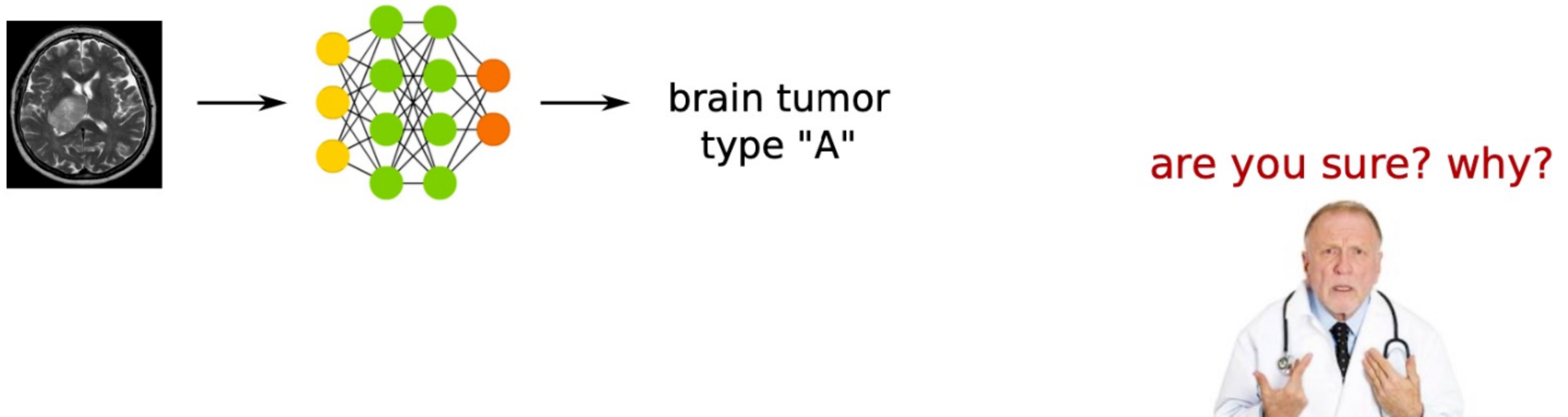


A Generalisation Analysis of Bayesian Machine Learning

Andres Masegosa

Aalborg University

Uncertainty in Machine Learning



- How confident is a ML model when making an specific prediction?
- Quite relevant in many domains.

Uncertainty in Machine Learning



Here's this patient's health record:
...
Could you summarise it for me?

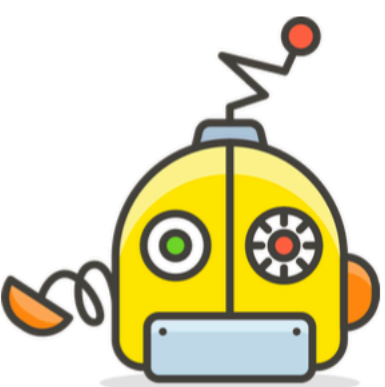
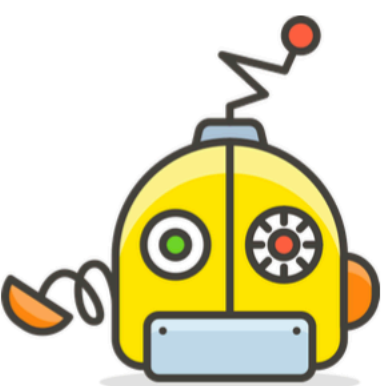
Here's a summary of the patient's health record you requested:
[point 1] with x% confidence
...



Here's the conditions of this construction site:
...
Could you tell me what the potential safety issues are?

Here's potential safety issues that need to be look after:
[point 1] with x% confidence
...

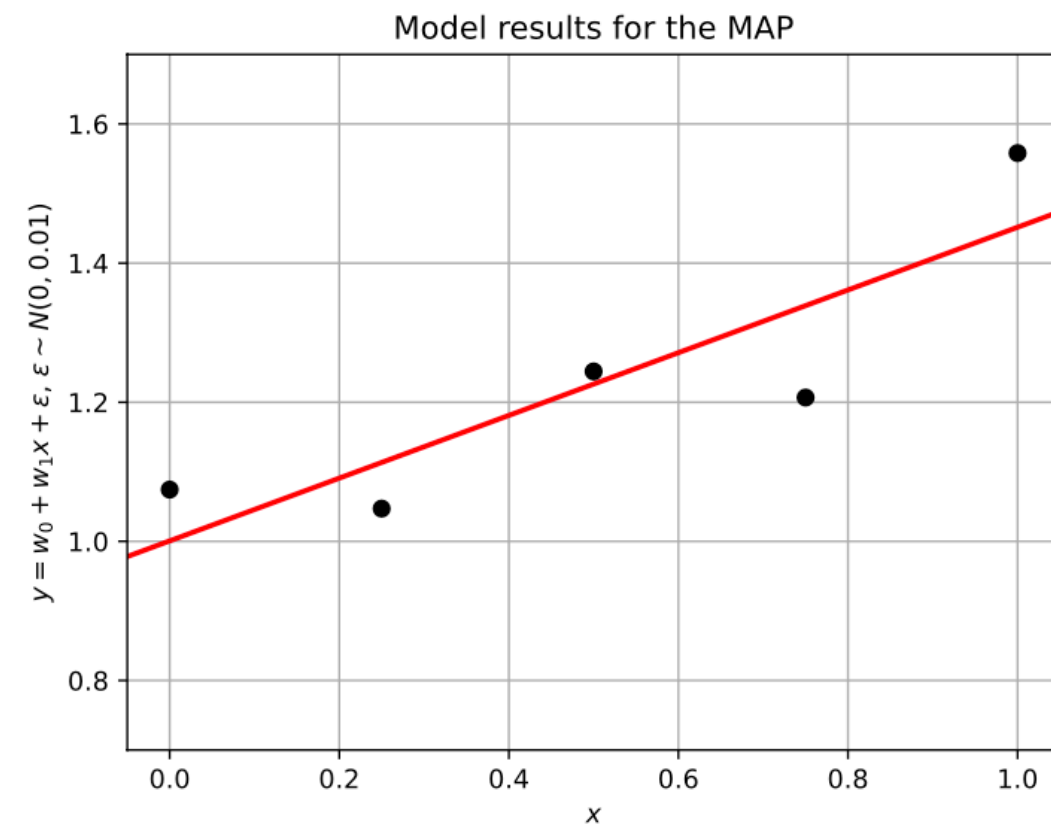
Desiderata



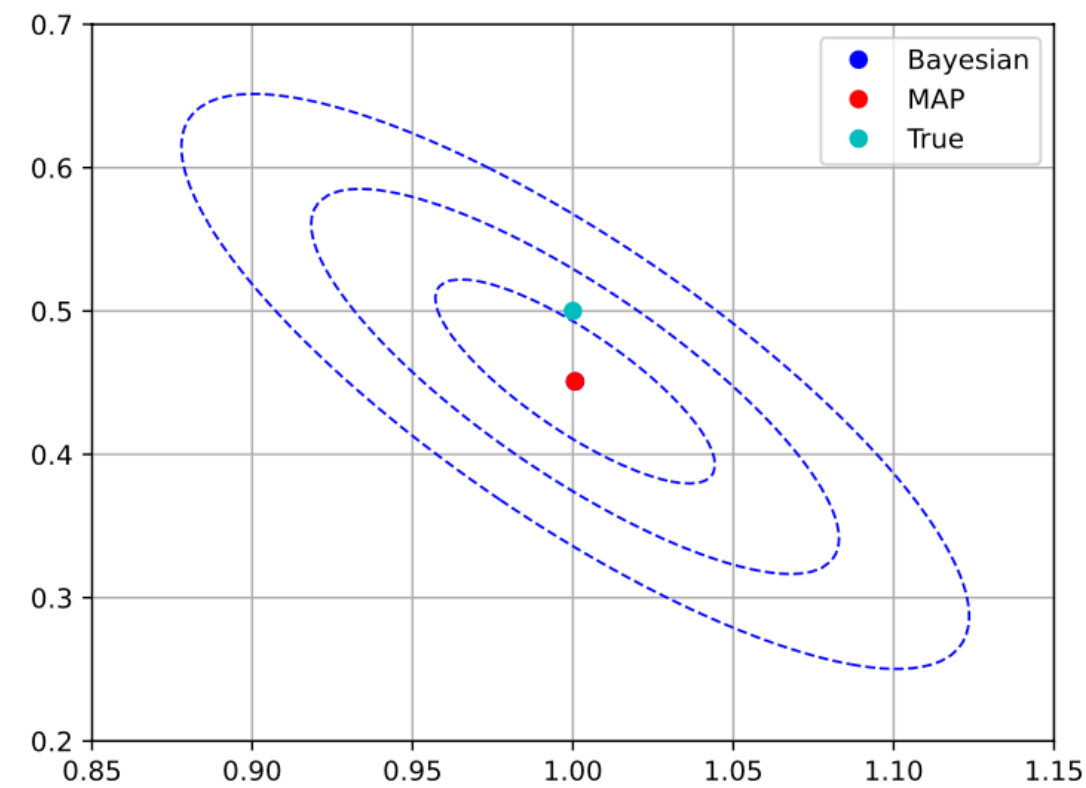
Bayesian in Machine Learning

- Bayesian methods are **widely used** in machine learning for uncertainty quantification.
- They provide well founded approach for dealing with **model/data uncertainty**.
- **Random variables + Probability Calculus.**
- They automatically account for **model complexity**.
- They allow to combine data with **prior knowledge**.

