

Black-box α -divergence Minimization

José Miguel Hernández-Lobato¹

Harvard Intelligent Probabilistic Systems Group
School of Engineering and Applied Sciences
Harvard University

<http://jmh1.org>, jmh@seas.harvard.edu

Joint work with
Yingzhen Li¹, Mark Rowland, Daniel Hernández-Lobato,
Thang D. Bui and Rich E. Turner.

¹Equal contributors.

$$p(y|\text{Data}) = \int p(y|\theta)p(\theta|\text{Data})d\theta$$

Inference algorithm

Model's
predictive
distribution

$$p(y|\text{Data}) = \int p(y|\theta)p(\theta|\text{Data})d\theta$$

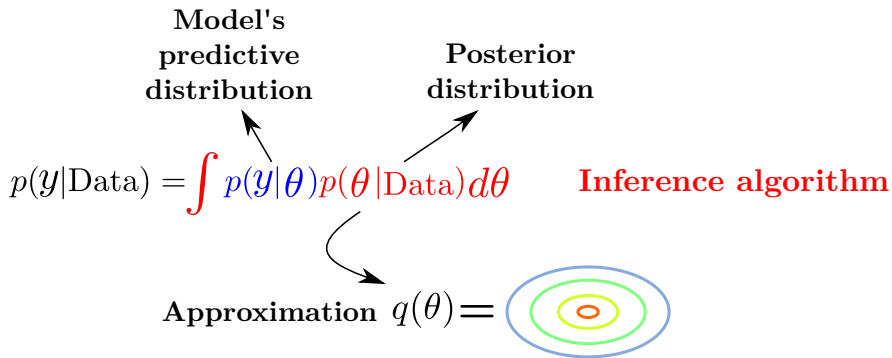

Inference algorithm

Model's
predictive
distribution

Posterior
distribution

$$p(y|\text{Data}) = \int p(y|\theta)p(\theta|\text{Data})d\theta$$

Inference algorithm



α -divergence

$$D_{\alpha}(p||q) = \frac{\int_{\theta} (\alpha p(\theta) + (1 - \alpha)q(\theta) - p(\theta)^{\alpha} q(\theta)^{1-\alpha}) d\theta}{\alpha(1 - \alpha)} .$$

[Amari, 1985].

α -divergence

$$D_{\alpha}(p||q) = \frac{\int_{\theta} (\alpha p(\theta) + (1 - \alpha)q(\theta) - p(\theta)^{\alpha} q(\theta)^{1-\alpha}) d\theta}{\alpha(1 - \alpha)}.$$

[Amari, 1985].

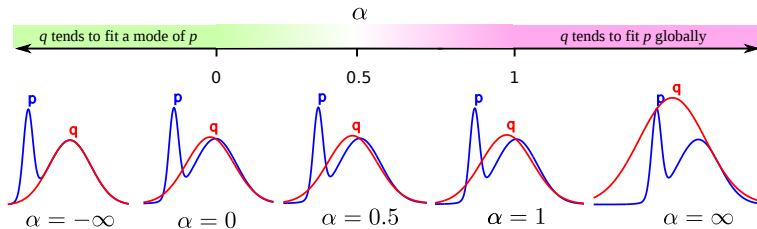


Figure source: [Minka, 2005].

α -divergence

$$D_{\alpha}(p||q) = \frac{\int_{\theta} (\alpha p(\theta) + (1 - \alpha)q(\theta) - p(\theta)^{\alpha} q(\theta)^{1-\alpha}) d\theta}{\alpha(1 - \alpha)}.$$

[Amari, 1985].

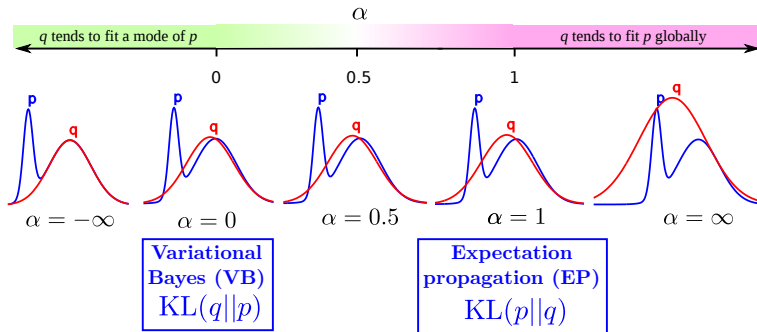


Figure source: [Minka, 2005].

α -divergence

$$D_{\alpha}(p||q) = \frac{\int_{\theta} (\alpha p(\theta) + (1 - \alpha)q(\theta) - p(\theta)^{\alpha} q(\theta)^{1-\alpha}) d\theta}{\alpha(1 - \alpha)}.$$

[Amari, 1985].

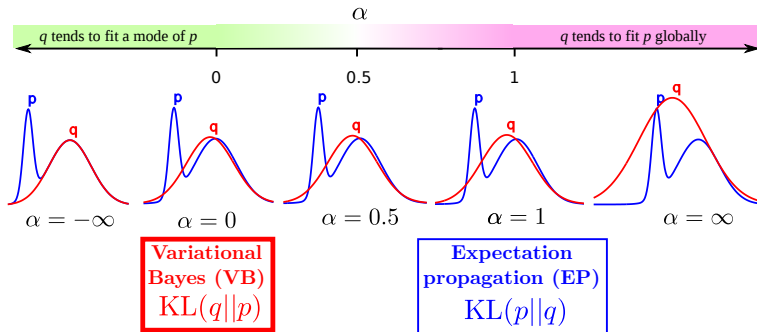


Figure source: [Minka, 2005].

Variational Bayes

$$\lim_{\alpha \rightarrow 0} D_{\alpha}(p||q) = KL(q||p) = -\mathbf{E}_q[\log p(\theta)] - \mathbf{H}[q]$$

Variational Bayes

$$\lim_{\alpha \rightarrow 0} D_{\alpha}(p||q) = KL(q||p) = -\mathbf{E}_q[\log p(\theta)] - \mathbf{H}[q]$$

Variational Bayes

$$\lim_{\alpha \rightarrow 0} D_{\alpha}(p||q) = KL(q||p) = -\mathbf{E}_q[\log p(\theta)] - \mathbf{H}[q]$$

Approximated
by Monte Carlo

A blue oval encircles the term $-\mathbf{E}_q[\log p(\theta)]$ in the equation. A blue arrow points from the text "Approximated by Monte Carlo" to this oval.

Variational Bayes

$$\lim_{\alpha \rightarrow 0} D_{\alpha}(p||q) = KL(q||p) = -\mathbf{E}_q[\log p(\theta)] - \mathbf{H}[q]$$

Approximated by Monte Carlo



There are automatic (black-box) tools for minimizing this objective using

- Stochastic optimization.
- Automatic differentiation.

[Kucukelbir et al., 2015, Ranganath et al., 2014, Salimans et al., 2013].

Variational Bayes

$$\lim_{\alpha \rightarrow 0} D_{\alpha}(p||q) = KL(q||p) = -\mathbf{E}_q[\log p(\theta)] - \mathbf{H}[q]$$

Approximated by Monte Carlo



There are automatic (black-box) tools for minimizing this objective using

- Stochastic optimization.
- Automatic differentiation.

[Kucukelbir et al., 2015, Ranganath et al., 2014, Salimans et al., 2013].

