, as there will always be a perturbation $\delta \in \Delta_{\epsilon_t}(x)$ causing $f(x + \delta) \leq \tau$ for any x .

Theorem 3.10. *Under the assumptions of Theorem 3.6 and Assumption 3.8, both the lenient regret and the robust satisficing regret of AdveRS-2 are bounded above by: $R_T^l, R_T^r \leq 4\beta_T \sqrt{T\gamma_T}$.*

While the lenient regret bound of AdveRS-2 requires further assumptions on the perturbation budget of the adversary, it is also stronger than the regret bound of AdveRS-1 and matches the standard regret bound $\tilde{O}(\gamma_T \sqrt{T})$ of GP-UCB [37]. When the kernel $k(\cdot, \cdot)$ is the RBF kernel or the Matérn- ν kernel, the regret bound reduces to $\tilde{O}(\sqrt{T})$ and $\tilde{O}(T^{\frac{2\nu+3m}{4\nu+2m}})$, respectively, where m is the dimension of the input space [39].

4 Unifying the RS Formulations

The reward guarantee of the RS-1 solution in (2) falls off linearly with the magnitude of the perturbation. On the other hand the reward guarantee of the RS-2 solution stays constant up to a perturbation

magnitude of ϵ_τ , then disappears. In practice one might want to be more flexible with the reward guarantee of their solution. For example one might desire the solution obtained to give a guarantee that stays close to τ for small perturbations, while not giving too much importance to larger perturbations as they may be deemed unlikely, depending on the structure of the problem. Motivated by this, we propose the following robust satisficing formulation RS-General (RS-G). For $p \geq 1$ define the p -fragility of an action as,

$$\kappa_{\tau,p}(x) := \min k \text{ s.t. } f(x + \delta) \geq \tau - [kd(x, x + \delta)]^p, \quad \forall \delta \in \Delta_\infty(x), \quad k \geq 0. \quad (10)$$

The RS-G action is $x^{\text{RS-G}} := \arg \min_{x \in \mathcal{X}} \kappa_{\tau,p}(x)$.

Proposition 4.1. *Let $x^{\text{RS-G}}$ denote the solution obtained from the RS-G formulation when applied with a power parameter p . It holds that $\lim_{p \rightarrow \infty} x^{\text{RS-G}} = x^{\text{RS-2}}$.*

The parameter p controls the trade-off in the RS-G formulation. Specifically, at $p = 1$, the reward guarantee decreases linearly with perturbation magnitude. As p increases, it enhances reward guarantees for small perturbations but diminishes for large perturbations. As p approaches infinity, the formulation converges to RS-2, offering robust guarantees for perturbations within $\Delta_{\epsilon_\tau}(x^{\text{RS-2}})$ but none for those outside it. A more detailed comparison of the RS-G formulation with RS-1 and RS-2 is available in Appendix B. Algorithm 1 can be readily adapted to the RS-G framework by replacing optimistic fragilities with optimistic p -fragilities $\bar{\kappa}_{\tau,p,t}(x)$, which are (10) calculated on $\text{ucb}_t(x)$, instead of $f(x)$. We call this general algorithm AdveRS-G. Similarly, we can generalize our RS regret to the RS-G formulation $R_T^{\text{RS-G}} := \sum_{t=1}^T (\tau - [\kappa_{\tau,p}\epsilon_t]^p - f(x_t))^+$. This measures, at each time step the discrepancy, between the played action, and the reward guarantee of the $x^{\text{RS-G}}$ solution, for a specific p . Finally, we bound the general RS regret in the next corollary.

Corollary 4.2. *Under the assumptions of Theorem 3.6, AdveRS-G satisfies $R_T^{\text{RS-G}} \leq 4\beta_T \sqrt{T\gamma_T}$.*

5 Experiments

In each experiment, we compare the performance of AdveRS-1 and AdveRS-2 against the RO baseline [22], which is run with an ambiguity ball of radius r that is equal to, smaller than, and greater than the true perturbation budget ϵ_t at each round and use the Euclidean distance as the perturbation metric. We use the following attack schemes in our experiment:

- **Random Attack:** $\delta_t \sim \text{Uniform}(\Delta_{\epsilon_t}(\tilde{x}_t))$
- **LCB Attack:** $\delta_t = \arg \min_{\delta \in \Delta_{\epsilon_t}(\tilde{x}_t)} \text{lcb}_t(\tilde{x}_t + \delta)$
- **Worst Case Attack:** $\delta_t = \arg \min_{\delta \in \Delta_{\epsilon_t}(\tilde{x}_t)} f(\tilde{x}_t + \delta)$

The results of all experiments are averaged over 100 runs, with error bars representing $\text{std}/2$. Additional experiments can be found in the Appendix E. Our code is available at: <https://github.com/Bilkent-CYBORG/AdveRS>.

Synthetic experiment In the first experiment, we use the proof-of-concept function in Figure 2, a modified version of the synthetic function from [22]. We set the threshold $\tau = -10$ and conduct two experiments: (a) $\epsilon_t = 0.5$ (Assumption 3.8 holds) and (b) $\epsilon_t = 1.5$ (Assumption 3.8 *fails* to hold). For the RO representative, we run STABLEOPT [22] with radius parameters $r = \epsilon_t$, $r = 4\epsilon_t$, and $r = 0.5\epsilon_t$. The observation noise follows $\eta_t \sim \mathcal{N}(0, 1)$, and the GP kernel is a polynomial kernel trained on 500 samples from the function. Figure 3a shows that STABLEOPT performs poorly when the adversary’s perturbation budget is misestimated, leading to linear regret, while AdveRS-2 consistently meets the threshold τ , achieving sublinear lenient regret. AdveRS-1, though not always reaching τ , still outperforms STABLEOPT when r is misspecified. As depicted in Figure 3b, when Assumption 3.8 fails, all algorithms experience linear regret, consistent with Remark 3.9. This figure also illustrates that AdveRS algorithms maintain robustness towards the goal τ even when τ is unattainable. Additionally, RS-G with $p = 2$ exhibits the smallest regret, suggesting that tuning the p parameter can be beneficial for certain problems. Figure 4 our algorithms consistently beat the RO baseline in RS regret, even when STABLEOPT is run with the perfect knowledge of the adversarial budget.

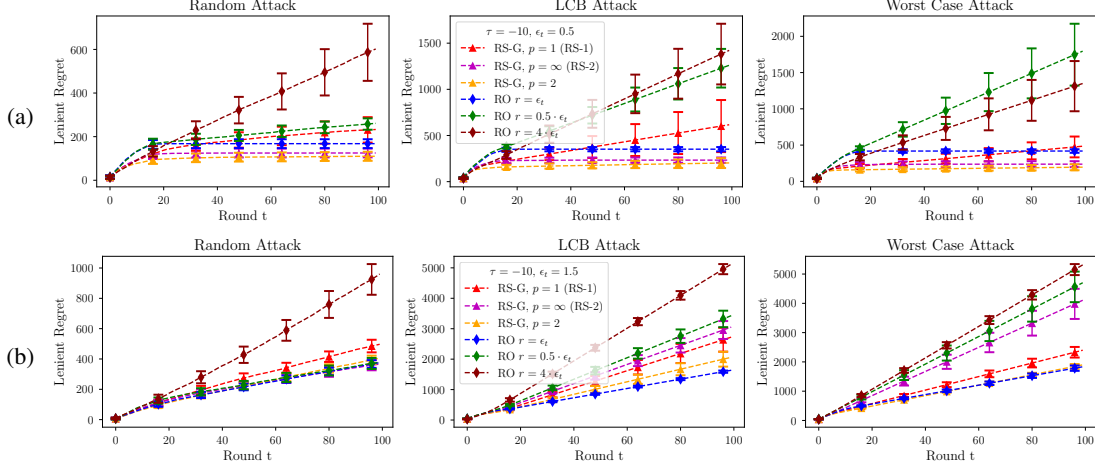


Figure 3: Lenient regret results for synthetic experiment shown in two scenarios: (a) satisfying and (b) failing to satisfy Assumption 3.8.

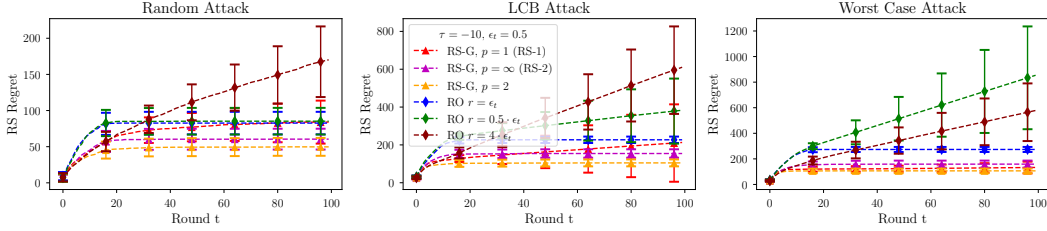


Figure 4: Robust satisficing regret results for synthetic experiment.

