



Weierstraß-Institut für
Angewandte Analysis und Stochastik

Premolab day Dec 2012

Vladimir Spokoiny

1 Linear parametric modeling

- quasi MLE
- Profile semiparametric estimation
- Penalized MLE

2 Nonparametric modeling

3 General MLE: modern view

- General parametric setup
- A large deviation bound
- Local bracketing

4 Further problems

- Penalized model selection
- Model selection as a parametric problem
- Bayes estimation
- Semiparametric estimation

5 Applied problems

1 Linear parametric modeling

- quasi MLE
- Profile semiparametric estimation
- Penalized MLE

2 Nonparametric modeling

3 General MLE: modern view

- General parametric setup
- A large deviation bound
- Local bracketing

4 Further problems

- Penalized model selection
- Model selection as a parametric problem
- Bayes estimation
- Semiparametric estimation

5 Applied problems

Data Y with DGP $Y \sim P$.

PA: $P \in (P_\theta)$ or $P = P_{\theta^*}$ for $\theta^* \in \Theta$.

Maximum likelihood approach:

$$\tilde{\theta} \stackrel{\text{def}}{=} \operatorname{argmax}_{\theta \in \Theta} L(\theta), \quad L(\theta) = L(Y, \theta) = \log \frac{dP_\theta}{d\mu}(Y)$$

If PA is probably wrong (PW), then $\tilde{\theta}$ is still meaningful (quasi MLE), the target is defined by

$$\theta^* \stackrel{\text{def}}{=} \operatorname{argmax}_{\theta \in \Theta} \mathbb{E} L(\theta).$$

Aim: inference on $\tilde{\theta}$.

A linear model:

$$Y = \Psi^\top \theta + \varepsilon$$

with $Y, \varepsilon \in \mathbb{R}^n$, $\theta \in \mathbb{R}^p$, and Ψ - $p \times n$ design (feature) matrix.

$\varepsilon \sim \mathcal{N}(0, \Sigma^{-1})$ leads to a quasi MLE

$$\tilde{\theta} = \underset{\theta}{\operatorname{argmax}} L(\theta) = \underset{\theta}{\operatorname{argmin}} (Y - \Psi^\top \theta)^\top \Sigma (Y - \Psi^\top \theta)$$

Gauss:

$$\tilde{\theta} = D^{-2} \Psi \Sigma Y, \quad D^2 \stackrel{\text{def}}{=} \Psi \Sigma \Psi^\top.$$

Consider PA $Y \sim \mathcal{N}(\Psi^\top \theta^*, \Sigma^{-1})$ yielding $L(\theta) = -(Y - \Psi^\top \theta)^\top \Sigma (Y - \Psi^\top \theta)/2 + R$

$$\tilde{\theta} = \operatorname{argmax} L(\theta) = D^{-2} \Psi \Sigma Y, \quad D^2 \stackrel{\text{def}}{=} \Psi \Sigma \Psi^\top.$$

Theorem (Gauss-Markov)

Assume a homogeneous noise $\mathbb{E}\varepsilon \equiv 0$, $\operatorname{Var}(\varepsilon) = \Sigma$. Then it holds for

$$\xi \stackrel{\text{def}}{=} D(\tilde{\theta} - \theta^*)$$

1. $\mathbb{E}\xi \equiv 0$ and $\operatorname{Var}(\xi) = \mathbf{I}_p$;
2. $\tilde{\theta}$ is linearly efficient;
3. $L(\tilde{\theta}) - L(\theta) \equiv \|D(\tilde{\theta} - \theta)\|^2/2$ for any θ ;
4. $L(\tilde{\theta}) - L(\theta^*) \equiv \|\xi\|^2/2$.

If PA is correct, i.e. $Y \sim \mathcal{N}(\theta^*, \mathbf{I}_p)$, then

1. ξ is standard normal;
2. $2L(\tilde{\theta}, \theta^*) \sim \chi_p^2$.

The first results 1-3 are entirely based on quadraticity of $L(\theta)$.

The results for the LR $L(\tilde{\theta}, \theta) = L(\tilde{\theta}) - L(\theta)$

$$2L(\tilde{\theta}, \theta) = \|D(\tilde{\theta} - \theta)\|^2, \quad 2L(\tilde{\theta}, \theta^*) \sim \chi_p^2$$

can be used to build likelihood-based **confidence ellipsoids** for the parameter θ^* .
Given $\mathfrak{z} > 0$, define

$$\mathcal{E}(\mathfrak{z}) = \{\theta : L(\tilde{\theta}, \theta) \leq \mathfrak{z}\} = \{\theta : \|D(\tilde{\theta} - \theta)\|^2 \leq 2\mathfrak{z}\}. \quad (1)$$

Theorem

Assume $Y = \Psi^\top \theta^ + \varepsilon$ with $\varepsilon \sim \mathcal{N}(0, \Sigma)$ and consider the MLE $\tilde{\theta}$. Define $\mathfrak{z}_\alpha$ by $P(\chi_p^2 > 2\mathfrak{z}_\alpha) = \alpha$. Then $\mathcal{E}(\mathfrak{z}_\alpha)$ from (1) is an α -confidence set for θ^* .*

- Large p , in particular, $p \gg n$
- Poor identification with $D^2 = \Psi \Sigma \Psi^\top$ badly conditioned
- Model misspecification
- Semiparametric (functional) estimation

v -setup:

$$Y = \Upsilon^\top v^* + \varepsilon = \Psi^\top \theta^* + \Phi^\top \eta^* + \varepsilon, \quad \mathbb{E}\varepsilon = 0, \text{Var}(\varepsilon) = \sigma^2 I_n.$$

Target: $\theta^* = P v^*$.

