

Today

- Admin
- Gaussians
- LDA
- Linear regression
- Pick up assignment

1. Assignment 1 **DUE NOW!**
(or Monday for 25% off)

2. No tutorial this week
(Ass 2 out this weekend)

(Tutorials are optional,
not excuse to leave classes early)

3. Piazza (closing soon, find friends)

4. End of class → pick up assignment to mark

DUE NOW!
25% off)

this week
(weekend)

1,
2 classes early)

find friends)

assignment to mark

Notation

I used:

D : data $\{(x_i, y_i)\}_{i=1}^N$

D_i : data sample (x_i, y_i)

D : dimension

θ : parameter of Bernoulli

$\mathcal{H}$: hypothesis for MLE $h \in \mathcal{H}$

X, x_i : design matrix, features of example i

Y, y_i : target vector, label of example i

? Should we use:

DS

DS_i

p

Naive

Issue

- div

- instr

pres

Naive Bayes

Issues:

- div by 0
- instructor badly presenting MLE (sorry)

Advantages

- simple
- fast (single pass) training
- parallel
- small # params
- not much data needed

Disadvantages

- strong mutual independence assumption.
 - ↳ limited modeling power

Course project idea:

- explore different ways to relax NB assumption

(TAN-Bayes, Bayesian Network Classifiers

Today:

- What if we have

$$x_i \in \mathbb{R}^D?$$

an Network Classifiers)

Gaussian Distribution

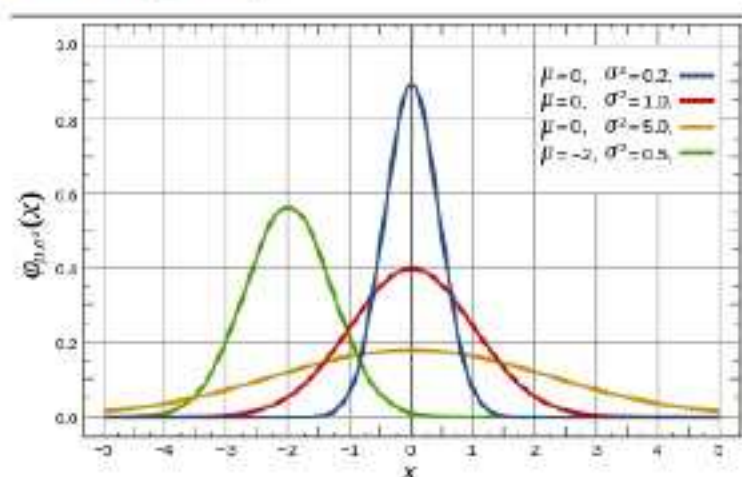
If x is distributed according to a Gaussian/normal distribution with mean μ and variance σ , we write this as

$$x \sim \mathcal{N}(\mu, \sigma^2),$$

and the probability density function (PDF) is

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

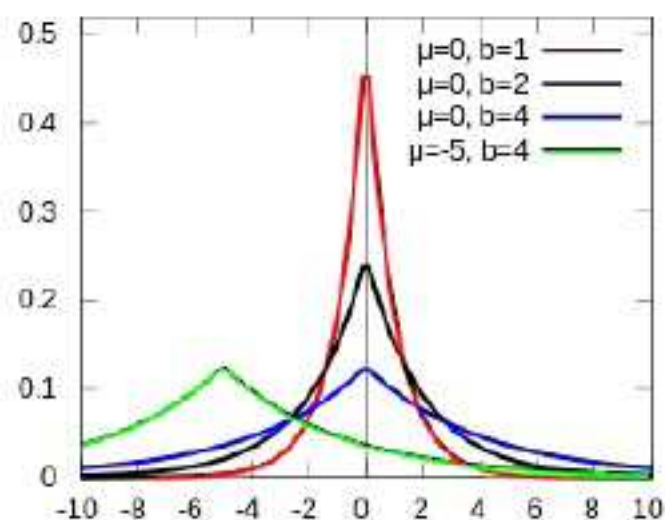
Here are some example PDFs from Wikipedia,



An alternative distribution is the Laplace distribution where we have

$$p(x|\mu, \sigma) = \frac{1}{2\sigma} \exp\left(-\frac{|x - \mu|}{\sigma}\right),$$

It's Wikipedia graphic looks like this:



The Laplace distribution has 'heavier-tails' than the Gaussian distribution, meaning that it assigns a higher probability to events that are far away from μ (while the Gaussian assumes everything is close to μ). A distribution with even heavier tails is the student's t distribution.

Motivation

- Central limit theorem.