Insulin dosage In adversarial contextual GP optimization, the learner selects an action $x_t \in \mathcal{X}$ after observing a reference contextual variable $c_t^{\text{ref}} \in \mathcal{C}$, which is then perturbed by an adversary. Then, the learner observes the perturbed context $c_t = c_t^{\text{ref}} + \delta_t$ along with the noisy observation $f(x_t, c_t) + \eta_t$. We apply this to an insulin dosage selection problem for Type 1 Diabetes Mellitus (T1DM) patients using the UVA/PADOVA T1DM simulator [40]. T1DM patients require bolus insulin administrations usually taken before a meal [41]. Contextual information about the carbohydrate intake is known to improve postprandial blood glucose (PBG) prediction [42].

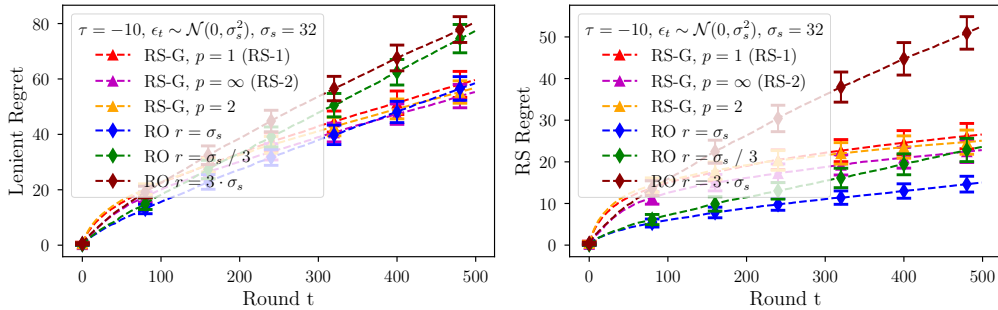


Figure 5: Lenient and RS regret results of insulin dosage experiment.

However, the carbohydrate announcements prior to meal consumption can deviate from the actual carbohydrate intake, hence robustness to contextual shifts are important in diabetes management. Maintaining PBG around a τ neighbourhood of a target level K as $[K - \tau, K + \tau]$ mg/dL represents a satisficing objective defined as $-|f(x) - K| \geq \tau$. The action space is units of insulin in the range $[0, 15]$, and the context is carbohydrate intake, perturbed by $\delta_t \sim \mathcal{N}(0, \sigma_s^2)$ with $\sigma_s = 32$,

representing roughly 2 slices of bread. As shown in Figure 10, although no algorithm achieves sublinear regret, both AdveRS-1 and AdveRS-2 outperform STABLEOPT when the radius r is suboptimal. While STABLEOPT performs best when $r = \sigma_s$, selecting this parameter can be difficult. In contrast, the parameter τ can be more easily chosen by a clinician, aligning with common clinical practice [41].

Robustness to worst case attack We consider the robustness of the solutions from Figure 2 (left), under worst case adversarial perturbations of increasing magnitude. Specifically we inspect $\min_{\delta \in \Delta_\epsilon(x^*)} f(x^* + \delta)$ where x^* corresponds to the points selected by each objective. To evaluate the robustness of a solution w.r.t. achieving τ , we calculate the following:

$$\text{Area}(\epsilon) = \int_0^\epsilon \max \left(0, \tau - \min_{\delta \in \Delta_{\epsilon'}(x^*)} f(x^* + \delta) \right) d\epsilon'.$$

The lower this value is, the better the robustness of the method to a range of different perturbations up to ϵ . A clear discrepancy can be seen between our methods and the RO baseline, with both RS-1 and RS-2 actions demonstrating greater robustness. Further, while RS-2 is better in achieving τ , RS-1 is more robust to discrepancies when τ is not attainable.

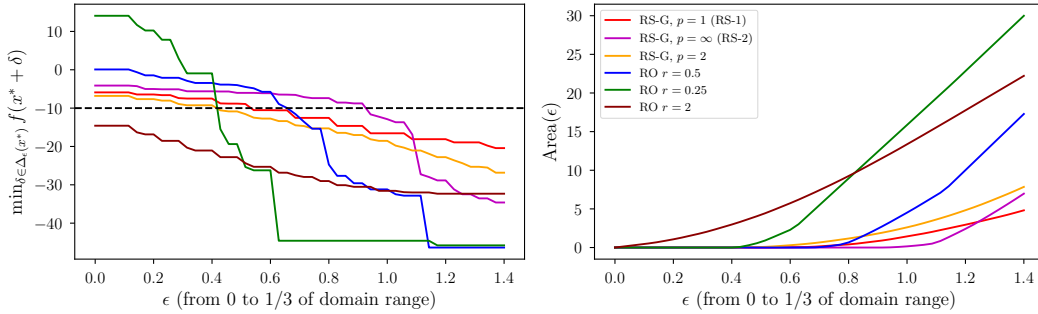


Figure 6: **(Left)** Rewards for the points selected by each algorithm, under worst case attack with perturbation budget ϵ . **(Right)** Area under the curves as a function of perturbation magnitudes.

6 Conclusion

In this paper, we addressed GP optimization under adversarial perturbations, where both the perturbation budget and strategy are unknown and variable. Traditional RO approaches rely on defining an ambiguity set, which is challenging with uncertain adversarial behavior. We introduced a robust satisfying framework with two formulations, RS-1 and RS-2, aimed at achieving a performance threshold τ under these conditions. We also showed these two formulations can be united under a more general RS formulation. The RS-1-based algorithm offers a lenient regret bound that scales with perturbation magnitude and a sublinear robust satisfying regret bound. The RS-2-based algorithm provides sublinear lenient and robust satisfying regret guarantees, assuming τ is achievable. Our experiments show that both algorithms outperform robust optimization methods, particularly when the ambiguity set in RO is misestimated. Despite its promise, our framework has limitations. Selecting τ can be challenging without domain expertise, and the regret bounds for AdveRS-2 assume that τ is achievable. Future work could focus on adaptive τ selection and applying the RS formulations to broader settings, such as distribution shifts and supervised learning. Additionally our framework focuses on the case with known adversarial perturbations. A promising future direction is to develop an RS approach in a setting where the perturbations are not observed. Another future direction would be to learn the parameter p in the RS-G formulation during algorithm run time.

Acknowledgements: This work was supported by the Scientific and Technological Research Council of Türkiye (TÜBİTAK) under Grant 124E065; TÜBİTAK 2024 Incentive Award; by the Turkish Academy of Sciences Distinguished Young Scientist Award Program (TÜBA-GEBİP-2023). Y. Cahit Yıldırım was supported by Turk Telekom as part of 5G and Beyond Joint Graduate Support Programme coordinated by Information and Communication Technologies Authority.