Profile qMLE 1: $\tilde{\theta} = P \tilde{v} = P(\Upsilon \Upsilon^\top)^{-1} \Upsilon Y = S Y, \quad S = P(\Upsilon \Upsilon^\top)^{-1} \Upsilon.$

Profile qMLE 2: $\tilde{\theta} \stackrel{\text{def}}{=} \underset{\theta}{\operatorname{argmax}} \check{L}(\theta), \quad \check{L}(\theta) \stackrel{\text{def}}{=} \sup_{v: P v = \theta} L(v).$

Theorem

$$\begin{aligned} \mathbb{E} \tilde{\theta} &= \theta^*, \\ \text{Var}(\tilde{\theta}) &= S \text{Var}(\varepsilon) S^\top = \sigma^2 S S^\top = \sigma^2 P(\Upsilon \Upsilon^\top)^{-1} P^\top. \end{aligned}$$

Model:

$$Y = \Psi^\top \theta^* + \Phi^\top \eta^* + \varepsilon \quad \mathbb{E}\varepsilon = 0, \text{Var}(\varepsilon) = \sigma^2 I_n.$$

Theorem

The profile MLE $\tilde{\theta}$ reads as

$$\tilde{\theta} = (\check{\Psi}\check{\Psi}^\top)^{-1}\check{\Psi}Y,$$

$$\check{\Psi} = \Psi - \Psi\Pi_\eta = \Psi - \Psi\Phi^\top(\Phi\Phi^\top)^{-1}\Phi.$$

Model:

$$Y = \Upsilon^\top v^* + \varepsilon = \Psi^\top \theta^* + \Phi^\top \eta^* + \varepsilon \quad \mathbb{E}\varepsilon = 0, \text{Var}(\varepsilon) = \sigma^2 I_n.$$

Theorem (Gauss-Markov)

1. $\tilde{\theta} = SY$ with $S = P(\Upsilon\Upsilon^\top)^{-1}\Upsilon$ fulfills

$$\mathbb{E}\tilde{\theta} = \theta^* = Pv^*,$$

$$\mathbb{E}\|\tilde{\theta} - \theta^*\|^2 = \text{Var}(\tilde{\theta}) = \sigma^2 P(\Upsilon\Upsilon^\top)^{-1}P^\top = \sigma^2(\check{\Psi}\check{\Psi}^\top)^{-1},$$

$$\check{\Psi} = \Psi - \Psi\Pi_\eta$$

$$\Pi_\eta = \Phi^\top(\Phi\Phi^\top)^{-1}\Phi.$$

2. This risk is minimal in the class of all unbiased linear estimates of θ^* .

Define a CS for θ^* as the level set of $\check{L}(\theta) = \sup_{v: Pv=\theta} L(v)$:

$$\mathcal{E}(\mathfrak{z}) \stackrel{\text{def}}{=} \left\{ \theta : \check{L}(\tilde{\theta}) - \check{L}(\theta) \leq \mathfrak{z} \right\}.$$

This definition can be rewritten as

$$\mathcal{E}(\mathfrak{z}) \stackrel{\text{def}}{=} \left\{ \theta : \sup_v L(v) - \sup_{v: Pv=\theta} L(v) \leq \mathfrak{z} \right\}.$$

In (θ, η) -setup the definition is even more transparent:

$$\mathcal{E}(\mathfrak{z}) \stackrel{\text{def}}{=} \left\{ \theta : \sup_{\theta, \eta} L(\theta, \eta) - \sup_{\eta} L(\theta, \eta) \leq \mathfrak{z} \right\}.$$

Model:

$$Y = \Psi^\top \theta^* + \Phi^\top \eta^* + \varepsilon \quad \mathbb{E}\varepsilon = 0, \text{Var}(\varepsilon) = \sigma^2 I_n.$$

Define

$$D^2 = \sigma^{-2} \check{\Psi} \check{\Psi}^\top, \quad \check{\Psi} = \Psi - \Psi \Pi_\eta.$$

Theorem

Let the matrix D be non-degenerated. It holds

$$\begin{aligned} 2\{\check{L}(\tilde{\theta}) - \check{L}(\theta^*)\} &= \|D(\tilde{\theta} - \theta^*)\|^2 = \|\xi\|^2, \\ \xi &= D(\tilde{\theta} - \theta^*), \quad \mathbb{E}\xi = 0, \text{Var}(\xi) = I_p. \end{aligned}$$

If $\varepsilon \sim \mathcal{N}(0, \sigma^2 I_n)$, then ξ is standard normal in $\mathbb{R}^p$ and

$$2\{\check{L}(\tilde{\theta}) - \check{L}(\theta^*)\} \sim \chi_p^2.$$

Let $\Sigma = I_n$ and $\tilde{\theta} = (\Psi\Psi^\top)^{-1}\Psi Y$.

