

Excursion Search for Constrained Bayesian Optimization under a Limited Budget of Failures

Alonso Marco¹ Alexander von Rohr^{1,2} Dominik Baumann¹ José Miguel Hernández-Lobato³
Sebastian Trimpe^{1,4}

Abstract

When learning to ride a bike, a child falls down a number of times before achieving the first success. As falling down usually has only mild consequences, it can be seen as a tolerable failure in exchange for a faster learning process, as it provides rich information about an undesired behavior. In the context of Bayesian optimization under unknown constraints (BOC), typical strategies for safe learning explore conservatively and avoid failures by all means. On the other side of the spectrum, non conservative BOC algorithms that allow failing may fail an unbounded number of times before reaching the optimum. In this work, we propose a novel decision maker grounded in control theory that controls the amount of risk we allow in the search as a function of a given budget of failures. Empirical validation shows that our algorithm uses the failures budget more efficiently in a variety of optimization experiments, and generally achieves lower regret, than state-of-the-art methods. In addition, we propose an original algorithm for unconstrained Bayesian optimization inspired by the notion of excursion sets in stochastic processes, upon which the failures-aware algorithm is built.

machine learning (ML) algorithms in real-world scenarios is an ongoing challenge. A key difficulty lies in the proper management of undesired outcomes, which are inevitable when learning under unknown or uncertain circumstances. As an extreme case, in applications like autonomous driving, a failure in the decision-making may lead to human casualties. Such safety-critical scenarios need conservative ML algorithms, which forbid any failures. On the other hand, there exist scenarios in which failures are still undesired, although might not come at a high cost. For example, when deploying ML algorithms to optimize the parameters of an industrial drilling machine to drill faster, a few configurations might break the drill bits, but in exchange, a faster drilling can be learned. In such non-safety-critical applications, failures shall be considered as a valuable source of knowledge, and one would tolerate a limited number of them in exchange for better learning performance.

When iteratively improving machine parameters directly from data, the mapping between a specific parameter configuration and the corresponding behavior of the machine is often unknown, and can only be revealed through experiments. Normally, such experiments are time-consuming, and thus, data collection is considered expensive. In order to learn the optima of expensive black box functions, Bayesian optimization (BO) has been established in the last decade as a promising probabilistic framework (Shahriari et al., 2016). Therein, the aim is to efficiently exploit the observed data in combination with prior probabilistic models to estimate the global optimum from a few trials. In the context of robot learning, BO has been used to mitigate the effort of manual controller tuning, see, e.g., (Calandra et al., 2016; von Rohr et al., 2018; Rai et al., 2018).

When the optimization is subject to unknown external restrictions, the goal is to solve a constrained optimization problem under multiple black box constraints. (Hernández-Lobato et al., 2016; Gelbart et al., 2014; Gardner et al., 2014; Gramacy & Lee, 2011; Schonlau et al., 1998; Picheny, 2014) propose different BO methods to estimate the constrained global optimum. In (Lam & Willcox, 2017), a variant of such problem is considered, where the total budget of evaluations is explicitly included in the decision-making, by

1. Introduction

Deploying machine learning (ML) algorithms in real-world scenarios has gained increasing interest during the last decade. Under some circumstances, lacking from sufficiently accurate models, or knowledge of the environment, such algorithms can lead to undesired outcomes. Deploying

¹Max Planck Institute for Intelligent Systems, Tübingen, Germany ²IAV GmbH, Germany ³Department of Engineering, University of Cambridge, Cambridge, UK ⁴Institute for Data Science in Mechanical Engineering, RWTH Aachen University, Aachen, Germany. Correspondence to: Alonso Marco <amarco@tue.mpg.de>.

formulating the problem as a dynamic programming instance. Because these methods do not have a limit on the number of incurred failures, they can fail many times. In other words, none of them inform the decision maker about the remaining budget of failures at each iteration.

From a different perspective, zero-budget strategies (Sui et al., 2015; Berkenkamp et al., 2016) are needed in safety-critical applications, where failures are not allowed. Such strategies avoid failures by conservatively expanding an initially given safe area, and never exploring beyond the learned safety boundaries. However, when applied in a context where failures are allowed, such strategies become suboptimal: they will ignore such budget and miss alternative, potentially more promising, safe areas, located outside the initial safe area.

In this work, we pose the problem of learning the constrained global optimum in settings where a non-zero budget of failures is given. In particular, we make two main contributions. Our first contribution is a failures-aware strategy for BOC that, in contrast to prior work, does not need to be initialized in a safe region and that makes decisions taking into account the budgets of remaining failures and evaluations.

Our second contribution is a novel acquisition function inspired by key notions of the geometry of excursion sets in stochastic processes. In (Adler & Taylor, 2009), an excursion set is defined over smooth manifolds as those points for which a process realization crosses upwards a given threshold. The larger the threshold, the more likely it is that an *upcrossing* will reveal the location of the global maximum. Based on this intuition, we derive an acquisition function, which can be written analytically, is cheap to evaluate, and explicitly includes the process derivative to make optimal decisions.

In the following, we explain and experimentally validate the aforementioned contributions. In Section 2, we characterize excursion sets in Gaussian processes (GP), and explain their benefits when used in BO. In Section 3, we formalize the proposed novel acquisition function to solve unconstrained problems. In Section 4, such acquisition is extended for the constrained case in the presence of a budget of failures. In Section 5, we validate both acquisition functions empirically on common benchmarks for global optimization and real-world applications. We conclude with a discussion in Section 6.

2. Excursion sets in Bayesian optimization

The proposed search strategy is inspired by the study of the differential and integral geometry of excursion sets in stochastic processes (Adler & Taylor, 2009). In the particular case of GPs, analytical expressions can be derived for

such sets. In the following, we provide the needed mathematical tools and intuition over which our search strategy is constructed.

2.1. Problem formulation

The main goal is to address the unconstrained optimization problem

$$x_* = \operatorname{argmin}_{x \in \mathcal{X}} f(x), \quad (1)$$

where the objective $f : \mathcal{X} \rightarrow \mathbb{R}$ is a black-box function, which evaluations are corrupted by noise and are expensive to collect (due to, e.g., energetic costs), and $\mathcal{X} \subset \mathbb{R}^D$.

2.2. Gaussian process (GP)

We model the objective as a Gaussian process, $f \sim \mathcal{GP}(0, k(x, \hat{x}))$, with covariance function $k : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$, and zero prior mean. Observations $y(x) = f(x) + \varepsilon$ are modeled using additive Gaussian noise $\varepsilon \sim \mathcal{N}(\varepsilon; 0, \sigma_n^2)$. After having collected t observations from the objective $\mathcal{D}_t^f = \{\mathbf{x}_t, \mathbf{y}_t\} = \{x_1, \dots, x_t, y_1, \dots, y_t\}$, its predictive distribution at a location x is given by $p(f|\mathcal{D}_t^f, x) = \mathcal{N}(f(x); \mu(x|\mathcal{D}_t^f), \sigma^2(x|\mathcal{D}_t^f))$, with predictive mean $\mu(x|\mathcal{D}_t^f) = \mathbf{k}_t^\top(x)[K_t + \sigma_n^2 I]^{-1} \mathbf{y}_t$, where the entries of vector $\mathbf{k}_t(x)$ are $[\mathbf{k}_t(x)]_i = k(x_i, x)$, the entries of the Gram matrix K_t are $[K_t]_{i,j} = k(x_i, x_j)$, and the entries of the vector of observations $\mathbf{y}_t$ are $[\mathbf{y}_t]_i = y_i$. The predictive variance is given by $\sigma^2(x|\mathcal{D}_t^f) = k(x, x) - \mathbf{k}_t^\top(x)[K_t + \sigma_n^2 I]^{-1} \mathbf{k}_t(x)$. In the remainder of the paper, we drop the dependency on the current data set $\mathcal{D}_t^f$ and write $\mu(x)$, $\sigma(x)$ to refer to $\mu(x|\mathcal{D}_t^f)$, $\sigma(x|\mathcal{D}_t^f)$, respectively.

