

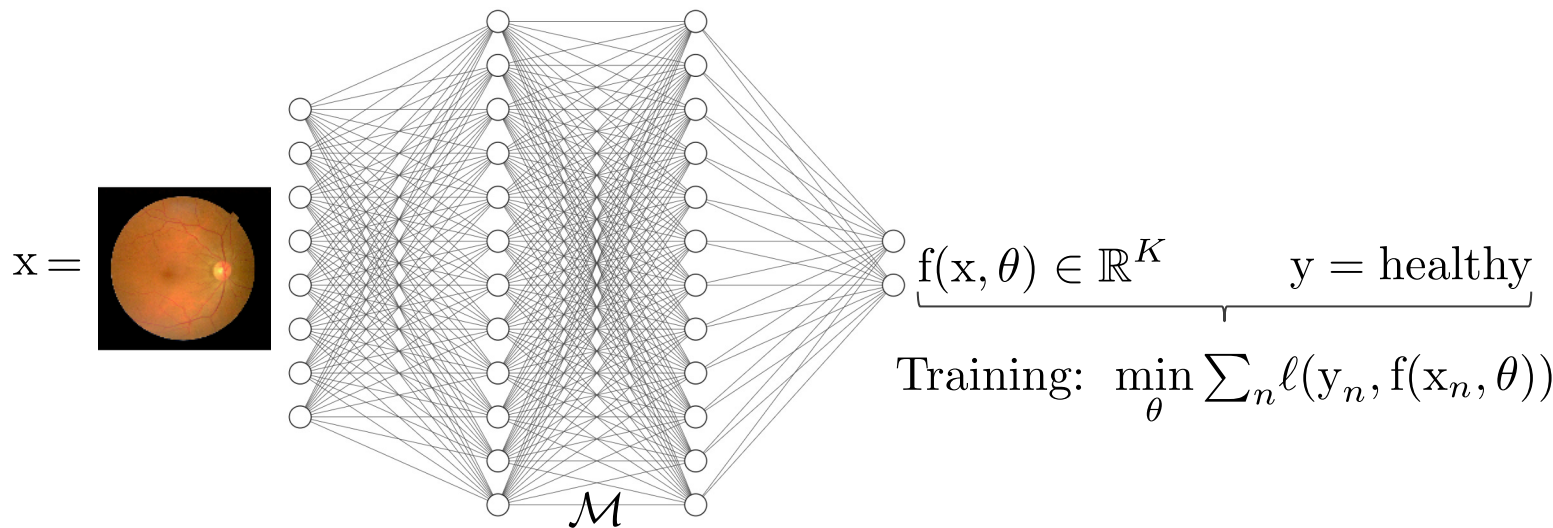
Laplace Approximations for Deep Learning

*A Bayesian Odyssey in Uncertainty: from Theoretical
Foundations to Real-World Applications*

Alexander Immer

Deep Learning

Data $\mathcal{D} = \{(\mathbf{x}_n, y_n)\}_{n=1}^N$, weights $\theta \in \mathbb{R}^P$, and model hyperparameters $\mathcal{M}$

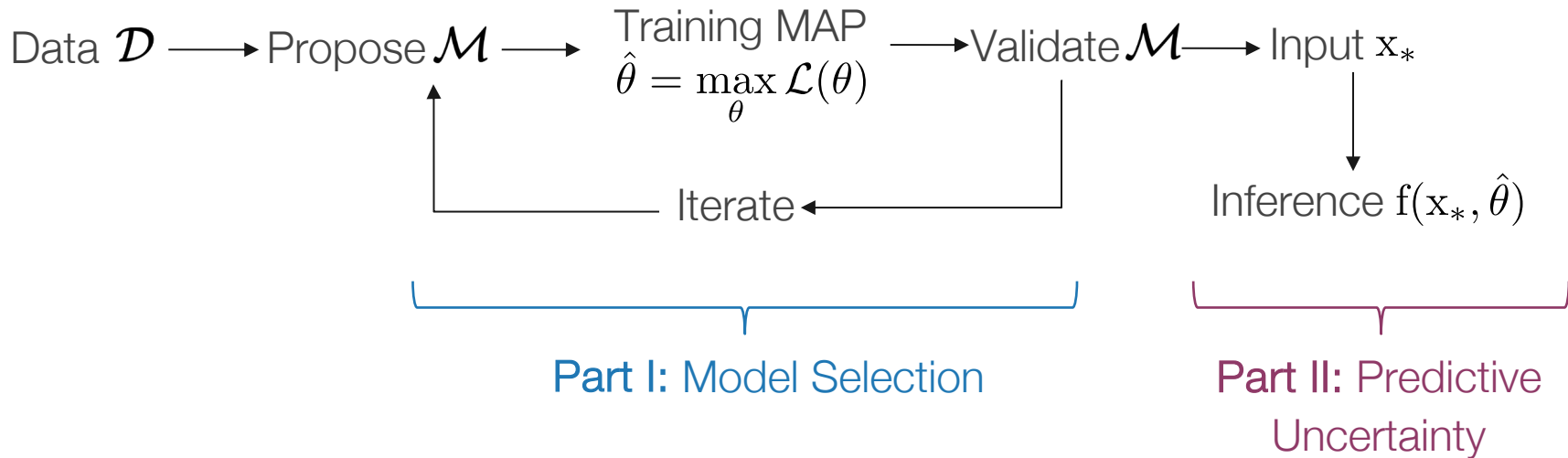


→ Successful at learning complex patterns from large data sets

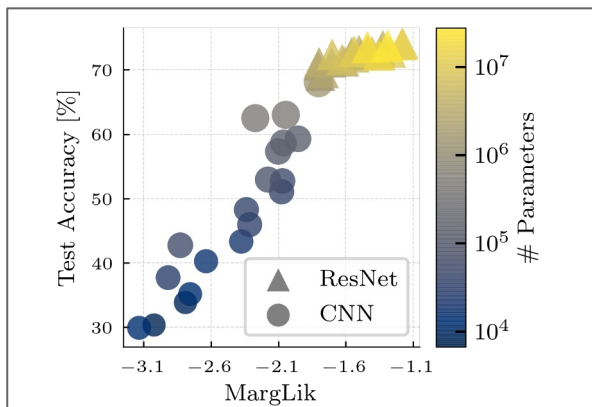
Probabilistic View

$$\begin{array}{ccc} \min_{\theta} \sum_n \ell(y_n, f(x_n, \theta)) + R(\theta) & & \\ \uparrow \text{Data-fit} & & \uparrow \text{Regularization} \\ \max_{\theta} \sum_n \log \underbrace{p(y_n | x_n, \theta, \mathcal{M})}_{\text{Likelihood}} + \log \underbrace{p(\theta | \mathcal{M})}_{\text{Prior}} & & \\ \underbrace{\hspace{15em}} & & \\ \mathcal{L}(\theta) & & \\ \text{Maximum a-posteriori objective} & & \end{array}$$

