

Approximate EP for Deep Gaussian Processes

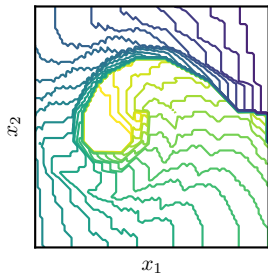
José Miguel Hernández-Lobato

Department of Engineering
University of Cambridge

Joint work with
Thang D. Bui, Daniel Hernández-Lobato,
Yingzhen Li and Rich E. Turner.

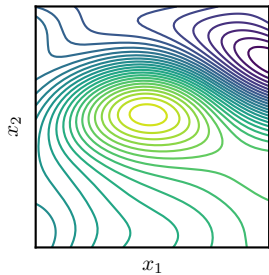
Motivation

Target function

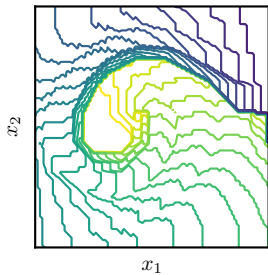


Motivation

GP fit

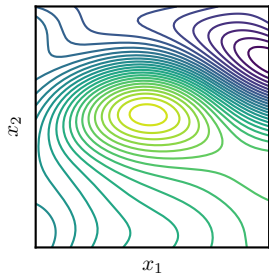


Target function

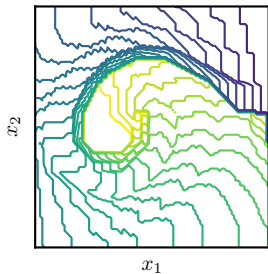


Motivation

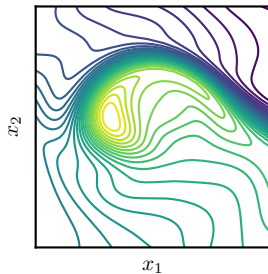
GP fit



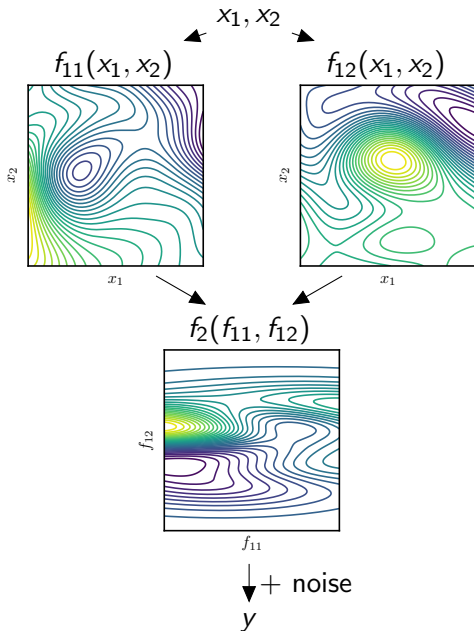
Target function



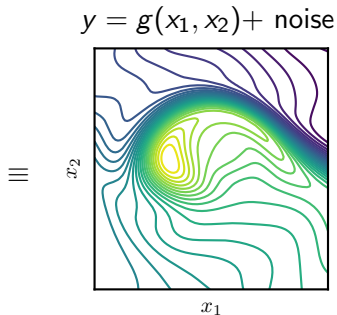
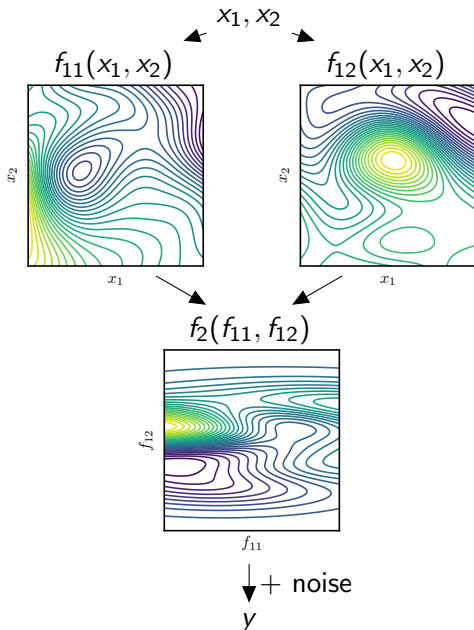
DGP fit



How do deep GPs work?



How do deep GPs work?



Why deep GPs?

Advantages:

- useful input warping: automatic, nonparametric kernel design
- repair damage done by sparse approximations to GPs
- inference exact for layer's output when layer's input is deterministic

Why deep GPs?