Can we have similar tools for other values of α ?

Local α -divergence minimization (Power EP)

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

[Minka, 2004]

Local α -divergence minimization (Power EP)

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

[Minka, 2004]

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


Local α -divergence minimization (Power EP)

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

[Minka, 2004]

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


The $\tilde{f}_n$ are tuned by minimizing the local α -divergences

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

The Power-EP approximation to the **evidence** is given by

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

The Power-EP approximation to the **evidence** is given by

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

The power-EP solution for q can be obtained by solving

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

The Power-EP approximation to the **evidence** is given by

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

The power-EP solution for q can be obtained by solving

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

Solved with **double-loop** algorithm [Heskes et al., 2002].

The Power-EP approximation to the **evidence** is given by

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

The power-EP solution for q can be obtained by solving

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

Solved with **double-loop** algorithm [Heskes et al., 2002]. **Too slow in practice!**

Main contribution

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


We tie the factor approximations



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


Main contribution

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


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!**

Main contribution

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

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!**
- **Memory saving** scales as $\mathcal{O}(N)$.

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

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



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

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



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

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



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



$$\log Z_q + \frac{N}{|\mathbf{S}|} \sum_{n \in \mathbf{S}} \frac{1}{\alpha} \log \frac{1}{K} \sum_{k=1}^K \left(\frac{f_n(\boldsymbol{\theta}_k)}{\tilde{f}(\boldsymbol{\theta}_k)} \right)^\alpha$$

$$\boldsymbol{\theta}_1, \dots, \boldsymbol{\theta}_K \sim q$$

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



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



$$\log Z_q + \frac{N}{|\mathbf{S}|} \sum_{n \in \mathbf{S}} \frac{1}{\alpha} \log \frac{1}{K} \sum_{k=1}^K \left(\frac{f_n(\boldsymbol{\theta}_k)}{\tilde{f}(\boldsymbol{\theta}_k)} \right)^\alpha$$

$$\boldsymbol{\theta}_1, \dots, \boldsymbol{\theta}_K \sim q$$

Stochastic estimate of the objective for **automatic, scalable** inference.

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



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



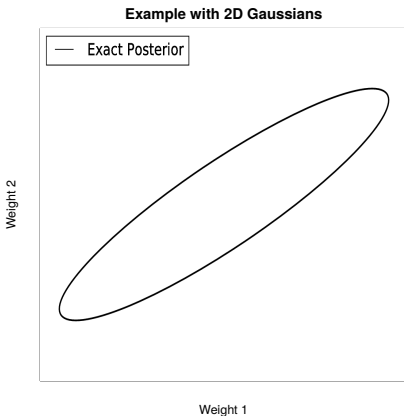
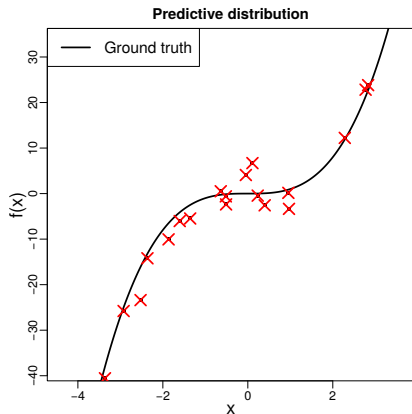
$$\log Z_q + \frac{N}{|\mathbf{S}|} \sum_{n \in \mathbf{S}} \frac{1}{\alpha} \log \frac{1}{K} \sum_{k=1}^K \left(\frac{f_n(\boldsymbol{\theta}_k)}{\tilde{f}(\boldsymbol{\theta}_k)} \right)^\alpha$$

$$\boldsymbol{\theta}_1, \dots, \boldsymbol{\theta}_K \sim q$$

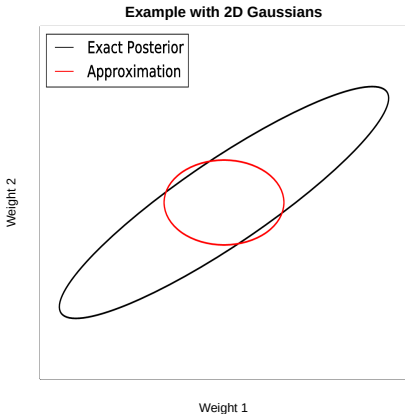
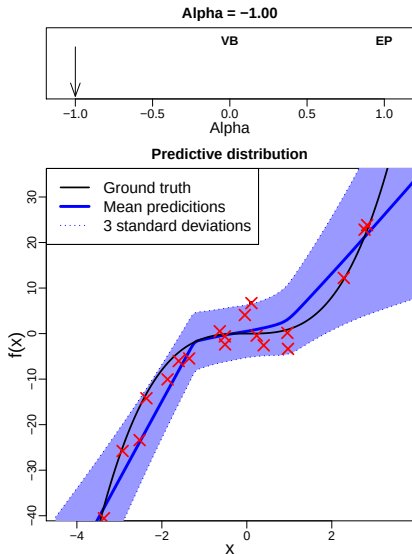
Stochastic estimate of the objective for **automatic, scalable** inference.

Biased estimator!

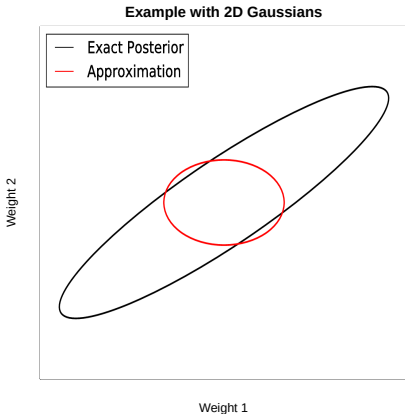
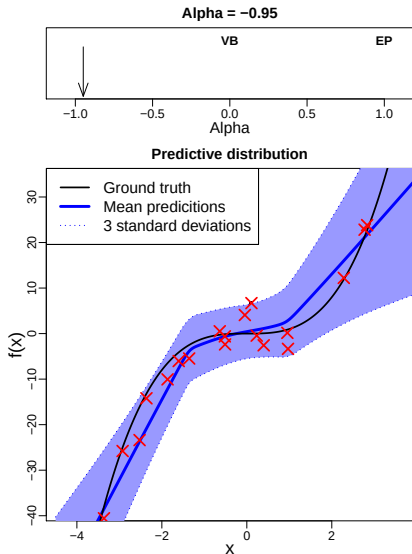
Regression with neural networks and 2D Gaussian



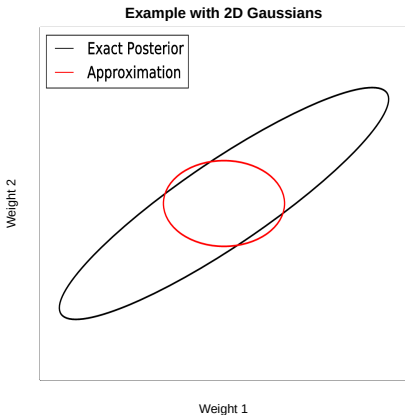
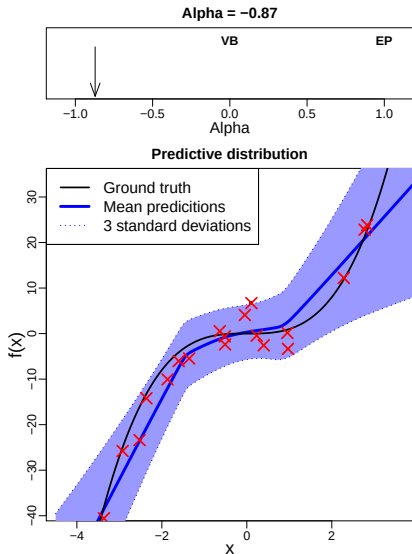
Regression with neural networks and 2D Gaussian