Let R be a positive symmetric $p \times p$ matrix. Then the sum $\Psi\Psi^\top + R$ is positive symmetric as well and can be inverted whatever Ψ is.

This suggests to replace $(\Psi\Psi^\top)^{-1}$ by $(\Psi\Psi^\top + R)^{-1}$ leading to the regularized least squares estimate $\tilde{\theta}_R$ of the parameter vector θ :

$$\tilde{\theta}_R \stackrel{\text{def}}{=} (\Psi\Psi^\top + R)^{-1}\Psi Y,$$

Such a method is also called *ridge regression*.

Tikhonov regularization: $R = \alpha I_p$ where $\alpha > 0$ and I_p stands for the unit matrix:

$$\tilde{\theta}_\alpha \stackrel{\text{def}}{=} (\Psi\Psi^\top + \alpha I_p)^{-1}\Psi Y.$$

MLE: with $L(\theta) = \sigma^{-2} \|Y - \Psi^\top \theta\|^2 / 2$

$$\tilde{\theta} = \operatorname{argmax}_{\theta} L(\theta) = \operatorname{arginf}_{\theta} \|Y - \Psi^\top \theta\|^2$$

Roughness or complexity penalization: add a quadratic penalty $\|G\theta\|^2/2$ for a given symmetric matrix G :

$$\begin{aligned} L_G(\theta) &\stackrel{\text{def}}{=} L(\theta) - \|G\theta\|^2/2 \\ &= -(2\sigma^2)^{-1} \|Y - \Psi^\top \theta\|^2 - \|G\theta\|^2/2 - R. \end{aligned} \tag{2}$$

The penalized ML problem reads as

$$\tilde{\theta}_G = \operatorname{argmax}_{\theta} L_G(\theta) = \operatorname{argmin}_{\theta} \{ (2\sigma^2)^{-1} \|Y - \Psi^\top \theta\|^2 + \|G\theta\|^2/2 \}.$$

Straightforward calculus yield

$$\tilde{\theta}_G \stackrel{\text{def}}{=} (\Psi\Psi^\top + \sigma^2 G^2)^{-1} \Psi Y = S_G Y. \quad (3)$$

$\tilde{\theta}_G$ estimates the value θ_G defined by

$$\begin{aligned} \theta_G &= \operatorname{argmax}_{\theta} \mathbb{E} L_G(\theta) \\ &= \operatorname{arginf}_{\theta} \mathbb{E} \{ \|Y - \Psi^\top \theta\|^2 + \sigma^2 \|G\theta\|^2 \} \\ &= (\Psi\Psi^\top + \sigma^2 G^2)^{-1} \Psi\Psi^\top \theta^* \end{aligned}$$

and $\theta_G \neq \theta^*$ unless $G = 0$. In other words, the penalized MLE $\tilde{\theta}_G$ is biased.

Theorem

Let $\tilde{\theta}_G$ be a penalized MLE from (3). Under noise homogeneity $\text{Var}(\varepsilon) = \sigma^2 \mathbf{I}_n$, it holds $\mathbb{E}\tilde{\theta}_G = \theta_G$ and

$$\text{Var}(\tilde{\theta}_G) = \sigma^2 S_G S_G^\top = \sigma^2 (\Psi \Psi^\top + \sigma^2 G^2)^{-1} \Psi \Psi^\top (\Psi \Psi^\top + \sigma^2 G^2)^{-1}.$$

The bias $\|\theta_G - \theta^*\|$ *monotonously increases* in G^2 while the variance *monotonously decreases* with the penalization G .

Finally

$$\mathbb{E}\|\tilde{\theta}_G - \theta^*\|^2 = \|\theta_G - \theta^*\|^2 + \sigma^2 \text{tr}(S_G S_G^\top).$$

Example (Tikhonov regularization): $\mathbb{E}\tilde{\theta}_\alpha = \theta_\alpha$ for $\theta_\alpha = (\Psi \Psi^\top + \alpha \mathbf{I}_p)^{-1} \Psi \Psi^\top \theta^*$, the bias $\|\theta_\alpha - \theta^*\|$ grows with the regularization parameter α .

Theorem

Let $L_G(\theta)$ be the penalized log-likelihood from (2). Then

$$\begin{aligned} 2L_G(\tilde{\theta}_G, \theta_G) &= (\tilde{\theta}_G - \theta_G)^\top (\sigma^{-2}\Psi\Psi^\top + G^2)(\tilde{\theta}_G - \theta_G) \\ &= \sigma^{-2}\varepsilon^\top \Pi_G \varepsilon \end{aligned}$$

with $\Pi_G = \Psi^\top (\Psi\Psi^\top + \sigma^2 G^2)^{-1} \Psi$.