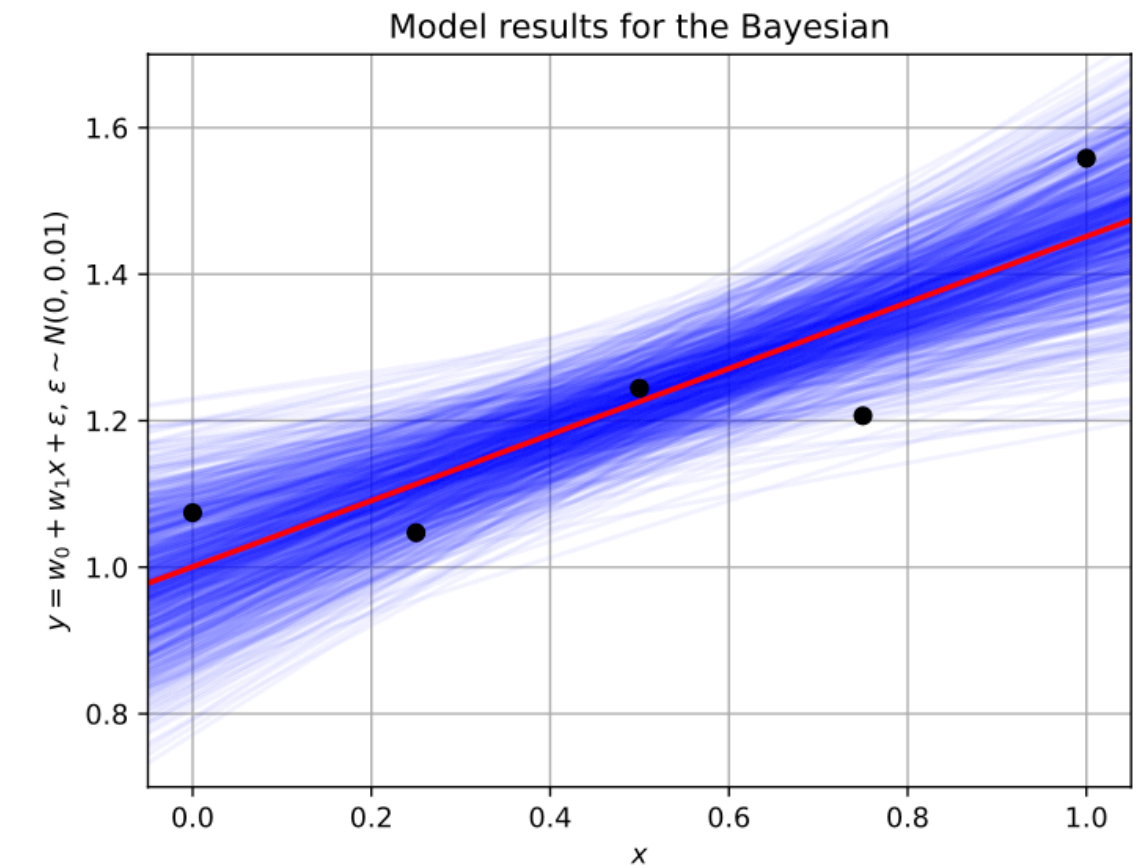
Bayesian Linear Regression



$$\theta_{MAP} = \arg \max_{\theta} p(D | \theta)$$



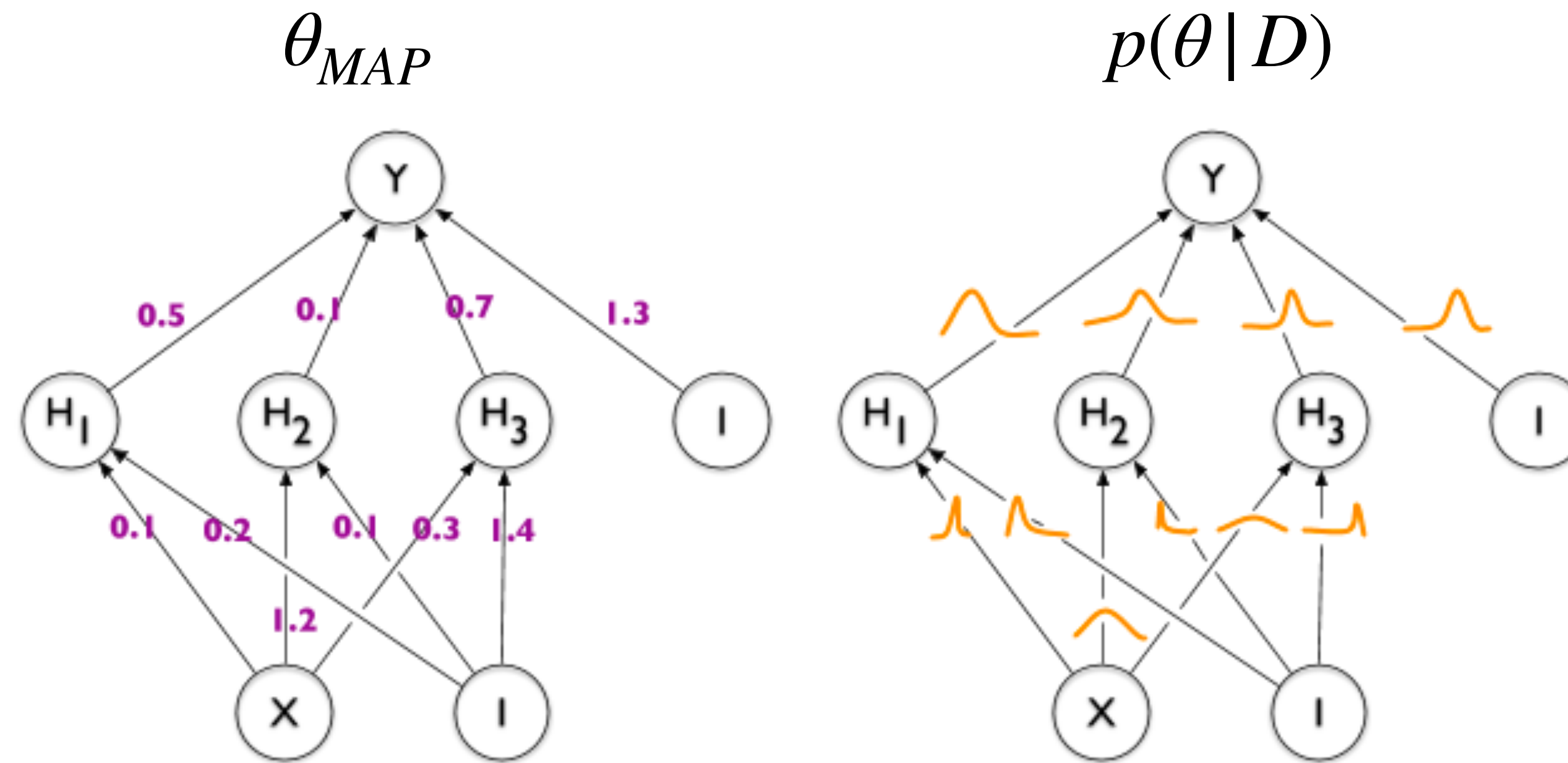
$$\theta_{MAP} = \arg \max_{\theta} p(\theta | D)$$



