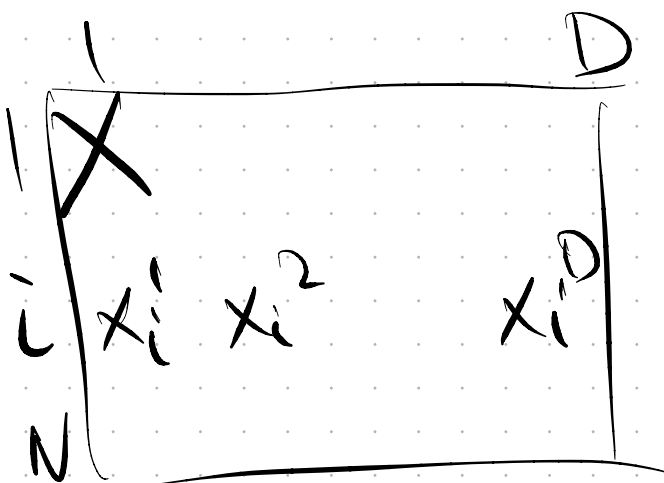


Lecture 6/6 Generative Models

HW 3: GDA, Naive Bayes etc

• Fit curve to data

≈ Fit density to data



want prob. density over data space

$\text{pr}(x_i | \text{param } \theta) = \text{prob that } x_i \text{ datapoint occurs}$

$$\text{pr}_\theta(x_i) = \frac{\# \text{datapoints } x_i \text{ or very similar}}{\# \text{total} = N}$$

curve = parametric with model P for density

ex: curve = Normal = Gaussian

$$\mathcal{N}(x | \mu, \sigma^2) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2}$$

Next Wed 6/11

Demos for HW2B NNets

Project Proposal June 24

Project page

Smaller Project ≈ HW
(20h)

• choices

$x_i = \text{patient}$

$\text{pr}(x_i) = \text{proportions of patients like this}$

Fit curve = Gaussian to data X
 (μ, σ^2)



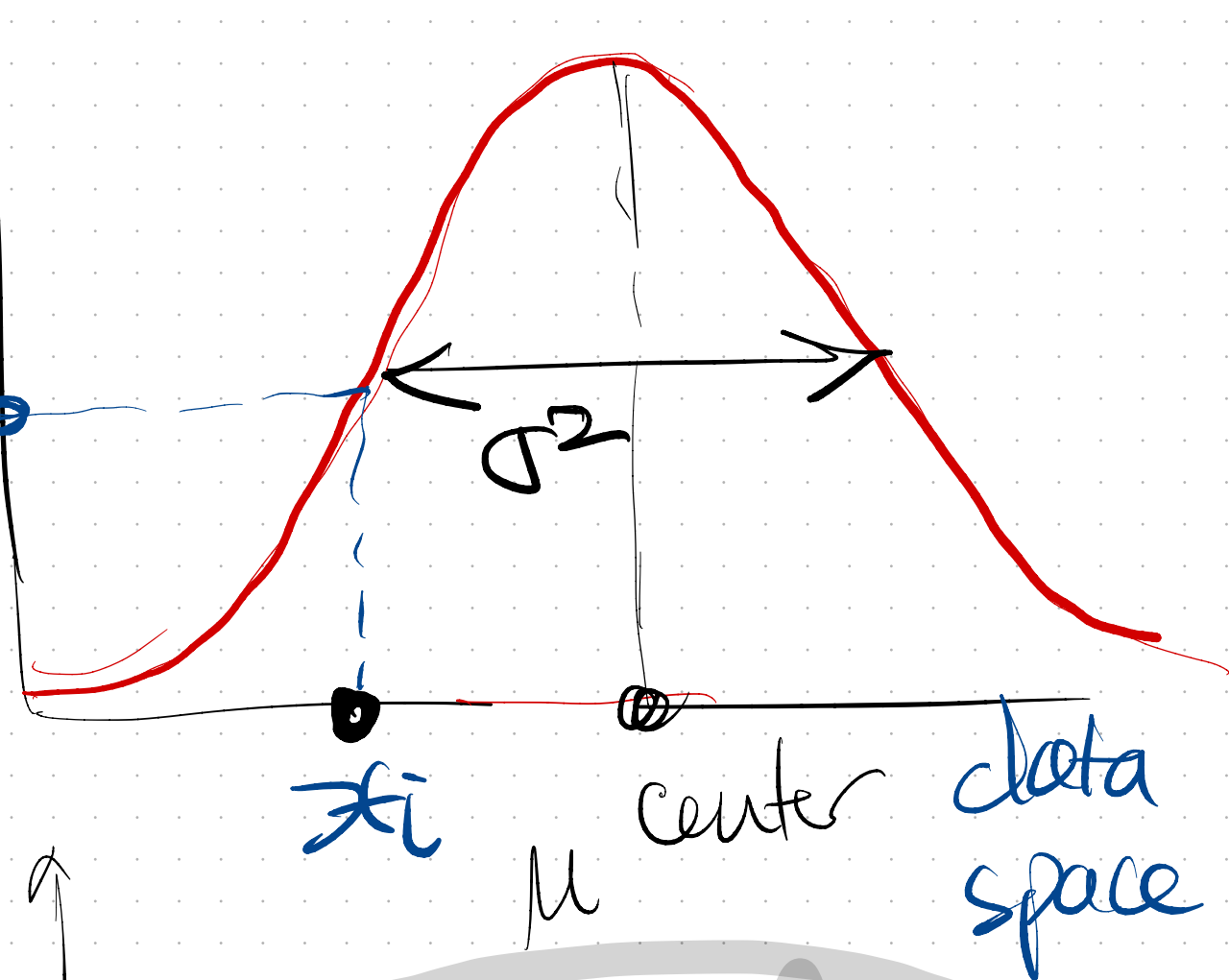
Find the optimal params μ, σ^2
to maximize Likelihood (data X)

