

ADAPTIVE BAYESIAN ESTIMATION IN INDIRECT GAUSSIAN SEQUENCE SPACE MODELS

Jan JOHANNES

Ruprecht-Karls-Universität Heidelberg

"Mathematical statistics and inverse problems"

February 2016, CIRM, Marseille

joint work with Anna Simoni and Rudolf Schenk



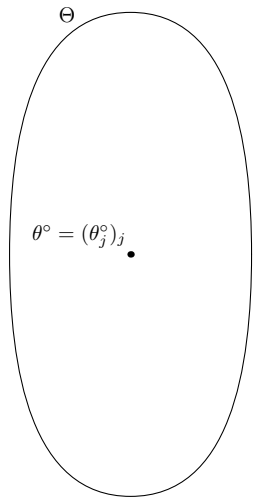
Outline

- Introduction
- Frequentist perspective reviewed
- Posterior concentration
- Adaptive posterior concentration
- Adaptive Bayes estimator

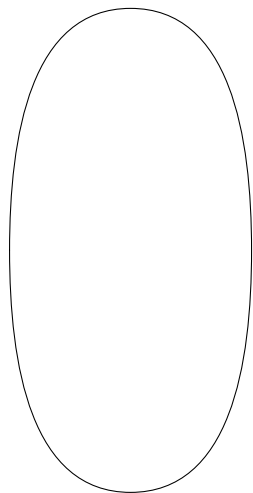
Outline

- Introduction
- Frequentist perspective reviewed
- Posterior concentration
- Adaptive posterior concentration
- Adaptive Bayes estimator

A glimpse to the essential:

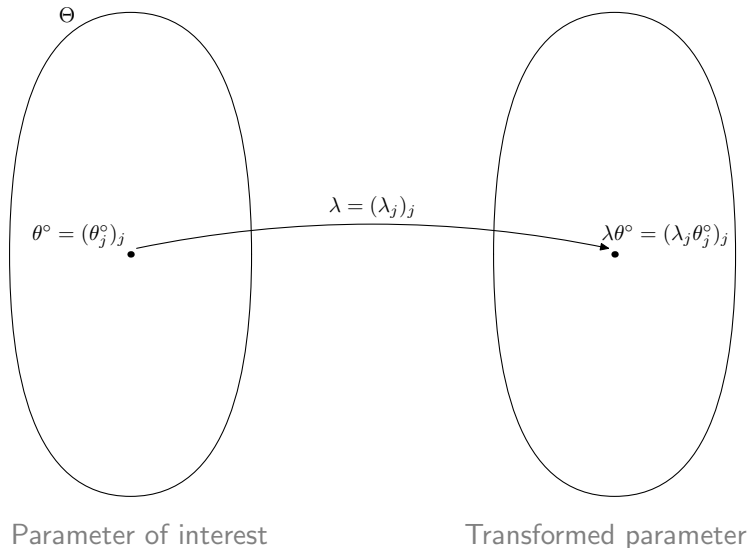


Parameter of interest

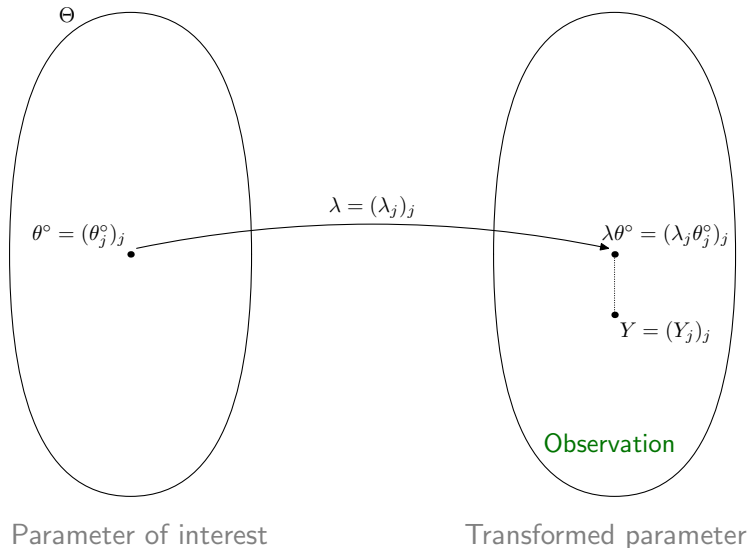


Transformed parameter

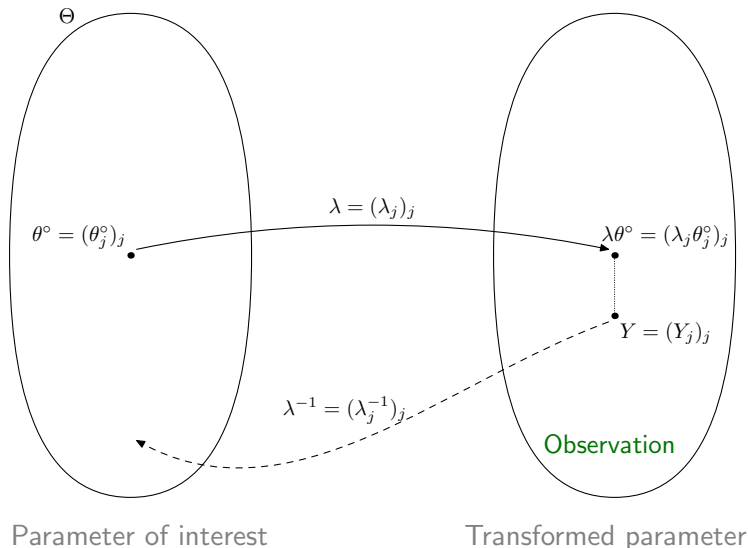
A glimpse to the essential:



A glimpse to the essential:



A glimpse to the essential: statistical inverse problem



Background: analytical inverse problem

Consider $(\mathcal{F}, \langle \cdot, \cdot \rangle)$, $(\mathcal{G}, \langle \cdot, \cdot \rangle)$ and $T : \mathcal{F} \rightarrow \mathcal{G}$ linear

$$f \xrightarrow{T} g$$

Background: analytical inverse problem

Consider $(\mathcal{F}, \langle \cdot, \cdot \rangle)$, $(\mathcal{G}, \langle \cdot, \cdot \rangle)$ and $T : \mathcal{F} \rightarrow \mathcal{G}$ linear

$$f \xrightarrow{T} g$$

special case T permits a singular value decomposition

- Singular values $\lambda := (\lambda_j)_j$
- Eigenfunctions $\{u_j\}_j$ and $\{v_j\}_j$ ONB of $\mathcal{F}$ and $\mathcal{G}$, resp.

Background: analytical inverse problem

Consider $(\mathcal{F}, \langle \cdot, \cdot \rangle)$, $(\mathcal{G}, \langle \cdot, \cdot \rangle)$ and $T : \mathcal{F} \rightarrow \mathcal{G}$ linear

$$\begin{array}{ccc} f & \xrightarrow{T} & g \\ \updownarrow & & \\ (\langle f, u_j \rangle)_j & & \end{array}$$

special case T permits a singular value decomposition

- Singular values $\lambda := (\lambda_j)_j$
- Eigenfunctions $\{u_j\}_j$ and $\{v_j\}_j$ ONB of $\mathcal{F}$ and $\mathcal{G}$, resp.

Representation

- $f \in \mathcal{F}$

Background: analytical inverse problem

Consider $(\mathcal{F}, \langle \cdot, \cdot \rangle)$, $(\mathcal{G}, \langle \cdot, \cdot \rangle)$ and $T : \mathcal{F} \rightarrow \mathcal{G}$ linear

$$\begin{array}{ccc} f & \xrightarrow{T} & g \\ \updownarrow & & \\ (\theta_j)_j & & \end{array}$$

special case T permits a singular value decomposition

- Singular values $\lambda := (\lambda_j)_j$
- Eigenfunctions $\{u_j\}_j$ and $\{v_j\}_j$ ONB of $\mathcal{F}$ and $\mathcal{G}$, resp.

Representation

- $f \in \mathcal{F} \leftrightarrow \theta \in \Theta := \ell^2$ via $\theta_j = \langle f, u_j \rangle$

Background: analytical inverse problem

Consider $(\mathcal{F}, \langle \cdot, \cdot \rangle)$, $(\mathcal{G}, \langle \cdot, \cdot \rangle)$ and $T : \mathcal{F} \rightarrow \mathcal{G}$ linear

$$\begin{array}{ccc} f & \xrightarrow{T} & g \\ \updownarrow & & \updownarrow \\ (\theta_j)_j & \xrightarrow{\lambda} & (\lambda_j \theta_j)_j \end{array}$$

special case T permits a singular value decomposition

- Singular values $\lambda := (\lambda_j)_j$
- Eigenfunctions $\{u_j\}_j$ and $\{v_j\}_j$ ONB of $\mathcal{F}$ and $\mathcal{G}$, resp.

Representation

- $f \in \mathcal{F} \leftrightarrow \theta \in \Theta := \ell^2$ via $\theta_j = \langle f, u_j \rangle$
- Operator $T \leftrightarrow$ Multiplication with λ

Framework: Gaussian sequence space models

indirect GSSM

$$Y_j = \lambda_j \theta_j^\circ + \sqrt{\varepsilon} Z_j,$$

- $\{Z_j\}_j \stackrel{iid}{\sim} \mathcal{N}(0, 1),$

Framework: Gaussian sequence space models

indirect GSSM

$$Y_j = \lambda_j \theta_j^\circ + \sqrt{\varepsilon} Z_j,$$

- $\{Z_j\}_j \stackrel{iid}{\sim} \mathcal{N}(0, 1),$
- $\lambda := (\lambda_j)_j \rightarrow 0.$

Framework: Gaussian sequence space models

indirect GSSM

$$Y_j = \lambda_j \theta_j^\circ + \sqrt{\varepsilon} Z_j,$$

- $\{Z_j\}_j \stackrel{iid}{\sim} \mathcal{N}(0, 1),$
- $\lambda := (\lambda_j)_j \rightarrow 0.$

direct GSSM

$$\lambda_j = 1$$

Framework: Gaussian sequence space models

indirect GSSM

$$Y_j = \lambda_j \theta_j^\circ + \sqrt{\varepsilon} Z_j,$$

- $\{Z_j\}_j \stackrel{iid}{\sim} \mathcal{N}(0, 1),$
- $\lambda := (\lambda_j)_j \rightarrow 0.$

direct GSSM

$$\lambda_j = 1$$

indirect Gaussian regression

$$dY = T f + \sqrt{\varepsilon} dW$$

Framework: Gaussian sequence space models

indirect GSSM

$$Y_j = \lambda_j \theta_j^\circ + \sqrt{\varepsilon} Z_j,$$

- $\{Z_j\}_j \stackrel{iid}{\sim} \mathcal{N}(0, 1),$
- $\lambda := (\lambda_j)_j \rightarrow 0.$

direct GSSM

$$\lambda_j = 1$$

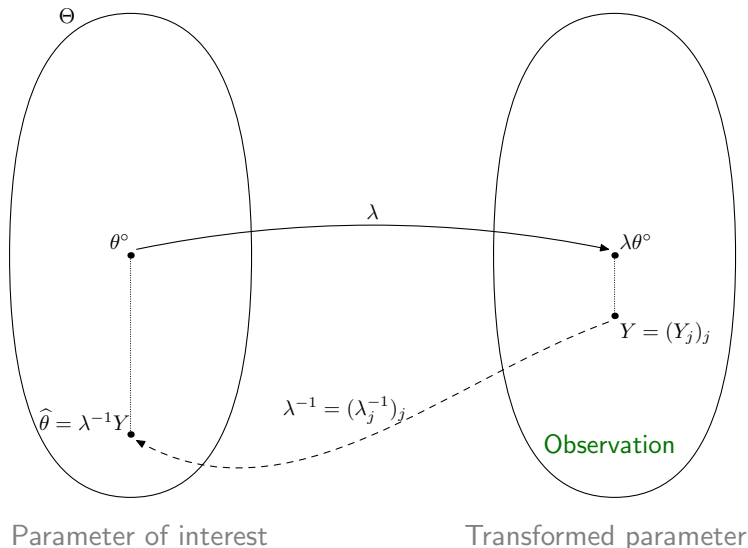
indirect Gaussian regression

$$dY = T f + \sqrt{\varepsilon} dW$$

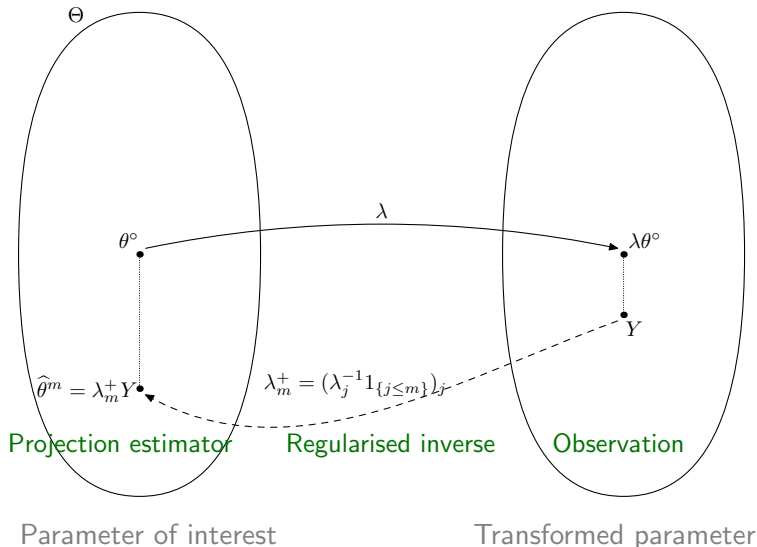
direct Gaussian regression

$$T = id$$

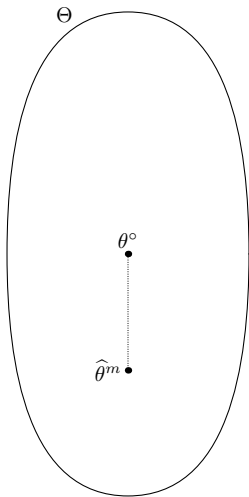
Frequentist point of view: statistical inverse problem



Frequentist point of view: regularised estimation



Frequentist point of view: oracle optimality

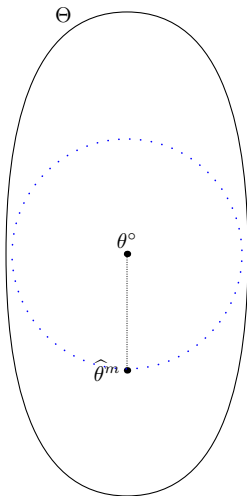


Parameter of interest

Frequentist point of view: oracle optimality

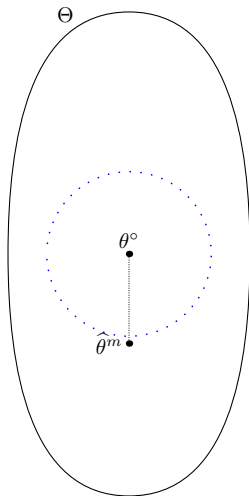
- Measure the performance – risk

$$\mathbb{E}_{\theta^\circ} [\|\hat{\theta}^m - \theta^\circ\|^2]$$



Parameter of interest

Frequentist point of view: oracle optimality



Parameter of interest

- Measure the performance – risk

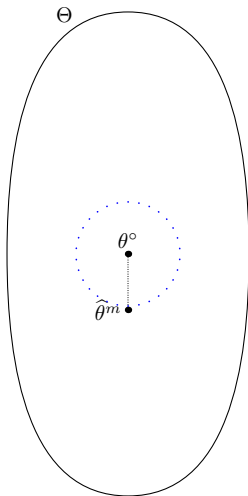
$$\mathbb{E}_{\theta^\circ} [\|\hat{\theta}^m - \theta^\circ\|^2]$$

- Complexity – lower bound

$$\inf_m \mathbb{E}_{\theta^\circ} [\|\hat{\theta}^m - \theta^\circ\|^2] \gtrsim \mathcal{R}_\varepsilon^\circ(\theta^\circ)$$

over a family $\{\hat{\theta}^m\}$ of estimators

Frequentist point of view: oracle optimality



Parameter of interest

- Measure the performance – risk

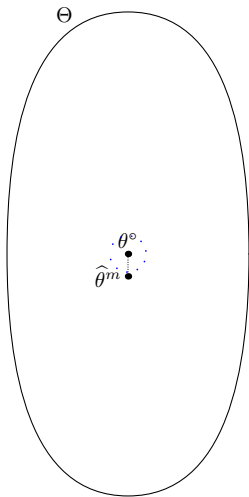
$$\mathbb{E}_{\theta^\circ} [\|\hat{\theta}^m - \theta^\circ\|^2]$$

