



Signal Processing

Lecture 14: Bayesian Inference

Konstantinos Chatzilygeroudis - costashatz@upatras.gr

Department of Electrical and Computer Engineering
University of Patras

Template made by Panagiotis Papagiannopoulos



Bayesian Inference: Core Idea

Unknown parameters as random variables

θ is uncertain $\Rightarrow$ treat θ as a random variable and encode belief with a prior $p(\theta)$.

Observations are random too

$\mathbf{Y}$ is generated from a model conditioned on $(\mathbf{X}, \theta)$,

so we need a probabilistic data model.

Goal: infer θ from data

Given $(\mathbf{X}, \mathbf{Y})$, compute a distribution over θ .

Why Bayesian?

- Gives *uncertainty* about parameters.
- Naturally incorporates *prior knowledge* / *regularization*.
- Enables *model comparison* via evidence.

Data model (likelihood)

$p(\mathbf{Y} \mid \mathbf{X}, \boldsymbol{\theta})$ describes how data is generated given $\boldsymbol{\theta}$.

Interpretation

- A function of $\mathbf{Y}$ for fixed $\boldsymbol{\theta}$: a probability model.
- A function of $\boldsymbol{\theta}$ for fixed data: a measure of *how compatible* $\boldsymbol{\theta}$ is with the observations.

Example (Gaussian noise)

$$y_i = f(x_i; \boldsymbol{\theta}) + \varepsilon_i, \quad \varepsilon_i \sim \mathcal{N}(0, \sigma^2)$$

$$\Rightarrow p(\mathbf{Y} \mid \mathbf{X}, \boldsymbol{\theta}) = \prod_{i=1}^N \mathcal{N}(y_i; f(x_i; \boldsymbol{\theta}), \sigma^2).$$

Bayes' Rule and Posterior

Bayes' rule: posterior distribution

$$p(\theta \mid \mathbf{Y}, \mathbf{X}) = \frac{p(\mathbf{Y} \mid \mathbf{X}, \theta) p(\theta)}{p(\mathbf{Y} \mid \mathbf{X})}.$$

Evidence / Marginal likelihood

$$p(\mathbf{Y} \mid \mathbf{X}) = \int p(\mathbf{Y} \mid \mathbf{X}, \theta) p(\theta) d\theta.$$

Key message

Posterior $\propto$ Likelihood $\times$ Prior.
--

Interpretation:

- **Prior** $p(\theta)$: knowledge / belief *before* seeing data.
- **Likelihood** $p(\mathbf{Y} | \mathbf{X}, \theta)$: how likely data is under θ .
- **Posterior** $p(\theta | \mathbf{Y}, \mathbf{X})$: updated belief after seeing data.
- **Evidence** $p(\mathbf{Y} | \mathbf{X})$: normalizes and measures model plausibility.

Intuition

- The likelihood pulls θ toward values that explain data.
- The prior pulls θ toward values we believed beforehand.
- The posterior balances the two.

Why do we care about evidence?

$$p(\mathbf{Y} | \mathbf{X}) = \int p(\mathbf{Y} | \mathbf{X}, \theta) p(\theta) d\theta.$$

Interpretation

- It is the probability of observing the data under the model *before* knowing parameters.
- It averages likelihood over the prior: good models explain data *for many plausible parameters*.

Model comparison (Bayes factor) For two models $\mathcal{M}_1, \mathcal{M}_2$:

$$\frac{p(\mathbf{Y} | \mathbf{X}, \mathcal{M}_1)}{p(\mathbf{Y} | \mathbf{X}, \mathcal{M}_2)} \quad \text{compares how well each model explains the data.}$$

Occam's razor effect: complex models are penalized unless strongly supported by data.

Bayesian inference gives a *distribution* over parameters. Often we want a single estimate $\hat{\theta}$.

Decision-theoretic view Choose $\hat{\theta}$ to minimize posterior expected loss:

$$\hat{\theta} = \arg \min_{\tilde{\theta}} \mathbb{E} \left[\mathcal{L}(\theta, \tilde{\theta}) \mid \mathbf{Y}, \mathbf{X} \right].$$

Different losses $\mathcal{L}$ lead to different estimators:

- 0–1 loss $\Rightarrow$ MAP.
- Squared error loss $\Rightarrow$ MMSE.
- Absolute error loss $\Rightarrow$ posterior median.