References

- [1] Roman Garnett. *Bayesian optimization*. Cambridge University Press, 2023.
- [2] Jia Wu, Xiu-Yun Chen, Hao Zhang, Li-Dong Xiong, Hang Lei, and Si-Hao Deng. Hyperparameter optimization for machine learning models based on Bayesian optimization. *Journal of Electronic Science and Technology*, 17(1):26–40, 2019.
- [3] Rémi Lam, Matthias Poloczek, Peter Frazier, and Karen E Willcox. Advances in Bayesian optimization with applications in aerospace engineering. In *2018 AIAA Non-Deterministic Approaches Conference*, page 1656, 2018.
- [4] Mahdi Imani and Seyede Fatemeh Ghoreishi. Bayesian optimization objective-based experimental design. In *2020 American Control Conference (ACC)*, pages 3405–3411. IEEE, 2020.
- [5] Ahmed Alaa and Mihaela van der Schaar. Autoprognosis: Automated clinical prognostic modeling via Bayesian optimization with structured kernel learning. In *International Conference on Machine Learning*, pages 139–148. PMLR, 2018.
- [6] Aharon Ben-Tal and Arkadi Nemirovski. Robust optimization—methodology and applications. *Mathematical Programming*, 92:453–480, 2002.
- [7] Hamed Rahimian and Sanjay Mehrotra. Distributionally robust optimization: A review. *arXiv preprint arXiv:1908.05659*, 2019.
- [8] Vishal Gupta. Near-optimal Bayesian ambiguity sets for distributionally robust optimization. *Management Science*, 65(9):4242–4260, 2019.
- [9] Herbert A Simon. A behavioral model of rational choice. *The Quarterly Journal of Economics*, 69(1):99–118, 1955.
- [10] Paul Reverdy, Vaibhav Srivastava, and Naomi Ehrich Leonard. Satisficing in multi-armed bandit problems. *IEEE Transactions on Automatic Control*, 62(8):3788–3803, 2016.
- [11] Nadav Merlis and Shie Mannor. Lenient regret for multi-armed bandits. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 35, pages 8950–8957, 2021.
- [12] Alihan Hüyük and Cem Tekin. Multi-objective multi-armed bandit with lexicographically ordered and satisficing objectives. *Machine Learning*, 110(6):1233–1266, 2021.
- [13] Daniel Russo and Benjamin Van Roy. Satisficing in time-sensitive bandit learning. *Mathematics of Operations Research*, 47(4):2815–2839, 2022.
- [14] Xu Cai, Selwyn Gomes, and Jonathan Scarlett. Lenient regret and good-action identification in Gaussian process bandits. In *International Conference on Machine Learning*, pages 1183–1192. PMLR, 2021.
- [15] Barry Schwartz, Yakov Ben-Haim, and Cliff Dacso. What makes a good decision? robust satisficing as a normative standard of rational decision making. *Journal for the Theory of Social Behaviour*, 41(2):209–227, 2011.
- [16] Andreas L Mogensen and David Thorstad. Tough enough? Robust satisficing as a decision norm for long-term policy analysis. *Synthese*, 200(1):36, 2022.
- [17] Daniel Zhuoyu Long, Melvyn Sim, and Minglong Zhou. Robust satisficing. *Operations Research*, 71(1):61–82, 2023.
- [18] Zhiyuan Lou, Zhi Chen, Melvyn Sim, Jingui Xie, and Peng Xiong. Estimation and prediction procedures for unified robust decision models. *Available at SSRN 4890089*, 2024.
- [19] Christian Szegedy, Wojciech Zaremba, Ilya Sutskever, Joan Bruna, Dumitru Erhan, Ian Goodfellow, and Rob Fergus. Intriguing properties of neural networks. *arXiv preprint arXiv:1312.6199*, 2013.
- [20] Shixiang Gu and Luca Rigazio. Towards deep neural network architectures robust to adversarial examples. *arXiv preprint arXiv:1412.5068*, 2014.

- [21] Ruben Martinez-Cantin, Kevin Tee, and Michael McCourt. Practical Bayesian optimization in the presence of outliers. In *International Conference on Artificial Intelligence and Statistics*, pages 1722–1731. PMLR, 2018.
- [22] Ilija Bogunovic, Jonathan Scarlett, Stefanie Jegelka, and Volkan Cevher. Adversarially robust optimization with Gaussian processes. *Advances in Neural Information Processing Systems*, 31, 2018.
- [23] Rafael Oliveira, Lionel Ott, and Fabio Ramos. Bayesian optimisation under uncertain inputs. In *The 22nd International Conference on Artificial Intelligence and Statistics*, pages 1177–1184. PMLR, 2019.
- [24] Lukas Fröhlich, Edgar Klenske, Julia Vinogradska, Christian Daniel, and Melanie Zeilinger. Noisy-input entropy search for efficient robust Bayesian optimization. In *International Conference on Artificial Intelligence and Statistics*, pages 2262–2272. PMLR, 2020.
- [25] Ilija Bogunovic, Andreas Krause, and Jonathan Scarlett. Corruption-tolerant Gaussian process bandit optimization. In *International Conference on Artificial Intelligence and Statistics*, pages 1071–1081. PMLR, 2020.
- [26] Ilija Bogunovic, Zihan Li, Andreas Krause, and Jonathan Scarlett. A robust phased elimination algorithm for corruption-tolerant Gaussian process bandits. *Advances in Neural Information Processing Systems*, 35:23951–23964, 2022.
- [27] Johannes Kirschner, Ilija Bogunovic, Stefanie Jegelka, and Andreas Krause. Distributionally robust Bayesian optimization. In *International Conference on Artificial Intelligence and Statistics*, pages 2174–2184. PMLR, 2020.
- [28] Artun Saday, Y Cahit Yıldırım, and Cem Tekin. Robust Bayesian satisficing. *Advances in Neural Information Processing Systems*, 36, 2024.
- [29] Sait Cakmak, Raul Astudillo Marban, Peter Frazier, and Enlu Zhou. Bayesian optimization of risk measures. *Advances in Neural Information Processing Systems*, 33:20130–20141, 2020.
- [30] Quoc Phong Nguyen, Zhongxiang Dai, Bryan Kian Hsiang Low, and Patrick Jaillet. Value-at-risk optimization with Gaussian processes. In *International Conference on Machine Learning*, pages 8063–8072. PMLR, 2021.
- [31] Quoc Phong Nguyen, Zhongxiang Dai, Bryan Kian Hsiang Low, and Patrick Jaillet. Optimizing conditional value-at-risk of black-box functions. *Advances in Neural Information Processing Systems*, 34:4170–4180, 2021.
- [32] Shogo Iwazaki, Yu Inatsu, and Ichiro Takeuchi. Mean-variance analysis in Bayesian optimization under uncertainty. In *International Conference on Artificial Intelligence and Statistics*, pages 973–981. PMLR, 2021.
- [33] I. Steinwart and A. Christmann. *Support Vector Machines*. Information Science and Statistics. Springer New York, 2008.
- [34] Niranjan Srinivas, Andreas Krause, Sham M Kakade, and Matthias Seeger. Gaussian process optimization in the bandit setting: No regret and experimental design. In *International Conference on Machine Learning*, 2010.
- [35] Pang Wei Koh, Shiori Sagawa, Henrik Marklund, Sang Michael Xie, Marvin Zhang, Akshay Balsubramani, Weihua Hu, Michihiro Yasunaga, Richard Lanus Phillips, Irena Gao, et al. Wilds: A benchmark of in-the-wild distribution shifts. In *International Conference on Machine Learning*, pages 5637–5664. PMLR, 2021.
- [36] Miriam Zacksenhouse, Simona Nemets, Mikhail A Lebedev, and Miguel AL Nicolelis. Robust satisficing linear regression: Performance/robustness trade-off and consistency criterion. *Mechanical systems and signal processing*, 23(6):1954–1964, 2009.
- [37] Sayak Ray Chowdhury and Aditya Gopalan. On kernelized multi-armed bandits. In *International Conference on Machine Learning*, pages 844–853. PMLR, 2017.

- [38] Christian Fiedler, Carsten W Scherer, and Sebastian Trimpe. Practical and rigorous uncertainty bounds for Gaussian process regression. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 35, pages 7439–7447, 2021.
- [39] Sattar Vakili, Kia Khezeli, and Victor Picheny. On information gain and regret bounds in Gaussian process bandits. In *International Conference on Artificial Intelligence and Statistics*, pages 82–90. PMLR, 2021.
- [40] Boris P Kovatchev, Marc Breton, Chiara Dalla Man, and Claudio Cobelli. In silico preclinical trials: a proof of concept in closed-loop control of type 1 diabetes, 2009.
- [41] Thomas Danne, Moshe Phillip, Bruce A Buckingham, Przemyslaw Jarosz-Chobot, Banshi Saboo, Tatsuhiko Urakami, Tadej Battelino, Ragnar Hanas, and Ethel Codner. Ispad clinical practice consensus guidelines 2018: Insulin treatment in children and adolescents with diabetes. *Pediatric Diabetes*, 19:115–135, 2018.
- [42] Chiara Zecchin, Andrea Facchinetti, Giovanni Sparacino, and Claudio Cobelli. How much is short-term glucose prediction in type 1 diabetes improved by adding insulin delivery and meal content information to cgm data? A proof-of-concept study. *Journal of Diabetes Science and Technology*, 10(5):1149–1160, 2016.
- [43] Alonso Marco, Philipp Hennig, Jeannette Bohg, Stefan Schaal, and Sebastian Trimpe. Automatic lqr tuning based on Gaussian process global optimization. In *2016 IEEE International Conference on Robotics and Automation (ICRA)*, pages 270–277. IEEE, 2016.
- [44] Alonso Marco, Felix Berkenkamp, Philipp Hennig, Angela P Schoellig, Andreas Krause, Stefan Schaal, and Sebastian Trimpe. 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)*, pages 1557–1563. IEEE, 2017.
- [45] William R Thompson. On the likelihood that one unknown probability exceeds another in view of the evidence of two samples. *Biometrika*, 25(3/4):285–294, 1933.
- [46] Daniel J Russo, Benjamin Van Roy, Abbas Kazerouni, Ian Osband, Zheng Wen, et al. A tutorial on Thompson sampling. *Foundations and Trends® in Machine Learning*, 11(1):1–96, 2018.