Advantages:

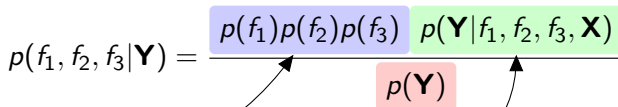
- useful input warping: automatic, nonparametric kernel design
- repair damage done by sparse approximations to GPs
- inference exact for layer's output when layer's input is deterministic

Drawbacks:

- require complicated approximate inference methods
- high computational cost

Bayesian inference

Posterior over latent functions:

$$p(f_1, f_2, f_3 | \mathbf{Y}) = \frac{p(f_1)p(f_2)p(f_3) \quad p(\mathbf{Y} | f_1, f_2, f_3, \mathbf{X})}{p(\mathbf{Y})}$$


- GP priors
- Likelihood function
- Marginal likelihood

But the posterior $p(f_1, f_2, f_3 | \mathbf{Y})$ is **intractable**.

Inducing inputs and outputs

Distribution on f given by GP with inducing inputs $\bar{\mathbf{X}}$ and outputs $\bar{\mathbf{f}}$.

$f(\mathbf{x}) \sim \mathcal{N}(m_{\mathbf{x}}, v_{\mathbf{x}})$, where

$$m_{\mathbf{x}} = \mathbf{k}_{\mathbf{x}, \bar{\mathbf{X}}} \mathbf{K}_{\bar{\mathbf{X}}, \bar{\mathbf{X}}}^{-1} \bar{\mathbf{f}},$$

$$v_{\mathbf{x}} = \mathbf{k}_{\mathbf{x}, \mathbf{x}} - \mathbf{k}_{\mathbf{x}, \bar{\mathbf{X}}} \mathbf{K}_{\bar{\mathbf{X}}, \bar{\mathbf{X}}}^{-1} \mathbf{k}_{\bar{\mathbf{X}}, \mathbf{x}}.$$

Inducing inputs and outputs

Distribution on f given by GP with inducing inputs $\bar{\mathbf{X}}$ and outputs $\bar{\mathbf{f}}$.

$f(\mathbf{x}) \sim \mathcal{N}(m_{\mathbf{x}}, v_{\mathbf{x}})$, where

$$m_{\mathbf{x}} = \mathbf{k}_{\mathbf{x}, \bar{\mathbf{X}}} \mathbf{K}_{\bar{\mathbf{X}}, \bar{\mathbf{X}}}^{-1} \bar{\mathbf{f}},$$

$$v_{\mathbf{x}} = \mathbf{k}_{\mathbf{x}, \mathbf{x}} - \mathbf{k}_{\mathbf{x}, \bar{\mathbf{X}}} \mathbf{K}_{\bar{\mathbf{X}}, \bar{\mathbf{X}}}^{-1} \mathbf{k}_{\bar{\mathbf{X}}, \mathbf{x}}.$$

$\bar{\mathbf{f}} \sim \mathcal{N}(\mathbf{m}_f, \mathbf{V}_f)$:

$$m_{\mathbf{x}} = \mathbf{k}_{\mathbf{x}, \bar{\mathbf{X}}} \mathbf{K}_{\bar{\mathbf{X}}, \bar{\mathbf{X}}}^{-1} \mathbf{m}_f,$$

$$v_{\mathbf{x}} = \mathbf{k}_{\mathbf{x}, \mathbf{x}} - \mathbf{k}_{\mathbf{x}, \bar{\mathbf{X}}} \mathbf{K}_{\bar{\mathbf{X}}, \bar{\mathbf{X}}}^{-1} \mathbf{k}_{\bar{\mathbf{X}}, \mathbf{x}} + \mathbf{k}_{\mathbf{x}, \bar{\mathbf{X}}} \mathbf{K}_{\bar{\mathbf{X}}, \bar{\mathbf{X}}}^{-1} \mathbf{V}_f \mathbf{K}_{\bar{\mathbf{X}}, \bar{\mathbf{X}}}^{-1} \mathbf{k}_{\bar{\mathbf{X}}, \mathbf{x}}.$$

Inducing inputs and outputs

Distribution on f given by GP with inducing inputs $\bar{\mathbf{X}}$ and outputs $\bar{\mathbf{f}}$.

$f(\mathbf{x}) \sim \mathcal{N}(m_{\mathbf{x}}, v_{\mathbf{x}})$, where

$$m_{\mathbf{x}} = \mathbf{k}_{\mathbf{x}, \bar{\mathbf{X}}} \mathbf{K}_{\bar{\mathbf{X}}, \bar{\mathbf{X}}}^{-1} \bar{\mathbf{f}},$$

$$v_{\mathbf{x}} = \mathbf{k}_{\mathbf{x}, \mathbf{x}} - \mathbf{k}_{\mathbf{x}, \bar{\mathbf{X}}} \mathbf{K}_{\bar{\mathbf{X}}, \bar{\mathbf{X}}}^{-1} \mathbf{k}_{\bar{\mathbf{X}}, \mathbf{x}}.$$