Define the confidence set $\mathcal{E}_G(\mathfrak{z})$ for θ_G as

$$\mathcal{E}_G(\mathfrak{z}) \stackrel{\text{def}}{=} \{\theta : L_G(\tilde{\theta}_G, \theta) \leq \mathfrak{z}\}.$$

The definition implies the following result for the coverage probability:

$$\mathbb{P}(\mathcal{E}_G(\mathfrak{z}) \not\ni \theta_G) \leq \mathbb{P}(L_G(\tilde{\theta}_G, \theta_G) > \mathfrak{z})$$

Theorem

Let $L_G(\theta)$ be the penalized log-likelihood from (2) and let $\varepsilon \sim \mathcal{N}(0, \sigma^2 \mathbf{I}_n)$. Then it holds with $\mathfrak{p}_G = \text{tr}(\Pi_G)$ and $v_G^2 = 2 \text{tr}(\Pi_G^2)$ that

$$\mathbb{P}(2L_G(\tilde{\theta}_G, \theta_G) > \mathfrak{p}_G + (2v_G \mathbf{x}^{1/2}) \vee (6\mathbf{x})) \leq \exp(-\mathbf{x}).$$

The key issue is an **identifiability**. Particularly challenging is any inference for

- inverse problems;
- poor design;
- dimension p too high;

More specific questions:

- Choice of a penalty class (given by G). Examples:
 - Projection (spectral cut-off)
 - Tikhonov
 - Roughness
 - diagonal (spectral weighting)
- Choice of a tuning parameters or aggregation
- Critical dimension (cardinality of the set of models)

Penalized MLE:

$$L(\theta, \lambda) \stackrel{\text{def}}{=} L(\theta) + \lambda t(\theta) \quad \tilde{\theta}(\lambda) \stackrel{\text{def}}{=} \underset{\theta}{\operatorname{argmax}} L(\theta, \lambda).$$

Complexity:

$$t(\theta) = \|\theta\|_0 = \#\{\theta_j \neq 0\}.$$

Sparse:

$$t(\theta) = \|\theta\|_1 = \sum_j |\theta_j|$$

Challenges:

- choice of penalty
- choice of λ
- constrained models under structural assumptions

1 Linear parametric modeling

- quasi MLE
- Profile semiparametric estimation
- Penalized MLE

2 Nonparametric modeling

3 General MLE: modern view

- General parametric setup
- A large deviation bound
- Local bracketing

4 Further problems

- Penalized model selection
- Model selection as a parametric problem
- Bayes estimation
- Semiparametric estimation

5 Applied problems

- Global parametric MLE $\tilde{\theta} = \operatorname{argmax} L(\theta)$, e.g. GLM, median and quantile regression, etc.
- Sieve parametric MLE
- Roughness penalty, regularization
- Complexity and sparse penalty

Issues:

- Choice of the basic model $L(\theta)$ and modeling bias;
- Choice of the sequence of models (sieves);
- Data-driven choice of a tuning parameter;
- Choice of the class of penalties $t(\theta)$;

Even a complex model structure admits a simple description in a proper vicinity of each point.

Local approach: fix a vicinity and approximate the model function by a simple parametric model in this vicinity.

Issues:

- Choice of a family for local approximation (e.g. polynomials of a fixed degree)
- Choice of the locality parameter (bandwidth):
 - global vs pointwise,
 - isotropic vs anisotropic
- Uniformity and confidence bands

Semiparametric structural assumptions:

- single index $f(x) = g(x^\top \theta)$
- multiple index $f(x) = g(x^\top \theta_1, \dots, x^\top \theta_m)$
- additive $f(x) = f_1(x_1) + \dots + f_d(x_d)$
- projection pursuit $f(x) = f_1(x^\top \theta_1) + \dots + f_d(x^\top \theta_d)$.

Local structures, manifold learning: $\theta = \theta(x)$.

1 Linear parametric modeling

- quasi MLE
- Profile semiparametric estimation
- Penalized MLE

2 Nonparametric modeling

3 General MLE: modern view

- General parametric setup
- A large deviation bound
- Local bracketing

4 Further problems

- Penalized model selection
- Model selection as a parametric problem
- Bayes estimation
- Semiparametric estimation

5 Applied problems

PA: $\mathbb{P} \in (\mathbb{P}_\theta, \theta \in \Theta \subseteq \mathbb{R}^p)$.

Classical parametric viewpoint:

- PA is exactly fulfilled
- sample size n large

Main results:

- ▶ MLE $\tilde{\theta}$ is **asymptotically efficient**.
- ▶ **Wilks phenomenon**:

$$2L(\tilde{\theta}) - 2L(\theta^*) \xrightarrow{w} \chi_p^2.$$

Classical assumptions:

PA1 $P \in (P_\theta, \theta \in \Theta \subset \mathbb{R}^p)$

PA2 $n \rightarrow \infty$ or n/p large.

Requirements to the “modern” theory:

1. addressing “model misspecification” and “small samples”
2. consistent with the classical asymptotic results including the efficiency bounds.
3. extendable to the nonparametric situation.

