

today : • MLE tied
• statistical decision theory

optimization comments about MLE

$$\min_{\theta \in \Theta} f(\theta)$$

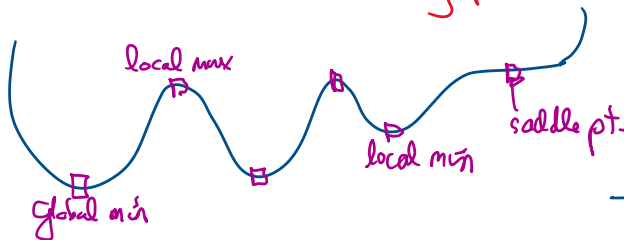
$$\nabla f(\theta^*) = 0$$

"stationary pts."

(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 $\text{Hessian}(f)(\theta^*) > 0$
for a local min

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

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

⊗ only local results in general

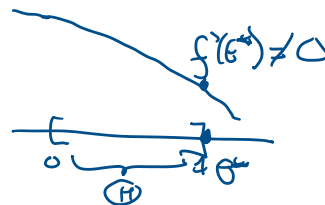
• but if $\text{Hessian}(f(\theta)) > 0 \quad \forall \theta \in \Theta$, $f(\theta)$ is said to be "convex"

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

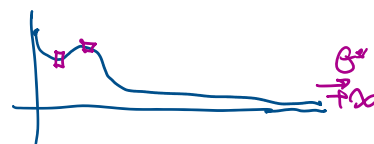
• otherwise, for smooth $f(\theta)$, looking at all zeros gradient pts and boundary pts
give you enough information to find global min ⊗⊗

⊗ be careful with boundary cases

ie. $\theta^* \in \text{boundary}(\Theta)$ e.g.



other example :



↑ example where MLE does not exist

✗ some notes about MLE

• does not always exist [$\theta^* \in \text{bd}(\Theta)$ but Θ is open] or when " $\theta^* = +\infty$ "

$$\Theta =]0, 1[$$

• it is not rec. unique [e.g. mixture models] 

• is not "admissible" in general [see later]
 $\exists$ strictly "better" estimators

example II: multinomial distribution

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

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

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

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

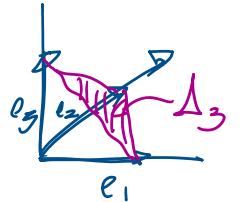
$$\pi \in \mathbb{R}^k \quad e_j = \begin{pmatrix} 0 \\ \vdots \\ 1 \text{ at } j\text{th coordinate} \\ \vdots \\ 0 \end{pmatrix}$$

parameter for discrete R.V. $\pi \in \Delta_k$ ($\Theta = \Delta_k$)

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

probability simplex on k choices

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



⊕ consider $X_i \stackrel{\text{iid}}{\sim} \text{Mult}(\pi)$

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

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

pmf for X :

$$p(x|\pi) = \binom{n}{(x_1, \dots, x_k)} \prod_{j=1}^k \pi_j^{x_j}$$

multinomial coeff

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

$$(x)_j \triangleq n_j$$

[vs. x_i of X_i]

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

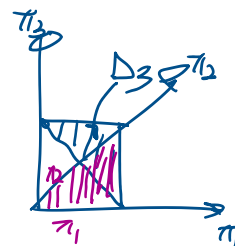
multinomial MLE

$$\text{log-likelihood } l(\pi) = \log p(x|\pi) = \underbrace{\log \binom{n}{n_1, \dots, n_k}}_{\substack{\text{constant during} \\ \rightarrow \text{ignore MLE}}} + \sum_{j=1}^k n_j \log \pi_j$$

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

$$\text{MLE: } \hat{\pi}_{\text{MLE}}(x) = \underset{\substack{\pi \in \mathbb{R}^k \\ \text{s.t. } \pi \in \Delta_K}}{\text{argmax}} l(\pi)$$

constraint



two options:

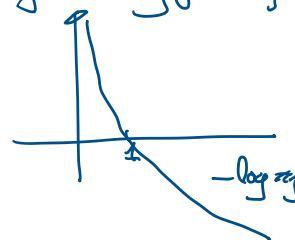
a) reparameterize problem so that Θ is full dimensional

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

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

here magic that $\log \pi_j$ acts as a barrier fct. away from $\pi_j \rightarrow 0$

can try unconstrained optimization on $\pi_1, \dots, \pi_{k-1}$
of $l(\pi_1, \dots, \pi_{k-1})$



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

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

$$\max f(\pi)$$

$$\text{s.t. } g(\pi) = 0$$

$$\left[1 - \sum_{j=1}^{k-1} \pi_j = 0 \right]$$

$$\triangleq g(\pi)$$

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

Lagrange multiplier

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

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

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

$$\Rightarrow g(\pi) = 0$$

necessary cond. for local opt.

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

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

(strictly concave
fct. in π_j)

$$\frac{\partial \ell}{\partial \pi_j} = 0$$

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

$$\text{want } g(\pi^*) = 0 \quad \text{i.e. } \sum_j \pi_j^* = 1$$

$$\sum_j \frac{n_j}{\lambda} = 1$$

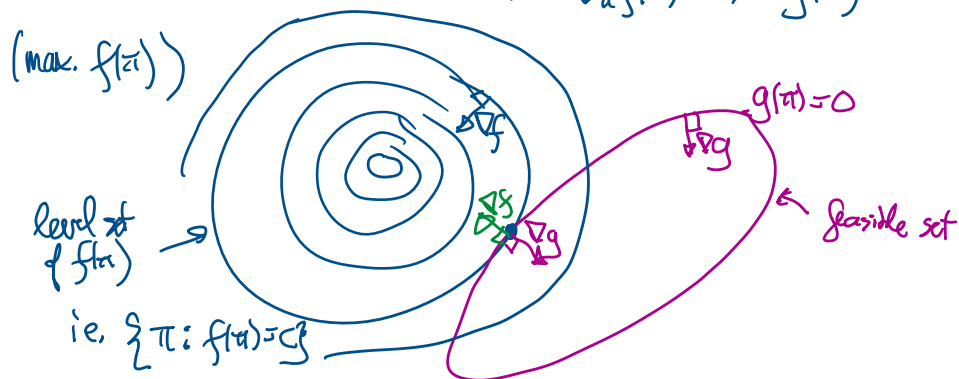
$$\Rightarrow \lambda^* = \sum_j n_j = n$$

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

$$\boxed{\pi_j^* = \frac{n_j}{n}} \quad \text{MLE for multinomial}$$

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

$$\Rightarrow \nabla_{\pi} f(\pi) = -\lambda \nabla_{\pi} g(\pi)$$



Statistical decision theory

A) Bias-variance decomposition for squared loss

estimator: fct. from data (observation) to parameter

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

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

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

* how do we evaluate these estimators?

estimator $\delta: \Omega_X \rightarrow \Theta$

$$\hat{\theta} = \delta(X)$$

most standard tool: frequentist risk of an estimator

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

$\frac{1}{n} \sum_{i=1}^n L(\theta, (x_i, y_i))$
 $\hookrightarrow$ average over possible data
 (statistical) loss fun.

Squared loss: $L(\theta, \theta') \triangleq \|\theta - \theta'\|_2^2$ $\theta = g(x)$

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

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

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

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

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

$X \rightarrow (x_i)_{i=1}^n$
 $\hat{\theta}_n$ (data of size n)

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