2.3. Excursion sets in Gaussian processes

Let us assume a zero-mean scalar Gaussian process f , with $\mathcal{X} = [0, 1]^D$, $D = 1$, and stationary covariance function $k(\tau) = k(\|x - \hat{x}\|_2)$. The excursion set $\{x \in \mathcal{X} : f(x) \geq u\}$ is defined as the set of locations where the process f is above the threshold u . In (Adler & Taylor, 2009, Part II. Geometry), such sets are characterized by the number of upcrossings of process samples through the level u , i.e., $N_u^+ = \#\{x \in \mathcal{X} : f(x) = u, f'(x) > 0\}$, where $f'(x)$ is the derivative of the process. Intuitively, large N_u^+ represents a high frequency of upcrossings, which is connected with having many areas in $\mathcal{X}$ where $f(x)$ lives above u . For a one-dimensional, stationary, almost surely continuous and mean-square differentiable Gaussian process, the expected number of upcrossings (Rasmussen & Williams, 2006, Sec. 4.1) is given by the well-known Rice's formula (Lindgren, 2006, Sec. 3.1.2)

$$\mathbb{E}[N_u^+] = \int_0^1 \mathbb{E}_{p(f, f'|x)} [f' : f = u, f' > 0] dx \quad (2)$$

$$\begin{aligned}
 &= \int_0^1 \int_{-\infty}^{+\infty} \int_0^{+\infty} f' \delta(f - u) p(f, f'|x) df' df dx \\
 &= \frac{1}{2\pi} \sqrt{\frac{-k''(0)}{k(0)}} \exp\left(-\frac{u^2}{2k(0)}\right),
 \end{aligned}$$

where $p(f, f'|x)$ is the joint density of the process and its derivative, both queried at location x , δ is the Dirac delta, and the second derivative of the covariance function k'' must exist. Interestingly, (2) can be used to approximate the probability of finding the supremum of a process realization above a high level u . The growth rate of the approximation error with respect to u is bounded

$$\left| \mathbb{E}[N_u^+] - \Pr\left(\sup_{x \in [0,1]} f(x) \geq u\right) \right| < O(e^{-\beta u^2/k(0)}), \quad (3)$$

as $u \rightarrow \infty$, with $O(\cdot)$ indicating the limiting behavior of the approximation error and $\beta > 1$ needs to be found (Adler & Taylor, 2009, Sec. 14). The intuitive reasoning behind this is simple: If f crosses a high level u , it is unlikely to do so more than once. Therefore, the probability that f meets its supremum above u is close to the probability that there is an upcrossing of u . Since the number of upcrossings of a high level will always be small, the probability of an upcrossing is well approximated by $\mathbb{E}[N_u^+]$.

While the bound in (3) does not hold for the general case $D > 1$, we use it as a starting point to build a new acquisition function for $D \geq 1$ (cf. Section 2.5), which shows empirically superior results than state-of-the-art BO methods. In the following section, we show, for $D = 1$, how $\mathbb{E}[N_u^+]$ can be leveraged to lead the search towards areas where the number of upcrossings is large, or equivalently, where the global maximum is more likely to be found. Thereafter, we extend the result for $D \geq 1$.

2.4. Practical interpretation for use in BO

The expected number of upcrossings (2) contains valuable information about the amount of times a sample realization of the process z “upcrosses” the level u . However, (2) cannot be used directly for decision-making because it is a global property of the process itself, rather than a local quantity at a specific location x . Next, we provide a practical interpretation that relaxes some of the assumptions made to obtain (2) and allows for its use in BO. To this end, we introduce three modifications.

First, when seeking for the optimum of the process, it is more useful to consider both, the up- and down-crossings through the level u , as both of them occur near the optimum when u is large. This quantity is defined in (Lindgren, 2006, Sec. 3.1.2) as the expected number of *crossings*

$$\mathbb{E}[N_u] = \int_0^1 \mathbb{E}_{p(f, f'|x)} [|f'| : f = u] dx \quad (4)$$

$$= \int_0^1 \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} |f'| \delta(f - u) p(f, f'|x) df' df dx,$$

with $N_u = \#\{x \in [0, 1] : f(x) = u\}$.

Second, BO uses pointwise information to decide on how interesting it is to explore a specific location x . (Lindgren, 2006, Theorem 3.1) proposes the *intensity of expected crossings* $\mathbb{E}[N_u(x)]$, which can be computed by simply removing the domain integral in (4)

$$\mathbb{E}[N_u(x)] = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} |f'| \delta(f - u) p(f, f'|x) df' df. \quad (5)$$

Third, when conditioning the Gaussian process f on observed data $\mathcal{D}_t^f$, it becomes non-stationary¹, and thus, the predictive distribution of a query $f(x)$ changes as a function of x . The dependency on $\mathcal{D}_t^f$ is introduced in (5) following (Lindgren, 2006, Remark 3.2), as

$$\mathbb{E}[N_u(x|\mathcal{D}_t^f)] = \int_{-\infty}^{+\infty} |f'| p(u, f'|x, \mathcal{D}_t^f) df', \quad (6)$$

where the joint density is evaluated at $f = u$ after resolving the integral over the Dirac delta. We next provide a brief analysis for solving (6).

Using the rule of conditional probability, we have $p(u, f'|x, \mathcal{D}_t^f) = p(u|x, \mathcal{D}_t^f) p(f'|u, x, \mathcal{D}_t^f)$. The first term, $p(u|x, \mathcal{D}_t^f) = \mathcal{N}(u; \mu(x), \sigma^2(x))$, is a Gaussian density² evaluated at u , with the predictive mean and variance of the GP model. The second term is also a Gaussian density over the process derivative, conditioned on $f = u$. This can be seen as adding a *virtual* observation u at location x to existing data set. Hence, $p(f'|x, u, \mathcal{D}_t^f) = p(f'|\mathcal{D}_t^f \cup \{x, u\}) = \mathcal{N}(f'; \mu'(x), \nu^2(x))$. Then, (6) can be rewritten as

$$\begin{aligned}
 &p(u|x, \mathcal{D}_t^f) \int_{-\infty}^{+\infty} |f'| p(f'|\mathcal{D}_t^f \cup \{x, u\}) df' = \\
 &\mathcal{N}(u; \mu(x), \sigma^2(x)) \left(2\nu(x) \phi(\gamma(x)) + \mu'(x) \operatorname{erf}\left(\frac{\gamma(x)}{\sqrt{2}}\right) \right),
 \end{aligned} \quad (7)$$

where $\gamma(x) = \mu'(x)/\nu(x)$, ϕ is the probability density function of a standard normal distribution, and $\operatorname{erf}(\cdot)$ is the error function (see Appendix A for a complete derivation). Fig. 1 shows $\mathbb{E}[N_u(x|\mathcal{D}_t^f)]$ for two different values of u , where the GP is conditioned on seven observations. As can be seen, different thresholds imply different intensity of crossings for the same process. When the threshold is near

¹Note that all GPs are non-stationary when conditioned on data, even if the covariance function that defines them is stationary.

²Using simplified notation, we write $p(u|x, \mathcal{D}_t^f)$ to refer to the density function $p_{f|x, \mathcal{D}_t^f}(\xi)$ evaluated at $\xi = u$. Similarly, we write $p(u, f'|x, \mathcal{D}_t^f)$ to refer to the joint density function $p_{f, f'|x, \mathcal{D}_t^f}(\xi, \zeta)$ evaluated at $\xi = u$ for some value ζ .

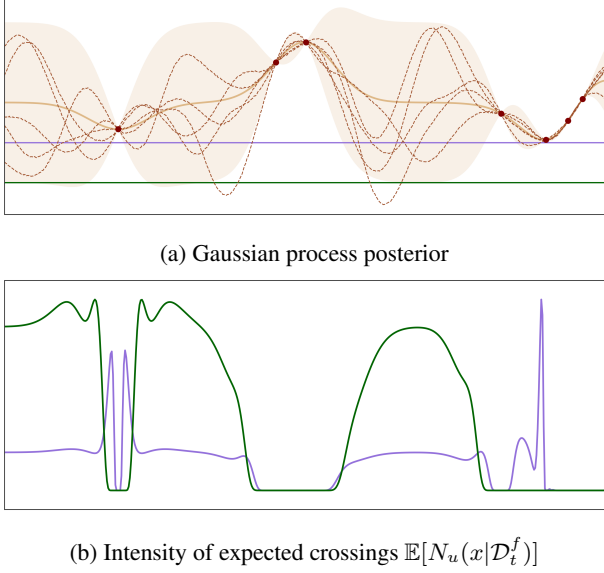


Figure 1: (a) Gaussian process posterior conditioned on a set of observations. Given a process realization (dashed lines), two choices for the threshold u (solid lines) determine two different excursion sets. (b) Intensity of expected crossings $\mathbb{E}[N_u(x|\mathcal{D}_t^f)]$ for each threshold u . Higher values correspond to areas where the boundaries of the excursion sets are likely to be, i.e., where the process is more likely to cross u . The curves are normalized to have the same maximum value.

collected evaluations, the largest intensity of crossings tends to be concentrated near the data. On the contrary, when it is far from the data, the largest intensity of crossings is found in areas of large variance.