$\bar{\mathbf{f}} \sim \mathcal{N}(\mathbf{m}_f, \mathbf{V}_f)$:

$$m_{\mathbf{x}} = \mathbf{k}_{\mathbf{x}, \bar{\mathbf{X}}} \mathbf{K}_{\bar{\mathbf{X}}, \bar{\mathbf{X}}}^{-1} \mathbf{m}_f,$$

$$v_{\mathbf{x}} = \mathbf{k}_{\mathbf{x}, \mathbf{x}} - \mathbf{k}_{\mathbf{x}, \bar{\mathbf{X}}} \mathbf{K}_{\bar{\mathbf{X}}, \bar{\mathbf{X}}}^{-1} \mathbf{k}_{\bar{\mathbf{X}}, \mathbf{x}} + \mathbf{k}_{\mathbf{x}, \bar{\mathbf{X}}} \mathbf{K}_{\bar{\mathbf{X}}, \bar{\mathbf{X}}}^{-1} \mathbf{V}_f \mathbf{K}_{\bar{\mathbf{X}}, \bar{\mathbf{X}}}^{-1} \mathbf{k}_{\bar{\mathbf{X}}, \mathbf{x}}.$$

$\mathbf{x} \sim \mathcal{N}(\mathbf{m}_{\mathbf{x}}, \mathbf{V}_{\mathbf{x}})$, $f(\mathbf{x})$ no longer Gaussian, but

$$m_{\mathbf{x}} = \mathbb{E}_{\mathbf{x}}[\mathbf{k}_{\mathbf{x}, \bar{\mathbf{X}}}] \mathbf{K}_{\bar{\mathbf{X}}, \bar{\mathbf{X}}}^{-1} \mathbf{m}_f,$$

$v_{\mathbf{x}} =$ complicated but doable, requires $\mathbb{E}_{\mathbf{x}}[\mathbf{k}_{\bar{\mathbf{X}}, \mathbf{x}} \mathbf{k}_{\mathbf{x}, \bar{\mathbf{X}}}]$.

Expectation propagation

Approximates $p(\boldsymbol{\theta}) \propto p_0(\boldsymbol{\theta}) \prod_{n=1}^N f_n(\boldsymbol{\theta})$ with $q(\boldsymbol{\theta}) \propto p_0(\boldsymbol{\theta}) \prod_{n=1}^N \tilde{f}_n(\boldsymbol{\theta})$

Expectation propagation

Approximates $p(\boldsymbol{\theta}) \propto p_0(\boldsymbol{\theta}) \prod_{n=1}^N f_n(\boldsymbol{\theta})$ with $q(\boldsymbol{\theta}) \propto p_0(\boldsymbol{\theta}) \prod_{n=1}^N \tilde{f}_n(\boldsymbol{\theta})$

$$p(\boldsymbol{\theta}) \propto p_0(\boldsymbol{\theta}) f_1(\boldsymbol{\theta}) f_2(\boldsymbol{\theta}) f_3(\boldsymbol{\theta}) \approx q(\boldsymbol{\theta}) \propto p_0(\boldsymbol{\theta}) \tilde{f}_1(\boldsymbol{\theta}) \tilde{f}_2(\boldsymbol{\theta}) \tilde{f}_3(\boldsymbol{\theta})$$


Expectation propagation

Approximates $p(\theta) \propto p_0(\theta) \prod_{n=1}^N f_n(\theta)$ with $q(\theta) \propto p_0(\theta) \prod_{n=1}^N \tilde{f}_n(\theta)$

$$p(\theta) \propto p_0(\theta) f_1(\theta) f_2(\theta) f_3(\theta) \approx q(\theta) \propto p_0(\theta) \tilde{f}_1(\theta) \tilde{f}_2(\theta) \tilde{f}_3(\theta)$$


The $\tilde{f}_n$ are tuned by minimizing the KL divergence

$$D_{\text{KL}}[p_n || q] \quad \text{for } n = 1, \dots, N, \quad \text{where} \quad \begin{aligned} p_n(\theta) &\propto f_n(\theta) \prod_{j \neq n} \tilde{f}_j(\theta) \\ q(\theta) &\propto \tilde{f}_n(\theta) \prod_{j \neq n} \tilde{f}_j(\theta) \end{aligned}$$

The EP approximation to the **evidence** $p(\mathbf{Y})$ is given by

$$\log Z_{\text{EP}} = \log Z_q + \sum_{n=1}^N \log \mathbf{E}_q \left[\left(\frac{f_n(\boldsymbol{\theta})}{\tilde{f}_n(\boldsymbol{\theta})} \right) \right],$$

The EP approximation to the **evidence** $p(\mathbf{Y})$ is given by

$$\log Z_{\text{EP}} = \log Z_q + \sum_{n=1}^N \log \mathbf{E}_q \left[\left(\frac{f_n(\boldsymbol{\theta})}{\tilde{f}_n(\boldsymbol{\theta})} \right) \right],$$

The EP solution for q can be obtained by solving

$$\max_q \min_{\tilde{f}_1, \dots, \tilde{f}_N} \log Z_{\text{EP}} \quad \text{subject to} \quad q(\boldsymbol{\theta}) = p_0(\boldsymbol{\theta}) \prod_{n=1}^N \tilde{f}_n(\boldsymbol{\theta}).$$

The EP approximation to the **evidence** $p(\mathbf{Y})$ is given by

$$\log Z_{\text{EP}} = \log Z_q + \sum_{n=1}^N \log \mathbf{E}_q \left[\left(\frac{f_n(\boldsymbol{\theta})}{\tilde{f}_n(\boldsymbol{\theta})} \right) \right],$$

The EP solution for q can be obtained by solving

$$\max_q \min_{\tilde{f}_1, \dots, \tilde{f}_N} \log Z_{\text{EP}} \quad \text{subject to} \quad q(\boldsymbol{\theta}) = p_0(\boldsymbol{\theta}) \prod_{n=1}^N \tilde{f}_n(\boldsymbol{\theta}).$$

Can be solved with **double-loop** algorithm.

