

Lec 11: Application, Statistical estimation



- Parametric distribution estimation

choose from a family of densities $p_x(y)$, indexed by a parameter x

- maximum likelihood estimation

$$\max_x p_x(y)$$

y is observed value

$\ell(x) = \log p_x(y)$ is called log-likelihood function

many densities are log-concave (in x vs. in y)

can add constraints $x \in C$ explicitly, or define $p_x(y) = 0$ for $x \notin C$
a convex optimization problem if $\log p_x(y)$ is concave in x for fixed y

eg. linear measurement with IID noise

$$y_i = a_i^T x + v_i$$

$x \in \mathbb{R}^n$ is the vector of unknown parameters

v_i is IID measurement noise, with density p_{v_i}

y_i is measurement: y_i has density $p_{y_i}(y) = \int_{\mathbb{R}^n} p_{v_i}(y - a_i^T x) p_x(x) dx$

$$\max_x \ell(x) = \sum \log p(y_i - a_i^T x) \quad (\text{need } p \text{ to be log-concave})$$

eg. Gaussian noise $\mathcal{N}(0, \sigma^2)$: $p(z) = \frac{1}{\sqrt{2\pi}\sigma} \cdot \exp\left(-\frac{z^2}{2\sigma^2}\right)$

$$\ell(x) = \sum_{i=1}^m \log p(y_i - a_i^T x)$$

$$= m \cdot \log \frac{1}{\sqrt{2\pi}\sigma} - \underbrace{\frac{1}{2\sigma^2} \cdot \sum_{i=1}^m (y_i - a_i^T x)^2}$$

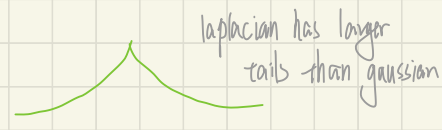
least-square is maximum likelihood estimation with Gaussian noise

eg. Laplacian noise $p(z) = \frac{1}{2a} \cdot \exp(-|z|/a)$

$$\ell(x) = \sum_{i=1}^m \log \left[\frac{1}{2a} \cdot \exp(-|y_i - a_i^T x|/a) \right]$$

$$= -m \cdot \log(2a) - \frac{1}{a} \sum_{i=1}^m |y_i - a_i^T x|$$

minimizing L_1 norm is mle with Laplacian noise



eg. uniform noise in $[-a, a]$

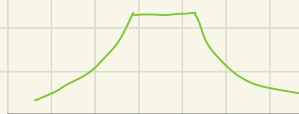
$$\ell(x) = \begin{cases} -m \log 2a & |a^T x - y_i| \leq a \\ \infty & \text{otherwise} \end{cases}$$

mle estimation is any x with $|a^T x - y_i| \leq a$ (feasibility problem)

eg. penalty function



penalty function

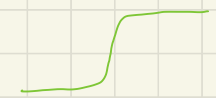


noise distribution
(uniform center and exponential tail)

eg. logistic regression

random variable $y \in \{0, 1\}$ with distribution

$$\hat{p} = \text{prob}(y=1) = \frac{\exp(a^T u + b)}{1 + \exp(a^T u + b)}$$



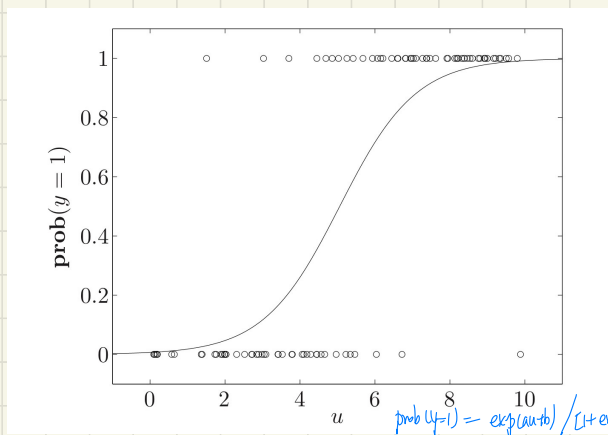
a, b are parameters, $u \in \mathbb{R}^n$ are observations

estimation problem: estimate a, b from (u_i, y_i)