↳ MLE

- data is Gaussian

- Analytic properties

- $(-\infty, \infty)$

↳ sampling

↳ symmetry

MLE

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

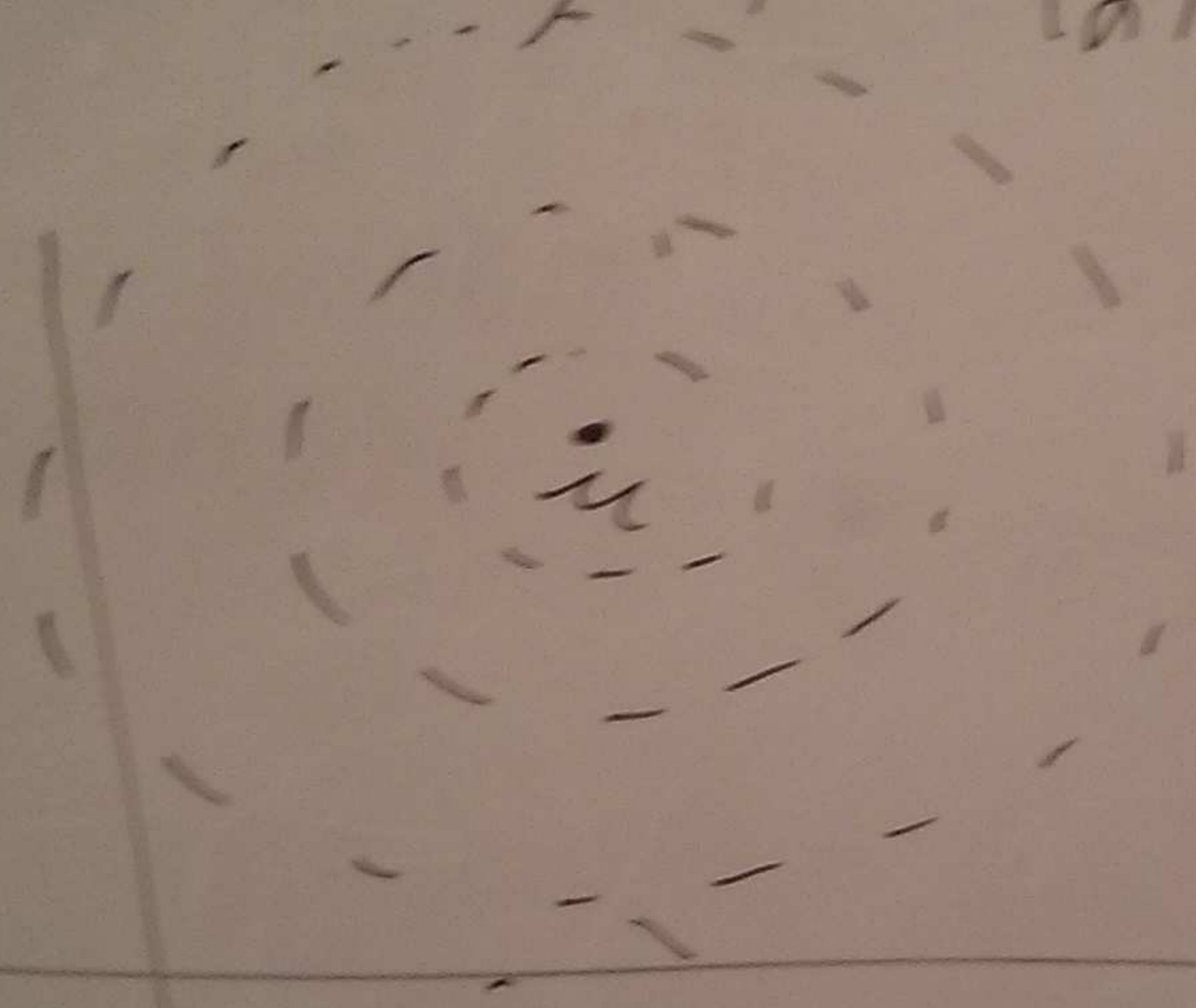
$$\hat{\sigma}^2 = \frac{1}{N} \sum_{i=1}^N (x_i - \hat{\mu})^2$$

$p(x^N | \mu, \sigma^2)$

$$\bar{x} \sim N(\mu, \Sigma)$$

$$\int_{-\infty}^{\infty} p(x | \mu, \Sigma) = 1$$

$$p(\bar{x} | \mu, \Sigma) = \frac{1}{(2\pi)^{d/2} |\Sigma|^{1/2}} \exp\left(-\frac{1}{2}(\bar{x} - \mu)^T \Sigma^{-1}(\bar{x} - \mu)\right)$$



$$\mu \in \mathbb{R}^d$$

$$\Sigma \in \mathbb{S}_{++}^d$$

$$\Sigma \succ 0$$

$$(x_i - \hat{\mu})^2$$

$$(x^i | y_i)$$

$$p(\bar{x}_i | y_i)$$

$$\bar{x}_i | y_i$$

$$\hat{\mu} =$$

$$\hat{\mu} =$$

$$\int_{-\infty}^{\infty} p(x|m, \Sigma) = 1$$

$$\frac{1}{(2\pi)^{d/2} |\Sigma|^{1/2}} \exp\left(-\frac{1}{2}(x-m)^T \Sigma^{-1}(x-m)\right)$$