The EP approximation to the **evidence** $p(\mathbf{Y})$ is given by

$$\log Z_{\text{EP}} = \log Z_q + \sum_{n=1}^N \log \mathbf{E}_q \left[\left(\frac{f_n(\boldsymbol{\theta})}{\tilde{f}_n(\boldsymbol{\theta})} \right) \right],$$

The EP solution for q can be obtained by solving

$$\max_q \min_{\tilde{f}_1, \dots, \tilde{f}_N} \log Z_{\text{EP}} \quad \text{subject to} \quad q(\boldsymbol{\theta}) = p_0(\boldsymbol{\theta}) \prod_{n=1}^N \tilde{f}_n(\boldsymbol{\theta}).$$

Can be solved with **double-loop** algorithm. **Too slow in practice!**

Approximate Expectation Propagation

$$p(\boldsymbol{\theta}) \propto p_0(\boldsymbol{\theta}) f_1(\boldsymbol{\theta}) f_2(\boldsymbol{\theta}) f_3(\boldsymbol{\theta}) \approx q(\boldsymbol{\theta}) \propto p_0(\boldsymbol{\theta}) \tilde{f}_1(\boldsymbol{\theta}) \tilde{f}_2(\boldsymbol{\theta}) \tilde{f}_3(\boldsymbol{\theta})$$


We tie the factor approximations

$$p(\boldsymbol{\theta}) \propto p_0(\boldsymbol{\theta}) f_1(\boldsymbol{\theta}) f_2(\boldsymbol{\theta}) f_3(\boldsymbol{\theta}) \approx q(\boldsymbol{\theta}) \propto p_0(\boldsymbol{\theta}) \tilde{f}(\boldsymbol{\theta})^N$$


Approximate Expectation Propagation

$$p(\theta) \propto p_0(\theta) f_1(\theta) f_2(\theta) f_3(\theta) \approx q(\theta) \propto p_0(\theta) \tilde{f}_1(\theta) \tilde{f}_2(\theta) \tilde{f}_3(\theta)$$


We tie the factor approximations



$$p(\theta) \propto p_0(\theta) f_1(\theta) f_2(\theta) f_3(\theta) \approx q(\theta) \propto p_0(\theta) \tilde{f}(\theta)^N$$


- $\max_q \min_{\tilde{f}_1, \dots, \tilde{f}_N}$ problem $\rightarrow$ $\max_q$ problem, **no double-loop needed!**

We only need to optimize

$$\log Z_{\text{EP}} = \log Z_q + \sum_{n=1}^N \log \mathbf{E}_q \left[\left(\frac{f_n(\boldsymbol{\theta})}{\tilde{f}(\boldsymbol{\theta})} \right) \right] .$$

We only need to optimize

$$\log Z_{\text{EP}} = \log Z_q + \sum_{n=1}^N \log \mathbf{E}_q \left[\left(\frac{f_n(\boldsymbol{\theta})}{\tilde{f}(\boldsymbol{\theta})} \right) \right] .$$

But this requires integrating the exact likelihood factors (**intractable**).

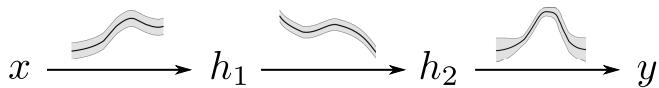
We only need to optimize

$$\log Z_{\text{EP}} = \log Z_q + \sum_{n=1}^N \log \mathbf{E}_q \left[\left(\frac{f_n(\boldsymbol{\theta})}{\tilde{f}(\boldsymbol{\theta})} \right) \right] .$$

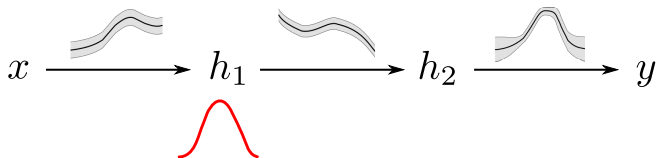
But this requires integrating the exact likelihood factors (**intractable**).

Solution: **moment match** each GP output to a Gaussian at each layer.