$$\theta_i \sim p(\theta | D)$$

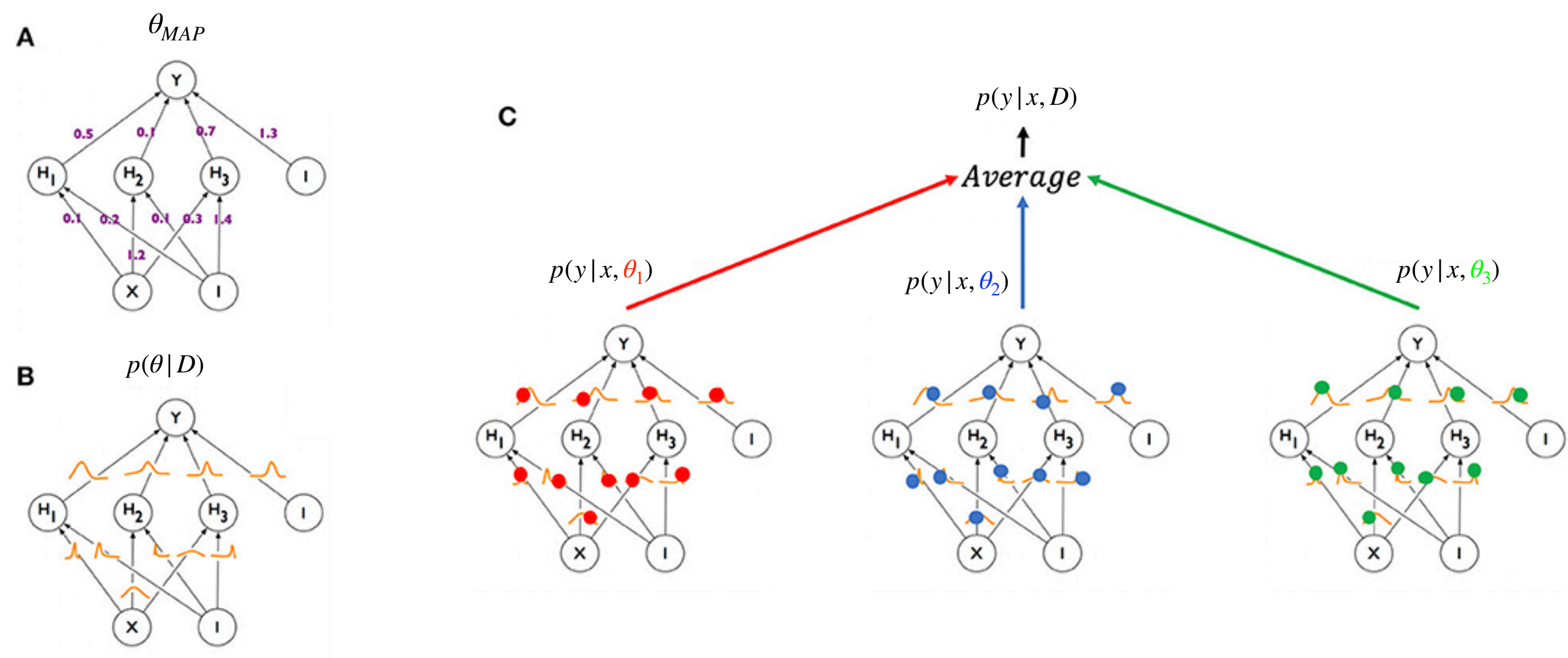
$$\underbrace{p(\boldsymbol{\theta} | D)}_{\text{Bayesian posterior}} = \frac{\overbrace{p(D | \boldsymbol{\theta})}^{\text{Likelihood}} \overbrace{\pi(\boldsymbol{\theta})}^{\text{Prior}}}{\underbrace{\int p(D | \boldsymbol{\theta}) \pi(\boldsymbol{\theta}) d\boldsymbol{\theta}}_{\text{Normalization Constant}}}$$

Bayesian Neural Networks



$$\underbrace{p(\boldsymbol{\theta} | D)}_{\text{Bayesian posterior}} = \frac{\overbrace{p(D | \boldsymbol{\theta})}^{\text{Likelihood}} \overbrace{\pi(\boldsymbol{\theta})}^{\text{Prior}}}{\underbrace{\int p(D | \boldsymbol{\theta}) \pi(\boldsymbol{\theta}) d\boldsymbol{\theta}}_{\text{Normalization Constant}}}$$

Bayesian Neural Networks



$$\theta_i \sim p(\theta|D)$$

Main Focus of Bayesian Machine Learning

- The main problem is how to compute the **Bayesian posterior**

$$\underbrace{p(\boldsymbol{\theta}|D)}_{\text{Bayesian posterior}} = \frac{\overbrace{p(D|\boldsymbol{\theta})}^{\text{Likelihood}} \overbrace{\pi(\boldsymbol{\theta})}^{\text{Prior}}}{\underbrace{\int p(D|\boldsymbol{\theta})\pi(\boldsymbol{\theta})d\boldsymbol{\theta}}_{\text{Normalization Constant}}}$$

- Intractable for models with a large number of parameters.
- There exist many approximate inference schemes.

Do Bayesian NNs have better performance?

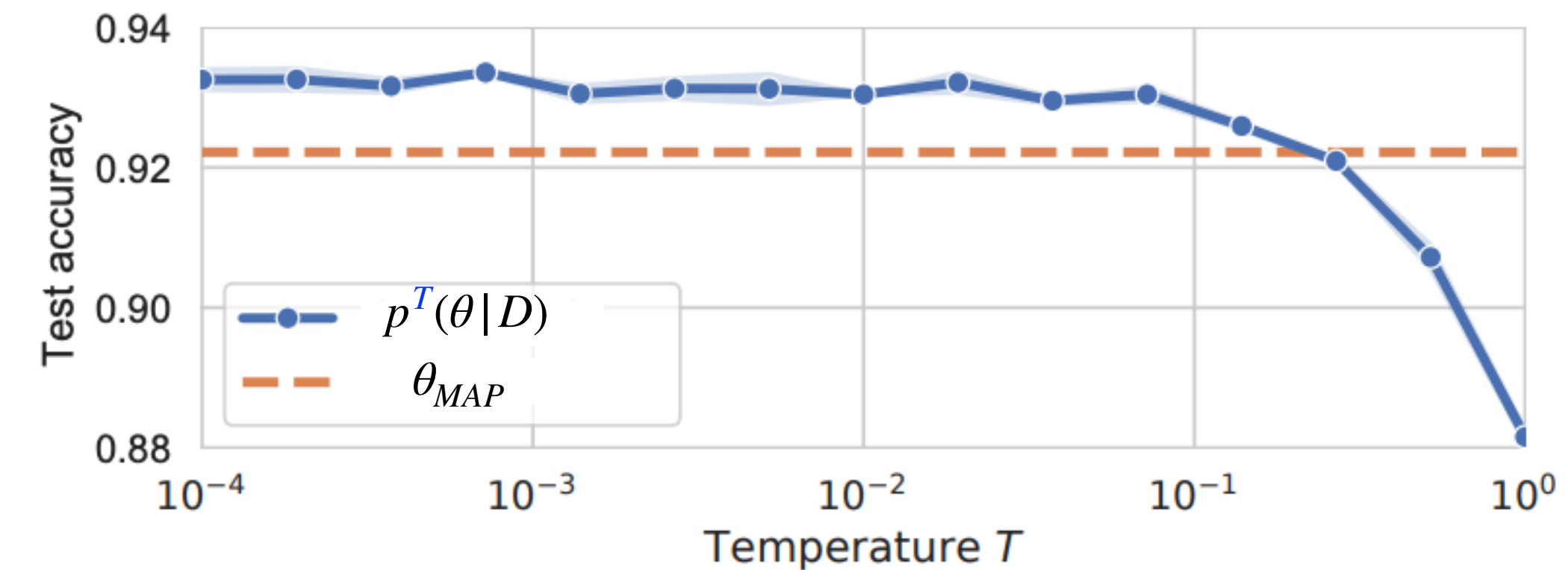
The answer is “not yet”.

