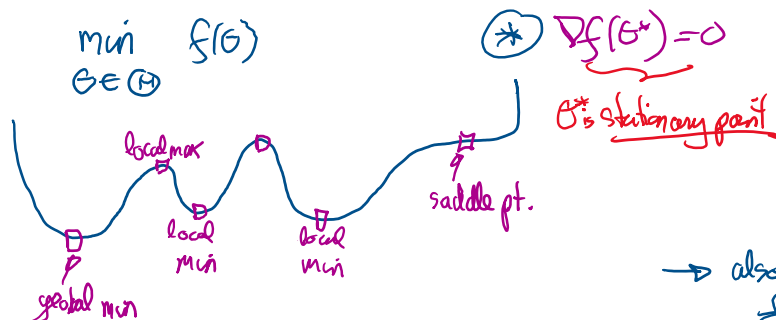


today: • MLE of θ
• Statistical decision theory

optimization comments about MLE



f is differentiable on Θ
is a necessary cond.
for θ^* to be a local min
when θ^* is in interior of Θ

→ also check that $\text{Hessian}(f)(\theta^*) \succ 0$
for a local min

$$H \succ 0 \Leftrightarrow u^T H u > 0 \quad \forall u \neq 0 \in \mathbb{R}^d$$

$$(f''(\theta^*) \succ 0)$$

(*) only local results in general

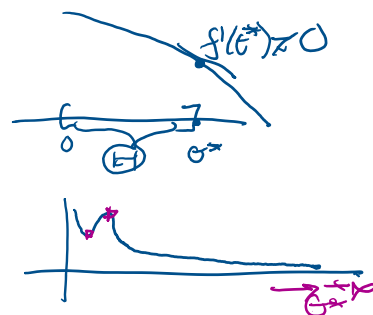
• but if $\text{Hessian}(f(\theta)) \succ 0 \quad \forall \theta \in \Theta$, $f(\cdot)$ is said to be "convex"

then $Df(\theta^*) = 0 \Rightarrow \theta^*$ is a global min

• otherwise, for smooth $f(\cdot)$, looking at zero gradient pts & boundary points
gives you enough info to find global min (**)

(*) be careful with boundary cases
i.e. $\theta^* \in \text{boundary}(\Theta)$

other example:



* some notes about MLE

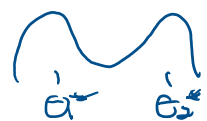
• does not always exist

$[\theta^* \in \text{bd}(\Theta) \text{ but } \Theta \text{ is open}]$ or when " $\theta^* = \pm \infty$ "

• it is not rec unique [ex. maxtime models]

• is not "admissible" in general [see later]

$\exists$ strictly "better" estimator



example II: multinomial distribution

suppose X_i is a discrete R.V. on k choices "multinomial"

(one could choose $\Omega_{X_i} = \{1, 2, 3, \dots, k\}$)

but instead, convenient to encode the k possibilities using unit basis in $\mathbb{R}^k$

i.e. $\Omega_{X_i} = \{e_1, e_2, \dots, e_k\}$ where $e_j \in \mathbb{R}^k$ "one hot encoding"

$$e_j = \begin{pmatrix} 0 \\ \vdots \\ 1 \\ \vdots \\ 0 \end{pmatrix} \text{ } j^{\text{th}} \text{ coordinate}$$

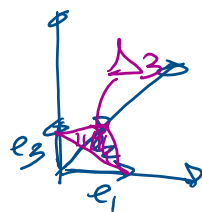
parameter for discrete RV $\pi \in \mathbb{R}^k$

$$(\Theta) = \Delta_k$$

$$\pi \in \Delta_k \quad \Delta_k = \left\{ \pi \in \mathbb{R}^k : \pi_j \geq 0; \sum_{j=1}^k \pi_j = 1 \right\}$$

probability simplex on k choices

$$\text{we will write } X_i \sim \text{Mult}(\pi) \quad \text{parameter} = \text{Mult}(1, \pi)$$



$$x_i \rightarrow \mathbb{R}^k$$

(*) consider $X_i \stackrel{\text{i.i.d.}}{\sim} \text{Mult}(\pi)$

$$\text{then } X = \sum_{i=1}^n X_i \sim \text{Mult}(n, \pi)$$

"multinomial distribution"

$$X \in \mathbb{N}^k \quad \Omega_X = \left\{ (n_1, n_2, \dots, n_k) : n_j \in \mathbb{N}; \sum_{j=1}^k n_j = n \right\}$$

pdf for X :

