

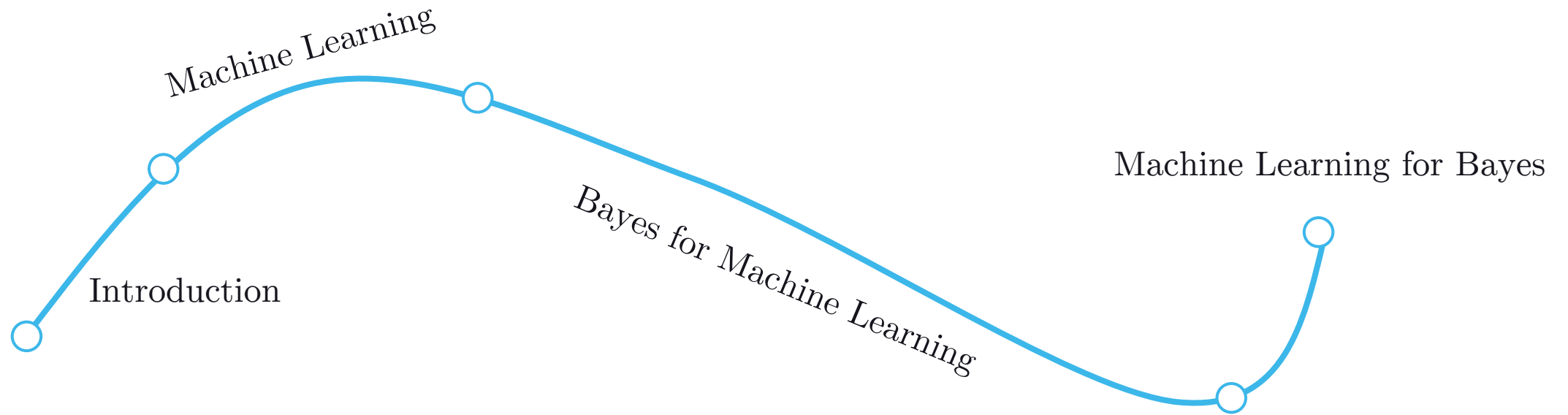
Machine Learning & Bayes

- Synergies and Challenges -

Martin Trapp
Aalto University

Learn Bayes Spring 2025 @ Karolinska Institutet

Outline



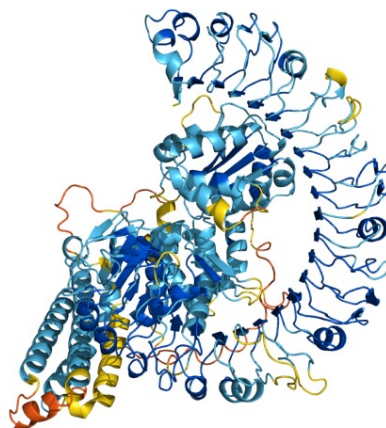
Success of Modern Machine Learning

Generative tasks



Image Generation

Predictive tasks



Protein Prediction ^[1]

“Reasoning” tasks

Chain-of-Thought Prompting

Model Input

Q: Roger has 5 tennis balls. He buys 2 more cans of tennis balls. Each can has 3 tennis balls. How many tennis balls does he have now?

A: Roger started with 5 balls. 2 cans of 3 tennis balls each is 6 tennis balls. $5 + 6 = 11$. The answer is 11.

Q: The cafeteria had 23 apples. If they used 20 to make lunch and bought 6 more, how many apples do they have?

Model Output

A: The cafeteria had 23 apples originally. They used 20 to make lunch. So they had $23 - 20 = 3$. They bought 6 more apples, so they have $3 + 6 = 9$. The answer is 9. ✓

Chain-of-Thought
in LLMs ^[2]

1: Jumper, J., et al. (2021). Highly accurate protein structure prediction with AlphaFold. Nature.
2: Wei, J., et al. (2022). Chain-of-thought prompting elicits reasoning in large language models. In NeurIPS.

Machine Learning in Healthcare



Personalised
treatment



Drug discovery
and development



ML-supported
diagnosis



ML-assisted
surgery



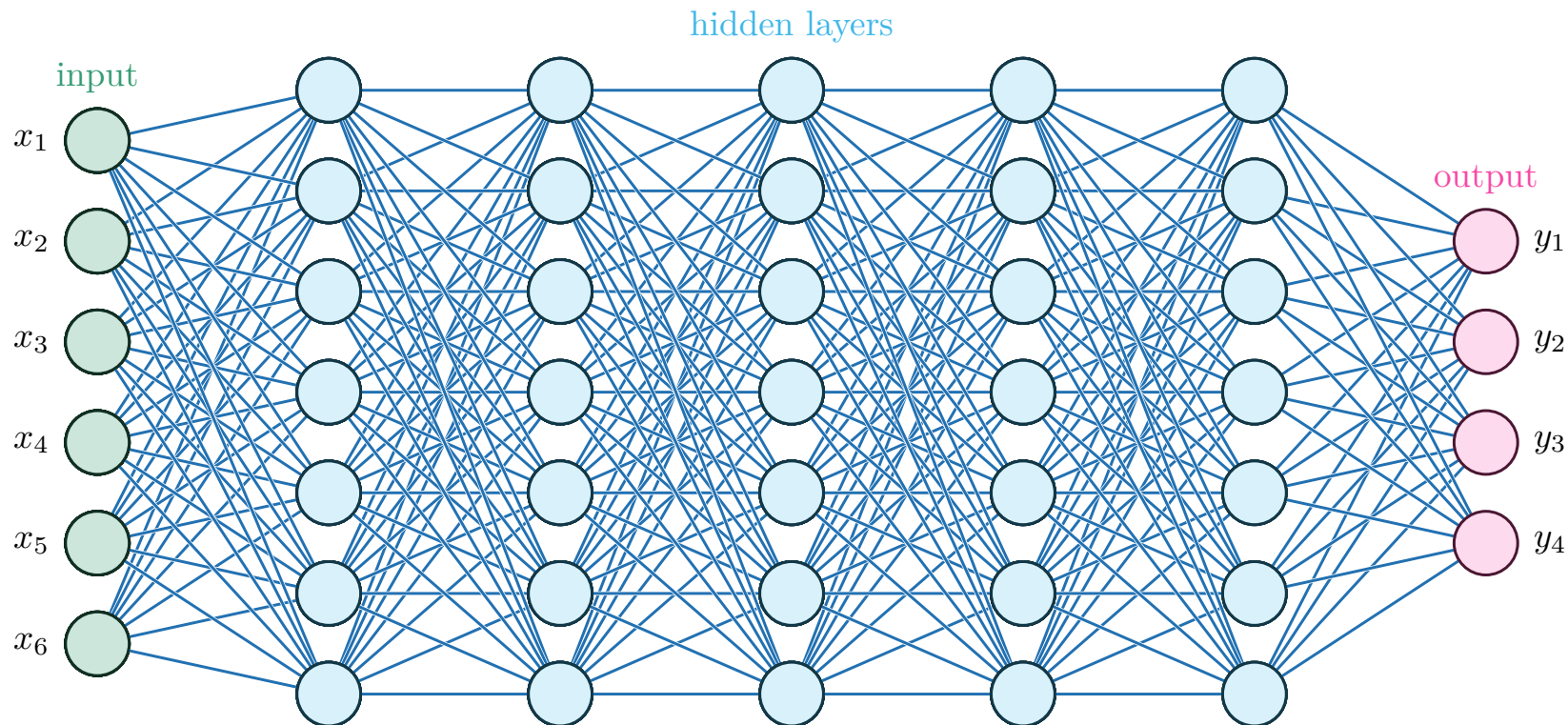
Automated image
analysis

Reminder on (modern) ML

... an incomplete and inaccurate picture of machine learning...

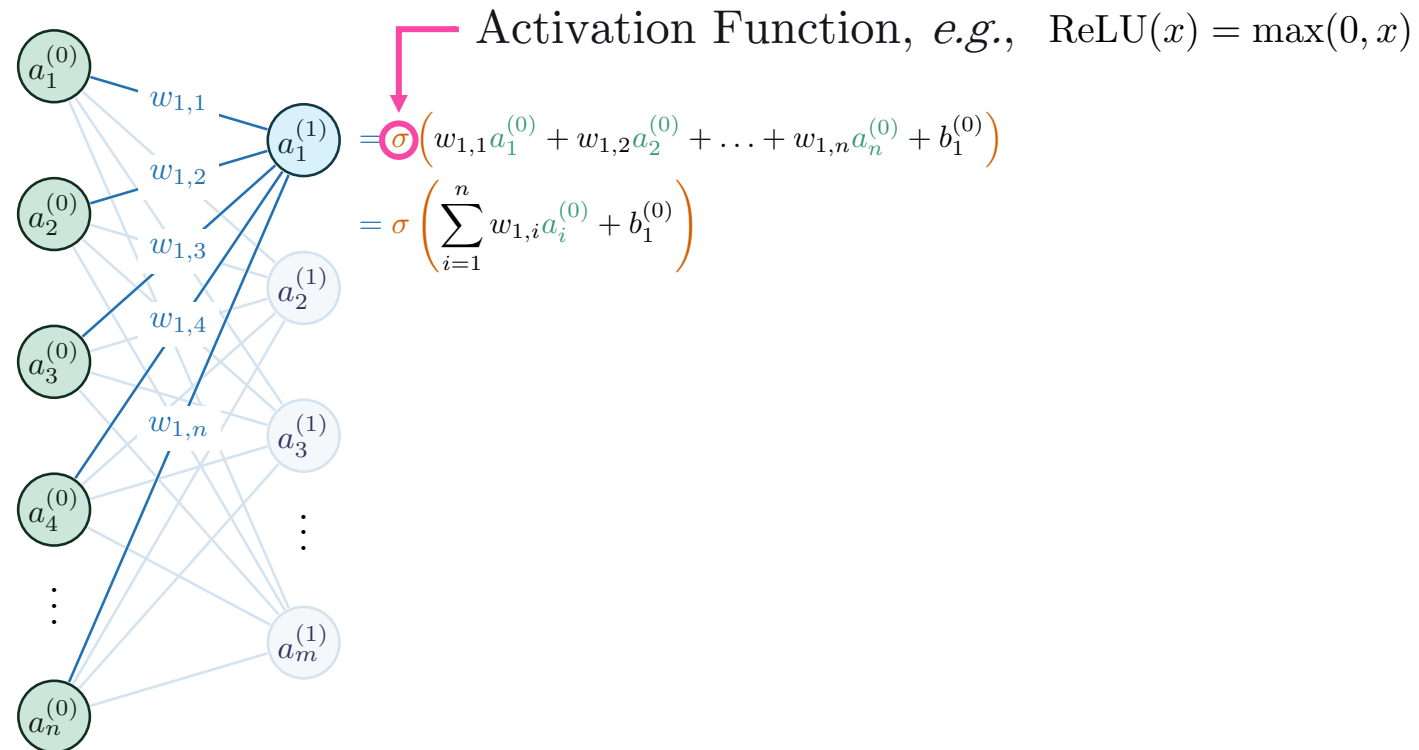
Reminder on Deep Learning

- The central objects in deep learning (modern ML) are **artificial neural networks or neural network architectures (NNs)**.



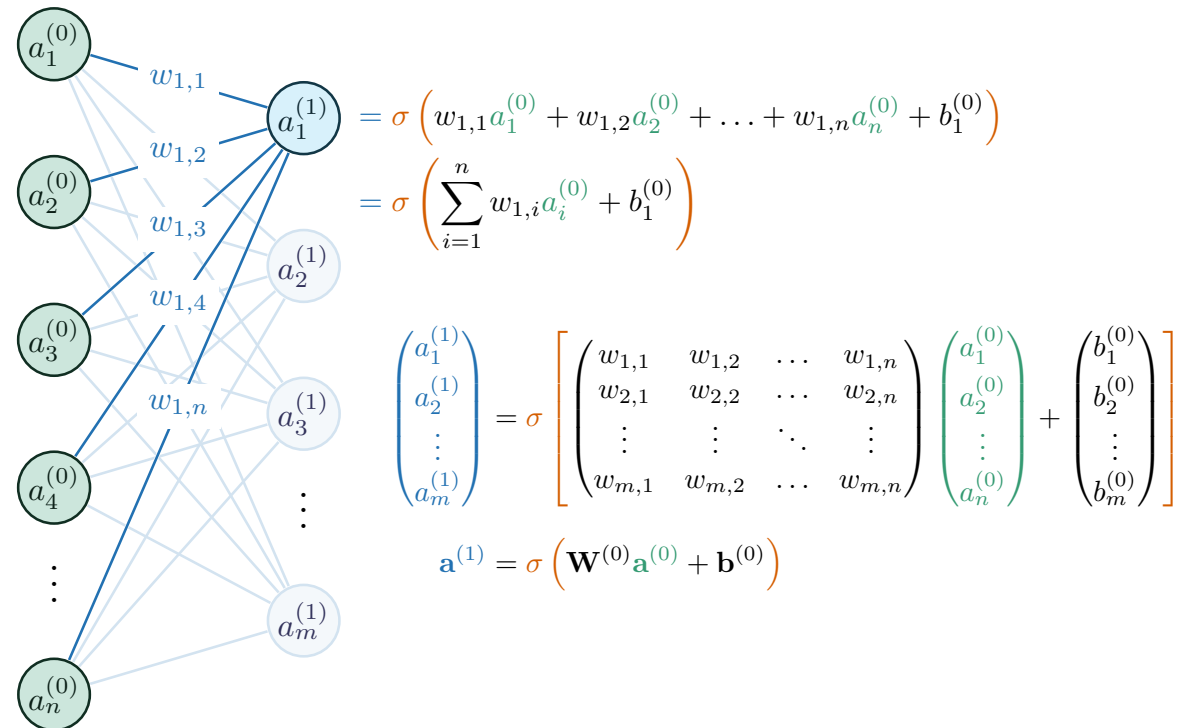
Reminder on Deep Learning

- The central objects in deep learning (modern ML) are **artificial neural networks** or **neural network architectures** (NNs).