Cold Posterior Effect (Wenzel et al., 2020)

- **Tempered Posteriors:**

$$p^T(\theta | D) \propto (P(D | \theta)\pi(\theta))^{1/T}$$

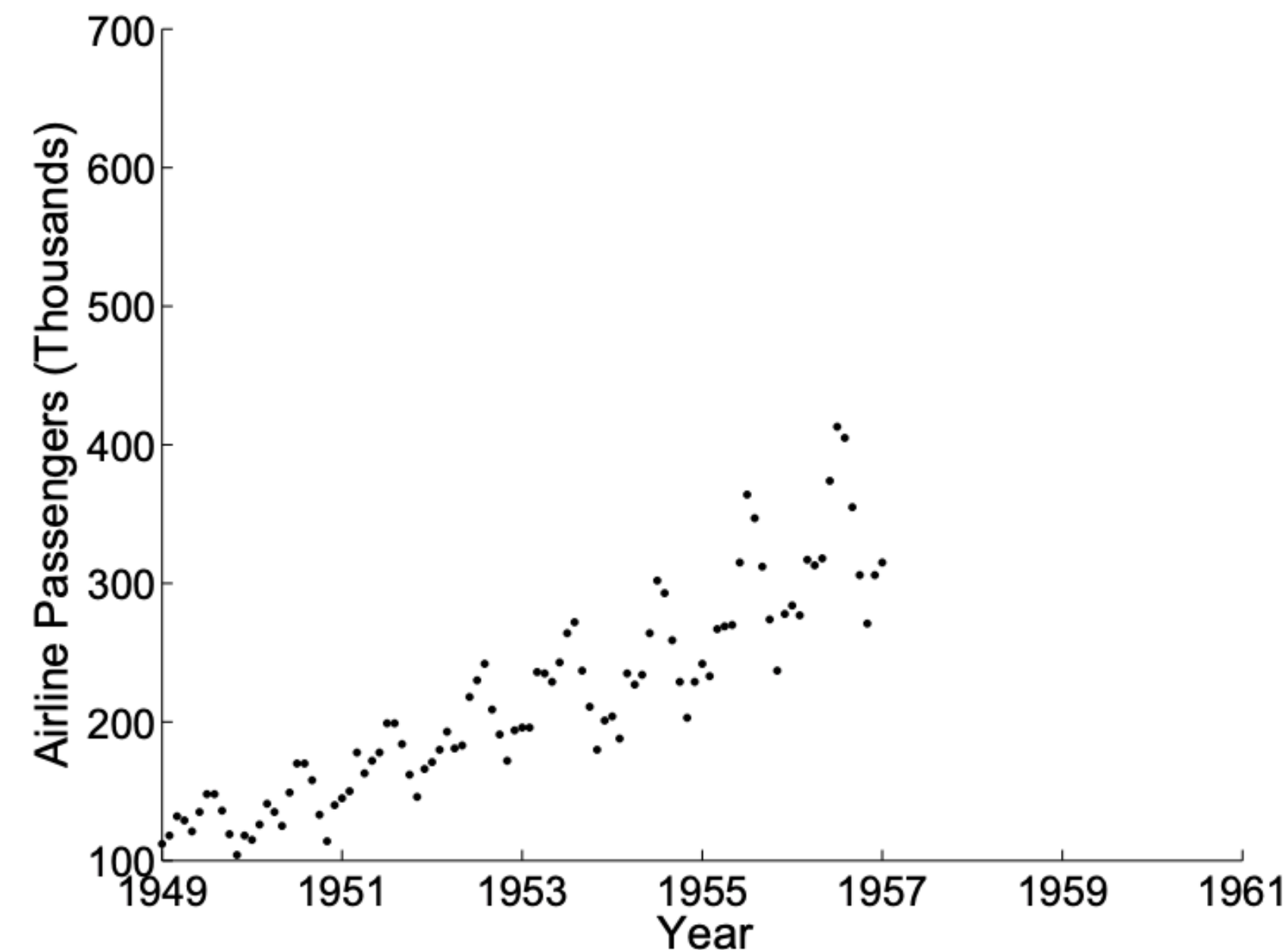
- $T=1$ corresponds to the Bayesian posterior.
- $T=1$ should be optimal for predictive performance.



- **Why do Cold Posteriors perform better?**

- Approximate inference methods (Izmailov et al., 2021).
- Model Misspecification issue (Aitchison, 2021; Kapoor et al., 2022; Fortuin et al., 2022).
- Underfitting issue (Zhang et al. 2023).

Model Selection



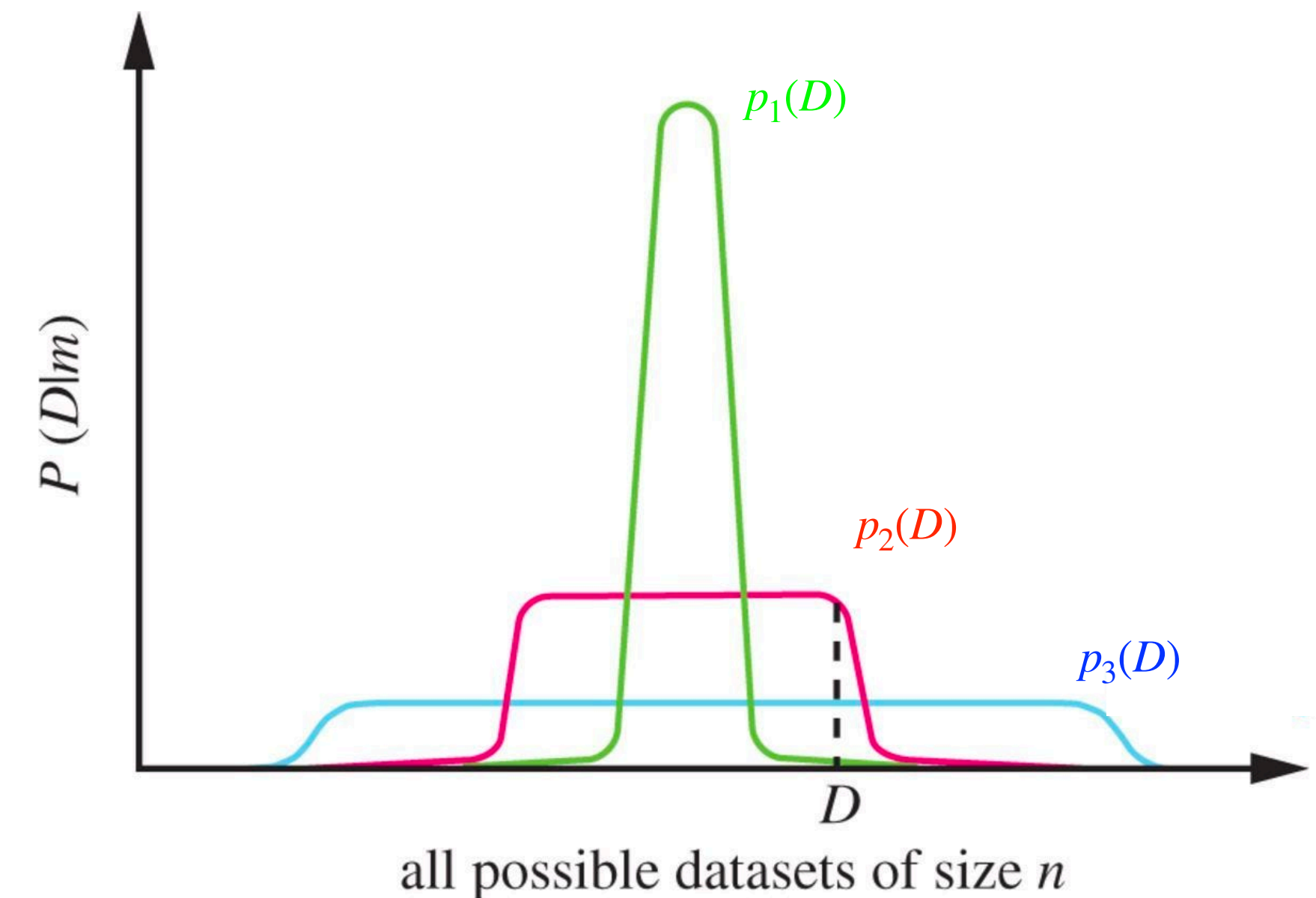
- Which model should we choose?

$$p_1(y|x, \theta) = \mathcal{N}(\theta_0 + \theta_1 x, \sigma)$$
$$p_2(y|x, \theta) = \mathcal{N}\left(\sum_{j=0}^3 \theta_j x^j, \sigma\right)$$
$$p_3(y|x, \theta) = \mathcal{N}\left(\sum_{j=0}^{10^4} \theta_j x^j, \sigma\right)$$

Bayesian Model Selection (Mackay, 2003)

