

Navigating the Unknown: Intelligent and Efficient Solution Discovery via Bayesian Optimization

Wang MA

24 Summer Seminar - Week 1

maw6@rpi.edu

July 22, 2024

Contents

1. Introduction
2. An Overview of Bayesian Optimization
3. A model of the function: Gaussian Process
4. Acquisition Functions
5. Conclusion

Problem Setup

Bayesian Optimization is a class of machine-learning-based optimization methods focused on solving the problem

$$\hat{x} = \max_{x \in A} f(x), \quad (1)$$

where the feasible set and objective function typically follow some assumptions/properties.

- The input x is in $\mathbb{R}^d$, where typically $d \leq 20$.
- The feasible set A is a simple set, e.g., box constraints (hyper-rectangle) or a simplex.
- f is continuous but lacks special structure, e.g., non-convex.
- f is derivative-free, evaluations do not give gradient info.
- f is expensive to evaluate.
- f may be noisy.

Scenarios Suitable for Bayesian Optimization

- ❶ **Expensive function evaluations:** the objective function is costly to evaluate, such as requiring significant computational resources or time
- ❷ **Black-box functions:** the objective function is a black-box with no explicit analytical form and no available gradient information.
- ❸ **Hyperparameter Optimization:** hyperparameter tuning in machine learning models, such as optimizing learning rates, regularization parameters, etc..

Optimization of expensive functions arises in

- fitting machine learning models
- tuning algorithms via backtesting
- optimizing physics-based models
- drug and materials discovery

Contents

1. Introduction
2. An Overview of Bayesian Optimization
3. A model of the function: Gaussian Process
4. Acquisition Functions
5. Conclusion

Bayesian Optimization

Algorithm 1 BayesOpt

Assume a Bayesian prior on f

(usually a Gaussian process prior, a **probabilistic model** of the function)

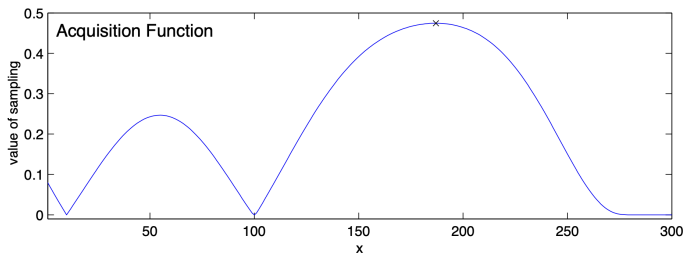
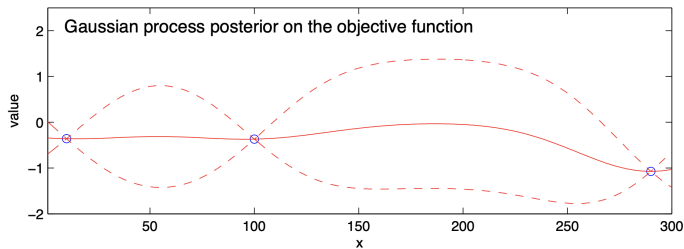
while *budget is not exhausted* **do**

 Find x that maximizes **acquisition function** ($x, \text{posterior}$)

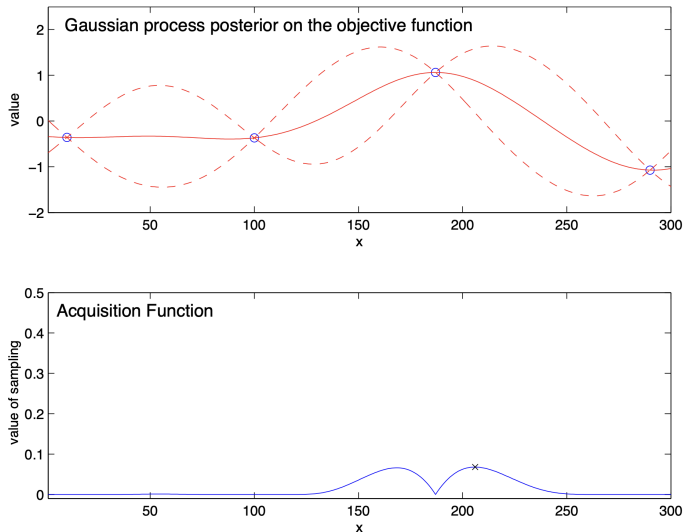
 Sample x and observe $f(x)$

 Update the posterior distribution on f

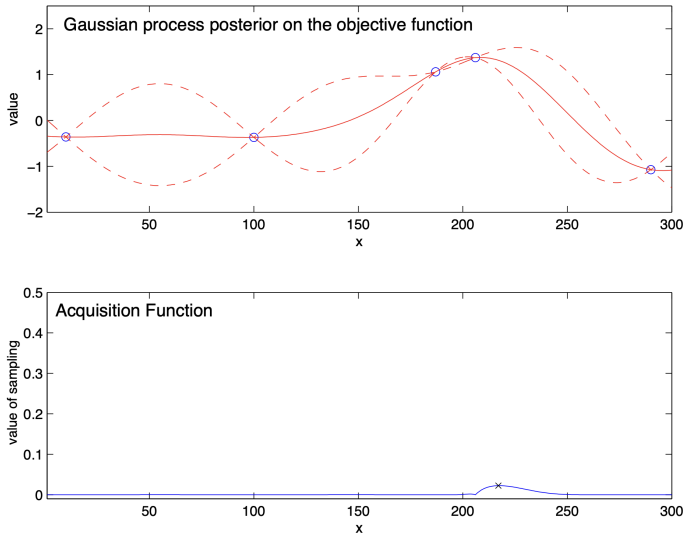
A 1-dim Example to Bayesian Optimization



