

MULTIDIMENSIONAL NORMAL DISTRIBUTION (GAUSSIAN)

SINGLE

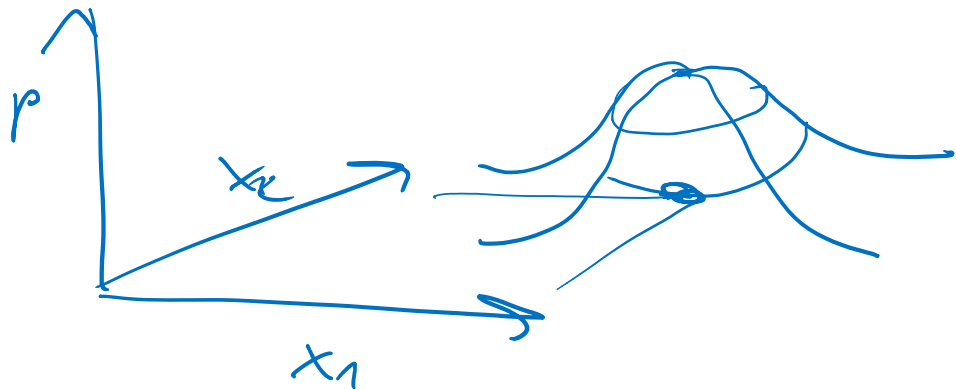
$$p(x) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(x-\mu)^2}{2\sigma^2}\right)$$

MULTI: $\mu \in \mathbb{R}^n$, $\Sigma \in \mathbb{R}^{n \times n}$

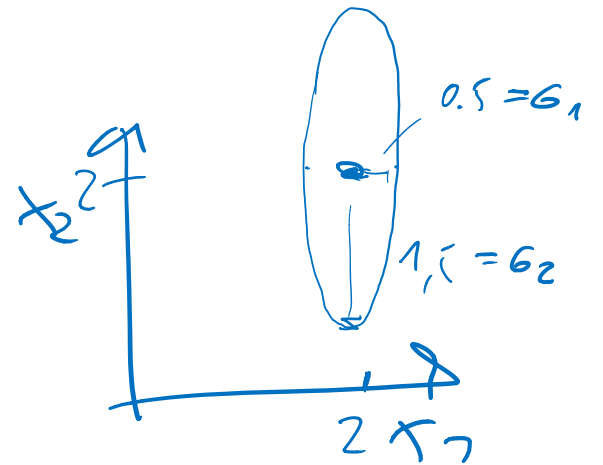
Σ is > 0 , SYMMETRIC

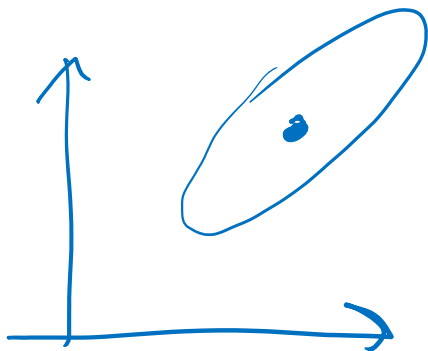
"COV. MATRIX"

$$p(x) = \frac{1}{\sqrt{\det(2\pi\Sigma)}} \cdot \exp\left(-\frac{(x-\mu)^T \Sigma^{-1} (x-\mu)}{2}\right)$$



$$\mu = \begin{pmatrix} 2 \\ 2 \end{pmatrix}, \Sigma = \begin{pmatrix} \sigma_1^2 & 0 \\ 0 & \sigma_2^2 \end{pmatrix} = \begin{pmatrix} 1/4 & 0 \\ 0 & 9/4 \end{pmatrix}$$

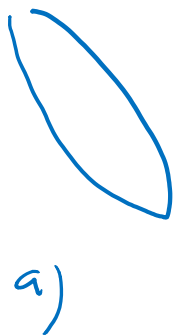




$$\Sigma = \begin{pmatrix} 1 & 0.5 \\ 0.5 & 1 \end{pmatrix}, \quad \mu = \begin{pmatrix} 2 \\ 2 \end{pmatrix}$$

$$\Sigma = V D V^T \quad \text{where} \quad D = \begin{pmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_n \end{pmatrix} \succeq 0$$

$V^T V = I$



$$\Sigma = \begin{pmatrix} G_1^2 & G_1 G_2 & \dots \\ & G_2^2 & \dots \\ & & \ddots & \\ & & & G_n^2 \end{pmatrix}$$

Correlation $[-1, 1]$

RANDOM VARS.