2.5. Extension to D dimensions

Although (7) was derived for $D = 1$, we can extend it to the case $D \geq 1$. Since (6) depends on $|f'|$, a natural extension is to consider the L-1 norm of the gradient of the process $\|\nabla f(x)\|_1 = \sum_{j=1}^D |\frac{\partial f(x)}{\partial x_j}|$. Following this, we extend (7) as $\mathbb{E}_{p(f(x), \nabla f(x))}[\|\nabla f(x)\|_1 : f(x) = u, \mathcal{D}_t^f]$,

$$\mathbb{E}[N_u(x|\mathcal{D}_t^f)] \simeq \mathcal{N}(u; \mu(x), \sigma^2(x)) \times \sum_{j=1}^D \left(2\nu_j(x) \phi(\gamma_j(x)) + \mu_j(x) \operatorname{erf}\left(\frac{\gamma_j(x)}{\sqrt{2}}\right) \right), \quad (8)$$

where $\gamma_j(x) = \mu_j(x)/\nu_j(x)$. The gradient $\nabla f(x) \sim \mathcal{N}(\nabla f(x); \nabla \mu(x), V(x))$ follows a multivariate Gaussian, and $\mu_j(x) = [\nabla \mu(x)]_j$ and $\nu_j(x) = ([V(x)]_{jj})^{1/2} = \sqrt{\partial^2 k(x_j, x_j) / \partial x_j^2}$. Note that $\nabla \mu(x)$ and $V(x)$ depend on the extended data set $\mathcal{D}_t^f \cup \{x, u\}$.

In the following sections, we propose two novel algorithms that build upon the quantity (8).

3. Excursion search algorithm

The modifications applied to (2), detailed above, allow extracting useful information about how likely is the process f to cross a certain level u at each location x . When u is a lower bound on the collected data, (8) reveals locations where the process is more likely to have a minimum. If we repeatedly evaluate at such locations, one would expect to approach faster the global minimum. In the following, we characterize (8) as an acquisition function for optimal decision-making.

3.1. Threshold of crossings as the global minimum

The choice of the threshold u in (8) is important when trying to find the global minimum. A hypothetically appropriate low value for u is right above the global minimum $f_* = f(x_*)$, i.e., $u = f_* + \epsilon$, where $\epsilon > 0$ is small. Then, if crossings through $u = f_* + \epsilon$ are likely to occur at a specific area, we know that such area is likely to contain the global minimum, and thus, will show a large $\mathbb{E}[N_u(x|\mathcal{D}_t^f)]$. However, in practice we do not have access to the true f_* of the objective function, and thus, cannot compute u in the aforementioned way. At most, we are able to assume a distribution over the global minimum $f_* \sim p(f_*)$, implied by the GP model on f . In the following, we assume that u follows such distribution, i.e., $u \sim p(u) = p(f_*)$.

It is well-known in extreme value theory (De Haan & Ferreira, 2007) that f_* follows one of the three extreme value distributions: Gumbel, Fréchet, or Weibull, which generally model tails distributions. For example, in (Wang & Jegelka, 2017), the Gumbel distribution is chosen to model $p(f_*)$. However, such distribution has infinite support, while in practice it is not useful to have any probability mass above the best observed evaluation $\eta = \min(y(x_1), \dots, y(x_T))$. Instead, we consider the Fréchet distribution as a more appropriate choice as it provides finite support $f_* \leq \eta$. For minimization problems, we can define it in terms of its survival function $\mathcal{F}_{s,q}(a) = \Pr(f_* \geq a)$, given by

$$\mathcal{F}_{s,q}(a) = \begin{cases} 0, & \text{if } a > \eta \\ \exp\left(-\left(\frac{\eta-a}{s}\right)^{-q}\right), & \text{if } a \leq \eta \end{cases} \quad (9)$$

where $\Pr(f_* \geq a) = \int_a^{+\infty} p(f_*) df_*$, and the parameters $s > 0$ and $q > 1$ can be estimated from data following the same approach as in (Wang & Jegelka, 2017, Appendix B). A thorough analysis on the advantage of using the Fréchet distribution, instead of the Gumbel distribution, for gathering samples of f_* can be found in Appendix B. Using the above definition, the stochastic threshold $u \sim p(u) = p(f_*)$, makes the quantity (8) also stochas-

tic. We propose to compute its expectation over u , i.e., $\mathbb{E}_{p(u)}[\mathbb{E}[N_u(x|\mathcal{D}_t^f)]] = \mathbb{E}_{p(f_*)}[\mathbb{E}[N_{f_*}(x|\mathcal{D}_t^f)]]$, which we explain next.

3.2. Acquisition function

We define the *excursion search* (Xs) acquisition function as

$$\alpha_X(x) = \mathbb{E}_{p(f_*)} \left[\mathbb{E} \left[N_{f_*}(x|\mathcal{D}_t^f) \right] \right] \quad (10)$$

$$\simeq \frac{1}{S} \sum_{l=1}^S \mathbb{E} \left[N_{f_*^l}(x|\mathcal{D}_t^f) \right],$$

where the outer expectation is intractable and is approximated via sampling. For each sample $f_*^l \sim p(f_*)$, (8) needs to be recomputed. The samples can be collected through the inverse of (9), $f_*^l = \mathcal{F}_{s,q}^{-1}(\xi^l)$. $\xi^l \sim U(0,1)$ follows a uniform distribution in the unit interval, and $\mathcal{F}_{s,q}^{-1}(\xi^l) = \eta - s(-\log(1 - \xi^l))^{-1/q}$.

Intuitively, the Xs acquisition function (10) reveals areas near the global maximum (i.e., where the gradient crosses the estimated f_* with large norm), instead of directly aiming at potential maximums, minimums, or saddle points. Furthermore, Xs inherently trades off exploration with exploitation: At early stages of the search, the estimated Fréchet distribution (9) reflects large uncertainty about f_* , which causes the samples f_*^l to lie far from the data. Hence, exploration is encouraged, as shown in Fig. 1 (green lines). At later stages, when more data is available, the Fréchet distribution (9) shrinks toward the lowest observations, which then encourages exploitation, as shown in Fig. 1 (violet lines).

The acquisition (10) is our first contribution, and can be used for unconstrained optimization problems, e.g., (1).

4. Bayesian optimization with a limited budget of failures

In the previous section, we introduced a new acquisition function (10) grounded in the connection between the true optimum of the process f and the expected number of crossings through its current estimate (cf. (3)). However, such acquisition does not explicitly have into account any budget of failures B or evaluations T . In the following, we propose an algorithm that makes use of B and T to balance the decision making between (i) safely exploring encountered safe areas, and (ii) searching outside the safe areas at the risk of failing, when safe areas contain no further information.

