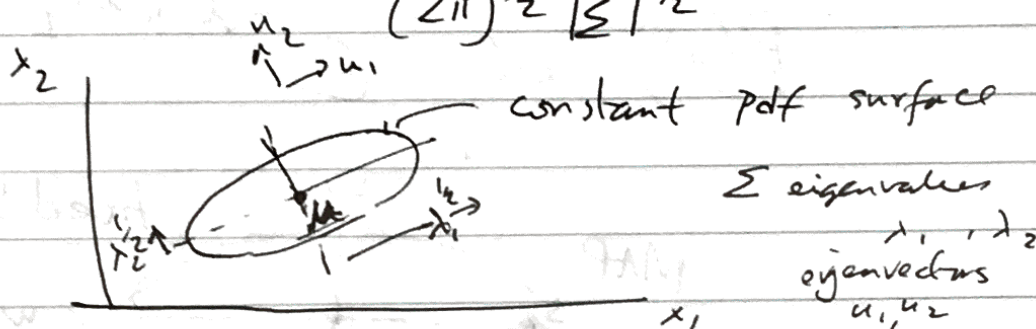


Gaussian Models & Classification

Multivariate Gaussian or multivariate normal (MVN) models.

$$\mathcal{N}(x|\mu, \Sigma) \triangleq \frac{1}{(2\pi)^{\frac{D}{2}} |\Sigma|^{\frac{1}{2}}} e^{-\frac{1}{2}(x-\mu)^T \Sigma^{-1}(x-\mu)}$$



If $\Sigma u_i = \lambda_i u_i$ $\Sigma^{-1} = \sum_i \frac{1}{\lambda_i} u_i u_i^T$

Define $y_i = u_i^T (x - \mu)$

$$(x-\mu)^T \Sigma^{-1} (x-\mu) = \sum_i \frac{y_i^2}{\lambda_i}$$

Mahalanobis distance.

In two-dim $\frac{y_1^2}{\lambda_1} + \frac{y_2^2}{\lambda_2} = \text{constant}$

→ Ellipses.

MLE for MVN

(lookup deriv in the back!)

$$\begin{aligned} \hat{\mu} &= \frac{1}{N} \sum_{i=1}^N x_i = \bar{x} \\ \hat{\Sigma} &= \frac{1}{N} \sum_{i=1}^N (x_i - \bar{x})(x_i - \bar{x})^T \\ &= \frac{1}{N} \sum_{i=1}^N x_i x_i^T - \bar{x} \bar{x}^T \end{aligned}$$

Gaussian discriminant analysis

$$P(X|y=c, \theta) = \mathcal{N}(X|\mu_c, \Sigma_c)$$

We will use this model for several supervised (in this lecture) and unsupervised tasks (in later ones).



Imagine we can estimate the parameters (say with labeled data)

Classifier:

$$\hat{y}(x) = \arg \max_c \left[\overset{\text{prior prob of class } c}{\log(\pi_c)} + \log P(x|\theta_c) \right]$$

$$\text{This is just } \hat{y}(x) = \arg \max_c \left\{ (x - \mu_c)^T \Sigma_c^{-1} (x - \mu_c) + \log |\Sigma_c| \right\}$$

for uniform π

The first term is the Mahalanobis distance from each centroid.

If Σ_c 's are different for each class then classes are decided by thresholding some

quadratic function $\Rightarrow$ Quadratic Discriminant analysis (QDA)

$$P(y=c) = \frac{\pi_c |\Sigma_c|^{-D/2} e^{-\frac{1}{2}(x-\mu_c)^T \Sigma_c^{-1} (x-\mu_c)}}{\sum_c \pi_c |\Sigma_c|^{-D/2} e^{-\frac{1}{2}(x-\mu_c)^T \Sigma_c^{-1} (x-\mu_c)}}$$

Life simplifies if we assume all $\Sigma_c = \Sigma$, namely, the covariance matrices are tied.