- **Marginal Likelihood**

$$p_k(D) = \int p_k(D | \theta) \pi_k(\theta) d\theta$$

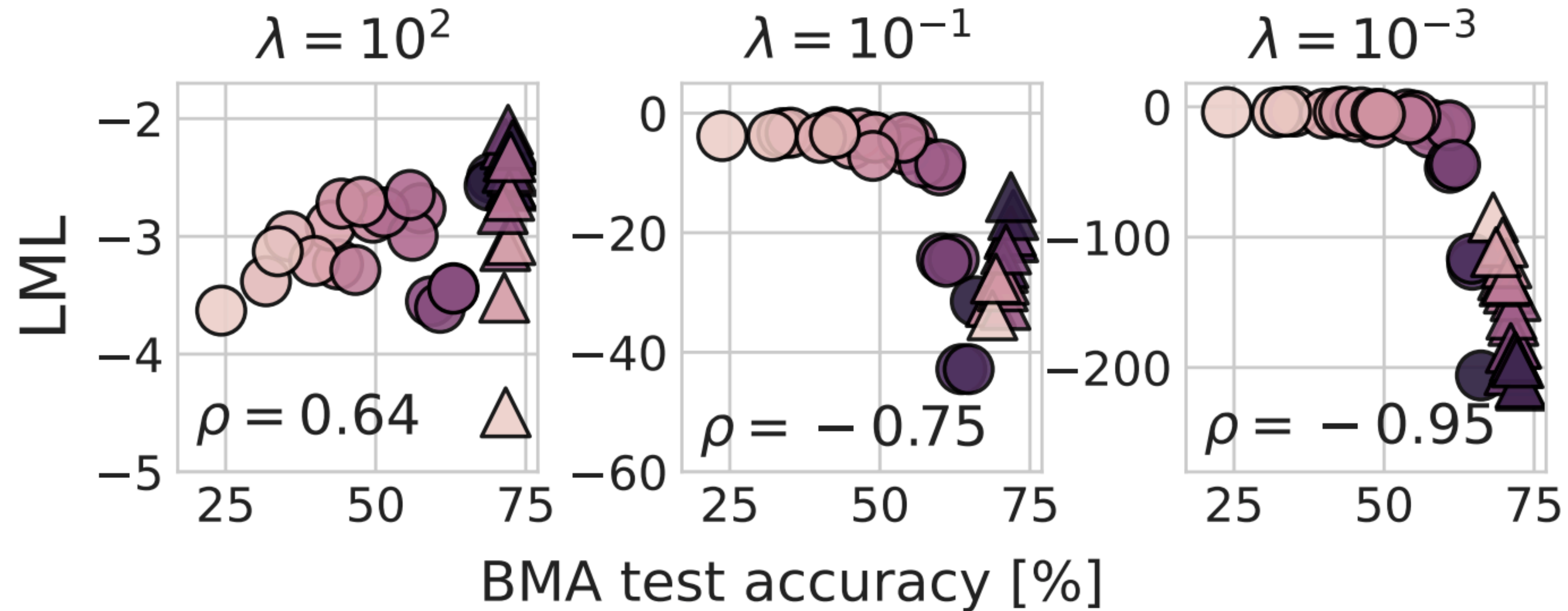


$$p_1(y|x, \theta) = \mathcal{N}(\theta_0 + \theta_1 x, \sigma) \quad p_2(y|x, \theta) = \mathcal{N}\left(\sum_{j=0}^3 \theta_j x, \sigma\right) \quad p_3(y|x, \theta) = \mathcal{N}\left(\sum_{j=0}^{10^4} \theta_j x, \sigma\right)$$

Does Marginal Likelihood work for NNs?

The answer is “not yet”.

Marginal Likelihood in NNs (Lotfi et al. 2023)



- Marginal Likelihood is poorly correlated with generalisation

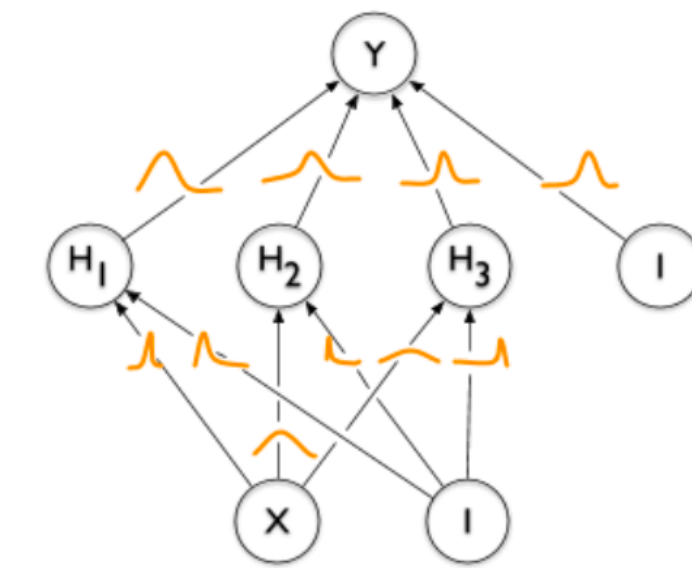
A Generalisation Analysis of Bayesian Machine Learning

PAC-Bayesian Analysis

PAC-Bayesian Bounds

(Alquier et al. 2016, Germain et al. 2016, Masegosa 2020)

$$\theta \sim \rho(\theta)$$



- Notation:

- $L(\theta)$ is the **expected log-loss**, $L(\theta) = \mathbb{E}_{\nu(\mathbf{y}, \mathbf{x})}[-\ln p(\mathbf{y}|\mathbf{x}, \theta)]$.
- $\hat{L}(\theta, D)$ is the **empirical log-loss**, $\hat{L}(\theta, D) = -\frac{1}{n} \ln p(D|\theta)$
- $B(\rho)$ is the **Bayes loss** of the posterior ρ , $B(\rho) = \mathbb{E}_{\nu(\mathbf{y}, \mathbf{x})}[-\ln \mathbb{E}_{\rho}[p(\mathbf{y}|\mathbf{x}, \theta)]]$

- For any prior $\pi(\theta)$ independent of D and any $\lambda > 0$, for all ρ simultaneously,

$$\underbrace{B(\rho)}_{\text{Bayes Loss}} \overset{\text{Jensen Inequality}}{\leq} \underbrace{\mathbb{E}_{\rho}[L(\theta)]}_{\text{Gibbs Loss}} \overset{\text{w.p. } (1-\delta)}{\lesssim} \underbrace{\mathbb{E}_{\rho}[\hat{L}(\theta, D)] + \frac{KL(\rho, \pi) + \phi_{\nu}^{\pi}(\lambda) + \ln \frac{1}{\delta}}{\lambda n}}_{\text{PAC-Bayes bound}}$$

- $\phi_{\nu}^{\pi}(\lambda)$ is a **log moment generating function**.

PAC-Bayesian Bounds

(Alquier et al. 2016, Germain et al. 2016, Masegosa 2020)

$$\theta \sim \rho(\theta)$$

- For any prior $\pi(\theta)$ independent of D and any $\lambda > 0$, for all ρ simultaneously,

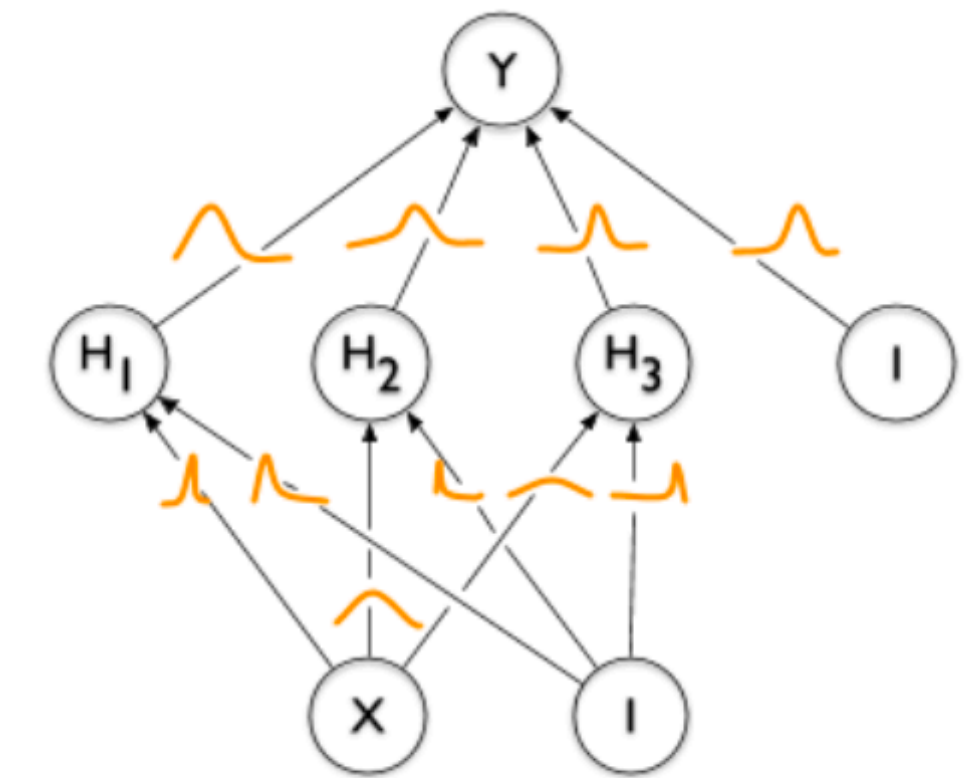
$$\underbrace{B(\rho)}_{\text{Bayes Loss}} \stackrel{\text{Jensen Inequality}}{\leq} \underbrace{\mathbb{E}_{\rho}[L(\theta)]}_{\text{Gibbs Loss}} \stackrel{\text{w.p. } (1-\delta)}{\lesssim} \underbrace{\mathbb{E}_{\rho}[\hat{L}(\theta, D)] + \frac{KL(\rho, \pi) + \phi_{\nu}^{\pi}(\lambda) + \ln \frac{1}{\delta}}{\lambda n}}_{\text{PAC-Bayes bound}}$$