4.1. Problem formulation

To the unconstrained problem (1), we add G black-box constraints, $g_j : \mathcal{X} \rightarrow \mathbb{R}$, $j = \{1, \dots, G\}$, also corrupted by noise and expensive to evaluate. Moreover, we assume a

non-safety critical scenario, where violating the constraints is allowed, but it is strictly forbidden to do so more than B times. Analogously, we allow only for a maximum number of $T \geq B$ evaluations. The case $T < B$ is not considered herein, as the budget of failures can simply be ignored. Under these conditions, we formulate the constrained optimization problem with limited budget of failures as

$$x_*^c = \underset{x \in \mathcal{X}}{\operatorname{argmin}} f(x), \text{ s.t. } g_1(x) \leq 0, \dots, g_G(x) \leq 0$$

$$\text{under failures } \sum_{t=1}^T \Gamma(x_t) \leq B, \quad (11)$$

where x_*^c is the location of the constrained minimum, and $\Gamma(x_t) = \mathbb{I}[g_1(x_t) > 0 \vee \dots \vee g_G(x_t) > 0]$ equals 1 if at least one of the constraints is violated at location x_t , and 0 otherwise. $\mathbb{I}$ is the indicator function, and $g(x_1), \dots, g(x_T)$ are the collected evaluations of the constraints at locations $x_1, \dots, x_T$. Since the constraints g_j are unknown, and modeled as independent Gaussian processes $g_j \sim \mathcal{GP}(0, k(x, \hat{x}))$, queries $f(x)$ and $g(x)$ are stochastic and (11) cannot be solved directly. Instead, we address the analogous probabilistic formulation from (Gelbart et al., 2014):

$$x_*^c \simeq \underset{x \in \mathcal{X}}{\operatorname{argmin}} \mu(x), \text{ s.t. } \prod_{j=1}^G \Pr(g_j(x) \leq 0) \geq \rho$$

$$\text{under failures } \sum_{t=1}^T \Gamma(x_t) \leq B, \quad (12)$$

where $\Pr(g_j(x) \leq 0) = \Phi(-\mu_j(x)/\sigma_j(x))$, Φ is the cumulative density function of a standard normal distribution, and $\rho \in (0, 1)$ is typically set close to one. The predictive mean μ_j and variance σ_j^2 conditioned on $\mathcal{D}_t^{g_j}$ of each g_j are computed as in Section 2.2. In the following, we provide a novel Bayesian optimization strategy to address (12).

4.2. Safe exploration with dynamic control

In order to include the probability of constraint satisfaction in the decision making, we propose a similar approach to (Gelbart et al., 2014) by explicitly adding a probabilistic constraint to the search of the next evaluation

$$x_{\text{next}} = \underset{x \in \mathcal{X}}{\operatorname{argmax}} \alpha_X(x)$$

$$\text{s.t. } \prod_{i=1}^K \Pr(g_j(x) \leq 0) \geq \rho_t, \quad (13)$$

where the parameter $\rho_t \in (0, 1)$ determines how much we are willing to tolerate constraint violation at each iteration t . This leads the search away from areas where the constraint is likely to be violated, as those areas get revealed during the search.

Contrary to (Gelbart et al., 2014), where ρ_t is fixed a priori, we propose to choose it at each iteration, depending on the remaining budget of failures $\Delta B_t = B - \sum_{j=1}^t \Gamma(x_j)$ and remaining iterations $\Delta T_t = T - t$. Intuitively, the more failures we have left (large ΔB_t), the more we are willing to tolerate constraint violation (large ρ_t). We achieve this by proposing an automatic control law to drive ρ_t , which we describe next.

Let us define a latent variable $z_t = \Phi^{-1}(\rho_t)$, $z_t \in \mathbb{R}$ that follows a deterministic process $z_{t+1} = z_t + u_t$, using a dynamic feedback controller $u_t = u_t(\Delta B_t, \Delta T_t)$. Such controller drives the process toward one of the two references: $z_{\text{safe}} = \Phi^{-1}(\rho_{\text{safe}})$ and $z_{\text{risk}} = \Phi^{-1}(\rho_{\text{risk}})$, where typical values are $\rho_{\text{safe}} = 0.99$ and $\rho_{\text{risk}} = 0.01$. We define a control law

$$u_t = (z_{\text{safe}} - z_t) \frac{\Gamma(x_t)}{\Delta B_t} + (z_{\text{risk}} - z_t) \frac{\Delta B_t}{2\Delta T_t}, \quad (14)$$

with $\Delta B_t > 0$, $\Delta T_t > 0$, and $\Delta B_t \leq \Delta T_t$. The first term drives the process toward z_{safe} when a failure occurs at iteration t , with intensity $1/\Delta B_t$. In this way, the fewer failures are left in the budget, the more urgently the process chases z_{safe} . The second term attempts to push z_t down to z_{risk} with an intensity proportional to the ratio between the remaining failures and iterations.

When $\Delta B_t = 0$, but $\Delta T_t > 0$, only a conservative safe exploration is allowed. To do so, we set $u_t = (z_{\text{safe}} - z_t)$ for the remaining iterations until $t = T$. Additionally, if there are more failures left than remaining iterations, i.e., $\Delta B_t > \Delta T_t$, the remaining budget of failures is not decisive for decision making, and thus, we set $u_t = (z_{\text{risk}} - z_t)$.

The resulting control strategy weights risky versus conservative decision-making by considering the budget of evaluations and iterations left: When no failures occur for a few consecutive iterations, ρ_t is slowly driven toward ρ_{risk} , and when a failure takes place, it lifts up ρ_t toward ρ_{safe} .

4.3. Risky search of new safe areas

The probabilistic constraint in (13) puts a hard constraint on the decision making by not allowing evaluations in regions that are known to be unsafe. When ρ_t is high, (13) will discard regions where no data has been collected and locally explore regions where safe evaluations are present. Such conservative decision making is desirable when $\Delta B_t \ll \Delta T_t$ because it avoids unsafe evaluations. The smaller the ρ_t , the more risky evaluations we can afford, which makes the constraint information less important in the decision making. However, when ρ_t is too low, the probabilistic constraint tends to be ignored, and the decisions are based on the information from the objective. Albeit this indeed counts as the wanted risky exploration strategy, completely ignoring the constraint information could result in repeated evaluations in unsafe areas. To avoid this, we follow the

approach from (Gelbart et al., 2014), where the acquisition function is aware of the constraint information, without this being a hard constraint. Therein, locations are chosen at

$$x_{\text{next}} = \operatorname{argmax}_{x \in \mathcal{X}} \alpha_{\mathcal{X}}(x) \prod_{j=1}^D \Pr(g_j(x) \leq 0). \quad (15)$$

This approach ‘‘jumps’’ outside the current safe areas at the risk of failing, while the multiplying term discourages exploration in areas revealed to be unsafe.

Trading off risky versus safe exploration depends on the remaining budget ΔB_t , and is quantified by ρ_t , as detailed in Section 4.2. We propose a user-defined decision boundary ρ_b , such that if $\rho_t \leq \rho_b$, the next location will be selected using (15), and (13) otherwise.

While (13) assumes that a safe area has already been found, this might not be the case at an early stage of the search. In such case, we collect observations using (15) and only resort to the risk versus safety trade-off once a safe area has been found.

Pseudocode for the overall framework, named *failures-aware excursion search* (XsF), and an analysis of its computational complexity can be found in Appendix C. XsF returns the estimated location of the constrained minimum x_*^c from (12), computed by setting $\rho = \rho_{\text{safe}}$.

5. Empirical analysis and validation

We empirically validate Xs and XsF by comparing their performance against state-of-the-art methods. We consider three different scenarios. In the first one, we validate each method on common challenging benchmarks for global optimization. In the second and third scenarios we compare XsF against state-of-the-art methods in constrained optimization problems. In the second, we optimize the hyperparameters of a neural network to achieve maximum compression without degrading its performance. In the third, we learn the state feedback controller of a cart-pole system. Both, Xs and XsF are implemented in Python. The code, which includes scripts to reproduce the results presented herein, is documented and publicly available at <https://github.com/alonrot/excursionsearch>.

5.1. Experimental setup

To assess the performance of all methods we use *simple regret* $r_T = f(x_{\text{bo}}) - \min_{x \in \mathcal{X}} f(x)$, where $x_{\text{bo}} = \arg \min_{t \in [1, T]} y(x_t)$ is the point that yielded the best observation so far. In the constrained case, such point is given by $x_{\text{bo}} = \min_{t \in [1, T]} y(x_t)$ s.t. $y^g(x_t) \leq 0$. We quantify how often safe evaluations are collected using $\Omega = 100N_{\text{safe}}/T$, where N_{safe} is the number of safe evaluations made at the end of each run.

In all cases, the domain is scaled to the unit hypercube. We set $\rho_{\text{safe}} = 0.99$, $\rho_{\text{risk}} = 0.01$, and $\rho_0 = 0.1$. The decision boundary was set at $\rho_b = 0.5$. Both, the objective function and the constraint are modeled with a zero-mean GP, with a squared exponential kernel. The lengthscales and the signal variance are fit to the data after each iteration. Further implementation details, such as hyperprior choices and number of random restarts, are reported in Appendix D.

5.2. Benchmarks for global optimization

We validate Xs and XsF in two challenging benchmarks for global optimization: Hartman 6D, and Michalewicz 10D (Jamil & Yang, 2013). We allow a budget of evaluations $T = 100$ in all cases and repeat all experiments 50 times for each function using a different seed. As in (Wang & Jegelka, 2017; Hernández-Lobato et al., 2016), we use the same initial evaluation (previously selected at random) across all repetitions.