Reminder on Deep Learning

- The central objects in deep learning (modern ML) are **artificial neural networks** or **neural network architectures** (NNs).



Reminder on Deep Learning

- The central objects in deep learning (modern ML) are **artificial neural networks or neural network architectures (NNs)**.
- NNs exploit function composition to define flexible non-linear function approximators with simple local operations.

$$f(\boldsymbol{x}) = f^{(L)} \circ \dots \circ f^{(2)} \circ f^{(1)}(\boldsymbol{x})$$

Reminder on Deep Learning

- The central objects in deep learning (modern ML) are **artificial neural networks or neural network architectures (NNs)**.
- NNs exploit function composition to define flexible non-linear function approximators with simple local operations.

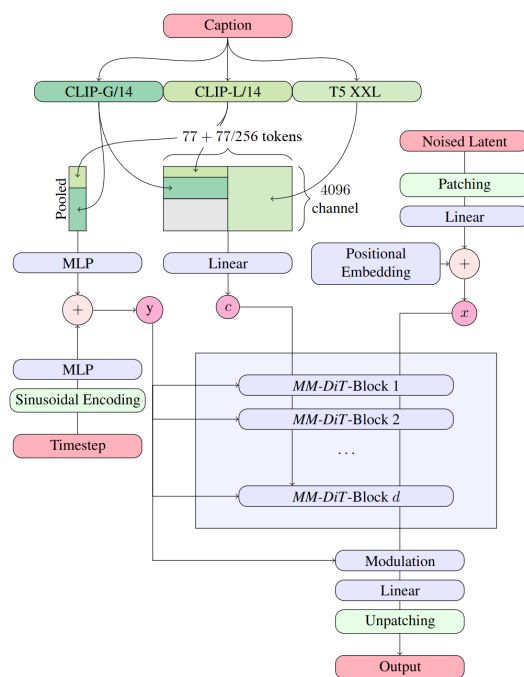
$$f(\boldsymbol{x}) = f^{(L)} \circ \dots \circ f^{(2)} \circ f^{(1)}(\boldsymbol{x})$$

- The network weights (parameters) are learned through backpropagation (chain rule) for a given loss function.

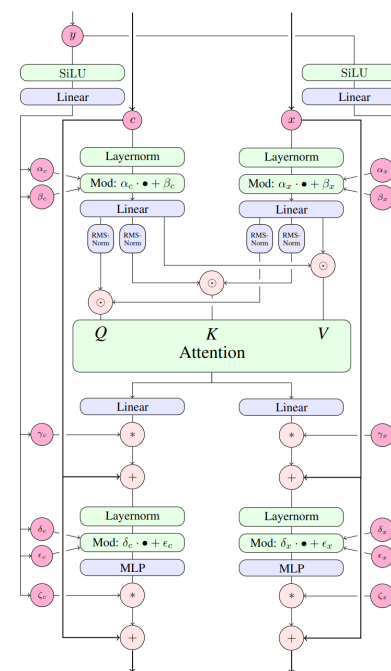


Reminder on Deep Learning

In practice, modern NNs are complex systems:



(a) Overview of all components.



(b) One *MM-DiT* block

Schematics of the Stable Diffusion 3.5 Model for Image Generation

Machine Learning in Healthcare



Personalised
treatment



Drug discovery
and development



ML-supported
diagnosis



ML-assisted
surgery



Automated image
analysis

Machine Learning in Healthcare

Many Applications are High-risk!



Personalised
treatment



Drug discovery
and development



ML-supported
diagnosis



ML-assisted
surgery



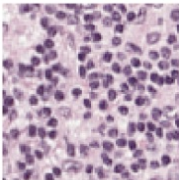
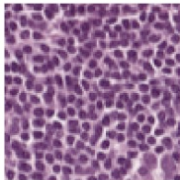
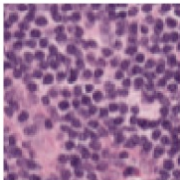
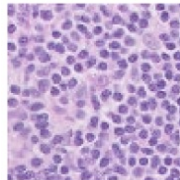
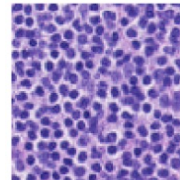
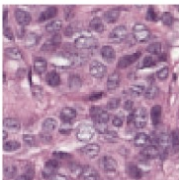
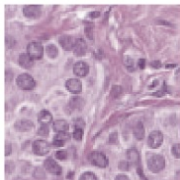
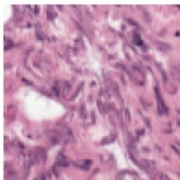
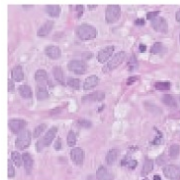
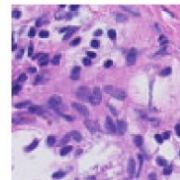
Automated image
analysis

The Real World is Messy



A World of Uncertainties

Distribution Shifts

Train			Val (OOD)	Test (OOD)	
	d = Hospital 1	d = Hospital 2	d = Hospital 3	d = Hospital 4	d = Hospital 5
y = Normal					
y = Tumor					

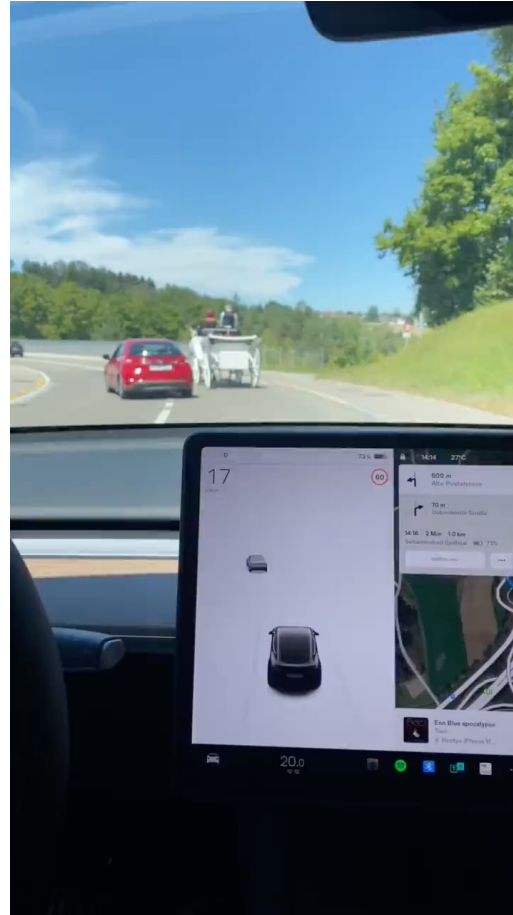
Koh, P W. et al. (2021) WILDS: A Benchmark of in-the-Wild Distribution Shifts. In *ICML*.

Various Sources of Noise



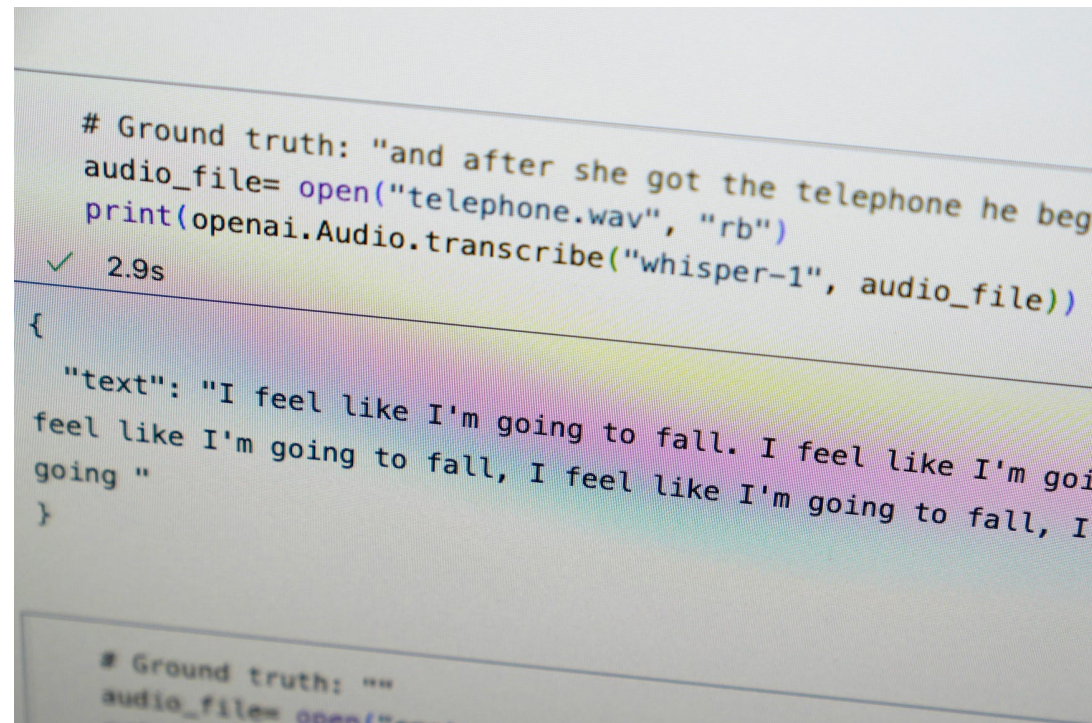
Sabour, S. et al. (2021) SpotlessSplats: Ignoring Distractors in 3D Gaussian Splatting. In *ICML*.

A World of Uncertainties



Out-of-domain data (Tesla)

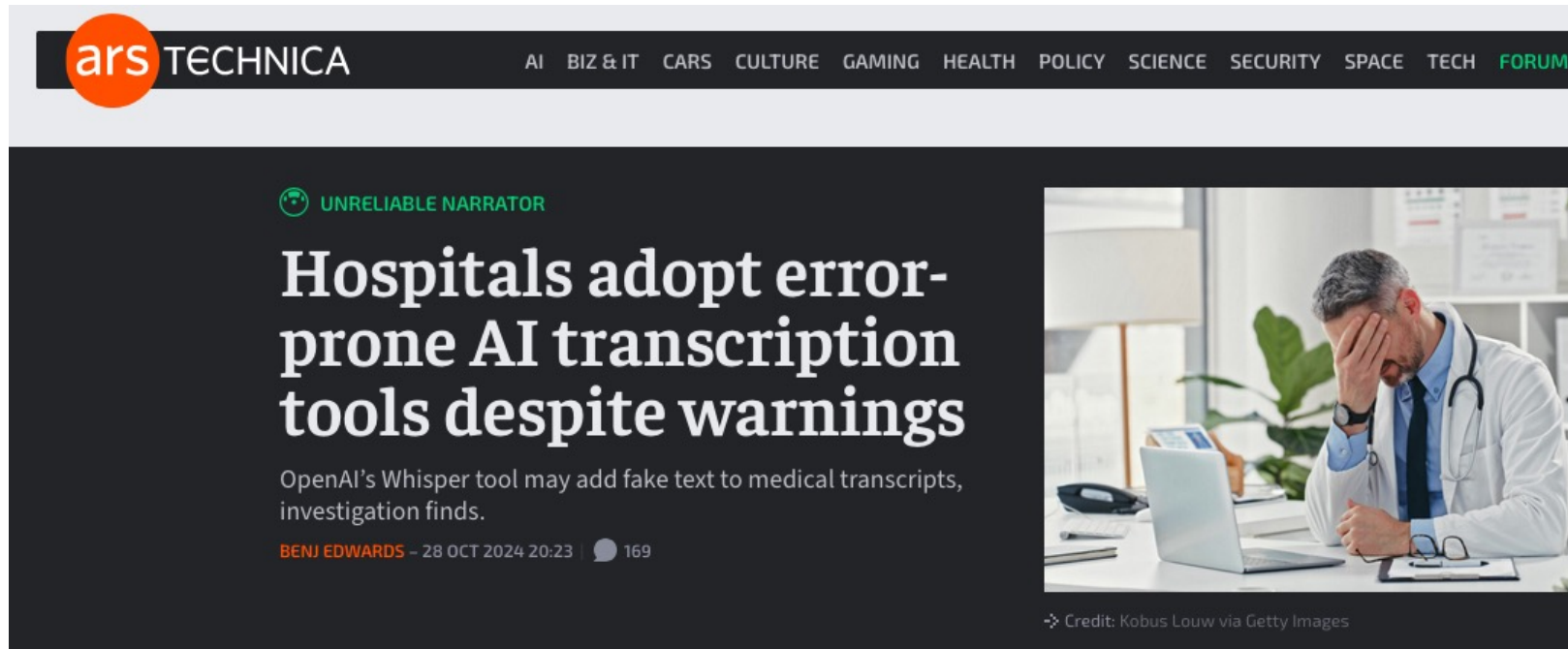
A World of Uncertainties

A screenshot of a code editor showing the output of an OpenAI Whisper transcription. The code includes a comment for the ground truth, file opening logic, and a transcription call. The output shows a checkmark and a duration of 2.9s, followed by a JSON object containing a text field with a hallucinated sentence. The text field contains the phrase "I feel like I'm going to fall" repeated multiple times.

```
# Ground truth: "and after she got the telephone he beg
audio_file= open("telephone.wav", "rb")
print(openai.Audio.transcribe("whisper-1", audio_file))
✓ 2.9s
{
  "text": "I feel like I'm going to fall. I feel like I'm goi
feel like I'm going to fall, I feel like I'm going to fall, I
going "
}
```

Hallucinations (OpenAI Whisper)

A World of Uncertainties



Expectations for ML Systems

Systems and models should be **safe, trustworthy and reliable**.