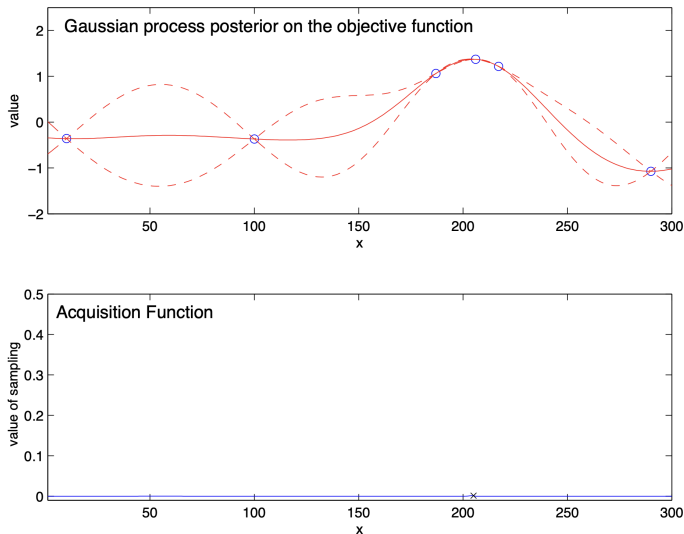
A 1-dim Example to Bayesian Optimization



A 1-dim Example to Bayesian Optimization



A 1-dim Example to Bayesian Optimization



Contents

1. Introduction
2. An Overview of Bayesian Optimization
3. A model of the function: Gaussian Process
4. Acquisition Functions
5. Conclusion

Gaussian Process: Definition

A GP is fully specified by its mean function $m(x)$ and covariance function $k(x, x')$. When we model our function as $f(x) \sim GP(m(x), k(x, x'))$, we are saying that

- Mean function: $m(x) = \mathbb{E}[f(x)]$
- Covariance function: $k(x, x') = \mathbb{E}[(f(x) - m(x))(f(x') - m(x'))]$

The mean function $m(x)$ often assumed to be constant or zero for simplicity, say, $m(x) = 0$. Covariance functions (kernels) decreases with $\|x - x'\|$, commonly used $k(x, x')$:

- Squared Exponential (RBF): $k(x, x') = \sigma_f^2 \exp\left(-\frac{(x-x')^2}{2\ell^2}\right)$
- Matern: $k(x, x') = \sigma_f^2 \frac{2^{1-\nu}}{\Gamma(\nu)} \left(\frac{\sqrt{2\nu}|x-x'|}{\ell}\right)^\nu K_\nu\left(\frac{\sqrt{2\nu}|x-x'|}{\ell}\right)$
- Rational Quadratic: $k(x, x') = \sigma_f^2 \left(1 + \frac{(x-x')^2}{2\alpha\ell^2}\right)^{-\alpha}$

Gaussian Process

Gaussian Process

A Gaussian Process is a collection of random variables, any finite number of which have a joint Gaussian distribution.

As a concrete example, let's choose:

$$m(\mathbf{x}) = 0 \quad (2)$$

$$k(\mathbf{x}, \mathbf{x}') = \exp\left(-\frac{1}{2}(\mathbf{x} - \mathbf{x}')^T(\mathbf{x} - \mathbf{x}')\right). \quad (3)$$

Given observations $\mathbf{f} = [f(\mathbf{x}_1), f(\mathbf{x}_2), \dots, f(\mathbf{x}_t)]$, we would like to make a **prediction** at a new point $\mathbf{x}^*$. According to the GP prior, $f(\mathbf{x}^*)$ is jointly normally distributed with $\mathbf{f}$ so that

$$\Pr\left(\begin{bmatrix} \mathbf{f} \\ f^* \end{bmatrix}\right) = \text{Norm}\left(0, \begin{bmatrix} \mathbf{K}[\mathbf{X}, \mathbf{X}] & \mathbf{K}[\mathbf{X}, \mathbf{x}^*] \\ \mathbf{K}[\mathbf{x}^*, \mathbf{X}] & \mathbf{K}[\mathbf{x}^*, \mathbf{x}^*] \end{bmatrix}\right). \quad (4)$$

Gaussian Process: Predicting

Given the jointly normal property, we have that

$$\Pr(\mathbf{f}^* | \mathbf{f}) = \text{Norm}(\mu[\mathbf{x}^*], \sigma^2[\mathbf{x}^*]), \quad (5)$$