Then we just need to compare $\log \pi_c - \frac{1}{2} (x-\mu_c)^T \Sigma^{-1} (x-\mu_c)$ for different classes [1/2 part is the same]

$$\begin{aligned} & (x-\mu_c)^T \Sigma^{-1} (x-\mu_c) \\ &= x^T \Sigma^{-1} x - 2 \mu_c^T \Sigma^{-1} x + \mu_c^T \Sigma^{-1} \mu_c + b \end{aligned}$$

↑
independent of class

Call $\gamma_c = \log \pi_c - \frac{1}{2} \mu_c^T \Sigma^{-1} \mu_c$

$$\beta_c = \Sigma^{-1} \mu_c$$

$$P(y=c | x, \theta) = \frac{e^{\beta_c^T x + \gamma_c}}{\sum_c e^{\beta_c^T x + \gamma_c}} = S(\eta_c)$$

$$\eta = [\beta_1^T x + \gamma_1, \dots; \beta_C^T x + \gamma_C]$$

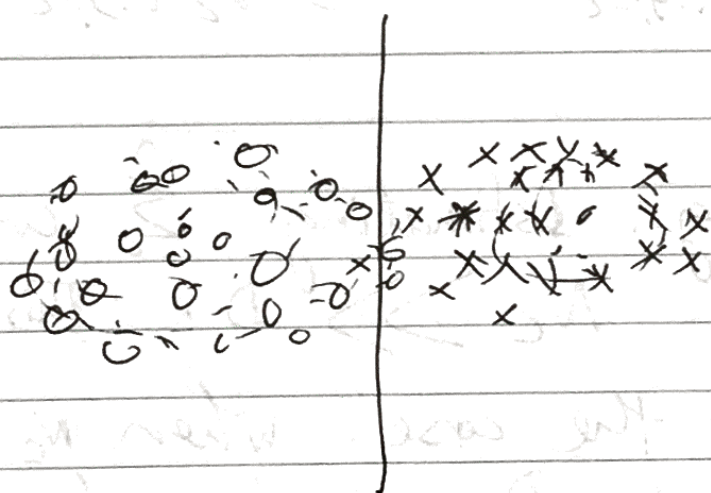
$$S(\eta)_c = \frac{e^{\eta_c}}{\sum_{c'} e^{\eta_{c'}}}$$

is the softmax function

If just have two classes

$C = 1, 2$

$$P(c=1|x, \theta) = \frac{e^{\beta_1^T x + \gamma_1}}{e^{\beta_1^T x + \gamma_1} + e^{\beta_2^T x + \gamma_2}} = \frac{e^{(\beta_1 - \beta_2)^T x + \gamma_1 - \gamma_2}}{1 + e^{(\beta_1 - \beta_2)^T x + \gamma_1 - \gamma_2}}$$