$$m \in \mathbb{R}^d$$

$$\Sigma \in \mathbb{S}_{++}^d$$

$$\Sigma \succ 0$$

$$p(\bar{x}_i | y_i)$$

$$\bar{x}_i | y_i \sim \mathcal{N}(m, \Sigma)$$

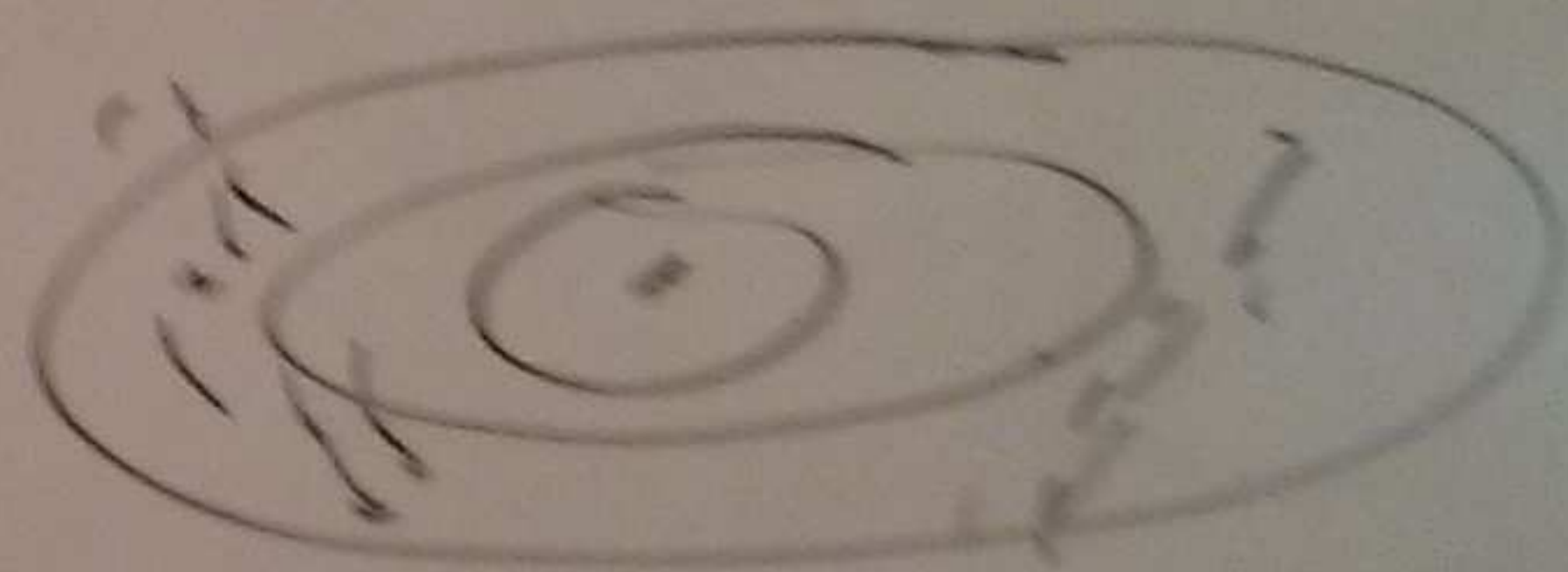
$$\hat{m} = \frac{1}{N} \sum_{i=1}^N \bar{x}_i$$

$$\hat{\Sigma} = \frac{1}{N} \sum_{i=1}^N (\bar{x}_i - \hat{m})(\bar{x}_i - \hat{m})^T \in \mathbb{S}_{++}^d$$

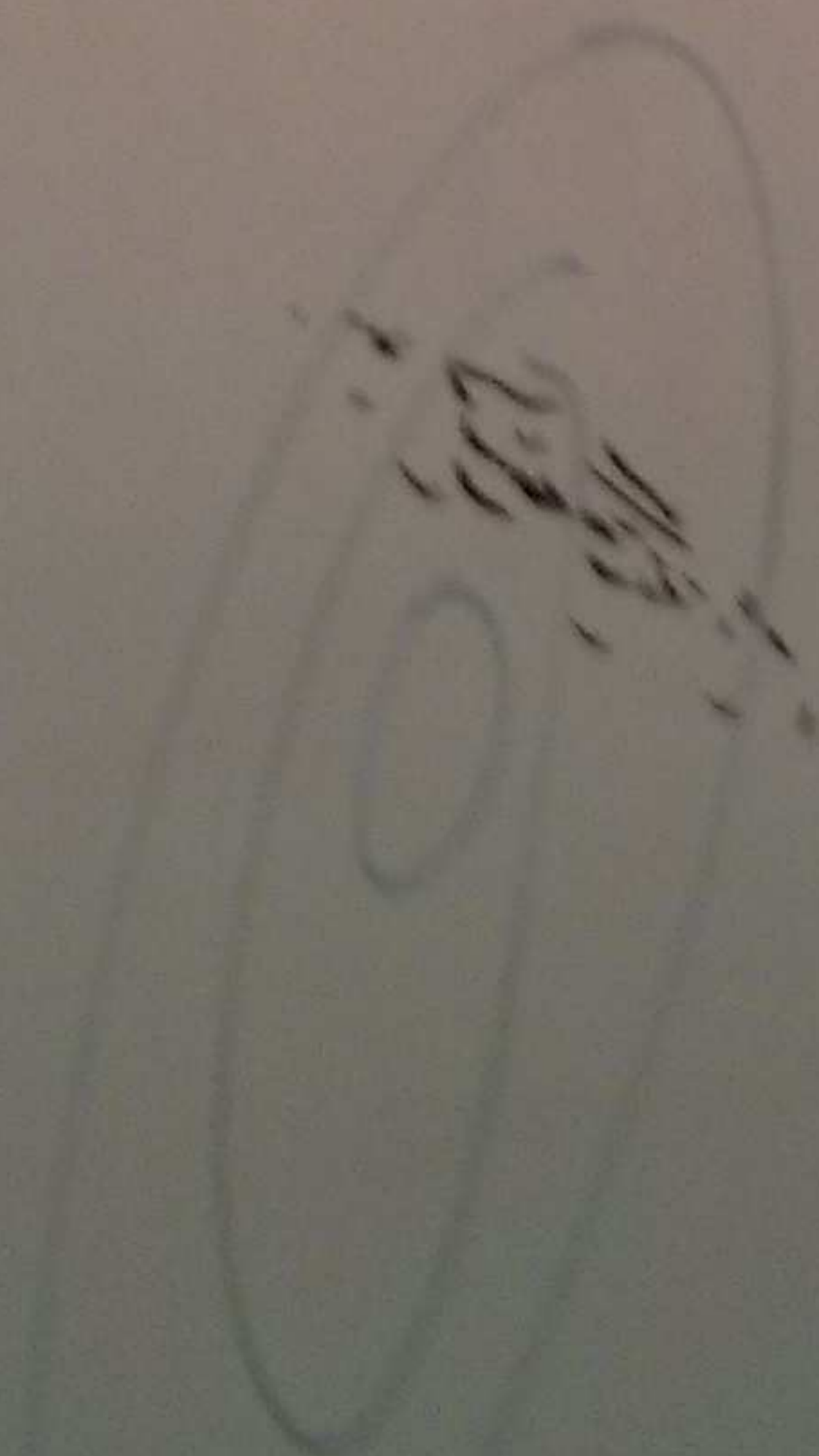
Problems:

1. $\hat{\Sigma}$ might not be valid

2. Unimodal



3. not "robust"



Today

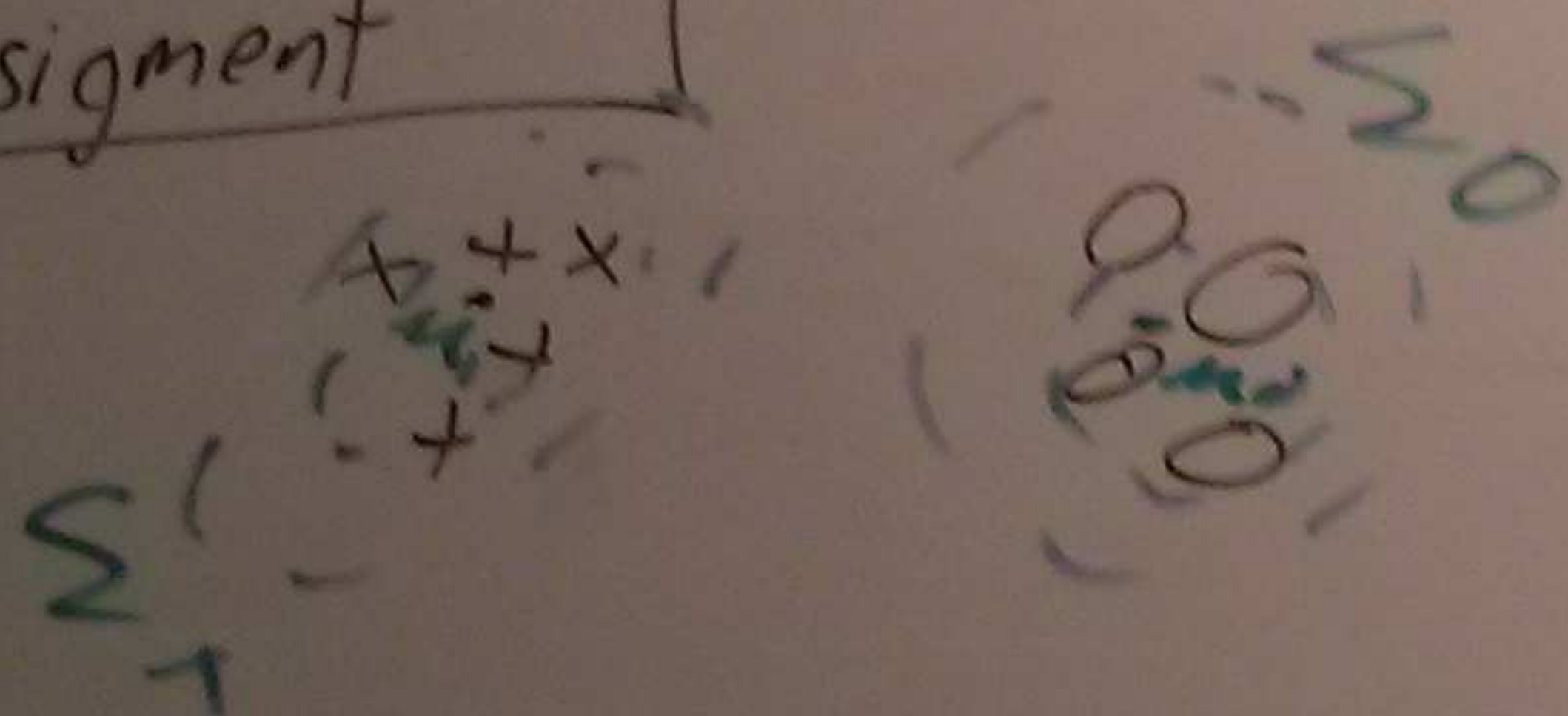
- Admin
- Gaussians
- LDA
- Linear regression
- Pick up assignment