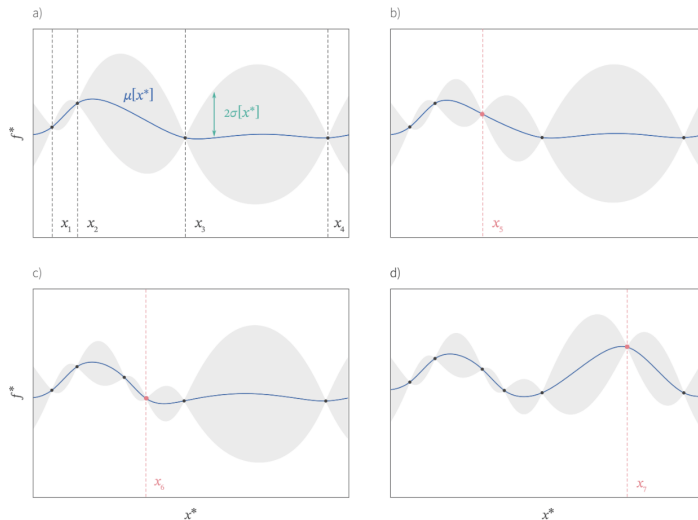
where

$$\mu[\mathbf{x}^*] = \mathbf{K}[\mathbf{x}^*, \mathbf{X}] \mathbf{K}[\mathbf{X}, \mathbf{X}]^{-1} \mathbf{f} \quad (6)$$

$$\sigma^2[\mathbf{x}^*] = \mathbf{K}[\mathbf{x}^*, \mathbf{x}^*] - \mathbf{K}[\mathbf{x}^*, \mathbf{X}] \mathbf{K}[\mathbf{X}, \mathbf{X}]^{-1} \mathbf{K}[\mathbf{X}, \mathbf{x}^*]. \quad (7)$$

Using above formula, we can estimate the distribution at any new point $\mathbf{x}^*$. The best estimate of the function value is given by the mean $\mu[\mathbf{x}]$, and the uncertainty is given by the variance $\sigma^2[\mathbf{x}]$.

Gaussian Process Model



Contents

1. Introduction
2. An Overview of Bayesian Optimization
3. A model of the function: Gaussian Process
- 4. Acquisition Functions**
5. Conclusion

Acquisition Functions

- Acquisition functions guide **the selection of the next point** to evaluate by balancing exploration and exploitation.
- Common acquisition functions include:
 - Probability of Improvement (PI)
 - Expected Improvement (EI)
 - Upper Confidence Bound (UCB)

Upper Confidence Bound (UCB)

- UCB selects points based on an upper confidence bound on the surrogate model's predictions.
- Mathematically defined as:

$$\text{UCB}[\mathbf{x}^*] = \mu[\mathbf{x}^*] + \kappa\sigma[\mathbf{x}^*], \quad (8)$$

where κ is a parameter that controls the trade-off between exploration and exploitation.

- This favors either (i) regions where $\mu[\mathbf{x}^*]$ is large (for exploitation) or (ii) regions where $\sigma[\mathbf{x}^*]$ is large (for exploration). The positive parameter κ trades off these two tendencies.

Probability of Improvement (PI)

- PI aims to maximize the probability that the next sample will improve over the current maximum.
- Mathematically defined as:

$$\text{PI}[\mathbf{x}^*] = \int_{f[\hat{\mathbf{x}}]}^{\infty} \text{Norm}_{f[\mathbf{x}^*]}[\mu[\mathbf{x}^*], \sigma[\mathbf{x}^*]] d\mathbf{f}[\mathbf{x}^*], \quad (9)$$

where $f[\hat{\mathbf{x}}]$ is the current maximum.

This acquisition function computes the likelihood that the function at $\mathbf{x}^*$ will return a result higher than current maximum. For each point $\mathbf{x}^*$, we integrate the part of the associated normal distribution that is above the current maximum.

Expected Improvement (EI)

- EI balances exploration and exploitation by considering both the magnitude and the probability of improvement.
- Mathematically defined as:

$$\text{EI}[\mathbf{x}^*] = \mathbb{E}_n[(f[\mathbf{x}^*] - f[\hat{\mathbf{x}}])^+] \quad (10)$$

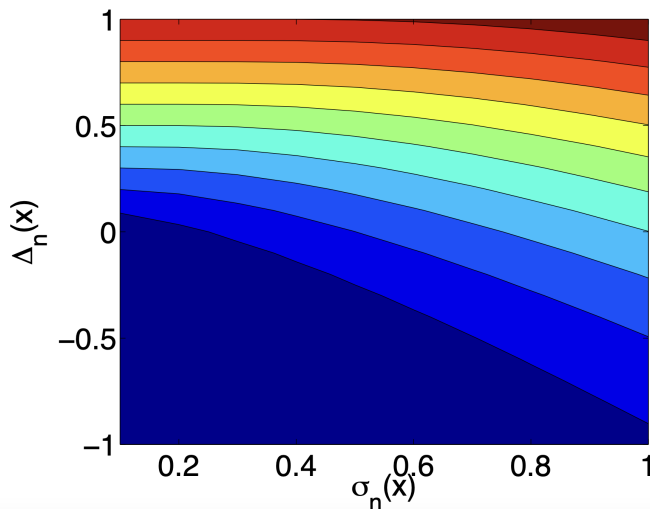
$$= \int_{f[\hat{\mathbf{x}}]}^{\infty} (f[\mathbf{x}^*] - f[\hat{\mathbf{x}}]) \text{Norm}_{f[\mathbf{x}^*]}[\mu[\mathbf{x}^*], \sigma[\mathbf{x}^*]] d\mathbf{f}[\mathbf{x}^*]. \quad (11)$$