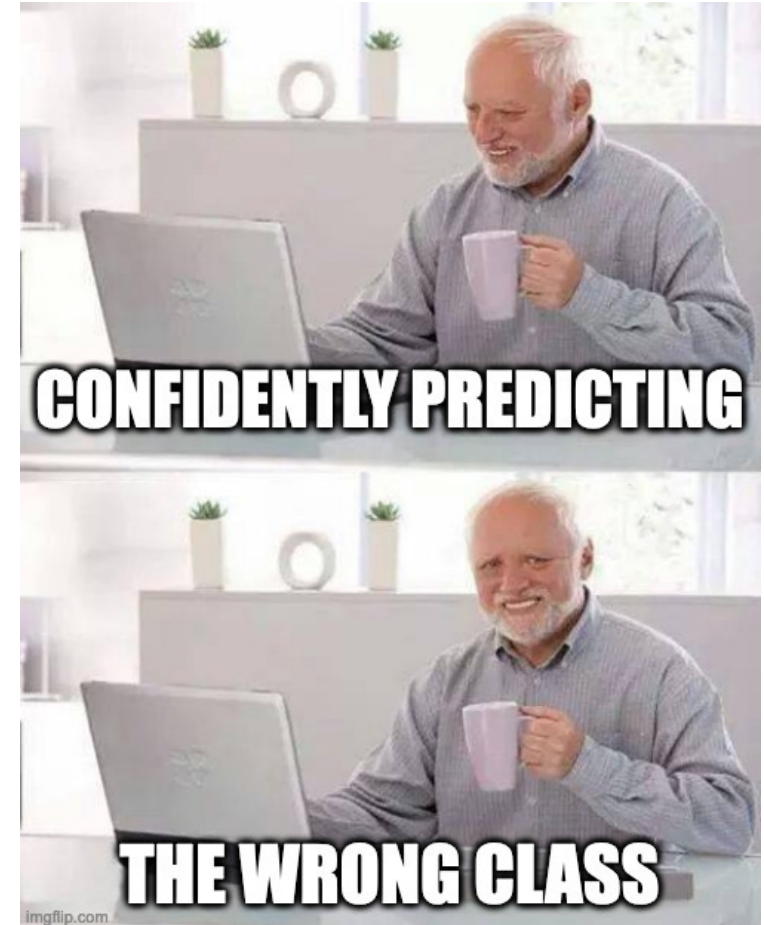
For this, we require that they are: (*incomplete list*)

- **Accurate** in their predictions
- **Robust** to input perturbations (noise, adversarial attacks)
- **"Know when they don't know"** (recognise out-of-domain data)
- **Act** if they are uncertain (*e.g.*, active learning, reasoning)

However, ...

- Accurate in their predictions
- Sensitive to perturbations^[1]
- Don't know when they don't know^[2]
- Overconfident in their predictions^[3]

Considered “unreliable”



1: Szegedy, C. et al. (2014). Intriguing properties of neural networks. In *ICLR*.

2: Nalisnick, E. et al. (2019). Do deep generative models know what they don't know? In *ICLR*.

3: Guo, C. et al. (2017). On calibration of modern neural networks. In *ICML*.

WELL, WHAT DO



WE DO NOW?

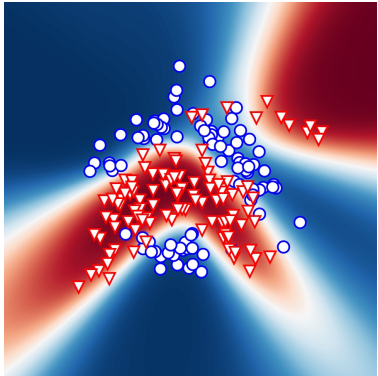
imgflip.com

Bayes for Machine Learning

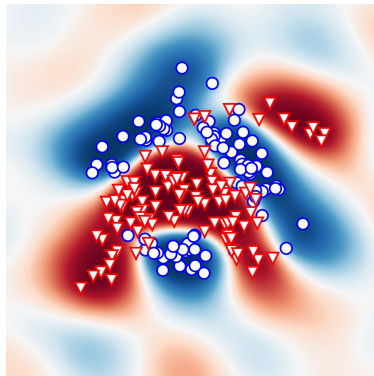
Bayesian Deep Learning

Functional Priors

Extrapolate



Conservative



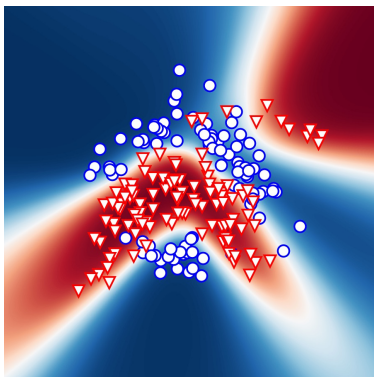
Meronen, L. et al. (2021). Periodic activation functions induce stationarity. In *NeurIPS*.

Bayesian Deep Learning

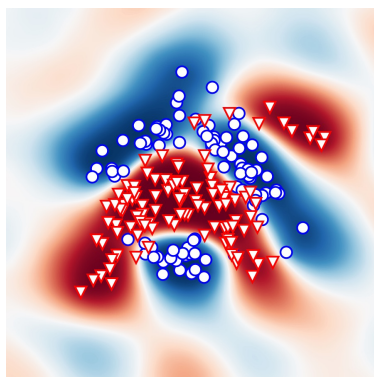
Functional Priors

Uncertainty Quant. &
Overconfidence

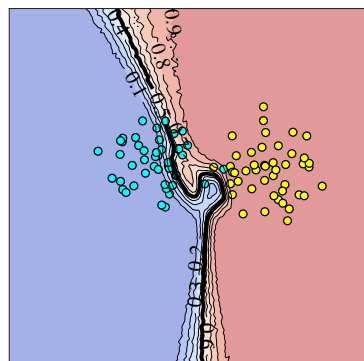
Extrapolate



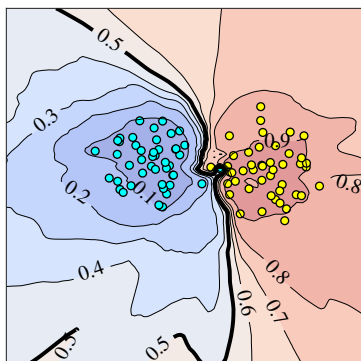
Conservative



Overconfident



Model Average



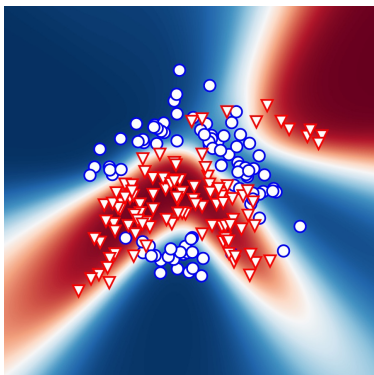
Meronen, L. et al. (2021). Periodic activation functions induce stationarity. In *NeurIPS*.

Kristiadi, A. et al. (2020). Being Bayesian, Even Just a Bit, Fixes Overconfidence in ReLU Networks. In *ICML*.

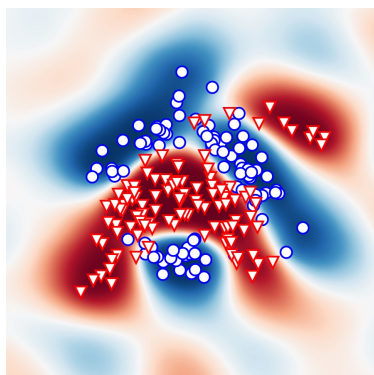
Bayesian Deep Learning

Functional Priors

Extrapolate



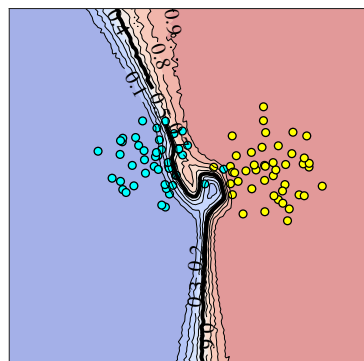
Conservative



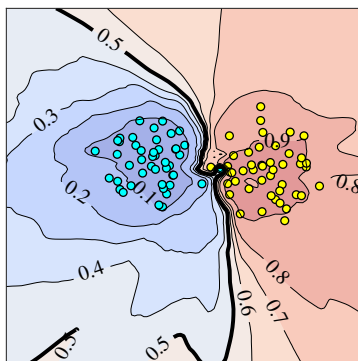
Meronen, L. et al. (2021). Periodic activation functions induce stationarity. In *NeurIPS*.

Uncertainty Quant. & Overconfidence

Overconfident

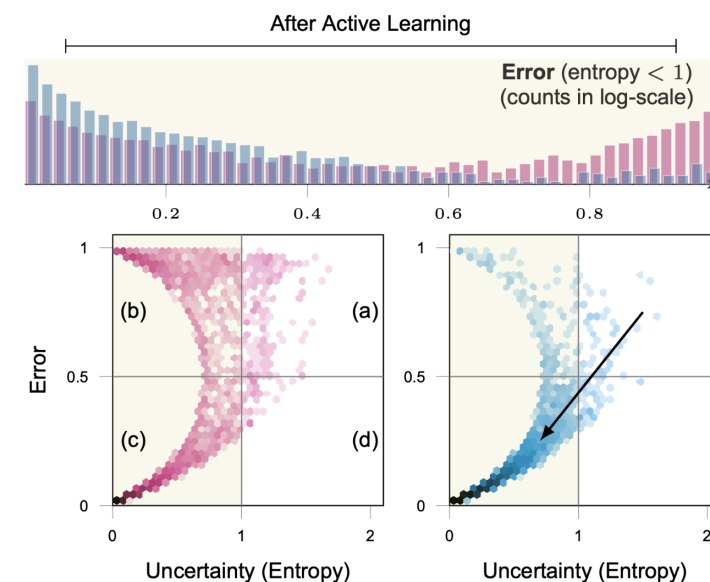


Model Average



Kristiadi, A. et al. (2020). Being Bayesian, Even Just a Bit, Fixes Overconfidence in ReLU Networks. In *ICML*.

Active Learning

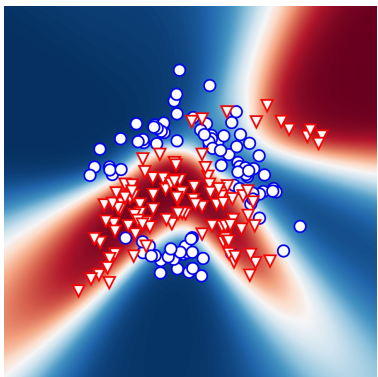


Baumann, A. et al. (2025). Post-hoc Probabilistic Vision-Language Models. ArXiv.

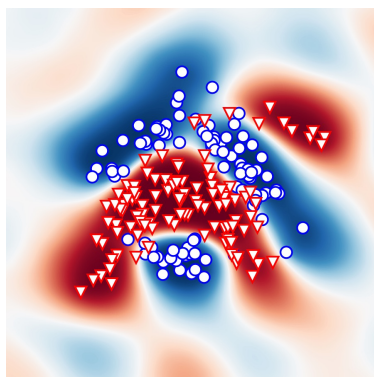
Bayesian Deep Learning

Functional Priors

Extrapolate



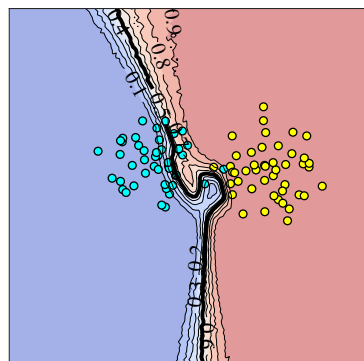
Conservative



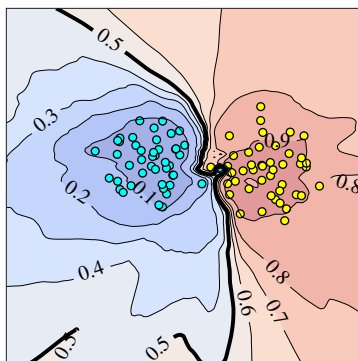
Meronen, L. et al. (2021). Periodic activation functions induce stationarity. In *NeurIPS*.