$(y_1 = \dots = y_k = 1, y_{k+1} = \dots = y_m = 0)$

$$\begin{aligned} \ell(a, b) &= \log \left(\prod_{i=1}^k \frac{\exp(a^T u_i + b)}{1 + \exp(a^T u_i + b)} \cdot \prod_{i=k+1}^m \frac{1}{1 + \exp(a^T u_i + b)} \right) \\ &= \sum_{i=1}^k (a^T u_i + b) - \sum_{i=1}^m \log[1 + \exp(a^T u_i + b)] \end{aligned}$$

can come in a, b



$$\text{prob}(y=1) = \frac{\exp(a^T u + b)}{1 + \exp(a^T u + b)}$$

b controls the natural point
 a controls the stretch

eg. Covariance estimation for Gaussian variables
 Suppose $y \in \mathbb{R}^n$ is a Gaussian random variable
 $E[y] = 0$ $R = \text{cov}(y) = E[yy^T]$

$$p_R(y) = \frac{1}{(2\pi)^{n/2} |R|^{1/2}} \exp\left(-\frac{y^T R^{-1} y}{2}\right) \quad R \in S_{++}^n$$

$$\ell(R) = \sum_{i=1}^m \log p(y_i)$$

$$= -m \cdot \left[\frac{n}{2} \cdot \log 2\pi + \frac{1}{2} \log \det(R) \right] - \frac{1}{2} \sum_{i=1}^m y_i^T R^{-1} y_i$$

$$= -m \left[\frac{n}{2} \cdot \log 2\pi + \frac{1}{2} \log \det(R) \right] - \frac{m}{2} \text{tr}(R^{-1} Y) \quad Y = \frac{1}{m} \sum_{i=1}^m y_i y_i^T \text{ is the sample covariance}$$

let $S = R^{-1}$ (the information matrix)

$$= -\frac{mn}{2} \cdot \log 2\pi + m/2 \log \det S - \frac{m}{2} \text{tr}(SY)$$

is concave in S

$$\max_S \log \det(S) - \text{tr}(SY)$$

1° no constraints on S

$$\max_{S \in S_{++}^n} \log \det(S) - \text{tr}(SY) \quad (S = R^{-1})$$

$$\nabla_S (\log \det(S) - \text{tr}(SY)) = S^{-1} - Y = 0$$

$$R = S^{-1} = Y \quad \text{if } Y \in S_{++}^n$$

the problem is unbounded below if $Y \notin S_{++}^n$

if there's no prior assumptions on R ,
 then mle estimation of R is sample covariance

• maximum a posteriori probability estimation.

(a bayesian version of MLE)

assume that x (to be estimated) and y (the observation) are random variables with a joint probability $p(x,y)$

the prior density is given by $p_x(x) = \int p(x,y) dy$ $p_y(y) = \int p(x,y) dx$

the conditional density is given by $p_{y|x}(y|x) = \frac{p(x,y)}{p_x(x)}$

$$\hat{x}_{\text{map}} = \underset{x}{\operatorname{argmax}} p_{x|y}(x|y)$$

$$= \underset{x}{\operatorname{argmax}} p_{x|y}(x|y) \cdot p(y)$$

$$= \underset{x}{\operatorname{argmax}} p(x,y)$$

$$= \underset{x}{\operatorname{argmax}} p_{y|x}(y|x) \cdot p_x(x) \quad \text{taking prior knowledge of } x \text{ into account compared to MLE}$$

eg. linear measurement with IID noise

$$y_i = a_i x + v_i$$

v_i are iid noise with density p_v on $\mathbb{R}$, x has prior density p_x on $\mathbb{R}^n$

$$p(x,y) = p_x(x) \cdot \prod_{i=1}^n p_v(y_i - a_i x)$$

$\Downarrow$

$$\max_x \log p_x(x) + \sum_{i=1}^n \log p_v(y_i - a_i x)$$

when v_i is uniform on $[-a, a]$, and the prior distribution of x is gaussian

$$\max_x \log \frac{1}{(2\pi)^{n/2} |\Sigma|^{1/2}} \cdot \exp\left(-\frac{1}{2} (x - \bar{x})^T \Sigma^{-1} (x - \bar{x})\right)$$

$$\text{s.t. } \|Ax - b\|_\infty \leq a$$

$\Downarrow$

$$\min_x (x - \bar{x})^T \Sigma^{-1} (x - \bar{x})$$

$$\text{s.t. } \|Ax - b\|_\infty \leq a$$

eg. MLE with perfect measurements

$x \in \mathbb{R}^n$ is a vector of parameters to be estimated, p_x have m perfect, noise-free measurements $y = Ax$

$$\max_x \log p_x(x)$$

$$\text{s.t. } Ax = y$$

• (Binary) hypothesis testing

1. detection (hypothesis testing) problem

given observation of a random variable $X \in \{1, \dots, n\}$ drawn between:

hypothesis 1: X was generated by distribution $p = (p_1, \dots, p_n)$ $\mathbb{1}^T p = 1$ $p \geq 0$
 hypothesis 2: X was generated by distribution $q = (q_1, \dots, q_n)$ $\mathbb{1}^T q = 1$ $q \geq 0$