Main tools:

- Local approximation
- Concentration of measures, probability bounds

$$\tilde{\theta} = \operatorname{argmax}_{\theta} L(\theta), \quad \theta^* = \operatorname{argmax}_{\theta} \mathbb{E}L(\theta).$$

- Concentration in a local neighborhood $\Theta_0(\mathbf{r})$ of θ^*

$$\mathbb{P}(\tilde{\theta} \notin \Theta_0(\mathbf{r})) \leq C_1 \exp\{-C_2 \mathbf{r}\}$$

- Local linear approximation:

Approximate $L(\theta)$ within $\Theta_0(\mathbf{r})$ by a quadratic (in θ) expansion:

$$\mathbb{E}L(\theta) - \mathbb{E}L(\theta^*) \approx -\frac{1}{2} \|D_0(\theta - \theta^*)\|^2, \quad D_0^2 \stackrel{\text{def}}{=} \nabla^2 \mathbb{E}L(\theta^*),$$

$$\zeta(\theta) - \zeta(\theta^*) \approx (\theta - \theta^*)^\top \nabla \zeta(\theta^*), \quad \zeta(\theta) \stackrel{\text{def}}{=} L(\theta) - \mathbb{E}L(\theta).$$

Yields

$$L(\theta, \theta^*) \approx (\theta - \theta^*)^\top \nabla \zeta(\theta^*) - \frac{1}{2} \|D_0(\theta - \theta^*)\|^2$$

Problem: error of approximation is too big.

Classical conditions:

- Θ is compact;
- identifiability:

$$\mathbb{E}L(\theta^*) - \mathbb{E}L(\theta) > 0.$$

- stochastic continuity of $L(\theta)$ and continuity of $\mathbb{E}L(\theta)$;
- large samples.

“Modern” approach:

- Local smoothness of $L(\theta)$: $\text{Var}\{\nabla L(\theta)\} \leq V^2$
- Qualified identifiability:

$$\mathbb{E}L(\theta^*) - \mathbb{E}L(\theta) \geq \begin{cases} b\|V(\theta - \theta^*)\|^2 \\ C_1 p \log(\|V(\theta - \theta^*)\|^2/p) \end{cases} \quad \text{or,} \quad .$$

Let $\zeta(\theta) = L(\theta) - \mathbb{E}L(\theta)$.

(ED) $\exists$ a positive symmetric matrix V^2 , and constants $g > 0$, $v_0 \geq 1$ s.t.
for all $\theta \in \Theta$

$$\sup_{\gamma \in \mathbb{S}^p} \log \mathbb{E} \exp \left\{ \lambda \frac{\gamma^\top \nabla \zeta(\theta)}{\|V\gamma\|} \right\} \leq v_0^2 \lambda^2 / 2, \quad |\lambda| \leq g.$$

(L) There is a constant $C_1 \geq 1$ s.t.

$$\mathfrak{M}^*(\theta, \theta^*) \geq C_1 p^* \log(\|V(\theta - \theta^*)\|^2 / p^*),$$

where

$$\mathfrak{M}^*(\theta, \theta^*) \stackrel{\text{def}}{=} \sup_{\mu \in \mathbb{M}} \mathfrak{M}(\mu, \theta, \theta^*).$$

Theorem

Suppose Then it holds for $\Theta_0(\mathbf{r}) \stackrel{\text{def}}{=} \{\theta : \|V_0(\theta - \theta^*)\| \leq \mathbf{r}\}$

$$\mathbb{P}(\tilde{\theta} \notin \Theta_0(\mathbf{r})) \leq C_1 \exp\{-C_2 \mathbf{r}\}.$$

- Ensures concentration of $\tilde{\theta}$ in a vicinity of θ^* of radius $\asymp \sqrt{p/n}$.
- The bound non-asymptotic, applies for small samples and nearly sharp. All constants implicit.
- The true measure $\mathbb{P}$ does not belong to $(\mathbb{P}_\theta)$.

Tools:

- Generic chaining in a local ball; Talagrand (2005), Bednorz (2006);
- Slicing arguments

(ED₀) $\exists$ a positive symmetric matrix V_0^2 , and constants $g > 0$, $v_0 \geq 1$ s.t.

$$\sup_{\gamma \in \mathbb{S}^p} \log \mathbb{E} \exp \left\{ \lambda \frac{\gamma^\top \nabla \zeta(\theta^*)}{\|V_0 \gamma\|} \right\} \leq v_0^2 \lambda^2 / 2, \quad |\lambda| \leq g.$$

(ED₁) For some R and each $r \leq R$, there exist a constant $\omega(r) \leq 1/2$ such that it holds for all $\theta \in \Theta_0(r) = \Theta_0(r) \stackrel{\text{def}}{=} \{\theta : \|V_0(\theta - \theta^*)\| \leq r\}$:

$$\sup_{\gamma \in \mathbb{S}^p} \log \mathbb{E} \exp \left\{ \lambda \frac{\gamma^\top \{\nabla \zeta(\theta) - \nabla \zeta(\theta^*)\}}{\omega(r) \|V_0 \gamma\|} \right\} \leq v_0^2 \lambda^2 / 2, \quad |\lambda| \leq g.$$

(L₀) There are a matrix D_0 and for each $r \leq R$ and a constant $\delta(r) \leq 1/2$, such that it holds on the set $\Theta_0(r) = \{\theta : \|V_0(\theta - \theta^*)\| \leq r\}$

$$\left| \frac{-2\mathbb{E}L(\theta, \theta^*)}{\|D_0(\theta - \theta^*)\|^2} - 1 \right| \leq \delta(r).$$

Yields

$$L(\theta, \theta^*) \approx (\theta - \theta^*)^\top \nabla \zeta(\theta^*) - \frac{1}{2} \|D_0(\theta - \theta^*)\|^2$$

Problem: error of approximation is too big.