Expected Improvement (EI) takes into account how much the improvement will be, so that we can find the favorable larger improvement. A closed form of EI:

$$\text{EI}[\mathbf{x}^*] = [\Delta(\mathbf{x}^*)]^+ + \sigma(\mathbf{x}^*) \varphi\left(\frac{\Delta(\mathbf{x}^*)}{\sigma(\mathbf{x}^*)}\right) - |\Delta(\mathbf{x}^*)| \Phi\left(-\frac{|\Delta(\mathbf{x}^*)|}{\sigma(\mathbf{x}^*)}\right), \quad (12)$$

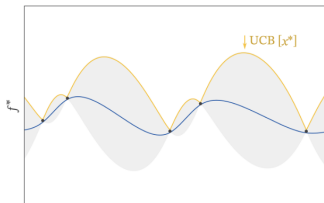
where $\Delta(\mathbf{x}^*) = \mu(\mathbf{x}^*) - f[\hat{\mathbf{x}}]$.

Expected Improvement: Exploration ($\sigma[x^*]$) vs. Exploitation ($\Delta[x^*]$)

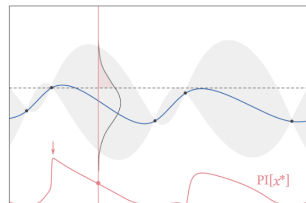


Acquisition Functions: Where should we sample next?

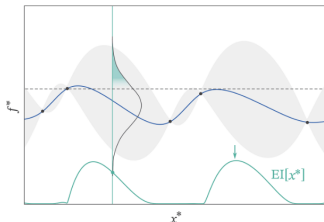
a) Upper confidence bound



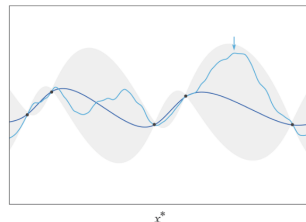
b) Probability of improvement



c) Expected improvement



d) Thompson sampling



Parallel Expected Improvement: Setting

We can parallelize EI:

$$\text{EI}[x_{1:q}] = \mathbb{E}_n[(\max(f[x_1], \dots, f[x_q]) - f[\hat{X}])^+] \quad (13)$$

How to maximize the parallel expected improvement?

- 1 Construct an unbiased estimator of the gradient of $\text{EI}[x_{1:q}]$.
- 2 Use **multistart stochastic gradient ascent** to approximately maximize $\text{EI}[x_{1:q}]$.

PEI: Estimate ∇EI

Here's how we estimate ∇EI :

- $Y = [f[x_1], \dots, f[x_q]]$ is multivariate normal (GP prior)
- Let $m = \mathbb{E}[Y]$ and $C = \text{Chol}(\text{Cov}[Y])$
- then $Y = m + CZ$, where Z is a vector of independent standard normals
- $\text{EI}[x_{1:q}] = \mathbb{E}[h(Y)]$, where $h(Y) = (\max[Y] - f[\hat{x}])^+$
- Assume the problem is well-behaved, then we can switch derivative and expectation:

$$\nabla \text{EI}[x_{1:q}] = \mathbb{E}[\nabla h(m + CZ)]$$

PEI: Estimate ∇EI

Here's how we estimate ∇EI :

- Simulate a vector Z of independent standard normals
- Calculate $m = \mathbb{E}[Y]$ and $C = \text{Chol}(\text{Cov}[Y])$
- Then the estimator of $\nabla EI[x_1, \dots, x_q]$ is $\nabla h(m + CZ)$, where $h(Y) = (\max[Y] - f[\hat{x}])^+$

Approximately Maximize $\text{EI}[x_{1:q}]$

We use the previous estimator of ∇EI in multistart stochastic gradient ascent:

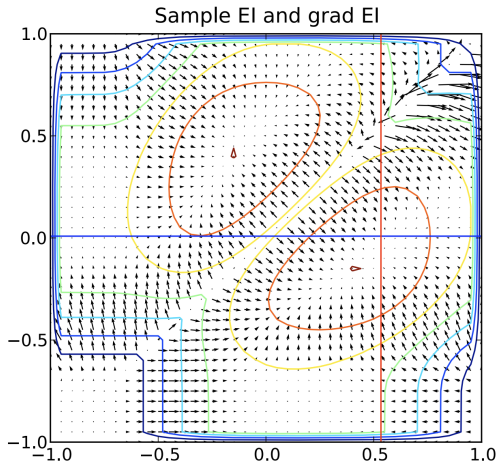
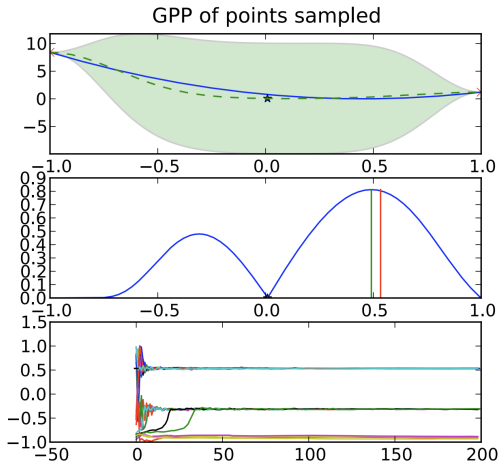
- 1 Select several starting points, uniformly at random
- 2 From each starting point, iterate using the stochastic gradient method until convergence.

$$(\vec{x}_1, \dots, \vec{x}_q) \leftarrow (\vec{x}_1, \dots, \vec{x}_q) + \alpha_n \mathbf{g}(\vec{x}_1, \dots, \vec{x}_q, \omega),$$

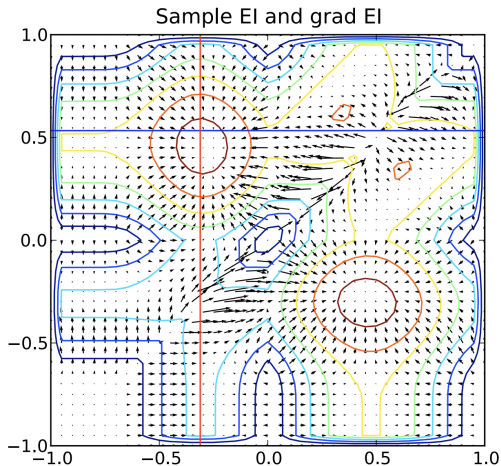
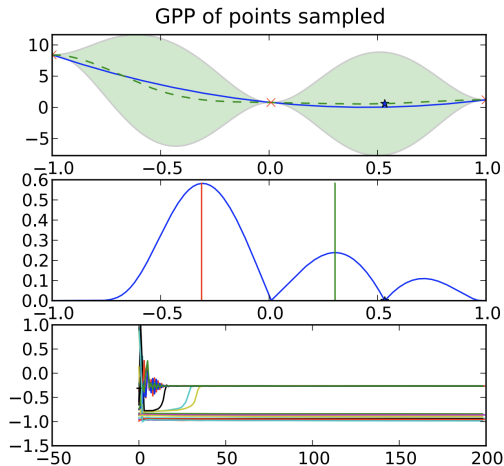
where (α_n) is a stepsize sequence,

- 3 For each starting point, average the iterates to get an estimated stationary point. (Polyak-Ruppert averaging)
- 4 Select the estimated stationary point with the best estimated value as the solution.

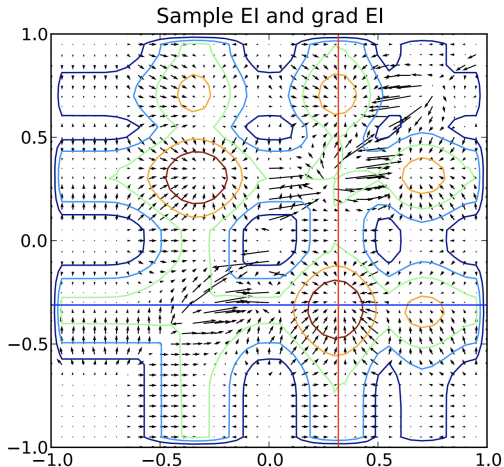
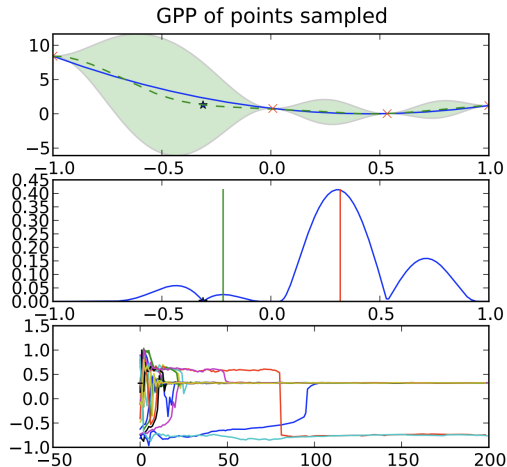
Parallel Expected Improvement: Illustration



Parallel Expected Improvement: Illustration

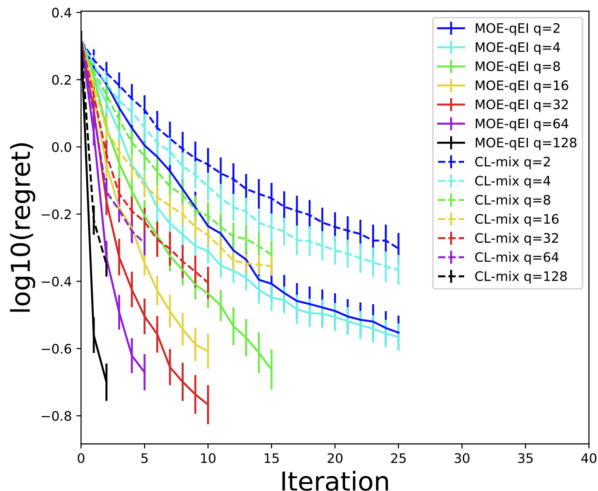


Parallel Expected Improvement: Illustration



Parallel Expected Improvement: Illustration

Here's a demonstration on a 6-dimensional Bayesian Optimization problem with up to 128 parallel evaluations.