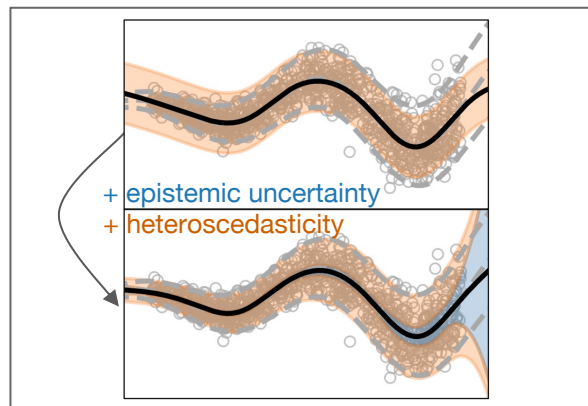
Deep Learning Workflow



Part I: Bayesian Model Selection with Laplace



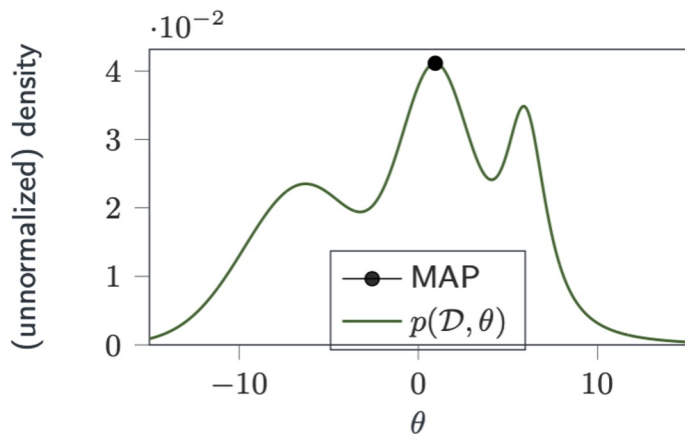
Part II: Predictive Uncertainties with Laplace



The Bayesian Approach

MAP

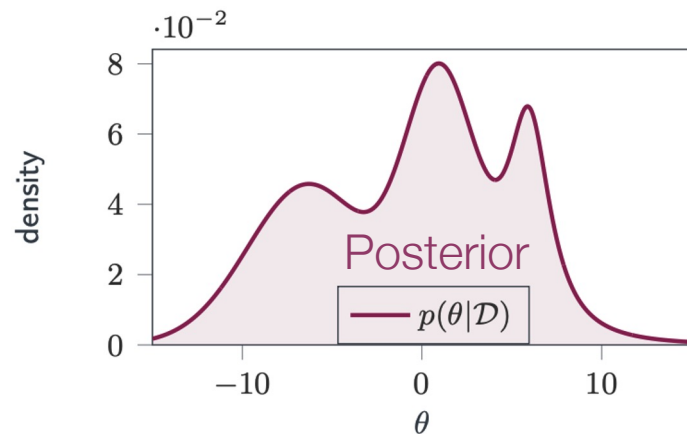
$$\max_{\theta} \underbrace{p(\mathcal{D}|\theta, \mathcal{M})}_{\text{Likelihood}} \underbrace{p(\theta|\mathcal{M})}_{\text{Prior}}$$



divide by $\rightarrow$

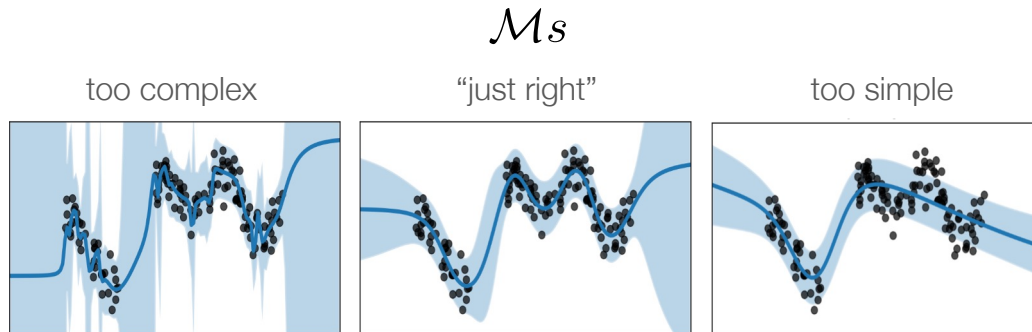
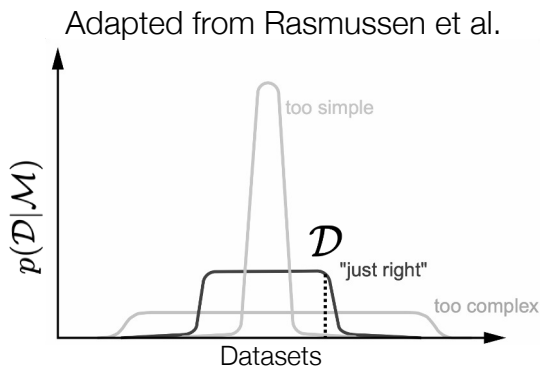
Bayesian

$$\frac{\int p(\mathcal{D}|\theta, \mathcal{M}) p(\theta|\mathcal{M}) d\theta}{\text{Marginal Likelihood}}$$



Advantage 1) Bayesian Model Selection

Optimize marginal likelihood: $p(\mathcal{D}|\mathcal{M}) = \int p(\mathcal{D}|\theta, \mathcal{M}) p(\theta|\mathcal{M}) d\theta$



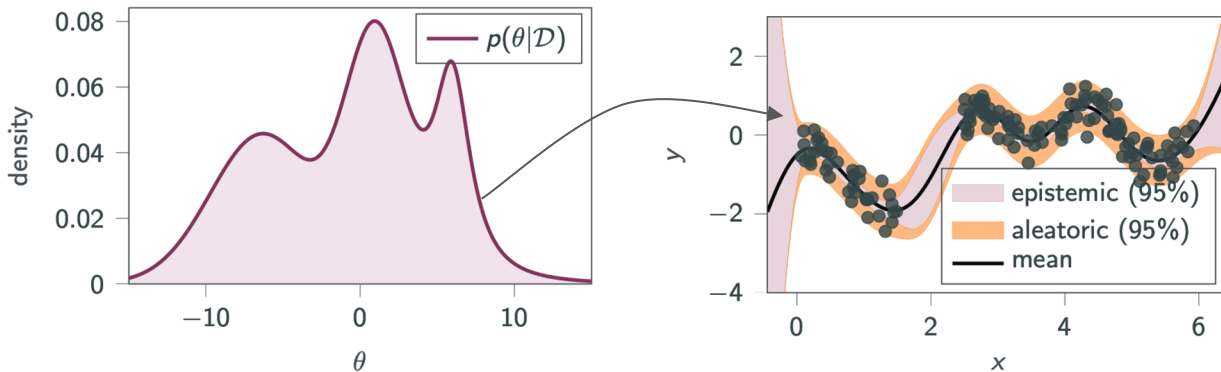
→ Occam's razor: balance data-fit and model complexity^{1,2}

¹Blumer, Ehrenfeucht, Haussler, Warmuth. *Occam's Razor*. Information processing letters 1987.