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: Sections 3, 4, and 5 support the claims made in the abstract and introduction.

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: We discuss the limitations of our work in the conclusion.

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: All assumptions are given in the main paper while proofs are kept in supplementary document.

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: The experimental setup is described in detail in Section 5 to support reproducibility.

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: [Yes]

Justification: We include our code in the supplementary material.

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: Important details are mentioned in Section 5, while minor implementation details are visible in the code.

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: [Yes]

Justification: Error bars are present where suitable.

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 used a modern computer with no special hardware for the experiments, as they do not require intensive computations. Therefore we did not see it necessary to disclose our compute sources.

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: Research conducted conforms to the NeurIPS 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: The authors does not foresee any societal impact of the conducted research.

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 work is theoretical in nature and does not have a foreseeable societal impact.

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: [Yes]

Justification: All copyrights are respected.

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: The 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: 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.
- 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: We did not use LLM's for the core method development.

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.

A Table of Notations

Notation	Description
$\mathcal{X}$	Action set
τ	satisficing level (threshold)
$d(\cdot, \cdot) : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$	Distance metric on $\mathcal{X}$
$\Delta_\epsilon(x)$	$\{x' - x : \mathbf{x}' \in \mathcal{X} \text{ and } d(x, x') \leq \epsilon\}$, set of perturbations in ϵ -ball
$\hat{x}$	$\arg \max_{x \in \mathcal{X}} f(x)$
$\tilde{x}_t$	Action selected by each algorithm
x_t	Point sampled after perturbation, $\tilde{x}_t + \delta_t$
δ_t	Perturbation of the adversary at round t
ϵ_t	True perturbation budget of the adversary at time t
$\kappa_\tau(x)$	$\begin{cases} \left(\max_{\delta \in \Delta_\infty(x) \setminus 0} \frac{\tau - f(x+\delta)}{d(x, x+\delta)} \right)^+ & \text{if } f(x) \geq \tau \\ +\infty & \text{if } f(x) < \tau \end{cases}$
κ_τ	$\min_{x \in \mathcal{X}} \kappa_\tau(x)$
$x^{\text{RS-1}}$	$\arg \min_{x \in \mathcal{X}} \kappa_{\tau,t}(x)$
$\bar{\kappa}_{\tau,t}(x)$	$\begin{cases} \left(\max_{\delta \in \Delta_\infty(x) \setminus 0} \frac{\tau - \text{ucb}_t(x+\delta)}{d(x, x+\delta)} \right)^+ & \text{if } \text{ucb}_t(x) \geq \tau \\ +\infty & \text{if } \text{ucb}_t(x) < \tau \end{cases}$
$\bar{\kappa}_{\tau,t}$	$\min_{x \in \mathcal{X}} \bar{\kappa}_{\tau,t}(x)$
$\underline{\kappa}_{\tau,t}(x)$	$\begin{cases} \left(\max_{\delta \in \Delta_\infty(x) \setminus 0} \frac{\tau - \text{lcb}_t(x+\delta)}{d(x, x+\delta)} \right)^+ & \text{if } \text{lcb}_t(x) \geq \tau \\ +\infty & \text{if } \text{lcb}_t(x) < \tau \end{cases}$
$\underline{\kappa}_{\tau,t}$	$\min_{x \in \mathcal{X}} \underline{\kappa}_{\tau,t}(x)$
$\epsilon_\tau(x)$	$\max_{x' \in \mathcal{X}} d(x, x') \text{ s.t. } f(x+\delta) \geq \tau, \forall \delta \in \Delta_\epsilon(x)$
ϵ_τ	$\max_{x \in \mathcal{X}} \epsilon_\tau(x)$
$x^{\text{RS-2}}$	$\arg \max_{x \in \mathcal{X}} \epsilon_\tau(x)$
$\bar{\epsilon}_{\tau,t}(x)$	$\max_{x' \in \mathcal{X}} d(x, x') \text{ s.t. } \text{ucb}_t(x+\delta) \geq \tau, \forall \delta \in \Delta_\epsilon(x)$
$\bar{\epsilon}_{\tau,t}$	$\max_{x \in \mathcal{X}} \bar{\epsilon}_{\tau,t}(x)$
$\underline{\epsilon}_{\tau,t}(x)$	$\max_{x' \in \mathcal{X}} d(x, x') \text{ s.t. } \text{lcb}_t(x+\delta) \geq \tau, \forall \delta \in \Delta_\epsilon(x)$
$\underline{\epsilon}_{\tau,t}$	$\max_{x \in \mathcal{X}} \underline{\epsilon}_{\tau,t}(x)$

Table 1: Table of notations

B Unified Formulation RS-G

The p -fragility $\kappa_{\tau,p}(x)$ of action $x \in \mathcal{X}$ is

$$\kappa_{\tau,p}(x) := \min k \text{ s.t. } f(x+\delta) \geq \tau - [kd(x, x+\delta)]^p, \forall \delta \in \Delta_\infty(x).$$

This means the RS-G action $x_p^{\text{RS-G}}$ gives the reward guarantee of $f(x+\delta) \geq \tau - [\kappa_{\tau,p} d(x_p^{\text{RS-G}}, x_p^{\text{RS-G}} + \delta)]^p$. When $p = 1$, this guarantee is a linear cone facing downward from τ , at the choosen point. Since the p -fragility is the minimum k that satisfies this inequality, this means the guarantee of the $x_p^{\text{RS-G}}$ action can be represented by the widest cone drawn from τ . When $p = 2$, instead of a linear cone, the guarantee takes quadratic form and so on. We generalize the notion of a cone for all p and define the *fragility-cone* of an action $x' \in \mathcal{X}$ as the following function: $\text{Cone}_{x'}^p(x) := \tau - [\kappa_{\tau,p} d(x', x)]^p$. Then we can say that the solution $x^{\text{RS-G}}$ which minimizes $\kappa_{\tau,p}$, is the solution that maximizes the *wideness* of the fragility-cone, as $\kappa_{\tau,p}$ controls how wide the cone is. Figure 1 gives a visual representation of RS-1, RS-2 and RS-G with $p = 2$ solutions, and their respective fragility-cones. Which formulation should be used depends on the structure of the problem and the goal of the optimizer. If, for example large perturbations are considered unlikely, RS-G with a larger p might be better suited.

C Proofs

C.1 Proof of Lemma 3.2

Assume the good event $\mathcal{E}$ holds. Consider the first inequality $\bar{\kappa}_{\tau,t}(x) \leq \kappa_\tau(x)$. When $\bar{\kappa}_{\tau,t}(x) = \infty$, then $\kappa_\tau(x) = \infty$ as well. When $\bar{\kappa}_{\tau,t}(x) < \infty$ and $\kappa_\tau(x) = \infty$, the inequality holds. Let $\bar{\delta}_{x,t} := \arg \max_{\delta \in \Delta_\infty(x) \setminus 0} \frac{\tau - \text{ucb}_t(x+\delta)}{d(x, x+\delta)}$, $\underline{\delta}_{x,t} := \arg \max_{\delta \in \Delta_\infty(x) \setminus 0} \frac{\tau - \text{lcb}_t(x+\delta)}{d(x, x+\delta)}$ and $\delta_{x,t} := \arg \max_{\delta \in \Delta_\infty(x) \setminus 0} \frac{\tau - f(x+\delta)}{d(x, x+\delta)}$. When $\bar{\kappa}_{\tau,t}(x), \kappa_\tau(x) < \infty$ we have

$$\begin{aligned} \bar{\kappa}_{\tau,t}(x) - \kappa_\tau(x) &= \frac{\tau - \text{ucb}_t(x + \bar{\delta}_{x,t})}{d(x, x + \bar{\delta}_{x,t})} - \frac{\tau - f(x + \delta_{x,t})}{d(x, x + \delta_{x,t})} \\ &\leq \frac{\tau - \text{ucb}_t(x + \bar{\delta}_{x,t})}{d(x, x + \bar{\delta}_{x,t})} - \frac{\tau - f(x + \bar{\delta}_{x,t})}{d(x, x + \bar{\delta}_{x,t})} \\ &= \frac{f(x + \bar{\delta}_{x,t}) - \text{ucb}_t(x + \bar{\delta}_{x,t})}{d(x, x + \bar{\delta}_{x,t})} \\ &\leq 0. \end{aligned}$$