Approach: majorate from above and from below by log-likelihoods of two different linear models:

With $\varepsilon = (\delta, \rho)$, define

$$\mathbb{L}_\varepsilon(\theta, \theta^*) \stackrel{\text{def}}{=} (\theta - \theta^*)^\top \nabla L(\theta^*) - \|D_\varepsilon(\theta - \theta^*)\|^2/2$$

$$\mathbb{L}_{\underline{\varepsilon}}(\theta, \theta^*) \stackrel{\text{def}}{=} (\theta - \theta^*)^\top \nabla L(\theta^*) - \|D_{\underline{\varepsilon}}(\theta - \theta^*)\|^2/2$$

where with $V_0^2 = \text{Var}(\nabla \zeta(\theta^*))$

$$D_\varepsilon^2 = D_0^2(1 - \delta) - \rho V_0^2,$$

$$D_{\underline{\varepsilon}}^2 = D_0^2(1 + \delta) + \rho V_0^2,$$

The construction depends upon the unknown quantities θ^* , D_0 , V_0 .

$$\mathbb{L}_{\varepsilon}(\theta, \theta^*) = (\theta - \theta^*)^\top \nabla L(\theta^*) - \|D_{\varepsilon}(\theta - \theta^*)\|^2/2$$

$$\mathbb{L}_{\underline{\varepsilon}}(\theta, \theta^*) = (\theta - \theta^*)^\top \nabla L(\theta^*) - \|D_{\underline{\varepsilon}}(\theta - \theta^*)\|^2/2,$$

$$D_{\varepsilon} = D_0^2(1 - \delta) - \rho V_0^2, \quad \varepsilon = (\delta, \rho), \quad \underline{\varepsilon} = -\varepsilon = (-\delta, -\rho).$$

Theorem

Assume (ED_1) and $(\mathcal{L}_0)$. Let for some $\mathbf{r}$, the values $\rho \geq \sqrt{3v_0} \omega(\mathbf{r})$ and $\delta \geq \delta(\mathbf{r})$ ensure $D_{\varepsilon} \geq 0$. Then

$$\mathbb{L}_{\underline{\varepsilon}}(\theta, \theta^*) - \diamond_{\underline{\varepsilon}}(\mathbf{r}) \leq L(\theta, \theta^*) \leq \mathbb{L}_{\varepsilon}(\theta, \theta^*) + \diamond_{\varepsilon}(\mathbf{r}), \quad \theta \in \Theta_0(\mathbf{r}),$$

where the random variables $\diamond_{\varepsilon}(\mathbf{r}), \diamond_{\underline{\varepsilon}}(\mathbf{r})$ fulfill

$$\log \mathbb{E} \exp\{\diamond_{\varepsilon}(\mathbf{r})/\omega(\mathbf{r}) - \mathbf{c}_1 p^*\} \leq 0.$$

Usually $\omega(\mathbf{r}) \asymp n^{-1/2} \mathbf{r}$ and $\delta(\mathbf{r}) \asymp n^{-1/2} \mathbf{r}$.

Tools:

- Generic chaining + slicing
- Sharp deviation bounds for quadratic forms of Gaussian and non-Gaussian r.v.

Issues:

- The bound is non-asymptotic, applies even for small samples
- The set $\Theta_0(r)$ should not be root-n small. It corresponds to a usual local vicinity.
- No normalization/scaling necessary, the bounds are self-normalizing.
- The difference $\mathbb{L}_\varepsilon(\theta, \theta^*) - \mathbb{L}_{\underline{\varepsilon}}(\theta, \theta^*)$ can be large. Matters “upper value” - “lower value”:

$$\sup_{\theta} \mathbb{L}_\varepsilon(\theta, \theta^*) - \sup_{\theta} \mathbb{L}_{\underline{\varepsilon}}(\theta, \theta^*).$$

Consider the confidence set

$$\mathcal{E}(\mathfrak{z}) = \{\theta : L(\tilde{\theta}, \theta) \leq \mathfrak{z}\}$$

Coverage probability:

$$\mathbb{P}\{\mathcal{E}(\mathfrak{z}) \not\ni \theta^*\} = \mathbb{P}\{L(\tilde{\theta}, \theta^*) > \mathfrak{z}\}.$$

Theorem

For any $\mathfrak{z} > 0$, it holds with $\xi_{\varepsilon} = D_{\varepsilon}^{-1} \nabla L(\theta^*)$

$$\mathbb{P}\{\mathcal{E}(\mathfrak{z}) \not\ni \theta^*, \tilde{\theta} \in \Theta_0(\mathbf{r})\} \leq \mathbb{P}\{\|\xi_{\varepsilon}\|^2 \geq 2\mathfrak{z} - 2\Diamond_{\varepsilon}(\mathbf{r})\}.$$

Define

$$\Delta_{\varepsilon}(\mathbf{r}) \stackrel{\text{def}}{=} \diamond_{\varepsilon}(\mathbf{r}) + \diamond_{\underline{\varepsilon}}(\mathbf{r}) + (\|\xi_{\varepsilon}\|^2 - \|\xi_{\underline{\varepsilon}}\|^2)/2.$$