2. randomized detector

a nonnegative matrix $T \in \mathbb{R}^{n \times 2}$, with $\mathbb{1}^T T = \mathbb{1}$

if we observe $X=k$, we choose hypo 1 with prob T_{1k} , hypo 2 with prob T_{2k}
 if T has all elements 0 or 1, it's a deterministic detector

3. detection probability matrix

$$D = T \begin{bmatrix} p \\ q \end{bmatrix} = \begin{bmatrix} T_{1p} & T_{1q} \\ T_{2p} & T_{2q} \end{bmatrix}$$

T_{1p} is probability of choosing hypothesis 1 when X generated by distribution 1

eg. multicriterion formulation of detector design.

$$\min. \text{ (w.r.t } \mathbb{R}^+) \quad (T_{1p}, T_{1q}) = (T_{1p}, T_{1q})$$

s.t.

$$T_{1k} + T_{2k} = 1$$

$$T_{ik} \geq 0$$

↓

$$\min. \quad (T_{1p}) + (T_{1q})$$

$$\text{s.t.} \quad T_{1k} + T_{2k} = 1$$

$$T_{ik} \geq 0$$

a LP with a simple analytic solution

$$(T_{1k}, T_{2k}) = \begin{cases} (1, 0) & p_k \geq \lambda \cdot q_k \\ (0, 1) & p_k < \lambda \cdot q_k \end{cases}$$

- a deterministic detector, given by a likelihood ratio test ($p_k/q_k \stackrel{?}{\geq} \lambda$)
- when $\lambda=1$, it's minimizing $P(\text{being wrong}) \Leftrightarrow \max_i P(\text{being right})$ (MLE)

eg. minimax detector

$$\min. \max \{T_{1p}, T_{1q}\} = \max \{T_{1p}, T_{1q}\}$$

$$\text{s.t.} \quad T_{1k} + T_{2k} = 1 \quad T_{ik} \geq 0$$

a LP; solution is usually not deterministic

• Experiment design

m linear measurements $y_i = a_i^T x + w_i$ of unknown $x \in \mathbb{R}^n$
 $w_i \sim \mathcal{N}(0, 1)$ iid

$$A = \begin{bmatrix} -a_1^T \\ \vdots \\ -a_m^T \end{bmatrix}$$

least-square estimate is: $\hat{x} = (\sum_i a_i a_i^T)^{-1} \sum_i y_i a_i$ $(A^T A)^{-1} A^T y$

error $e = \hat{x} - x$ has zero mean and covariance: $E = E[ee^T] = (\sum_i a_i a_i^T)^{-1}$

$$\begin{aligned} \hat{x} &= (A^T A)^{-1} A^T y \\ &= (A^T A)^{-1} A^T (Ax + w) \\ &= (A^T A)^{-1} A^T A x + (A^T A)^{-1} A^T w \\ &= x + (A^T A)^{-1} A^T w \end{aligned}$$

$$E[\hat{x}] = x \quad (w \sim \mathcal{N}(0, 1))$$

$$E[(\hat{x} - x)(\hat{x} - x)^T] = E[(A^T A)^{-1} A^T w w^T A (A^T A)^{-1}]$$

$$= (A^T A)^{-1} A^T E[ww^T] A (A^T A)^{-1}$$

$$E = (A^T A)^{-1} \quad (\text{if } A \text{ large, then it has high signal-to-noise ratio, so covariance small})$$

confidence ellipsoid is given by: $\{x \mid (x - \bar{x})^T E^{-1} (x - \bar{x}) \leq \beta\}$

experiment design: choose a set of possible tests to make E small
 $\{ \text{choose } a_i \text{ for } m_i \text{ times} \quad \sum_i m_i = m \}$

eg sensor a_1 has the highest signal-to-noise ratio,
if choose on for 100 times and no other a_i .
then $\sum a_i a_i^T$ is not invertible (rank 1)

↓

$$\text{min. (w.r.t } S) \quad E = \left(\sum_{k=1}^K m_k \cdot v_k v_k^T \right)^{-1}$$

$$\text{s.t.} \quad m_k \geq 0 \quad m_1 + \dots + m_p = m$$

$$m_k \in \mathbb{Z}$$

(hard to solve due to the integer constraint.)