- Complexity – lower bound

$$\inf_m \mathbb{E}_{\theta^\circ} [\|\hat{\theta}^m - \theta^\circ\|^2] \gtrsim \mathcal{R}_\varepsilon^\circ(\theta^\circ)$$

over a family $\{\hat{\theta}^m\}$ of estimators

Frequentist point of view: oracle optimality



Parameter of interest

- Measure the performance – risk

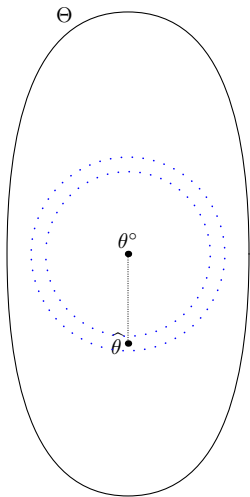
$$\mathbb{E}_{\theta^\circ} [\|\hat{\theta}^m - \theta^\circ\|^2]$$

- Complexity – lower bound

$$\inf_m \mathbb{E}_{\theta^\circ} [\|\hat{\theta}^m - \theta^\circ\|^2] \gtrsim \mathcal{R}_\varepsilon^\circ(\theta^\circ)$$

over a family $\{\hat{\theta}^m\}$ of estimators

Frequentist point of view: oracle optimality



Parameter of interest

- Measure the performance – **risk**

$$\mathbb{E}_{\theta^\circ} [\|\hat{\theta}^m - \theta^\circ\|^2]$$

- Complexity – **lower bound**

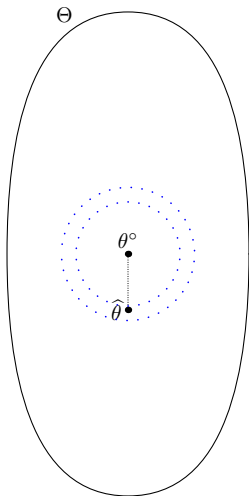
$$\inf_m \mathbb{E}_{\theta^\circ} [\|\hat{\theta}^m - \theta^\circ\|^2] \gtrsim \mathcal{R}_\varepsilon^\circ(\theta^\circ)$$

over a family $\{\hat{\theta}^m\}$ of estimators

- $\hat{\theta} \in \{\hat{\theta}^m\}$ is called an **oracle** if

$$\mathbb{E}_{\theta^\circ} [\|\hat{\theta} - \theta^\circ\|^2] \lesssim \mathcal{R}_\varepsilon^\circ(\theta^\circ)$$

Frequentist point of view: oracle optimality



Parameter of interest

- Measure the performance – **risk**

$$\mathbb{E}_{\theta^\circ} [\|\hat{\theta}^m - \theta^\circ\|^2]$$

- Complexity – **lower bound**

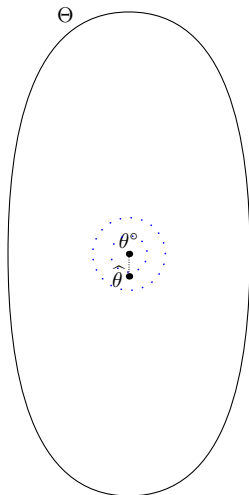
$$\inf_m \mathbb{E}_{\theta^\circ} [\|\hat{\theta}^m - \theta^\circ\|^2] \gtrsim \mathcal{R}_\varepsilon^\circ(\theta^\circ)$$

over a family $\{\hat{\theta}^m\}$ of estimators

- $\hat{\theta} \in \{\hat{\theta}^m\}$ is called an **oracle** if

$$\mathbb{E}_{\theta^\circ} [\|\hat{\theta} - \theta^\circ\|^2] \lesssim \mathcal{R}_\varepsilon^\circ(\theta^\circ)$$

Frequentist point of view: oracle optimality



Parameter of interest

- Measure the performance – **risk**

$$\mathbb{E}_{\theta^\circ} [\|\hat{\theta}^m - \theta^\circ\|^2]$$

- Complexity – **lower bound**

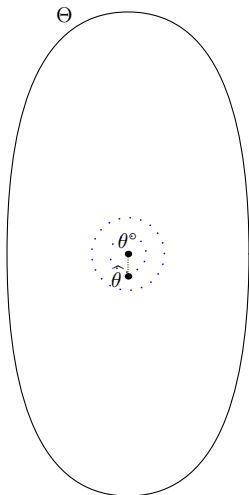
$$\inf_m \mathbb{E}_{\theta^\circ} [\|\hat{\theta}^m - \theta^\circ\|^2] \gtrsim \mathcal{R}_\varepsilon^\circ(\theta^\circ)$$

over a family $\{\hat{\theta}^m\}$ of estimators

- $\hat{\theta} \in \{\hat{\theta}^m\}$ is called an **oracle** if

$$\mathbb{E}_{\theta^\circ} [\|\hat{\theta} - \theta^\circ\|^2] \lesssim \mathcal{R}_\varepsilon^\circ(\theta^\circ)$$

Frequentist point of view: oracle optimality



Parameter of interest

- Measure the performance – **risk**

$$\mathbb{E}_{\theta^\circ} [\|\hat{\theta}^m - \theta^\circ\|^2]$$

- Complexity – **lower bound**

$$\inf_m \mathbb{E}_{\theta^\circ} [\|\hat{\theta}^m - \theta^\circ\|^2] \gtrsim \mathcal{R}_\varepsilon^\circ(\theta^\circ)$$

over a family $\{\hat{\theta}^m\}$ of estimators

- $\hat{\theta} \in \{\hat{\theta}^m\}$ is called an **oracle** if

$$\mathbb{E}_{\theta^\circ} [\|\hat{\theta} - \theta^\circ\|^2] \lesssim \mathcal{R}_\varepsilon^\circ(\theta^\circ)$$

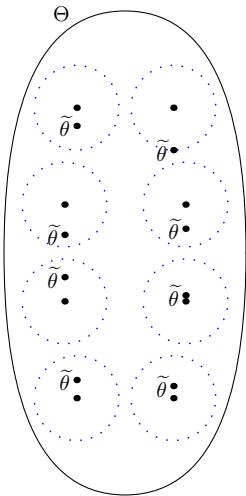
and **adaptive** if $\hat{\theta}$ does not depend on θ° .

Frequentist point of view: minimax optimality

- Measure the performance – maximal risk

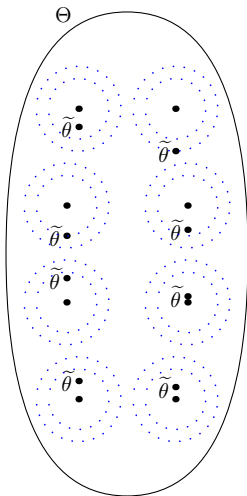
$$\sup_{\theta^\circ \in \Theta_a} \mathbb{E}_{\theta^\circ} [\|\tilde{\theta} - \theta^\circ\|^2]$$

over a class Θ_a of parameters



Parameter of interest

Frequentist point of view: minimax optimality



Parameter of interest

- Measure the performance – maximal risk

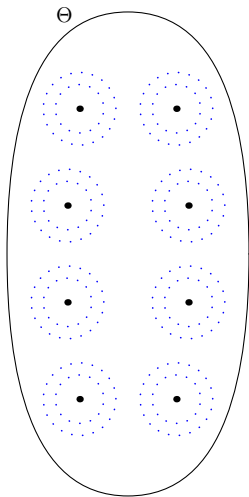
$$\sup_{\theta^\circ \in \Theta_{\mathfrak{a}}} \mathbb{E}_{\theta^\circ} [\|\tilde{\theta} - \theta^\circ\|^2]$$

over a class $\Theta_{\mathfrak{a}}$ of parameters

- Complexity – lower bound

$$\inf_{\tilde{\theta}} \sup_{\theta^\circ \in \Theta_{\mathfrak{a}}} \mathbb{E}_{\theta^\circ} [\|\tilde{\theta} - \theta^\circ\|^2] \gtrsim \mathcal{R}_\varepsilon^*(\mathfrak{a})$$

Frequentist point of view: minimax optimality



Parameter of interest

- Measure the performance – maximal risk

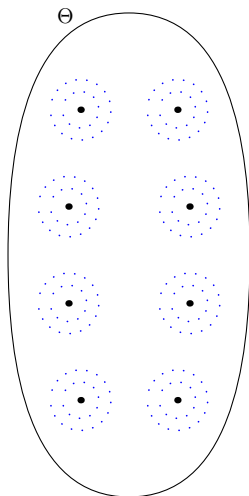
$$\sup_{\theta^\circ \in \Theta_\alpha} \mathbb{E}_{\theta^\circ} [\|\tilde{\theta} - \theta^\circ\|^2]$$

over a class Θ_α of parameters

- Complexity – lower bound

$$\inf_{\tilde{\theta}} \sup_{\theta^\circ \in \Theta_\alpha} \mathbb{E}_{\theta^\circ} [\|\tilde{\theta} - \theta^\circ\|^2] \gtrsim \mathcal{R}_\varepsilon^*(\alpha)$$

Frequentist point of view: minimax optimality



Parameter of interest

- Measure the performance – maximal risk

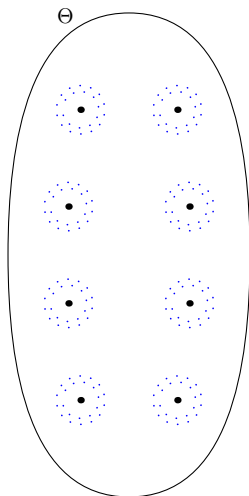
$$\sup_{\theta^\circ \in \Theta_a} \mathbb{E}_{\theta^\circ} [\|\tilde{\theta} - \theta^\circ\|^2]$$

over a class Θ_a of parameters

- Complexity – lower bound

$$\inf_{\tilde{\theta}} \sup_{\theta^\circ \in \Theta_a} \mathbb{E}_{\theta^\circ} [\|\tilde{\theta} - \theta^\circ\|^2] \gtrsim \mathcal{R}_\varepsilon^*(a)$$

Frequentist point of view: minimax optimality



Parameter of interest

- Measure the performance – maximal risk

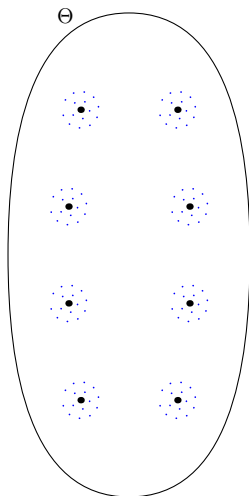
$$\sup_{\theta^\circ \in \Theta_{\mathfrak{a}}} \mathbb{E}_{\theta^\circ} [\|\tilde{\theta} - \theta^\circ\|^2]$$

over a class $\Theta_{\mathfrak{a}}$ of parameters

- Complexity – lower bound

$$\inf_{\tilde{\theta}} \sup_{\theta^\circ \in \Theta_{\mathfrak{a}}} \mathbb{E}_{\theta^\circ} [\|\tilde{\theta} - \theta^\circ\|^2] \gtrsim \mathcal{R}_\varepsilon^*(\mathfrak{a})$$

Frequentist point of view: minimax optimality



Parameter of interest

- Measure the performance – maximal risk

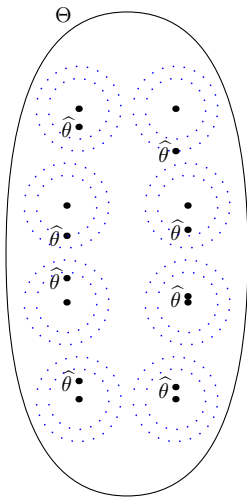
$$\sup_{\theta^\circ \in \Theta_a} \mathbb{E}_{\theta^\circ} [\|\tilde{\theta} - \theta^\circ\|^2]$$

over a class Θ_a of parameters

- Complexity – lower bound

$$\inf_{\tilde{\theta}} \sup_{\theta^\circ \in \Theta_a} \mathbb{E}_{\theta^\circ} [\|\tilde{\theta} - \theta^\circ\|^2] \gtrsim \mathcal{R}_\varepsilon^*(a)$$

Frequentist point of view: minimax optimality



Parameter of interest

- Measure the performance – **maximal risk**

$$\sup_{\theta^\circ \in \Theta_a} \mathbb{E}_{\theta^\circ} [\|\tilde{\theta} - \theta^\circ\|^2]$$

over a class Θ_a of parameters

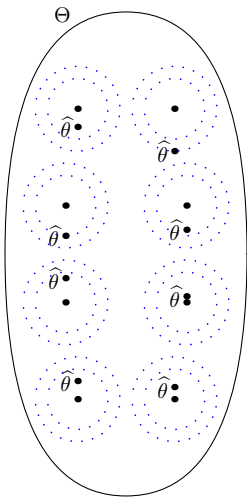
- Complexity – **lower bound**

$$\inf_{\tilde{\theta}} \sup_{\theta^\circ \in \Theta_a} \mathbb{E}_{\theta^\circ} [\|\tilde{\theta} - \theta^\circ\|^2] \gtrsim \mathcal{R}_\varepsilon^*(a)$$

- $\hat{\theta}$ is called **minimax** rate optimal if

$$\sup_{\theta^\circ \in \Theta_a} \mathbb{E}_{\theta^\circ} [\|\hat{\theta} - \theta^\circ\|^2] \lesssim \mathcal{R}_\varepsilon^*(a)$$

Frequentist point of view: minimax optimality



Parameter of interest

- Measure the performance – **maximal risk**

$$\sup_{\theta^\circ \in \Theta_a} \mathbb{E}_{\theta^\circ} [\|\tilde{\theta} - \theta^\circ\|^2]$$

over a class Θ_a of parameters

- Complexity – **lower bound**

$$\inf_{\tilde{\theta}} \sup_{\theta^\circ \in \Theta_a} \mathbb{E}_{\theta^\circ} [\|\tilde{\theta} - \theta^\circ\|^2] \gtrsim \mathcal{R}_\varepsilon^*(a)$$

- $\hat{\theta}$ is called **minimax** rate optimal if

$$\sup_{\theta^\circ \in \Theta_a} \mathbb{E}_{\theta^\circ} [\|\hat{\theta} - \theta^\circ\|^2] \lesssim \mathcal{R}_\varepsilon^*(a)$$

and **adaptive** if $\hat{\theta}$ does not depend on a .

A glimpse to the essential: Bayesian perspective

- ▶ θ outcome of Θ -valued r.v. $\mathcal{V}$

A glimpse to the essential: Bayesian perspective

- ▶ θ outcome of Θ -valued r.v. ϑ
- ▶ likelihood: $P_{Y|\vartheta}$ with density $p_{Y|\vartheta}$

A glimpse to the essential: Bayesian perspective

- ▶ θ outcome of Θ -valued r.v. ϑ
- ▶ likelihood: $P_{Y|\vartheta}$ with density $p_{Y|\vartheta}$
- ▶ prior distribution: P_{ϑ} on Θ with density p_{ϑ}

A glimpse to the essential: Bayesian perspective

- ▶ θ outcome of Θ -valued r.v. ϑ
- ▶ likelihood: $P_{Y|\vartheta}$ with density $p_{Y|\vartheta}$
- ▶ prior distribution: P_{ϑ} on Θ with density p_{ϑ}
- ▶ posterior distribution $P_{\vartheta|Y}$ with density:

$$p_{\vartheta|Y}(\theta|y) = \frac{p_{Y|\vartheta}(y|\theta)p_{\vartheta}(\theta)}{\int_{\Theta} p_{Y|\vartheta}(y|\theta)p_{\vartheta}(\theta)d\theta}$$

A glimpse to the essential: Bayesian perspective

- ▶ θ outcome of Θ -valued r.v. ϑ
- ▶ likelihood: $P_{Y|\vartheta}$ with density $p_{Y|\vartheta}$
- ▶ prior distribution: P_{ϑ} on Θ with density p_{ϑ}
- ▶ posterior distribution $P_{\vartheta|Y}$ with density:

$$p_{\vartheta|Y}(\theta|y) = \frac{p_{Y|\vartheta}(y|\theta)p_{\vartheta}(\theta)}{\int_{\Theta} p_{Y|\vartheta}(y|\theta)p_{\vartheta}(\theta)d\theta}$$

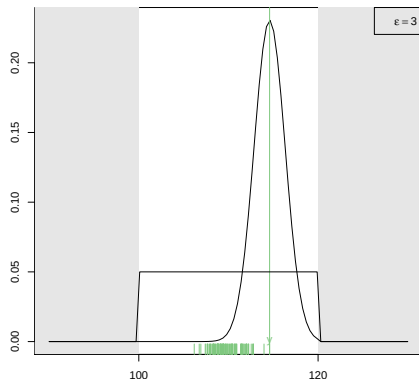
- ▶ “posterior is proportional to likelihood times prior”