- Which is the optimal density $\rho(\theta)$?,

$$\rho^*(\theta|D) = \arg \min_{\rho} \underbrace{\mathbb{E}_{\rho}[\hat{L}(\theta, D)] + \frac{KL(\rho, \pi) + \phi_{\nu}^{\pi}(\lambda) + \ln \frac{1}{\delta}}{\lambda n}}_{\text{PAC-Bayes bound}}$$

- $\rho^*(\theta|D)$ is the **Tempered Bayesian posterior**,

$$\rho^*(\theta|D) \propto p(D|\theta)^{\lambda} \pi(\theta)$$



PAC-Bayesian Bounds

(Alquier et al. 2016, Germain et al. 2016, Masegosa 2020)

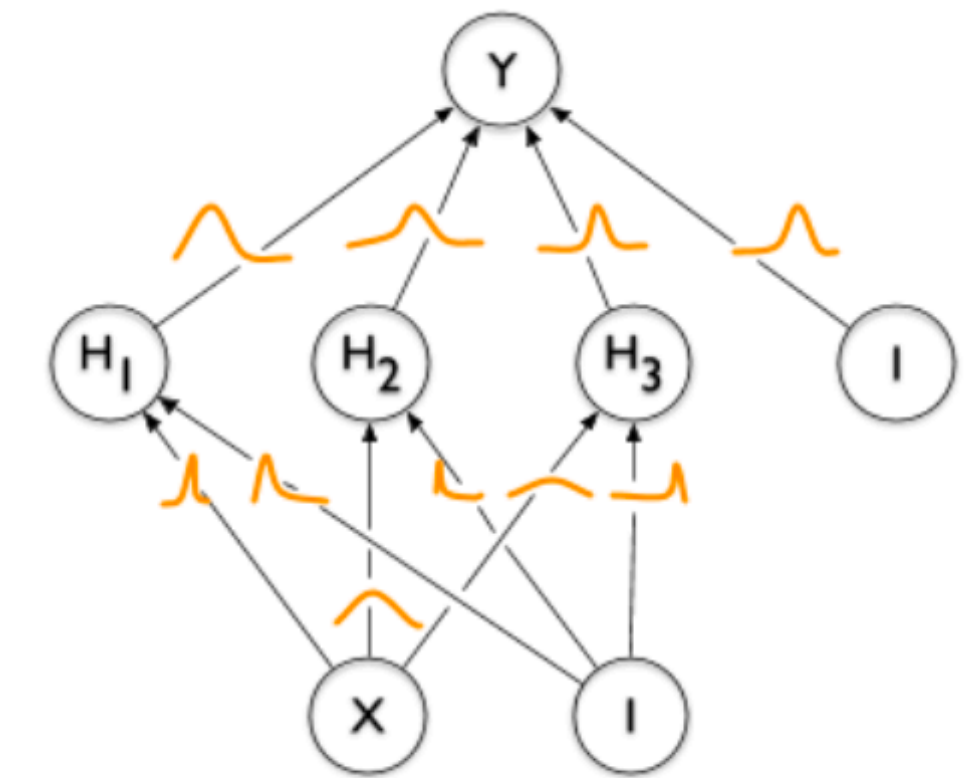
- For any prior $\pi(\theta)$ independent of D and any $\lambda > 0$, for all ρ simultaneously,

$$\underbrace{B(\rho)}_{\text{Bayes Loss}} \stackrel{\text{Jensen Inequality}}{\leq} \underbrace{\mathbb{E}_{\rho}[L(\theta)]}_{\text{Gibbs Loss}} \stackrel{\text{w.p. } (1-\delta)}{\lesssim} \underbrace{\mathbb{E}_{\rho}[\hat{L}(\theta, D)] + \frac{KL(\rho, \pi) + \phi_{\nu}(n, \lambda) + \ln \frac{1}{\delta}}{\lambda n}}_{\text{PAC-Bayes bound}}$$

- Which is the optimal density $\rho(\theta)$?,

$$\rho^*(\theta|D) = \arg \min_{\rho, \lambda \in \Lambda} \underbrace{\mathbb{E}_{\rho}[\hat{L}(\theta, D)] + \frac{KL(\rho, \pi) + \phi_{\nu}^{\pi}(\lambda) + \ln \frac{|\Lambda|}{\delta}}{\lambda n}}_{\text{PAC-Bayes bound}}$$

$$\theta \sim \rho(\theta)$$



PAC-Bayesian Bounds

(Alquier et al. 2016, Germain et al. 2016, Masegosa 2020)

- For any prior $\pi(\theta)$ independent of D and any $\lambda > 0$, for all ρ simultaneously,

$$\underbrace{B(\rho)}_{\text{Bayes Loss}} \stackrel{\text{Jensen Inequality}}{\leq} \underbrace{\mathbb{E}_{\rho}[L(\theta)]}_{\text{Gibbs Loss}} \stackrel{\text{w.p. } (1-\delta)}{\lesssim} \underbrace{\mathbb{E}_{\rho}[\hat{L}(\theta, D)] + \frac{KL(\rho, \pi) + \phi_{\nu}(n, \lambda) + \ln \frac{1}{\delta}}{\lambda n}}_{\text{PAC-Bayes bound}}$$

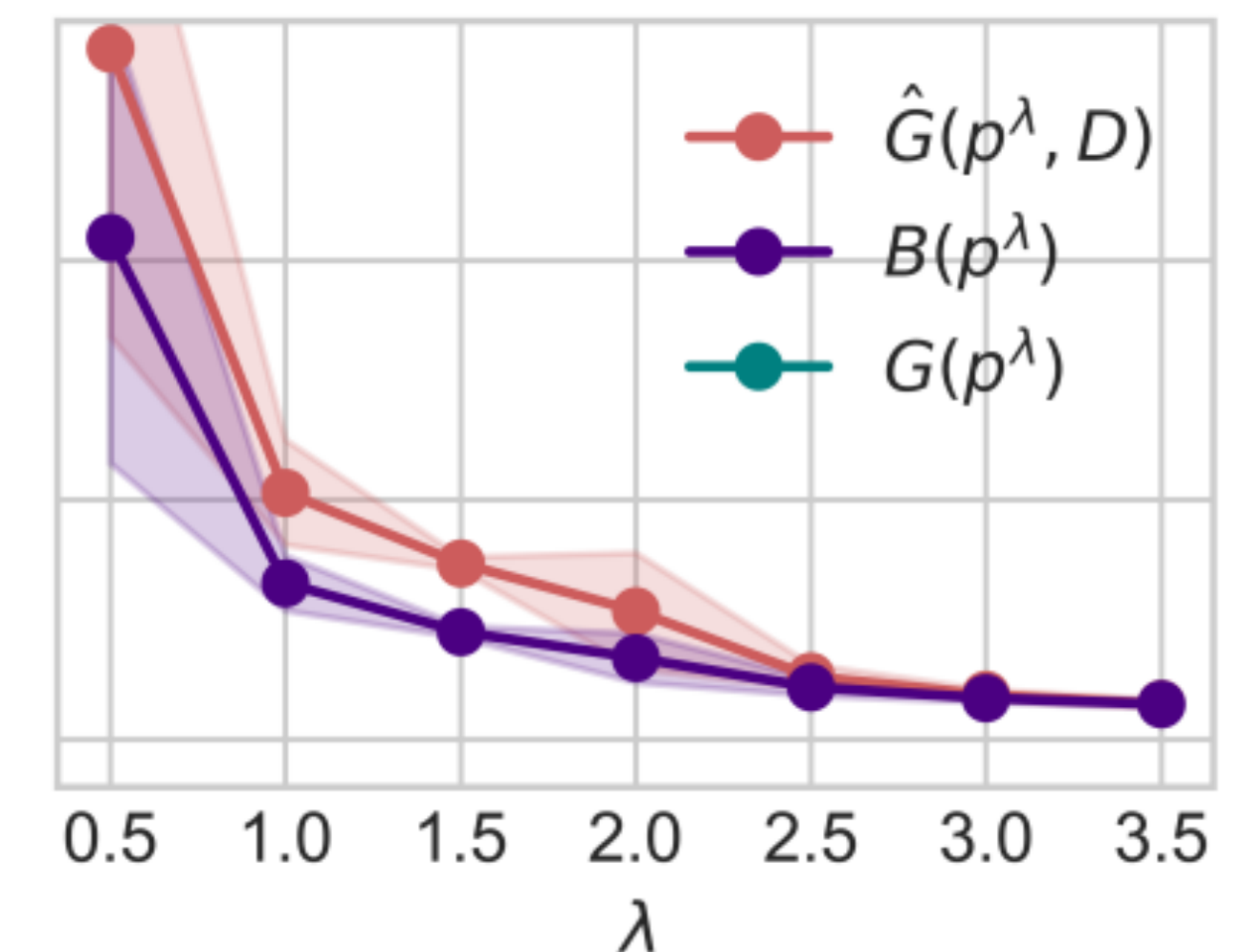
- Which is the optimal density $\rho(\theta)$?,

$$\rho^*(\theta|D) = \arg \min_{\lambda \in \Lambda} \underbrace{\mathbb{E}_{p^{\lambda}}[\hat{L}(\theta, D)] + \frac{KL(p^{\lambda}, \pi) + \phi_{\nu}^{\pi}(\lambda) + \ln \frac{|\Lambda|}{\delta}}{\lambda n}}_{\text{PAC-Bayes bound}}$$