5.2.1. EXCURSION SEARCH (Xs)

We assess the performance of Xs by comparing against popular BO methods: Expected improvement (EI) (Moćkus, 1975), Probability of improvement (PI) (Kushner, 1964), Min-Value Entropy Search (MES) (Wang & Jegelka, 2017), and Gaussian process upper confidence bound (UCB) (Srinivas et al., 2010). Our implementations are based on those used by (Wang & Jegelka, 2017), available online³

Fig. 2a shows the evolution of the simple regret over iterations in the Michalewicz 10D benchmark. Xs reaches the lowest regret, and none of the methods is able to achieve a regret close to zero, which is not surprising given high dimensionality of the problem and the number of allowed evaluations. Table 1 (top) shows statistics on the regret value for both benchmarks. While all methods report a generally high regret in Michalewicz 10D, Xs clearly outperforms all the other methods in Hartman 6D, as it finds a near-zero regret.

5.2.2. FAILURES-AWARE EXCURSION SEARCH (XsF)

To validate XsF, we propose a constrained optimization problem under a limited budget of failures. For this, we simply impose a constraint to the aforementioned benchmarks $g(x) = \prod_{i=1}^D \sin(x_i) - 2^{-D}$. Such function uniformly divides the volume in 2^D sub-hypercubes, and places 2^{D-1} convex disjoint unsafe areas in each one of the sub-hypercubes, so that they are never adjacent to each other. We allow $T = 100$ and a considerably small budget of failures $B = 10$ to all methods. We compare XsF against expected improvement with constraints (EIC) (Gelbart et al., 2014) and predictive entropy search with constraints (PESC)

³<https://github.com/zi-w/Max-value-Entropy-Search>

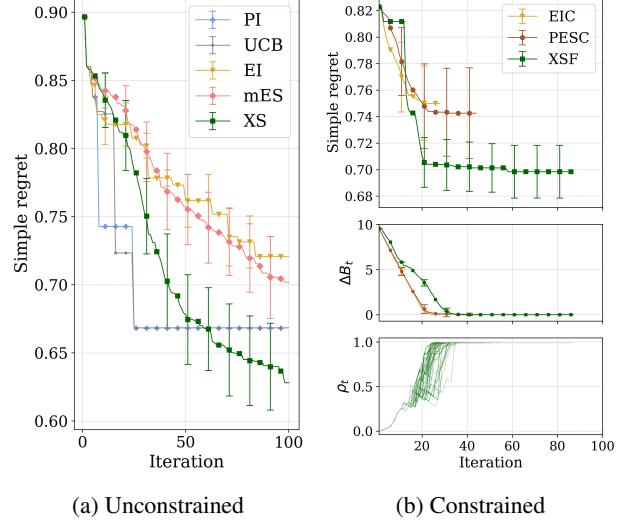


Figure 2: Performance assesment of Xs and XsF on the Michalewicz 10-dimensional benchmark.

(Hernández-Lobato et al., 2016). EIC and PESC are terminated when their budget is depleted. Although individual experiments rarely finish at the same iteration (i.e., some may deplete the budget of failures earlier than others), we use in our results the last regret reported by each algorithm. For EIC, we use our own implementation, while for PESC we use the available open source implementation, included in Spearmint⁴.

In Fig. 2b, we see that XsF reaches a higher number of total evaluations and consistently achieves lower regret than EIC and PESC. Fig. 2b (middle) shows the evolution of the remaining budget of failures ΔB_t over iterations (mean and standard deviation). As can be seen, EIC and PESC deplete

⁴<https://github.com/HIPS/Spearmint/tree/PESC>

Table 1: Constrained (top) and unconstrained benchmarks (bottom). Simple regret r_T (mean $\pm$ std) and percentage of safe evaluations Ω .

	HARTMAN 6D		MICHALEWICZ 10D	
	r_T		r_T	
EI	0.75 $\pm$ 0.00		0.67 $\pm$ 0.00	
MES	0.47 $\pm$ 0.00		0.67 $\pm$ 0.00	
PI	0.34 $\pm$ 0.11		0.72 $\pm$ 0.03	
UCB	0.39 $\pm$ 0.18		0.70 $\pm$ 0.06	
Xs	0.02 $\pm$ 0.01		0.63 $\pm$ 0.06	
	r_T	Ω (%)	r_T	Ω (%)
EIC	0.33 $\pm$ 0.35	68 $\pm$ 30	0.75 $\pm$ 0.06	15 $\pm$ 3
PESC	0.14 $\pm$ 0.22	61 $\pm$ 29	0.74 $\pm$ 0.07	16 $\pm$ 5
XsF	0.09 $\pm$ 0.14	90 $\pm$ 16	0.70 $\pm$ 0.04	28 $\pm$ 7

the budget faster than XsF. Finally, Fig. 2b (bottom) shows the evolution of the ρ_t parameter used to switch between risky and safe strategies in XsF, and also as a threshold for probabilistic constraint satisfaction (cf. Section 4.2). We differentiate two stages: During the initial iterations ρ_t is low, and thus, risky exploration is preferred, which allows XsF to quickly discover better safe areas. At the last iterations, when the budget is depleted, XsF keeps exploring conservatively the discovered safe areas, with $\rho_t = \rho_{\text{safe}}$.

Table 1 (bottom) shows the regret for both, the Michalewicz 10D and the Hartman 6D functions in the constrained case. While the regret comparison is similar to the 10D case, the 6D case shows that Xs clearly outperforms the other methods. The quantity Ω confirms that XsF visits safe evaluations more often than the other methods.

Generally, hyperparameter learning influences the performance of the algorithms. In Appendix E, we show experiments with fixed hyperparameters and a correct GP model, where Xs and XsF outperform the aforementioned methods.

5.3. Compressing a deep neural network

Applying modern deep neural networks (NNs) to large amounts of data typically results in large memory requirements to store the learned weights. Therefore, finding ways of reducing model size without degrading the NN performance has become an important goal in deep learning, for example, to meet storage requirements or to reduce energy consumption. Bayesian compression has been recently proposed as a mean to reduce the NN size: Given an NN architecture, an approximate posterior distribution q on the NN weights is obtained by maximizing the evidence lower bound (ELBO), which balances the expected log-likelihood of samples from q and the theoretical compression size, as given by the KL divergence between q and a prior distribution p (Havasi et al., 2018). A penalization factor β can be used to scale the KL divergence to control the final size of the NN. Finding the value of β that achieves the lowest compression size without significantly degrading NN performance is a challenging and expensive tuning problem. To alleviate the effort of tuning hyperparameters, Bayesian optimization is commonly used. Herein, we propose to minimize the validation error of the NN while keeping its size below a threshold, using constrained Bayesian optimization under a limited budget of failures. While in this example failing to comply with the size requirements is not catastrophic, collecting many failures is undesirable.

We use a LeNet-5 on the MNIST dataset, and a required size below 15 kB. The parameters to tune are β , the learning rate χ , and a scaling factor κ on the the number of neurons of all layers. As a reference for our implementation, we used

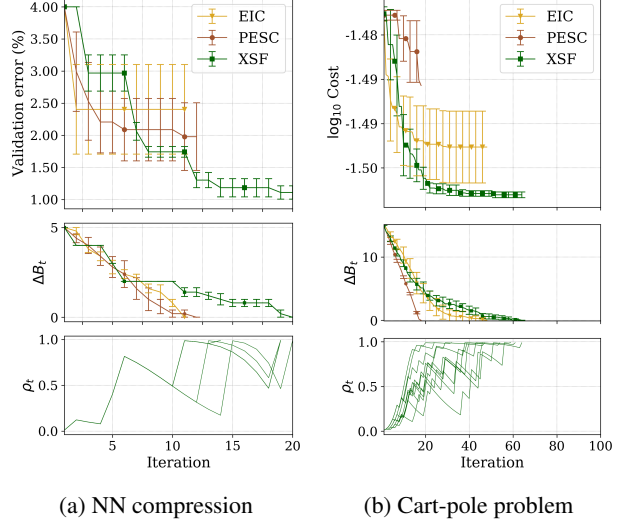


Figure 3: Performance comparison of XsF against EIC and PESC

the open source implementation of MIRACLE⁵ (Havasi et al., 2018). We allow $T = 20$ and $B = 5$ and repeat the experiments 5 times. We fix the training epochs to 20000 for each evaluation (about 25 min in wall-clock time). As shown in Fig. 3a, XsF achieves the lowest regret and standard deviation. The best safe observation is reported by XsF, with validation error 0.76% and theoretical NN size of 12.4 kB (x553 compression). The learned parameters are $\beta = 6.56 \times 10^{-7}$, $\chi = 1.35 \times 10^{-3}$ and $\kappa = 10$.