$$x = (n_1, \dots, n_k)$$

$$(x)_j \triangleq n_j$$

$$p(x|\pi) = \binom{n}{(x)_1, (x)_2, \dots, (x)_k} \prod_{j=1}^k \pi_j^{(x)_j}$$

multinomial coeff

$$\binom{n}{n_1, n_2, \dots, n_k} \triangleq \frac{n!}{n_1! n_2! \dots n_k!}$$

15h15 multinomial MLE

$$\text{log-likelihood } l(\pi) = \log p(x|\pi) = \log \binom{n}{n_1, \dots, n_k} + \sum_{j=1}^k n_j \log \pi_j$$

$$x = (n_1, \dots, n_k)$$

constant $\rightarrow$ ignore

$$\text{MLE } \hat{\pi}_{\text{MLE}}(x) = \underset{\pi \in \mathbb{R}^k}{\text{argmax}} l(\pi)$$

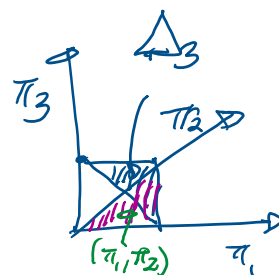
s.t. $\pi \in \Delta_k$ } constraint

two options:

a) reparameterize problem so that Θ is full dimensional

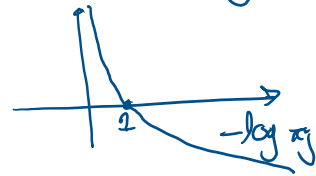
$$\pi_k = 1 - \sum_{j=1}^{k-1} \pi_j$$

$$\rightarrow \pi_1, \dots, \pi_{k-1} \in [0, 1] \text{ with constraint } \sum_{j=1}^{k-1} \pi_j \leq 1$$



$\rightarrow \pi_1, \dots, \pi_{k-1} \in [0, 1]$ with constraint $\sum_{j=1}^{k-1} \pi_j \leq 1$
 here magic that $\log \pi_j$ acts as a barrier fct. away from $\pi_j = 0$

can try unconstrained opt. on $\pi_1, \dots, \pi_{k-1}$
 of $l(\pi_1, \dots, \pi_{k-1}, 1 - \sum_{j=1}^{k-1} \pi_j)$



hoping sol'n is in interior of constraint set
 (and it is usually case for log-type problems)

b) use Lagrange multiplier approach to handle equality constraint on Δ
 [and still ignoring $\pi_j \in [0, 1]$]

$$\begin{aligned}
 &\max f(\pi) \\
 &\text{st. } g(\pi) = 0 \\
 &[g(\pi) \triangleq 1 - \sum_{j=1}^k \pi_j]
 \end{aligned}$$

$$J(\pi, \lambda) \triangleq f(\pi) + \lambda g(\pi)$$

Lagrange multiplier

method: look at stationary pt. of $J(\pi, \lambda)$
 (0-gradient)

$$\text{ie. } \nabla_{\pi} J(\pi, \lambda) = 0$$

$$\nabla_{\lambda} J(\pi, \lambda) = 0 \Rightarrow g(\pi) = 0$$

Necessary cond. for local opt.

(check "bordered Hessian" to get local or max)

(see [Wikipedia](https://en.wikipedia.org/wiki/Bordered_Hessian))

$$l(\pi) = \sum_j \pi_j \log \pi_j$$

(strictly concave fct. in π_j)

$$\text{want } \frac{\partial J}{\partial \pi_j} = 0 \Rightarrow \frac{n_j}{\pi_j} - \lambda = 0 \Rightarrow \pi_j^* = \frac{n_j}{\lambda} \quad \left\{ \begin{array}{l} \text{want} \\ \text{scaling constant} \end{array} \right.$$

$$\begin{aligned}
 &\text{want } g(\pi^*) = 0 \text{ i.e. } \sum_j \pi_j^* = 1 \\
 &\text{i.e. } \sum_j \frac{n_j}{\lambda} = 1 \\
 &\Rightarrow \lambda^* = \sum_{j=1}^n n_j = n
 \end{aligned}$$