²Rasmussen, Ghahramani. *Occam's Razor*. NIPS 2001.

Advantage 2) Predictive Uncertainty

Posterior predictive: $p(y_*|x_*, \mathcal{D}) = \int p(y_*|x_*, \theta) \underbrace{p(\theta|\mathcal{D})}_{\text{Posterior}} d\theta$



→ Additionally provide epistemic uncertainty

Aleatoric vs. Epistemic Uncertainty

Potential definition of uncertainties as variance decomposition:

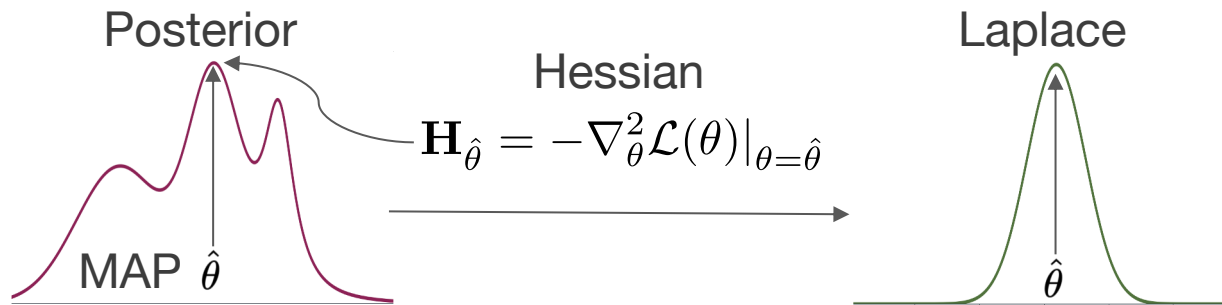
$$\mathbb{V}_{p(y|x_*, \mathcal{D})}[y] = \underbrace{\int \mathbb{V}_{p(y|x_*, \theta)}[y] p(\theta|\mathcal{D}) d\theta}_{\text{aleatoric}} + \underbrace{\int (\mathbb{E}_{p(y|x_*, \theta)}[y] - \mathbb{E}_{p(y|x_*, \mathcal{D})}[y])^2 p(\theta|\mathcal{D}) d\theta}_{\text{epistemic}}$$

Variance in labels for same input $\rightarrow$ aleatoric uncertainty

No similar input to x_* $\rightarrow$ epistemic uncertainty

Laplace Approximations

Laplace Approximation



- 1) MargLik: $\log q(\mathcal{D}|\mathcal{M}) = \log p(\mathcal{D}, \hat{\theta}|\mathcal{M}) - \frac{1}{2} \log \left| \frac{1}{2\pi} \mathbf{H}_{\hat{\theta}}(\mathcal{M}) \right|$
- 2) Posterior: $q(\theta) = \mathcal{N}(\theta; \hat{\theta}, \mathbf{H}_{\hat{\theta}}^{-1})$
- 3) Predictive: $q(y_* | \mathbf{x}_*, \mathcal{D}) = \int p(y_* | \mathbf{x}_*, \theta) \mathcal{N}(\theta; \hat{\theta}, \mathbf{H}_{\hat{\theta}}^{-1}) d\theta$

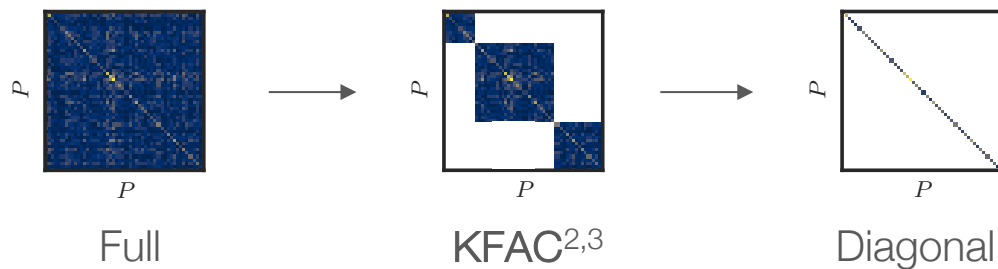
But: Hessian in $O(P^2)$ not tractable for Deep Learning

¹Laplace. *Mémoire sur les probabilités*. Mémoires de l'Académie royale des sciences de Paris 1778.

²MacKay. *A practical Bayesian framework for backpropagation networks*. Neural Computation 1992.

Efficient Hessian Approximations

Replace full $O(P^2)$ by structured positive-definite approximations $O(P)^{1,2}$



→ Tractable for extremely large models

¹Ritter, Botev, Barber. *A Scalable Laplace Approximation for Neural Networks*. ICLR 2018.

²Martens, Grosse. *Optimizing Neural Networks with Kronecker-factored Approximate Curvature*. ICML 2015.

Efficient Hessian Approximations

- 1 Generalized Gauss-Newton (GGN) instead of Hessian:

$$\nabla_{\theta}^2 \log p(\mathcal{D}|\theta) = \nabla_{\theta}^2 \sum_n \log p(y_n | f(x_n, \theta))$$

$$\approx \sum_n \underbrace{\mathbf{J}_{\theta}(x_n)}_{\text{NN Jacobian } O(PK)}^{\top} \underbrace{\nabla_f^2 \log p(y_n | f_n)}_{\text{Log loss Hessian w.r.t. NN output function } O(K^2)} \underbrace{\mathbf{J}_{\theta}(x_n)}_{\text{NN Jacobian } O(PK)} + \cancel{\nabla_{\theta}^2 f(x_n, \theta)^{\top} \nabla_f \log p(y_n | f_n)}$$

NN Jacobian $O(PK)$

Log loss Hessian w.r.t.
NN output function $O(K^2)$

→ Positive-semidefinite and reduced cost

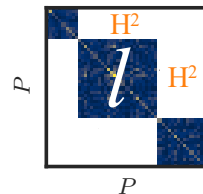
Efficient Hessian Approximations

② Decompositions and sampling for Jacobian-vector products:

$$\begin{aligned}
 & \mathbf{J}_{\theta}(\mathbf{x})^{\top} \underbrace{\nabla_{\mathbf{f}}^2 \log p(\mathbf{y}|\mathbf{f})}_{\text{Hessian}} \mathbf{J}_{\theta}(\mathbf{x}) \\
 &= \sum_k \mathbf{v}_k \mathbf{v}_k^T && \text{GGN} && O(PK) \\
 &= \mathbb{E}_{p(\mathbf{y}|\mathbf{f})} [\nabla_{\mathbf{f}} \log p(\mathbf{y}|\mathbf{f})^2] && \text{Fisher} && O(PS) \quad S \geq 1 \\
 &\approx \nabla_{\mathbf{f}} \log p(\mathbf{y}|\mathbf{f})^2 && \text{Empirical Fisher} && O(P)
 \end{aligned}$$