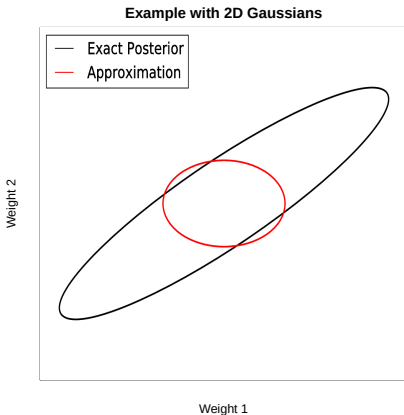
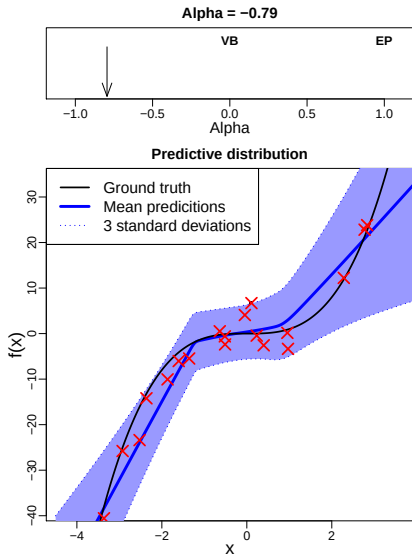
Regression with neural networks and 2D Gaussian



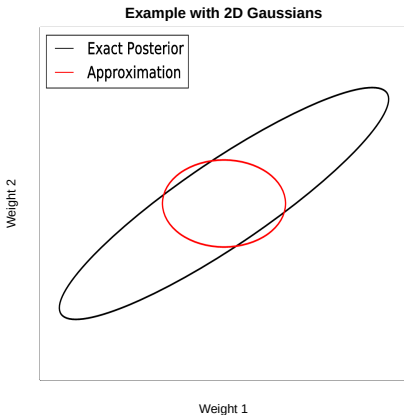
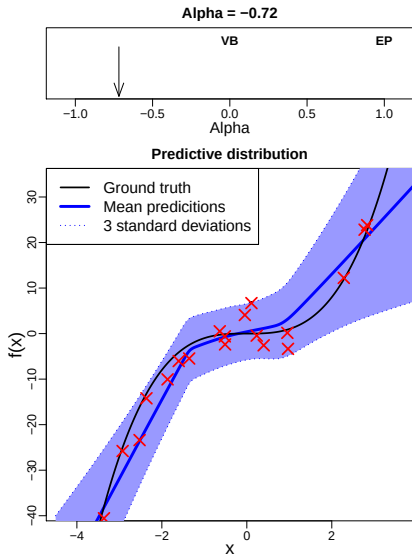
Regression with neural networks and 2D Gaussian



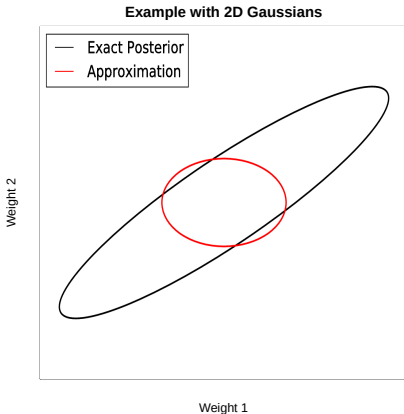
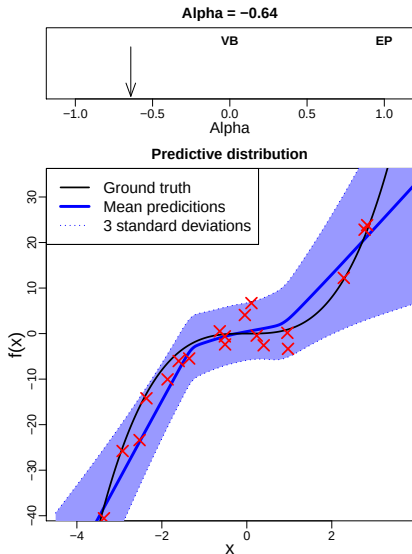
Regression with neural networks and 2D Gaussian



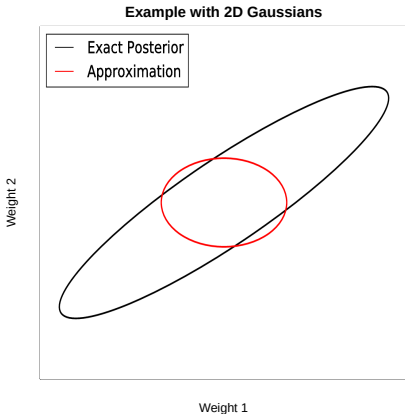
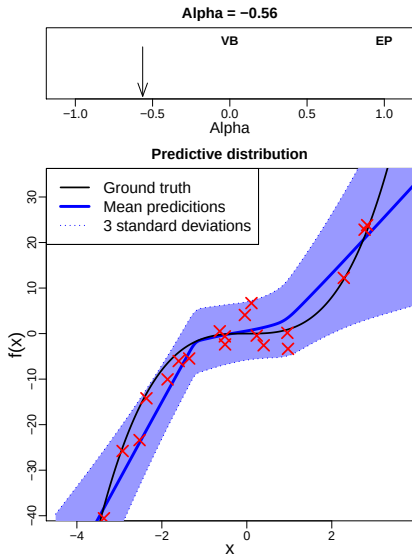
Regression with neural networks and 2D Gaussian



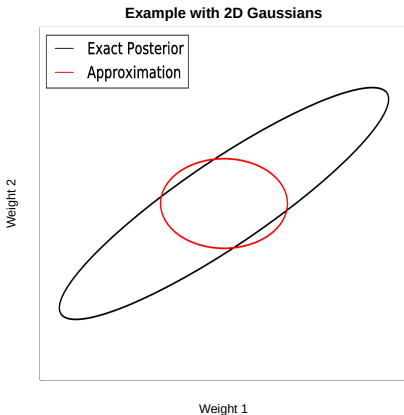
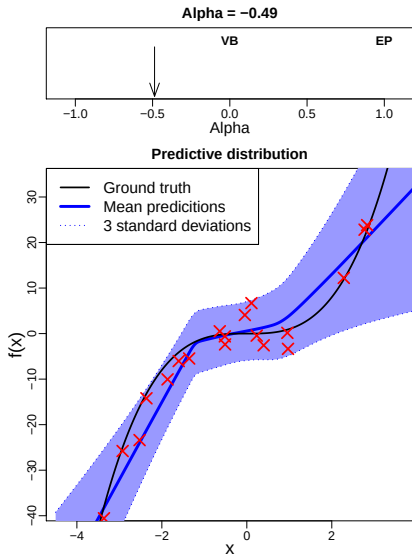
Regression with neural networks and 2D Gaussian