- $p^{\lambda}(\theta|D)$ is the **Tempered Bayesian posterior**,

$$p^{\lambda}(\theta|D) \propto p(D|\theta)^{\lambda} \pi(\theta)$$

Cold Posterior Effect



(Zhang et al. 2023)

Marginal Likelihood and Generalization

- The Log Marginal Likelihood decomposes as:

$$-\frac{1}{n} \ln p(D) = \mathbb{E}_{\mathbf{p}}[\hat{L}(D, \boldsymbol{\theta})] + \frac{KL(\mathbf{p}, \pi)}{n}$$

- For the Bayesian posterior we have $\mathbf{p}(\boldsymbol{\theta}|D)$,

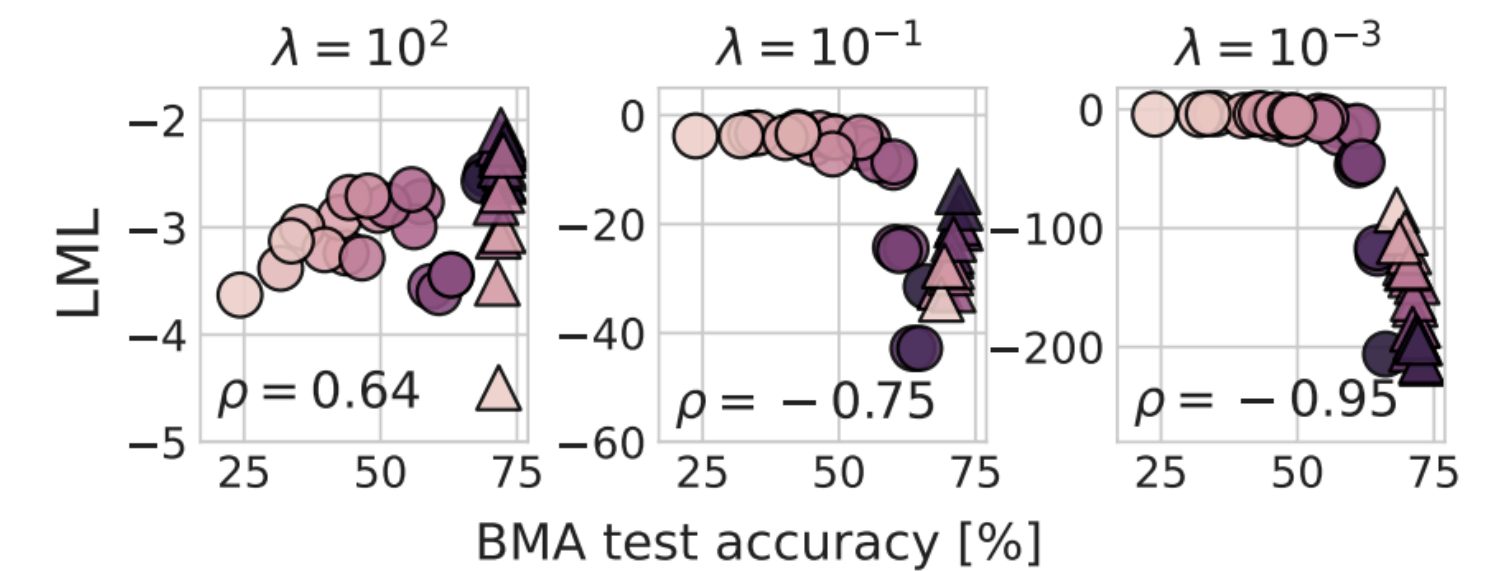
$$\underbrace{B(\mathbf{p})}_{\text{Bayes Loss}} \stackrel{\text{Jensen Inequality}}{\leq} \underbrace{\mathbb{E}_{\mathbf{p}}[L(\boldsymbol{\theta})]}_{\text{Gibbs Loss}} \stackrel{\text{w.p. } (1-\delta)}{\lesssim} \underbrace{\mathbb{E}_{\mathbf{p}}[\hat{L}(\boldsymbol{\theta}, D)] + \phi^*\left(\frac{KL(\mathbf{p}, \pi) + \ln \frac{|\Lambda|}{\delta}}{n}\right)}_{\text{PAC-Bayes bound}}$$

- $\phi^*(a)$ is an approximation of the the inverse Legendre transform of ϕ_{ν}^{π} .

- The Log Marginal would be a good proxy of generalization when $\phi^*(a) = a$.

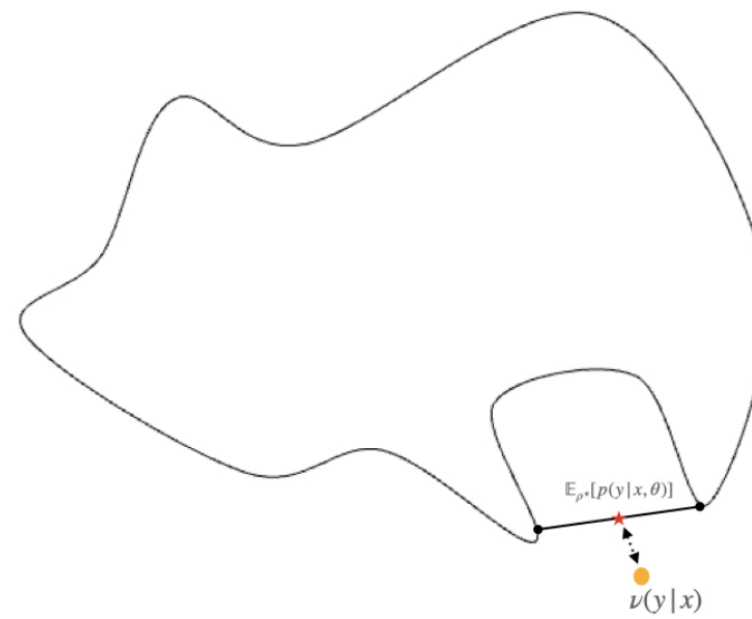
$$\underbrace{B(\mathbf{p})}_{\text{Bayes Loss}} \stackrel{\text{Jensen Inequality}}{\leq} \underbrace{\mathbb{E}_{\mathbf{p}}[L(\boldsymbol{\theta})]}_{\text{Gibbs Loss}} \stackrel{\text{w.p. } (1-\delta)}{\lesssim} \underbrace{\mathbb{E}_{\mathbf{p}}[\hat{L}(\boldsymbol{\theta}, D)] + \frac{KL(\mathbf{p}, \pi) + \ln \frac{|\Lambda|}{\delta}}{n}}_{\text{PAC-Bayes bound}}$$

LML vs Generalisation



Model Misspecification (Masegosa, 2020)

$$\arg \min_{\rho} \underbrace{B(\rho)}_{\text{Bayes Loss}} \neq \arg \min_{\rho} \underbrace{\mathbb{E}_{\rho}[L(\theta)]}_{\text{Gibbs Loss}}$$



- For any prior $\pi(\theta)$ independent of D and any $\lambda > 0$, for all ρ simultaneously,

$$\underbrace{B(\rho)}_{\text{Bayes Loss}} \stackrel{\text{Jensen Inequality}}{\leq} \underbrace{\mathbb{E}_{\rho}[L(\theta)]}_{\text{Gibbs Loss}} \stackrel{\text{w.p. } (1-\delta)}{\lesssim} \underbrace{\mathbb{E}_{\rho}[\hat{L}(\theta, D)] + \frac{KL(\rho, \pi) + \phi_{\nu}(n, \lambda) + \ln \frac{1}{\delta}}{\lambda n}}_{\text{PAC-Bayes bound}}$$