Similarly for $\kappa_\tau(x) \leq \underline{\kappa}_{\tau,t}(x)$ we have

$$\begin{aligned} \kappa_\tau(x) - \underline{\kappa}_{\tau,t}(x) &= \frac{\tau - f(x + \delta_{x,t})}{d(x, x + \delta_{x,t})} - \frac{\tau - \text{lcb}_t(x + \underline{\delta}_{x,t})}{d(x, x + \underline{\delta}_{x,t})} \\ &\leq \frac{\tau - f(x + \delta_{x,t})}{d(x, x + \delta_{x,t})} - \frac{\tau - \text{lcb}_t(x + \delta_{x,t})}{d(x, x + \delta_{x,t})} \\ &= \frac{\text{lcb}_t(x + \delta_{x,t}) - f(x + \delta_{x,t})}{d(x, x + \delta_{x,t})} \\ &\leq 0. \end{aligned}$$

□

C.2 Proof of Lemma 3.4

Assume the good event $\mathcal{E}$ holds. Consider the inequality $\epsilon_\tau(x) \leq \bar{\epsilon}_{\tau,t}(x)$. Assume that this inequality is false for some x , then $\exists \delta$ s.t. $\delta \in \Delta_{\epsilon_\tau(x)}(x)$ and $\delta \notin \Delta_{\bar{\epsilon}_{\tau,t}(x)}(x)$ since $\Delta_{\bar{\epsilon}_{\tau,t}(x)}(x) \subset \Delta_{\epsilon_\tau(x)}(x)$, and $f(x + \delta) \geq \tau > \text{ucb}_t(x + \delta)$, which is a contradiction. Similarly assume $\underline{\epsilon}_{\tau,t}(x) \leq \epsilon_\tau(x)$ is false for some x . Then $\exists \delta$ s.t. $\delta \in \Delta_{\underline{\epsilon}_{\tau,t}(x)}(x)$ and $\delta \notin \Delta_{\epsilon_\tau(x)}(x)$ since $\Delta_{\epsilon_\tau(x)}(x) \subset \Delta_{\underline{\epsilon}_{\tau,t}(x)}(x)$, and $\text{lcb}_t(x + \delta) \geq \tau > f(x + \delta)$, which is a contradiction. Therefore Lemma 3.4 is true.

C.3 Proof of Theorem 3.6

Lenient Regret Bound of AdveRS-1 The regret bound of AdveRS-1 follows a similar structure to the bound in [28], but instead of the MMD metric, we utilize the Lipschitz continuity of f . Define the instantaneous regret at round t as $r_t^l := \tau - f(\tilde{x}_t + \delta_t)$. The cumulative regret is $R_T^l = \sum_{t=1}^T (r_t^l)^+$. Assume that the good event $\mathcal{E}$ holds. Then the instantaneous regret can be bounded as

$$r_t^l = \tau - f(\tilde{x}_t + \delta_t) \leq \tau - \text{ucb}_t(\tilde{x}_t + \delta_t) + 2\beta_t \sigma_t(\tilde{x}_t + \delta_t). \quad (11)$$

If the adversary does not perturb the action, i.e., $\delta_t = 0$, then we have

$$r_t^l \leq \tau - \text{ucb}_t(\tilde{x}_t) + 2\beta_t \sigma_t(\tilde{x}_t)$$

by Assumption 2.2 and the selection rule of the algorithm we have $\text{ucb}_t(\tilde{x}_t) \geq \tau$, hence

$$r_t^l \leq 2\beta_t \sigma_t(\tilde{x}_t) = 2\beta_t \sigma_t(x_t).$$

If $\delta_t \neq 0$, continuing from (11),

$$r_t^l \leq \tau - \text{ucb}_t(\tilde{x}_t + \delta_t) + 2\beta_t \sigma_t(\tilde{x}_t + \delta_t) \quad (12)$$

$$\leq d(\tilde{x}_t, \tilde{x}_t + \delta_t) \frac{\tau - \text{ucb}_t(\tilde{x}_t + \delta_t)}{d(\tilde{x}_t, \tilde{x}_t + \delta_t)} + 2\beta_t \sigma_t(\tilde{x}_t + \delta_t) \quad (13)$$

$$\leq d(\tilde{x}_t, \tilde{x}_t + \delta_t) \bar{\kappa}_{\tau,t}(\tilde{x}_t) + 2\beta_t \sigma_t(\tilde{x}_t + \delta_t) \quad (14)$$

$$\leq d(\tilde{x}_t, \tilde{x}_t + \delta_t) \bar{\kappa}_{\tau,t}(\hat{x}) + 2\beta_t \sigma_t(\tilde{x}_t + \delta_t) \quad (15)$$

Define $\bar{\delta}_{x,t} := \arg \max_{\delta \in \Delta_\infty(x) \setminus 0} \frac{\tau - \text{ucb}_t(x+\delta)}{d(x, x+\delta)}$. Then we can write above as,

$$= d(\tilde{x}_t, \tilde{x}_t + \delta_t) \frac{\tau - \text{ucb}_t(\hat{x} + \bar{\delta}_{\hat{x},t})}{d(\hat{x}, \hat{x} + \bar{\delta}_{\hat{x},t})} + 2\beta_t \sigma_t(\tilde{x}_t + \delta_t) \quad (16)$$

$$\leq d(\tilde{x}_t, \tilde{x}_t + \delta_t) \frac{f(\hat{x}) - f(\hat{x} + \bar{\delta}_{\hat{x},t})}{d(\hat{x}, \hat{x} + \bar{\delta}_{\hat{x},t})} + 2\beta_t \sigma_t(\tilde{x}_t + \delta_t) \quad (17)$$

$$\leq d(\tilde{x}_t, \tilde{x}_t + \delta_t) \frac{Bd(\hat{x}, \hat{x} + \bar{\delta}_{\hat{x},t})}{d(\hat{x}, \hat{x} + \bar{\delta}_{\hat{x},t})} + 2\beta_t \sigma_t(\tilde{x}_t + \delta_t) \quad (18)$$

$$\leq 2\beta_t \sigma_t(\tilde{x}_t + \delta_t) + B\epsilon_t = 2\beta_t \sigma_t(x_t) + B\epsilon_t. \quad (19)$$

(14) comes from the definition (7) and (15) follows from $\tilde{x}_t$ being the minimizer of $\bar{\kappa}_{\tau,t}(x)$. (17) follows from Assumption 2.2 and the confidence bounds. Finally (18) follows from the Lipschitz continuity of f with respect to the kernel metric d , and (19) from $d(\tilde{x}_t, \tilde{x}_t + \delta_t) \leq \epsilon_t$. From this point on, we bound the lenient regret by following standard steps for bounding regret of GP bandits. First note that since (19) is nonnegative, it is also a bound for $(r_t^l)^+$. By plugging this upper bound to (6), and using monotonicity of β_t , we obtain

$$R_T^l \leq 2\beta_T \sum_{t=1}^T \sigma_t(x_t) + B \sum_{t=1}^T \epsilon_t \quad (20)$$

$$\leq 2\beta_T \sqrt{T \sum_{t=1}^T \sigma_t(x_t)^2} + B \sum_{t=1}^T \epsilon_t \quad (21)$$

where (21) uses the Cauchy-Schwartz inequality. Next, to relate the sum of variances that appears in (21) to the maximum information gain, we use $\alpha \leq 2\log(1 + \alpha)$ for all $\alpha \in [0, 1]$ to obtain

$$\begin{aligned} \sum_{t=1}^T \sigma_t(x_t)^2 &\leq \sum_{t=1}^T 2\log(1 + \sigma_t(x_t)^2) \\ &\leq 4\gamma_T, \end{aligned} \quad (22)$$

where (22) follows from [37, Lemma 3]. Putting (22) in the regret definition we get the final bound

$$R_T^l \leq 4\beta_T \sqrt{T\gamma_T} + B \sum_{t=1}^T \epsilon_t.$$

□

Robust Satisficing Regret Bound of AdveRS-1 Define the instantaneous regret at round t as $r_t^{\text{rs}} := \tau - \kappa_\tau \epsilon_t - f(\tilde{x}_t + \delta_t)$. The cumulative regret is $R_T^{\text{rs}} = \sum_{t=1}^T (r_t^{\text{rs}})^+$. Assume that the good event $\mathcal{E}$ holds. Then the instantaneous regret can be bounded as

$$r_\tau^{\text{rs}} \leq \tau - \kappa_\tau \epsilon_t - \text{ucb}_t(\tilde{x}_t + \delta_t) + 2\beta_t \sigma_t(\tilde{x}_t + \delta_t) \quad (23)$$

$$\leq \tau - \kappa_\tau \epsilon_t - (\tau - \bar{\kappa}_{\tau,t}(\tilde{x}_t) d(\tilde{x}_t, \tilde{x}_t + \delta_t)) + 2\beta_t \sigma_t(\tilde{x}_t + \delta_t) \quad (24)$$

$$\leq \tau - \kappa_\tau \epsilon_t - (\tau - \bar{\kappa}_{\tau,t}(\tilde{x}_t) \epsilon_t) + 2\beta_t \sigma_t(\tilde{x}_t + \delta_t) \quad (25)$$

$$= \epsilon_t (\bar{\kappa}_{\tau,t} - \kappa_\tau) + 2\beta_t \sigma_t(\tilde{x}_t + \delta_t) \quad (26)$$

$$\leq 2\beta_t \sigma_t(\tilde{x}_t + \delta_t) = 2\beta_t \sigma_t(x_t). \quad (27)$$