product
because
indep. datap

$$\prod_{i=1}^N \text{pr}(x_i)$$

$\text{pr}(x_i) =$

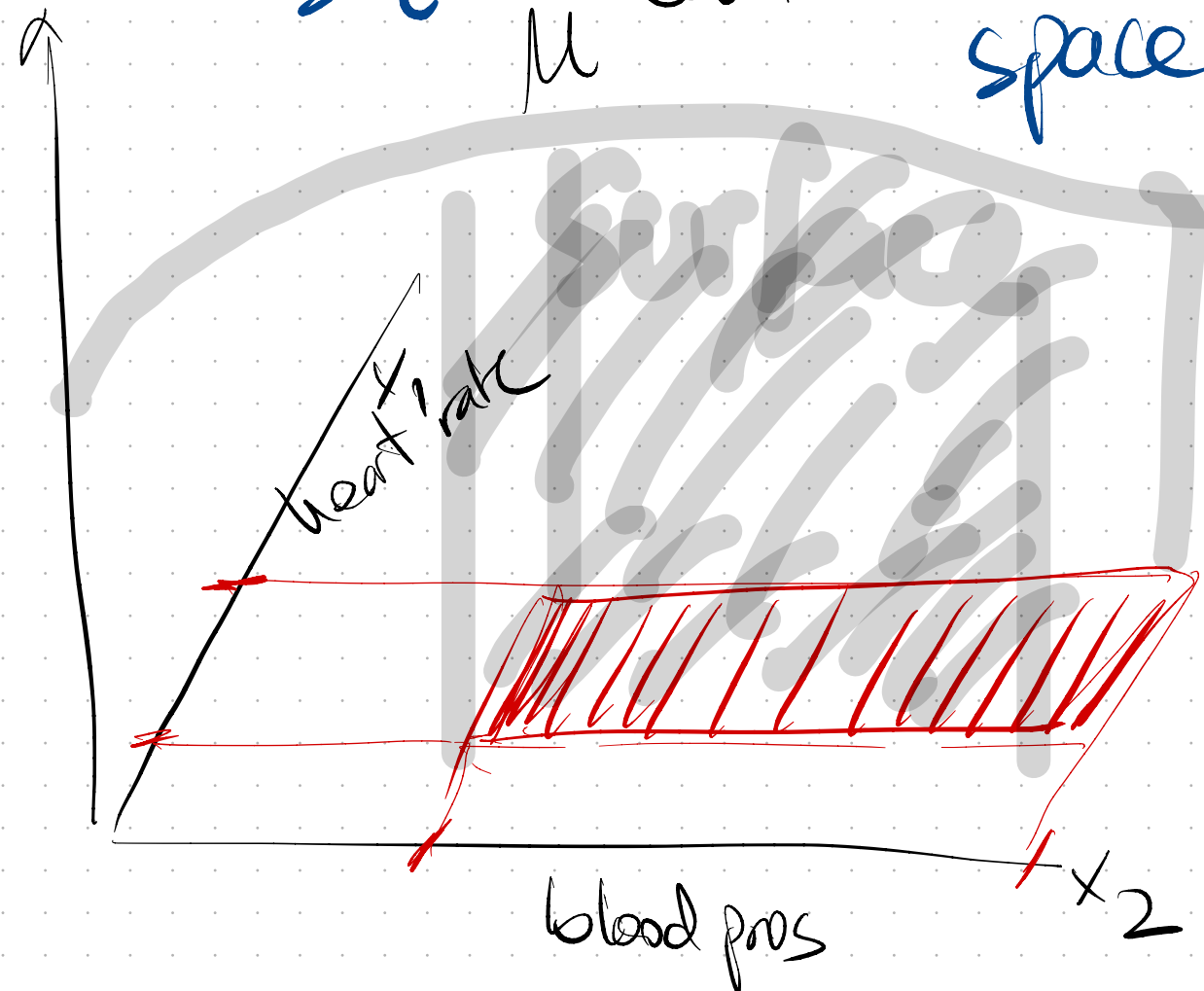
$N(x_i)$



Generative formulation:

Find curve = Gaussian (μ, Σ)

curve = generator of data $\Rightarrow$ data X
is most
likely.



1-dim Gaussian: $P(x | \mu, \sigma^2) = \frac{1}{\sqrt{2\pi}} \cdot e^{-1/2 (x - \mu)^2 / \sigma^2}$

$\sigma^2 = \text{VAR}$ (indicated by a red arrow pointing to σ^2)

$\text{std} = \sqrt{\text{VAR}}$ (indicated by a red arrow pointing to σ)

data center most likely x (indicated by a red arrow pointing to μ)

$X_i = D\text{-dim} \Rightarrow W \left(\begin{matrix} x \\ D\text{-dim} \end{matrix} \middle| \begin{matrix} \mu \\ D\text{-dim} \\ \text{vector} \end{matrix}, \begin{matrix} \Sigma \\ \text{sigma (capital)} \\ D \times D \text{ COVAR}() \end{matrix} \right)$

$$pr(x_i | \mu, \Sigma) = \frac{1}{(2\pi)^{-D/2}} |\Sigma|^{-1/2} \exp \left\{ -\frac{1}{2} (x_i - \mu)^T \Sigma^{-1} (x_i - \mu) \right\}$$

Σ = covariance matrix
 Σ^{-1} = inverse of Σ