→ Reduce cost from $O(PK)$ to $O(P)$ and not require Jacobian

Efficient Hessian Approximations



- ③ KFAC¹ is a layer-wise block-diagonal structured approximation:

$$[\sum_n \mathbf{J}_\theta(\mathbf{x}_n)^\top \mathbf{\Lambda}(\mathbf{f}_n) \mathbf{J}_\theta(\mathbf{x}_n)]_l = \sum_n [\mathbf{a}_{l,n} \otimes \mathbf{g}_{l,n}] \mathbf{\Lambda}(\mathbf{f}_n) [\mathbf{a}_{l,n} \otimes \mathbf{g}_{l,n}]^\top$$

Factors vary with layer type²

$$= \underbrace{\sum_n [\mathbf{a}_{l,n} \mathbf{a}_{l,n}^\top]}_{O(H^4)} \otimes \underbrace{[\sum_n \mathbf{g}_{l,n} \mathbf{\Lambda}(\mathbf{f}_n) \mathbf{g}_{l,n}^\top]}_{O(H^2)} \approx \frac{1}{N} \underbrace{[\sum_n \mathbf{a}_{l,n} \mathbf{a}_{l,n}^\top]}_{O(H^2)} \otimes \underbrace{[\sum_n \mathbf{g}_{l,n} \mathbf{\Lambda}(\mathbf{f}_n) \mathbf{g}_{l,n}^\top]}_{O(H^2)}$$

→ Reduce storage $O(P^2)$ to $O(P)$ despite having off-diagonal entries!

¹Martens, Grosse. *Optimizing Neural Networks with Kronecker-factored Approximate Curvature*. ICML 2015.

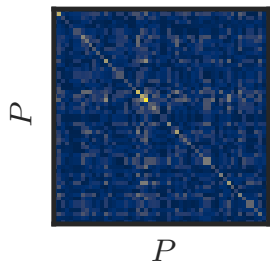
²Eschenhagen, Immer, Turner, Schneider, Hennig. *Kronecker-Factored Approximate Curvature for Modern Neural Network Architectures*. NeurIPS 2023.

Part I: Bayesian Model Selection for Deep Learning

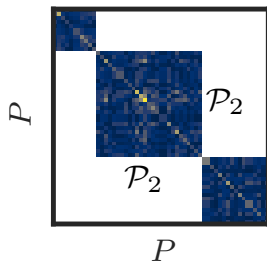
Structured Laplace Approximations

Structured Hessian approximations yield lower bounds to the Laplace marglik¹:

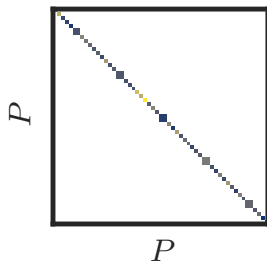
$$\mathcal{L} - \frac{1}{2} \log |\mathbf{H}| \geq \mathcal{L} - \frac{1}{2} \log |\mathbf{H}_{\mathcal{P}}| \geq \mathcal{L} - \frac{1}{2} \log |\mathbf{H}_{\mathcal{P}'}|$$



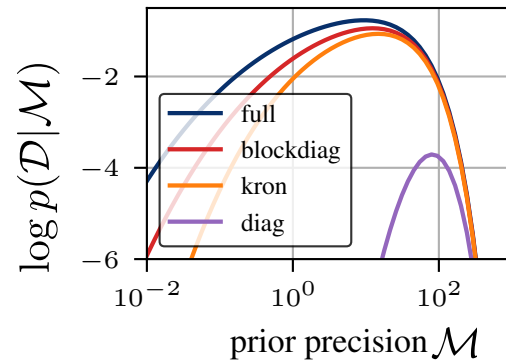
Full



Block-Diag
~KFAC



Diag



→ Theoretical justification for approximations

Practical Bayesian Model Selection with Laplace

Online Laplace approximation every F steps during training¹:

1) Gradient-based optimization of differentiable $\mathcal{M}$

$$\mathcal{M} \leftarrow \mathcal{M} + \gamma \nabla_{\mathcal{M}} \log q(\mathcal{D}|\mathcal{M})$$