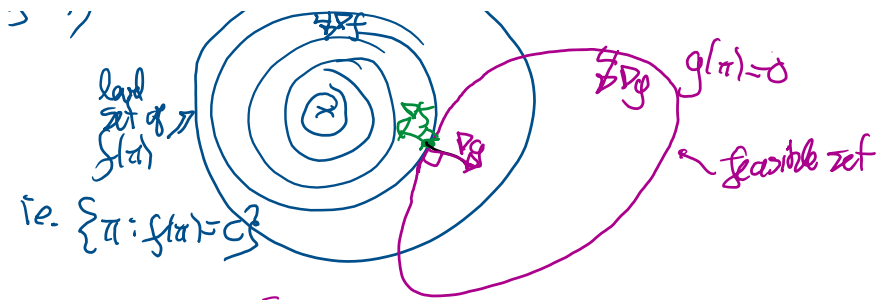
notice: $\pi_j^* \in [0, 1]$ $\pi_j^* = \frac{n_j}{n}$ MLE for multinomial

$$\nabla_{\pi} L(\pi, \lambda) = 0 \Rightarrow \nabla_{\pi} f(\pi) + \lambda \nabla_{\pi} g(\pi) = 0$$

$$\Rightarrow \nabla f(\pi) = -\lambda \nabla g(\pi)$$

(max $f(\pi)$)





Statistical decision theory

A) Bias-variance decomposition for squared loss

estimator: fct. from data (observation) to parameter

$$\text{MLE: } \hat{\theta}_{\text{MLE}}(x) = \underset{\theta \in \Theta}{\text{argmax}} p(x|\theta)$$

$$\text{MAP: } \hat{\theta}_{\text{MAP}}(x) = \underset{\theta \in \Theta}{\text{argmax}} p(\theta|x) = \underset{\theta \in \Theta}{\text{argmax}} \underbrace{p(x|\theta)}_{\text{likelihood}} \underbrace{p(\theta)}_{\text{prior}}$$

$\log p(x|\theta) + \log p(\theta)$
regulation

* how do we evaluate these estimators? estimator $\delta: \mathcal{X} \rightarrow \Theta$
 $\hat{\theta} = \delta(x)$

most standard fct: frequentist risk of an estimator

$$R(\theta, \delta) \triangleq \mathbb{E}_x [L(\theta, \delta(x))]$$

average over possible data / training set
"statistical loss fct."

squared loss, $L(\theta, \hat{\theta}) \triangleq \|\theta - \hat{\theta}\|_2^2$

$$\begin{aligned} \mathbb{E}_x [\|\theta - \hat{\theta}\|_2^2] &= \mathbb{E} [\|\underbrace{\theta - \mathbb{E}\hat{\theta}}_0 + \underbrace{\mathbb{E}\hat{\theta} - \hat{\theta}}_0\|_2^2] \quad \hat{\theta} = \delta(x) \\ &= \mathbb{E} [\|\theta - \mathbb{E}\hat{\theta}\|_2^2] + \mathbb{E} [\|\mathbb{E}\hat{\theta} - \hat{\theta}\|_2^2] \quad \|a+b\|^2 = \|a\|^2 + \|b\|^2 + 2\langle a, b \rangle \\ &\quad + 2\mathbb{E} [\langle \theta - \mathbb{E}\hat{\theta}, \mathbb{E}\hat{\theta} - \hat{\theta} \rangle] \\ &= 2\langle \theta - \mathbb{E}\hat{\theta}, \mathbb{E}(\mathbb{E}\hat{\theta} - \hat{\theta}) \rangle \quad \text{constant} \\ &= 2\langle \theta - \mathbb{E}\hat{\theta}, \mathbb{E}(\mathbb{E}\hat{\theta} - \hat{\theta}) \rangle \end{aligned}$$

$$R(\theta, \delta) = \mathbb{E}_x [\|\theta - \hat{\theta}\|_2^2] = \underbrace{\|\theta - \mathbb{E}\hat{\theta}\|_2^2}_{\text{bias}^2} + \underbrace{\mathbb{E} [\|\mathbb{E}\hat{\theta} - \hat{\theta}\|_2^2]}_{\text{variance}}$$

$$\text{bias} \triangleq \|\theta - \mathbb{E}\hat{\theta}\|_2$$

(freq.) risk for squared loss = bias² + variance. "bias-variance."

$$\text{bias} \equiv \|E - E(\hat{\theta})\|_2$$

(freq.) risk for squared loss = bias² + variance.

"bias-variance tradeoff"

⊗ consistency: informally "do right thing as $n \rightarrow \infty$ " where n is the training set size

$$\hat{\theta}_n \xrightarrow{P} \theta$$

$$X \sim (X_i)_{i=1}^n$$

$\hat{\theta}_n$ (data of size n)

assignment: if $\text{bias}(\hat{\theta}_n) \xrightarrow{n \rightarrow \infty} 0$
and $\text{variance}(\hat{\theta}_n) \xrightarrow{n \rightarrow \infty} 0 \Rightarrow R(\theta, \hat{\theta}_n) \xrightarrow{n \rightarrow \infty} 0 \Rightarrow (\hat{\theta}_n \xrightarrow{P} \theta)$
i.e. $\hat{\theta}_n$ is consistent