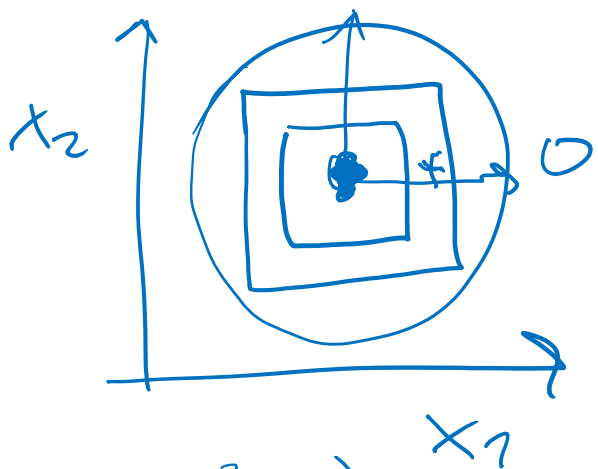
CORRELATION $\Rightarrow$ DEPENDENCE

MEAN = EXPECTATION

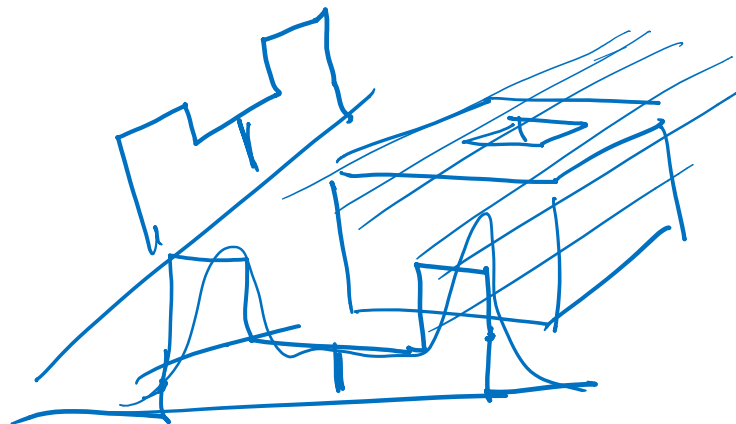
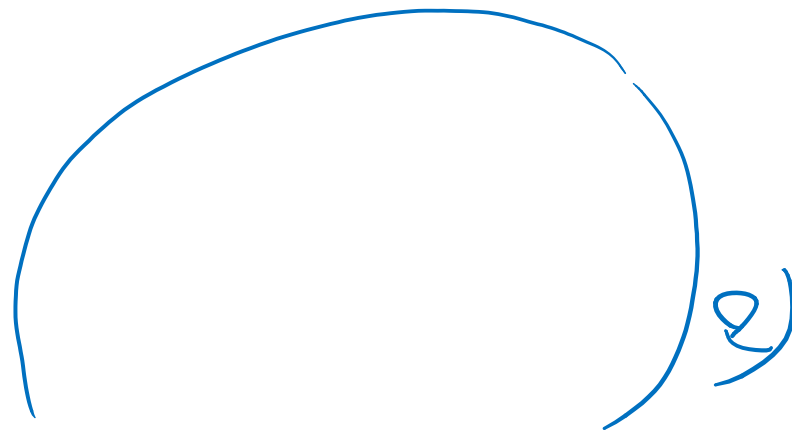
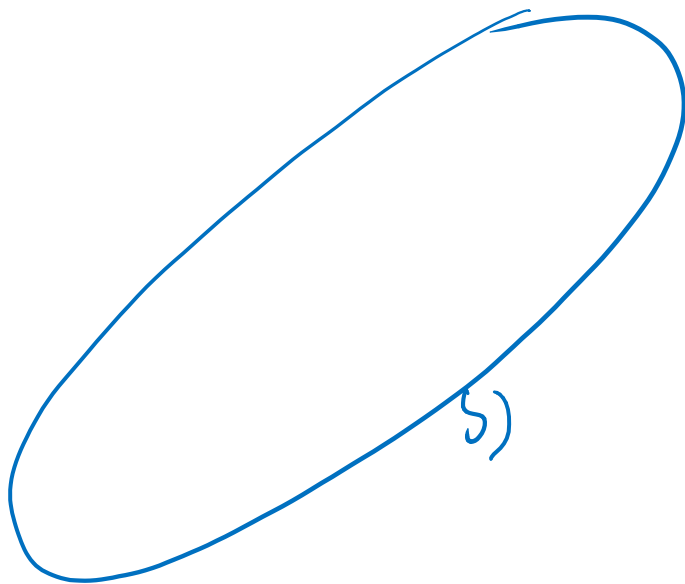
VARIANCE