 Σ = D x D COVAR()
 Σ^{-1} = D x D
 $(x_i - \mu)^T$ = 1 x D
 $(x_i - \mu)$ = D x 1
 $\Sigma^{-1} (x_i - \mu)$ = D x 1
 $(x_i - \mu)^T \Sigma^{-1} (x_i - \mu)$ = scalar
 $\exp$ = exponential function
 $|\Sigma|$ = determinant of Σ
 $(2\pi)^{-D/2}$ = normalization constant
 $\sum_{i=1}^N$ = summation

Later we show

after we show

x_i

N

$\mu =$

μ^1 μ^2 μ^D

arg of each
wt = feature

x_i^1 x_i^2 x_i^D

μ D

$x^T x = \sum \text{COVAR} = D \times D$

$\mu = 0$

$$\left[\begin{array}{cc} \langle \omega_l^1, \omega_l^1 \rangle & \langle \omega_l^1, \omega_l^0 \rangle \\ \langle \omega_l^0, \omega_l^1 \rangle & \langle \omega_l^0, \omega_l^0 \rangle \end{array} \right]$$

• Fit D-Dim Gaussian to data X (D-features)
 $\Rightarrow \mu = \text{mean}(X)$ D-dim vector $(\mu^1 \mu^2 \dots \mu^D)$ avg for each feat

$$\Rightarrow \Sigma = \text{COVAR}(X) = (X - \mu)^T (X - \mu)$$

(μ, Σ) param that makes generator-Gaussian best for data X
 $\downarrow$
when used to generate
data $\Rightarrow$ max-chance to
obtain $X_{1:N}$