Uncertainty Quant. & Overconfidence

Overconfident

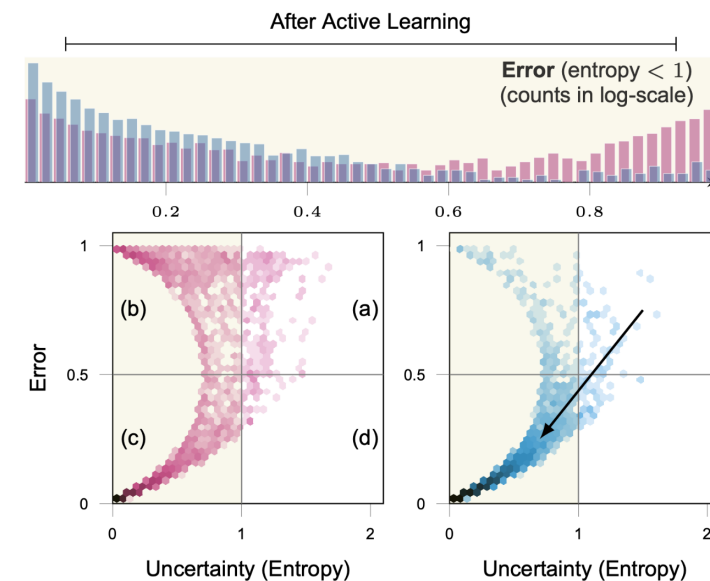


Model Average



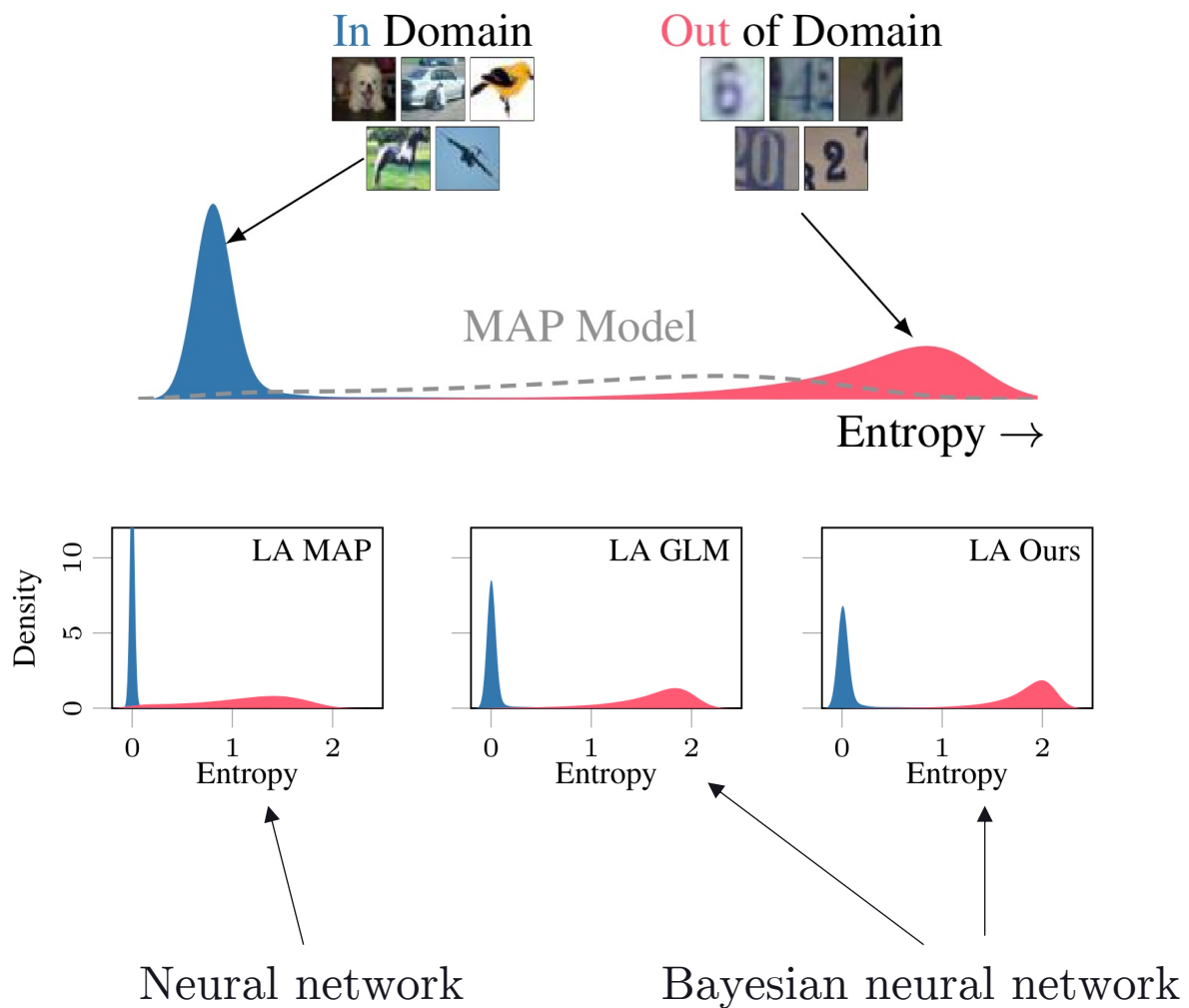
Kristiadi, A. et al. (2020). Being Bayesian, Even Just a Bit, Fixes Overconfidence in ReLU Networks. In *ICML*.

Active Learning

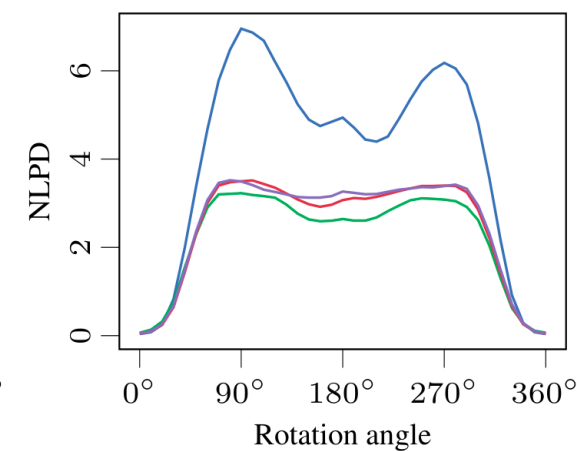
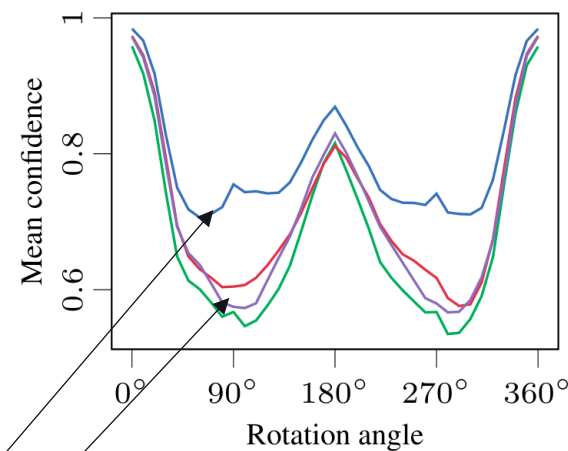
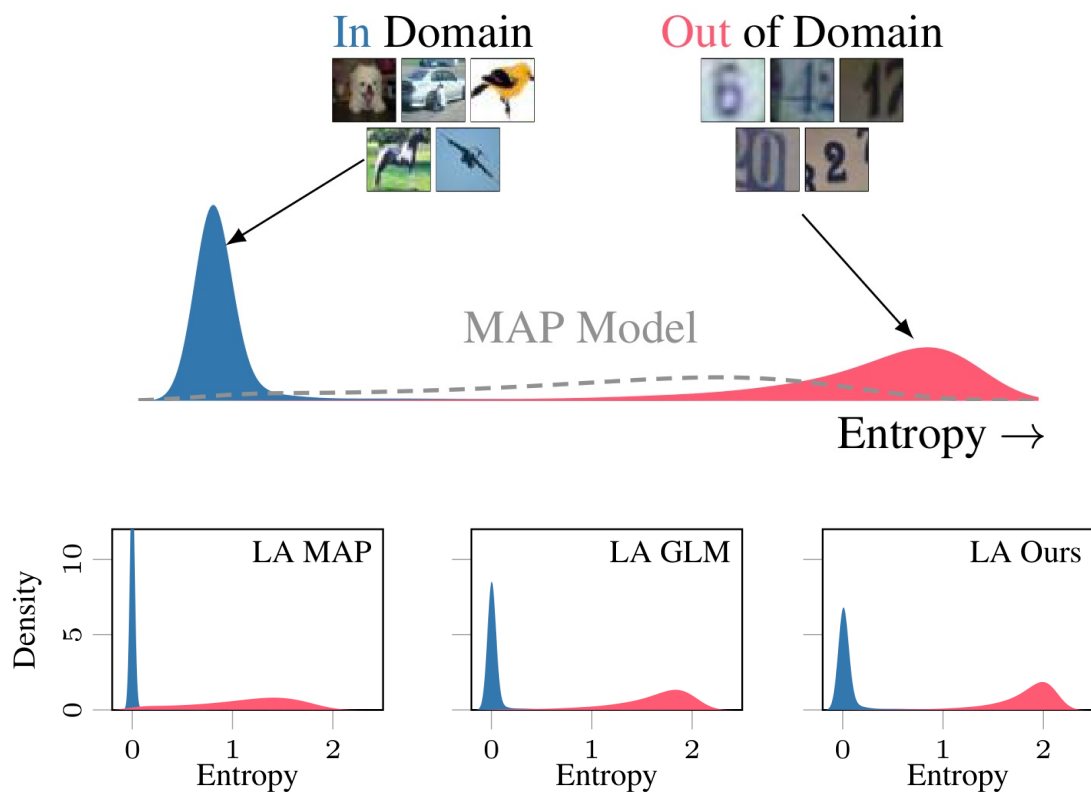


Baumann, A. et al. (2025). Post-hoc Probabilistic Vision-Language Models. ArXiv.