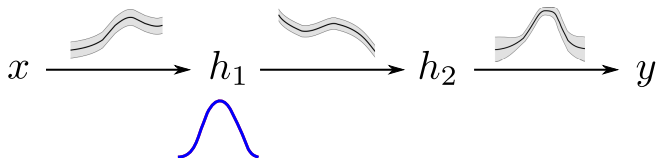
Approximate Gaussian message propagation



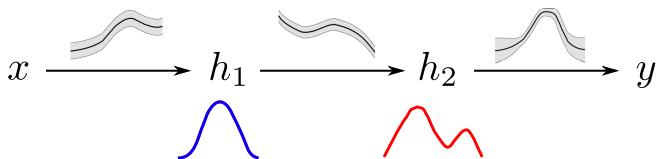
Approximate Gaussian message propagation



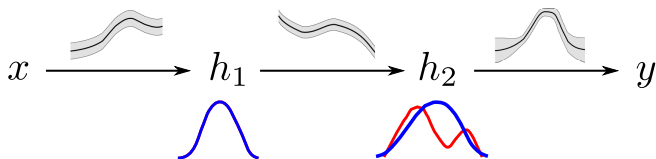
Approximate Gaussian message propagation



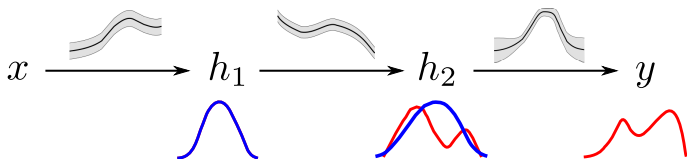
Approximate Gaussian message propagation



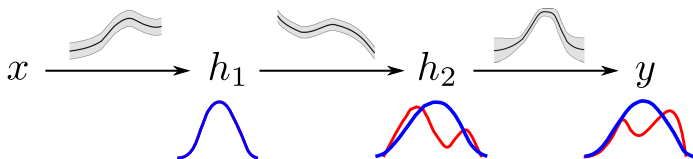
Approximate Gaussian message propagation



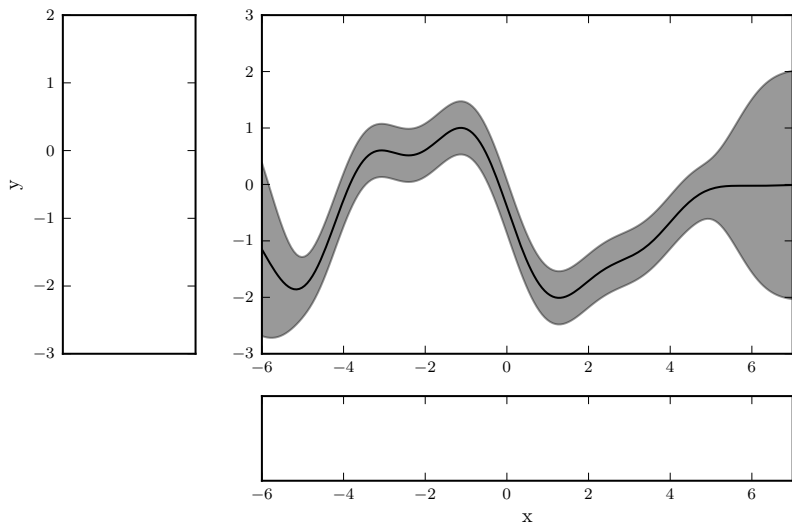
Approximate Gaussian message propagation



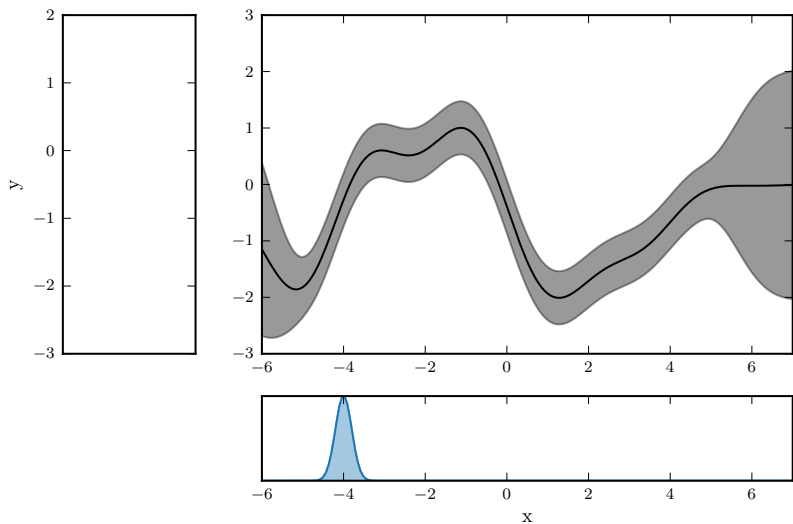
Approximate Gaussian message propagation