Given the posterior

$$p(\theta \mid \mathbf{Y}, \mathbf{X}) \propto p(\mathbf{Y} \mid \mathbf{X}, \theta) p(\theta),$$

1) MAP (Maximum A Posteriori): posterior mode

$$\hat{\theta}_{\text{MAP}} = \arg \max_{\theta} p(\theta \mid \mathbf{Y}, \mathbf{X}) = \arg \max_{\theta} p(\mathbf{Y} \mid \mathbf{X}, \theta) p(\theta).$$

Equivalent minimization form

$$\hat{\theta}_{\text{MAP}} = \arg \min_{\theta} \left(-\log p(\mathbf{Y} \mid \mathbf{X}, \theta) - \log p(\theta) \right).$$

Interpretation: maximize posterior $\Leftrightarrow$ minimize *data misfit + regularization*.

The optimization over the MAP criterion results to the parameters, $\hat{\theta}_{\text{MAP}}$, that are the *posterior distribution's mode value*. In other words, we find *the peak of the posterior probability density function* (in general, we may have multiple peaks!).

MAP, Priors, and Regularization

Example: Gaussian prior

$$p(\boldsymbol{\theta}) = \mathcal{N}(\mathbf{0}, \alpha^{-1}I) \quad \Rightarrow \quad -\log p(\boldsymbol{\theta}) \propto \frac{\alpha}{2} \|\boldsymbol{\theta}\|_2^2.$$

With Gaussian noise likelihood:

$$-\log p(\mathbf{Y} \mid \mathbf{X}, \boldsymbol{\theta}) \propto \frac{1}{2\sigma^2} \|\mathbf{Y} - f(\mathbf{X}; \boldsymbol{\theta})\|_2^2$$

MAP becomes regularized least squares

$$\hat{\boldsymbol{\theta}}_{\text{MAP}} = \arg \min_{\boldsymbol{\theta}} \frac{1}{2\sigma^2} \|\mathbf{Y} - f(\mathbf{X}; \boldsymbol{\theta})\|_2^2 + \frac{\alpha}{2} \|\boldsymbol{\theta}\|_2^2.$$

So: priors correspond to regularizers.

MLE as a Special Case of MAP

If $p(\boldsymbol{\theta}) \propto \text{const}$ (flat / uninformative prior), then

$$\hat{\boldsymbol{\theta}}_{\text{MAP}} = \arg \max_{\boldsymbol{\theta}} p(\mathbf{Y} \mid \mathbf{X}, \boldsymbol{\theta}) = \hat{\boldsymbol{\theta}}_{\text{MLE}}.$$

Interpretation

- MAP uses *prior + data*.
- MLE uses *data only*.
- With enough data, MAP and MLE often get close.

2) MMSE (Minimum Mean Square Error): posterior mean

Definition (squared-error loss):

$$\hat{\theta}_{\text{MMSE}} = \arg \min_{\tilde{\theta}} \mathbb{E} \left[\|\theta - \tilde{\theta}\|_2^2 \mid \mathbf{Y}, \mathbf{X} \right].$$

where θ is the *true (latent) parameter* treated as a random vector, and $\tilde{\theta}$ is a *deterministic estimate* we choose to minimize posterior risk.

Result: the minimizer is the posterior mean

$$\hat{\theta}_{\text{MMSE}} = \mathbb{E}[\theta \mid \mathbf{Y}, \mathbf{X}] = \int \theta p(\theta \mid \mathbf{Y}, \mathbf{X}) d\theta.$$

Interpretation:

- MMSE averages all plausible parameters, weighted by posterior probability.
- Optimal when we care about *expected squared deviation* from the truth.
- More robust than MAP when the posterior is skewed or multimodal.