Knowledge Gradient

- KG measures the expected increase in the maximum value of the objective function after a new observation.
- We will not select the previously evaluated point as the final solution: the exploration can be "unsuccessful"
- Formally, for a candidate point x , KG is defined as:

$$KG(x) = \mathbb{E}[\mu_{n+1}^* - \mu_n^* | \text{observing } x_{n+1}]$$

- Here, $\mu_n^* = \max_x \mu_n[x] = \max_x \mathbb{E}_n[f(x)]$ is the best expected value of our objective under the posterior at time step n .

Knowledge Gradient: Formula

- Let $\mu(\mathbf{x})$ and $\sigma(\mathbf{x})$ be the posterior mean and standard deviation of the objective function at $\mathbf{x}$.
- The KG acquisition function can be expressed as:

$$KG(\mathbf{x}) = (\mu(\mathbf{x}) - \mu_n^*)\Phi\left(\frac{\Delta(\mathbf{x})}{\sigma(\mathbf{x})}\right) + \sigma(\mathbf{x})\phi\left(\frac{\Delta(\mathbf{x})}{\sigma(\mathbf{x})}\right)$$

Where:

- μ_n^* is the current best observed value.
- Φ is the cumulative distribution function of the standard normal distribution.
- ϕ is the probability density function of the standard normal distribution.

Our approach for optimizing **parallel EI** also works for **optimizing KG**

1. Estimate $\nabla \text{KG}(\mathbf{x}_{1:q})$ **using infinitesimal perturbation analysis (IPA) & the envelope theorem:**

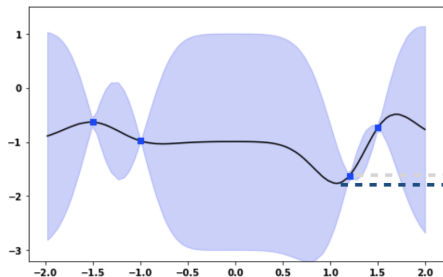
$$\nabla \text{KG} = \nabla \mu_{n+1}(\mathbf{x}^*; m(\mathbf{x}_{1:q}) + \mathbf{CZ}, \mathbf{x}_{1:q}),$$

calculating the gradient holding $\mathbf{x}^*$ fixed.

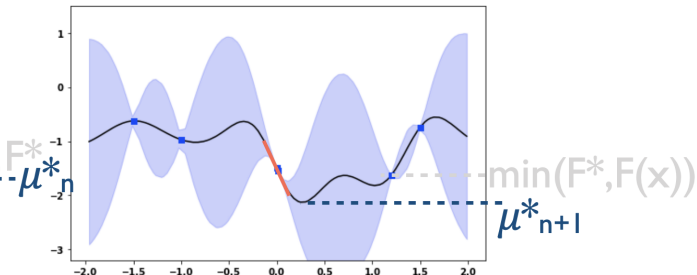
2. Use multistart stochastic gradient ascent to maximize $\text{KG}(\mathbf{x}_{1:q})$

EI May Make Poor Decisions

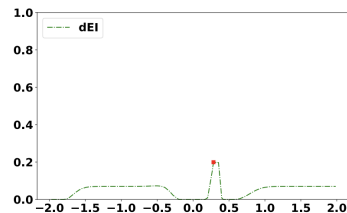
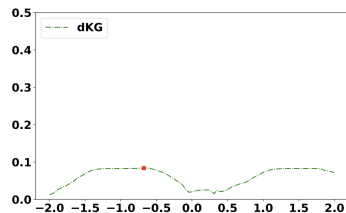
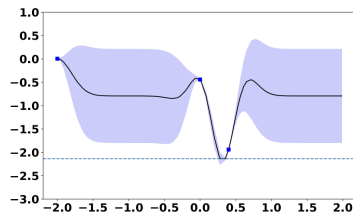
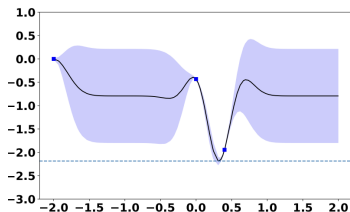
Posterior @ time n