Gaussian Discriminant Analysis

$$p(\bar{x}_i | y_i) = p(\bar{x}_i | y_i) p(y_i)$$

$$\bar{x}_i \in \mathbb{R}^d$$

$$y_i \in \{1, 2, 3, 4\} \quad \bar{x}_i | y_i = c \sim \mathcal{N}(\mu_c)$$



$$\begin{array}{cc} \mu_1 & \Sigma_1 \\ \mu_2 & \Sigma_2 \\ \mu_3 & \Sigma_3 \\ \mu_4 & \Sigma_4 \end{array}$$

$$p(y_i | \bar{x}_i) = \frac{p(\bar{x}_i | y_i) p(y_i)}{p(\bar{x}_i)} \propto p(\bar{x}_i | y_i) p(y_i)$$

$$p(\bar{x}_i) = \sum_{c=1}^C p(\bar{x}_i | y_i = c) p(y_i = c)$$

$p(\bar{x}_i | y_i = c)$

analysis

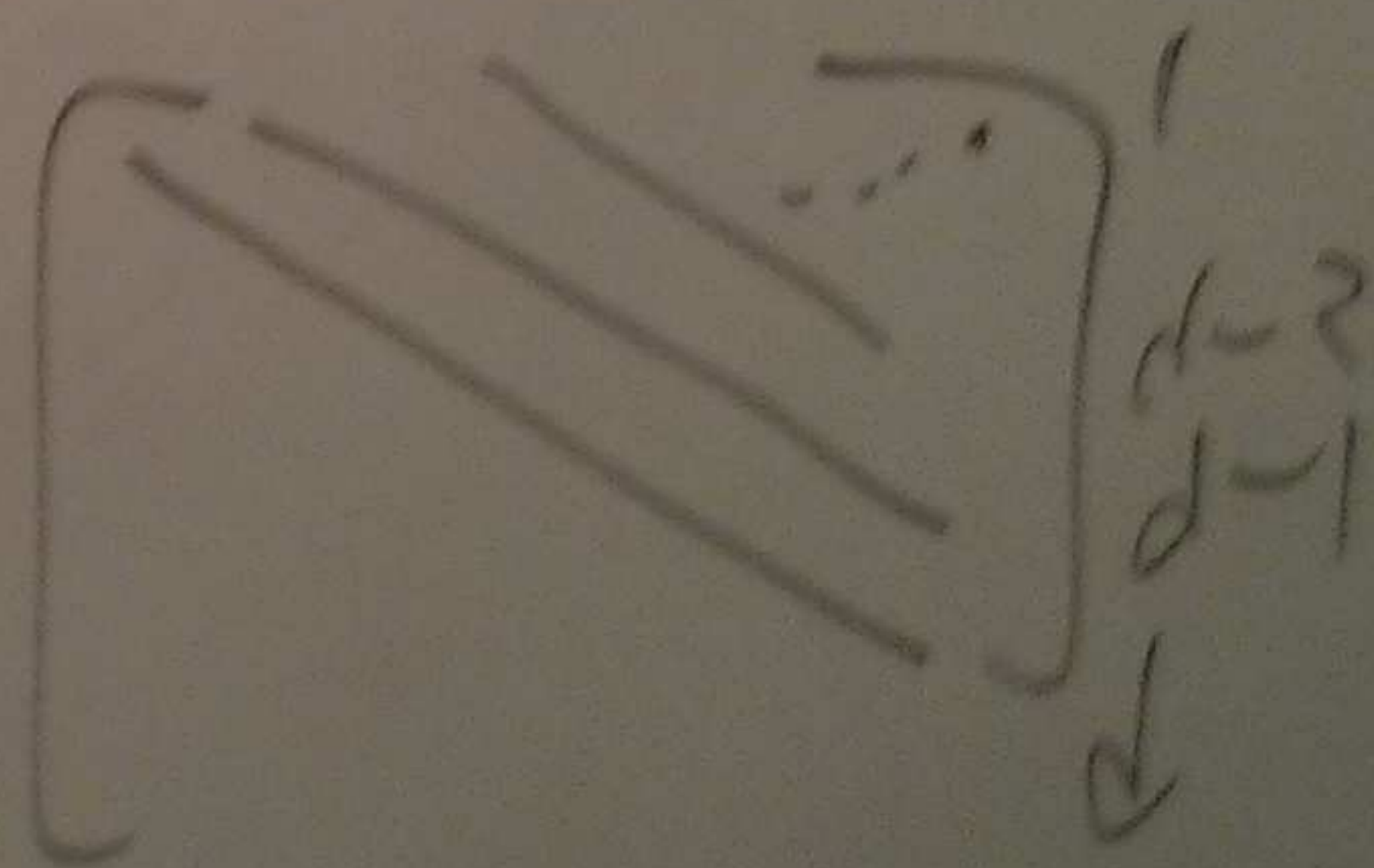
(y_i)

(u_d)