COVARIANCE MATRIX

LINEAR (AFFINE) MAP: MEAN? COV. M.?



$$\Sigma = \begin{pmatrix} G^2 & 0 \\ 0 & G^2 \end{pmatrix}$$



UNCORRELATED ? YES / NO

INDEPENDENT ? YES / NO

$$p(x_1, x_2) = p(x_1) \cdot p(x_2)$$

BAYES THM

$$p(x_2|x_1) = \frac{\overbrace{p(x_1|x_2)}^{\text{A-POSTERIORI}} \cdot \overbrace{p(x_2)}^{\text{COND. PROB. OF } x_2 \text{ GIVEN } x_1}}{\underbrace{p(x_1)}_{\text{A-PRIORI}}}$$

PAR → OBS.

CONDITIONAL PDF: $p(x_1|x_2) = \frac{p(x_1, x_2)}{p(x_2)} \Leftrightarrow p(x_1, x_2) = p(x_1|x_2) \cdot p(x_2)$

ESTIMATOR: GIVEN DATA $x_N \in \mathbb{R}^N$, EST. SOME $\theta \in \mathbb{R}^1$. y_N IS RANDOM

MAP FROM MEAS. TO POINT ESTIMATE

$$\hat{\theta}_N(y_N)$$

EX: "SAMPLE MEAN": $\hat{\theta} = \frac{\sum y(i)}{N} =: \mu(y_N)$

"SAMPLE VARIANCE":

$$\frac{\sum_i (y(i) - \mu(y_N))^2}{(N-1)}$$

~~$(N-1)$~~
 ~~(N)~~

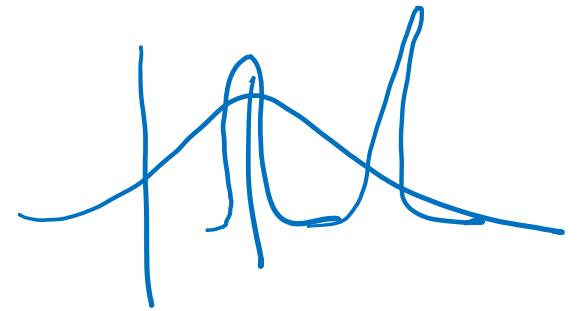
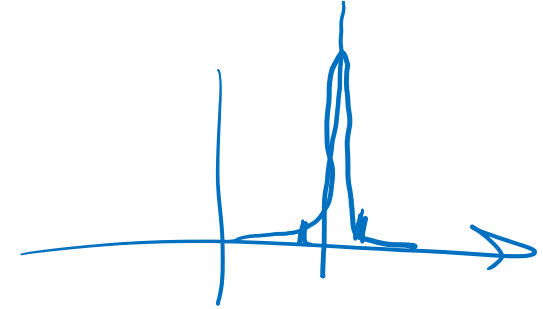
UNBIASED ESTIMATORS FOR (TRUE) MEAN AND VARIANCE OF $y(i)$ (i.i.d.)

UNBIASED ESTIMATOR: $E[\hat{\theta}] = \text{TRUE VAL.}$

ASYMPTOTIC UNBIASED EST.: $E[\lim_{N \rightarrow \infty} \hat{\theta}_N] = \text{TRUE VAL.}$

CONSISTENCY: $\forall \epsilon > 0: \lim_{N \rightarrow \infty} P(\hat{\theta}_N \in [\theta_0 - \epsilon, \theta_0 + \epsilon]) = 1$

$\uparrow$
AS UNBIASEDNESS & $\text{COV}(\hat{\theta}) \rightarrow 0$
 $N \rightarrow \infty$