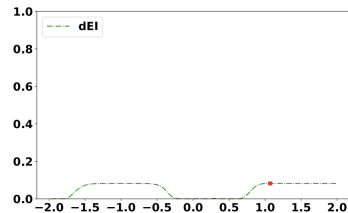
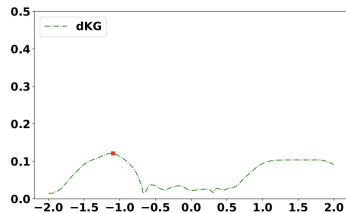
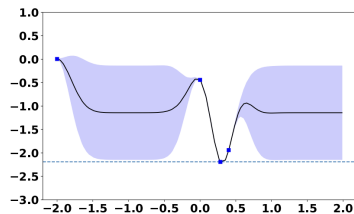
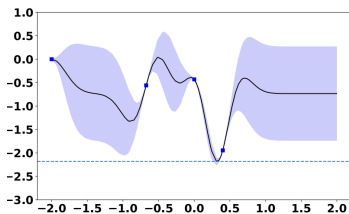
Posterior @ time n+1



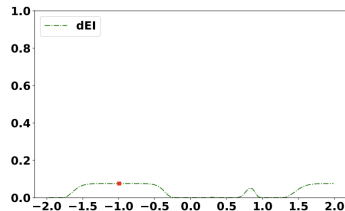
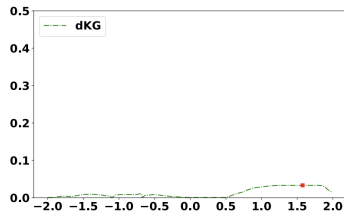
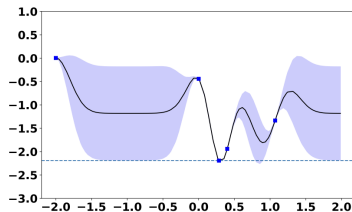
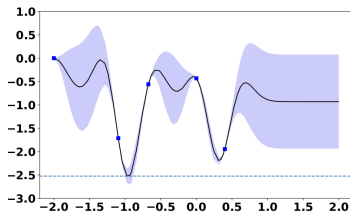
Knowledge Gradient is Efficient



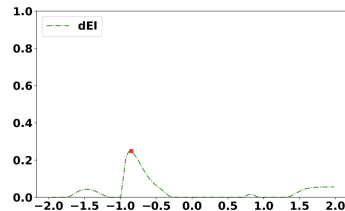
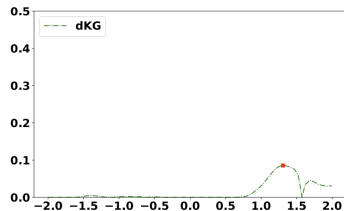
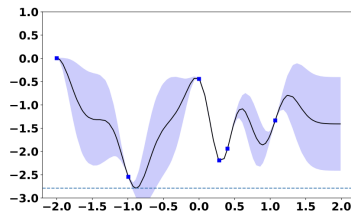
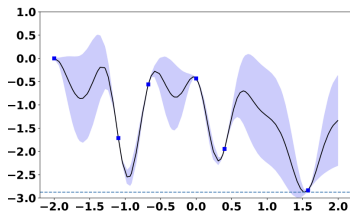
Knowledge Gradient is Efficient



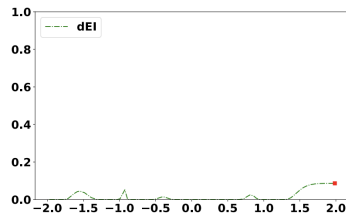
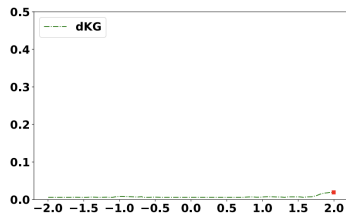
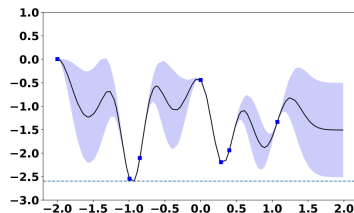
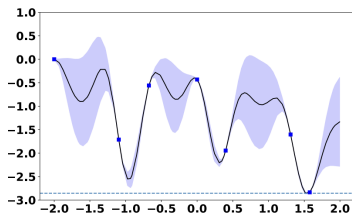
Knowledge Gradient is Efficient



Knowledge Gradient is Efficient



Knowledge Gradient is Efficient



Incorporating Noisy Measurements

We add an extra noise term to the expression for the Gaussian process covariance

$$\mathbb{E}[(y(x) - m(x))(y(x') - m(x')))] = k(x, x') + \sigma_n^2, \quad (14)$$

where $y(x) = f(x) + \epsilon$ is a noisy observation. And the joint normal becomes

$$\Pr \left(\begin{bmatrix} \mathbf{f} \\ \mathbf{f}^* \end{bmatrix} \right) = \text{Norm} \left(\mathbf{0}, \begin{bmatrix} \mathbf{K}[\mathbf{X}, \mathbf{X}] + \sigma_n^2 \mathbf{I} & \mathbf{K}[\mathbf{X}, \mathbf{x}^*] \\ \mathbf{K}[\mathbf{x}^*, \mathbf{X}] & \mathbf{K}[\mathbf{x}^*, \mathbf{x}^*] \end{bmatrix} \right). \quad (15)$$