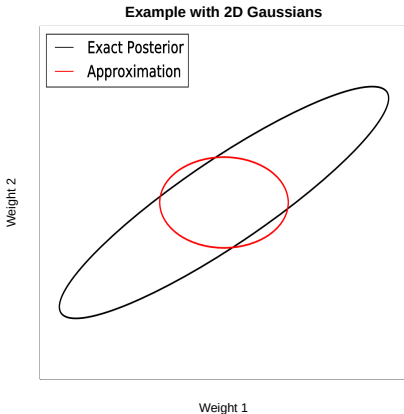
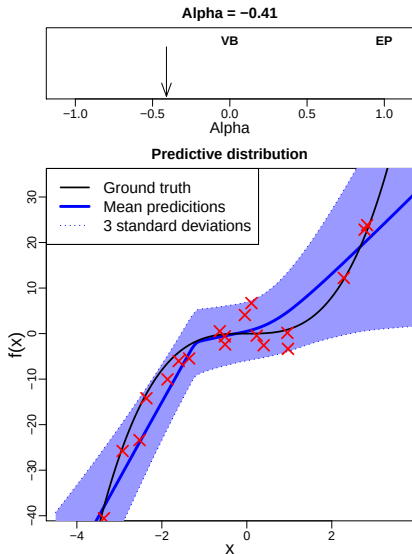
Regression with neural networks and 2D Gaussian



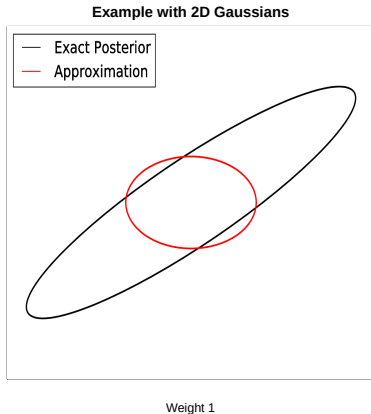
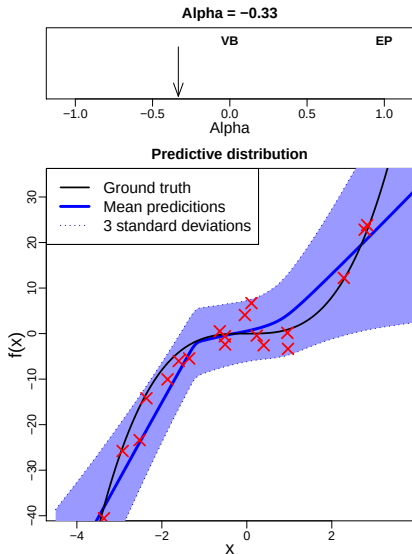
Regression with neural networks and 2D Gaussian



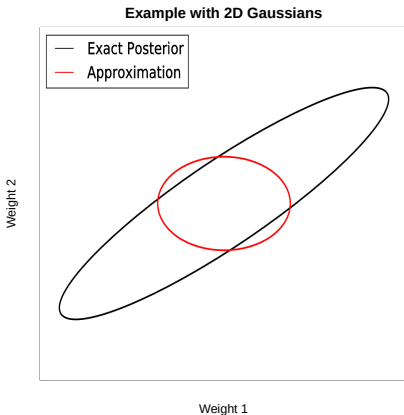
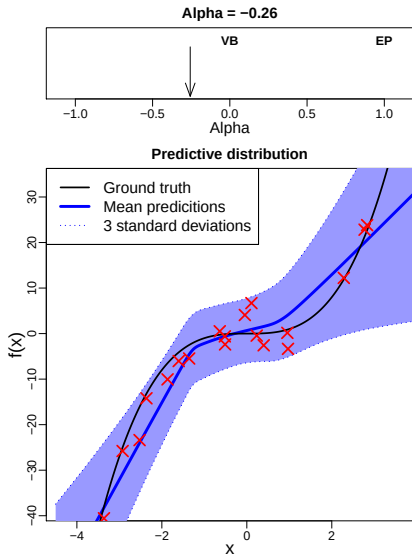
Regression with neural networks and 2D Gaussian



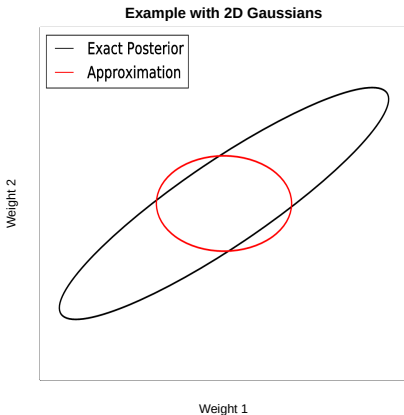
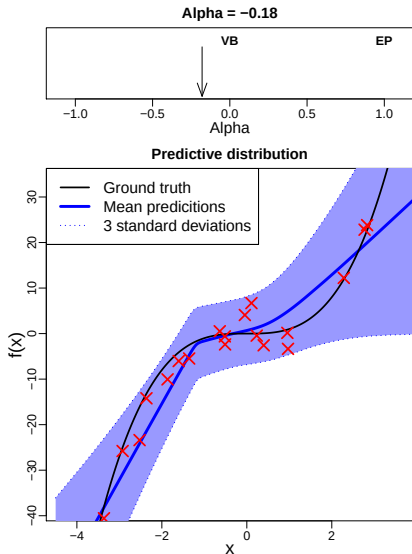
Regression with neural networks and 2D Gaussian



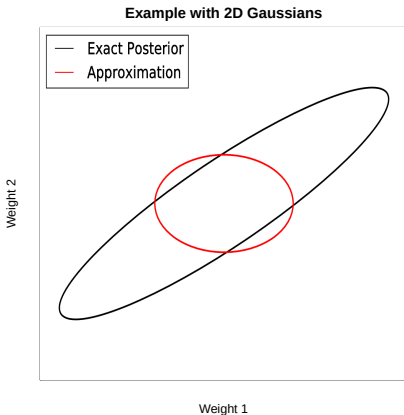
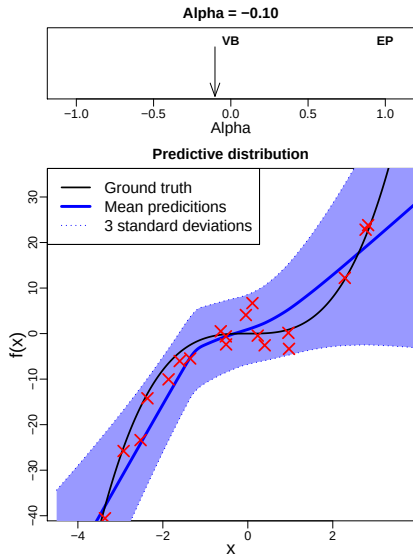
Regression with neural networks and 2D Gaussian