LS-ESTIMATORS (LINEAR LS) $\swarrow$ MODEL PREDICTION

GIVEN MODEL

$$y(i) = \underbrace{\varphi(i)^T}_{\substack{\uparrow \\ \text{REGRESSION VECTORS}}} \cdot \theta_0 + \varepsilon(i)$$

i.i.d.,
ZERO MEAN

$$\hat{\theta}_{LS} = \arg \min_{\theta} \|y_N - \Phi_N \cdot \theta\|_2^2$$

$$\varphi(i) = \begin{bmatrix} * \\ \vdots \\ * \end{bmatrix} \in \mathbb{R}^d$$

$$y_N = \begin{bmatrix} y(1) \\ \vdots \end{bmatrix}, \Phi_N = \begin{bmatrix} \varphi(1)^T \\ \vdots \end{bmatrix}$$

$$\hat{\theta}_{LS} = \phi_N^+ \cdot y_N$$

ϕ_N^+ : PSEUDO-INVERSE =

$$\text{pinv}(\phi)$$

$$\left\{ \begin{array}{ll} (\phi_N^T \phi_N)^{-1} \phi_N^T & \text{IF } \phi_N^T \phi_N \text{ INV } (\Leftrightarrow \text{RANK } \phi_N = d) \\ \text{ELSE} \\ \text{IF } \phi_N = U S V^T & \text{(SVD)} \\ \text{ELSE} \\ V \cdot S^+ \cdot U^T & \end{array} \right.$$

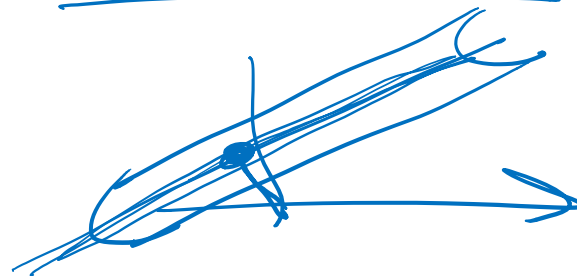
$$S = \begin{pmatrix} \sigma_1 & \dots & \sigma_r & 0 & \dots \\ \hline 0 & \dots & 0 & \dots & 0 \end{pmatrix}$$

$$S^+ = \begin{pmatrix} \sigma_1^{-1} & \dots & \sigma_r^{-1} & 0 & \dots \\ \hline 0 & \dots & 0 & \dots & 0 \end{pmatrix}$$

Ass. $\text{RANK}(\phi_N) < d$

$$\hat{\theta}_{\text{pinv}} = \phi_N^+ \cdot y_N = \arg \min_{\theta} \|\theta\|_2^2 \quad \text{s.t.} \quad \underline{\phi_N^T \phi_N \theta = \phi_N^T y_N}$$

MINIMUM NORM SOLUTION AMONG ALL
OPTIMIZERS OF LS-PROBLEM



WLS : $W \succeq 0$

$$\begin{aligned}\hat{\theta}_{\text{WLS}} &= \arg\min_{\theta} \|\gamma_N - \phi_N \cdot \theta\|_W^2 \\ &= \arg\min_{\theta} \|W^{\frac{1}{2}}(\gamma_N - \phi_N \cdot \theta)\|_2^2 \\ &= \arg\min_{\theta} \|\underbrace{W^{\frac{1}{2}}\gamma_N}_{\tilde{\gamma}_N} - \underbrace{(W^{\frac{1}{2}}\phi_N)\theta}_{\tilde{\phi}_N}\|_2^2 \\ &= (W^{\frac{1}{2}}\phi_N)^+ \cdot W^{\frac{1}{2}}\gamma_N\end{aligned}$$

(W)LS GIVES UNBIASED ESTIMATE IF ε IS ZERO MEAN

OPTIMAL WEIGHTING W : $W^* = \Sigma_{\varepsilon}^{-1}$
 $\downarrow$ LEADS TO SMALLEST $\text{COV}(\hat{\theta}_{\text{WLS}})$
(= ML EST.)

WITH $\Sigma_{\varepsilon} = \text{COV}(\varepsilon)$

FOR i.i.d. noise : LS IS OPT.
(NO WEIGHTING NEE.)

CLOSED BOOK EXAM

(TWO SHEETS WITH 4 PAGES OF HANDWRITTEN OR NOT-TOO
PEN SMALL FONT
(8 pts))

DATE: MARCH 16, 14:00 - ~~16:00~~
17:00

MEASURE OF FIT

R-SQUARE

$$R^2 = 1 - \frac{\| \epsilon \|^2}{\| y \|^2} = \frac{\| \hat{y} \|^2}{\| y \|^2}$$

RES. AFTER FIT RES.