Bayesian Deep Learning



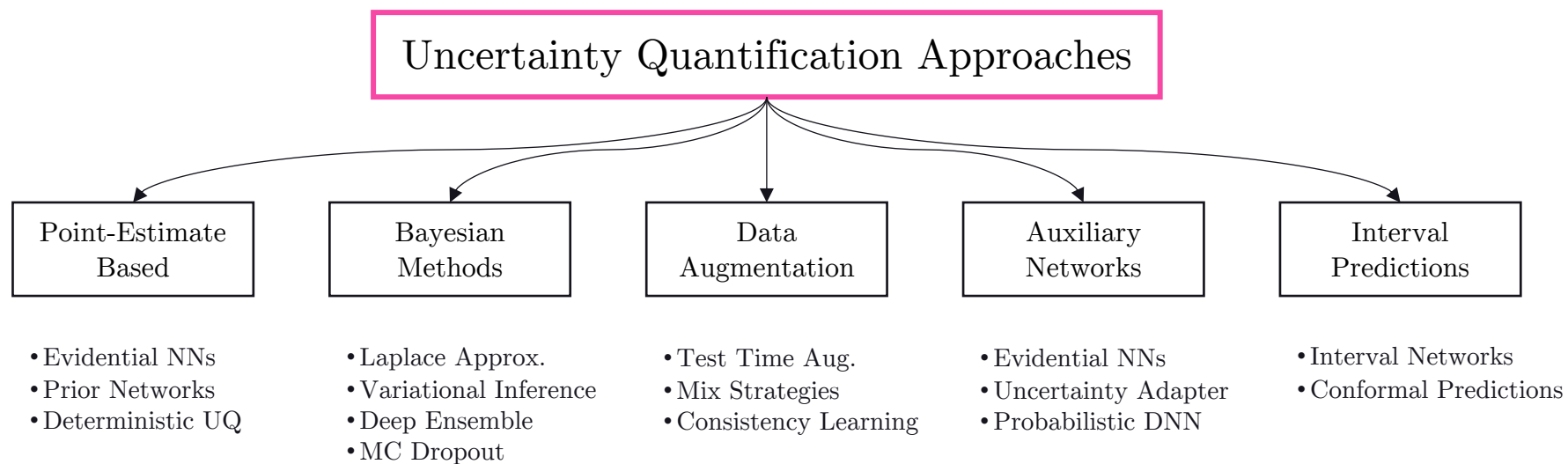
Bayesian Deep Learning



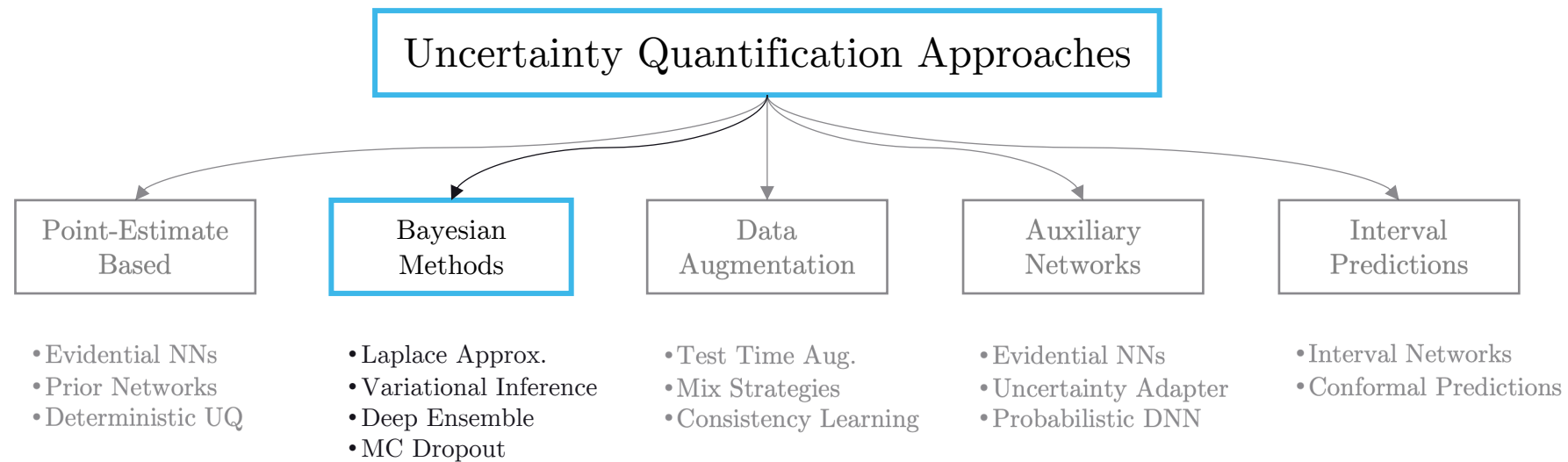
Neural network

Bayesian neural network

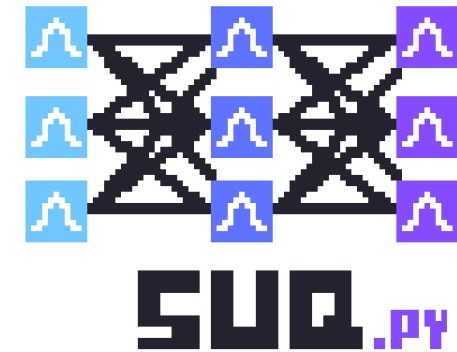
Uncertainty Quantification in ML



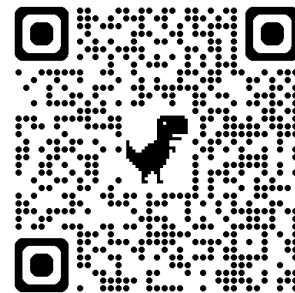
Uncertainty Quantification in ML



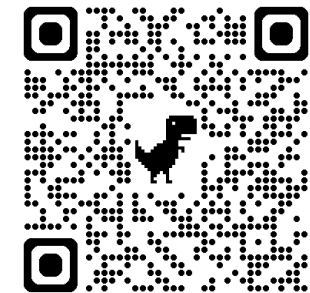
Uncertainty Quantification in ML



<https://torch-uncertainty.github.io/>



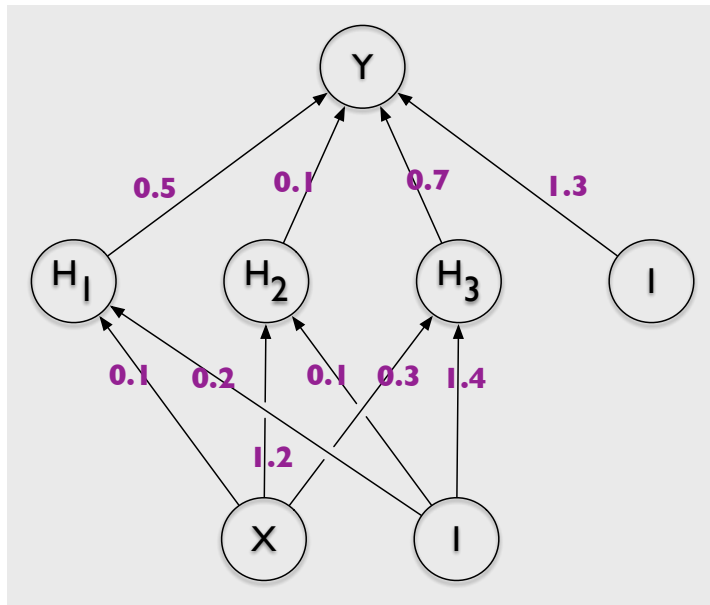
<https://aleximmer.com/Laplace/>



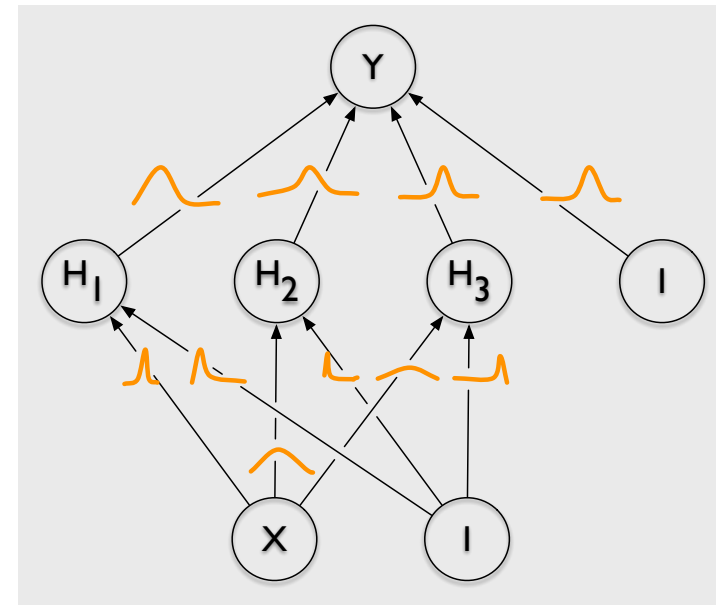
<https://github.com/AaltoML/SUQ>

Bayesian Deep Learning

(Deterministic) Neural Network

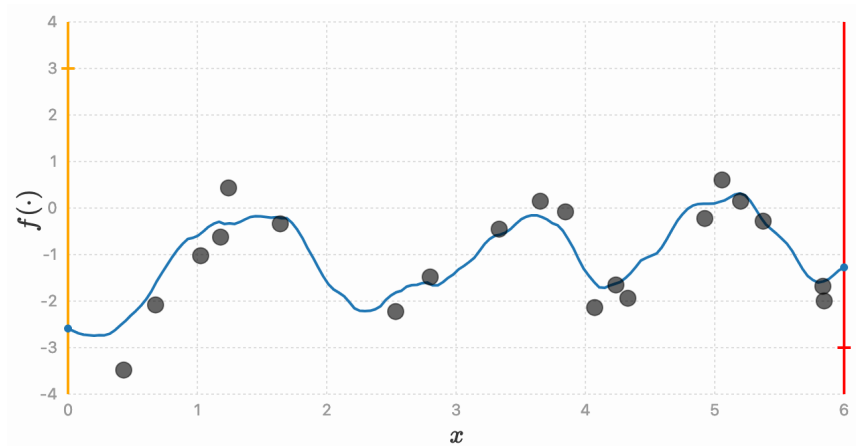


Bayesian Neural Network

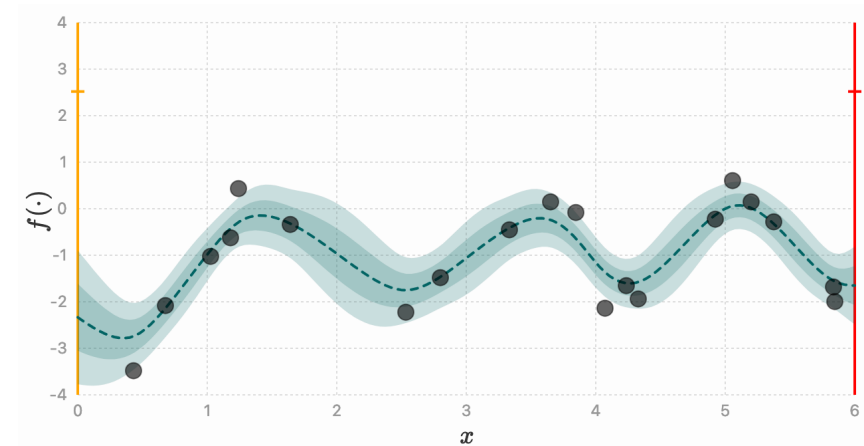


Bayesian Deep Learning

(Deterministic) Neural Network



Bayesian Neural Network



Bayesian Deep Learning

Computations in Bayesian Learning are notoriously **hard!** *in general.

Challenges in Bayesian Deep Learning:

- Hard to specify functional priors.
- Intractable high-dimensional multi-modal posterior.
(billions of parameters)
- Models can be difficult to train from scratch.
(high compute resource demand)

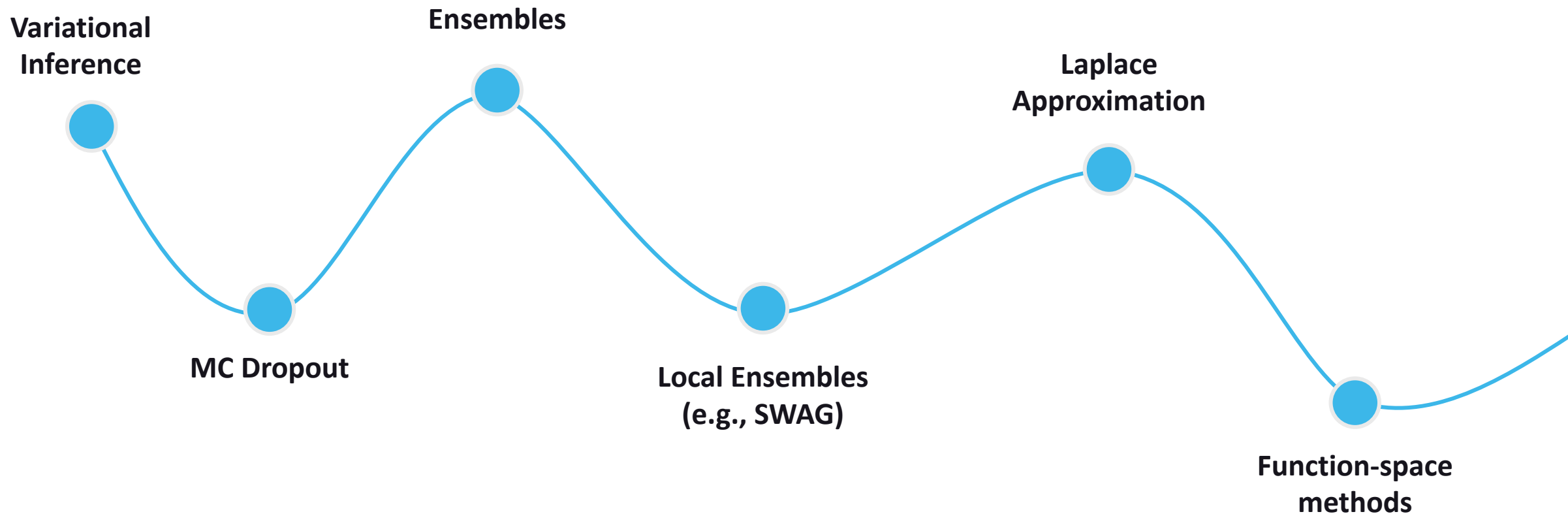
Bayesian Deep Learning

Computations in Bayesian Learning are notoriously **hard!** *in general.

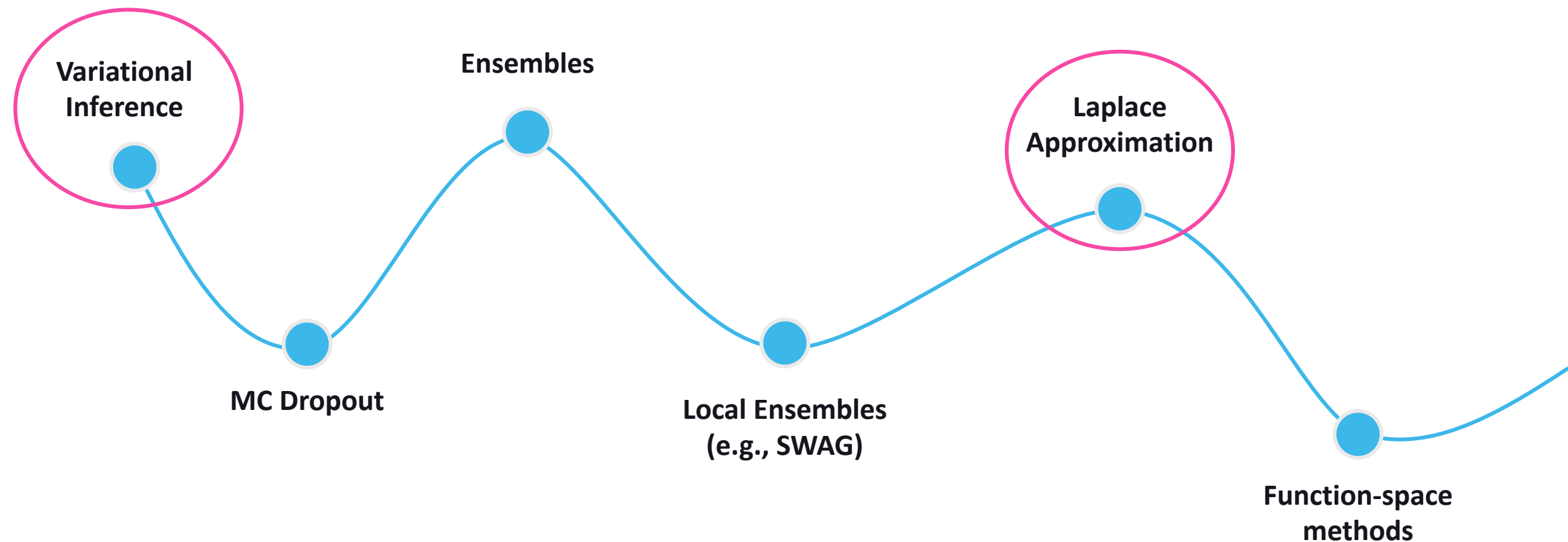
Challenges in Bayesian Deep Learning:

- Hard to specify functional priors.
- **Intractable high-dimensional multi-modal posterior.**
(billions of parameters)
- **Models can be difficult to train from scratch.**
(high compute resource demand)