Illustration

► likelihood $Y|\vartheta = \theta \sim \mathcal{N}(\theta, \varepsilon)$

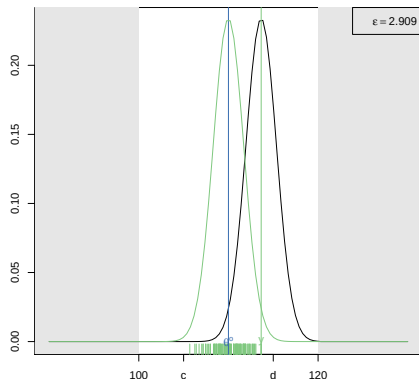
► prior $\vartheta \sim \mathcal{U}[a, b]$

► posterior density: $p_{\vartheta|Y}(\theta|y) = \frac{\varepsilon^{-1/2} \phi\left(\frac{\theta-y}{\sqrt{\varepsilon}}\right)}{\Phi\left(\frac{b-y}{\sqrt{\varepsilon}}\right) - \Phi\left(\frac{a-y}{\sqrt{\varepsilon}}\right)} \mathbb{1}_{[a,b]}(\theta)$



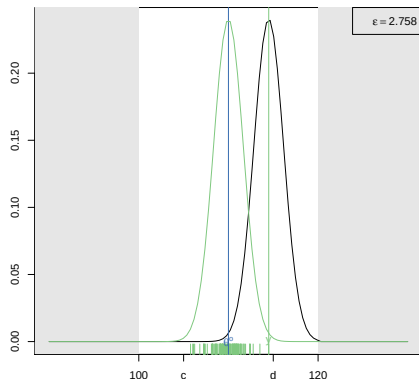
Illustration

- ▶ likelihood $Y|\vartheta = \theta \sim \mathcal{N}(\theta, \varepsilon)$
- ▶ prior $\vartheta \sim \mathcal{U}[a, b]$
- ▶ posterior density: $p_{\vartheta|Y}(\theta|y) = \frac{\varepsilon^{-1/2} \phi\left(\frac{\theta-y}{\sqrt{\varepsilon}}\right)}{\Phi\left(\frac{b-y}{\sqrt{\varepsilon}}\right) - \Phi\left(\frac{a-y}{\sqrt{\varepsilon}}\right)} \mathbb{1}_{[a,b]}(\theta)$



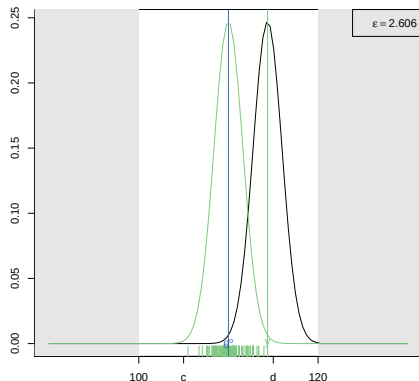
Illustration

- ▶ likelihood $Y|\vartheta = \theta \sim \mathcal{N}(\theta, \varepsilon)$
- ▶ prior $\vartheta \sim \mathcal{U}[a, b]$
- ▶ posterior density: $p_{\vartheta|Y}(\theta|y) = \frac{\varepsilon^{-1/2} \phi\left(\frac{\theta-y}{\sqrt{\varepsilon}}\right)}{\Phi\left(\frac{b-y}{\sqrt{\varepsilon}}\right) - \Phi\left(\frac{a-y}{\sqrt{\varepsilon}}\right)} \mathbb{1}_{[a,b]}(\theta)$



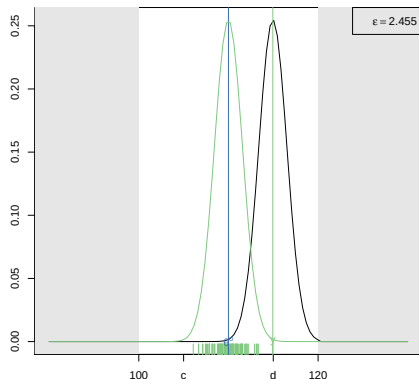
Illustration

- ▶ likelihood $Y|\vartheta = \theta \sim \mathcal{N}(\theta, \varepsilon)$
- ▶ prior $\vartheta \sim \mathcal{U}[a, b]$
- ▶ posterior density: $p_{\vartheta|Y}(\theta|y) = \frac{\varepsilon^{-1/2} \phi\left(\frac{\theta-y}{\sqrt{\varepsilon}}\right)}{\Phi\left(\frac{b-y}{\sqrt{\varepsilon}}\right) - \Phi\left(\frac{a-y}{\sqrt{\varepsilon}}\right)} \mathbb{1}_{[a,b]}(\theta)$



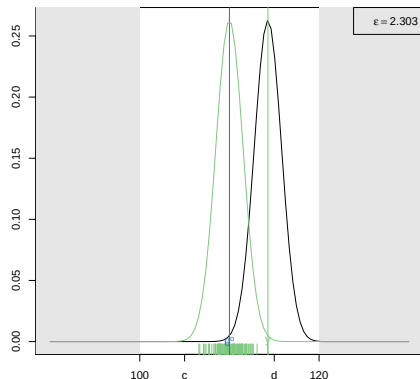
Illustration

- ▶ likelihood $Y|\boldsymbol{\vartheta} = \theta \sim \mathcal{N}(\theta, \varepsilon)$
- ▶ prior $\boldsymbol{\vartheta} \sim \mathcal{U}[a, b]$
- ▶ posterior density: $p_{\boldsymbol{\vartheta}|Y}(\theta|y) = \frac{\varepsilon^{-1/2} \phi\left(\frac{\theta-y}{\sqrt{\varepsilon}}\right)}{\Phi\left(\frac{b-y}{\sqrt{\varepsilon}}\right) - \Phi\left(\frac{a-y}{\sqrt{\varepsilon}}\right)} \mathbb{1}_{[a,b]}(\theta)$



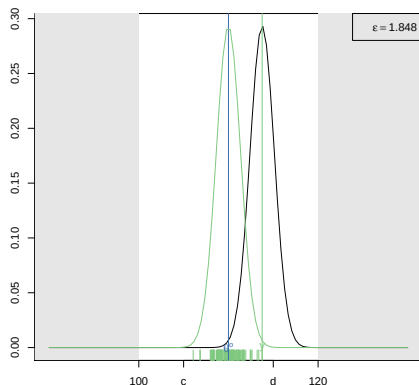
Illustration

- ▶ likelihood $Y|\vartheta = \theta \sim \mathcal{N}(\theta, \varepsilon)$
- ▶ prior $\vartheta \sim \mathcal{U}[a, b]$
- ▶ posterior density: $p_{\vartheta|Y}(\theta|y) = \frac{\varepsilon^{-1/2} \phi\left(\frac{\theta-y}{\sqrt{\varepsilon}}\right)}{\Phi\left(\frac{b-y}{\sqrt{\varepsilon}}\right) - \Phi\left(\frac{a-y}{\sqrt{\varepsilon}}\right)} \mathbb{1}_{[a,b]}(\theta)$



Illustration

- ▶ likelihood $Y|\vartheta = \theta \sim \mathcal{N}(\theta, \varepsilon)$
- ▶ prior $\vartheta \sim \mathcal{U}[a, b]$
- ▶ posterior density: $p_{\vartheta|Y}(\theta|y) = \frac{\varepsilon^{-1/2} \phi\left(\frac{\theta-y}{\sqrt{\varepsilon}}\right)}{\Phi\left(\frac{b-y}{\sqrt{\varepsilon}}\right) - \Phi\left(\frac{a-y}{\sqrt{\varepsilon}}\right)} \mathbb{1}_{[a,b]}(\theta)$

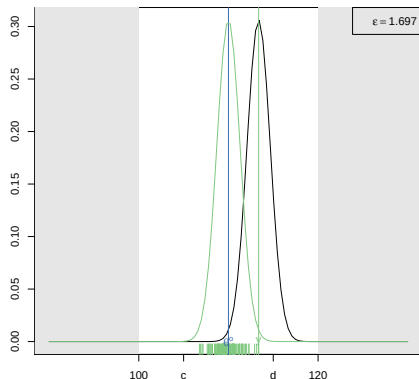


Illustration

► likelihood $Y|\vartheta = \theta \sim \mathcal{N}(\theta, \varepsilon)$

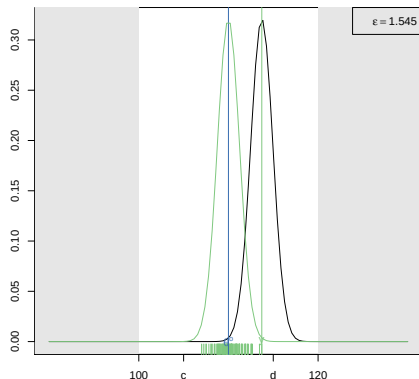
► prior $\vartheta \sim \mathcal{U}[a, b]$

► posterior density: $p_{\vartheta|Y}(\theta|y) = \frac{\varepsilon^{-1/2} \phi\left(\frac{\theta-y}{\sqrt{\varepsilon}}\right)}{\Phi\left(\frac{b-y}{\sqrt{\varepsilon}}\right) - \Phi\left(\frac{a-y}{\sqrt{\varepsilon}}\right)} \mathbb{1}_{[a,b]}(\theta)$



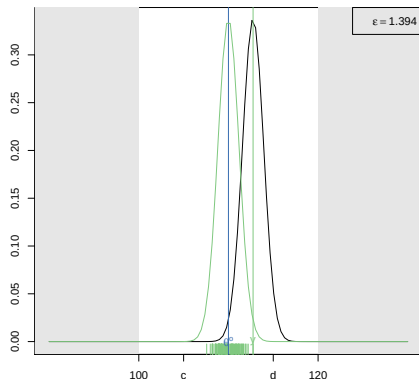
Illustration

- ▶ likelihood $Y|\vartheta = \theta \sim \mathcal{N}(\theta, \varepsilon)$
- ▶ prior $\vartheta \sim \mathcal{U}[a, b]$
- ▶ posterior density: $p_{\vartheta|Y}(\theta|y) = \frac{\varepsilon^{-1/2} \phi\left(\frac{\theta-y}{\sqrt{\varepsilon}}\right)}{\Phi\left(\frac{b-y}{\sqrt{\varepsilon}}\right) - \Phi\left(\frac{a-y}{\sqrt{\varepsilon}}\right)} \mathbb{1}_{[a,b]}(\theta)$



Illustration

- ▶ likelihood $Y|\boldsymbol{\vartheta} = \theta \sim \mathcal{N}(\theta, \varepsilon)$
- ▶ prior $\boldsymbol{\vartheta} \sim \mathcal{U}[a, b]$
- ▶ posterior density: $p_{\boldsymbol{\vartheta}|Y}(\theta|y) = \frac{\varepsilon^{-1/2} \phi\left(\frac{\theta-y}{\sqrt{\varepsilon}}\right)}{\Phi\left(\frac{b-y}{\sqrt{\varepsilon}}\right) - \Phi\left(\frac{a-y}{\sqrt{\varepsilon}}\right)} \mathbb{1}_{[a,b]}(\theta)$

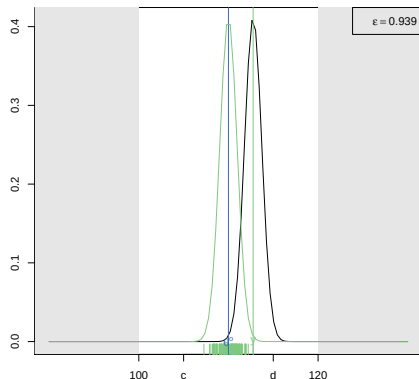


Illustration

► likelihood $Y|\vartheta = \theta \sim \mathcal{N}(\theta, \varepsilon)$

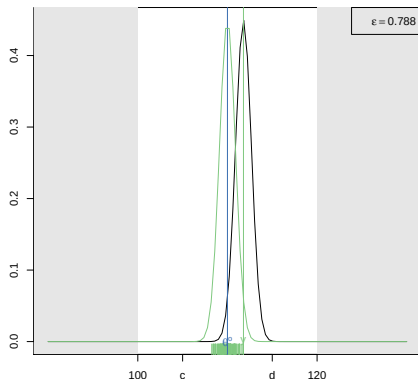
► prior $\vartheta \sim \mathcal{U}[a, b]$

► posterior density: $p_{\vartheta|Y}(\theta|y) = \frac{\varepsilon^{-1/2} \phi\left(\frac{\theta-y}{\sqrt{\varepsilon}}\right)}{\Phi\left(\frac{b-y}{\sqrt{\varepsilon}}\right) - \Phi\left(\frac{a-y}{\sqrt{\varepsilon}}\right)} \mathbb{1}_{[a,b]}(\theta)$



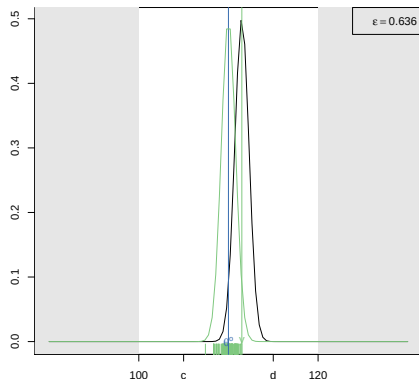
Illustration

- ▶ likelihood $Y|\vartheta = \theta \sim \mathcal{N}(\theta, \varepsilon)$
- ▶ prior $\vartheta \sim \mathcal{U}[a, b]$
- ▶ posterior density: $p_{\vartheta|Y}(\theta|y) = \frac{\varepsilon^{-1/2} \phi\left(\frac{\theta-y}{\sqrt{\varepsilon}}\right)}{\Phi\left(\frac{b-y}{\sqrt{\varepsilon}}\right) - \Phi\left(\frac{a-y}{\sqrt{\varepsilon}}\right)} \mathbb{1}_{[a,b]}(\theta)$



Illustration

- ▶ likelihood $Y|\vartheta = \theta \sim \mathcal{N}(\theta, \varepsilon)$
- ▶ prior $\vartheta \sim \mathcal{U}[a, b]$
- ▶ posterior density: $p_{\vartheta|Y}(\theta|y) = \frac{\varepsilon^{-1/2} \phi\left(\frac{\theta-y}{\sqrt{\varepsilon}}\right)}{\Phi\left(\frac{b-y}{\sqrt{\varepsilon}}\right) - \Phi\left(\frac{a-y}{\sqrt{\varepsilon}}\right)} \mathbb{1}_{[a,b]}(\theta)$

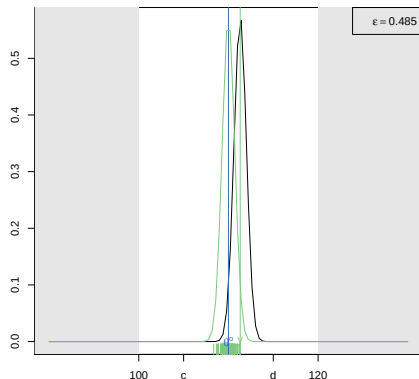


Illustration

► likelihood $Y|\vartheta = \theta \sim \mathcal{N}(\theta, \varepsilon)$

► prior $\vartheta \sim \mathcal{U}[a, b]$

► posterior density: $p_{\vartheta|Y}(\theta|y) = \frac{\varepsilon^{-1/2} \phi\left(\frac{\theta-y}{\sqrt{\varepsilon}}\right)}{\Phi\left(\frac{b-y}{\sqrt{\varepsilon}}\right) - \Phi\left(\frac{a-y}{\sqrt{\varepsilon}}\right)} \mathbb{1}_{[a,b]}(\theta)$

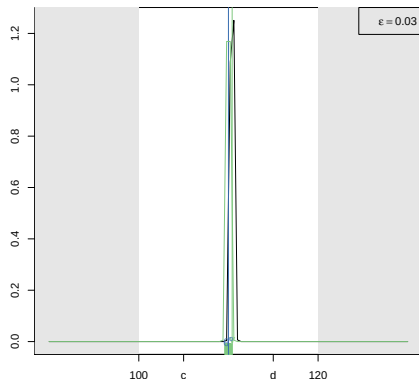


Illustration

► likelihood $Y|\vartheta = \theta \sim \mathcal{N}(\theta, \varepsilon)$

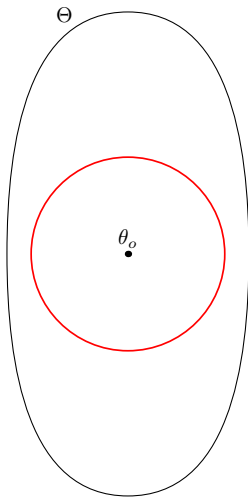
► prior $\vartheta \sim \mathcal{U}[a, b]$

► posterior density: $p_{\vartheta|Y}(\theta|y) = \frac{\varepsilon^{-1/2} \phi\left(\frac{\theta-y}{\sqrt{\varepsilon}}\right)}{\Phi\left(\frac{b-y}{\sqrt{\varepsilon}}\right) - \Phi\left(\frac{a-y}{\sqrt{\varepsilon}}\right)} \mathbb{1}_{[a,b]}(\theta)$