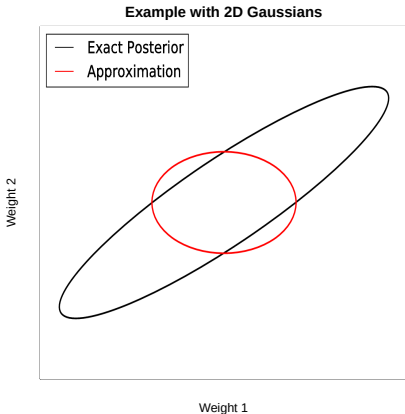
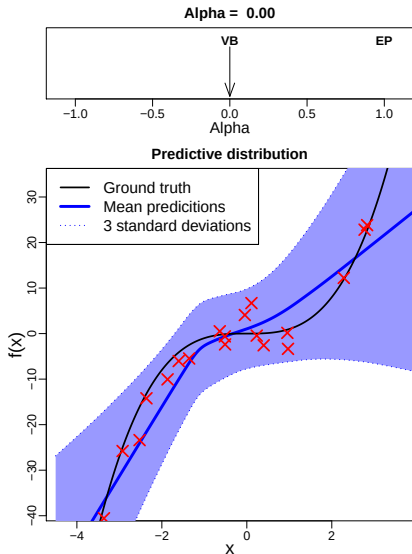
Regression with neural networks and 2D Gaussian



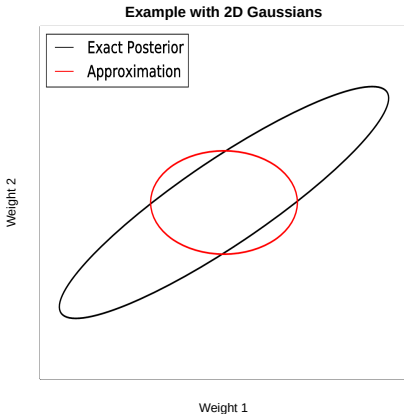
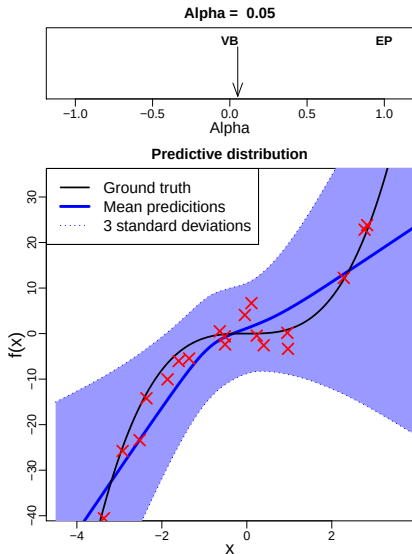
Regression with neural networks and 2D Gaussian



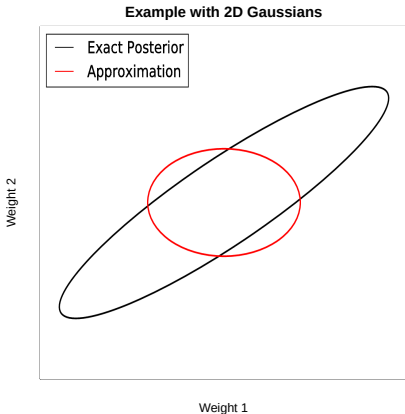
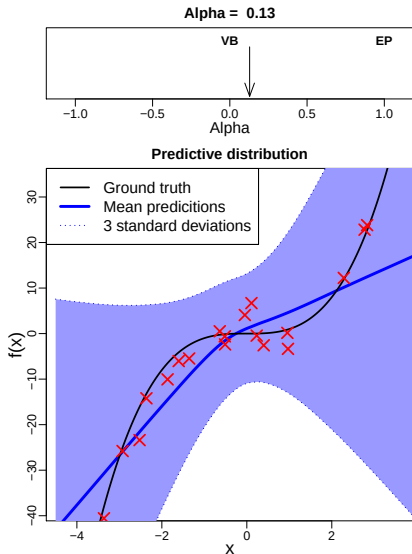
Regression with neural networks and 2D Gaussian



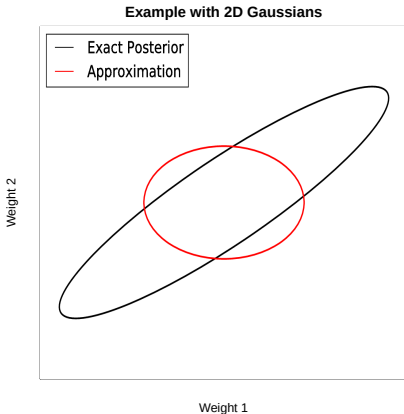
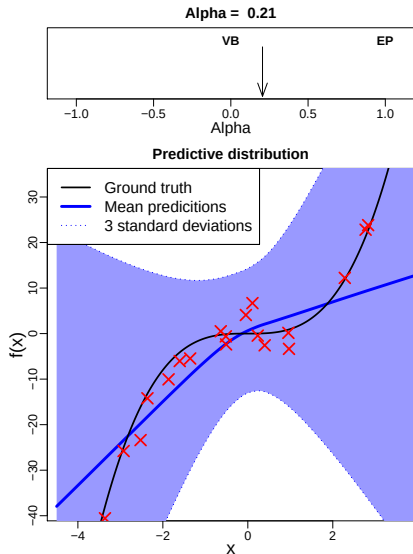
Regression with neural networks and 2D Gaussian



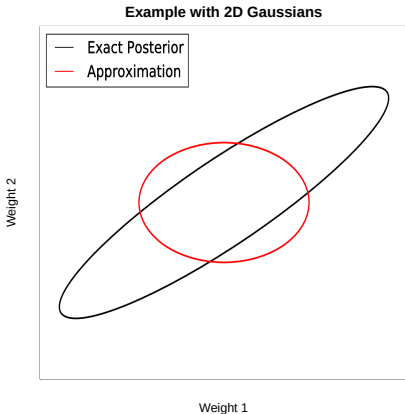
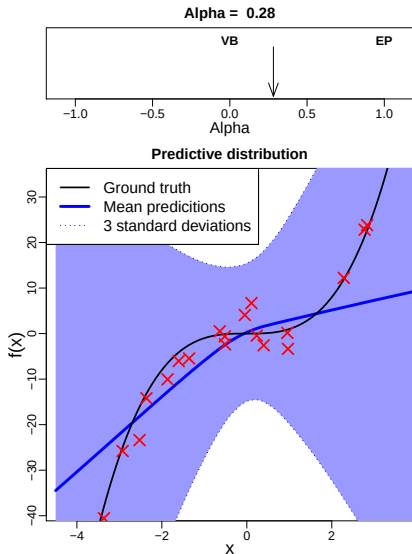
Regression with neural networks and 2D Gaussian



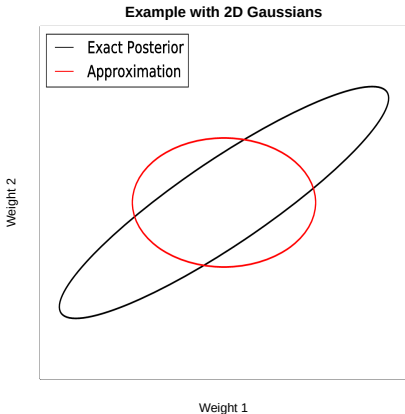
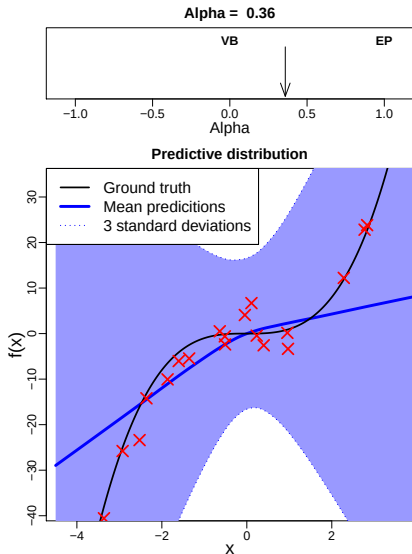
Regression with neural networks and 2D Gaussian