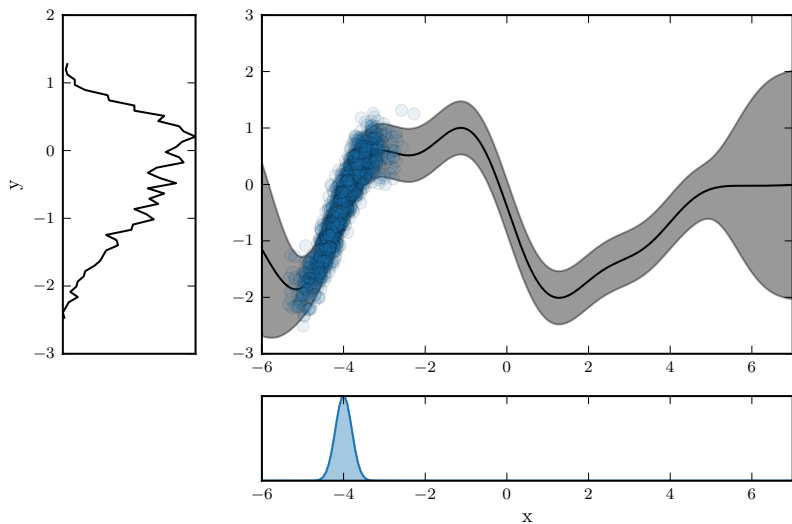
Gaussian projection



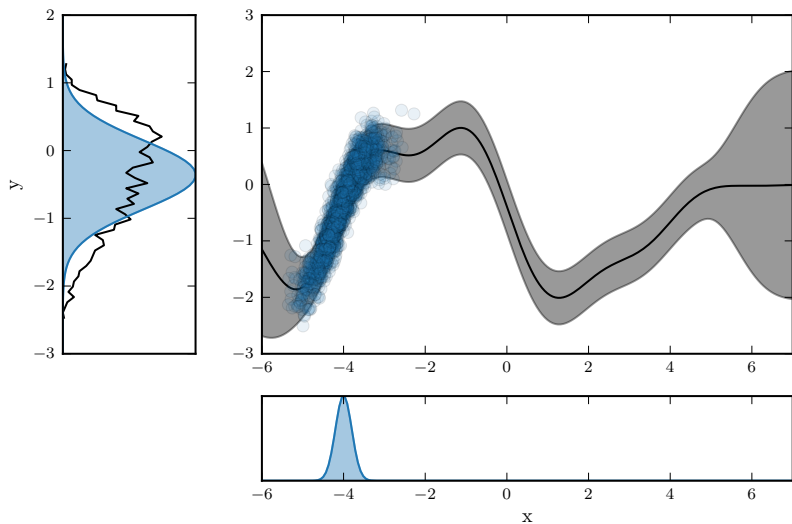
Gaussian projection



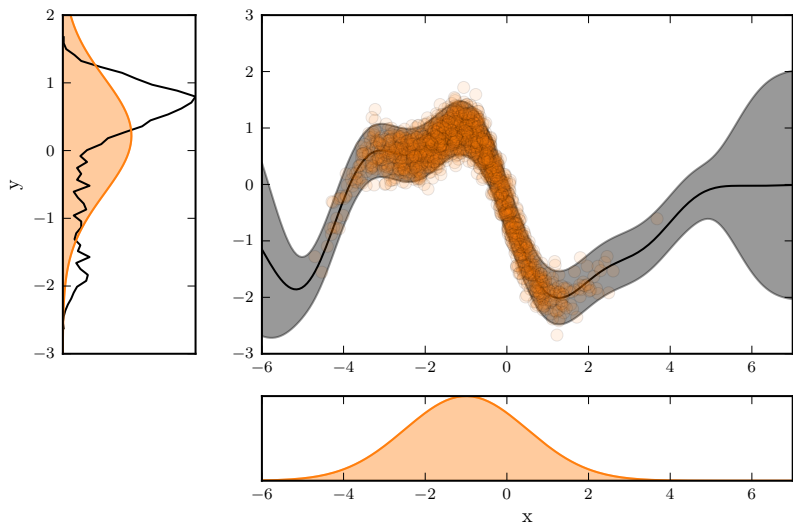
Gaussian projection



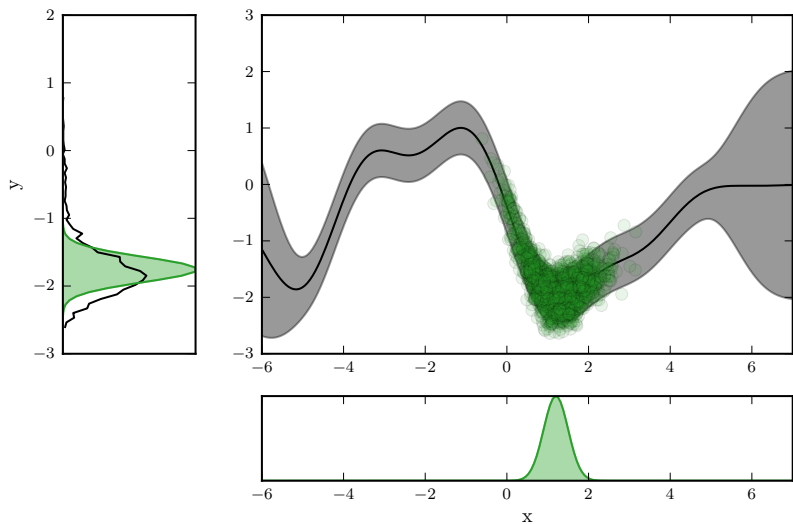
Gaussian projection