$$\Delta \beta^T x + \Delta \gamma > 0$$

$$< 0$$

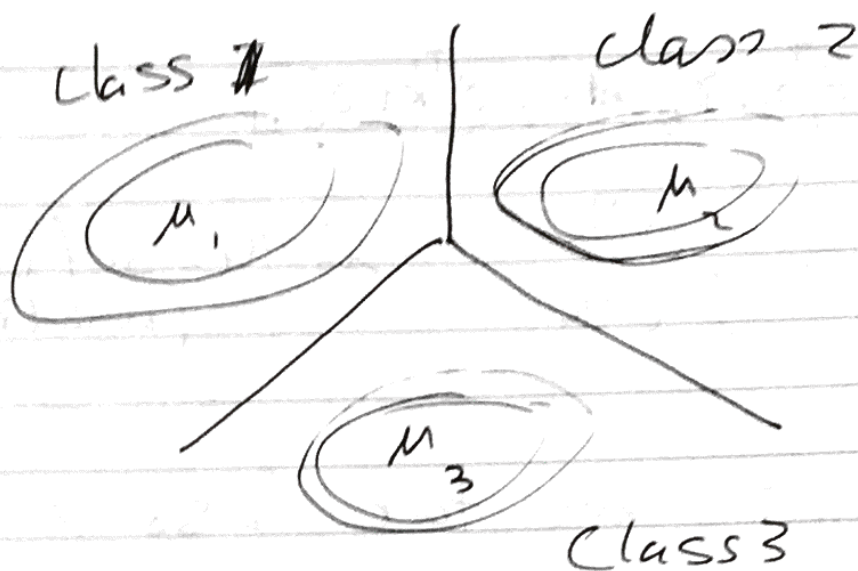
decides
class

$$\Delta \beta^T x + \Delta \gamma = 0$$

is the decision
boundary.

Linear ~~discriminant~~ discriminant analysis (LDA)

Even in the multiclass
case, we have



For MLE parameter fit with supervised learning for full-fledged gaussian model.

$$\hat{\mu}_c = \frac{1}{N_c} \sum_{i: y_i = c} x_i$$

$$\hat{\Sigma}_c = \frac{1}{N_c} \sum_{i: y_i = c} (x_i - \hat{\mu}_c)(x_i - \hat{\mu}_c)^T$$

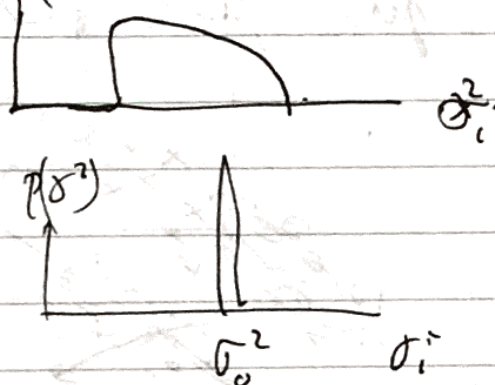
Overfilling: Estimating $\hat{\Sigma}_c$ well requires $N_c \gg D$. This is not always the case. When N_c is comparable to D we do not estimate $\hat{\Sigma}_c$ well.

For example, if I give you $x_i \sim \mathcal{N}(0, \sigma^2 I_D)$, and you estimate $i=1, \dots, N$, x_i independent

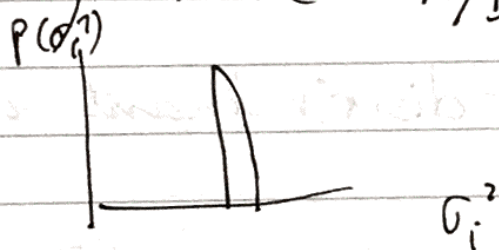
$$\hat{\Sigma} = \frac{1}{N} \sum_i (x_i - \bar{x})(x_i - \bar{x})^T$$

$\hat{\Sigma}$ would have eigenvalues distributed according to the ~~Mar~~ Marchenko Pastur distribution for large N, D with N/D fixed

$$\text{True } \Sigma = \sigma_0^2 I_D$$



Only when $N/D \rightarrow \infty$ do you get

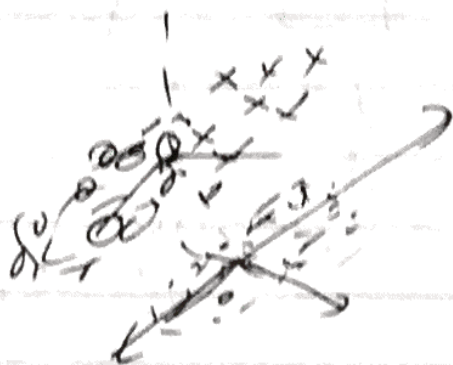


One more we could be fitting noise if $N \sim D$.

- Various solutions:
- a) Parameter tying : $\Sigma_c = \Sigma$ (accept bias reduce variance)
 - b) Insist on diagonal Σ (like naive Bayes)
 - c) Integrate over Σ .
 - d) Project data to lower dim spaces.

The last idea: projecting to the maximally discriminative subspace, is an important topic in itself.