Why MMSE Equals the Posterior Mean? (Sketch)

We want to minimize

$$\mathbb{E}[\|\boldsymbol{\theta} - \tilde{\boldsymbol{\theta}}\|^2 \mid \mathbf{Y}, \mathbf{X}].$$

Expand around the posterior mean $\boldsymbol{\mu} = \mathbb{E}[\boldsymbol{\theta} \mid \mathbf{Y}, \mathbf{X}]$:

$$\|\boldsymbol{\theta} - \tilde{\boldsymbol{\theta}}\|^2 = \|(\boldsymbol{\theta} - \boldsymbol{\mu}) + (\boldsymbol{\mu} - \tilde{\boldsymbol{\theta}})\|^2 = \|\boldsymbol{\theta} - \boldsymbol{\mu}\|^2 + \|\boldsymbol{\mu} - \tilde{\boldsymbol{\theta}}\|^2 + 2(\boldsymbol{\theta} - \boldsymbol{\mu})^\top (\boldsymbol{\mu} - \tilde{\boldsymbol{\theta}}).$$

Take posterior expectation:

$$\mathbb{E}[\boldsymbol{\theta} - \boldsymbol{\mu} \mid \mathbf{Y}, \mathbf{X}] = \mathbf{0} \quad \Rightarrow \quad \text{cross term vanishes.}$$

Thus

$$\mathbb{E}[\|\boldsymbol{\theta} - \tilde{\boldsymbol{\theta}}\|^2 \mid \cdot] = \underbrace{\mathbb{E}[\|\boldsymbol{\theta} - \boldsymbol{\mu}\|^2 \mid \cdot]}_{\text{independent of } \tilde{\boldsymbol{\theta}}} + \|\boldsymbol{\mu} - \tilde{\boldsymbol{\theta}}\|^2,$$

minimized when $\tilde{\boldsymbol{\theta}} = \boldsymbol{\mu}$.

3) MMAE (Minimum Mean Absolute Error): posterior median

Definition (absolute-error loss):

$$\hat{\theta}_{\text{MMAE}} = \arg \min_{\tilde{\theta}} \mathbb{E} \left[\|\theta - \tilde{\theta}\|_1 \mid \mathbf{Y}, \mathbf{X} \right].$$

where θ is the *true (latent) parameter* treated as a random vector, and $\tilde{\theta}$ is a *deterministic estimate* we choose to minimize posterior risk.

Result: the minimizer is the posterior median. We assume $\|\cdot\|_1 = \sum_i |\cdot|$,

$$\hat{\theta}_{\text{MMAE}} = \left[\text{median}(p(\theta_1 \mid \mathbf{Y}, \mathbf{X})), \dots, \text{median}(p(\theta_d \mid \mathbf{Y}, \mathbf{X})) \right]^\top.$$

i.e., the estimator is the *componentwise posterior median*.

Interpretation:

- MMAE picks the most “central” posterior value in terms of *absolute deviation*.
- Optimal when we care about *expected absolute error* rather than squared error.
- More robust to outliers/heavy tails than MMSE (mean can be pulled by extremes).
- For skewed posteriors, median and mean can differ substantially.

MAP vs MMSE vs MMAE

MAP:

$$\hat{\theta}_{\text{MAP}} = \arg \max_{\theta} p(\theta \mid \mathbf{Y}, \mathbf{X}) \quad (\text{posterior mode})$$

MMSE:

$$\hat{\theta}_{\text{MMSE}} = \mathbb{E}[\theta \mid \mathbf{Y}, \mathbf{X}] \quad (\text{posterior mean})$$

MMAE:

$$\hat{\theta}_{\text{MMAE}} = \arg \min_{\tilde{\theta}} \mathbb{E}[\|\theta - \tilde{\theta}\|_1 \mid \mathbf{Y}, \mathbf{X}] \quad (\text{posterior median})$$

When are they equal? If $p(\theta \mid \mathbf{Y}, \mathbf{X})$ is Gaussian (or any symmetric unimodal pdf):

$$\text{mean} = \text{median} = \text{mode} \quad \Rightarrow \quad \hat{\theta}_{\text{MMSE}} = \hat{\theta}_{\text{MMAE}} = \hat{\theta}_{\text{MAP}}.$$

Practical intuition

- MAP picks the *single most probable* parameter.
- MMSE averages all plausible parameters with posterior weights.
- MMAE selects the *central* posterior value, robust to outliers.
- For skewed / heavy-tailed posteriors:

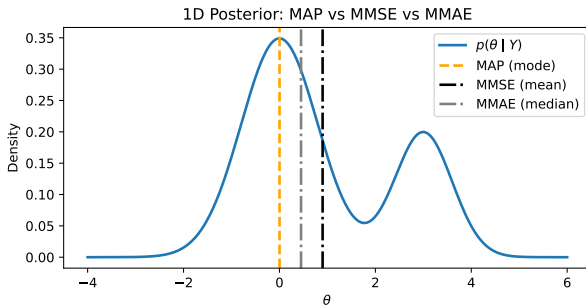
$$\text{mode} \neq \text{median} \neq \text{mean},$$

so the three estimators can differ substantially.