(24) follows from the guarantee of the RS formulation $\text{ucb}_t(\tilde{x}_t + \delta) \geq \tau - \bar{\kappa}_{\tau,t} d(\tilde{x}_t, \tilde{x}_t + \delta)$ for any $\delta \in \Delta_\infty(\tilde{x}_t)$. (27) follows from $\bar{\kappa}_{\tau,t} \leq \kappa_{\tau,t}$ since $\text{ucb}_t(x) \geq f(x)$ for all x . Then we bound the cumulative robust satisficing regret, following the same standard steps as in the proof of the cumulative lenient regret bound, to obtain an upper bound of:

$$R_T^{\text{rs}} \leq 4\beta_T \sqrt{T\gamma_T}.$$

□

C.4 Proof of Theorem 3.10

Under the good event $\mathcal{E}$, Assumption 2.2 and Assumption 3.8, we can bound the lenient regret of AdveRS-2 as:

$$r_\tau^l := \tau - f(\tilde{x}_t + \delta_t) \quad (28)$$

$$\leq \min_{\delta \in \Delta_{\bar{\epsilon}_{\tau,t}}(\tilde{x}_t)} \text{ucb}_t(\tilde{x}_t + \delta) - f(\tilde{x}_t + \delta_t) \quad (29)$$

$$\leq \min_{\delta \in \Delta_{\epsilon_\tau}(\tilde{x}_t)} \text{ucb}_t(\tilde{x}_t + \delta) - f(\tilde{x}_t + \delta_t) \quad (30)$$

$$\leq \min_{\delta \in \Delta_{\epsilon_t}(\tilde{x}_t)} \text{ucb}_t(\tilde{x}_t + \delta) - f(\tilde{x}_t + \delta_t) \quad (31)$$

$$\leq \text{ucb}_t(\tilde{x}_t + \delta_t) - f(\tilde{x}_t + \delta_t) \quad (32)$$

$$\leq 2\beta_t \sigma_t(x_t). \quad (33)$$

Note that from Assumption 2.2, under $\mathcal{E}$ we have $\text{ucb}_t(\hat{x}) \geq \tau$ which means $\bar{\epsilon}_{\tau,t}(\hat{x}) \geq 0$. Then by definition, for any $\delta \in \Delta_{\bar{\epsilon}_t}(\tilde{x}_t)$, $\text{ucb}_t(\tilde{x}_t + \delta_t) \geq \tau$, which gives (29). (30) is true because $\epsilon_\tau \leq \bar{\epsilon}_{\tau,t}(x^{\text{RS-2}})$ by Lemma 3.4 and $\bar{\epsilon}_{\tau,t}(x^{\text{RS-2}}) \leq \bar{\epsilon}_{\tau,t}$ by the equation (8). Hence $\Delta_{\epsilon_\tau}(\tilde{x}_t) \subseteq \Delta_{\bar{\epsilon}_{\tau,t}}(\tilde{x}_t)$. Similarly, (31) follows from Assumption 3.8. For (32), observe that $\delta_t \in \Delta_{\epsilon_t}(\tilde{x}_t)$, no matter the attacking strategy of the adversary. Finally (33) holds under the good event $\mathcal{E}$. The rest of the proof is the same as the previous ones. Noting that $R^{\text{rs}} \leq R^l$ by definition, and hence shares the same upper bound completes the proof of Theorem 3.10. $\square$

C.5 Proof of Corollary 3.3

By the definition of $\kappa_{\tau,t}(x)$, we have $\text{lcb}_t(\tilde{x}_t + \delta) \geq \tau - \kappa_{\tau,t} d(\tilde{x}_t, \tilde{x}_t + \delta) \quad \forall \delta \in \Delta_\infty(\tilde{x}_t)$. Noting that under $\mathcal{E}$ $f(x) \geq \text{lcb}_t(x) \quad \forall x \in \mathcal{X}$, we obtain the first inequality. For the second inequality observe that from the definition of $\bar{\kappa}_{\tau,t}(x)$ we know that $\exists \delta' \in \Delta_\infty(\tilde{x}_t)$ such that $\text{ucb}_t(\tilde{x}_t + \delta') = \tau - \bar{\kappa}_{\tau,t} d(\tilde{x}_t, \tilde{x}_t + \delta')$. Noting that $f(x) \leq \text{ucb}_t(x)$ for all x under $\mathcal{E}$ concludes the proof.

C.6 Proof of Corollary 3.5

By Lemma 3.4 we have $\Delta_{\epsilon_{\tau,t}}(\tilde{x}_t) \subseteq \Delta_{\epsilon_\tau}(\tilde{x}_t)$, then for any $\delta \in \Delta_{\epsilon_{\tau,t}}(\tilde{x}_t)$, we have $f(\tilde{x}_t + \delta) \geq \text{lcb}_t(\tilde{x}_t + \delta) \geq \tau$ under the selection rule and the good event $\mathcal{E}$. Conversely again by Lemma 3.4 we have $\Delta_{\epsilon_\tau}(\tilde{x}_t) \subseteq \Delta_{\bar{\epsilon}_{\tau,t}}(\tilde{x}_t)$. Assuming $\mathcal{X}$ is continuous, $\exists \delta' \in \Delta_{\epsilon_\tau}$ such that $f(\tilde{x}_t + \delta') = \tau$. Noting that $\delta' \in \Delta_{\bar{\epsilon}_{\tau,t}}$ completes the proof.

C.7 Proof of Proposition 3.7

Consider the following instance of our problem. Let $\mathcal{X} = [0, 1]$ and $f(x) = x$, noting that this is an element of RKHS with a linear kernel. The Lipschitz constant of this function is 1. Let the satisficing goal $\tau = 1$ which is the function maximum. If $\epsilon_t < 1$, no matter what action is chosen by the learner, the adversary can choose a perturbation $\delta_t = -\epsilon_t$ to obtain a reward $f(x_t + \delta_t) \leq \tau - \epsilon_t$, making the instantaneous lenient regret $r^l = (\tau - f(x_t + \delta_t))^+ \geq \epsilon_t$. This constitute a worst case lower bound, hence the bound is tight and the linear penalty term is not avoidable in general.

C.8 Proof of Proposition 4.1

WLOG assume that $\exists \delta \in \Delta_\infty(x)$ such that $f(x + \delta) \leq \tau$. Let $\Delta'(x) = \{\delta \in \Delta_\infty(x) \mid f(x + \delta) < \tau\}$. Notice that (10) is then equivalent to

$$\kappa_{\tau,p}(x) := \min k \text{ s.t. } f(x + \delta) \geq \tau - [kd(x, x + \delta)]^p, \quad \forall \delta \in \Delta'(x),$$

since for any $\delta \notin \Delta'(x)$, the inequality constraint already holds for any $k \geq 0$.