A glimpse to the essential: **exact posterior concentration**

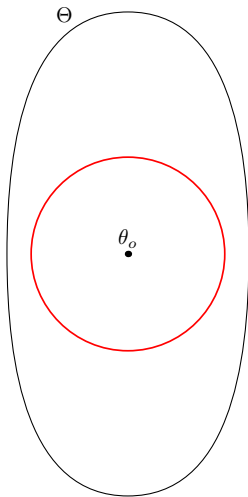
Given prior P_{θ} find $\mathcal{R}_{\varepsilon}$



A glimpse to the essential: exact posterior concentration

Given prior P_{ϑ} find $\mathcal{R}_{\varepsilon} \rightarrow 0$ as $\varepsilon \rightarrow 0$ such that for $P_{\theta^{\circ}} = P_{Y|\vartheta=\theta^{\circ}}$

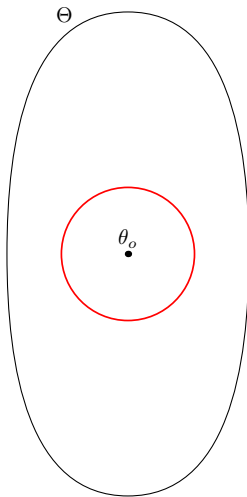
$$\blacktriangleright P_{\vartheta|Y}(\|\vartheta - \theta^{\circ}\| \lesssim \mathcal{R}_{\varepsilon}) \xrightarrow{P_{\theta^{\circ}}} 1$$



A glimpse to the essential: exact posterior concentration

Given prior P_{ϑ} find $\mathcal{R}_\varepsilon \rightarrow 0$ as $\varepsilon \rightarrow 0$ such
that for $P_{\theta^\circ} = P_{Y|\vartheta=\theta^\circ}$

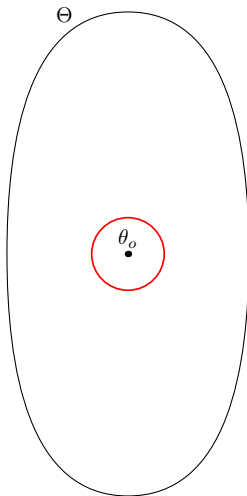
$$\blacktriangleright P_{\vartheta|Y}(\|\vartheta - \theta^\circ\| \lesssim \mathcal{R}_\varepsilon) \xrightarrow{P_{\theta^\circ}} 1$$



A glimpse to the essential: exact posterior concentration

Given prior P_{ϑ} find $\mathcal{R}_\varepsilon \rightarrow 0$ as $\varepsilon \rightarrow 0$ such
that for $P_{\theta^\circ} = P_{Y|\vartheta=\theta^\circ}$

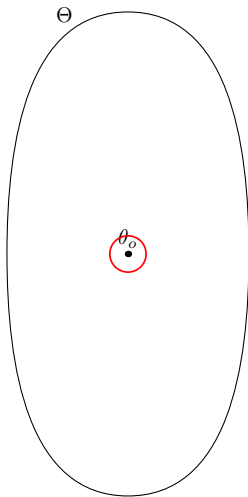
$$\blacktriangleright P_{\vartheta|Y}(\|\vartheta - \theta^\circ\| \lesssim \mathcal{R}_\varepsilon) \xrightarrow{P_{\theta^\circ}} 1$$



A glimpse to the essential: exact posterior concentration

Given prior P_{ϑ} find $\mathcal{R}_{\varepsilon} \rightarrow 0$ as $\varepsilon \rightarrow 0$ such that for $P_{\theta^{\circ}} = P_{Y|\vartheta=\theta^{\circ}}$

$$\blacktriangleright P_{\vartheta|Y}(\|\vartheta - \theta^{\circ}\| \lesssim \mathcal{R}_{\varepsilon}) \xrightarrow{P_{\theta^{\circ}}} 1$$

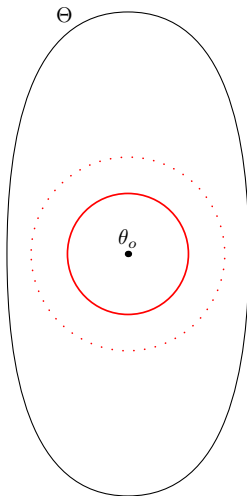


A glimpse to the essential: exact posterior concentration

Given prior P_{ϑ} find $\mathcal{R}_\varepsilon \rightarrow 0$ as $\varepsilon \rightarrow 0$ such that for $P_{\theta^\circ} = P_{Y|\vartheta=\theta^\circ}$

$$\blacktriangleright P_{\vartheta|Y}(\|\vartheta - \theta^\circ\| \lesssim \mathcal{R}_\varepsilon) \xrightarrow{P_{\theta^\circ}} 1$$

$$\blacktriangleright P_{\vartheta|Y}(\|\vartheta - \theta^\circ\| \gtrsim \mathcal{R}_\varepsilon) \xrightarrow{P_{\theta^\circ}} 0$$

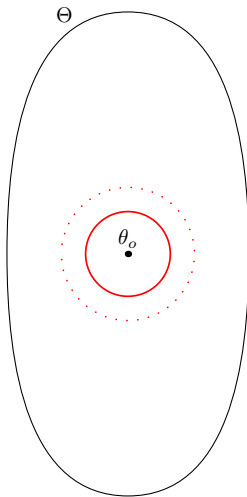


A glimpse to the essential: exact posterior concentration

Given prior P_{ϑ} find $\mathcal{R}_\varepsilon \rightarrow 0$ as $\varepsilon \rightarrow 0$ such that for $P_{\theta^\circ} = P_{Y|\vartheta=\theta^\circ}$

$$\blacktriangleright P_{\vartheta|Y}(\|\vartheta - \theta^\circ\| \lesssim \mathcal{R}_\varepsilon) \xrightarrow{P_{\theta^\circ}} 1$$

$$\blacktriangleright P_{\vartheta|Y}(\|\vartheta - \theta^\circ\| \gtrsim \mathcal{R}_\varepsilon) \xrightarrow{P_{\theta^\circ}} 0$$

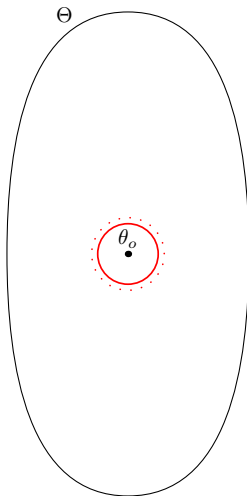


A glimpse to the essential: exact posterior concentration

Given prior P_{ϑ} find $\mathcal{R}_\varepsilon \rightarrow 0$ as $\varepsilon \rightarrow 0$ such that for $P_{\theta^\circ} = P_{Y|\vartheta=\theta^\circ}$

$$\blacktriangleright P_{\vartheta|Y}(\|\vartheta - \theta^\circ\| \lesssim \mathcal{R}_\varepsilon) \xrightarrow{P_{\theta^\circ}} 1$$

$$\blacktriangleright P_{\vartheta|Y}(\|\vartheta - \theta^\circ\| \gtrsim \mathcal{R}_\varepsilon) \xrightarrow{P_{\theta^\circ}} 0$$

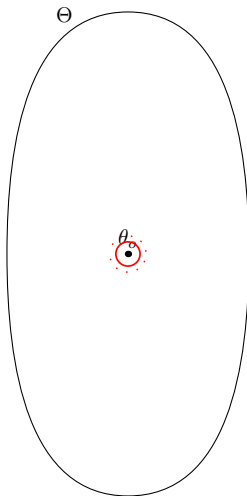


A glimpse to the essential: exact posterior concentration

Given prior P_{ϑ} find $\mathcal{R}_\varepsilon \rightarrow 0$ as $\varepsilon \rightarrow 0$ such that for $P_{\theta^\circ} = P_{Y|\vartheta=\theta^\circ}$

$$\blacktriangleright P_{\vartheta|Y}(\|\vartheta - \theta^\circ\| \lesssim \mathcal{R}_\varepsilon) \xrightarrow{P_{\theta^\circ}} 1$$

$$\blacktriangleright P_{\vartheta|Y}(\|\vartheta - \theta^\circ\| \gtrsim \mathcal{R}_\varepsilon) \xrightarrow{P_{\theta^\circ}} 0$$



A glimpse to the essential: exact posterior concentration

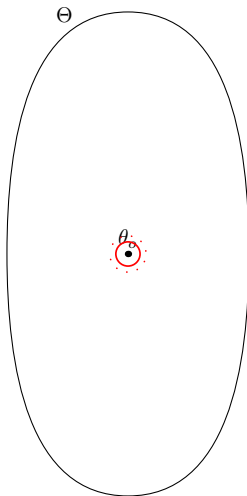
Given prior P_{ϑ} find $\mathcal{R}_{\varepsilon} \rightarrow 0$ as $\varepsilon \rightarrow 0$ such that for $P_{\theta^{\circ}} = P_{Y|\vartheta=\theta^{\circ}}$

$$\blacktriangleright P_{\vartheta|Y}(\|\vartheta - \theta^{\circ}\| \lesssim \mathcal{R}_{\varepsilon}) \xrightarrow{P_{\theta^{\circ}}} 1$$

$$\blacktriangleright P_{\vartheta|Y}(\|\vartheta - \theta^{\circ}\| \gtrsim \mathcal{R}_{\varepsilon}) \xrightarrow{P_{\theta^{\circ}}} 0$$

Remarks:

- $\mathcal{R}_{\varepsilon}$ depends on the prior



A glimpse to the essential: exact posterior concentration

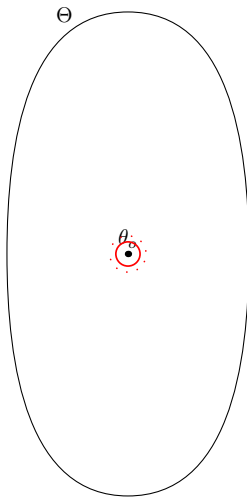
Given prior P_{ϑ} find $\mathcal{R}_{\varepsilon} \rightarrow 0$ as $\varepsilon \rightarrow 0$ such that for $P_{\theta^{\circ}} = P_{Y|\vartheta=\theta^{\circ}}$

$$\triangleright P_{\vartheta|Y}(\|\vartheta - \theta^{\circ}\| \lesssim \mathcal{R}_{\varepsilon}) \xrightarrow{P_{\theta^{\circ}}} 1$$

$$\triangleright P_{\vartheta|Y}(\|\vartheta - \theta^{\circ}\| \gtrsim \mathcal{R}_{\varepsilon}) \xrightarrow{P_{\theta^{\circ}}} 0$$

Remarks:

- $\mathcal{R}_{\varepsilon}$ depends on the prior
- $\mathcal{R}_{\varepsilon}$ might be arbitrarily slow



A glimpse to the essential: exact posterior concentration

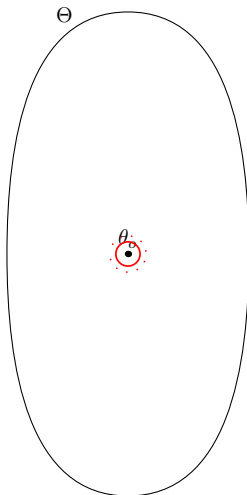
Given prior P_{ϑ} find $\mathcal{R}_{\varepsilon} \rightarrow 0$ as $\varepsilon \rightarrow 0$ such that for $P_{\theta^{\circ}} = P_{Y|\vartheta=\theta^{\circ}}$

$$\triangleright P_{\vartheta|Y}(\|\vartheta - \theta^{\circ}\| \lesssim \mathcal{R}_{\varepsilon}) \xrightarrow{P_{\theta^{\circ}}} 1$$

$$\triangleright P_{\vartheta|Y}(\|\vartheta - \theta^{\circ}\| \gtrsim \mathcal{R}_{\varepsilon}) \xrightarrow{P_{\theta^{\circ}}} 0$$

Remarks:

- ▶ $\mathcal{R}_{\varepsilon}$ depends on the prior
- ▶ $\mathcal{R}_{\varepsilon}$ might be arbitrarily slow
- ▶ consistency could fail



-  S. Ghosal, J.K. Ghosh and A.W. Van Der Vaart (2000) *Convergence rates of posterior distributions*. The Annals of Statistics 28(2):500–531
-  X. Shen and L. Wasserman. (2001) *Rates of convergence of posterior distributions*. The Annals of Statistics, 29:687–714
-  B. Knapik, A. Van der Vaart, and J. Van Zanten (2011) *Bayesian inverse problems with gaussian priors*. The Annals of Statistics, 39:2626–2657
-  J. Arbel, G. Gayraud and J. Rousseau (2013) *Bayesian optimal adaptive estimation using a sieve prior*. Scandinavian Journal of Statistics, 40:549–570
-  K. Ray (2013) *Bayesian inverse problems with non-conjugate priors*. Electronic Journal of Statistics, 7: 2516–2549
-  M. Hoffmann, J. Rousseau and J. Schmidt-Hieber (2015) *On adaptive posterior concentration rates*. The Annals of Statistics, 43:5,2259–2295
-  I. Castillo (2008) *Lower bounds for posterior rates with Gaussian process priors*. Electronic Journal of Statistics, 2:1281–1299

Objective: exact posterior concentration

- Construct a family $\{P_{\vartheta^m}\}_m$ of prior distributions with **exact** posterior concentration $\mathcal{R}_\varepsilon^{m_\varepsilon} := \mathcal{R}_\varepsilon^{m_\varepsilon}(\theta^\circ)$ for $\theta^\circ \in \Theta$, i.e.,

$$\lim_{\varepsilon \rightarrow 0} \mathbb{E}_{\theta^\circ} P_{\vartheta^{m_\varepsilon}} | \mathcal{Y} \left(\mathcal{R}_\varepsilon^{m_\varepsilon} \lesssim \|\vartheta^{m_\varepsilon} - \theta^\circ\|^2 \lesssim \mathcal{R}_\varepsilon^{m_\varepsilon} \right) = 1;$$

Objective: exact posterior concentration

- Construct a family $\{P_{\vartheta^m}\}_m$ of prior distributions with **exact** posterior concentration $\mathcal{R}_\varepsilon^{m_\varepsilon} := \mathcal{R}_\varepsilon^{m_\varepsilon}(\theta^\circ)$ for $\theta^\circ \in \Theta$, i.e.,

$$\lim_{\varepsilon \rightarrow 0} \mathbb{E}_{\theta^\circ} P_{\vartheta^{m_\varepsilon}} | \mathcal{Y} \left(\mathcal{R}_\varepsilon^{m_\varepsilon} \lesssim \|\vartheta^{m_\varepsilon} - \theta^\circ\|^2 \lesssim \mathcal{R}_\varepsilon^{m_\varepsilon} \right) = 1;$$

- Consider for a given $\theta^\circ \in \Theta$ the **oracle** rate

$$\mathcal{R}_\varepsilon^\circ := \mathcal{R}_\varepsilon^\circ(\theta^\circ) := \inf_m \mathcal{R}_\varepsilon^m(\theta^\circ).$$

Objective: exact posterior concentration

- Construct a family $\{P_{\vartheta^m}\}_m$ of prior distributions with **exact** posterior concentration $\mathcal{R}_\varepsilon^{m_\varepsilon} := \mathcal{R}_\varepsilon^{m_\varepsilon}(\theta^\circ)$ for $\theta^\circ \in \Theta$, i.e.,

$$\lim_{\varepsilon \rightarrow 0} \mathbb{E}_{\theta^\circ} P_{\vartheta^{m_\varepsilon}} | \mathcal{Y} \left(\mathcal{R}_\varepsilon^{m_\varepsilon} \lesssim \|\vartheta^{m_\varepsilon} - \theta^\circ\|^2 \lesssim \mathcal{R}_\varepsilon^{m_\varepsilon} \right) = 1;$$

- Consider for a given $\theta^\circ \in \Theta$ the **oracle** rate

$$\mathcal{R}_\varepsilon^\circ := \mathcal{R}_\varepsilon^\circ(\theta^\circ) := \inf_m \mathcal{R}_\varepsilon^m(\theta^\circ).$$

- Consider for a given $\Theta_a \subset \Theta$ the **minimax** rate

$$\mathcal{R}_\varepsilon^\star := \mathcal{R}_\varepsilon^\star(a) := \inf_m \sup_{\theta^\circ \in \Theta_a} \mathcal{R}_\varepsilon^m(\theta^\circ).$$

Objective: optimal prior choice

- A sub-family $\{P_{\vartheta^{m_\varepsilon^\circ}}\}_{m_\varepsilon^\circ}$ with **exact** posterior concentration $\mathcal{R}_\varepsilon^\circ$, i.e.

$$\lim_{\varepsilon \rightarrow 0} \mathbb{E}_{\theta^\circ} P_{\vartheta^{m_\varepsilon^\circ}} | \mathcal{Y} \left(\mathcal{R}_\varepsilon^\circ \lesssim \|\vartheta^{m_\varepsilon^\circ} - \theta^\circ\|^2 \lesssim \mathcal{R}_\varepsilon^\circ \right) = 1,$$

is called **oracle** prior and **adaptive** if it does not depend on θ° .