Gradient-based Selection

→ Can optimize high-dimensional hyperparameters without re-training

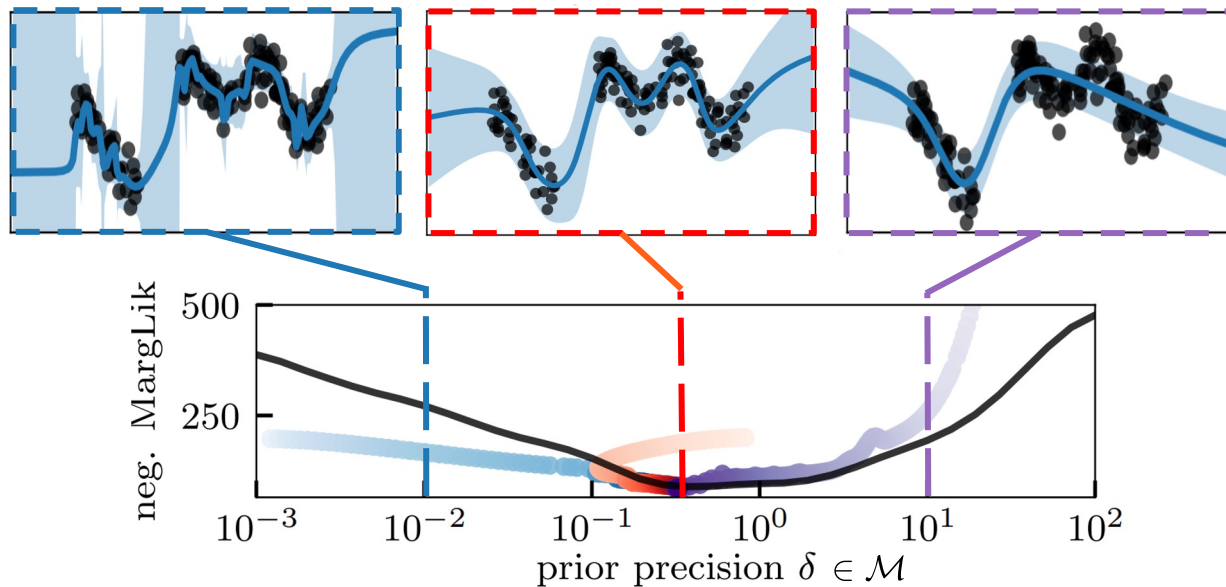
2) Comparison of models and checkpoints

$$\log q(\mathcal{D}|\mathcal{M})$$

Discrete Selection

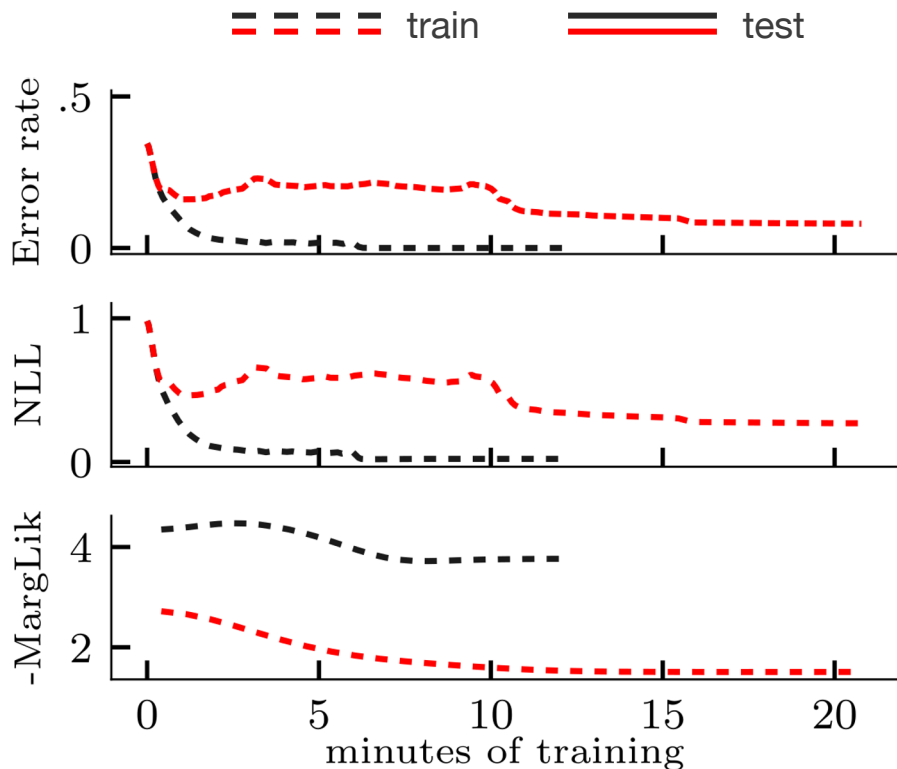
Example: Gradient-Based Model Selection

Optimize prior precision (regularization) with parameters in one training run



→ Validation-based selection requires many iterations

CIFAR-10 Classification



$\mathcal{M}$ = regularization per layer

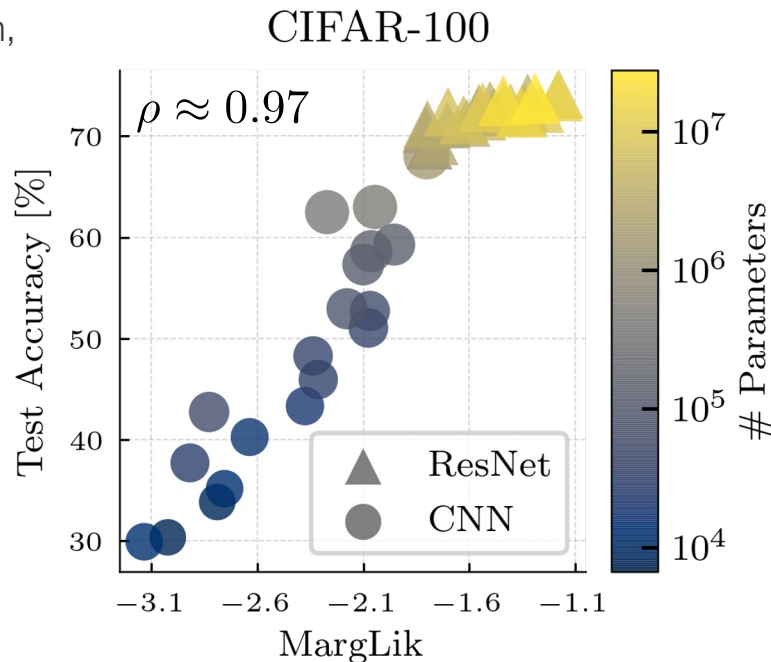
Compare train performance:
Standard vs **MargLik**

Compare test performance:
→ **MargLik** generalizes better

→ Optimize thousands of regularization parameters in only 2x runtime

Marginal Likelihood for Architecture Comparison

$\mathcal{M}$ = layer-wise prior precision,
architecture



→ Discrete selection possible after optimizing the prior

Examples of Bayesian Model Selection

- Regularization: weight decay/prior precision

Immer, Bauer, Fortuin, Rätsch, Khan. *Scalable Bayesian Model Selection for Deep Learning*. ICML 2021.

Antoran, Barbano, ..., Jin. *Uncertainty Estimation for Computed Tomography with a Linearised Deep Image Prior*. TMLR 2023.

Dhahri, Immer, Charpentier, Günnemann, Fortuin. *Shaving Weights with Occam's Razor*. Preprint 2024.

- Selecting architecture, structure, variables

Zhou*, Yang*, Wang, Pan. *BayesNAS: A Bayesian Approach for Neural Architecture Search*. CVPR 2019.

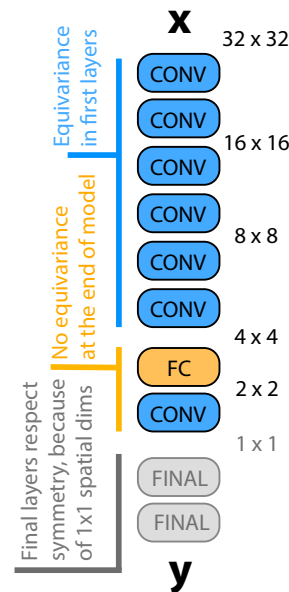
van der Ouderaa, Immer, van der Wilk. *Learning Layer-wise Equivariances Automatically using Gradients*. NeurIPS 2023.

Bouchiat, Immer, Yeché, Rätsch, Fortuin. *Improving Neural Additive Models with Bayesian Principles*. ICML 2024.