Inference Methods in BDL



Inference Methods in BDL



Variational Inference: Big Picture

Recipe for approximating an **intractable** distribution $p \in \mathcal{P}$

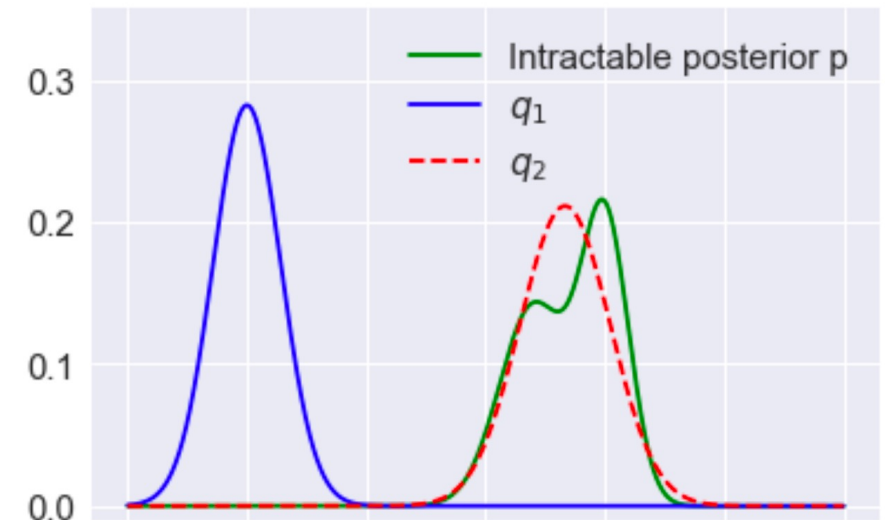
1. Define a **tractable** family of distributions $\mathcal{Q}$

Variational Inference: Big Picture

Recipe for approximating an **intractable** distribution $p \in \mathcal{P}$

1. Define a **tractable** family of distributions $\mathcal{Q}$
2. Define a way to compute a “distance” between distributions

$$D(p||q_1) > D(p||q_2)$$



Variational Inference: Big Picture

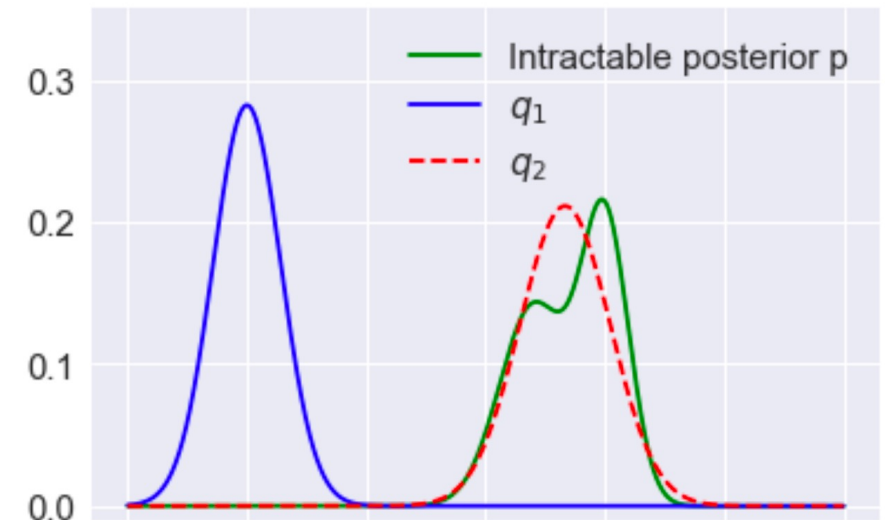
Recipe for approximating an **intractable** distribution $p \in \mathcal{P}$

1. Define a **tractable** family of distributions $\mathcal{Q}$
2. Define a way to compute a “distance” between distributions

$$D(p||q_1) > D(p||q_2)$$

3. Search for best approximation

$$q^* = \arg \min_{q \in \mathcal{Q}} D(p||q)$$



Reminder on KL divergence

Let P and Q be two probability distributions with p.d.f. denoted as p and q , respectively.

Then their **Kullback–Leibler divergence** is given as:

$$D_{\text{KL}}(P||Q) = \int_x p(x) \log \left(\frac{p(x)}{q(x)} \right) dx$$

Reminder on KL divergence

Let P and Q be two probability distributions with p.d.f. denoted as p and q , respectively.

Then their **Kullback–Leibler divergence** is given as:

$$D_{\text{KL}}(P||Q) = \int_x p(x) \log \left(\frac{p(x)}{q(x)} \right) dx$$

$$D_{\text{KL}}(P||Q) \geq 0 \quad \text{and} \quad D_{\text{KL}}(P||Q) = 0 \iff P = Q$$

Reminder on ELBO

Reverse KL

$$\boxed{D_{\text{KL}}(q(\theta) || p(\theta \mid \mathbf{x}_{1:n}))} = \mathbb{E}_{\theta \sim q} \left[\log \left(\frac{q(\theta)}{p(\theta \mid \mathbf{x}_{1:n})} \right) \right]$$

Reminder on ELBO

$$\begin{aligned} D_{\text{KL}}(q(\theta) || p(\theta \mid \mathbf{x}_{1:n})) &= \mathbb{E}_{\theta \sim q} \left[\log \left(\frac{q(\theta)}{p(\theta \mid \mathbf{x}_{1:n})} \right) \right] \\ &= \mathbb{E}_{\theta \sim q} \left[\log \left(\frac{q(\theta)}{\frac{p(\mathbf{x}_{1:n} \mid \theta) p(\theta)}{p(\mathbf{x}_{1:n})}} \right) \right] \end{aligned}$$

Reminder on ELBO

$$\begin{aligned} D_{\text{KL}}(q(\theta) || p(\theta \mid \mathbf{x}_{1:n})) &= \mathbb{E}_{\theta \sim q} \left[\log \left(\frac{q(\theta)}{p(\theta \mid \mathbf{x}_{1:n})} \right) \right] \\ &= \mathbb{E}_{\theta \sim q} \left[\log \left(\frac{q(\theta)}{\frac{p(\mathbf{x}_{1:n} \mid \theta) p(\theta)}{p(\mathbf{x}_{1:n})}} \right) \right] \\ &= \mathbb{E}_{\theta \sim q} \left[\log \left(\frac{q(\theta)}{p(\theta)} \right) + \log \left(\frac{p(\mathbf{x}_{1:n})}{p(\mathbf{x}_{1:n} \mid \theta)} \right) \right] \end{aligned}$$

Reminder on ELBO

$$\begin{aligned} D_{\text{KL}}(q(\theta) || p(\theta | \mathbf{x}_{1:n})) &= \mathbb{E}_{\theta \sim q} \left[\log \left(\frac{q(\theta)}{p(\theta | \mathbf{x}_{1:n})} \right) \right] \\ &= \mathbb{E}_{\theta \sim q} \left[\log \left(\frac{q(\theta)}{\frac{p(\mathbf{x}_{1:n} | \theta) p(\theta)}{p(\mathbf{x}_{1:n})}} \right) \right] \\ &= \mathbb{E}_{\theta \sim q} \left[\log \left(\frac{q(\theta)}{p(\theta)} \right) + \log \left(\frac{p(\mathbf{x}_{1:n})}{p(\mathbf{x}_{1:n} | \theta)} \right) \right] \\ &= D_{\text{KL}}(q(\theta) || p(\theta)) - \mathbb{E}_{\theta \sim q} [\log p(\mathbf{x}_{1:n} | \theta)] + \log p(\mathbf{x}_{1:n}) \end{aligned}$$

Reminder on ELBO

$$\begin{aligned} D_{\text{KL}}(q(\theta) || p(\theta \mid \mathbf{x}_{1:n})) &= \mathbb{E}_{\theta \sim q} \left[\log \left(\frac{q(\theta)}{p(\theta \mid \mathbf{x}_{1:n})} \right) \right] \\ &\vdots \\ &= D_{\text{KL}}(q(\theta) || p(\theta)) - \mathbb{E}_{\theta \sim q} [\log p(\mathbf{x}_{1:n} \mid \theta)] + \log p(\mathbf{x}_{1:n}) \end{aligned}$$

Reminder on ELBO

$$D_{\text{KL}}(q(\theta) || p(\theta | \mathbf{x}_{1:n})) = \mathbb{E}_{\theta \sim q} \left[\log \left(\frac{q(\theta)}{p(\theta | \mathbf{x}_{1:n})} \right) \right]$$

$\vdots$

$$= D_{\text{KL}}(q(\theta) || p(\theta)) - \mathbb{E}_{\theta \sim q} [\log p(\mathbf{x}_{1:n} | \theta)] + \log p(\mathbf{x}_{1:n})$$

$$\log p(\mathbf{x}_{1:n}) = -D_{\text{KL}}(q(\theta) || p(\theta)) + \mathbb{E}_{\theta \sim q} [\log p(\mathbf{x}_{1:n} | \theta)] + \underbrace{D_{\text{KL}}(q(\theta) || p(\theta | \mathbf{x}_{1:n}))}_{>0}$$

Reminder on ELBO

$$D_{\text{KL}}(q(\theta) || p(\theta | \mathbf{x}_{1:n})) = \mathbb{E}_{\theta \sim q} \left[\log \left(\frac{q(\theta)}{p(\theta | \mathbf{x}_{1:n})} \right) \right]$$

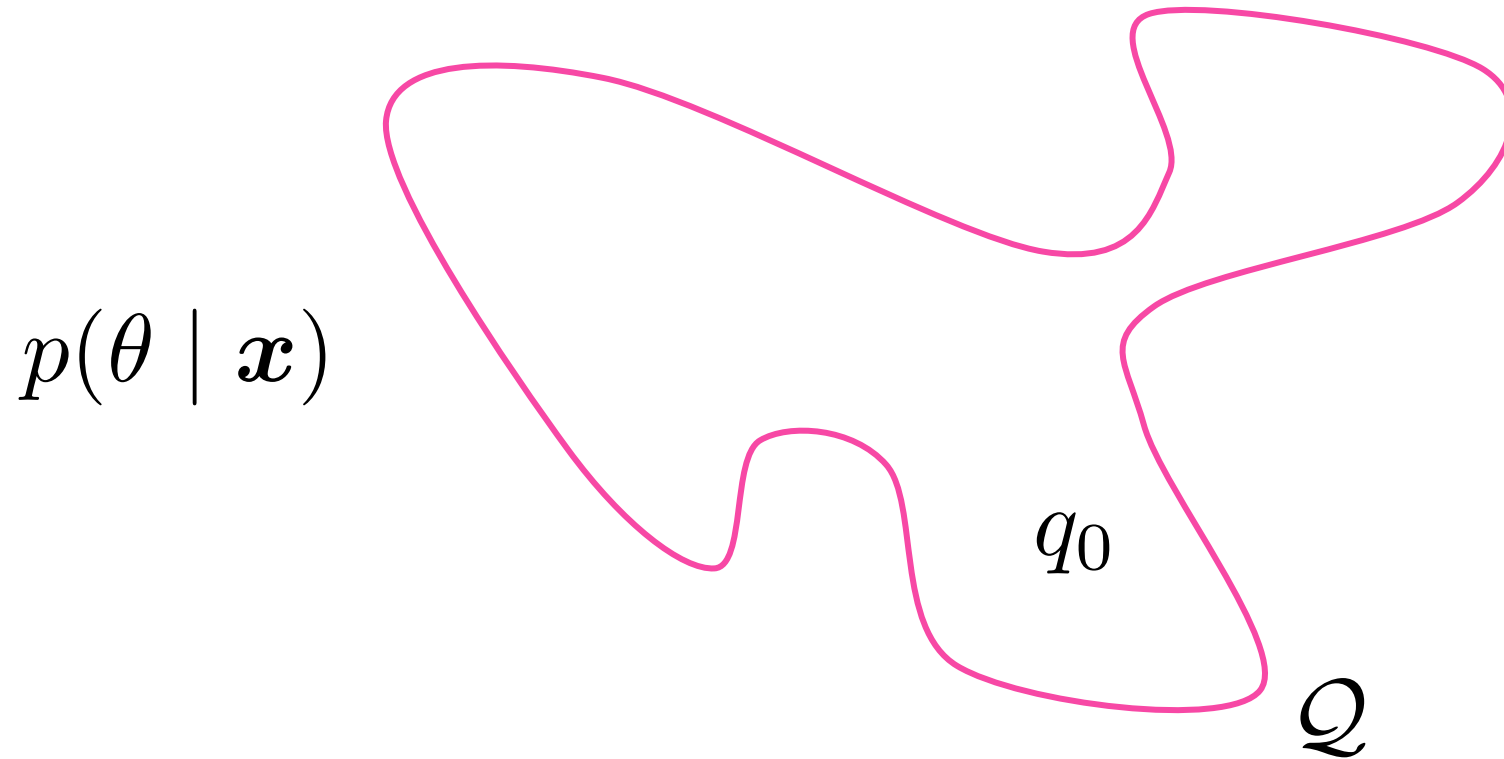
$\vdots$

$$= D_{\text{KL}}(q(\theta) || p(\theta)) - \mathbb{E}_{\theta \sim q} [\log p(\mathbf{x}_{1:n} | \theta)] + \log p(\mathbf{x}_{1:n})$$