1D Illustration

Consider a posterior $p(\theta | Y)$ in 1D.

- **MAP** $\hat{\theta}_{\text{MAP}}$ is at the *peak*.
- **MMSE** $\hat{\theta}_{\text{MMSE}}$ is the *center of mass*.
- **MMAE** $\hat{\theta}_{\text{MMAE}}$ is the *median*.



Key idea: different estimators reflect different optimality criteria.

Maximum Likelihood Estimation (MLE)

Setup: we observe data $\mathbf{y}$ generated by a model parameterized by θ with likelihood

$$p(\mathbf{y} | \theta).$$

Goal (MLE): choose parameters that make the observed data most probable:

$$\hat{\theta}_{\text{MLE}} = \arg \max_{\theta} p(\mathbf{y} | \theta).$$

Log-likelihood (equivalent objective):

$$\mathcal{L}(\theta) = \log p(\mathbf{y} | \theta), \quad \hat{\theta}_{\text{MLE}} = \arg \max_{\theta} \mathcal{L}(\theta).$$

(Log is used because it turns products into sums and is numerically stable.)

If samples are i.i.d.:

$$p(\mathbf{y} | \theta) = \prod_{k=1}^N p(y_k | \theta) \quad \Rightarrow \quad \mathcal{L}(\theta) = \sum_{k=1}^N \log p(y_k | \theta).$$

How do we solve it?

- If $\mathcal{L}(\theta)$ is concave and simple $\Rightarrow$ closed form.
- Otherwise $\Rightarrow$ iterative optimization (Newton / Gauss-Newton methods / gradient descent).

Key idea: MLE depends on the assumed noise/data distribution in $p(\mathbf{y} | \theta)$.

Maximum Likelihood Estimation: i.i.d. Data

Data: N i.i.d. input–output samples

$$\{(\mathbf{x}_i, \mathbf{y}_i)\}_{i=1}^N, \quad \mathbf{X} = \begin{bmatrix} \mathbf{x}_1^T \\ \vdots \\ \mathbf{x}_N^T \end{bmatrix}, \quad \mathbf{Y} = \begin{bmatrix} \mathbf{y}_1^T \\ \vdots \\ \mathbf{y}_N^T \end{bmatrix}.$$

Model: conditional distribution of outputs given inputs and parameters with **i.i.d. assumption**:

$$p(\mathbf{Y} | \mathbf{X}, \boldsymbol{\theta}) = \prod_{i=1}^N p(\mathbf{y}_i | \mathbf{x}_i, \boldsymbol{\theta}).$$

Log-likelihood: (turn product into sum)

$$\mathcal{L}(\boldsymbol{\theta}) = \log p(\mathbf{Y} | \mathbf{X}, \boldsymbol{\theta}) = \sum_{i=1}^N \log p(\mathbf{y}_i | \mathbf{x}_i, \boldsymbol{\theta}).$$

MLE objective:

$$\hat{\boldsymbol{\theta}}_{\text{MLE}} = \arg \max_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}) = \arg \max_{\boldsymbol{\theta}} \sum_{i=1}^N \log p(\mathbf{y}_i | \mathbf{x}_i, \boldsymbol{\theta}).$$

Note: the form of $p(\mathbf{y}_i | \mathbf{x}_i, \boldsymbol{\theta})$ depends on the assumed noise / data model.

MLE for Linear Models with Gaussian Noise

Data: N i.i.d. pairs $\{(\mathbf{x}_i, y_i)\}_{i=1}^N$ with $\mathbf{x}_i \in \mathbb{R}^P$, $y_i \in \mathbb{R}$.