5.4. Tuning a feedback controller

Bayesian optimization has been used for learning robot controllers to alleviate manual tuning (Calandra et al., 2016; Rai et al., 2018). Herein, we propose to tune a 4D state feedback controller on a cart-pole system, where unstable controllers found during the search are undesirable, as human intervention is required to reset the platform, but not catastrophic. In this setting, allowing a limited budget of failures might increase chances of finding a better optimum. In practice, a constraint can be placed in the cart position to trigger an emergency stop when it grows large (Marco et al., 2016). Controllers that surpass such limit at any moment during the experiment are considered a failure. We use the simulated cart-pole system⁶ from openAI gym (Brockman et al., 2016), implemented in the MuJoCo physics engine (Todorov et al., 2012). The tasks consists on, first stabilizing the pendulum starting from random initial conditions, and second, disturbing the cart position with a small step. We consider a budget $B = 15$ and $T = 100$, and repeat all

⁵<https://github.com/cambridge-mlg/miracle>

⁶<https://gym.openai.com/envs/InvertedPendulum-v2/>

experiments 10 times. Fig. 3b shows that XsF finds a better controller than the other methods.

6. Conclusions

In this paper, we have presented two novel algorithms for BO: Excursion search (Xs), which is based on the study of excursion sets in Gaussian processes, and failures-aware excursion search (XsF), which trades off risky and safe exploration as a function of the remaining budget of failures through a dynamic feedback controller. Our empirical validation shows that both algorithms outperform state-of-the-art methods. Specifically, in situations in which failing is permitted, but undesirable, XsF makes better use of a given budget of failures.

References

- Adler, R. J. and Taylor, J. E. *Random fields and geometry*. Springer Science and Business Media, 2009.
- Berkenkamp, F., Schoellig, A. P., and Krause, A. Safe controller optimization for quadrotors with Gaussian processes. In *IEEE International Conference on Robotics and Automation (ICRA)*, pp. 491–496, 2016.
- Brockman, G., Cheung, V., Pettersson, L., Schneider, J., Schulman, J., Tang, J., and Zaremba, W. OpenAI Gym, 2016.
- Byrd, R. H., Lu, P., Nocedal, J., and Zhu, C. A limited memory algorithm for bound constrained optimization. *SIAM Journal on scientific computing*, 16(5):1190–1208, 1995.
- Calandra, R., Seyfarth, A., Peters, J., and Deisenroth, M. P. Bayesian optimization for learning gaits under uncertainty. *Annals of Mathematics and Artificial Intelligence*, 76(1-2):5–23, 2016.
- De Haan, L. and Ferreira, A. *Extreme value theory: An introduction*. Springer Science & Business Media, 2007.
- Gardner, J. R., Kusner, M. J., Xu, Z., Weinberger, K. Q., and Cunningham, J. P. Bayesian optimization with inequality constraints. In *International Conference on Machine Learning (ICML)*, pp. 937–945, 2014.
- Gelbart, M. A., Snoek, J., and Adams, R. P. Bayesian optimization with unknown constraints. In *Conference on Uncertainty in Artificial Intelligence*, pp. 250–259, 2014.
- Gramacy, R. B. and Lee, H. Optimization under unknown constraints. *Bayesian Statistics 9*, 2011.
- Havasi, M., Peharz, R., and Hernández-Lobato, J. M. Minimal random code learning: Getting bits back from compressed model parameters. *arXiv preprint arXiv:1810.00440*, 2018.
- Hernández-Lobato, J. M., Gelbart, M. A., Adams, R. P., Hoffman, M. W., and Ghahramani, Z. A general framework for constrained bayesian optimization using information-based search. *The Journal of Machine Learning Research*, 17(1):5549–5601, 2016.
- Jamil, M. and Yang, X.-S. A literature survey of benchmark functions for global optimization problems. *International Journal of Mathematical Modelling and Numerical Optimisation*, 4(2):150–194, 2013.
- Jawitz, J. W. Moments of truncated continuous univariate distributions. *Advances in water resources*, 27(3):269–281, 2004.
- Kushner, H. J. A new method of locating the maximum point of an arbitrary multipeak curve in the presence of noise. *Journal of Basic Engineering*, 86(1):97–106, 1964.
- Lam, R. and Willcox, K. Lookahead Bayesian optimization with inequality constraints. In *Advances in Neural Information Processing Systems (NeurIPS)*, pp. 1890–1900, 2017.
- Lindgren, G. Lectures on stationary stochastic processes. *PhD course of Lunds University*, 2006.
- Marco, A., Hennig, P., Bohg, J., Schaal, S., and Trimpe, S. Automatic LQR tuning based on Gaussian process global optimization. In *IEEE international conference on robotics and automation (ICRA)*, pp. 270–277, 2016.
- Moćkus, J. On bayesian methods for seeking the extremum. In *Optimization Techniques IFIP Technical Conference*, pp. 400–404. Springer, 1975.
- Picheny, V. A stepwise uncertainty reduction approach to constrained global optimization. In *International Conference on Artificial Intelligence and Statistics (AISTATS)*, pp. 787–795, 2014.
- Powell, M. J. A direct search optimization method that models the objective and constraint functions by linear interpolation. In *Advances in optimization and numerical analysis*, pp. 51–67. Springer, 1994.
- Rai, A., Antonova, R., Song, S., Martin, W., Geyer, H., and Atkeson, C. Bayesian optimization using domain knowledge on the ATRIAS biped. In *IEEE International Conference on Robotics and Automation (ICRA)*, pp. 1771–1778, 2018.
- Rasmussen, C. E. and Williams, C. K. *Gaussian Processes for Machine Learning*. MIT Press, 2006.
- Schonlau, M., Welch, W. J., and Jones, D. R. Global versus local search in constrained optimization of computer models. *Lecture Notes-Monograph Series*, pp. 11–25, 1998.
- Shahriari, B., Swersky, K., Wang, Z., Adams, R. P., and De Freitas, N. Taking the human out of the loop: A review of Bayesian optimization. *IEEE*, 104(1):148–175, 2016.
- Srinivas, N., Krause, A., Kakade, S. M., and Seeger, M. Gaussian process optimization in the bandit setting: No regret and experimental design. In *International Conference on Machine Learning (ICML)*, 2010.
- Sui, Y., Gotovos, A., Burdick, J., and Krause, A. Safe exploration for optimization with Gaussian processes. In *International Conference on Machine Learning (ICML)*, pp. 997–1005, 2015.

- Todorov, E., Erez, T., and Tassa, Y. MuJoCo: A physics engine for model-based control. In *2012 IEEE/RSJ International Conference on Intelligent Robots and Systems*, pp. 5026–5033, 2012.
- von Rohr, A., Trimpe, S., Marco, A., Fischer, P., and Palagi, S. Gait learning for soft microrobots controlled by light fields. In *IEEE International Conference on Intelligent Robots and Systems (IROS)*, pp. 6199–6206, 2018.
- Wang, Z. and Jegelka, S. Max-value entropy search for efficient Bayesian optimization. In *International Conference on Machine Learning (ICML)*, pp. 3627–3635, 2017.
- Wu, J., Poloczek, M., Wilson, A. G., and Frazier, P. Bayesian optimization with gradients. In *Advances in Neural Information Processing Systems (NeurIPS)*, pp. 5267–5278, 2017.

A. Additional details to Sec. 2.4

Herein, the derivation of (7) is complemented with two additional insights. First, in Appendix A.1, we show how the integral from (7) resolves into an analytical expression. Then, in Appendix A.2, we reason about adding $\{x, u\}$ to the dataset $\mathcal{D}_t^f$ as a *virtual* observation.