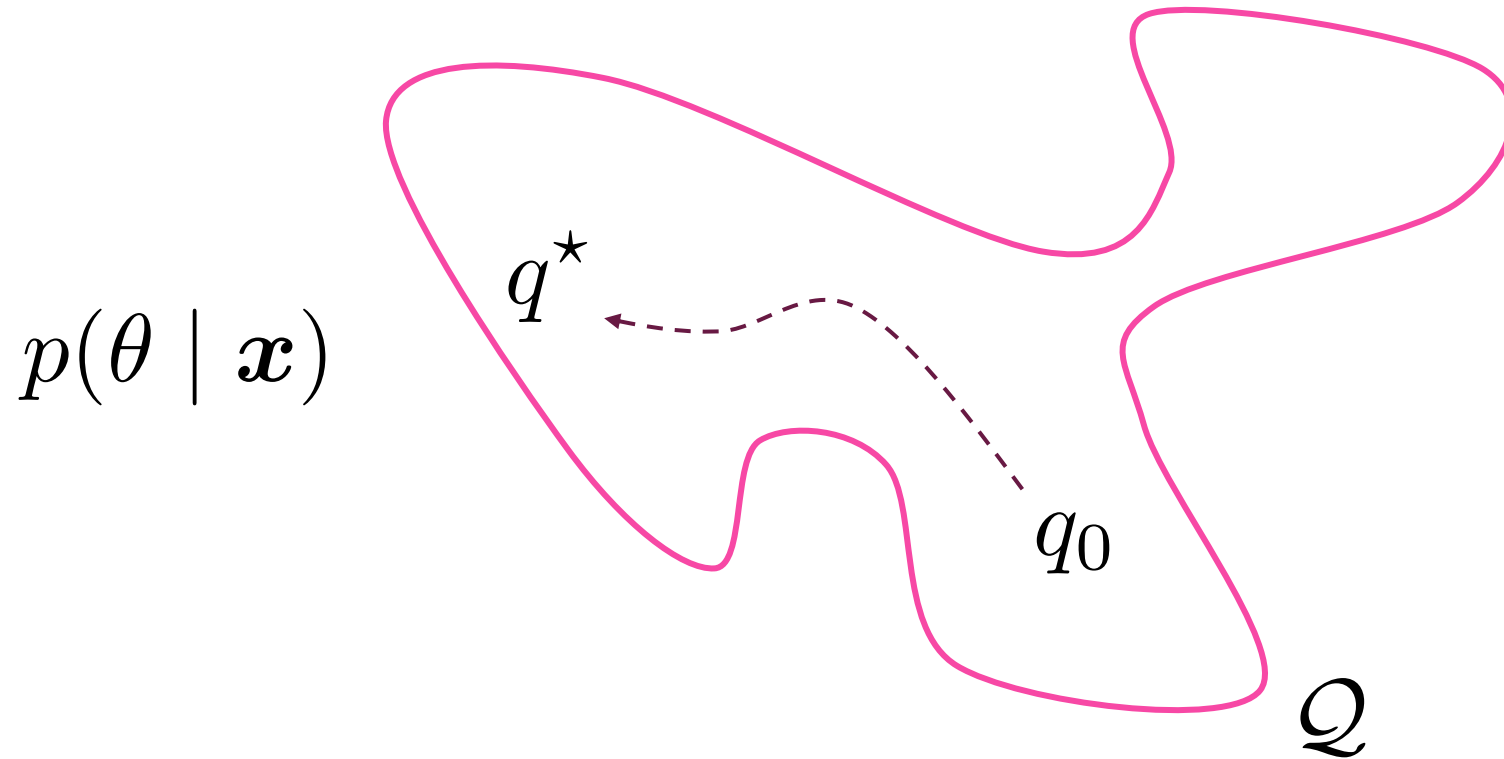
$$\log p(\mathbf{x}_{1:n}) = -D_{\text{KL}}(q(\theta) || p(\theta)) + \mathbb{E}_{\theta \sim q} [\log p(\mathbf{x}_{1:n} | \theta)] + \underbrace{D_{\text{KL}}(q(\theta) || p(\theta | \mathbf{x}_{1:n}))}_{\geq 0}$$

$$\geq -D_{\text{KL}}(q(\theta) || p(\theta)) + \mathbb{E}_{\theta \sim q} [\log p(\mathbf{x}_{1:n} | \theta)] = \text{ELBO (Evidence Lower Bound)}$$

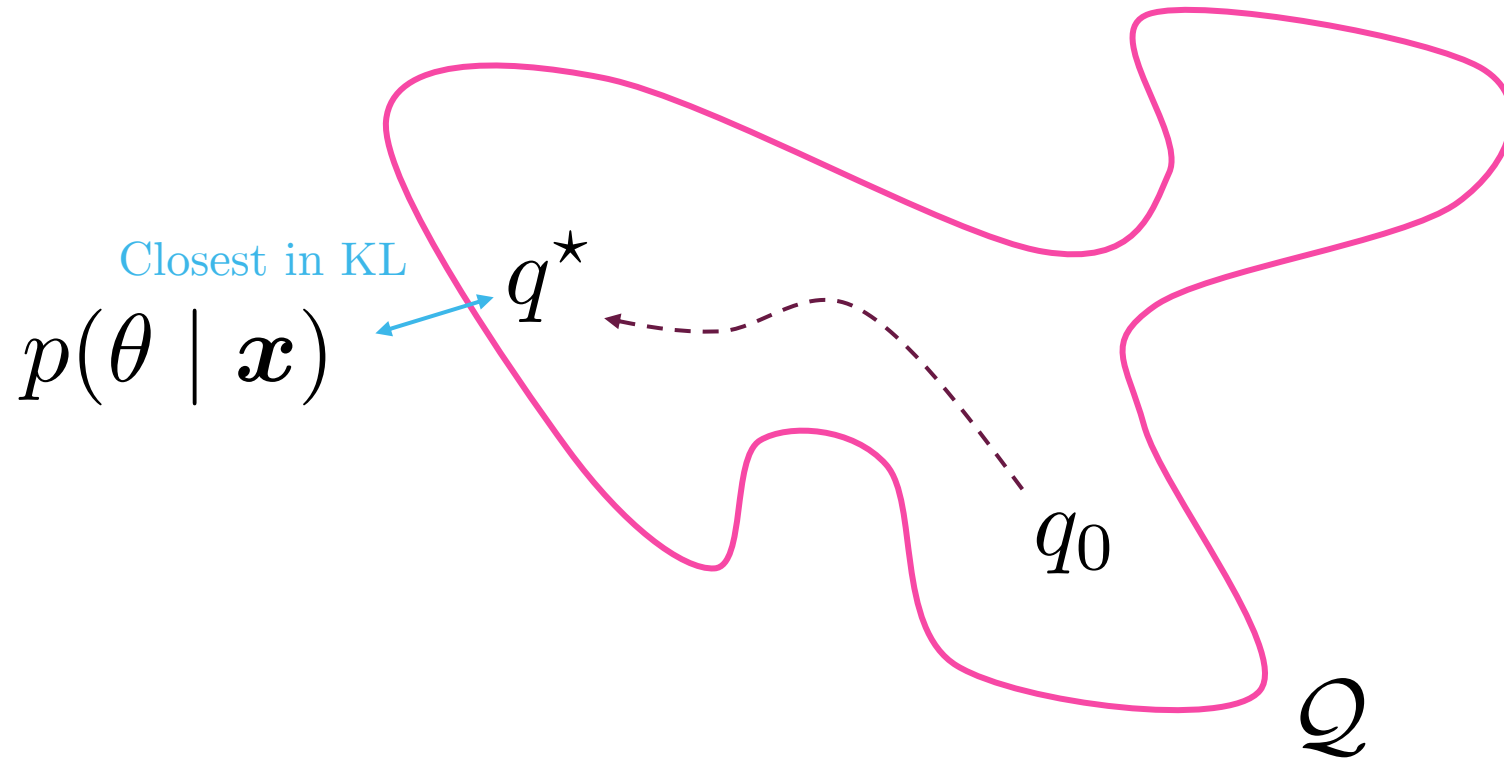
Variational Inference



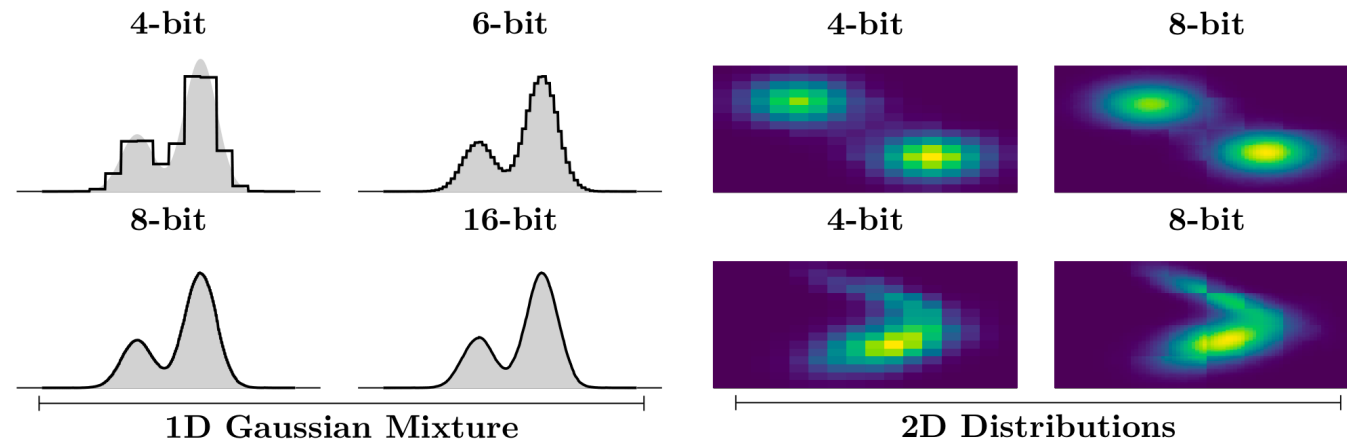
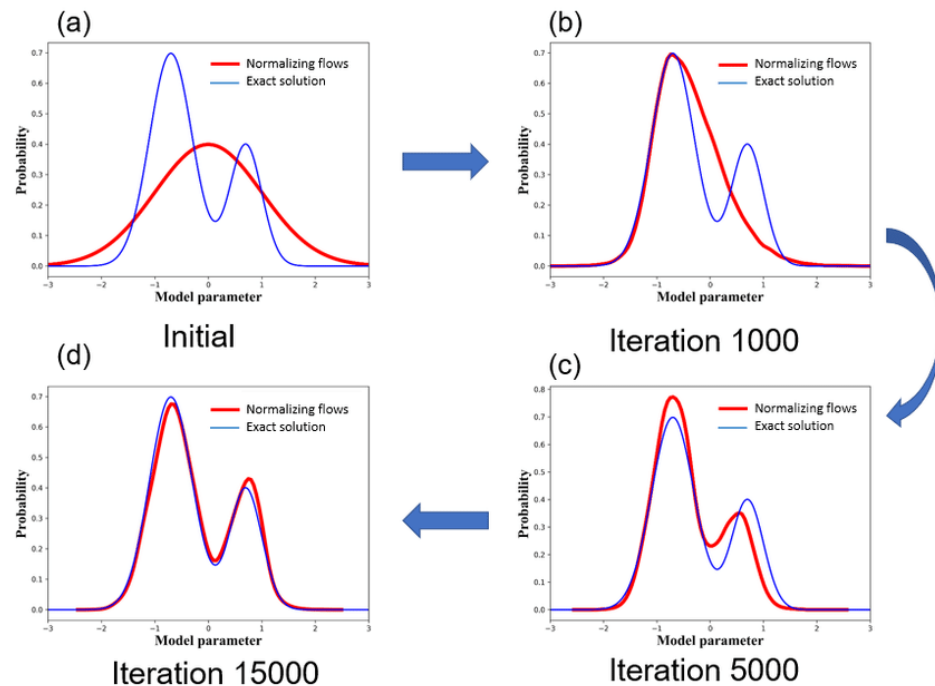
Variational Inference



Variational Inference



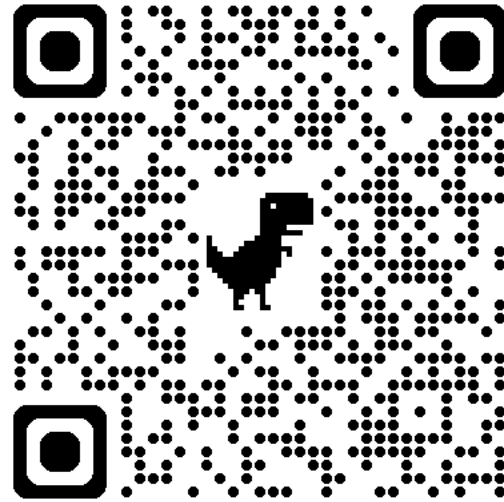
Variational Inference



Sladek, A. et al. (2025). Approximate Bayesian Inference via Bitstring Representations. In *UAI*.

Rezende, D. et al. (2015). Variational Inference with Normalizing Flows. In *ICML*.

Variational Inference



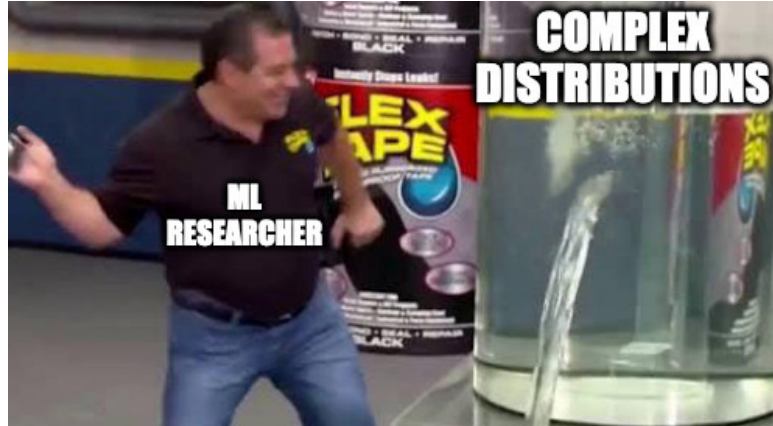
<https://lacerbi.github.io/blog/2024/vi-is-inference-is-optimization/>

Variational Inference: Takeaways

- VI turns inference into an optimisation problem.
 - Can scale to large data sets, but is challenging to optimise
 - Can exhibit high variance in stochastic gradients
 - Tendency to underestimate the variance if reverse KL is used
-
- Recent work has shown that VI can be effective for large-scale models.