$\in [0, 1]$

1 : PERFECT FIT
0 : BAD FIT

COVARIANCE FROM SINGLE EXP. (LS)

$$\text{cov}(\hat{\theta}) \approx G_{\epsilon}^2 \cdot (\Phi_N^T \Phi_N)^{-1}$$

$$G_{\epsilon}^2 = \frac{\| \epsilon \|^2}{N-d}$$

FOR NLS

$$\Phi_N \approx J = \frac{\partial \pi}{\partial \theta}$$

$$\theta_i = \hat{\theta}_i \pm \sqrt{G_i^2}$$

G_i^2 : DIAGONAL ENTRIES OF $\text{cov}(\hat{\theta})$

ALWAYS : GIVE CONFD. INTERVALS FOR ESTIMATES

MAX. LIKELIHOOD (ML) EST.

GIVEN $\checkmark$ PREDICTION
MODEL

$$M_N(\theta) + \epsilon_N = Y_N$$

• NOISE ASSUMPTIONS ON ϵ_N (ZERO MEAN,)

$$\text{PDF : } p(Y_N | \theta)$$

"LIKELIHOOD FUN"

"NEG. LOG. LIKELIHOOD" $-\log(p(Y_N | \theta))$

ϵ_N GAUSSIAN $\epsilon_N \sim \mathcal{N}(0, \Sigma)$

$$p(Y_N | \theta) = \frac{1}{\sqrt{2\pi}} \exp\left(-\frac{(Y_N - M(\theta))^T \Sigma^{-1} (Y_N - M(\theta))}{2}\right)$$

$$-\log p = \dots + \frac{1}{2} \|(Y_N - M(\theta))\|_{\Sigma^{-1}}^2$$

DAVIESIAN EST. : ADD "A-PRIORI KNOWLEDGE" $p(\theta)$

$$\text{MAP : } \hat{\theta}_{\text{MAP}} = \arg \max_{\theta} p(Y_N | \theta) \cdot p(\theta)$$

$$\text{ML : } \hat{\theta}_{\text{ML}} = \arg \max_{\theta} p(Y_N | \theta) = \arg \min_{\theta} -\log p(Y_N | \theta)$$

NONLIN. LEAST SQ.
(lsqnonlin)
(GAUSS-NEURON)

CRAMER-RAO INEQUALITY: (LOWER BOUND)

$$L(\theta, y) = -\log(y|\theta)$$

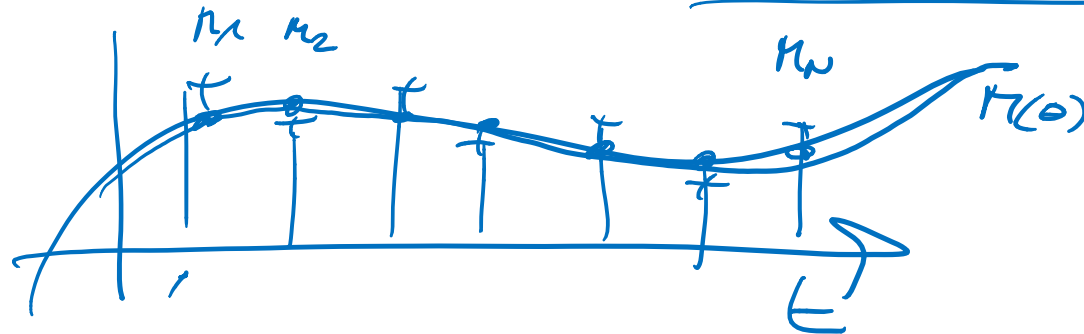
$$E \{ \nabla_{\theta}^2 L(\theta, y) \} =: M = \text{FISHER INF. MATRIX}$$

FOR ANY UNBIASED ESTIM. $\hat{\theta}$:

$$\boxed{\text{COV}(\hat{\theta}) \geq M^{-1}}$$

FOR $N \rightarrow \infty$, ML COV $\rightarrow$ CR. (ML ASYMPTOT. OPT.)

$M(\theta)$ CAN BE "FORWARD SIMULATION MAT."



NOISE: OUTPUT ERRORS
INPUT ERRORS
 $\rightarrow$ OR. ERRORS (LS)

MODELS: L18, ARX

OPE SIMULATION
L11 "