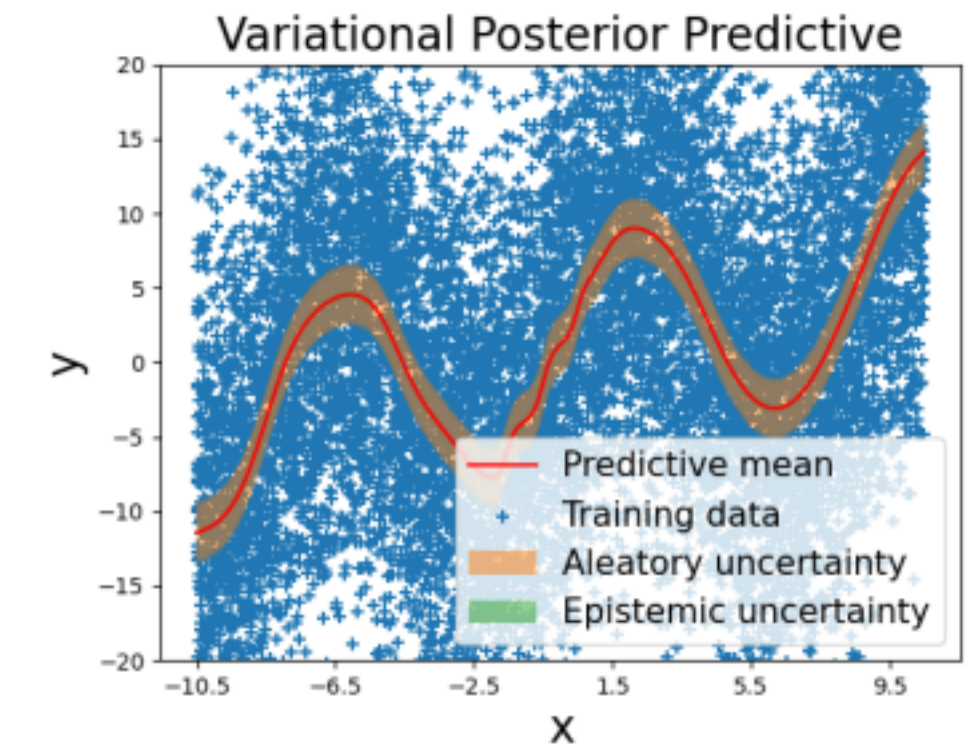
The Bayesian Posterior is not optimal under model misspecification

Bayesian Learning under Model Misspecification

- For any prior $\pi(\theta)$ independent of D and any $\lambda > 0$, for all ρ simultaneously,

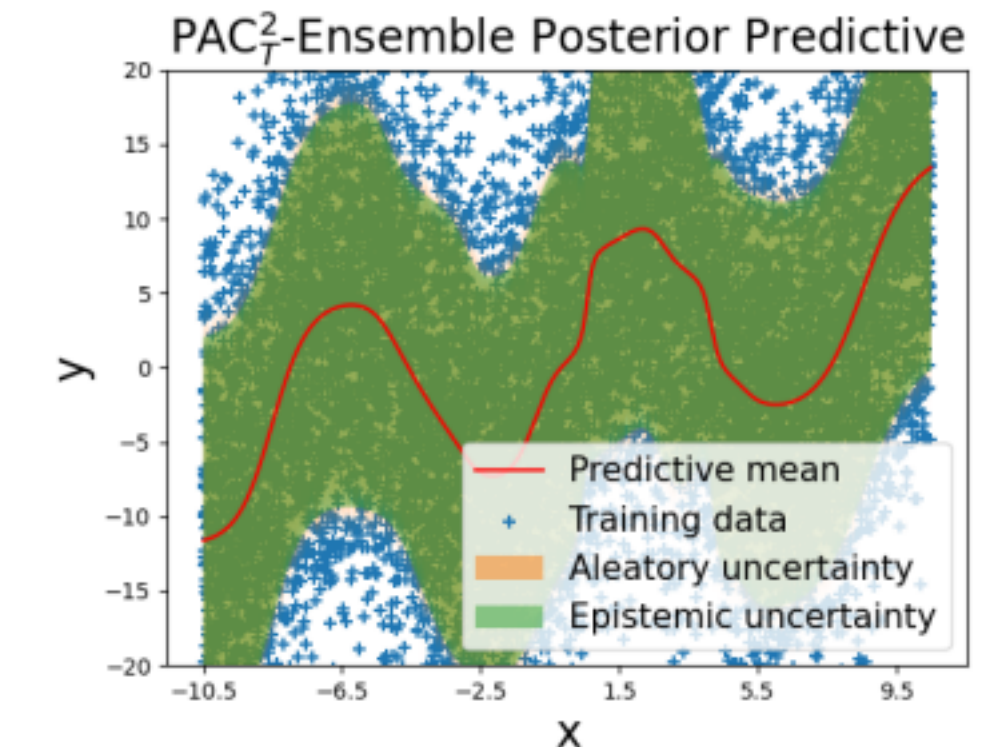
$$\underbrace{B(\rho)}_{\text{Bayes Loss}} \leq \underbrace{\mathbb{E}_{\theta^{(m)} \sim \rho}[L_P(\theta^{(m)})]}_{\text{Extended Gibbs Loss}} \stackrel{\text{w.p. } (1-\delta)}{\lesssim} \underbrace{\mathbb{E}_{\theta^{(m)} \sim \rho}[\hat{L}_P(\theta^{(m)}, D)] + m \frac{KL(\rho, \pi)}{\lambda n} + \frac{\phi(\cdot)}{n}}_{\text{PAC}^m\text{-Bayes bound}}$$

Empirical Loss



- $L_P(\theta^{(m)})$ with $m > 1$ induces tighter bounds:

$$\underbrace{B(\rho)}_{\text{Bayes Loss}} \leq \underbrace{\mathbb{E}_{\theta^{(m)} \sim \rho}[L_P(\theta^{(m)})]}_{\text{Extended Gibbs Loss}} \leq \underbrace{\mathbb{E}_{\rho}[L(\theta)]}_{\text{Gibbs Loss}}$$

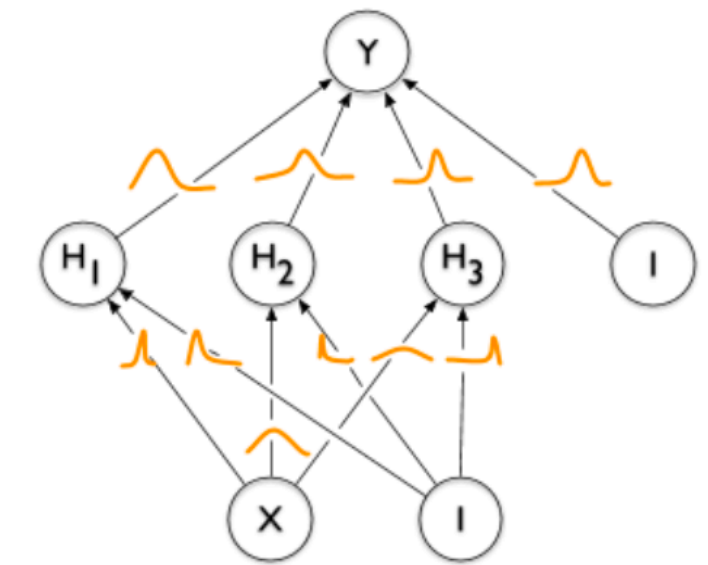


- (Masegosa, 2020) use a loss based on **second-order Jensen inequalities**.
- (Futami et al., 2021) use an even **tighter second-order Jensen inequality**.
- (Morningstar et al., 2022) use a **multi-sample bound** which is **arbitrary tight**.
- (Zechin et al., 2022) adapts this scheme to **robust** log-losses (t-logarithm).

PAC-Bayesian Bounds

$$\theta \sim p(\theta | D)$$

- Many open questions in Bayesian Machine learning:
 - What is the Cold Posterior Effect?
 - When does the marginal likelihood correlate with generalization?
 - How to design better Bayesian priors?
 - How to build better learning algorithms?
- PAC-Bayesian bounds provide a generalization perspective of Bayesian methods.
 - Potential to achieve a better understanding of BML.
 - Potential to design better learning algorithms.
- Open Issues with PAC-Bayes:
 - The log-loss used in Bayesian learning is unbounded.
 - Are current PAC-Bayes bounds tight?
 - Do we need distribution-dependent PAC-Bayes bounds?
 - Should we move to information theoretic bounds?
 - Derive a *killer* learning algorithm.



Thanks for you attention!! :D