$\sum_c$

d # parameters

$$\underbrace{cd}_n + \underbrace{c(d+1)d}_2 = O(\underline{cd^2})$$



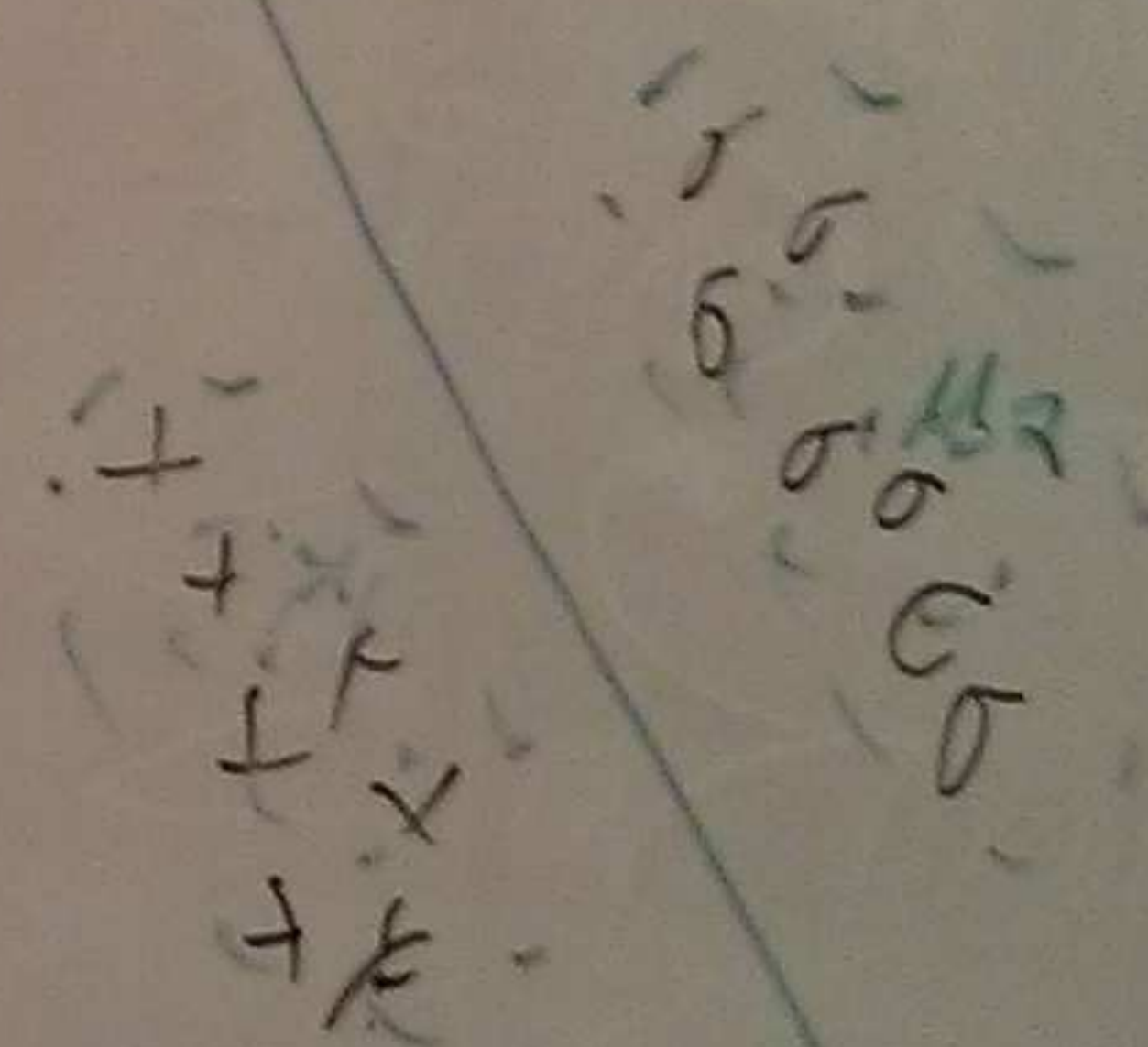
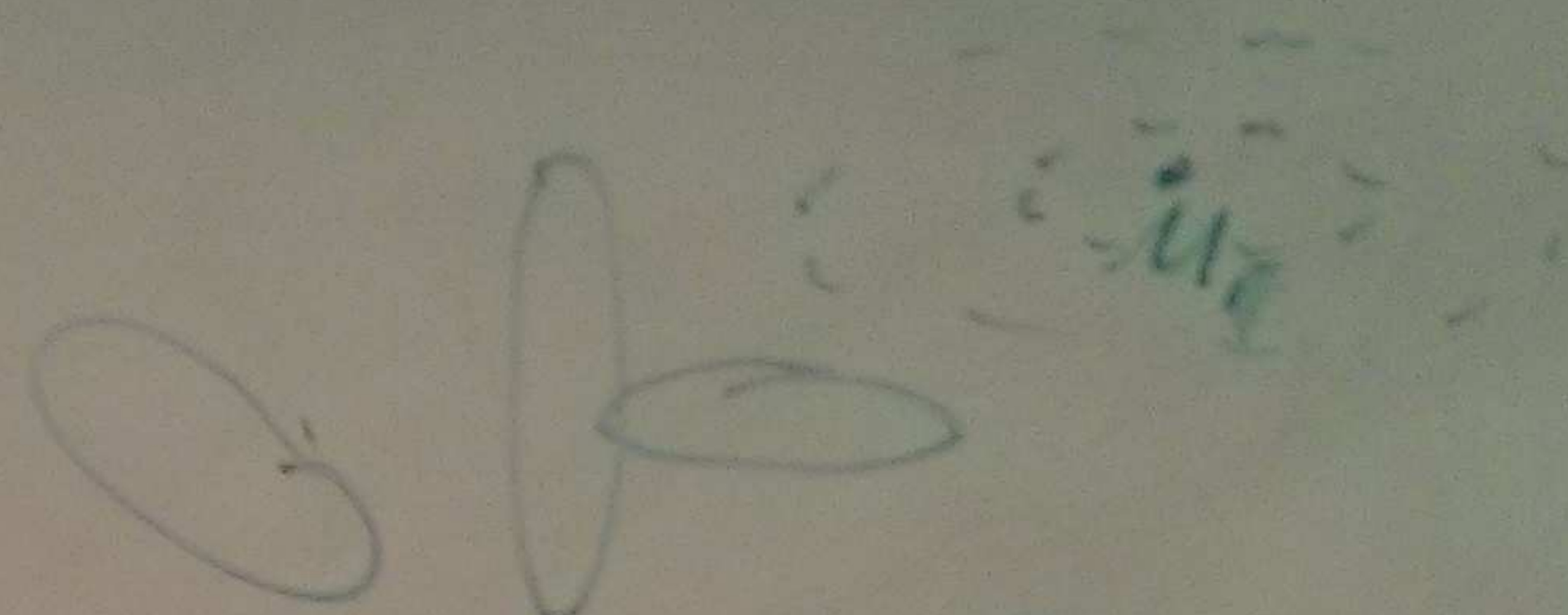
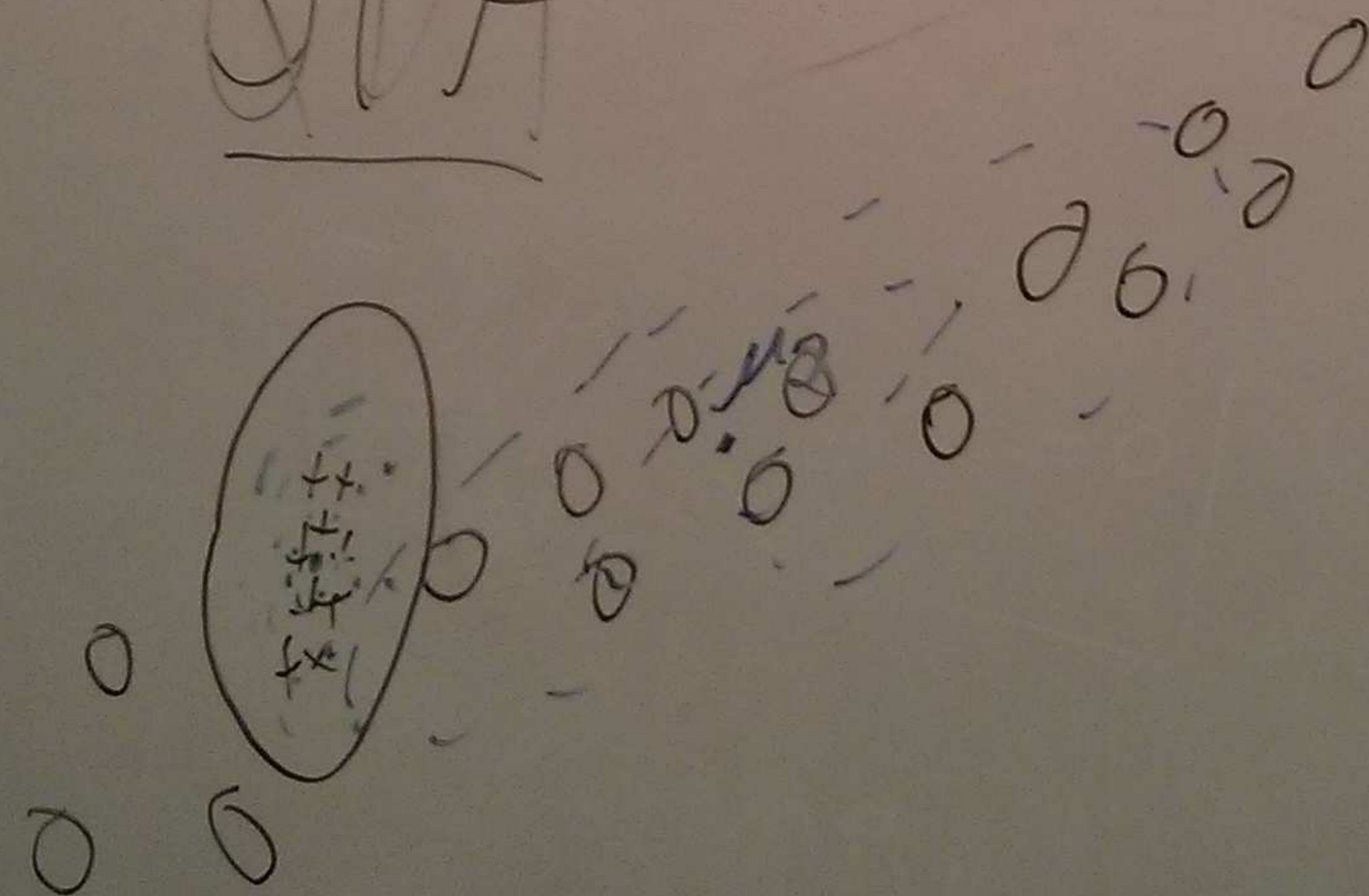
$O(c_d^2)$