Laplace Approximation

What if training a model “from scratch” is too expensive?



Laplace Approximation

What if training a model “from scratch” is too expensive?



Laplace Approximation

$$\log p(\theta \mid \mathbf{x}) = \overbrace{\underbrace{\log p(\theta)}_{\text{L2 regularization}} + \underbrace{\log p(\mathbf{x} \mid \theta)}_{\text{Cross-Entropy}}}^{\text{Negated Loss}} - \underbrace{C}_{\text{Unknown}}$$

Laplace Approximation

$$\log p(\theta \mid \mathbf{x}) = \overbrace{\underbrace{\log p(\theta)}_{\text{L2 regularization}} + \underbrace{\log p(\mathbf{x} \mid \theta)}_{\text{Cross-Entropy}}}_{\text{Negated Loss}} - \underbrace{C}_{\text{Unknown}}$$

$$\underbrace{\log p(\theta) + \log p(\mathbf{x} \mid \theta)}_{\ell(\theta)} \approx \ell(\theta^*) + \mathbf{J}_{\ell|_{\theta=\theta^*}} (\theta - \theta^*) - \frac{1}{2} \mathbf{H}_{\ell|_{\theta=\theta^*}} (\theta - \theta^*)^2$$

Laplace Approximation

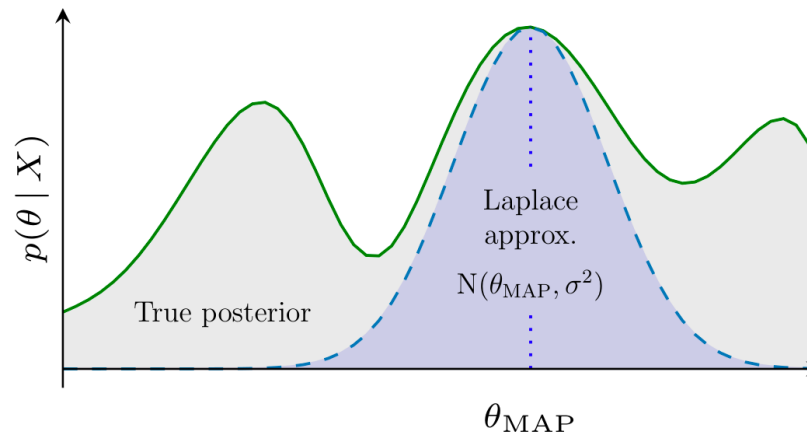
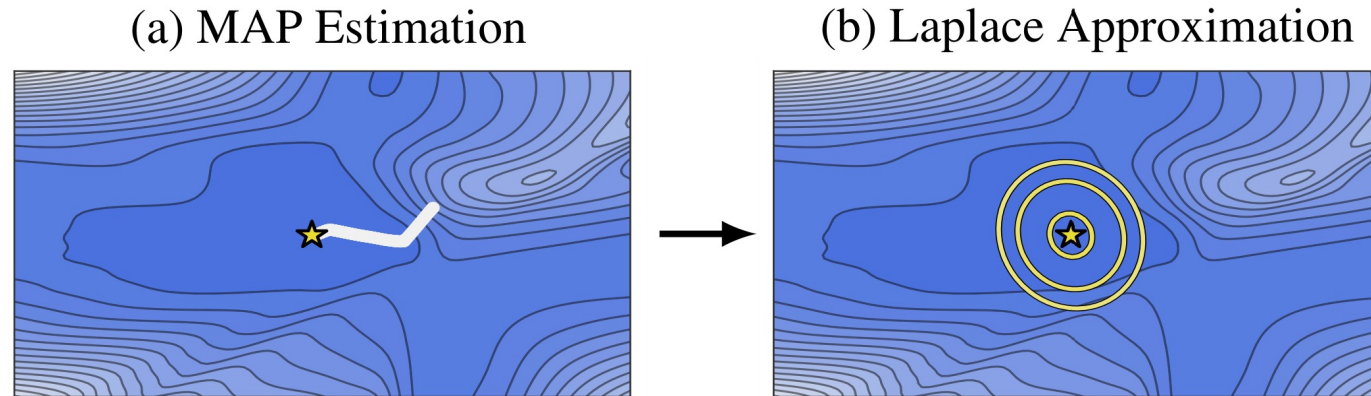
$$\log p(\theta \mid \mathbf{x}) = \overbrace{\underbrace{\log p(\theta)}_{\text{L2 regularization}} + \underbrace{\log p(\mathbf{x} \mid \theta)}_{\text{Cross-Entropy}}}^{\text{Negated Loss}} - \underbrace{C}_{\text{Unknown}}$$

$$\underbrace{\log p(\theta) + \log p(\mathbf{x} \mid \theta)}_{\ell(\theta)} \approx \ell(\theta^*) + \mathbf{J}_{\ell|_{\theta=\theta^*}} (\theta - \theta^*) - \frac{1}{2} \mathbf{H}_{\ell|_{\theta=\theta^*}} (\theta - \theta^*)^2$$

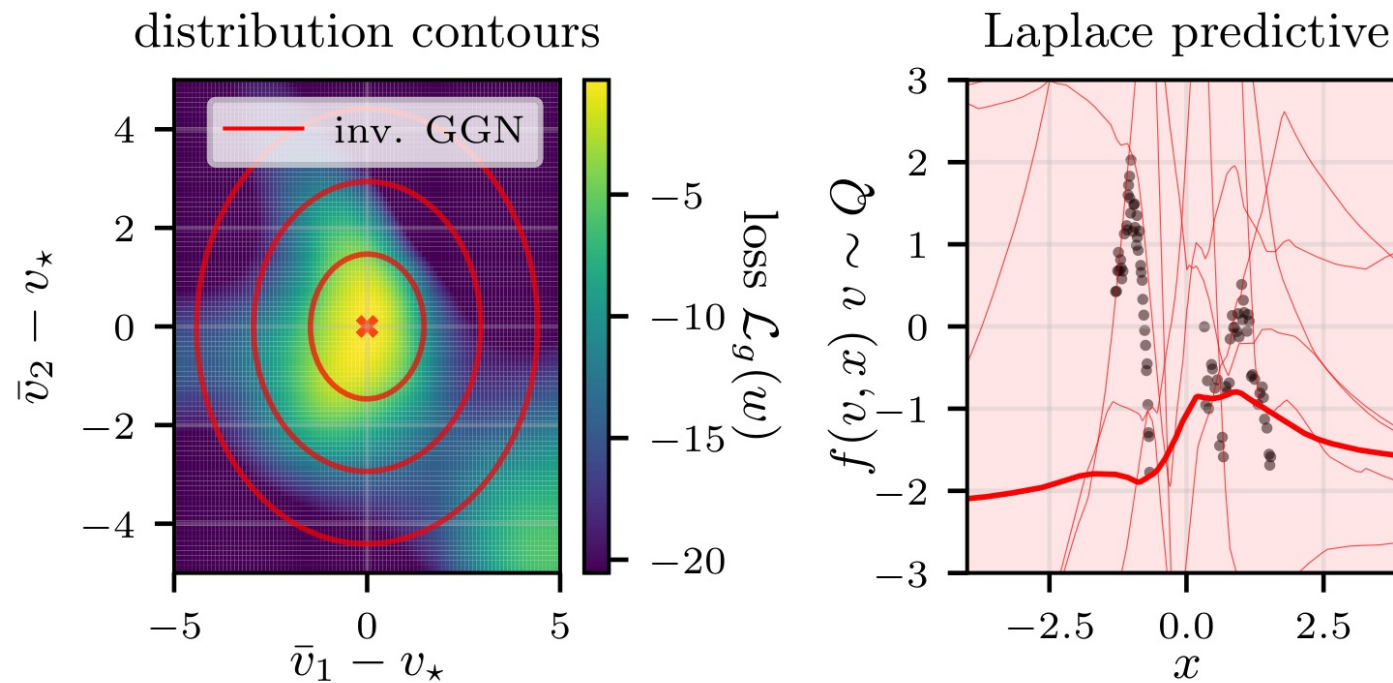
$$= \ell(\theta^*) - \frac{1}{2} \mathbf{H}_{\ell|_{\theta=\theta^*}} (\theta - \theta^*)^2$$

$$\propto \log \mathcal{N}(\theta^*, \mathbf{H}_{\ell|_{\theta=\theta^*}}^{-1})$$

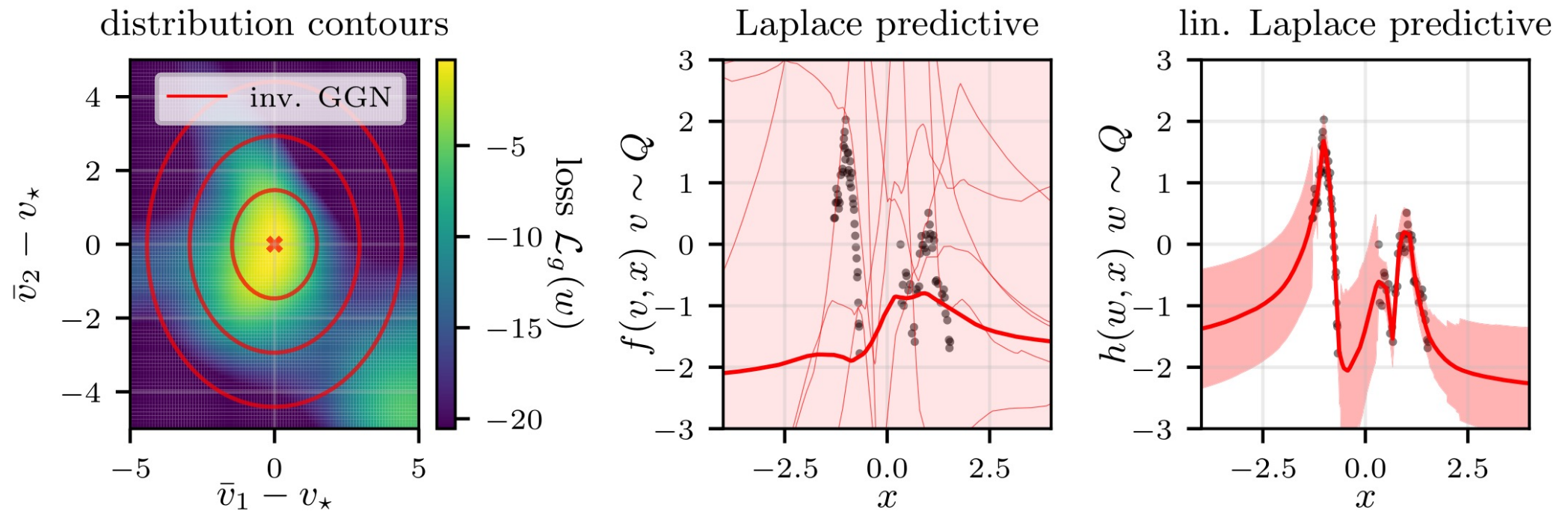
Laplace Approximation



Linearised Laplace: Problem Setting



Linearised Laplace: Problem Setting



Linearised Laplace

The linearised Laplace approximates the predictive distribution:

$$f(x; \theta) \approx f(x; \theta_{\text{MAP}}) + \nabla f(x; \theta) \big|_{\theta_{\text{MAP}}}^{\top} (\theta - \theta_{\text{MAP}}) = \hat{f}(x; \theta)$$

Linearised Laplace

The linearised Laplace approximates the predictive distribution:

$$f(x; \theta) \approx f(x; \theta_{\text{MAP}}) + \nabla f(x; \theta) \big|_{\theta_{\text{MAP}}}^{\top} (\theta - \theta_{\text{MAP}}) = \hat{f}(x; \theta)$$

$$p(y^* \mid x^*, \mathbf{x}, \mathbf{y}) \approx \int \mathcal{N}(y^* \mid \hat{f}(x; \theta), \sigma) \mathcal{N}(\theta \mid \theta_{\text{MAP}}, \Sigma) d\theta$$

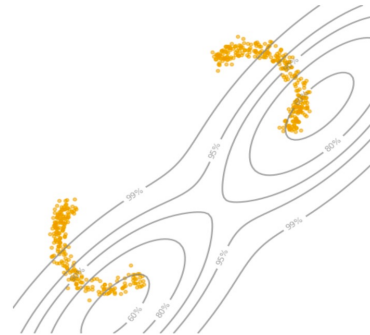
Laplace Approximation: Takeaways

- Does not require retraining (post-hoc).
 - Requires linearisation in the case of deep learning models.
 - Crude and only local approximation.
 - Estimation of the Hessian can be difficult.
-
- Crude but useful (easy to use) “tool” for Bayesian inference in deep learning.

Machine Learning for Bayes

Amortized Inference

- NNs are powerful function approximators and generators (e.g., diffusion models)
- How can we use NNs for approximate Bayesian inference?
- → Emerging research on amortized inference with NNs (condition and predict)



(a) Prior v samples



(b) PriorGuide v samples



(c) PriorGuide v retrained

Recap

- Deep learning methods can be problematic in high-risk domains.
- The Bayesian approach to deep learning **can help reduce overconfidence, quantify uncertainties**, and utilise uncertainties in decision-making.
- **Prior specification is challenging** and an open question.
- Performing **computations is challenging**, but **promising avenues exist**.
- **Machine learning can help scaling Bayesian inference**, e.g., through amortisation.

Thanks for your attention!