Data likelihood (as generated by Gaussian with param, μ, Σ)

$$L(x) = \prod_{i=1}^N \text{pr}(x_i) = \prod_{i=1}^N \mathcal{N}(x_i | \mu, \Sigma_i)$$

all datap
 $i=1 \dots N$

$$\prod_{i=1}^N \left[(2\pi)^{-d/2} |\Sigma_i|^{-1/2} \exp \left\{ -\frac{1}{2} (x_i - \mu)^T \Sigma_i^{-1} (x_i - \mu) \right\} \right]$$

$$\text{log likelihood} = \text{LL}(x) = \log \left[\prod_{i=1}^N \mathcal{N}(x_i | \mu, \Sigma_i) \right]$$

$$= \sum_{i=1}^N \left[\log(2\pi)^{-d/2} + \frac{1}{2} \log |\Sigma_i^{-1}| - \frac{1}{2} (x_i - \mu)^T \Sigma_i^{-1} (x_i - \mu) \right]$$

$$\text{want } \max \text{LL}(x) \Rightarrow \frac{\partial \text{LL}(x)}{\partial \mu} ; \quad \frac{\partial \text{LL}(x)}{\partial \Sigma_i^{-1}}$$

μ, Σ_i
are indep.
(no constraints)

$$\frac{\partial \log L}{\partial \mu} = \sum_{i=1}^N \cancel{2} \cancel{\frac{1}{2}} \cancel{\Sigma^{-1}} (x_i - \mu) \stackrel{\text{want}}{=} 0 \Rightarrow \sum_{i=1}^N (x_i - \mu) = 0$$

$$\Rightarrow \boxed{\mu = \frac{1}{N} \sum_{i=1}^N x_i}$$

$$\frac{\partial \log L}{\partial \Sigma^{-1}} = \sum_{i=1}^N \left[\cancel{\frac{1}{2}} \cancel{\Sigma}^{-1} - \frac{1}{2} (x_i - \mu)^T (x_i - \mu) \right] \stackrel{\text{want}}{=} 0$$

$$\sum_{i=1}^N \frac{(x_i - \mu)^T (x_i - \mu)}{(x_i - \mu)^2} \dots \sum_{i=1}^N (x_i - \mu)^T (x_i - \mu^T)$$

$$N \cdot \Sigma = \sum_{i=1}^N (x_i - \mu)^T (x_i - \mu)$$

COVAR

$$\Sigma = \frac{1}{N} \sum_{i=1}^N \underbrace{(x_i - \mu)^T}_{D \times 1} \underbrace{(x_i - \mu)}_{1 \times D}$$

sum of N

$$\left[\begin{matrix} D \times D \\ \text{mat.} \end{matrix} \right]$$

$$\sum_{i=1}^N (x_i^h - \mu^h) (x_i^T - \mu^T)$$

$$D \sum_{i=1}^N (x_i^D - \mu^D) (x_i^1 - \mu^1) \dots \sum_{i=1}^N (x_i^D - \mu^D)^2$$

$$\mu=0 \Leftrightarrow X \text{ centered} = X - \mu$$

$$\text{COVAR } \Sigma = \frac{1}{N} X^T X$$

Classification with gen. models (Gaussian fit to data)

want $P(Y|X)$ ← label

$$P(Y|X) = \frac{P(X|Y) \cdot P(Y)}{P(X)}$$

→ prior for each $y = \text{class } Y_1, Y_2, \dots, Y_L$
prior dist over $Y_1, \dots, Y_L$

→ normalization of numerators
For L classes

density of x in certain
class $y = y_c$
fit curve to $y = 1$ points

y binary $\in \{1, 0\}$
 $x = \text{test point}$

$$P(Y=1|X) = \frac{P(X|Y=1) \cdot P(Y=1)}{P(X)}$$

$$P(Y=0|X) = \frac{P(X|Y=0) \cdot P(Y=0)}{P(X)}$$

fit curve to $y=0$ points

ignore

GDA = Gauss Discriminant Analysis

- $\mathcal{N}_1(\underline{\mu}_1, \underline{\Sigma}_1)$ fit on X data with $y = +1$
subset

- separate
 $\mathcal{N}_0(\underline{\mu}_0, \underline{\Sigma}_0)$ fit on X data
subset $y = 0$

- predict: compute
test point z

- predict $y = \arg\max$

$$P(Y=+1|z) = \mathcal{N}_1(z|\underline{\mu}_1, \underline{\Sigma}_1) \cdot P(Y=+1) / P(z)$$

$$P(Y=0|z) = \mathcal{N}_0(z|\underline{\mu}_0, \underline{\Sigma}_0) \cdot P(Y=0) / P(z)$$

ignore

- can do this for $y=1, y=2, \dots, y=L$ (L classes)

- can use common Σ for all classes.