Linear model (scalar output):

$$y_i = \boldsymbol{\theta}^T \mathbf{x}_i + w_i, \quad \boldsymbol{\theta} \in \mathbb{R}^P.$$

Gaussian noise assumption:

$$w_i \sim \mathcal{N}(0, \sigma^2) \quad \Rightarrow \quad p(y_i | \mathbf{x}_i, \boldsymbol{\theta}) = \mathcal{N}(\boldsymbol{\theta}^T \mathbf{x}_i, \sigma^2).$$

Log-likelihood:

$$\mathcal{L}(\boldsymbol{\theta}) = \log p(\mathbf{Y} | \mathbf{X}, \boldsymbol{\theta}) = \sum_{i=1}^N \log p(y_i | \mathbf{x}_i, \boldsymbol{\theta}).$$

For Gaussian noise:

$$\mathcal{L}(\boldsymbol{\theta}) = \sum_{i=1}^N \left[-\frac{1}{2} \log(2\pi\sigma^2) - \frac{1}{2\sigma^2} (y_i - \boldsymbol{\theta}^T \mathbf{x}_i)^2 \right].$$

MLE (maximize log-likelihood):

$$\hat{\boldsymbol{\theta}}_{\text{MLE}} = \arg \max_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}) \iff \underbrace{\arg \min_{\boldsymbol{\theta}} \sum_{i=1}^N (y_i - \boldsymbol{\theta}^T \mathbf{x}_i)^2}_{\text{Linear Least Squares!}}$$

Bayesian Linear Regression (BLR): Problem Setting

We observe training pairs $\{(\mathbf{x}_i, y_i)\}_{i=1}^N$ and assume

$$y_i = \boldsymbol{\theta}^\top \boldsymbol{\phi}(\mathbf{x}_i) + w_i, \quad w_i \sim \mathcal{N}(0, \sigma_w^2).$$

Feature map:

$$\boldsymbol{\phi} : \mathbb{R}^d \rightarrow \mathbb{R}^P \quad (\text{could be polynomial, Fourier, RBF, etc.})$$

Matrix form:

$$\mathbf{y} = \boldsymbol{\Phi} \boldsymbol{\theta} + \mathbf{w}, \quad \mathbf{w} \sim \mathcal{N}(\mathbf{0}, \sigma_w^2 \mathbf{I}), \quad \boldsymbol{\Phi} = \begin{bmatrix} \boldsymbol{\phi}(\mathbf{x}_1)^\top \\ \vdots \\ \boldsymbol{\phi}(\mathbf{x}_N)^\top \end{bmatrix}.$$

Goal: infer $\boldsymbol{\theta}$ and predict y_* with uncertainty.

Parameters are random variables

$\theta \sim p(\theta)$ (encode uncertainty / prior knowledge).

Gaussian prior (conjugate choice):

$$p(\theta) = \mathcal{N}(\mu_0, \Sigma_0).$$

Interpretation

- μ_0 : prior guess of weights.
- Σ_0 : prior uncertainty / correlation between weights.
- If $\Sigma_0 = \alpha^{-1}I$: isotropic prior $\Rightarrow$ shrink weights toward zero.

Given θ , the data likelihood is Gaussian:

$$p(\mathbf{y} \mid \Phi, \theta) = \mathcal{N}(\Phi\theta, \sigma_w^2 I).$$

Log-likelihood

$$\log p(\mathbf{y} \mid \Phi, \theta) = -\frac{N}{2} \log(2\pi\sigma_w^2) - \frac{1}{2\sigma_w^2} \|\mathbf{y} - \Phi\theta\|^2.$$

Connection

Gaussian noise $\Rightarrow$ Least Squares $\Leftrightarrow$ MLE.

BLR: Prior Predictive Distribution (Before Data)

Prior prediction at new input $\mathbf{x}_*$:

$$p(y_* | \mathbf{x}_*) = \int p(y_* | \mathbf{x}_*, \boldsymbol{\theta}) p(\boldsymbol{\theta}) d\boldsymbol{\theta}.$$

Since (likelihood)

$$p(y_* | \mathbf{x}_*, \boldsymbol{\theta}) = \mathcal{N}(\phi_*^\top \boldsymbol{\theta}, \sigma_w^2), \quad \phi_* \equiv \phi(\mathbf{x}_*),$$

introduce the **noiseless latent prediction**

$$z_* \equiv \phi_*^\top \boldsymbol{\theta}.$$

Because $\boldsymbol{\theta} \sim \mathcal{N}(\boldsymbol{\mu}_0, \boldsymbol{\Sigma}_0)$ and z_* is a linear function of $\boldsymbol{\theta}$,

$$z_* \sim \mathcal{N}(\phi_*^\top \boldsymbol{\mu}_0, \phi_*^\top \boldsymbol{\Sigma}_0 \phi_*).$$

Now $y_* = z_* + w_*$ with $w_* \sim \mathcal{N}(0, \sigma_w^2)$ independent, so the sum of Gaussians is Gaussian:

$$p(y_* | \mathbf{x}_*) = \mathcal{N}(\phi_*^\top \boldsymbol{\mu}_0, \phi_*^\top \boldsymbol{\Sigma}_0 \phi_* + \sigma_w^2).$$

Intuition:

- We predict by averaging over all plausible $\boldsymbol{\theta}$ under the prior.
- Predictive variance = parameter uncertainty + measurement noise.
- No data yet: predictions are driven only by the prior.