A.1. Analytical expression for the integral in (7)

The integral from (7) can be split in two parts

$$\begin{aligned} \int_{-\infty}^{+\infty} |f'| p(f'|\tilde{\mathcal{D}}) df' &= - \int_{-\infty}^0 f' p(f'|\tilde{\mathcal{D}}) df' \\ &+ \int_0^{+\infty} f' p(f'|\tilde{\mathcal{D}}) df', \end{aligned}$$

where the placeholder $\tilde{\mathcal{D}} = \mathcal{D}_t^f \cup \{x, u\}$ is used for simplicity, and the dependency of f' on x is implicit, and also omitted. Since $f' \sim \mathcal{N}(f'; \mu'(x), \nu^2(x))$ is Gaussian distributed, each of the integrals above can be seen as the expected value of an unnormalized truncated normal distribution with support $[-\infty, 0]$, and $[0, +\infty]$, respectively. These expectations are given by (Jawitz, 2004)

$$\begin{aligned} \int_{-\infty}^0 f' p(f'|\tilde{\mathcal{D}}) df' &= \mu'(x) Z_u(x) - \nu(x) \phi\left(-\frac{\mu'(x)}{\nu(x)}\right) \\ \int_0^{+\infty} f' p(f'|\tilde{\mathcal{D}}) df' &= \mu'(x) Z_l(x) + \nu(x) \phi\left(-\frac{\mu'(x)}{\nu(x)}\right), \end{aligned}$$

where $Z_l(x) = \Phi\left(\frac{\mu'(x)}{\nu(x)}\right)$, $Z_u(x) = \Phi\left(-\frac{\mu'(x)}{\nu(x)}\right)$, ϕ is the density of a standard normal distribution and Φ is its cumulative density function. We make use of the definition $\Phi(a) = \frac{1}{2}(1 + \operatorname{erf}(a/\sqrt{2}))$, where $\operatorname{erf}(\cdot)$ is the error function, to compute $\Phi(a) - \Phi(-a) = \operatorname{erf}(a/\sqrt{2})$. Then, $Z_l(x) - Z_u(x) = \operatorname{erf}\left(\frac{\mu'(x)}{\sqrt{2}\nu(x)}\right)$, and the integral can be solved analytically as

$$\begin{aligned} \int_{-\infty}^{+\infty} |f'| p(f'|\tilde{\mathcal{D}}) df' &= \mu'(x)(Z_l(x) - Z_u(x)) + 2\nu(x) \phi\left(\frac{\mu'(x)}{\nu(x)}\right) \\ &= \mu'(x) \operatorname{erf}\left(\frac{\mu'(x)}{\sqrt{2}\nu(x)}\right) + 2\nu(x) \phi\left(\frac{\mu'(x)}{\nu(x)}\right). \end{aligned}$$

Then, (7) follows.

A.2. Virtual observation $\{x, u\}$

The posterior of the process derivative $p(f'|x, u, \mathcal{D}_t^f)$ is a Gaussian density and can be seen as conditioning $f'(x)$ on an extended dataset that includes $\{x, u\}$ as a virtual observation. In the following, we briefly discuss this.

Since differentiation is a linear operation, the derivative of a GP remains a GP (Rasmussen & Williams, 2006, Sec. 9.4).

Furthermore, the joint density between a process value $f(x)$, its derivative $f'(x)$ and the dataset $\{X, y\}$ is Gaussian (Wu et al., 2017)

$$p(y, f, f'|x, X) = \mathcal{N}\left(\begin{bmatrix} y \\ f \\ f' \end{bmatrix}; \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix}, \begin{bmatrix} \tilde{K}(X, X) & K(X, x) & K'(X, x) \\ K(x, X) & K(x, x) & K'(x, x) \\ K'(x, X) & K'(x, x) & K''(x, x) \end{bmatrix}\right),$$

where $\tilde{K}(X, X) = K(X, X) + \sigma_n^2 I$, $K'(X, x) = \partial K(X, x)/\partial x$, $K''(x, x) = \partial^2 K(x, x)/\partial x^2$, and the prior mean of the GP is assumed to be zero. Then, the conditional $p(f'|f, x, \mathcal{D}_t^f) = \mathcal{N}(f'; \mu'(x; f), \nu^2(x))$ is also Gaussian, and can be obtained using Gaussian algebra (Rasmussen & Williams, 2006, A. 2). The mean $\mu'(x; f)$ depends on the random variable f as

$$\begin{aligned} \mu'(x; f) &= [K'(x, X) \quad K'(x, x)] \begin{bmatrix} \tilde{K}(X, X) & K(X, x) \\ K(x, X) & K(x, x) \end{bmatrix}^{-1} \begin{bmatrix} y \\ f \end{bmatrix}. \end{aligned} \quad (16)$$

The sought Gaussian density $\mathcal{N}(f'; \mu'(x; u), \nu^2(x))$ is obtained by replacing the value f in the expression for the mean (16). Thereby, $\{x, u\}$ appears in (16) as an additional *virtual* observation at location x added to the existing dataset $\{X, y\}$, in shorthand notation: $p(f'|u, x, \mathcal{D}_t^f) = p(f'|\mathcal{D}_t^f \cup \{x, u\})$.

B. Fréchet distribution

In this section, we present a brief analysis on why assuming a Fréchet distribution is more error prone in practice than using the Gumbel distribution, when it comes to model the distribution over the global minimum $p(f_*)$. This analysis complements Sec. 3.1 in the paper.

When modeling $p(f_*)$ with the Gumbel distribution and sampling from it, some samples of the global minimum can lie above η , with non-zero probability, which is unrealistic. This can be explicitly avoided by using the Fréchet distribution which, contrary to Gumbel, has zero probability mass near η . We illustrate this with an example, in which a GP with zero mean, unit variance, and squared exponential kernel is considered, conditioned on 20 observations sampled from the GP prior. We discretize the domain in 200 points and sample the resulting GP posterior at them. In Fig. 4, we see that a portion of the Gumbel samples lie above η . To show consistency, we sample the posterior GP 100 times and average the number of times that Gumbel exceeds η , i.e., $1.60 \pm 1.22\%$ of the cases, while the Fréchet distribution exceeds η in 0% of the cases.

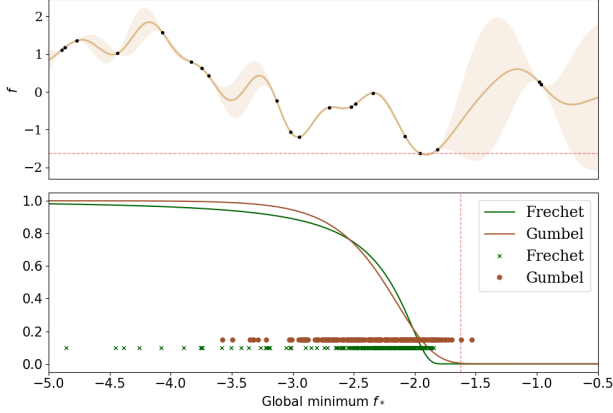


Figure 4: (top) Gaussian process, and η (red dashed line). (bottom) Survival functions for both Gumbel, and Fréchet distributions. Samples from the Gumbel (crosses) and from the Fréchet (circles) distribution are shown.

C. Algorithm and complexity

Herein, we discuss pseudocode for XsF and its computational complexity.

C.1. XsF algorithm

Pseudocode for XsF is shown in Algorithm 1. The decision boundary ρ_b is used to switch between safe search (cf. (13)) and risky search (cf. (15)). The algorithm returns the location where the mean of the posterior GP is minimized without violating the probabilistic constraints. To abbreviate, we have used the placeholder $\varphi(x) = \prod_{i=1}^K \Pr(g_j(x) \leq 0)$.

We do not explicitly discuss Xs, as it simply comprises a standard Bayesian optimization loop, which involves (i) computing samples of the global minimum, and (ii) maximizing the acquisition function (10).

C.2. Complexity

At each iteration, the most expensive operations required to obtain (13) and (15) are: (a) obtaining samples from the global minimum $p(f_*)$ and (b) maximizing the acquisition function using local optimization with random restarts.

As explained in (Wang & Jegelka, 2017), obtaining S samples from $p(f_*)$ involves discretizing the input domain and performing a binary search, which has a total cost of $\mathcal{O}(S + N_d \log(1/\kappa))$, where N_d is the size of the discretization grid, and κ is the accuracy of the binary search.