Then we can write the equivalent formulation for an action that satisfies $f(x) \geq \tau$

$$\kappa_{\tau,p}(x) := \sup_{\delta \in \Delta'(x)} \frac{[\tau - f(x + \delta)]^{1/p}}{d(x, x + \delta)}.$$

As $p \rightarrow \infty$, note that for each fixed x and $\delta \in \Delta'(x)$, $[\tau - f(x + \delta)]^{1/p} \downarrow 1$ monotonically. Hence, by monotone convergence,

$$\begin{aligned}\kappa_{\tau,p}(x) &= \sup_{\delta \in \Delta'(x)} \frac{[\tau - f(x + \delta)]^{1/p}}{d(x, x + \delta)} \downarrow \sup_{\delta \in \Delta'(x)} \frac{1}{d(x, x + \delta)}, \\ \frac{1}{\kappa_{\tau,p}(x)} &= \min_{\delta \in \Delta'(x)} d(x, x + \delta) = \epsilon_\tau(x).\end{aligned}$$

Since $\kappa_{\tau,p}(x)$ decreases pointwise to $1/\epsilon_\tau(x)$ for all x , we can exchange the limit and the minimization by uniform convergence:

$$\lim_{p \rightarrow \infty} \min_{x \in \mathcal{X}} \kappa_{\tau,p}(x) = \min_{x \in \mathcal{X}} \lim_{p \rightarrow \infty} \kappa_{\tau,p}(x) = \min_{x \in \mathcal{X}} \frac{1}{\epsilon_\tau(x)} = \max_{x \in \mathcal{X}} \epsilon_\tau(x) = x^{\text{RS-2}}.$$

C.9 Proof of Corollary 4.2

The proof is almost identical to the RS regret proof for AdveRS-1, we include it for completeness. Define the instantaneous regret at round t as $r_t^{\text{rs-g}} := \tau - [\kappa_\tau \epsilon_t]^p - f(\tilde{x}_t + \delta_t)$. The cumulative regret is $R_T^{\text{rs-g}} = \sum_{t=1}^T (r_t^{\text{rs-g}})^+$. Assume that the good event $\mathcal{E}$ holds. Then the instantaneous regret can be bounded as

$$r_t^{\text{rs-g}} \leq \tau - [\kappa_\tau \epsilon_t]^p - \text{ucb}_t(\tilde{x}_t + \delta_t) + 2\beta_t \sigma_t(\tilde{x}_t + \delta_t) \quad (34)$$

$$\leq \tau - [\kappa_\tau \epsilon_t]^p - (\tau - [\bar{\kappa}_{\tau,p,t}(\tilde{x}_t) d(\tilde{x}_t, \tilde{x}_t + \delta_t)]^p) + 2\beta_t \sigma_t(\tilde{x}_t + \delta_t) \quad (35)$$

$$\leq \tau - [\kappa_\tau \epsilon_t]^p - (\tau - [\bar{\kappa}_{\tau,p,t}(\tilde{x}_t) \epsilon_t]^p) + 2\beta_t \sigma_t(\tilde{x}_t + \delta_t) \quad (36)$$

$$= \epsilon_t^p (\bar{\kappa}_{\tau,p,t}^p - \kappa_\tau^p) + 2\beta_t \sigma_t(\tilde{x}_t + \delta_t) \quad (37)$$

$$\leq 2\beta_t \sigma_t(\tilde{x}_t + \delta_t) = 2\beta_t \sigma_t(x_t). \quad (38)$$

(35) follows from the guarantee of the RS-G formulation $\text{ucb}_t(\tilde{x}_t + \delta) \geq \tau - [\bar{\kappa}_{\tau,p,t} d(\tilde{x}_t, \tilde{x}_t + \delta)]^p$ for any $\delta \in \Delta_\infty(\tilde{x}_t)$. (37) follows from $\bar{\kappa}_{\tau,p,t} \leq \kappa_{\tau,p,t}$ since $\text{ucb}_t(x) \geq f(x)$ for all x . Then we bound the cumulative robust satisficing regret, following the same standard steps to obtain an upper bound of:

$$R_T^{\text{rs-g}} \leq 4\beta_T \sqrt{T\gamma_T}.$$

□

D Implementation details of the algorithms

We note that maximizing UCB and LCB over a compact domain is non-trivial due to the non-convex nature of the posterior mean and variance. Like foundational works (e.g., [34]), we focus on theoretical guarantees and assume an oracle optimizer, abstracting computational details. This challenge is common across all GP-based methods, including the RO algorithms we compare against, and in practical BO, various optimization (e.g., gradient descent or quasi-Newton) techniques are used to address it.

In our experiments, we work with discretized domains and the implementation of the acquisition functions of our algorithms have complexity $\mathcal{O}(N^2)$ where $N = |\mathcal{X}|$. Notably, this complexity is the same as that of RO-based algorithms. Specifically, in all algorithms that we implement (RS or RO based), action-specific calculations are performed in an inner loop, while an outer loop is used the selection of best action.

E Additional Experimental Results

E.1 Inverted Pendulum Experiment

In this experiment, we focus on optimizing a parameterized controller to achieve robust performance in controlling an inverted pendulum. The state vector $\mathbf{s}_k$ consists of four variables: cart position, pendulum angle, and their derivatives. The simulation is conducted using OpenAI's Gym environment.

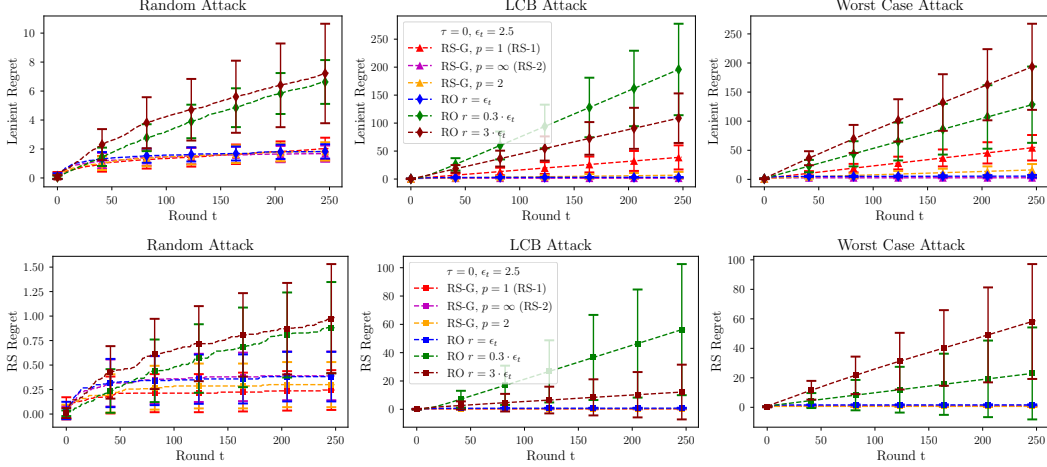


Figure 7: Lenient and RS regret results of the inverted pendulum experiment. Perturbation budget is $\epsilon_t = 2.5$ and $\tau = 0$ for all t . STABLEOPT run with: $r = \epsilon_t$, $r = 3\epsilon_t$, and $r = 0.3\epsilon_t$. Attack schemes are indicated in the subtitles.

The system dynamics are described by $\mathbf{s}_{k+1} = h(\mathbf{s}_k, u_k)$, where $\mathbf{s}_k$ is the 4-dimensional state vector, and $u_k \in [-10, 10]$ is the control action representing the voltage applied to the controller at step k . We implement a static feedback controller similar to the one in [43], defined as $u_k = \mathbf{F}\mathbf{s}_k$, where $\mathbf{F} \in \mathbb{R}^{1 \times 4}$ is the gain matrix. We structure $\mathbf{F}$ as a Linear Quadratic Regulator (LQR) using the linearized system dynamics $(\mathbf{A}, \mathbf{B})$ around the equilibrium point. The parameterization, which has been shown to be learnable by GP's [43] is given by $\mathbf{F} = \text{dlqr}(\mathbf{A}, \mathbf{B}, \mathbf{W}_s(\theta), \mathbf{W}_u(\theta))$ with the following parameterization:

$$\begin{aligned} \mathbf{W}_s(\theta) &= \text{diag}(10^{\theta_1}, 10^{\theta_2}, 10^{\theta_3}, 0.1), & \theta_{1,2,3} &\in [-3, 2], \\ \mathbf{W}_u(\theta) &= 10^{-\theta_4}, & \theta_4 &\in [1, 5]. \end{aligned}$$

We discretize the parameter space into 4096 points and compute the cost for each parameter by simulating it over 2000 seconds, using the same cost function as in [44]. After the simulation, the cost function is standardized and clipped to the range $[-2, 2]$. For the GP kernel, we use a Radial Basis Function (RBF) kernel with Automatic Relevance Determination (ARD), with hyperparameters selected using 400 samples prior to the experiment.

For adversarial perturbations, we consider an adversary that damages the learning process by perturbing the selected parameter. The unknown perturbation budget is constant with $\epsilon_t = 2$. The satisfying goal is set to $\tau = 0$ which corresponds to ~ 25 percentile of the objective function. Figure 7 shows the regrets of the algorithms over a time horizon of $T = 250$. AdveRS family of algorithms matches the STABLEOPT that has the perfect knowledge of the adversarial budget, and they outperform STABLEOPT otherwise. While AdveRS-1 achieves linear regret, it can still perform better than STABLEOPT when the perturbation budget is misspecified.