BLR: Generating Synthetic Data from the Prior

Prior sampling idea: before observing any data, we can *sample plausible models* from the prior and generate synthetic observations.

Step 1: sample parameters

$$\boldsymbol{\theta}^{(s)} \sim \mathcal{N}(\boldsymbol{\mu}_0, \boldsymbol{\Sigma}_0), \quad s = 1, \dots, S.$$

Step 2: form noiseless functions

$$z_*^{(s)} = \boldsymbol{\phi}(\mathbf{x}_*)^\top \boldsymbol{\theta}^{(s)}.$$

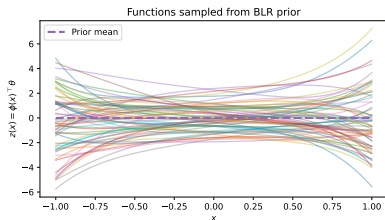
Step 3: add observation noise

$$y_*^{(s)} = z_*^{(s)} + w_*^{(s)}, \quad w_*^{(s)} \sim \mathcal{N}(0, \sigma_w^2).$$

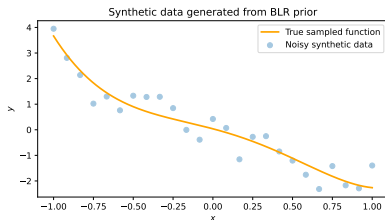
Why useful?

- Visualize what functions the prior considers plausible.
- Create toy datasets for demos / assignments.
- Stress-test inference: does BLR recover $\boldsymbol{\theta}$ under its own assumptions?
- Debug priors/features before collecting real data.

BLR: Prior Sampling and Synthetic Data



Prior function samples



One synthetic dataset

- Left: plausible functions before seeing data.
- Right: sample one function, add noise to generate training points.

BLR: Posterior is Gaussian (Conjugacy)

Posterior

$$p(\boldsymbol{\theta} \mid \mathbf{y}, \boldsymbol{\Phi}) \propto p(\mathbf{y} \mid \boldsymbol{\Phi}, \boldsymbol{\theta}) p(\boldsymbol{\theta}).$$

Because both prior and likelihood are Gaussian in $\boldsymbol{\theta}$,

$$p(\boldsymbol{\theta} \mid \mathbf{y}, \boldsymbol{\Phi}) = \mathcal{N}(\boldsymbol{\mu}_N, \boldsymbol{\Sigma}_N).$$

Meaning: data updates our belief, but keeps a Gaussian form.

BLR: Derivation (Complete the Square)

Start from log posterior up to constants:

$$\log p(\boldsymbol{\theta} \mid \mathbf{y}, \boldsymbol{\Phi}) = -\frac{1}{2\sigma_w^2} \|\mathbf{y} - \boldsymbol{\Phi}\boldsymbol{\theta}\|^2 - \frac{1}{2}(\boldsymbol{\theta} - \boldsymbol{\mu}_0)^\top \boldsymbol{\Sigma}_0^{-1}(\boldsymbol{\theta} - \boldsymbol{\mu}_0).$$

Expand the quadratic terms in $\boldsymbol{\theta}$:

$$= -\frac{1}{2} \left[\boldsymbol{\theta}^\top \left(\boldsymbol{\Sigma}_0^{-1} + \frac{1}{\sigma_w^2} \boldsymbol{\Phi}^\top \boldsymbol{\Phi} \right) \boldsymbol{\theta} - 2\boldsymbol{\theta}^\top \left(\boldsymbol{\Sigma}_0^{-1} \boldsymbol{\mu}_0 + \frac{1}{\sigma_w^2} \boldsymbol{\Phi}^\top \mathbf{y} \right) \right].$$

Identify Gaussian form $\Rightarrow$ read off posterior mean/covariance.