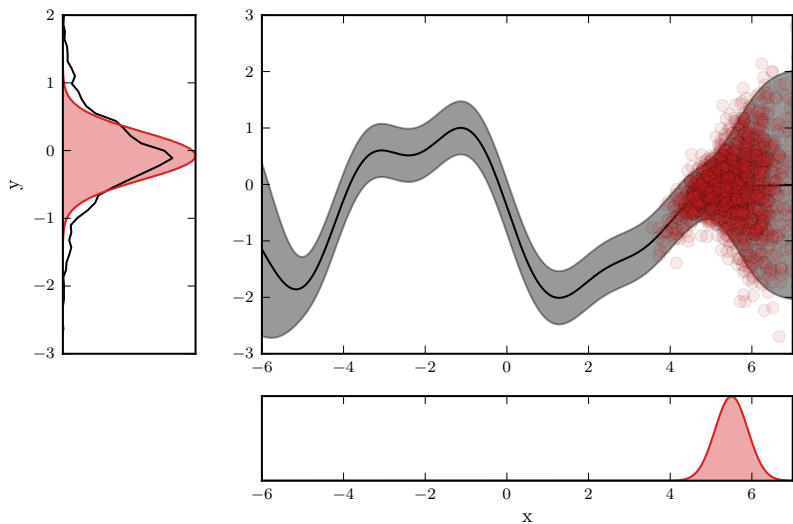
Gaussian projection



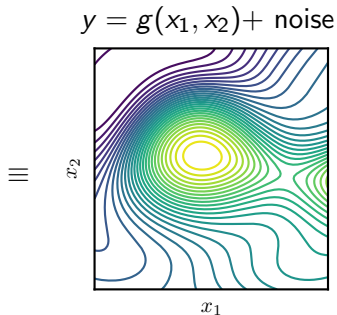
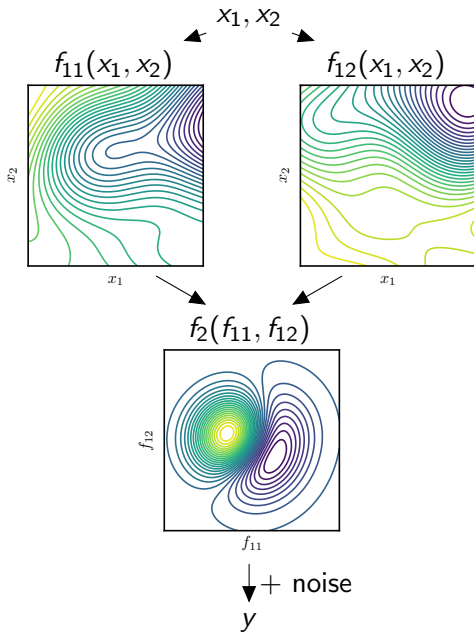
Gaussian projection



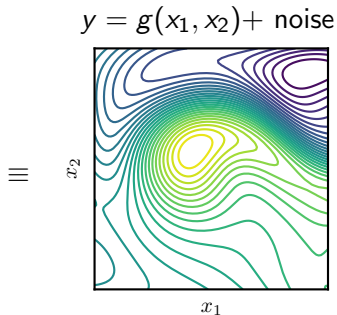
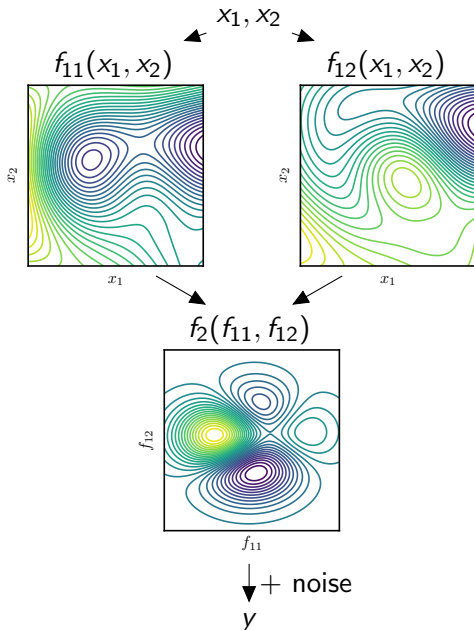
Gaussian projection