- Learning invariances/data augmentation

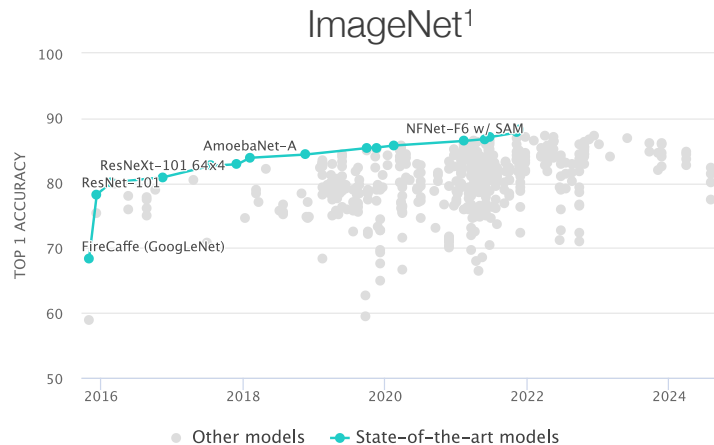
van der Wilk, Bauer, John, Hensman. *Learning Invariances using the Marginal Likelihood*. NeurIPS 2018.

Immer*, van der Ouderaa*, Fortuin, Rätsch, van der Wilk. *Invariance Learning in Deep Neural Networks [...]*. NeurIPS 2022.



Limitations of Bayesian Model Selection

- Required probabilistic model
- Can still require manual tuning or validation
- No benefit on saturated benchmarks

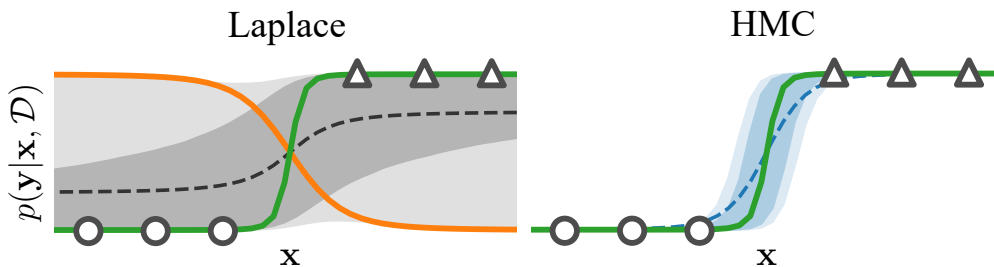


¹<https://paperswithcode.com/sota/image-classification-on-imagenet> accessed on 18th Sept 2024.

Part 2: Predictive Uncertainties

Laplace Posterior Predictive Underfitting

$$q(y_* | \mathbf{x}_*, \mathcal{D}) = \int p(y_* | \mathbf{x}_*, \theta) \mathcal{N}(\theta; \hat{\theta}, \mathbf{H}_{\hat{\theta}}^{-1}) d\theta$$



Problem: Laplace posterior predictive underfits¹

¹Lawrence. *Variational Inference in Probabilistic Models*. PhD Thesis 2001.

Linearized Laplace

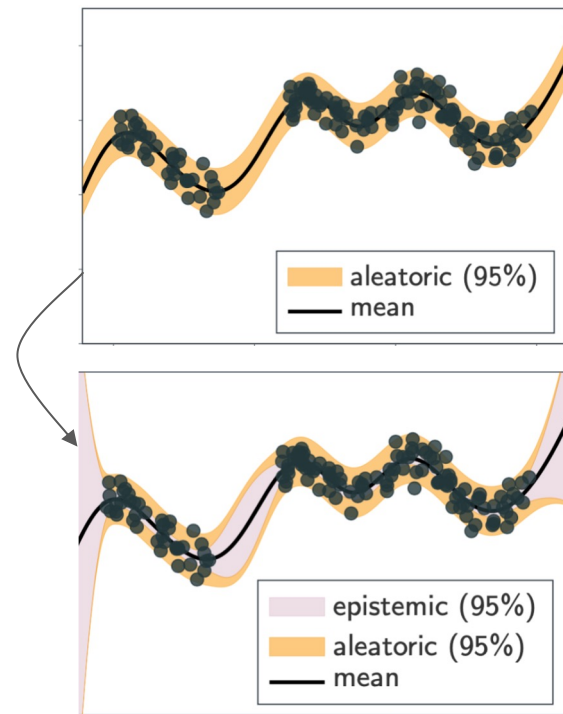
Hessian approximations linearize implicitly²

$$f_{\hat{\theta}}^{\text{lin}}(\mathbf{x}, \theta) = f(\mathbf{x}, \hat{\theta}) + \mathbf{J}_{\hat{\theta}}(\mathbf{x})(\theta - \hat{\theta})$$

We should keep this in mind for predictions:

$$q(y_* | \mathbf{x}_*, \mathcal{D}) \approx \int p(y_* | f_{\hat{\theta}}^{\text{lin}}(\mathbf{x}_*, \theta)) \mathcal{N}(\theta; \hat{\theta}, \mathbf{H}_{\hat{\theta}}^{-1}) d\theta$$

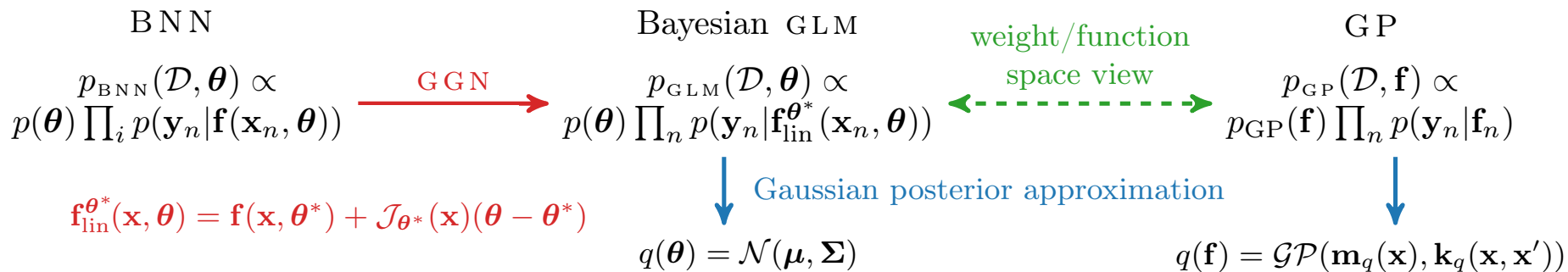
→ Stabilizes and improves Laplace predictive



¹Lawrence. *Variational Inference in Probabilistic Models*. PhD Thesis 2001.

²Bottou, Curtis, Nocedal. *Optimization methods for large-scale machine learning*. SIAM Review 2018.