Objective: optimal prior choice

- A sub-family $\{P_{\vartheta^{m_\varepsilon^\circ}}\}_{m_\varepsilon^\circ}$ with **exact** posterior concentration $\mathcal{R}_\varepsilon^\circ$, i.e.

$$\lim_{\varepsilon \rightarrow 0} \mathbb{E}_{\theta^\circ} P_{\vartheta^{m_\varepsilon^\circ}}|_Y \left(\mathcal{R}_\varepsilon^\circ \lesssim \|\vartheta^{m_\varepsilon^\circ} - \theta^\circ\|^2 \lesssim \mathcal{R}_\varepsilon^\circ \right) = 1,$$

is called **oracle** prior and **adaptive** if it does not depend on θ° .

- A sub-family $\{P_{\vartheta^{m_\varepsilon^\star}}\}_{m_\varepsilon^\star}$ with **exact** posterior concentration $\mathcal{R}_\varepsilon^\star$, i.e.

$$\lim_{\varepsilon \rightarrow 0} \inf_{\theta^\circ \in \Theta_\alpha} \mathbb{E}_{\theta^\circ} P_{\vartheta^{m_\varepsilon^\star}}|_Y \left(\mathcal{R}_\varepsilon^\star \lesssim \|\vartheta^{m_\varepsilon^\star} - \theta^\circ\|^2 \lesssim \mathcal{R}_\varepsilon^\star \right) = 1.$$

is called **minimax** prior and **adaptive** if it depends not on Θ_α .

Outline

- Introduction
- Frequentist perspective reviewed
- Posterior concentration
- Adaptive posterior concentration
- Adaptive Bayes estimator

Projection estimator reviewed: oracle optimality

Observations $Y_j = \lambda_j \theta_j^\circ + \sqrt{\varepsilon} Z_j, j \geq 1$

- $\hat{\theta}^m = (\hat{\theta}_j^m)_{j \geq 1} = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator

Projection estimator reviewed: oracle optimality

Observations $Y_j = \lambda_j \theta_j^\circ + \sqrt{\varepsilon} Z_j, j \geq 1$

■ $\hat{\theta}^m = (\hat{\theta}_j^m)_{j \geq 1} = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator

► Risk: $\mathbb{E}_{\theta^\circ} \|\hat{\theta}^m - \theta^\circ\|^2 = \sum_{j=1}^m \text{Var} \hat{\theta}_j^m + \sum_{j>m} (\theta_j^\circ)^2$

Projection estimator reviewed: oracle optimality

Observations $Y_j = \lambda_j \theta_j^\circ + \sqrt{\varepsilon} Z_j, j \geq 1$

■ $\hat{\theta}^m = (\hat{\theta}_j^m)_{j \geq 1} = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator

► Risk: $\mathbb{E}_{\theta^\circ} \|\hat{\theta}^m - \theta^\circ\|^2 = \varepsilon \sum_{j=1}^m \lambda_j^{-2} + \sum_{j>m} (\theta_j^\circ)^2$

Projection estimator reviewed: oracle optimality

Observations $Y_j = \lambda_j \theta_j^\circ + \sqrt{\varepsilon} Z_j, j \geq 1$

■ $\hat{\theta}^m = (\hat{\theta}_j^m)_{j \geq 1} = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator

■ $\Lambda_j := \lambda_j^{-2},$

► Risk: $\mathbb{E}_{\theta^\circ} \|\hat{\theta}^m - \theta^\circ\|^2 = \varepsilon \sum_{j=1}^m \Lambda_j + \sum_{j>m} (\theta_j^\circ)^2$

Projection estimator reviewed: oracle optimality

Observations $Y_j = \lambda_j \theta_j^\circ + \sqrt{\varepsilon} Z_j, j \geq 1$

■ $\hat{\theta}^m = (\hat{\theta}_j^m)_{j \geq 1} = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator

■ $\Lambda_j := \lambda_j^{-2}, \quad \bar{\Lambda}_m := \frac{1}{m} \sum_{j=1}^m \Lambda_j$

► Risk: $\mathbb{E}_{\theta^\circ} \|\hat{\theta}^m - \theta^\circ\|^2 = \varepsilon m \bar{\Lambda}_m + \sum_{j > m} (\theta_j^\circ)^2$

Projection estimator reviewed: oracle optimality

Observations $Y_j = \lambda_j \theta_j^\circ + \sqrt{\varepsilon} Z_j, j \geq 1$

■ $\hat{\theta}^m = (\hat{\theta}_j^m)_{j \geq 1} = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator

■ $\Lambda_j := \lambda_j^{-2}, \quad \bar{\Lambda}_m := \frac{1}{m} \sum_{j=1}^m \Lambda_j \quad \text{and} \quad \mathfrak{b}_m^2 := \mathfrak{b}_m^2(\theta^\circ) := \sum_{j>m} (\theta_j^\circ)^2$

► Risk: $\mathbb{E}_{\theta^\circ} \|\hat{\theta}^m - \theta^\circ\|^2 = \varepsilon m \bar{\Lambda}_m + \mathfrak{b}_m^2$

Projection estimator reviewed: oracle optimality

Observations $Y_j = \lambda_j \theta_j^\circ + \sqrt{\varepsilon} Z_j, j \geq 1$

■ $\hat{\theta}^m = (\hat{\theta}_j^m)_{j \geq 1} = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator

■ $\Lambda_j := \lambda_j^{-2}, \quad \bar{\Lambda}_m := \frac{1}{m} \sum_{j=1}^m \Lambda_j \quad \text{and} \quad \mathfrak{b}_m^2 := \mathfrak{b}_m^2(\theta^\circ) := \sum_{j>m} (\theta_j^\circ)^2$

► Risk: $[\varepsilon m \bar{\Lambda}_m \vee \mathfrak{b}_m^2] \leq \mathbb{E}_{\theta^\circ} \|\hat{\theta}^m - \theta^\circ\|^2 = \varepsilon m \bar{\Lambda}_m + \mathfrak{b}_m^2$

Projection estimator reviewed: oracle optimality

Observations $Y_j = \lambda_j \theta_j^\circ + \sqrt{\varepsilon} Z_j, j \geq 1$

■ $\hat{\theta}^m = (\hat{\theta}_j^m)_{j \geq 1} = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator

■ $\Lambda_j := \lambda_j^{-2}, \quad \bar{\Lambda}_m := \frac{1}{m} \sum_{j=1}^m \Lambda_j \quad \text{and} \quad \mathfrak{b}_m^2 := \mathfrak{b}_m^2(\theta^\circ) := \sum_{j>m} (\theta_j^\circ)^2$

► Risk: $[\varepsilon m \bar{\Lambda}_m \vee \mathfrak{b}_m^2] \leq \mathbb{E}_{\theta^\circ} \|\hat{\theta}^m - \theta^\circ\|^2 \lesssim [\varepsilon m \bar{\Lambda}_m \vee \mathfrak{b}_m^2] =: \mathcal{R}_\varepsilon^m(\theta^\circ)$

Projection estimator reviewed: oracle optimality

Observations $Y_j = \lambda_j \theta_j^\circ + \sqrt{\varepsilon} Z_j, j \geq 1$

■ $\hat{\theta}^m = (\hat{\theta}_j^m)_{j \geq 1} = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator

■ $\Lambda_j := \lambda_j^{-2}, \quad \bar{\Lambda}_m := \frac{1}{m} \sum_{j=1}^m \Lambda_j \quad \text{and} \quad \mathfrak{b}_m^2 := \mathfrak{b}_m^2(\theta^\circ) := \sum_{j>m} (\theta_j^\circ)^2$

► Risk: $[\varepsilon m \bar{\Lambda}_m \vee \mathfrak{b}_m^2] \leq \mathbb{E}_{\theta^\circ} \|\hat{\theta}^m - \theta^\circ\|^2 \lesssim [\varepsilon m \bar{\Lambda}_m \vee \mathfrak{b}_m^2] =: \mathcal{R}_\varepsilon^m(\theta^\circ)$

■ oracle cut-off $m_\varepsilon^\circ := m_\varepsilon^\circ(\theta^\circ) := \arg \min_{m \geq 1} \{\mathcal{R}_\varepsilon^m(\theta^\circ)\}$

Projection estimator reviewed: oracle optimality

Observations $Y_j = \lambda_j \theta_j^\circ + \sqrt{\varepsilon} Z_j, j \geq 1$

■ $\hat{\theta}^m = (\hat{\theta}_j^m)_{j \geq 1} = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator

■ $\Lambda_j := \lambda_j^{-2}, \quad \bar{\Lambda}_m := \frac{1}{m} \sum_{j=1}^m \Lambda_j \quad \text{and} \quad \mathfrak{b}_m^2 := \mathfrak{b}_m^2(\theta^\circ) := \sum_{j>m} (\theta_j^\circ)^2$

► Risk: $[\varepsilon m \bar{\Lambda}_m \vee \mathfrak{b}_m^2] \leq \mathbb{E}_{\theta^\circ} \|\hat{\theta}^m - \theta^\circ\|^2 \lesssim [\varepsilon m \bar{\Lambda}_m \vee \mathfrak{b}_m^2] =: \mathcal{R}_\varepsilon^m(\theta^\circ)$

■ oracle cut-off $m_\varepsilon^\circ := m_\varepsilon^\circ(\theta^\circ) := \arg \min_{m \geq 1} \{\mathcal{R}_\varepsilon^m(\theta^\circ)\}$

■ oracle rate $\mathcal{R}_\varepsilon^\circ := \mathcal{R}_\varepsilon^\circ(\theta^\circ) := [\varepsilon m_\varepsilon^\circ \bar{\Lambda}_{m_\varepsilon^\circ} \vee \mathfrak{b}_{m_\varepsilon^\circ}^2] = \min_{m \geq 1} \mathcal{R}_\varepsilon^m(\theta^\circ)$

Projection estimator reviewed: oracle optimality

Observations $Y_j = \lambda_j \theta_j^\circ + \sqrt{\varepsilon} Z_j, j \geq 1$

- $\hat{\theta}^m = (\hat{\theta}_j^m)_{j \geq 1} = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator
- $\Lambda_j := \lambda_j^{-2}, \quad \bar{\Lambda}_m := \frac{1}{m} \sum_{j=1}^m \Lambda_j \quad \text{and} \quad \mathfrak{b}_m^2 := \mathfrak{b}_m^2(\theta^\circ) := \sum_{j>m} (\theta_j^\circ)^2$

► Risk: $[\varepsilon m \bar{\Lambda}_m \vee \mathfrak{b}_m^2] \leq \mathbb{E}_{\theta^\circ} \|\hat{\theta}^m - \theta^\circ\|^2 \lesssim [\varepsilon m \bar{\Lambda}_m \vee \mathfrak{b}_m^2] =: \mathcal{R}_\varepsilon^m(\theta^\circ)$

- oracle cut-off $m_\varepsilon^\circ := m_\varepsilon^\circ(\theta^\circ) := \arg \min_{m \geq 1} \{\mathcal{R}_\varepsilon^m(\theta^\circ)\}$
- oracle rate $\mathcal{R}_\varepsilon^\circ := \mathcal{R}_\varepsilon^\circ(\theta^\circ) := [\varepsilon m_\varepsilon^\circ \bar{\Lambda}_{m_\varepsilon^\circ} \vee \mathfrak{b}_{m_\varepsilon^\circ}^2] = \min_{m \geq 1} \mathcal{R}_\varepsilon^m(\theta^\circ)$
- oracle estimator $\hat{\theta}^{m_\varepsilon^\circ}$, i.e., $\mathbb{E}_{\theta^\circ} \|\hat{\theta}^{m_\varepsilon^\circ} - \theta^\circ\|^2 \lesssim \mathcal{R}_\varepsilon^\circ$

Illustration: typical smoothness condition

Consider a non-increasing weight sequence $\mathbf{a} = (a_j)_{j \geq 1}$ and let

$$\Theta_{\mathbf{a}} := \left\{ \theta \in \Theta : \sum_{j=1}^{\infty} \theta_j^2 / a_j =: \|\theta\|_{\mathbf{a}}^2 \leq r \right\}$$

Illustration: typical smoothness condition

Consider a non-increasing weight sequence $\mathbf{a} = (a_j)_{j \geq 1}$ and let

$$\Theta_{\mathbf{a}} := \left\{ \theta \in \Theta : \sum_{j=1}^{\infty} \theta_j^2 / a_j =: \|\theta\|_{\mathbf{a}}^2 \leq r \right\}$$

Illustration

- ▶ $(u_j)_j$ trigonometric basis
- ▶ $a_1 = 1, a_{2j} = (2j)^{-2p}, a_{2j+1} = (2j)^{-2p}, p > 0$

Illustration: typical smoothness condition

Consider a non-increasing weight sequence $\mathbf{a} = (a_j)_{j \geq 1}$ and let

$$\Theta_{\mathbf{a}} := \left\{ \theta \in \Theta : \sum_{j=1}^{\infty} \theta_j^2 / a_j =: \|\theta\|_{\mathbf{a}}^2 \leq r \right\}$$

Illustration

- ▶ $(u_j)_j$ trigonometric basis
- ▶ $a_1 = 1, a_{2j} = (2j)^{-2p}, a_{2j+1} = (2j)^{-2p}, p > 0$
 - Sobolev ellipsoid of p -times differentiable, absolutely continuous, periodic functions

Illustration: typical smoothness condition

Consider a non-increasing weight sequence $\mathbf{a} = (a_j)_{j \geq 1}$ and let

$$\Theta_{\mathbf{a}} := \left\{ \theta \in \Theta : \sum_{j=1}^{\infty} \theta_j^2 / a_j =: \|\theta\|_{\mathbf{a}}^2 \leq r \right\}$$

Illustration

- ▶ $(u_j)_j$ trigonometric basis
- ▶ $a_1 = 1, a_{2j} = (2j)^{-2p}, a_{2j+1} = (2j)^{-2p}, p > 0$
 - Sobolev ellipsoid of p -times differentiable, absolutely continuous, periodic functions
- ▶ $a_j = \exp(1 - |j|^{2p}), p > \frac{1}{2}$
 - analytical functions

Projection estimator reviewed: minimax optimality

Observations $Y_j = \lambda_j \theta_{\circ j} + \sqrt{\varepsilon} Z_j, j \geq 1$

■ $\hat{\theta}^m = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator

■ $\Lambda_j := \lambda_j^{-2}, \bar{\Lambda}_m := \frac{1}{m} \sum_{j=1}^m \Lambda_j$ and $\mathfrak{b}_m^2 := \mathfrak{b}_m^2(\theta^\circ) := \sum_{j>m} (\theta_j^\circ)^2$

► Risk: $\mathbb{E}_{\theta^\circ} \|\hat{\theta}^m - \theta^\circ\|^2 \lesssim [\varepsilon m \bar{\Lambda}_m \vee \mathfrak{b}_m^2]$

Projection estimator reviewed: minimax optimality

Observations $Y_j = \lambda_j \theta_{\circ j} + \sqrt{\varepsilon} Z_j, j \geq 1$

■ $\hat{\theta}^m = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator

■ $\Lambda_j := \lambda_j^{-2}, \bar{\Lambda}_m := \frac{1}{m} \sum_{j=1}^m \Lambda_j$ and $\mathfrak{b}_m^2 := \mathfrak{b}_m^2(\theta^\circ) := \sum_{j>m} (\theta_j^\circ)^2$

► Risk: $\mathbb{E}_{\theta^\circ} \|\hat{\theta}^m - \theta^\circ\|^2 \lesssim [\varepsilon m \bar{\Lambda}_m \vee \mathfrak{b}_m^2] \leq (1 \vee r) [\varepsilon m \bar{\Lambda}_m \vee \mathfrak{a}_m]$

Projection estimator reviewed: minimax optimality

Observations $Y_j = \lambda_j \theta_{\circ j} + \sqrt{\varepsilon} Z_j, j \geq 1$

■ $\hat{\theta}^m = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator

■ $\Lambda_j := \lambda_j^{-2}, \bar{\Lambda}_m := \frac{1}{m} \sum_{j=1}^m \Lambda_j$ and $\mathfrak{b}_m^2 := \mathfrak{b}_m^2(\theta^\circ) := \sum_{j>m} (\theta_j^\circ)^2$

► Maximal risk: $\sup_{\theta^\circ \in \Theta_a} \mathbb{E}_{\theta^\circ} \|\hat{\theta}^m - \theta^\circ\|^2 \lesssim [\varepsilon m \bar{\Lambda}_m \vee \mathfrak{a}_m]$