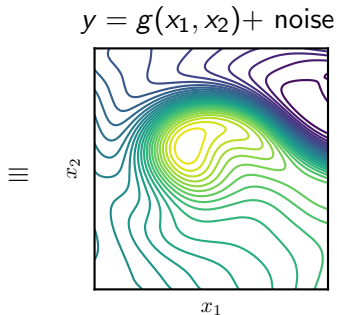
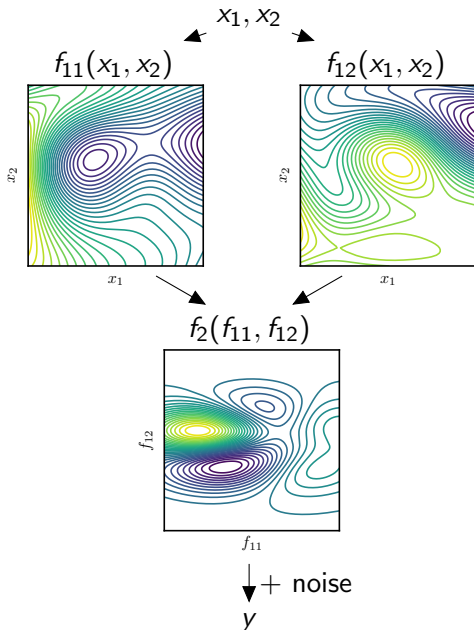
Training deep GPs: 1 epochs



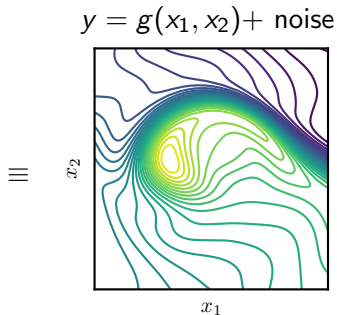
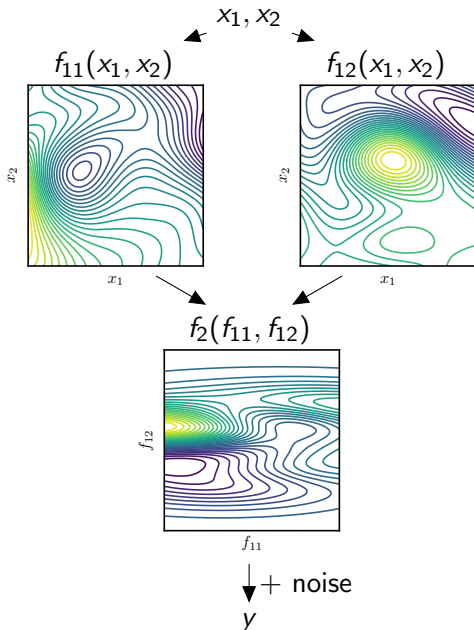
Training deep GPs: 2 epochs



Training deep GPs: 6 epochs

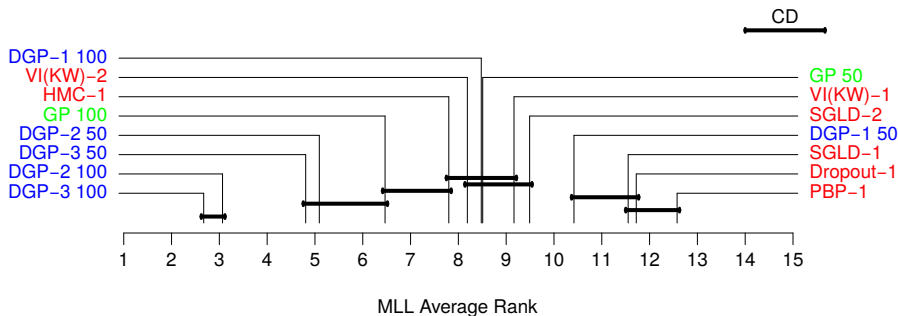


Training deep GPs: 25 epochs



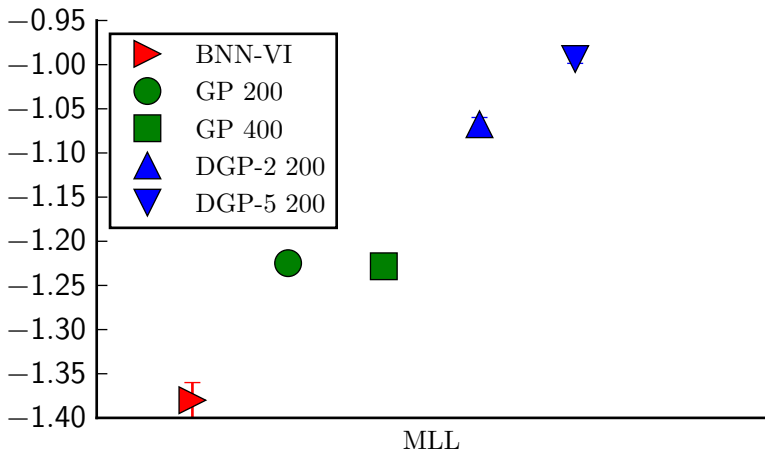
Comparison to Bayesian neural networks

Rankings of all methods across ten different regression datasets



Efficiency of organic photovoltaic molecules

Dataset: 50k/10k training/test points, 512-dim. binary input features



Summary

Novel approach for regression with Deep GPs.

- Inference based on approximate EP energy minimization.
- Extensive comparison to GPs and Bayesian neural networks.
- We obtain very good predictive performance.