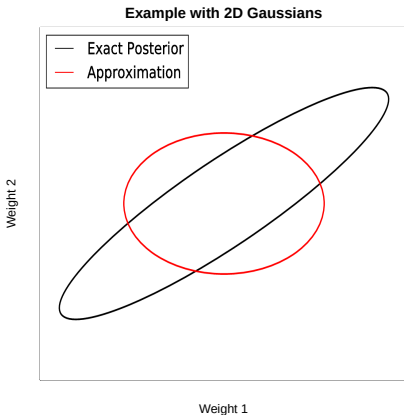
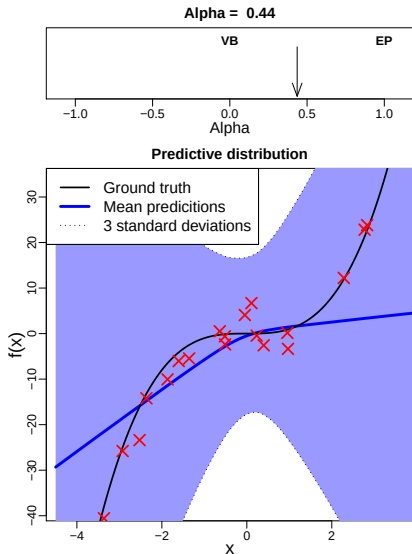
Regression with neural networks and 2D Gaussian



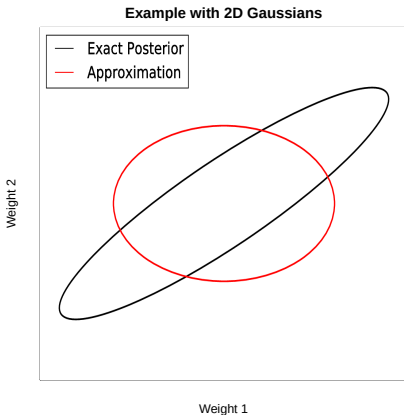
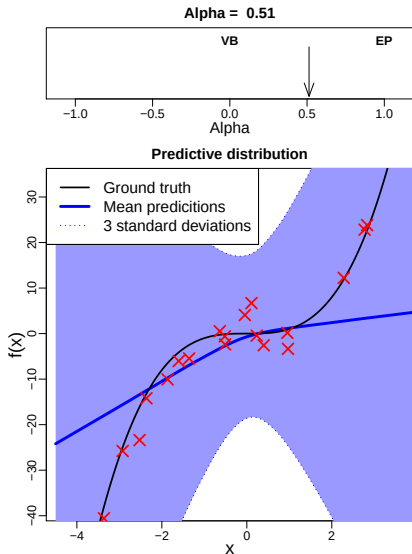
Regression with neural networks and 2D Gaussian



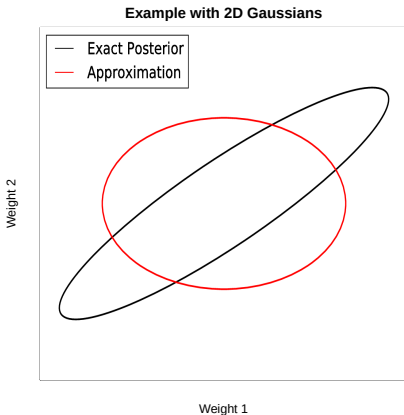
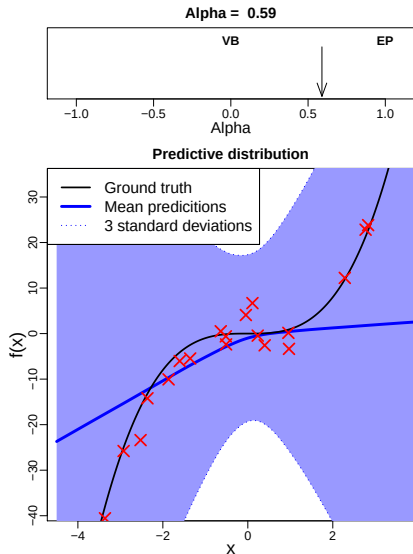
Regression with neural networks and 2D Gaussian



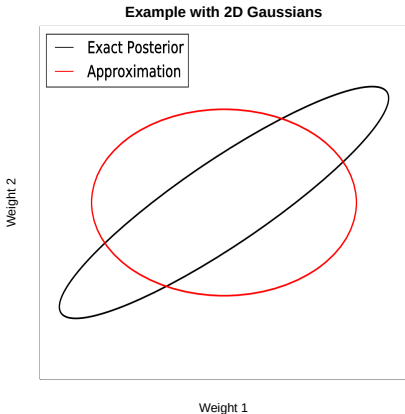
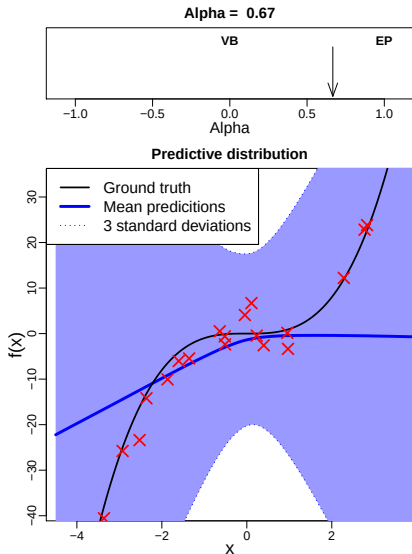
Regression with neural networks and 2D Gaussian



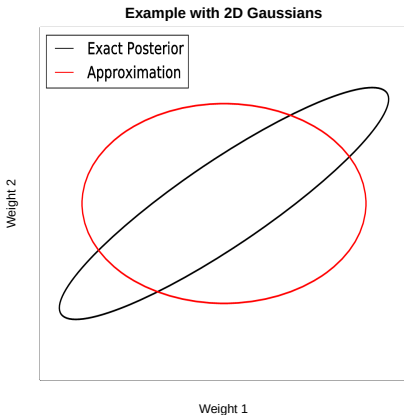
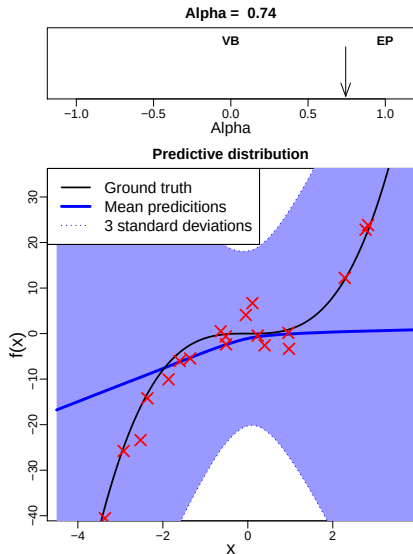
Regression with neural networks and 2D Gaussian