Projection estimator reviewed: minimax optimality

Observations $Y_j = \lambda_j \theta_{\circ j} + \sqrt{\varepsilon} Z_j, j \geq 1$

■ $\hat{\theta}^m = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator

■ $\Lambda_j := \lambda_j^{-2}, \bar{\Lambda}_m := \frac{1}{m} \sum_{j=1}^m \Lambda_j$ and $\mathfrak{b}_m^2 := \mathfrak{b}_m^2(\theta^\circ) := \sum_{j>m} (\theta_j^\circ)^2$

► Maximal risk: $\sup_{\theta^\circ \in \Theta_{\mathfrak{a}}} \mathbb{E}_{\theta^\circ} \|\hat{\theta}^m - \theta^\circ\|^2 \lesssim [\varepsilon m \bar{\Lambda}_m \vee \mathfrak{a}_m]$

■ minimax cut-off $m_\varepsilon^* := m_\varepsilon^*(\mathfrak{a}) := \arg \min_{m \geq 1} \{[\varepsilon m \bar{\Lambda}_m \vee \mathfrak{a}_m]\}$

Projection estimator reviewed: minimax optimality

Observations $Y_j = \lambda_j \theta_{\circ j} + \sqrt{\varepsilon} Z_j, j \geq 1$

■ $\hat{\theta}^m = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator

■ $\Lambda_j := \lambda_j^{-2}, \bar{\Lambda}_m := \frac{1}{m} \sum_{j=1}^m \Lambda_j$ and $\mathfrak{b}_m^2 := \mathfrak{b}_m^2(\theta^\circ) := \sum_{j>m} (\theta_j^\circ)^2$

► Maximal risk: $\sup_{\theta^\circ \in \Theta_{\mathfrak{a}}} \mathbb{E}_{\theta^\circ} \|\hat{\theta}^m - \theta^\circ\|^2 \lesssim [\varepsilon m \bar{\Lambda}_m \vee \mathfrak{a}_m]$

■ minimax cut-off $m_\varepsilon^* := m_\varepsilon^*(\mathfrak{a}) := \arg \min_{m \geq 1} \{[\varepsilon m \bar{\Lambda}_m \vee \mathfrak{a}_m]\}$

■ minimax rate $\mathcal{R}_\varepsilon^* := \mathcal{R}_\varepsilon^*(\mathfrak{a}) := [\varepsilon m_\varepsilon^* \bar{\Lambda}_{m_\varepsilon^*} \vee \mathfrak{a}_{m_\varepsilon^*}]$

Projection estimator reviewed: minimax optimality

Observations $Y_j = \lambda_j \theta_{\circ j} + \sqrt{\varepsilon} Z_j, j \geq 1$

- $\hat{\theta}^m = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator
- $\Lambda_j := \lambda_j^{-2}, \bar{\Lambda}_m := \frac{1}{m} \sum_{j=1}^m \Lambda_j$ and $\mathfrak{b}_m^2 := \mathfrak{b}_m^2(\theta^\circ) := \sum_{j>m} (\theta_j^\circ)^2$

► Maximal risk: $\sup_{\theta^\circ \in \Theta_a} \mathbb{E}_{\theta^\circ} \|\hat{\theta}^m - \theta^\circ\|^2 \lesssim [\varepsilon m \bar{\Lambda}_m \vee \mathfrak{a}_m]$

- minimax cut-off $m_\varepsilon^* := m_\varepsilon^*(\mathfrak{a}) := \arg \min_{m \geq 1} \{[\varepsilon m \bar{\Lambda}_m \vee \mathfrak{a}_m]\}$
- minimax rate $\mathcal{R}_\varepsilon^* := \mathcal{R}_\varepsilon^*(\mathfrak{a}) := [\varepsilon m_\varepsilon^* \bar{\Lambda}_{m_\varepsilon^*} \vee \mathfrak{a}_{m_\varepsilon^*}]$
- minimax estimator $\hat{\theta}^{m_\varepsilon^*}$, i.e., $\sup_{\theta^\circ \in \Theta_a} \mathbb{E}_{\theta^\circ} \|\hat{\theta}^{m_\varepsilon^*} - \theta^\circ\|^2 \lesssim \mathcal{R}_\varepsilon^*$

Projection estimator reviewed: minimax optimality

Observations $Y_j = \lambda_j \theta_{\circ j} + \sqrt{\varepsilon} Z_j, j \geq 1$

■ $\hat{\theta}^m = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator

■ $\Lambda_j := \lambda_j^{-2}, \bar{\Lambda}_m := \frac{1}{m} \sum_{j=1}^m \Lambda_j$ and $\mathfrak{b}_m^2 := \mathfrak{b}_m^2(\theta^\circ) := \sum_{j>m} (\theta_j^\circ)^2$

► Maximal risk: $\sup_{\theta^\circ \in \Theta_a} \mathbb{E}_{\theta^\circ} \|\hat{\theta}^m - \theta^\circ\|^2 \lesssim [\varepsilon m \bar{\Lambda}_m \vee \mathfrak{a}_m]$

■ minimax cut-off $m_\varepsilon^* := m_\varepsilon^*(\mathfrak{a}) := \arg \min_{m \geq 1} \{[\varepsilon m \bar{\Lambda}_m \vee \mathfrak{a}_m]\}$

■ minimax rate $\mathcal{R}_\varepsilon^* := \mathcal{R}_\varepsilon^*(\mathfrak{a}) := [\varepsilon m_\varepsilon^* \bar{\Lambda}_{m_\varepsilon^*} \vee \mathfrak{a}_{m_\varepsilon^*}]$

■ minimax estimator $\hat{\theta}^{m_\varepsilon^*}$, i.e., $\sup_{\theta^\circ \in \Theta_a} \mathbb{E}_{\theta^\circ} \|\hat{\theta}^{m_\varepsilon^*} - \theta^\circ\|^2 \lesssim \mathcal{R}_\varepsilon^*$

(JJ & Schwarz, 2013)

► Lower bound: $\inf_{\hat{\theta}} \sup_{\theta^\circ \in \Theta_a} \mathbb{E} \|\hat{\theta} - \theta^\circ\|^2 \gtrsim \mathcal{R}_\varepsilon^*$, if

$$\inf_\varepsilon [\varepsilon m_\varepsilon^* \bar{\Lambda}_{m_\varepsilon^*} \wedge \mathfrak{a}_{m_\varepsilon^*}] / [\varepsilon m_\varepsilon^* \bar{\Lambda}_{m_\varepsilon^*} \vee \mathfrak{a}_{m_\varepsilon^*}] > 0$$

Projection estimator reviewed: adaptation

Observations $Y_j = \lambda_j \theta_{\circ j} + \sqrt{\varepsilon} Z_j, j \geq 1$

■ $\hat{\theta}^m = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator

■ $\Lambda_j := \lambda_j^{-2}, \bar{\Lambda}_m := \frac{1}{m} \sum_{j=1}^m \Lambda_j$ and $\mathfrak{b}_m^2 := \mathfrak{b}_m^2(\theta^\circ) := \sum_{j>m} (\theta_j^\circ)^2$

► Select m by using a penalised minimum contrast criterion

Projection estimator reviewed: adaptation

Observations $Y_j = \lambda_j \theta_{\circ j} + \sqrt{\varepsilon} Z_j, j \geq 1$

■ $\hat{\theta}^m = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator

■ $\Lambda_j := \lambda_j^{-2}, \bar{\Lambda}_m := \frac{1}{m} \sum_{j=1}^m \Lambda_j$ and $\mathfrak{b}_m^2 := \mathfrak{b}_m^2(\theta^\circ) := \sum_{j>m} (\theta_j^\circ)^2$

► Select $\hat{m}$ by using a penalised minimum contrast criterion

$$\hat{m} := \arg \min_{1 \leq m \leq M} \{ -\|\hat{\theta}^m\| + 10\varepsilon m \bar{\Lambda}_m \}$$

$$M := \max \{ 1 \leq m \leq \lfloor \varepsilon^{-1} \rfloor : \varepsilon m \bar{\Lambda}_m \leq \Lambda_1 \}$$

Projection estimator reviewed: adaptation

Observations $Y_j = \lambda_j \theta_{\circ j} + \sqrt{\varepsilon} Z_j, j \geq 1$

■ $\hat{\theta}^m = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator

■ $\Lambda_j := \lambda_j^{-2}, \bar{\Lambda}_m := \frac{1}{m} \sum_{j=1}^m \Lambda_j$ and $\mathfrak{b}_m^2 := \mathfrak{b}_m^2(\theta^\circ) := \sum_{j>m} (\theta_j^\circ)^2$

► Select $\hat{m}$ by using a penalised minimum contrast criterion

$$\hat{m} := \arg \min_{1 \leq m \leq M} \{ -\|\hat{\theta}^m\| + 10\varepsilon m \bar{\Lambda}_m \}$$

$$M := \max \{ 1 \leq m \leq \lfloor \varepsilon^{-1} \rfloor : \varepsilon m \bar{\Lambda}_m \leq \Lambda_1 \}$$

■ minimax optimal data-driven estimator, if $1 \leq \Lambda_{(m)}/\bar{\Lambda}_m \leq L_\lambda$.

$$\sup_{\theta^\circ \in \Theta_\alpha} \mathbb{E}_{\theta^\circ} \|\hat{\theta}^{\hat{m}} - \theta^\circ\|^2 \lesssim \mathcal{R}_\varepsilon^*(\alpha) = [\varepsilon m_\varepsilon^* \bar{\Lambda}_{m_\varepsilon^*} \vee \alpha_{m_\varepsilon^*}]$$

Projection estimator reviewed: adaptation

Observations $Y_j = \lambda_j \theta_{\circ j} + \sqrt{\varepsilon} Z_j, j \geq 1$

■ $\hat{\theta}^m = (Y_1/\lambda_1, \dots, Y_m/\lambda_m, 0, \dots)$ projection estimator

■ $\Lambda_j := \lambda_j^{-2}, \bar{\Lambda}_m := \frac{1}{m} \sum_{j=1}^m \Lambda_j$ and $\mathfrak{b}_m^2 := \mathfrak{b}_m^2(\theta^\circ) := \sum_{j>m} (\theta_j^\circ)^2$

► Select $\hat{m}$ by using a penalised minimum contrast criterion

$$\hat{m} := \arg \min_{1 \leq m \leq M} \{ -\|\hat{\theta}^m\| + 10\varepsilon m \bar{\Lambda}_m \}$$

$$M := \max \{ 1 \leq m \leq \lfloor \varepsilon^{-1} \rfloor : \varepsilon m \bar{\Lambda}_m \leq \Lambda_1 \}$$

■ minimax optimal data-driven estimator, if $1 \leq \Lambda_{(m)}/\bar{\Lambda}_m \leq L_\lambda$.

$$\sup_{\theta^\circ \in \Theta_\alpha} \mathbb{E}_{\theta^\circ} \|\hat{\theta}^{\hat{m}} - \theta^\circ\|^2 \lesssim \mathcal{R}_\varepsilon^*(\alpha) = [\varepsilon m_\varepsilon^* \bar{\Lambda}_{m_\varepsilon^*} \vee \alpha_{m_\varepsilon^*}]$$

Outline

- Introduction
- Frequentist perspective reviewed
- Posterior concentration
- Adaptive posterior concentration
- Adaptive Bayes estimator

Objective: optimal prior choice

- A sub-family $\{P_{\vartheta^{m_\varepsilon^\circ}}\}_{m_\varepsilon^\circ}$ with **exact** posterior concentration $\mathcal{R}_\varepsilon^\circ$, i.e.

$$\lim_{\varepsilon \rightarrow 0} \mathbb{E}_{\theta^\circ} P_{\vartheta^{m_\varepsilon^\circ}} | \mathcal{Y} \left(\mathcal{R}_\varepsilon^\circ \lesssim \|\vartheta^{m_\varepsilon^\circ} - \theta^\circ\|^2 \lesssim \mathcal{R}_\varepsilon^\circ \right) = 1,$$

is called **oracle** prior and **adaptive** if it does not depend on θ° .

- A sub-family $\{P_{\vartheta^{m_\varepsilon^\star}}\}_{m_\varepsilon^\star}$ with **exact** posterior concentration $\mathcal{R}_\varepsilon^\star$, i.e.

$$\lim_{\varepsilon \rightarrow 0} \inf_{\theta^\circ \in \Theta_\alpha} \mathbb{E}_{\theta^\circ} P_{\vartheta^{m_\varepsilon^\star}} | \mathcal{Y} \left(\mathcal{R}_\varepsilon^\star \lesssim \|\vartheta^{m_\varepsilon^\star} - \theta^\circ\|^2 \lesssim \mathcal{R}_\varepsilon^\star \right) = 1.$$

is called **minimax** prior and **adaptive** if it depends not on Θ_α .

Sieve family of prior distributions

- Gaussian prior distribution for $\boldsymbol{\vartheta} = (\vartheta_j)_{j \geq 1}$
 - ▶ $\{\vartheta_j\}_{j \geq 1}$ are independent, normally distributed
 - ▶ prior means 0 and prior variances $\varsigma = (\varsigma_j)_{j \geq 1}$:
 - $\vartheta_j \sim \mathcal{N}(0, \varsigma_j)$, independent, $j \in \mathbb{N}$.

Sieve family of prior distributions

- Gaussian prior distribution for $\boldsymbol{\vartheta} = (\vartheta_j)_{j \geq 1}$
 - ▶ $\{\vartheta_j\}_{j \geq 1}$ are independent, normally distributed
 - ▶ prior means 0 and prior variances $\varsigma = (\varsigma_j)_{j \geq 1}$:
 - $\vartheta_j \sim \mathcal{N}(0, \varsigma_j)$, independent, $j \in \mathbb{N}$.
- Sieve family of prior distributions $\{P_{\boldsymbol{\vartheta}^m}\}_m$ depending on a hyperparameter m :
 - ▶ first m random parameters $\{\vartheta_j\}_{j=1}^m$ non-degenerate
 - ▶ independent random variables $\{\vartheta_j^m\}_{j \geq 1}$ with marginals:
 - $\vartheta_j^m \sim \mathcal{N}(0, \varsigma_j)$ for $1 \leq j \leq m$ and
 - $\vartheta_j^m \sim \delta_0$ for $j > m$.

Posterior distribution

- posterior distribution of $\boldsymbol{\vartheta} = (\vartheta_j)_{j \geq 1}$ given $Y = (Y_j)_{j \geq 1}$
 - ▶ $\{\vartheta_j\}_{j \geq 1}$ are conditionally independent, normally distributed
 - posterior variance $\sigma_j := \text{Var}[\vartheta_j | Y] = (\lambda_j^2 \varepsilon^{-1} + \varsigma_j^{-1})^{-1}$
 - posterior mean $\theta_j^Y := \mathbb{E}[\vartheta_j | Y] = \sigma_j (\lambda_j \varepsilon^{-1} Y_j)$

Posterior distribution

- posterior distribution of $\boldsymbol{\vartheta} = (\vartheta_j)_{j \geq 1}$ given $Y = (Y_j)_{j \geq 1}$
 - ▶ $\{\vartheta_j\}_{j \geq 1}$ are conditionally independent, normally distributed
 - posterior variance $\sigma_j := \text{Var}[\vartheta_j | Y] = (\lambda_j^2 \varepsilon^{-1} + \varsigma_j^{-1})^{-1}$
 - posterior mean $\theta_j^Y := \mathbb{E}[\vartheta_j | Y] = \sigma_j (\lambda_j \varepsilon^{-1} Y_j)$
- posterior distribution of $\boldsymbol{\vartheta}^m = (\vartheta_j^m)_{j \geq 1}$ given Y
 - ▶ $\{\vartheta_j^m\}_{j \geq 1}$ are conditionally independent, normally distributed
 - posterior mean θ_j^Y and variance σ_j for $1 \leq j \leq m$
 - degenerate on 0 for $j > m$

Exact posterior concentration

- Let $\Lambda_j := \lambda_j^{-2}$, $\bar{\Lambda}_m := \frac{1}{m} \sum_{j=1}^m \Lambda_j$, $\Lambda_{(m)} := \max_{1 \leq j \leq m} \Lambda_j$ and $\mathfrak{b}_m^2 := \sum_{j>m} (\theta_j^\circ)^2$.

ASSUMPTION A1. (Prior variances)

Let $G_\varepsilon := \max\{1 \leq m \leq \lfloor \varepsilon^{-1} \rfloor : \varepsilon \Lambda_{(m)} \leq \Lambda_1\}$.

There exists $d > 0$ such that $\varsigma_j \geq d[(\varepsilon \Lambda_j)^{1/2} \vee \varepsilon \Lambda_j]$ for all $1 \leq j \leq G_\varepsilon$.