Naive Bayes: Σ_c diagonal

Linear Discrim Analysis: $\Sigma_c = \Sigma, \forall c$

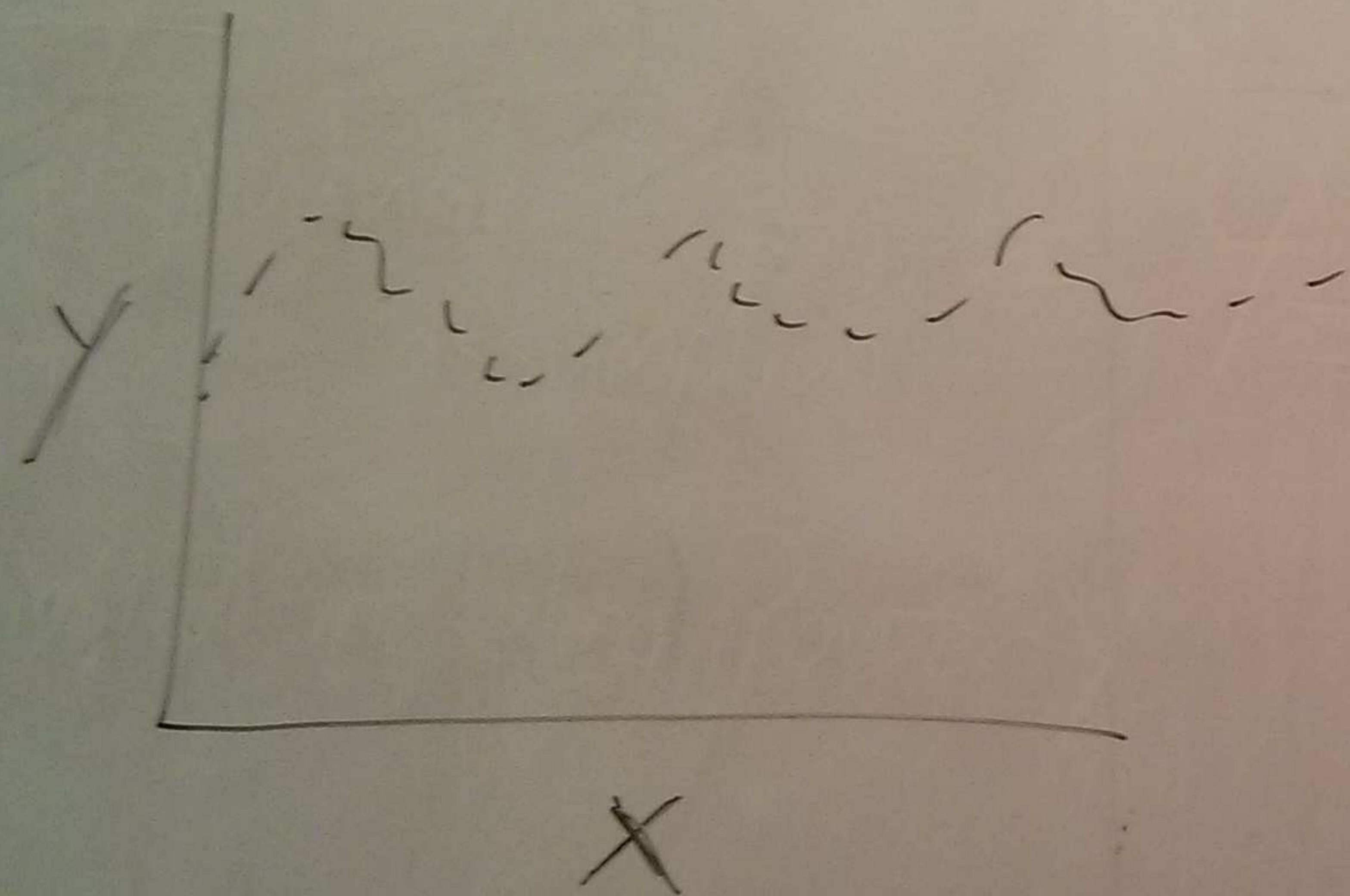
Quad Disc. Anal: general Σ_c

QDA:



Gaussian Regression

$$y_i = w_1 x_i^1 + w_2 x_i^2 + \dots$$



"linear model"

"linear regression"

$$y_i = \bar{w}^T \bar{x}_i + \epsilon_i$$

$$\epsilon_i \sim N(0, 1)$$

$$y_i | \bar{x}_i, w \sim N(w^T \bar{x}_i, 1)$$

$$\bar{X} = \begin{bmatrix} \text{---} \\ \vdots \end{bmatrix} \quad \bar{Y} = \begin{bmatrix} y_1 \\ \vdots \end{bmatrix}$$

$$w_{MLE} = \underset{\bar{w} \in \mathbb{R}^d}{\operatorname{argmin}} \sum_{i=1}^N \log p(y_i | \bar{x}_i, \bar{w})$$

$$= \underset{\bar{w} \in \mathbb{R}^d}{\operatorname{argmin}} \frac{1}{2} \sum_{i=1}^N (y_i - \bar{w}^T \bar{x}_i)^2$$

"Least Squares"

$$X^T w = \begin{bmatrix} \bar{w}^T \bar{x}_1 \\ \bar{w}^T \bar{x}_2 \\ \vdots \\ \bar{w}^T \bar{x}_N \end{bmatrix}$$

$$NLL(w) = \arg \min_{\bar{w} \in \mathbb{R}^d} \frac{1}{2} (Y - X\bar{w})^T (Y - X\bar{w})$$

+ cons.

scalar

$$NLL(w) = \frac{1}{2} Y^T Y - w^T X^T Y + \frac{1}{2} \bar{w}^T X^T X w$$

d-dimensional
vector

$$\nabla NLL(w) = -X^T Y + X^T X w$$

stationary

$$= 0$$

d-dimensional
matrix

$$X^T X w = X^T Y$$

$$\nabla^2 NLL(w) = X^T X$$