$$\text{Let } A = \boldsymbol{\Sigma}_0^{-1} + \frac{1}{\sigma_w^2} \boldsymbol{\Phi}^\top \boldsymbol{\Phi}, \quad b = \boldsymbol{\Sigma}_0^{-1} \boldsymbol{\mu}_0 + \frac{1}{\sigma_w^2} \boldsymbol{\Phi}^\top \mathbf{y}.$$

$$-\frac{1}{2} \left[\boldsymbol{\theta}^\top A \boldsymbol{\theta} - 2\boldsymbol{\theta}^\top b \right] = -\frac{1}{2} (\boldsymbol{\theta} - A^{-1}b)^\top A (\boldsymbol{\theta} - A^{-1}b) + \text{const.}$$

$$\Rightarrow p(\boldsymbol{\theta} \mid \mathbf{y}, \boldsymbol{\Phi}) = \mathcal{N}(\boldsymbol{\mu}_N, \boldsymbol{\Sigma}_N), \quad \boldsymbol{\Sigma}_N = A^{-1}, \quad \boldsymbol{\mu}_N = A^{-1}b.$$

BLR: Posterior Mean and Covariance

$$p(\theta \mid \mathbf{y}, \Phi) = \mathcal{N}(\mu_N, \Sigma_N).$$

$$\Sigma_N = \left(\Sigma_0^{-1} + \frac{1}{\sigma_w^2} \Phi^T \Phi \right)^{-1}$$

$$\mu_N = \Sigma_N \left(\Sigma_0^{-1} \mu_0 + \frac{1}{\sigma_w^2} \Phi^T \mathbf{y} \right).$$

Precision form (often cleaner):

$$\Sigma_N^{-1} = \Sigma_0^{-1} + \frac{1}{\sigma_w^2} \Phi^T \Phi.$$

Interpretation

- Prior precision + data precision.
- More / cleaner data $\Rightarrow$ larger precision $\Rightarrow$ smaller variance.

MAP estimator (posterior mode):

$$\hat{\theta}_{\text{MAP}} = \arg \max_{\theta} p(\theta \mid \mathbf{y}, \Phi).$$

Using the log posterior:

$$\hat{\theta}_{\text{MAP}} = \arg \min_{\theta} \frac{1}{2\sigma_w^2} \|\mathbf{y} - \Phi\theta\|^2 + \frac{1}{2}(\theta - \mu_0)^\top \Sigma_0^{-1}(\theta - \mu_0).$$

Special case $\mu_0 = \mathbf{0}$, $\Sigma_0 = \alpha^{-1}I$:

$$\hat{\theta}_{\text{MAP}} = \arg \min_{\theta} \|\mathbf{y} - \Phi\theta\|^2 + \lambda \|\theta\|^2, \quad \lambda = \sigma_w^2 \alpha,$$

i.e. ridge regression / Tikhonov regularization.

BLR: Posterior Predictive (After Data)

Predict at $\mathbf{x}_*$ after observing data:

$$p(y_* | \mathbf{x}_*, \mathbf{y}, \Phi) = \int p(y_* | \mathbf{x}_*, \theta) p(\theta | \mathbf{y}, \Phi) d\theta.$$

With

$$p(y_* | \mathbf{x}_*, \theta) = \mathcal{N}(\phi_*^\top \theta, \sigma_w^2), \quad p(\theta | \mathbf{y}, \Phi) = \mathcal{N}(\boldsymbol{\mu}_N, \boldsymbol{\Sigma}_N),$$

we get:

$$p(y_* | \mathbf{x}_*, \mathbf{y}, \Phi) = \mathcal{N}\left(\phi_*^\top \boldsymbol{\mu}_N, \phi_*^\top \boldsymbol{\Sigma}_N \phi_* + \sigma_w^2\right).$$

Posterior Predictive: Mean vs Variance

Predictive mean:

$$\mathbb{E}[y_* | \cdot] = \phi_*^T \mu_N \quad (\text{use posterior mean weights}).$$

Predictive variance:

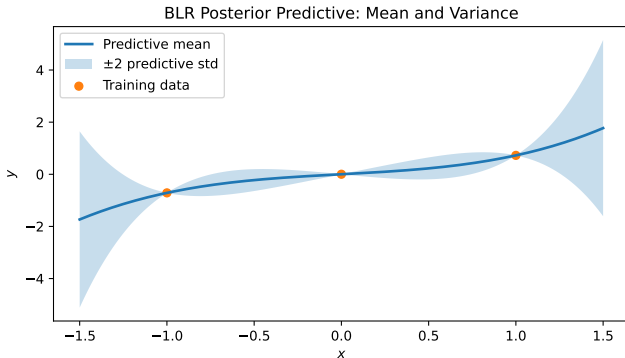
$$\text{Var}(y_* | \cdot) = \underbrace{\phi_*^T \Sigma_N \phi_*}_{\text{parameter uncertainty}} + \underbrace{\sigma_w^2}_{\text{measurement noise}}.$$

Key messages

- Even with zero noise ($\sigma_w^2 \rightarrow 0$), uncertainty remains if $\Sigma_N \neq 0$.
- As $N \rightarrow \infty$, $\Sigma_N \rightarrow 0$: uncertainty bands tighten.

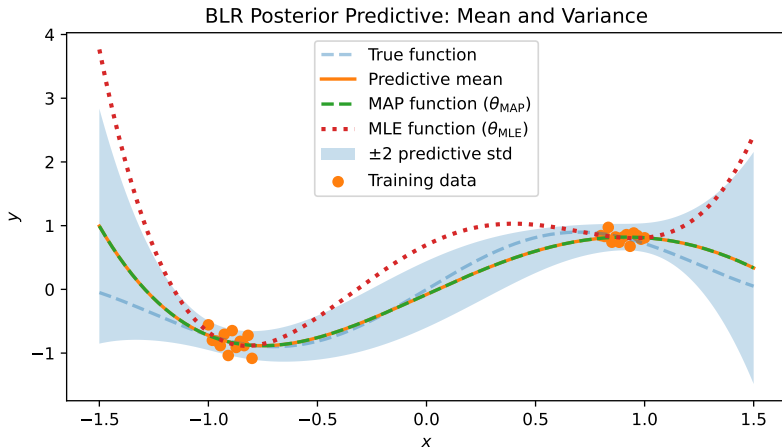
Posterior Predictive: Mean vs Variance (Example)

Example: BLR with polynomial features on noisy 1D data.



- Solid line: predictive mean $\mathbb{E}[y_* | \cdot] = \phi_*^\top \mu_N$.
- Shaded band: ± 2 predictive std, $\sqrt{\phi_*^\top \Sigma_N \phi_* + \sigma_w^2}$.
- Uncertainty is small near data, larger far away (parameter uncertainty).

Posterior Predictive: BLR vs MLE vs MAP



BLR: Sequential / Recursive Update

BLR can be updated one data point at a time. Assume current posterior after k samples:

$$p(\boldsymbol{\theta} \mid \mathcal{D}_k) = \mathcal{N}(\boldsymbol{\mu}_k, \boldsymbol{\Sigma}_k).$$

Given new sample $(\mathbf{x}_{k+1}, y_{k+1})$ with feature ϕ_{k+1} :

$$\boldsymbol{\Sigma}_{k+1}^{-1} = \boldsymbol{\Sigma}_k^{-1} + \frac{1}{\sigma_w^2} \phi_{k+1} \phi_{k+1}^\top,$$

$$\boldsymbol{\mu}_{k+1} = \boldsymbol{\Sigma}_{k+1} \left(\boldsymbol{\Sigma}_k^{-1} \boldsymbol{\mu}_k + \frac{1}{\sigma_w^2} \phi_{k+1} y_{k+1} \right).$$

BLR: What We Gain vs LS/MLE

- A *distribution* over parameters (uncertainty on θ).
- Built-in regularization via the prior.
- Predictive uncertainty:

$$p(y_* | \mathbf{x}_*, \mathbf{y}) = \int p(y_* | \mathbf{x}_*, \theta) p(\theta | \mathbf{y}) d\theta.$$

- With enough data, posterior concentrates and approaches the MLE.

Takeaway: BLR = $\underbrace{\text{LS/MLE} + \text{prior knowledge}}_{\text{MAP}} + \text{uncertainty quantification}.$

Thank you

- **Any Questions?**
- **Office Hours:**
 - **Tue & Thu (09:00-11:00)**
 - 24/7 by email (costashatz@upatras.gr, subject: *ECE_SP_AM*)
- **Material and Announcements**



Laboratory of Automation & Robotics