↓

Relaxed experiment design

$$\min. (\text{w.r.t } S) \quad E = \frac{1}{n} \left(\sum_{k=1}^p \lambda_k \cdot v_k v_k^T \right)^{-1}$$

$$\text{s.t.} \quad \lambda \geq 0 \\ \mathbf{1}^T \lambda = 1$$



Common scalarization: $\min. \log \det(E), \text{tr}(E), \lambda_{\max}(E)$

eg. D-optimal design

$$\min. \log \det \left(\sum_{k=1}^p \lambda_k v_k v_k^T \right)^{-1}$$

$$\text{s.t.} \quad \lambda \geq 0 \quad \mathbf{1}^T \lambda = 1$$

(minimize the volume of confidence ellipsoid)

||.

$$\min. \log \det(X^{-1})$$

$$\text{s.t.} \quad X = \sum_{k=1}^p \lambda_k v_k v_k^T \quad (2) \quad (v_k \in \mathbb{R}^n)$$

$$\lambda \geq 0 \quad \mathbf{1}^T \lambda = 1$$

(a) (b)

$$\begin{aligned} \mathcal{L}(X, \lambda, z, \beta) &= \log \det(X^{-1}) + \text{tr} \left\{ z \left(X - \sum_{k=1}^p \lambda_k v_k v_k^T \right) \right\} - \alpha^T \lambda + \beta (\mathbf{1}^T \lambda - 1) \\ &= \log \det(X^{-1}) + \text{tr}(zX) - \text{tr} \left\{ z \cdot \left(\sum_{k=1}^p \lambda_k v_k v_k^T \right) \right\} - \alpha^T \lambda + \beta (\mathbf{1}^T \lambda - 1) \\ &= \log \det(X^{-1}) + \text{tr}(zX) - \sum_{k=1}^p \lambda_k \cdot v_k^T z v_k - \sum_{i=1}^n \alpha_i \lambda_i + \sum_{k=1}^p \beta \cdot \lambda_k - \beta \\ &= \log \det(X^{-1}) + \text{tr}(zX) + \sum_{k=1}^p \lambda_k (-v_k^T z v_k - \alpha_k + \beta) - \beta \end{aligned}$$

$$\begin{aligned} g(\lambda, z, \beta) &= \inf_{\lambda \geq 0, \mathbf{1}^T \lambda = 1} \left\{ \mathcal{L}(X, \lambda, z, \beta) \right\} \\ &= \begin{cases} \log \det(X^{-1}) + \text{tr}(zX) - \beta & \text{if } -v_i^T z v_i - \alpha_i + \beta \geq 0 \\ -\infty & \text{otherwise} \end{cases} \end{aligned}$$

$$\nabla_z \left(\log \det(X^{-1}) + \text{tr}(zX) \right) = 0 \\ -X^{-1} + Z = 0 \quad Z = X^{-1}$$

$$\therefore g(\lambda, z, \beta) = \begin{cases} \log \det Z + n - \beta & \text{if } -v_i^T z v_i - \alpha_i + \beta \geq 0 \\ -\infty & \text{otherwise} \end{cases}$$

↓

$$\max. \log \det Z + n - \beta$$

$$\text{s.t.} \quad \beta - v_i^T z v_i \geq 0$$

$$\downarrow \quad W = (\beta) \cdot Z$$

$$\max. \quad \log \det W + n \log \beta + n - \beta$$

$$\text{s.t.} \quad v_i^T W v_i \leq 1$$

$$\Downarrow \quad \max \quad n \log \beta - \beta \quad \text{over } \beta : \beta = n$$

dual of D-optimal

$$\max. \quad \log \det W + n \log n$$

$$\text{s.t.} \quad v_i^T W v_i \leq 1$$

$\{x | x^T W x \leq 1\}$ is the minimum volume ellipsoid at origin that includes all test vector v_i

complementary slackness:

for λ, W that's a primal-optimal - dual-optimal pair

$$\lambda_i (1 - v_i^T W v_i) = 0$$

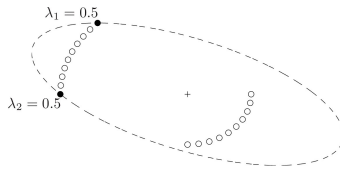


Figure 7.9 Experiment design example. The 20 candidate measurement vectors are indicated with circles. The D -optimal design uses the two measurement vectors indicated with solid circles, and puts an equal weight $\lambda_i = 0.5$ on each of them. The ellipsoid is the minimum volume ellipsoid centered at the origin, that contains the points v_i .