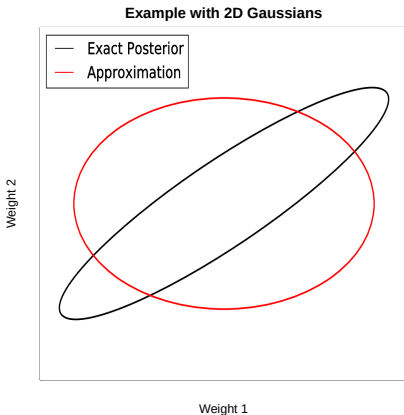
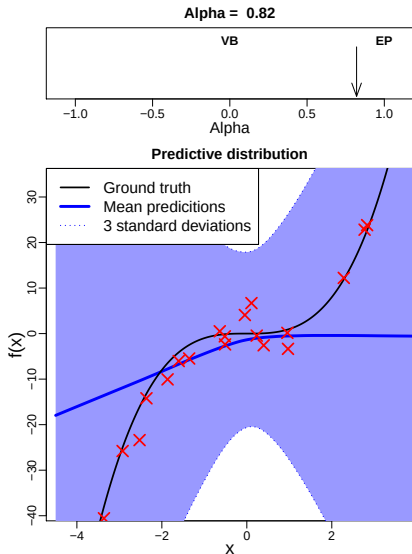
Regression with neural networks and 2D Gaussian



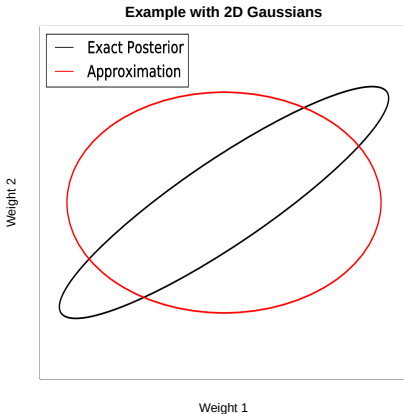
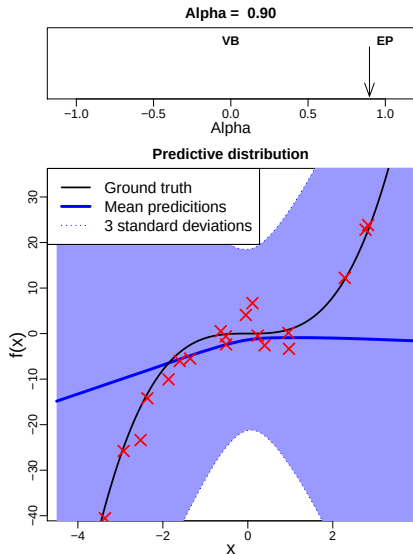
Regression with neural networks and 2D Gaussian



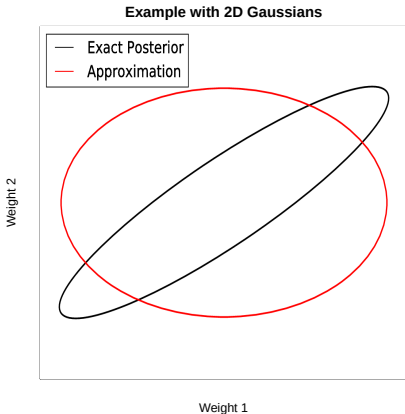
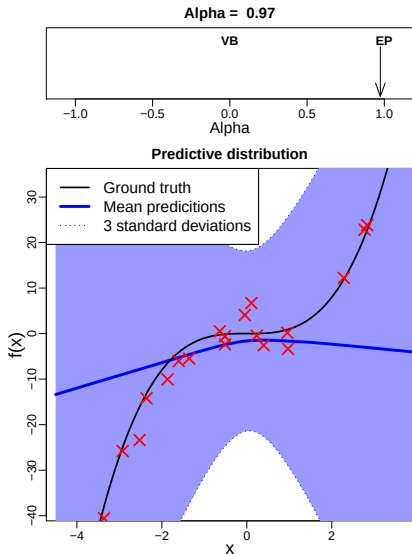
Regression with neural networks and 2D Gaussian



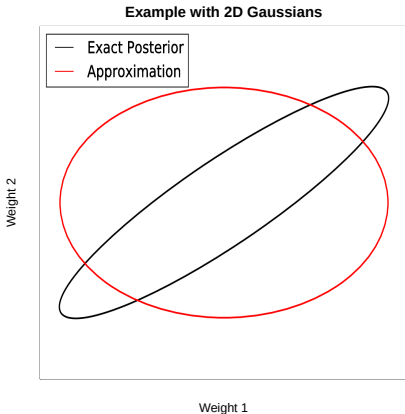
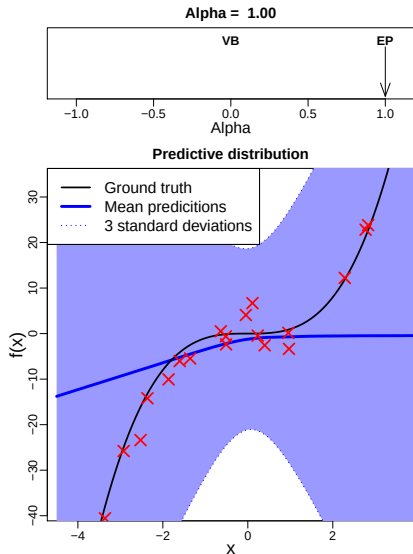
Regression with neural networks and 2D Gaussian



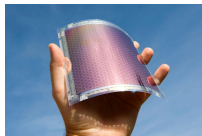
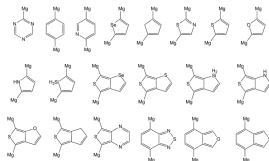
Regression with neural networks and 2D Gaussian



Regression with neural networks and 2D Gaussian

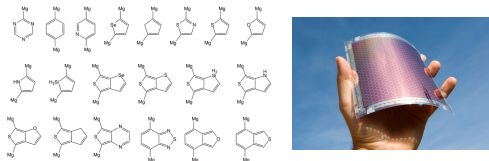


Harvard clean energy project. Data from 60.000 **organic photovoltaics**.



Bayesian neural networks, 2 hidden layers with 400 units each.

Harvard clean energy project. Data from 60.000 **organic photovoltaics**.

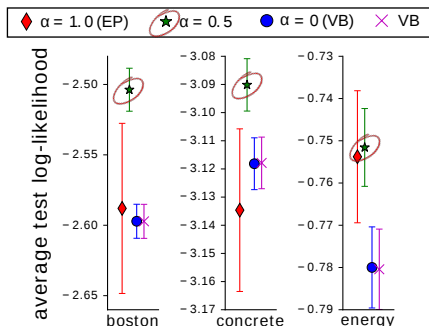


Bayesian neural networks, 2 hidden layers with 400 units each.

Table: Average Test Error and Test Log-likelihood.