Linearized Model as a Gaussian Process



→ Enables different posterior approximation techniques

Laplace Posterior Predictive

- Last-layer is very cheap and effective baseline

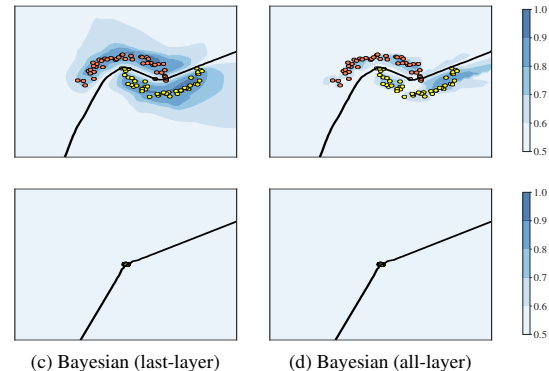
Kristiadi, Hein, Hennig. *Being Bayesian, even just a bit, fixes overconfidence in relu networks*. ICML 2020.
Ober, Rasmussen. *Benchmarking the neural linear model for regression*. AABI 2019.

- Laplace can improve ensembling further

Eschenhagen, Daxberger, Hennig, Kristiadi. *Mixtures of Laplace Approximations [...]*. BDL@NeurIPS 2021.

- Linearized predictive can have a closed-form (approximation)

Immer, Korzepa, Bauer. *Improving predictions of Bayesian neural nets via local linearization*. AISTATS 2021.
Deng, Zhou, Zhu. *Accelerated Linearized Laplace Approximation for Bayesian Deep Learning*. NeurIPS 2022.



Limitations

- Scalability issues in some cases¹
- Linearization can be expensive²
- Post-hoc cannot fix badly trained network

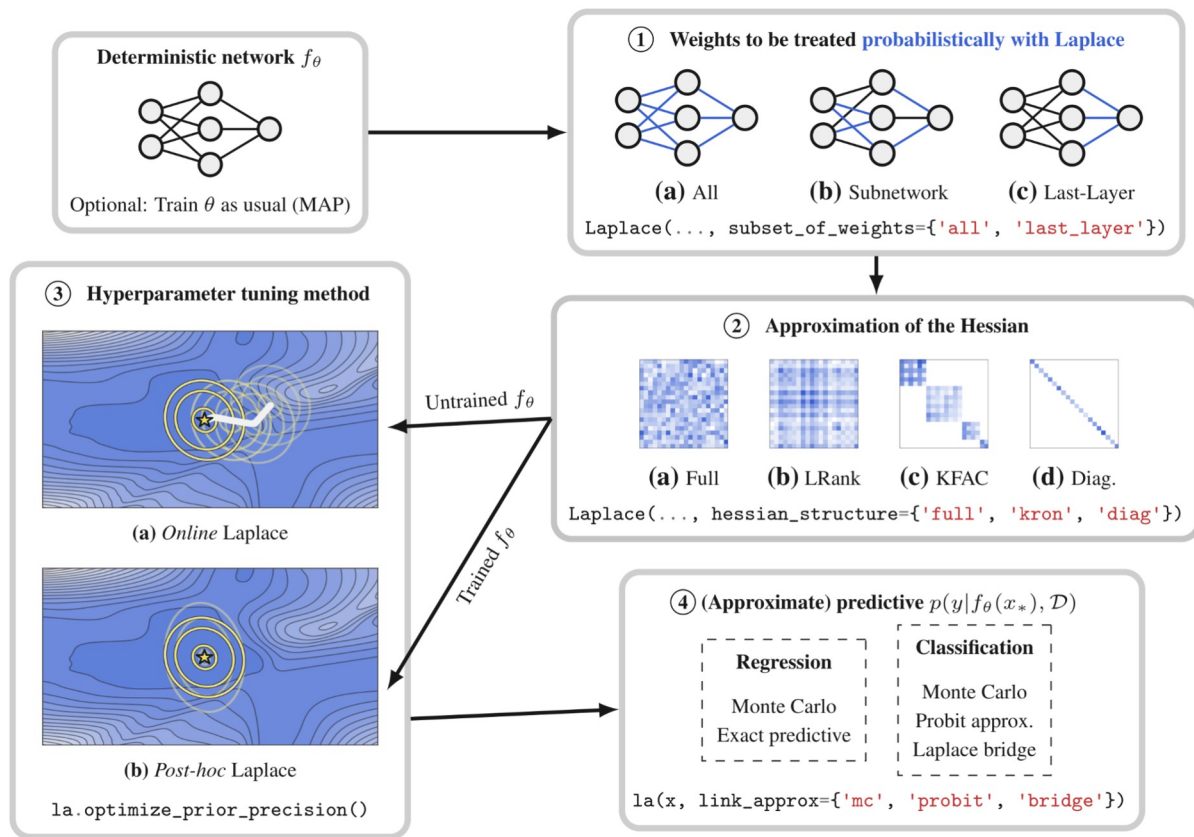
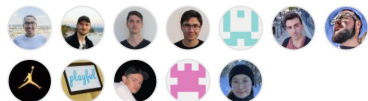
¹Kepf, Wanna, Miani, Moore, ..., Warburg. *Laplacian Segmentation Networks [...]*. MICCAI 2024.

²Antoran, Janz, Allingham, ..., Hernandez-Lobato. *Adapting the Linearised Laplace Model Evidence for Modern Deep Learning*. ICML 2022.

laplace-torch¹ on github.com/AlexImmer/Laplace

Fork 70 Star 447

Contributors 12



¹Daxberger*, Kristiadi*, Immer*, Eschenhagen*, Bauer, Hennig. *Laplace Redux - Effortless Bayesian Deep Learning*. NeurIPS 2021.

Thank you!

Co-Authors

Matthias Bauer

Kouroche Bouchiat

Peter Bühlmann

Bertrand Charpentier

Ryan Cotterell

Erik Daxberger

Rayen Dhahri

Runa Eschenhagen

Vincent Fortuin

Stephan Günnemann

Philipp Hennig

Lucas Torroba Hennigen

Francis Jacob

Mohammad Emtiyaz Khan

Maciej Korzepa

Agustinus Kristiadi

Kjong-Van Lehmann

Alexander Marx

Emanuele Palumbo

Gunnar Rätsch

Frank Schneider

Bernhard Schölkopf

Christoph Schultheiss

Stefan Stark

Richard Turner

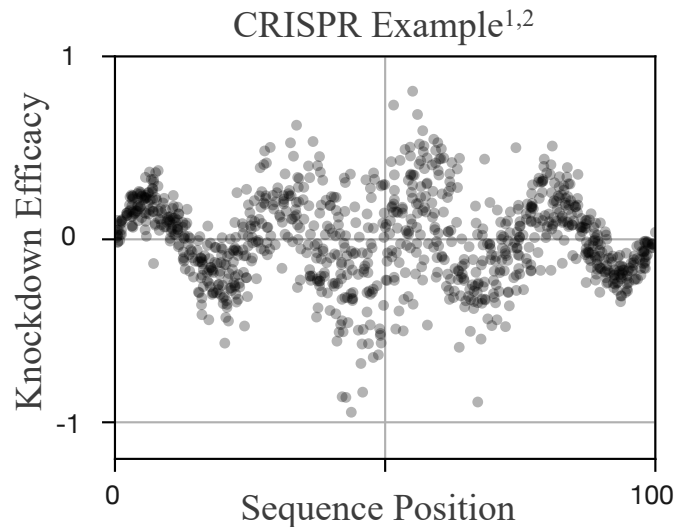
Tycho van der Ouderaa

Mark van der Wilk

Julia Vogt

Hugo Yèche

Estimating Heteroscedastic Aleatoric Uncertainty



Model mean $\mu(\mathbf{x})$ and variance $\sigma^2(\mathbf{x})$ of a Gaussian likelihood:

$$\log p(y|\mathbf{x}, \mu, \sigma^2) \propto -\frac{1}{2} \log \sigma^2(\mathbf{x}) - \frac{(y - \mu(\mathbf{x}))^2}{2\sigma^2(\mathbf{x})}$$

Objective balances mean and variance

→ Hard to optimize and regularize^{2,3}

¹Wessels, Mendez-Mancilla, ..., Sanjana. *Massively parallel Cas13 screens reveal principles for guide RNA design*. Nature Biotechnology 2020.

²Stirn, Wessels, ..., Knowles. *Faithful Heteroscedastic Regression with Neural Networks*. AISTATS 2023.

³Seitzer, Tavakoli, Antic, Martius. *On the Pitfalls of Heteroscedastic Uncertainty Estimation with Probabilistic Neural Networks*. ICLR 2022.

Problem with Mean-Variance Parameterization

Mean-variance parameterization does not yield concave log likelihood:

$$\frac{\partial^2 \log \mathcal{N}(y|\mu, \sigma^2)}{\partial \sigma^2 \partial \sigma^2} = \frac{\sigma^2 - 2(\mu - y)^2}{2\sigma^6} \leftarrow \begin{array}{l} \text{Can be positive or negative.} \\ \rightarrow \text{potentially indefinite Hessian} \\ \rightarrow \text{not concave} \end{array}$$

→ Cannot apply Hessian or Laplace approximations naively

Revisiting the Natural Parameterization

Natural parameterization always yields a concave log likelihood^{1,2} with

$$\eta_1 = \frac{\mu}{\sigma^2} \quad \text{and} \quad \eta_2 = -\frac{1}{2\sigma^2}$$

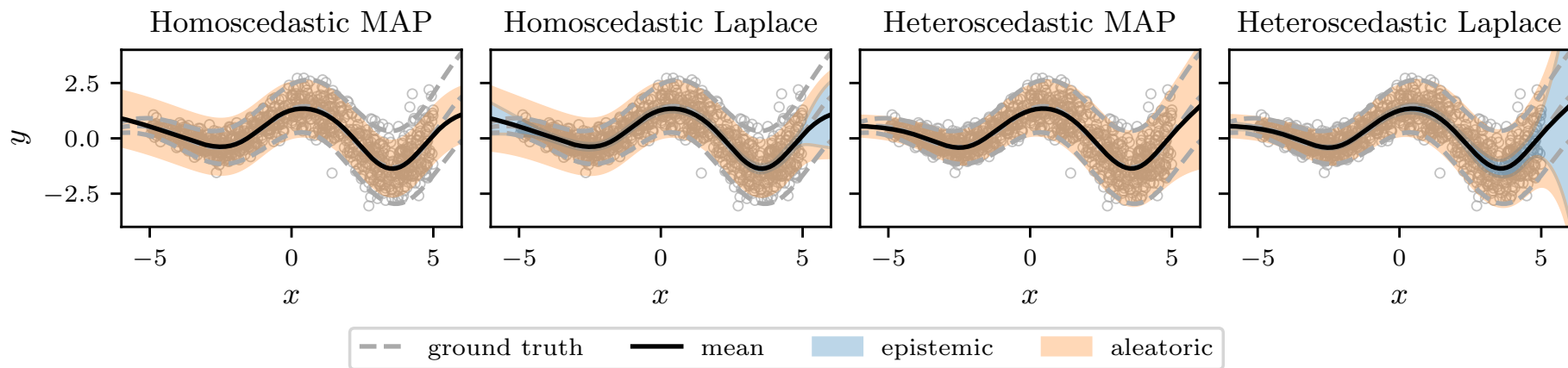
→ Hessian and Laplace approximations straightforward

→ Neural network models both natural parameters

¹Le, Smola, Canu. *Heteroscedastic Gaussian process regression*. ICML 2005.

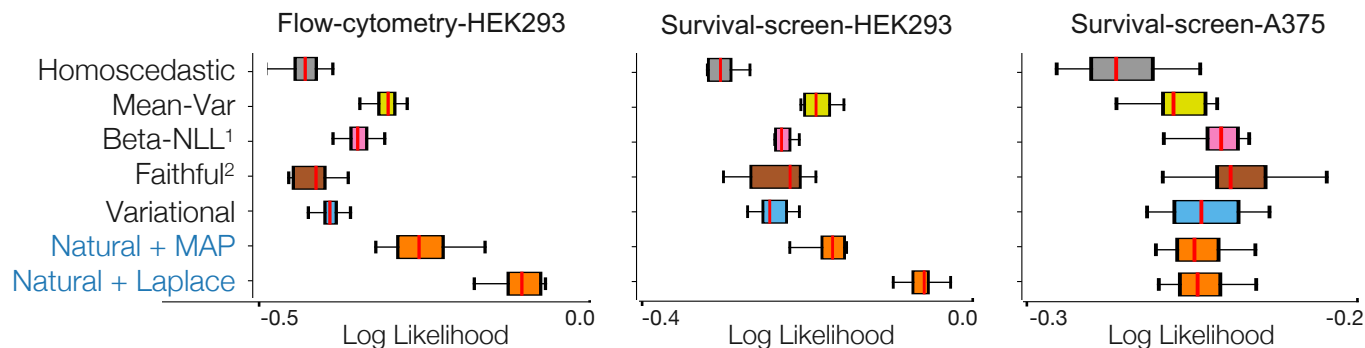
²Martens. *New insights and perspectives on the natural gradient method*. JMLR 2020.

Heteroscedastic Regression Illustration



Epistemic Uncertainty in Heteroscedastic Regression

Problem: predict gene knockdown efficacy of guides in CRISPR system²



→ Improvements due to both natural parameterization and Laplace

¹Seitzer, Tavakoli, Antic, Martius. *On the Pitfalls of Heteroscedastic Uncertainty Estimation with Probabilistic Neural Networks*. ICLR 2022.

²Stirn, Wessels, ..., Knowles. *Faithful Heteroscedastic Regression with Neural Networks*. AISTATS 2023.