Exact posterior concentration

► Let $\Lambda_j := \lambda_j^{-2}$, $\bar{\Lambda}_m := \frac{1}{m} \sum_{j=1}^m \Lambda_j$, $\Lambda_{(m)} := \max_{1 \leq j \leq m} \Lambda_j$ and $\mathfrak{b}_m^2 := \sum_{j>m} (\theta_j^\circ)^2$.

ASSUMPTION A1. (Prior variances)

Let $G_\varepsilon := \max\{1 \leq m \leq \lfloor \varepsilon^{-1} \rfloor : \varepsilon \Lambda_{(m)} \leq \Lambda_1\}$.

There exists $d > 0$ such that $\varsigma_j \geq d[(\varepsilon \Lambda_j)^{1/2} \vee \varepsilon \Lambda_j]$ for all $1 \leq j \leq G_\varepsilon$.

ASSUMPTION A2. (Regular sub-family $\{P_{\vartheta^{m_\varepsilon}}\}_{m_\varepsilon}$)

$$\sup_\varepsilon (\varepsilon m_\varepsilon \Lambda_{(m_\varepsilon)}) [\varepsilon m_\varepsilon \bar{\Lambda}_{m_\varepsilon} \vee \mathfrak{b}_{m_\varepsilon}^2]^{-1} < \infty$$

Exact posterior concentration

PROPOSITION. Under Assumption A1. and A.2 holds

$$\mathbb{E}_{\theta^\circ} P_{\mathfrak{y}^{m_\varepsilon}} | Y \left(\mathcal{R}_\varepsilon^{m_\varepsilon}(\theta^\circ) \lesssim \|\mathfrak{y}^{m_\varepsilon} - \theta^\circ\|^2 \lesssim \mathcal{R}_\varepsilon^{m_\varepsilon}(\theta^\circ) \right) \geq 1 - ce^{-cm_\varepsilon}$$

Exact posterior concentration and oracle prior

PROPOSITION. Under Assumption A1. and A.2 holds

$$\mathbb{E}_{\theta^\circ} P_{\mathfrak{y}^{m_\varepsilon}} | \mathcal{Y} \left(\mathcal{R}_\varepsilon^{m_\varepsilon}(\theta^\circ) \lesssim \|\mathfrak{y}^{m_\varepsilon} - \theta^\circ\|^2 \lesssim \mathcal{R}_\varepsilon^{m_\varepsilon}(\theta^\circ) \right) \geq 1 - c e^{-c m_\varepsilon}$$

- oracle cut-off $m_\varepsilon^\circ := m_\varepsilon^\circ(\theta^\circ) := \arg \min_{m \geq 1} \{\mathcal{R}_\varepsilon^m(\theta^\circ)\}$
- oracle rate $\mathcal{R}_\varepsilon^\circ := \mathcal{R}_\varepsilon^\circ(\theta^\circ) := [\varepsilon m_\varepsilon^\circ \bar{\Lambda}_{m_\varepsilon^\circ} \vee \mathfrak{b}_{m_\varepsilon^\circ}^2] = \min_{m \geq 1} \mathcal{R}_\varepsilon^m(\theta^\circ)$

Exact posterior concentration and oracle prior

PROPOSITION. Under Assumption A1. and A.2 holds

$$\mathbb{E}_{\theta^\circ} P_{\mathfrak{y}^{m_\varepsilon}} | \mathcal{Y} \left(\mathcal{R}_\varepsilon^{m_\varepsilon}(\theta^\circ) \lesssim \|\mathfrak{y}^{m_\varepsilon} - \theta^\circ\|^2 \lesssim \mathcal{R}_\varepsilon^{m_\varepsilon}(\theta^\circ) \right) \geq 1 - ce^{-cm_\varepsilon}$$

- oracle cut-off $m_\varepsilon^\circ := m_\varepsilon^\circ(\theta^\circ) := \arg \min_{m \geq 1} \{\mathcal{R}_\varepsilon^m(\theta^\circ)\}$
- oracle rate $\mathcal{R}_\varepsilon^\circ := \mathcal{R}_\varepsilon^\circ(\theta^\circ) := [\varepsilon m_\varepsilon^\circ \bar{\Lambda}_{m_\varepsilon^\circ} \vee \mathfrak{b}_{m_\varepsilon^\circ}^2] = \min_{m \geq 1} \mathcal{R}_\varepsilon^m(\theta^\circ)$

THEOREM. Under A1., A2. and in addition $m_\varepsilon^\circ \rightarrow \infty$ then

$$\lim_{\varepsilon \rightarrow 0} \mathbb{E}_{\theta^\circ} P_{\mathfrak{y}^{m_\varepsilon^\circ}} | \mathcal{Y} \left(\mathcal{R}_\varepsilon^\circ \lesssim \|\mathfrak{y}^{m_\varepsilon^\circ} - \theta^\circ\|^2 \lesssim \mathcal{R}_\varepsilon^\circ \right) = 1$$

Exact posterior concentration and minimax prior

- minimax cut-off $m_\epsilon^* := m_\epsilon^*(\mathfrak{a}) := \arg \min_{m \geq 1} \{[\epsilon m \bar{\Lambda}_m \vee \mathfrak{a}_m]\}$
- minimax rate $\mathcal{R}_\epsilon^* := \mathcal{R}_\epsilon^*(\mathfrak{a}) := [\epsilon m_\epsilon^* \bar{\Lambda}_{m_\epsilon^*} \vee \mathfrak{a}_{m_\epsilon^*}]$

Exact posterior concentration and minimax prior

- minimax cut-off $m_\epsilon^* := m_\epsilon^*(\mathfrak{a}) := \arg \min_{m \geq 1} \{[\epsilon m \bar{\Lambda}_m \vee \mathfrak{a}_m]\}$
- minimax rate $\mathcal{R}_\epsilon^* := \mathcal{R}_\epsilon^*(\mathfrak{a}) := [\epsilon m_\epsilon^* \bar{\Lambda}_{m_\epsilon^*} \vee \mathfrak{a}_{m_\epsilon^*}]$

ASSUMPTION A3. (Regular parameter space)

$$\inf_\epsilon [\epsilon m_\epsilon^* \bar{\Lambda}_{m_\epsilon^*} \wedge \mathfrak{a}_{m_\epsilon^*}] [\epsilon m_\epsilon^* \bar{\Lambda}_{m_\epsilon^*} \vee \mathfrak{a}_{m_\epsilon^*}]^{-1} > 0$$

Exact posterior concentration and minimax prior

- minimax cut-off $m_\epsilon^* := m_\epsilon^*(\mathbf{a}) := \arg \min_{m \geq 1} \{[\epsilon m \bar{\Lambda}_m \vee \mathbf{a}_m]\}$
- minimax rate $\mathcal{R}_\epsilon^* := \mathcal{R}_\epsilon^*(\mathbf{a}) := [\epsilon m_\epsilon^* \bar{\Lambda}_{m_\epsilon^*} \vee \mathbf{a}_{m_\epsilon^*}]$

ASSUMPTION A3. (Regular parameter space)

$$\inf_\epsilon [\epsilon m_\epsilon^* \bar{\Lambda}_{m_\epsilon^*} \wedge \mathbf{a}_{m_\epsilon^*}] [\epsilon m_\epsilon^* \bar{\Lambda}_{m_\epsilon^*} \vee \mathbf{a}_{m_\epsilon^*}]^{-1} > 0$$

THEOREM. Under A1, A2 and A3 holds

$$\lim_{\epsilon \rightarrow \infty} \inf_{\theta^\circ \in \Theta_{\mathbf{a}}} \mathbb{E}_{\theta^\circ} P_{\boldsymbol{\vartheta}^{m_\epsilon^*}}|_{\mathcal{Y}} \left(\mathcal{R}_\epsilon^* \lesssim \|\boldsymbol{\vartheta}^{m_\epsilon^*} - \theta^\circ\|^2 \lesssim \mathcal{R}_\epsilon^* \right) = 1$$

Remark:

- oracle prior sub-family $\{P_{\vartheta_{m_\varepsilon^\circ}}\}_{m_\varepsilon^\circ}$ attains oracle rate $\mathcal{R}_\varepsilon^\circ$ as exact posterior concentration rate

Remark:

- **oracle** prior sub-family $\{P_{\vartheta_{m_\epsilon^\circ}}\}_{m_\epsilon^\circ}$ attains **oracle** rate $\mathcal{R}_\epsilon^\circ$ as **exact** posterior concentration rate
- **minimax** prior sub-family $\{P_{\vartheta_{m_\epsilon^\star}}\}_{m_\epsilon^\star}$ attains **minimax** rate $\mathcal{R}_\epsilon^\star$ as **exact** posterior concentration rate

Remark:

- **oracle** prior sub-family $\{P_{\vartheta^{m_\epsilon^\circ}}\}_{m_\epsilon^\circ}$ attains **oracle** rate $\mathcal{R}_\epsilon^\circ$ as **exact** posterior concentration rate
- **minimax** prior sub-family $\{P_{\vartheta^{m_\epsilon^\star}}\}_{m_\epsilon^\star}$ attains **minimax** rate $\mathcal{R}_\epsilon^\star$ as **exact** posterior concentration rate
- choice of m_ϵ° and $m_\epsilon^\star$ depends, respectively, on θ° and Θ_a

Remark:

- **oracle** prior sub-family $\{P_{\vartheta^{m_\epsilon^\circ}}\}_{m_\epsilon^\circ}$ attains **oracle** rate $\mathcal{R}_\epsilon^\circ$ as **exact** posterior concentration rate
- **minimax** prior sub-family $\{P_{\vartheta^{m_\epsilon^\star}}\}_{m_\epsilon^\star}$ attains **minimax** rate $\mathcal{R}_\epsilon^\star$ as **exact** posterior concentration rate
- choice of m_ϵ° and $m_\epsilon^\star$ depends, respectively, on θ° and Θ_α

Objective:

- put prior on hyperparameter m as outcome of a r.v. M

Outline

- Introduction
- Frequentist perspective reviewed
- Posterior concentration
- Adaptive posterior concentration
- Adaptive Bayes estimator

Hierarchical Prior

Suppose $\Lambda_{(k)} \leq C \min_{j>k} \Lambda_j$ and let $G_\varepsilon := \max\{1 \leq m \leq \lfloor \varepsilon^{-1} \rfloor : \varepsilon \Lambda_{(m)} \leq \Lambda_1\}$.

Hierarchical Prior

Suppose $\Lambda_{(k)} \leq C_\lambda \min_{j>k} \Lambda_j$ and let $G_\varepsilon := \max\{1 \leq m \leq \lfloor \varepsilon^{-1} \rfloor : \varepsilon \Lambda_{(m)} \leq \Lambda_1\}$.

■ random parameter M taking its values in $\{1, \dots, G_\varepsilon\}$

► prior distribution P_M given for $1 \leq m \leq G_\varepsilon$ by

$$p_M(m) = P_M(M = m) \propto \exp(-3C_\lambda m/2) \prod_{j=1}^m (\varsigma_j/\sigma_j)^{1/2}$$

► and $p_M(m) = 0$ otherwise

Hierarchical Prior

Suppose $\Lambda_{(k)} \leq C_\lambda \min_{j>k} \Lambda_j$ and let $G_\varepsilon := \max\{1 \leq m \leq \lfloor \varepsilon^{-1} \rfloor : \varepsilon \Lambda_{(m)} \leq \Lambda_1\}$.

■ random parameter M taking its values in $\{1, \dots, G_\varepsilon\}$

► prior distribution P_M given for $1 \leq m \leq G_\varepsilon$ by

$$p_M(m) = P_M(M = m) \propto \exp(-3C_\lambda m/2) \prod_{j=1}^m (\varsigma_j/\sigma_j)^{1/2}$$

► and $p_M(m) = 0$ otherwise

■ random variables $Y = (Y_j)_{j \geq 1}$ and $\vartheta^M = (\vartheta_j^M)_{j \geq 1}$ determined by

► $Y_j = \vartheta_j^M + \sqrt{\varepsilon} Z_j$ and $\vartheta_j^M = \sqrt{\varsigma_j} V_j \mathbb{1}\{1 \leq j \leq M\}$

where $\{Z_j, V_j\}_{j \geq 1}$ are iid. standard normally distributed and independent of M .

Posterior distribution

■ posterior distribution $P_{M|Y}$ of M given Y

► given for $1 \leq m \leq G_\epsilon$ by

$$p_{M|Y}(m) = P_{M|Y}(M = m) \propto \exp \left(\frac{1}{2} \left\{ \sum_{j=1}^m \sigma_j^{-1} (\theta_j^Y)^2 - 3C_\lambda m \right\} \right)$$

► and $p_{M|Y}(m) = 0$ otherwise

Posterior distribution

- posterior distribution $P_{M|Y}$ of M given Y

► given for $1 \leq m \leq G_\epsilon$ by

$$p_{M|Y}(m) = P_{M|Y}(M = m) \propto \exp \left(\frac{1}{2} \left\{ \sum_{j=1}^m \sigma_j^{-1} (\theta_j^Y)^2 - 3C_\lambda m \right\} \right)$$

► and $p_{M|Y}(m) = 0$ otherwise

- posterior distribution $P_{\vartheta^M|Y}$ of $\vartheta^M = (\vartheta_j^M)_{j \geq 1}$ given Y

► weighted mixture of posterior distributions of $\{\vartheta^m\}_{m=1}^{G_\epsilon}$, i.e.,

$$P_{\vartheta^M|Y} = \sum_{m=1}^{G_\epsilon} p_{M|Y}(m) P_{\vartheta^m|Y}$$

Adaptive oracle posterior concentration

ASSUMPTION A4.

There exists $L_\lambda \geq 1$ such that $1 \leq \Lambda_{(k)}/\bar{\Lambda}_k \leq L_\lambda$ and $\Lambda_{(kl)} \leq \Lambda_{(k)}\Lambda_{(l)}$.

Adaptive oracle posterior concentration

ASSUMPTION A4.

There exists $L_\lambda \geq 1$ such that $1 \leq \Lambda_{(k)}/\bar{\Lambda}_k \leq L_\lambda$ and $\Lambda_{(kl)} \leq \Lambda_{(k)}\Lambda_{(l)}$.

LEMMA. Under A1 and A4 there exist $1 \leq G_\epsilon^- \leq m_\epsilon^\circ \leq G_\epsilon^+ \leq G_\epsilon$ s.t.

$$\mathbb{E}_{\theta^\circ} P_{\mathbf{M}|\mathbf{Y}}(G_\epsilon^- \leq \mathbf{M} \leq G_\epsilon^+) \geq 1 - 4 \exp(-C_\lambda m_\epsilon^\circ/5 + \log G_\epsilon).$$

Adaptive oracle posterior concentration

ASSUMPTION A4.

There exists $L_\lambda \geq 1$ such that $1 \leq \Lambda_{(k)}/\bar{\Lambda}_k \leq L_\lambda$ and $\Lambda_{(kl)} \leq \Lambda_{(k)}\Lambda_{(l)}$.