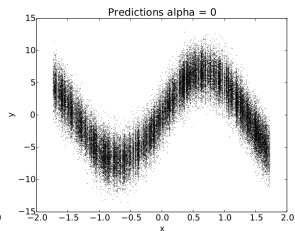
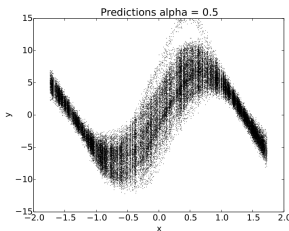
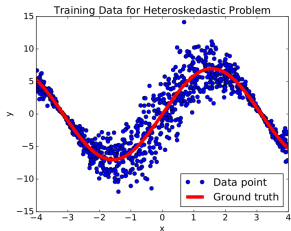
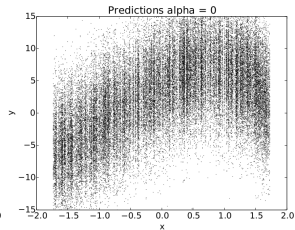
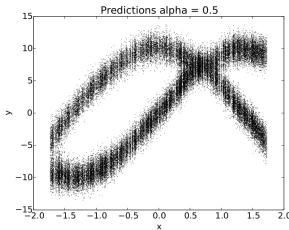
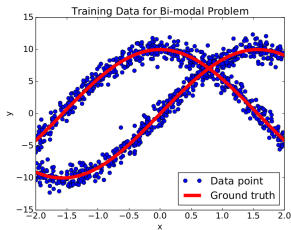
	$\alpha=1.0$ (EP)	$\alpha=0.5$	$\alpha=0$	VB
Avg. Error	1.28 ± 0.01	1.08 ± 0.01	1.13 ± 0.01	1.14 ± 0.01
Avg. Log-likelihood	-0.93 ± 0.01	-0.74 ± 0.01	-1.39 ± 0.03	-1.38 ± 0.02

The value $\alpha = 0.5$ seems to be consistently better than $\alpha = 1$ (EP) or $\alpha = 0$ (VB).



$\alpha = 0.5$

Variational Bayes



$$\frac{1}{\alpha} \log \mathbf{E}_{q^{\text{cav}}} [f_n(\theta)^\alpha]$$

$$\mathbf{E}_q [\log f_n(\theta)]$$

Take home message

Black-box α -divergence minimization...

- 1 generalizes VB and an EP-like algorithm.
- 2 better than any of them in prediction problems with neural nets.
- 3 straightforward to apply in very complex models.

- Important applications in **probabilistic programming** to go beyond MCMC and VB.



References I

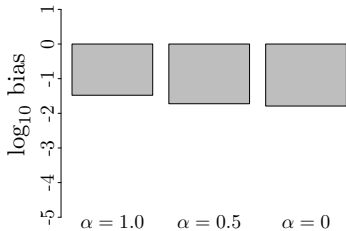
- S. Amari. *Differential-geometrical methods in statistics*, volume 28. Springer Science & Business Media, 1985.
- S. Depeweg, J. M. Hernández-Lobato, F. Doshi-Velez, and S. Udluft. Learning and policy search in stochastic dynamical systems with bayesian neural networks. *arXiv preprint arXiv:1605.07127*, 2016.
- Heskes et al. Expectation propagation for approximate inference in dynamic bayesian networks. In *Proceedings of the Eighteenth conference on Uncertainty in artificial intelligence*, 2002.
- A. Kucukelbir, R. Ranganath, A. Gelman, and D. Blei. Automatic variational inference in stan. In C. Cortes, N. Lawrence, D. Lee, M. Sugiyama, and R. Garnett, editors, *Advances in Neural Information Processing Systems 28*, pages 568–576. Curran Associates, Inc., 2015.
- T. Minka. Divergence measures and message passing. Technical report, Microsoft Research, 2005.

References II

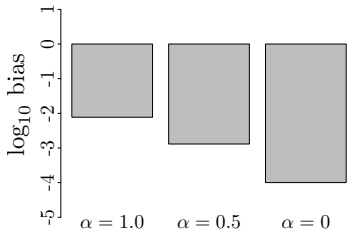
- T. P. Minka. Power ep. Technical report, Technical report, Microsoft Research, Cambridge, 2004.
- R. Ranganath, S. Gerrish, and D. Blei. Black box variational inference. In *AISTATS*, pages 814–822, 2014.
- T. Salimans, D. A. Knowles, et al. Fixed-form variational posterior approximation through stochastic linear regression. *Bayesian Analysis*, 8(4):837–882, 2013.

Average Bias in the Gradient.

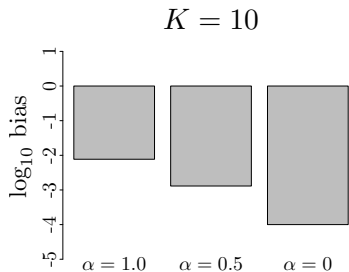
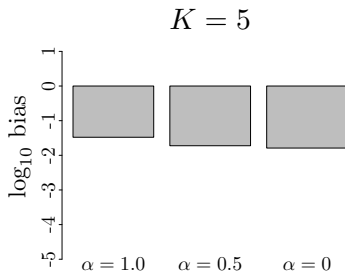
$K = 5$



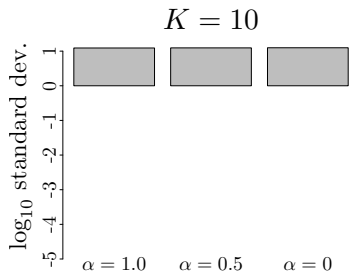
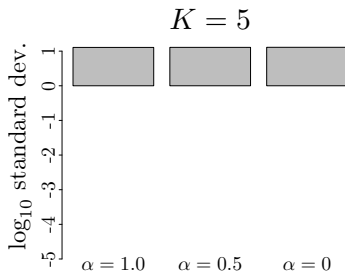
$K = 10$



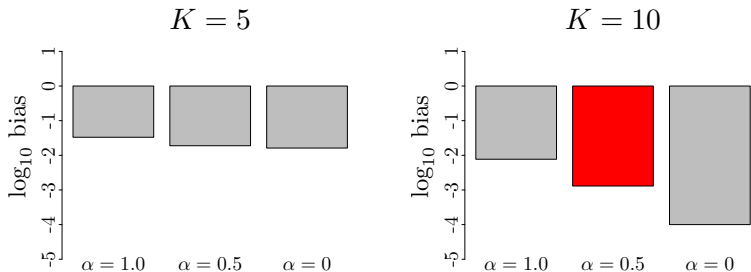
Average Bias in the Gradient.



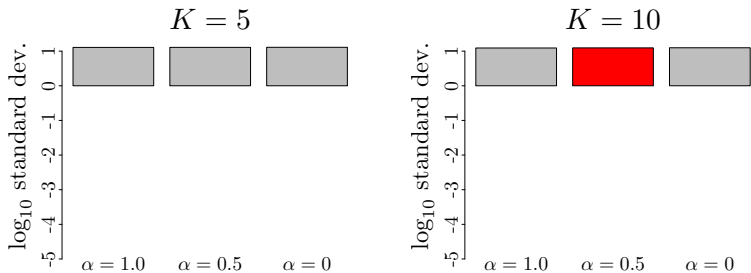
Average Standard Deviation in the Gradient.



Average Bias in the Gradient.



Average Standard Deviation in the Gradient.



Multi-class classification

Bayesian neural networks, 2 hidden layers with 400 units each.

Table: Average Test Error and Test Log-likelihood in MNIST.

MNIST	$\alpha=1.0$	$\alpha=0.5$	$\alpha=0$	VB	$\alpha=-1.0$
Error	0.0151 ± 0.0001	0.0144 ± 0.0001	0.0136 ± 0.0001	0.0136 ± 0.0001	0.0133 ± 0.0001
Log-likelihood	-0.0551 ± 0.0004	-0.0509 ± 0.0003	-0.0468 ± 0.0002	-0.0468 ± 0.0002	-0.0447 ± 0.0002