Theorem

For any $z > 0$, it holds

$$\mathbb{P}\{\|D_{\varepsilon}(\tilde{\theta} - \theta^*)\| > z, \tilde{\theta} \in \Theta_0(\mathbf{r})\} \leq \mathbb{P}\{\|\xi_{\varepsilon}\| > z - \sqrt{2\Delta_{\varepsilon}(\mathbf{r})}\}.$$

Theorem (Wilks phenomenon)

The following approximation holds on the random set $C_\varepsilon(\mathbf{r}) = \{\tilde{\boldsymbol{\theta}} \in \boldsymbol{\Theta}_0(\mathbf{r}), \|\xi_\varepsilon\| \leq \mathbf{r}\}$:

$$\|\xi_\varepsilon\|^2/2 - \diamond_{\varepsilon}(\mathbf{r}) \leq L(\tilde{\boldsymbol{\theta}}, \boldsymbol{\theta}^*) \leq \|\xi_\varepsilon\|^2/2 + \diamond_{\varepsilon}(\mathbf{r}).$$

Theorem (Fisher expansion)

On the same random set $C_\varepsilon(\mathbf{r})$

$$\|D_\varepsilon(\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}^*) - \xi_\varepsilon\|^2 \leq 2\Delta_\varepsilon(\mathbf{r}).$$

- ▶ The results are useless if the error Δ_ε is too big.
- ▶ The value Δ_ε depends on the dimension p .
- ▶ Critical dimension p_n ?

Theorem

In the regular i.i.d. case the Wilks result holds if

p_n/n is small.

1 Linear parametric modeling

- quasi MLE
- Profile semiparametric estimation
- Penalized MLE

2 Nonparametric modeling

3 General MLE: modern view

- General parametric setup
- A large deviation bound
- Local bracketing

4 Further problems

- Penalized model selection
- Model selection as a parametric problem
- Bayes estimation
- Semiparametric estimation

5 Applied problems

Modern applications:

- ▶ large parameter space,
- ▶ lack of identifiability.

Approach: penalization:

$$\tilde{\theta} \stackrel{\text{def}}{=} \operatorname{argmax}_{\theta} \{L(\theta) - \mathfrak{t}(\theta)\}$$

- smooth (roughness) penalty $\mathfrak{t}(\theta) = \lambda \|G\theta\|^2/2$
- complexity penalty $\mathfrak{t}(\theta) = \lambda \|\theta\|_0$
- sparcity penalty $\mathfrak{t}(\theta) = \lambda \|\theta\|_1$

New target:

$$\theta_G^* \stackrel{\text{def}}{=} \operatorname{argmax}_{\theta} \{\mathbb{E}L(\theta) - \mathfrak{t}(\theta)\}.$$

The general theory extends easily to the smooth and complexity penalization.

Challenge: a general theory for the sparse case.

$$\tilde{\theta} \stackrel{\text{def}}{=} \operatorname{argmax}_{\theta} \{L(\theta) - t(\theta)\} = \operatorname{argmax}_{\theta} \{L(\theta) - \lambda \|G\theta\|^2/2\}$$

Choice of λ crucial (central problem in nonparametric statistics).

Approach: include λ in the list of parameters.

Naive definition

$$(\tilde{\theta}, \tilde{\lambda}) \stackrel{\text{def}}{=} \operatorname{argmax}_{\theta, \lambda} \{L(\theta) - \lambda \|G\theta\|^2/2\}$$

Does not work (yields **overfitting**). Better

$$(\tilde{\theta}, \tilde{\lambda}) \stackrel{\text{def}}{=} \operatorname{argmax}_{\theta, \lambda} \{L(\theta) - \lambda \|G\theta\|^2/2 - t(\lambda)\}$$

New challenge: choice of $t(\lambda)$.

Examples: AIC, BIC, SCAD, ...

$$Y \mid \theta \sim p(y \mid \theta),$$

$$\vartheta \sim \pi(\cdot).$$

ϑ random from the **prior** π .

Finally the *posteriori (conditional) density* $p(\vartheta \mid y)$ of ϑ given y :

$$p(\theta \mid y) = \frac{p(y, \theta)}{p(y)} = \frac{p(y \mid \theta) \pi(\theta)}{\int_{\Theta} p(y \mid \theta) \pi(\theta) \lambda(d\theta)}.$$

or

$$\vartheta \mid Y \propto \exp\{L(\theta)\} \Pi(d\theta).$$

Similarly to our general qMLE approach

$$\theta^* \stackrel{\text{def}}{=} \operatorname{argmax}_{\theta \in \Theta} \mathbb{E} L(\theta).$$

Aim: concentration of the posterior, near Gaussianity.

- **Non-informative prior:** $\pi(d\theta) = d\theta$.

Posterior $\vartheta \mid Y$ fulfills:

$$\mathbb{E}(\vartheta \mid Y) \approx \tilde{\theta} \quad \text{Var}(\vartheta \mid Y) \approx D_0^{-2}, \quad D_0(\vartheta \mid Y) \approx \xi \sim \mathcal{N}(0, \mathbf{I}_p)$$

- **Gaussian priors:** $d\pi_G(\theta) \propto \exp(-\|G\theta\|^2/2)$.