LEMMA. Under A1 and A4 there exist $1 \leq G_\epsilon^- \leq m_\epsilon^\circ \leq G_\epsilon^+ \leq G_\epsilon$ s.t.

$$\mathbb{E}_{\theta^\circ} P_{\mathbf{M}}|Y(G_\epsilon^- \leq \mathbf{M} \leq G_\epsilon^+) \geq 1 - 4 \exp(-C_\lambda m_\epsilon^\circ/5 + \log G_\epsilon).$$

LEMMA. Under A1 and A4 there exist $1 \leq G_\epsilon^{\star-} \leq m_\epsilon^\star \leq G_\epsilon^{\star+} \leq G_\epsilon$ s.t.

$$\inf_{\theta^\circ \in \Theta_\alpha^r} \mathbb{E}_{\theta^\circ} P_{\mathbf{M}}|Y(G_\epsilon^{\star-} \leq \mathbf{M} \leq G_\epsilon^{\star+}) \geq 1 - 4 \exp(-C_\lambda m_\epsilon^\star/5 + \log G_\epsilon).$$

Adaptive oracle posterior concentration

ASSUMPTION A5. (Regular parameter)

$$\inf_{\varepsilon} [\varepsilon m_{\varepsilon}^{\circ} \bar{\wedge}_{m_{\varepsilon}^{\circ}} \wedge \mathfrak{b}_{m_{\varepsilon}^{\circ}}] / \mathcal{R}_{\varepsilon}^{\circ} > 0$$

Adaptive oracle posterior concentration

ASSUMPTION A5. (Regular parameter)

$$\inf_{\varepsilon} [\varepsilon m_{\varepsilon}^{\circ} \bar{\Lambda}_{m_{\varepsilon}^{\circ}} \wedge \mathfrak{b}_{m_{\varepsilon}^{\circ}}] / \mathcal{R}_{\varepsilon}^{\circ} > 0$$

LEMMA. Under A1, A4 and A5 holds

$$\sum_{m=G_{\varepsilon}^{-}}^{G_{\varepsilon}^{+}} \mathbb{E}_{\theta^{\circ}} P_{\mathfrak{g}^m} | \mathcal{Y} \left(\mathcal{R}_{\varepsilon}^{\circ} \lesssim \| \mathfrak{g}^m - \theta^{\circ} \|^2 \lesssim \mathcal{R}_{\varepsilon}^{\circ} \right) \leq 1 - c e^{-c G_{\varepsilon}^{-}}.$$

Adaptive oracle posterior concentration

ASSUMPTION A5. (Regular parameter)

$$\inf_{\varepsilon} [\varepsilon \bar{m}_{\varepsilon}^{\circ} \wedge \mathfrak{b}_{m_{\varepsilon}^{\circ}}] / \mathcal{R}_{\varepsilon}^{\circ} > 0$$

LEMMA. Under A1, A4 and A5 holds

$$\sum_{m=G_{\varepsilon}^{-}}^{G_{\varepsilon}^{+}} \mathbb{E}_{\theta^{\circ}} P_{\mathfrak{g}^m} |_{\mathcal{Y}} \left(\mathcal{R}_{\varepsilon}^{\circ} \lesssim \|\mathfrak{g}^m - \theta^{\circ}\|^2 \lesssim \mathcal{R}_{\varepsilon}^{\circ} \right) \leq 1 - ce^{-cG_{\varepsilon}^{-}}.$$

THEOREM. Under A1, A4, A5 and $(\log G_{\varepsilon})/m_{\varepsilon}^{\circ} = o(1)$ holds

$$\lim_{\varepsilon \rightarrow \infty} \mathbb{E}_{\theta^{\circ}} P_{\mathfrak{g}^M} |_{\mathcal{Y}} \left(\mathcal{R}_{\varepsilon}^{\circ} \lesssim \|\mathfrak{g}^M - \theta^{\circ}\|^2 \lesssim \mathcal{R}_{\varepsilon}^{\circ} \right) = 1$$

Adaptive minimax posterior concentration

ASSUMPTION A3. (Regular parameter space)

$$\inf_{\varepsilon} [\varepsilon m_{\varepsilon}^* \bar{\Lambda}_{m_{\varepsilon}^*} \wedge \mathfrak{a}_{m_{\varepsilon}^*}] [\varepsilon m_{\varepsilon}^* \bar{\Lambda}_{m_{\varepsilon}^*} \vee \mathfrak{a}_{m_{\varepsilon}^*}]^{-1} > 0$$

Adaptive minimax posterior concentration

ASSUMPTION A3. (Regular parameter space)

$$\inf_{\varepsilon} [\varepsilon \bar{m}_{\varepsilon}^* \bar{\Lambda}_{m_{\varepsilon}^*} \wedge \mathfrak{a}_{m_{\varepsilon}^*}] [\varepsilon \bar{m}_{\varepsilon}^* \bar{\Lambda}_{m_{\varepsilon}^*} \vee \mathfrak{a}_{m_{\varepsilon}^*}]^{-1} > 0$$

LEMMA. Under A1, A3 and A4 holds for all $\theta^{\circ} \in \Theta_{\mathfrak{a}}$

$$\sum_{m=G_{\varepsilon}^{\star-}}^{G_{\varepsilon}^{\star+}} \mathbb{E}_{\theta^{\circ}} R_{\mathfrak{g}^m} | Y \left(\|\mathfrak{g}^m - \theta^{\circ}\|^2 \lesssim \mathcal{R}_{\varepsilon}^{\star} \right) \geq 1 - c e^{-c G_{\varepsilon}^{\star-}}.$$

Adaptive minimax posterior concentration

ASSUMPTION A3. (Regular parameter space)

$$\inf_{\varepsilon} [\varepsilon \mathbf{m}_{\varepsilon}^* \bar{\Lambda}_{\mathbf{m}_{\varepsilon}^*} \wedge \mathbf{a}_{\mathbf{m}_{\varepsilon}^*}] [\varepsilon \mathbf{m}_{\varepsilon}^* \bar{\Lambda}_{\mathbf{m}_{\varepsilon}^*} \vee \mathbf{a}_{\mathbf{m}_{\varepsilon}^*}]^{-1} > 0$$

LEMMA. Under A1, A3 and A4 holds for all $\theta^{\circ} \in \Theta_{\mathbf{a}}$

$$\sum_{m=G_{\varepsilon}^{\star-}}^{G_{\varepsilon}^{\star+}} \mathbb{E}_{\theta^{\circ}} P_{\mathbf{y}^m} \left(\|\mathbf{y}^m - \theta^{\circ}\|^2 \lesssim \mathcal{R}_{\varepsilon}^{\star} \right) \geq 1 - c e^{-c G_{\varepsilon}^{\star-}}.$$

THEOREM. Under A1, A3, A4 and $(\log G_{\varepsilon})/\mathbf{m}_{\varepsilon}^* = o(1)$ holds

$$\blacktriangleright \lim_{\varepsilon \rightarrow \infty} \mathbb{E}_{\theta^{\circ}} P_{\mathbf{y}^M} \left(\|\mathbf{y}^M - \theta^{\circ}\|^2 \lesssim \mathcal{R}_{\varepsilon}^{\star} \right) = 1 \text{ for all } \theta^{\circ} \in \Theta_{\mathbf{a}};$$

Adaptive minimax posterior concentration

ASSUMPTION A3. (Regular parameter space)

$$\inf_{\varepsilon} [\varepsilon \mathbf{m}_{\varepsilon}^* \bar{\Lambda}_{\mathbf{m}_{\varepsilon}^*} \wedge \mathbf{a}_{\mathbf{m}_{\varepsilon}^*}] [\varepsilon \mathbf{m}_{\varepsilon}^* \bar{\Lambda}_{\mathbf{m}_{\varepsilon}^*} \vee \mathbf{a}_{\mathbf{m}_{\varepsilon}^*}]^{-1} > 0$$

LEMMA. Under A1, A3 and A4 holds for all $\theta^{\circ} \in \Theta_{\mathbf{a}}$

$$\sum_{m=G_{\varepsilon}^{*-}}^{G_{\varepsilon}^{*+}} \mathbb{E}_{\theta^{\circ}} P_{\mathbf{y}^m} \left(\|\mathbf{y}^m - \theta^{\circ}\|^2 \lesssim \mathcal{R}_{\varepsilon}^* \right) \geq 1 - c e^{-c G_{\varepsilon}^{*-}}.$$

THEOREM. Under A1, A3, A4 and $(\log G_{\varepsilon})/\mathbf{m}_{\varepsilon}^* = o(1)$ holds

- ▶ $\lim_{\varepsilon \rightarrow \infty} \mathbb{E}_{\theta^{\circ}} P_{\mathbf{y}^M} \left(\|\mathbf{y}^M - \theta^{\circ}\|^2 \lesssim \mathcal{R}_{\varepsilon}^* \right) = 1$ for all $\theta^{\circ} \in \Theta_{\mathbf{a}}$;
- ▶ $\lim_{\varepsilon \rightarrow \infty} \inf_{\theta^{\circ} \in \Theta_{\mathbf{a}}} \mathbb{E}_{\theta^{\circ}} P_{\mathbf{y}^M} \left(\|\mathbf{y}^M - \theta^{\circ}\|^2 \lesssim K_{\varepsilon} \mathcal{R}_{\varepsilon}^* \right) = 1$ for any $K_{\varepsilon} \rightarrow \infty$;

Outline

- Introduction
- Frequentist perspective reviewed
- Posterior concentration
- Adaptive posterior concentration
- Adaptive Bayes estimator

Adaptive data-driven Bayes estimator

Common Bayes estimator $\hat{\theta} = (\hat{\theta}_j)_{j \geq 1} = \mathbb{E}[\boldsymbol{\vartheta}^M | Y]$ given by

$$\hat{\theta}_j = \theta_j^Y P_{M|Y}(j \leq M \leq G_\varepsilon) \mathbb{1}_{\{1 \leq j \leq G_\varepsilon\}}$$

THEOREM. Under A1, A4, A5 and $\log(G_\varepsilon/\mathcal{R}_\varepsilon^\circ)/m_\varepsilon^\circ = o(1)$ holds

$$\mathbb{E}_{\theta^\circ} \|\hat{\theta} - \theta^\circ\|^2 \lesssim \mathcal{R}_\varepsilon^\circ(\theta^\circ) = \min_{m \geq 1} \mathcal{R}_\varepsilon^m(\theta^\circ).$$

Adaptive data-driven Bayes estimator

Common Bayes estimator $\hat{\theta} = (\hat{\theta}_j)_{j \geq 1} = \mathbb{E}[\boldsymbol{\vartheta}^M | Y]$ given by

$$\hat{\theta}_j = \theta_j^Y P_{M|Y}(j \leq M \leq G_\varepsilon) \mathbb{1}_{\{1 \leq j \leq G_\varepsilon\}}$$

THEOREM. Under A1, A4, A5 and $\log(G_\varepsilon/\mathcal{R}_\varepsilon^\circ)/m_\varepsilon^\circ = o(1)$ holds

$$\mathbb{E}_{\theta^\circ} \|\hat{\theta} - \theta^\circ\|^2 \lesssim \mathcal{R}_\varepsilon^\circ(\theta^\circ) = \min_{m \geq 1} \mathcal{R}_\varepsilon^m(\theta^\circ).$$

THEOREM. Under A1, A3, A4 and $\log(G_\varepsilon/\mathcal{R}_\varepsilon^\star)/m_\varepsilon^\star = o(1)$ holds

$$\sup_{\theta^\circ \in \Theta_a} \mathbb{E}_{\theta^\circ} \|\hat{\theta} - \theta^\circ\|^2 \lesssim \mathcal{R}_\varepsilon^\star = \min_{m \geq 1} \sup_{\theta^\circ \in \Theta_a} \mathcal{R}_\varepsilon^m(\theta^\circ)$$

Diffuse hierarchical prior

The posterior weights under a diffuse prior for ϑ satisfy

$$p_{\mathbf{M}|\mathbf{Y}}(m) = \frac{\exp\left(\frac{1}{2}\left\{\sum_{j=1}^m (\varepsilon \Lambda_j)^{-1} (Y_j / \lambda_j)^2 - 3C_{\lambda} m\right\}\right)}{\sum_{k=1}^{G_{\varepsilon}} \exp\left(\frac{1}{2}\left\{\sum_{j=1}^k (\varepsilon \Lambda_j)^{-1} (Y_j / \lambda_j)^2 - 3C_{\lambda} k\right\}\right)}$$

Diffuse hierarchical prior

The posterior weights under a diffuse prior for ϑ satisfy

$$p_{\mathbf{M}|\mathbf{Y}}(\mathbf{m}) = \frac{\exp\left(-\frac{1}{2}\left\{-\|(Y/\lambda)^{\mathbf{m}}\|_{\varepsilon\Lambda}^2 + 3C_{\lambda}\mathbf{m}\right\}\right)}{\sum_{k=1}^{G_{\varepsilon}} \exp\left(-\frac{1}{2}\left\{-\|(Y/\lambda)^k\|_{\varepsilon\Lambda}^2 + 3C_{\lambda}k\right\}\right)}$$

where $\|\theta\|_{\varepsilon\Lambda}^2 := \sum_{j \geq 1} (\varepsilon\Lambda_j)^{-1} (\theta_j)^2$ and $(Y/\lambda)^k := (Y_j/\lambda_j \mathbb{1}_{\{1 \leq j \leq k\}})_{j \in \mathbb{N}}$.

Diffuse hierarchical prior

The posterior weights under a diffuse prior for ϑ satisfy

$$p_{\mathbf{M}|Y}(m) = \frac{\exp\left(-\frac{1}{2}\left\{-\|(Y/\lambda)^m\|_{\varepsilon\Lambda}^2 + 3C_{\lambda}m\right\}\right)}{\sum_{k=1}^{G_{\varepsilon}} \exp\left(-\frac{1}{2}\left\{-\|(Y/\lambda)^k\|_{\varepsilon\Lambda}^2 + 3C_{\lambda}k\right\}\right)}$$

where $\|\theta\|_{\varepsilon\Lambda}^2 := \sum_{j \geq 1} (\varepsilon\lambda_j)^{-1}(\theta_j)^2$ and $(Y/\lambda)^k := (Y_j/\lambda_j \mathbb{1}_{\{1 \leq j \leq k\}})_{j \in \mathbb{N}}$.

The adaptive Bayes estimator $\hat{\theta} = (\hat{\theta}_j)_{j \geq 1} = \mathbb{E}[\vartheta^{\mathbf{M}} | Y]$ is given by

$$\hat{\theta}_j = \left(1 - \sum_{k=1}^{j-1} p_{\mathbf{M}|Y}(m)\right) \times \frac{Y_j}{\lambda_j} \times \mathbb{1}_{\{1 \leq j \leq G_{\varepsilon}\}}.$$

Data-driven shrunked projection estimator

Consider weights

$$w_m^Y := \frac{\exp\left(-\frac{1}{2\varepsilon}\left\{-\|(Y/\lambda)^m\|^2 + 3C_{\lambda\varepsilon}m\bar{\Lambda}_m\right\}\right)}{\sum_{k=1}^{G_\varepsilon} \exp\left(-\frac{1}{2\varepsilon}\left\{-\|(Y/\lambda)^k\|^2 + 3C_{\lambda\varepsilon}k\bar{\Lambda}_k\right\}\right)}$$

Data-driven shrinked projection estimator

Consider weights

$$w_m^Y := \frac{\exp\left(-\frac{1}{2\varepsilon}\left\{-\|(Y/\lambda)^m\|^2 + 3C_{\lambda\varepsilon}m\bar{\Lambda}_m\right\}\right)}{\sum_{k=1}^{G_\varepsilon} \exp\left(-\frac{1}{2\varepsilon}\left\{-\|(Y/\lambda)^k\|^2 + 3C_{\lambda\varepsilon}k\bar{\Lambda}_k\right\}\right)}$$

and the data-driven shrinked projection estimator $\hat{\theta} = (\hat{\theta}_j)_{j \geq 1}$ given by

$$\hat{\theta}_j = \left(1 - \sum_{k=1}^{j-1} w_k^Y\right) \times \frac{Y_j}{\lambda_j} \times \mathbb{1}_{\{1 \leq j \leq G_\varepsilon\}}.$$

Data-driven shrunked projection estimator

Consider weights

$$w_m^Y := \frac{\exp\left(-\frac{1}{2\varepsilon}\left\{-\|(Y/\lambda)^m\|^2 + 3C_{\lambda\varepsilon}m\bar{\Lambda}_m\right\}\right)}{\sum_{k=1}^{G_\varepsilon} \exp\left(-\frac{1}{2\varepsilon}\left\{-\|(Y/\lambda)^k\|^2 + 3C_{\lambda\varepsilon}k\bar{\Lambda}_k\right\}\right)}$$

and the data-driven shrunked projection estimator $\hat{\theta} = (\hat{\theta}_j)_{j \geq 1}$ given by

$$\hat{\theta}_j = \left(1 - \sum_{k=1}^{j-1} w_k^Y\right) \times \frac{Y_j}{\lambda_j} \times \mathbb{1}_{\{1 \leq j \leq G_\varepsilon\}}.$$

Open question: oracle/minimax optimality up to a constant?

Conclusion

■ Summary

- **oracle** and **minimax** optimal concentration rates in an indirect sequence space model

Conclusion

■ Summary

- **oracle** and **minimax** optimal concentration rates in an indirect sequence space model
- hierarchical prior leading to **oracle** and **minimax** optimal concentration rates and a fully-data driven Bayes estimate that is **oracle** and **minimax** optimal in an indirect sequence space model

Conclusion

■ Summary

- **oracle** and **minimax** optimal concentration rates in an indirect sequence space model
- hierarchical prior leading to **oracle** and **minimax** optimal concentration rates and a fully-data driven Bayes estimate that is **oracle** and **minimax** optimal in an indirect sequence space model

■ Perspectives

- additional noise: $X_j = \lambda_j + \text{noise}$

Conclusion

■ Summary

- **oracle** and **minimax** optimal concentration rates in an indirect sequence space model
- hierarchical prior leading to **oracle** and **minimax** optimal concentration rates and a fully-data driven Bayes estimate that is **oracle** and **minimax** optimal in an indirect sequence space model

■ Perspectives

- additional noise: $X_j = \lambda_j + \text{noise}$
- non-diagonal inverse problems with noise in the operator

References



J.J., A. Simoni and R. Schenk (2015) *Adaptive Bayesian estimation in indirect Gaussian sequence space models*. Discussion paper (arXiv:1502.00184)

Thank you for your attention.

ADAPTIVE BAYESIAN ESTIMATION IN INDIRECT GAUSSIAN SEQUENCE SPACE MODELS

Jan JOHANNES

Ruprecht-Karls-Universität Heidelberg

"Mathematical statistics and inverse problems"

February 2016, CIRM, Marseille

joint work with Anna Simoni and Rudolf Schenk