Then We have the conditional distribution with noise:

$$\Pr(\mathbf{f}^* | \mathbf{f}) = \text{Norm}(\mu[\mathbf{x}^*], \sigma^2[\mathbf{x}^*]), \quad (16)$$

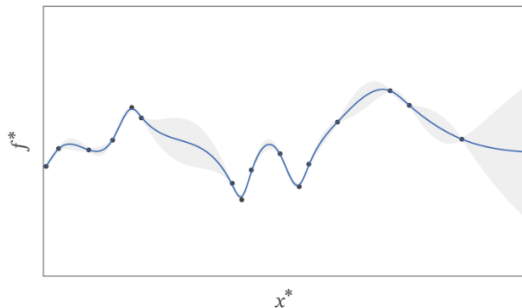
where

$$\mu[\mathbf{x}^*] = \mathbf{K}[\mathbf{x}^*, \mathbf{X}] [\mathbf{K}[\mathbf{X}, \mathbf{X}] + \sigma_n^2 \mathbf{I}]^{-1} \mathbf{f} \quad (17)$$

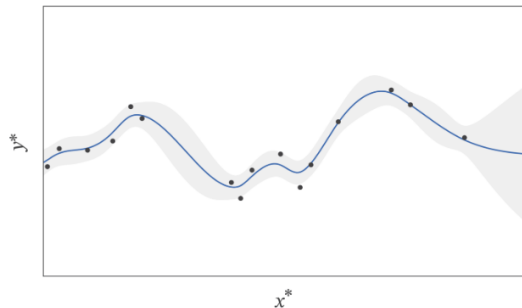
$$\sigma^2[\mathbf{x}^*] = \mathbf{K}[\mathbf{x}^*, \mathbf{x}^*] - \mathbf{K}[\mathbf{x}^*, \mathbf{X}] [\mathbf{K}[\mathbf{X}, \mathbf{X}] + \sigma_n^2 \mathbf{I}]^{-1} \mathbf{K}[\mathbf{x}^*, \mathbf{x}^*]. \quad (18)$$

Noisy Measurements Illustration

a) No measurement noise



b) Measurement noise



Further Reading

- Learning GP parameters: MLE, full Bayesian Approach
- Different Probabilistic Models: Random Forest (SMAC)
- Discrete Variables: Beta-Bernoulli Bandit
- Kernel Choices
- Tips, tricks, and limitations
 - Inducing points
 - Decomposing the kernel
 - Using random projections

Contents

1. Introduction
2. An Overview of Bayesian Optimization
3. A model of the function: Gaussian Process
4. Acquisition Functions
5. Conclusion

Conclusion: Bayesian Optimization

Algorithm 2 BayesOpt

Assume a Bayesian prior on f

(usually a Gaussian process prior, a **probabilistic model** of the function)

while *budget is not exhausted* **do**

 Find x that maximizes **acquisition function** ($x, \text{posterior}$)

 Sample x and observe $f(x)$

 Update the posterior distribution on f

Conclusion: Gaussian Process Regression

$$\Pr(\mathbf{f}^* | \mathbf{f}) = \text{Norm}(\mu[\mathbf{x}^*], \sigma^2[\mathbf{x}^*]), \quad (19)$$

where

$$\mu[\mathbf{x}^*] = \mathbf{K}[\mathbf{x}^*, \mathbf{X}] \mathbf{K}[\mathbf{X}, \mathbf{X}]^{-1} \mathbf{f} \quad (20)$$

$$\sigma^2[\mathbf{x}^*] = \mathbf{K}[\mathbf{x}^*, \mathbf{x}^*] - \mathbf{K}[\mathbf{x}^*, \mathbf{X}] \mathbf{K}[\mathbf{X}, \mathbf{X}]^{-1} \mathbf{K}[\mathbf{x}^*, \mathbf{x}^*]. \quad (21)$$

Conclusion: Acquisition Functions

- Acquisition functions guide **the selection of the next point** to evaluate by balancing exploration and exploitation.

- Common acquisition functions include:

- Probability of Improvement (PI)

$$\text{PI}[x^*] = \int_{f[\hat{x}]}^{\infty} \text{Norm}_{f[x^*]}[\mu[x^*], \sigma[x^*]] df[x^*], \quad (22)$$

- Expected Improvement (EI)

$$\text{EI}[x^*] = \mathbb{E}_n[(f[x^*] - f[\hat{x}])^+] \quad (23)$$

$$= \int_{f[\hat{x}]}^{\infty} (f[x^*] - f[\hat{x}]) \text{Norm}_{f[x^*]}[\mu[x^*], \sigma[x^*]] df[x^*]. \quad (24)$$

- Upper Confidence Bound (UCB)

$$\text{UCB}[x^*] = \mu[x^*] + \kappa\sigma[x^*], \quad (25)$$

Review

1. Introduction
2. An Overview of Bayesian Optimization
3. A model of the function: Gaussian Process
4. Acquisition Functions
5. Conclusion

Bayesian Optimization: END

Thank you!

Questions and Opinions are Welcome!