Posterior $\vartheta_G \mid Y$ fulfills:

$$\mathbb{E}(\vartheta_G \mid Y) \approx \tilde{\theta}_G \quad \text{Var}(\vartheta \mid Y) \approx D_G^{-2}, \quad D_G(\vartheta_G \mid Y) \approx \xi \sim \mathcal{N}(0, \mathbf{I}_p)$$

- **Critical dimension:** the BvM result applies if “ p/n is small”.

Challenge 1: Sparse priors: (with heavy tails). Only few special cases studied.

Challenge 2: Bayesian model selection and Bernstein - von Mises Theorem for hyperpriors (priors on prior parameters).

Data Y with DGP $Y \sim \mathbb{P}$.

SPA: $\mathbb{P} \in (\mathbb{P}_{\theta, \eta})$, probably misspecified.

θ , target, η , nuisance.

Goal: inference on θ .

Examples in mind:

- an inverse problem with error in operator;
observed Y and $\hat{A}$ following $Y = A\theta + \varepsilon$ with θ and A unknown;
- IV problem; see e.g. Hall and Horowitz (2005), Horowitz (2011)
- dimension reduction in single- and multi-index models; e.g. Xia et al (2000), Hristache et al (2001a,b).

SPA : $Y \sim P \in (\mathcal{P}_{\theta, \eta}, \theta \in \Theta, \eta \in H)$

Log-likelihood: $L(\theta, \eta) = \frac{dP_{\theta, \eta}}{d\mu_0}(Y)$

Profile MLE: $\tilde{\theta} = \operatorname{argmax}_{\theta} \max_{\eta} L(\theta, \eta) = \operatorname{argmax}_{\theta} \check{L}(\theta), \quad \check{L}(\theta) = \max_{\eta} L(\theta, \eta).$

Murphy, van der Vaart (2000), Kosorok (2005, 2008): Let $P = P_{\theta^*, \eta^*}$. Then $\tilde{\theta}$ is

- root- n consistent and normal
- semiparametrically efficient
- $2\check{L}(\tilde{\theta}) - 2\check{L}(\theta^*) \xrightarrow{w} \chi_p^2$, where $p = \dim(\Theta)$.

Limitations:

- hard optimization problem, often unfeasible
- SPA is crucial but questionable
- large sample asymptotics

- Alternating or EM-algorithms: iterate two steps until convergence;
- backfitting / projection pursuit;
- MCMC, Gibbs, ...

Problems:

- convergence, asymptotic properties, efficiency only addressed in few cases;
- no general theory available;

Challenges: a general study for an alternating method focusing on:

- finite samples;
- model misspecification;
- unified approach for direct and inverse problems;
- range of applicability, critical dimension of the nuisance;

1 Linear parametric modeling

- quasi MLE
- Profile semiparametric estimation
- Penalized MLE

2 Nonparametric modeling

3 General MLE: modern view

- General parametric setup
- A large deviation bound
- Local bracketing

4 Further problems

- Penalized model selection
- Model selection as a parametric problem
- Bayes estimation
- Semiparametric estimation

5 Applied problems

Observed $Y \in \mathbb{R}^1, X \in \mathbb{R}^d$, d large.

Aim: regression $Y = f(X) + \varepsilon$. **Unfeasible** for $d \geq 3$.

EDR assumptions:

► **Single-index:**

$$f(X) = g(X^\top \theta).$$

► **Multi-index:**

$$f(X) = g(TX), \quad T : \mathbb{R}^d \rightarrow \mathbb{R}^m$$

► **Additive**

$$f(X) = g_1(X_1) + \dots + g_d(X_d)$$

► **Projection-pursuit**

$$f(X) = g_1(X^\top \theta_1) + \dots + g_d(X^\top \theta_d)$$

$X_1, \dots, X_n \in \mathbb{R}^d$, i.i.d. from $p(\cdot)$

Aim: features of $p(\cdot)$ like multimodality or heavy tails.

Non-Gaussian components:

$$p(x) = g(TX)\phi_\Gamma(x)$$

T , a non-Gaussian subspace.

Applications in **clustering** (unsupervised learning) for **multimodality** and in **risk management** in **heavy tails**

Linear inverse problem:

$$Y = A\theta + \varepsilon$$

A unknown, instead a pilot $\hat{A}$ is available.

Aim: estimation of θ , the operator A is the **nuisance** parameter.

Applications in Instrumental Variable regression, functional data analysis, imaging, and many others

CMR modeling:

$$\mathbb{E}[g(Z, \theta) \mid X] = 0 \quad \text{a.s.} \quad \Leftrightarrow \quad \theta = \theta^*.$$

Approach:

$$\tilde{\theta} = \tilde{\theta}_h \stackrel{\text{def}}{=} \underset{\theta \in \Theta}{\operatorname{argmin}} M(\theta)$$

with

$$M(Z, \theta) \stackrel{\text{def}}{=} \sum_{i \neq j} g_i(\theta) g_j(\theta) w_{ij}$$

where w_{ij} are localizing weights with the bandwidth h .

Challenge: Theory for quadratic forms.

- ▶ Adaptive clustering
- ▶ Adaptive classification
- ▶ Random design and time series
- ▶ Structure adaptive manifold smoothing (locally adaptive dimension reduction).
- ▶ Nonparametric change point detection
- ▶ Hidden Markov chain modeling
- ▶ Resampling methods