One ~~idea~~ approach could be to just find PCA components before classifying and keep a small number of principal components

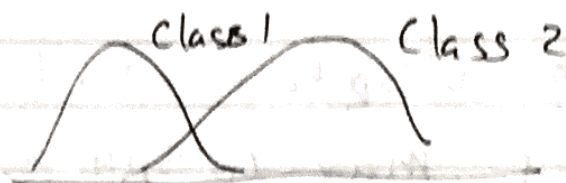


However, that approach does not pay attention to class labels.

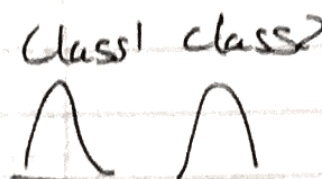
Fisher linear discriminant analysis (FLDA)

For two classes, ~~it~~ it is just ~~a~~ a matter of finding a $w (= \mu_1 - \mu_2)$ weight vector that does the best discrimination.

For x_i from a particular class we expect $z_i = w^T x_i$ would be normally distributed



we should choose w s.t.



$$m_c = \frac{1}{N_c} \sum_{i: y_i = c} x_i$$

$$S_c^2 = \sum_{i: y_i = c} (x_i - m_c)^2$$

↑
sum sq. error

Construct a measure

$$J(w) = \frac{(m_2 - m_1)^2}{S_1^2 + S_2^2} = \frac{(\text{centroid dist})^2}{(\text{intraclass variability})}$$

This is proportional to the pooled t -statistic

$$J(w) = \frac{w^T S_B w}{w^T S_W w}$$

$$S_B = (\mu_2 - \mu_1)(\mu_2 - \mu_1)^T$$

$$S_W = \sum_{i: y_i = 1} (x_i - \mu_1)(x_i - \mu_1)^T + \sum_{i: y_i = 2} (x_i - \mu_2)(x_i - \mu_2)^T$$

$$\arg \max_{\substack{w \\ \|w\|=1}} J(w)$$

We need the direction,
So, we could choose different
normalizations of w .

$$\begin{aligned} & \text{Max } W^T S_B W \quad \text{s.t. } W^T S_W W = 1 \\ & \text{max } W^T S_B W - \lambda (W^T S_W W - 1) \\ & \text{max } W^T S_B W - \lambda (W^T S_W W - 1) \end{aligned}$$

With the lagrange multiplier adjusted to satisfy the constraint

$$S_B W = \lambda S_W W$$

Generalized eigenvalue problem

If S_W is invertible.

$S_W^{-1} S_B W = \lambda W$ is a conventional eigenvalue problem

Find the largest eigenvalue.

Because S_B is only rank one for two classes, making further simplification possible

$$\begin{aligned} S_B W &= (\mu_2 - \mu_1)(\mu_2 - \mu_1)^T W = (\mu_2 - \mu_1)(m_2 - m_1) \\ \lambda W &= S_W^{-1} (\mu_2 - \mu_1)(\mu_2 - \mu_1)^T W \leftarrow \text{scalar} \quad \text{ } \quad \text{ } \\ W &\propto S_W^{-1} (\mu_2 - \mu_1) \end{aligned}$$

This method could be generalized to more classes. Since S_B is a rank $C-1$ matrix. We could only find $L \leq C-1$ dimensions. This could be a limitation in some applications.

For completeness: $z_i = W x_i$
 [Rows of W are like w^T] $z_i \in \mathbb{R}^L, x_i \in \mathbb{R}^D$

$$m_c = \frac{1}{N_c} \sum_{i: y_i = c} z_i$$

$$m_c \in \mathbb{R}^L$$

$$m = \frac{1}{N} \sum_c N_c m_c$$

over all mean

$$J(W) = \frac{|\sum_c N_c (m_c - m)(m_c - m)^T|}{|\sum_{c=1}^C \sum_{i: y_i = c} (z_i - m_c)(z_i - m_c)^T|}$$