E.2 Misspecified Kernel

In this experiment, we evaluate the performance of our algorithms when the GP is run with a kernel that is not aligned with the true function. Specifically, we use the function shown in Figure 2, but unlike Experiment 1, we replace the polynomial kernel with an RBF kernel, using a lengthscale of 0.1 and a variance of 10. All algorithms are run with the same parameters as in Experiment 1 with $\tau = -10$ and $\epsilon_t = 0.5$. Figure 8 shows that AdveRS-2 maintains strong performance even when the kernel is misspecified, while AdveRS-1 outperforms STABLEOPT when the ambiguity ball is not correctly estimated. Figure 9 shows the RS regrets and we see that our algorithms achieve sublinear results, while STABLEOPT with misspecified ambiguity radius can achieve linear RS regret. All plots show the mean results over 100 independent runs with error bars showing std/2.

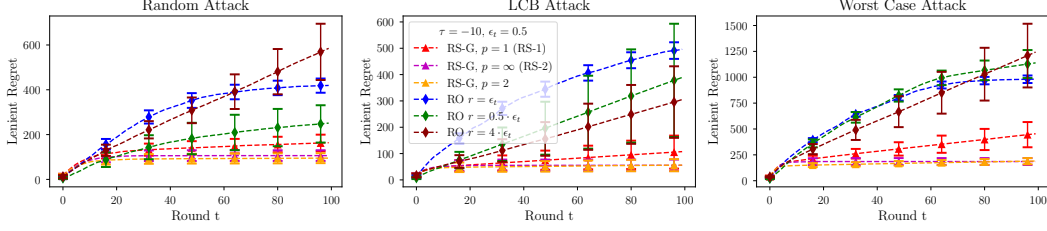


Figure 8: Lenient regret results for synthetic experiment using a misspecified kernel, with a perturbation budget of $\epsilon_t = 0.5$ and $\tau = -10$ for all t . STABLEOPT run with: $r = \epsilon_t$, $r = 4\epsilon_t$, and $r = 0.5\epsilon_t$.

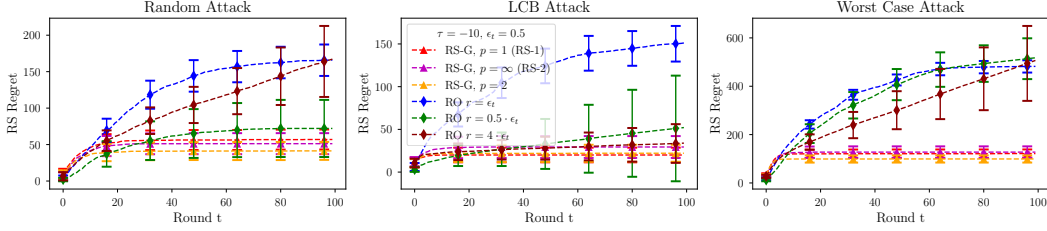


Figure 9: Robust satisficing regret results for synthetic experiment using a misspecified kernel, with a perturbation budget of $\epsilon_t = 0.5$ and $\tau = -10$ for all t . STABLEOPT run with: $r = \epsilon_t$, $r = 4\epsilon_t$, and $r = 0.5\epsilon_t$.

E.3 Robust Satisficing with Thompson sampling

Thompson Sampling. Thompson sampling [45, 46] selects the next point by sampling a function $\tilde{f}$ from the GP posterior and then choosing

$$x_t \in \arg \max_{x \in \mathcal{X}} \tilde{f}(x).$$

This randomized policy implicitly balances exploration and exploitation by occasionally sampling from regions of high uncertainty.

Algorithm 3 AdveRS-G-TS

- 1: **Input:** Kernel function k , $\mathcal{X}$, τ , p , time horizon T
 - 2: **Initialize:** $\mathcal{D}_0 = \emptyset$ (empty dataset), and $\mu_0(x) = 0$, $\sigma_0(x) = 1 \quad \forall x \in \mathcal{X}$
 - 3: **for** $t = 1$ to T **do**
 - 4: Sample $\tilde{f}_t(x) \sim \mathcal{GP}(\mu_{t-1}(x), \sigma_{t-1}^2(x))$, $\forall x \in \mathcal{X}$
 - 5: Compute $\tilde{\kappa}_{\tau,p,t}(x)$ using (39), $\forall x \in \mathcal{X}$
 - 6: Select $\tilde{x}_t = \arg \min_{x \in \mathcal{X}} \tilde{\kappa}_{\tau,t}(x)$
 - 7: Adversary selects perturbation $\delta_t \in \Delta_{\epsilon_t}(\tilde{x}_t)$
 - 8: Sample $x_t = \tilde{x}_t + \delta_t$, observe $y_t = f(x_t) + \eta_t$
 - 9: $\mathcal{D}_t = \mathcal{D}_{t-1} \cup \{(x_t, y_t)\}$
 - 10: Update GP posterior using (5)
 - 11: **end for**
-

Algorithm for RS-G using Thompson sampling. In order to perform Bayesian optimization with the objective as RS-G, building on the algorithm family of AdveRS, we propose *Adversarially Robust Satisficing-General-Thompson Sampling* (AdveRS-G-TS) algorithm (see Algorithm 3).

At the beginning of each round t , AdveRS-G-TS samples a random representative $\tilde{f}_t$ from the posterior distribution of GP. Then, replacing the f in (10) instead with $\tilde{f}_t$ and rearranging as in , it

computes the Thompson p -fragility defined as:

$$\bar{\kappa}_{\tau,p,t}(x) := \left(\max_{\delta \in \Delta_{\infty}(x) \setminus 0} \frac{\tau - \tilde{f}_t(x + \delta)^{1/p}}{d(x, x + \delta)} \right)^+, \quad (39)$$

if $\tilde{f}_t(x) \geq \tau$, and otherwise $\bar{\kappa}_{\tau,p,t}(x) := \infty$. For each action, the Thompson p -fragility $\bar{\kappa}_{\tau,p,t}(x)$ provides a random sample of the true fragility $\kappa_{\tau,p}(x)$. AdveRS-G-TS then selects the action with the smallest random fragility, $\tilde{x}_t = \arg \min_{x \in \mathcal{X}} \bar{\kappa}_{\tau,p,t}(x)$, to guide exploration and exploitation.

Insulin dosage experiment using AdveRS-G-TS

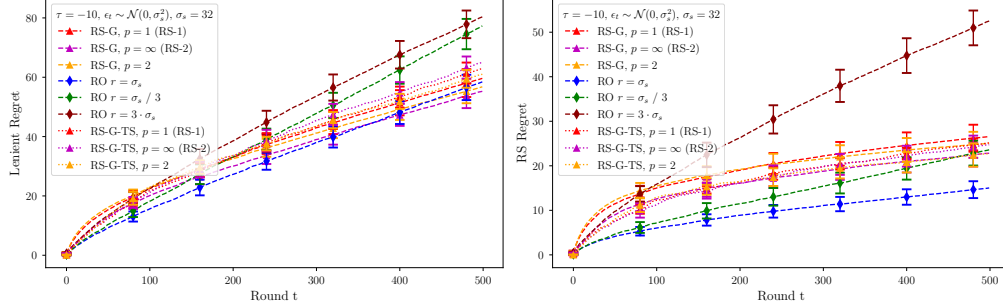


Figure 10: Lenient and RS regret results of insulin dosage experiment. Note that the results for algorithms other than AdveRS-G-TS are the same as in 10 and are averaged over 100 random runs, while AdveRS-G-TS results are averaged over 20 runs.

E.4 Effect of τ and r

Table 2 shows the varying effects of the parameters. The lenient regret of all algorithms are calculated on the synthetic function from Experiment 1. The perturbations are sampled i.i.d. from a normal distribution $\mathcal{N}(0, 0.3^2)$. For this experiment, we report results from 20 random runs. Results show that the target oriented nature of RS approaches makes them better suited for the satisficing objective under unknown perturbations.

Alg. \ τ	-30.00	-20.00	-10.00	0.00	10.00
RS1	103.83	201.97	287.81	651.38	1478.62
RS2	47.09	104.01	185.39	434.30	626.48
RSG	52.70	113.03	177.90	619.17	1505.98
RO, $r=0.10$	112.29	207.81	335.10	458.96	645.07
RO, $r=0.57$	94.20	180.79	293.17	537.80	1489.18
RO, $r=1.05$	74.30	138.19	224.45	723.06	1695.44
RO, $r=1.52$	55.05	126.99	258.48	1053.13	2121.94
RO, $r=2.00$	60.05	284.83	722.58	1523.14	2721.97

Table 2: Average lenient regrets for synthetic experiment for RS-1, RS-2 and RO algorithms with varying τ and r values. Best values are in bold for each τ value.