Each call to the acquisition function α_X (10), has a cost of $\mathcal{O}(SD)$ where D is the dimensionality of the input space. Then, assuming R random restarts, and M maximum number of function calls, the total cost of XsF in per iteration

Algorithm 1 Failures-aware Excursion Search (XsF)

Input: $T, B, \mathcal{D}_0^f, \mathcal{D}_0^g, \rho_{\text{safe}}, \rho_{\text{risk}}, \rho_b, \rho_0$

for $t = 1$ **to** T **do**

$\rho_t \leftarrow \text{UPDATEDECISIONBOUNDARY}(\rho_{t-1})$

$f_* \leftarrow \text{SAMPLEGLOBALMINIMUM}(S)$

if $\rho_t > \rho_b$ **then**

$x_t \leftarrow \arg \max_{x \in \mathcal{X}} \alpha_X(x; f_*)$ s.t. $\varphi(x) \geq \rho_t$ (13)

else

$x_t \leftarrow \arg \max_{x \in \mathcal{X}} \alpha_X(x; f_*) \varphi(x)$ (15)

end if

$\text{EVALUATEANDUPDATEGPS}(x_t)$

end for

$x_*^c \leftarrow \arg \min_{x \in \mathcal{X}} \mu(x)$ s.t. $\varphi(x) \geq \rho_{\text{safe}}$

Return: x_*^c

function $\text{UPDATEDECISIONBOUNDARY}(\rho_t)$

$z_t \leftarrow \Phi^{-1}(\rho_t)$

$u_t \leftarrow u_t(\Delta B_t, \Delta T_t)$ Controller update (14)

$z_t \leftarrow z_t + u_t$ Process update

Return: $\Phi(z_t)$

end function

function $\text{SAMPLEGLOBALMINIMUM}(S)$

Estimate Fréchet distribution $\mathcal{F}_{s,q}$ following (Wang & Jegelka, 2017, Appendix B)

for $l = 1$ **to** S **do**

$f_*^l = \mathcal{F}_{s,q}^{-1}(\xi^l)$. $\xi^l \sim U(0, 1)$

end for

Return: $f_*^1, \dots, f_*^S$

end function

function $\text{EVALUATEANDUPDATEGPS}(x_t)$

$y = f(x_t)$, $y_j = g_j(x_t)$ $j = \{1, \dots, G\}$

$\mathcal{D}_t^f \leftarrow \{y, x_t\}$, $\mathcal{D}_t^{g_j} \leftarrow \{y_j, x_t\}$ $j = \{1, \dots, G\}$

Update hyperparameters of GP models

end function

the worst case scenario is given by $\mathcal{O}(MRD(S+1) + N_d \log(1/\kappa) + (G+1)(N_{\text{obs}}+1)^3)$. The last term is the cost of inverting the Gram matrix, needed for GP predictions (cf. (16)), after having collected N_{obs} observations, and having G constraints. When setting $G = 0$, we obtain the computational cost of Xs, as it also requires gathering samples from $p(f_*)$ and local optimization with random restarts.

D. Implementation details

Both, Xs and XsF are developed using BoTorch⁷, a Python library for Bayesian optimization that serves as a low-level API for building and optimizing new acquisition functions and fitting GP models. It makes use of SCIPY

⁷<https://botorch.org/docs/introduction.html>

Python optimizers⁸ for estimating the GP hyperparameters and optimizing the acquisition function through local optimization with random restarts. In all cases we allow 10 random restarts and use L-BFGS-B (Byrd et al., 1995) as local optimization algorithm. Currently, BoTorch does not support optimization under non-linear constraints, which is needed to solve (13). To overcome this, we use the implementation of COBYLA (Powell, 1994) from NLOPT⁹.

In all experiments, the noise of the likelihood is fixed to $\sigma_n = 0.01$ for all GPs. The chosen hyperpriors on the lengthscales and the signal variance are reported in Table 2, where $\mathcal{U}(a, b)$ refers to a uniform prior on the interval $[a, b]$, $\mathcal{G}(a, b)$ refers to a Gamma prior with concentration a and rate b , and $\mathcal{N}(a, b^2)$ refers to a normal distribution with mean a and standard deviation b .

In Sec. 5.2., both, the Michalewicz and the Hartman functions are normalized to have zero mean and unit variance. The true minimum is known for both functions, which allows to compute the regret.

In Sec. 5.4, the goal is to find the state feedback gain $x \in \mathbb{R}^{4 \times 1}$ for the cart-pole problem that minimizes a quadratic cost $f(x)$, which penalizes deviations of the pendulum states $s_k = [\varphi_k, \dot{\varphi}_k, l_k, \dot{l}_k]^\top$ from an equilibrium point s^* . The pole angle is φ_k , the pole angular velocity is $\dot{\varphi}_k$, the cart displacement is l_k , and the cart velocity is $\dot{l}_k$. The input to the system is the cart acceleration a_k , which is given by $a_k = x^\top (s_k - s^*) + 0.01 \sum_1^{N_{\text{simu}}} (l_k - l^*)$, where an integrator, with gain 0.01, is added to eliminate the steady-state the error. For each parametrization x , the constraint value is computed as the maximum displacement of the cart over a simulation of $N_{\text{simu}} = 800$ steps, i.e., $g(x) = \max(l_k)$, $k = \{1, \dots, N_{\text{simu}}\}$. Constraint violation is quantified as $g(x) > l_{\text{max}}$, where l_{max} is the physical limit of the rail in which the cart moves. To allow the system to dissipate energy, the damping value of the simulated cart-pole in MuJoCo was increased from 1.0 to 1.5.

E. Additional results

In this section, we present complementary results to Sec. 5.

To decouple the influence of the hyperparameter learning from the performance of the acquisition function itself, we fix the GP hyperparameters and sample the true objective f and the true constraint g from the corresponding GP priors. To obtain such samples we follow the same approach as in (Hernández-Lobato et al., 2016): First, the input domain is discretized to an irregular grid of 8000 points. Second,

Table 2: Hyperprior choices for the GP model hyperparameters for all experiments.

		LENGTHSCALE λ	VARIANCE σ^2
Michalewicz 10D	f	$\mathcal{U}(0.01, 0.3)$	$\mathcal{N}(0.5, 0.25^2)$
	g	$\mathcal{U}(0.01, 0.3)$	$\mathcal{N}(0.5, 0.25^2)$
Hartman 6D	f	$\mathcal{G}(1.0, 5.0)$	$\mathcal{N}(0.5, 0.25^2)$
	g	$\mathcal{G}(1.0, 5.0)$	$\mathcal{N}(0.5, 0.25^2)$
NN compression	f	$\mathcal{U}(0.01, 0.3)$	$\mathcal{N}(0.5, 0.2^2)$
	g	$\mathcal{U}(0.01, 0.3)$	$\mathcal{N}(7.5, 2.0^2)$
Pendulum	f	$\mathcal{U}(0.01, 0.3)$	$\mathcal{N}(1.0, 0.25^2)$
	g	$\mathcal{U}(0.01, 0.3)$	$\mathcal{N}(0.5, 0.25^2)$

Table 3: Constrained (top) and unconstrained in-model comparisons (bottom). Simple regret r_T (mean $\pm$ std) and percentage of safe evaluations Ω .

	3D SYNTHETIC FUNCTION	
	r_T	
EI	1.03 ± 0.50	
MES	1.03 ± 0.43	
PI	0.86 ± 0.41	
UCB	1.00 ± 0.43	
Xs	0.19 ± 0.34	
	r_T	Ω (%)
EIC	0.71 ± 0.61	21 ± 19
PESC	1.32 ± 0.62	14 ± 6
XsF	0.30 ± 0.51	52 ± 15

function evaluations are randomly sampled from the corresponding GP prior at such locations. Finally, the GP is conditioned on those evaluations and the resulting posterior mean is used as true objective. The lengthscales were fixed to 0.1 and the signal variance to 1.0.

The simple regret cannot be computed because the true minimum of the GP sample is unknown a priori. Instead, we report results assuming a very conservative lower bound on all the possible sampled functions, i.e., $\min_{x \in \mathcal{X}} f(x) = -4.0$. We allow a maximum of $T = 100$ iterations, and a budget of failures $B = 15$ in the constrained case. The experiments were repeated 50 times for all algorithms. At each repetition, a new function is sampled from the GP priors.

In Table 3, we show a performance comparison of both, Xs and XsF in optimizing a 3D input space. Without the influence of hyperparameter optimization, the proposed methods reach lower observations than state-of-the-art methods.

⁸<https://docs.scipy.org/doc/scipy/reference/tutorial/optimize.html>

⁹<https://nlopt.readthedocs.io/en/latest/>