

Multivariate Data

d -columns - d variables

(inputs / features / attributes)

N -rows - i.i.d.

observation / examples / instances.

$$X = \begin{bmatrix} X_1^1 & X_2^1 & \dots & X_d^1 \\ X_1^2 & & & X_d^2 \\ \vdots & & & \vdots \\ X_1^N & \dots & \dots & X_d^N \end{bmatrix}$$

EG. Loan application -

customer observation vector - age, marital status, income, etc.

Typically, these variables are correlated
If not, great simplification.

Goals: simplification / summarization - explain data by few parameters.
exploratory, - generate hypothesis about data.

Parameter Estimation

$$E[X] = \mu = [\mu_1 \dots \mu_d]$$

Variance of X_i is σ_i^2 , covariance $\sigma_{ij} = \text{Cov}(X_i, X_j)$
 $= E[(X_i - \mu_i)(X_j - \mu_j)]$.

Σ is a matrix, $\Sigma_{ii} = \sigma_i^2$, $\Sigma_{ij} = \sigma_{ij}$.

$$\Sigma = E[(X - \mu)(X - \mu)^T] = E[XX^T] - \mu\mu^T$$

Correlation

$$\text{Corr}(X_i, X_j) = \rho_{ij} = \frac{\sigma_{ij}}{\sigma_i \sigma_j}$$

If variables are indep, then correlation and covariance is zero

(2)

Parameter Estimation (cont)

Given multivariate sample
Sample mean. $\underline{m} = \frac{1}{N} \sum_{t=1}^N \underline{x}^t$, $m_i = \frac{1}{N} \sum_{t=1}^N x_i^t$

Sample covariance $\sum$ with entries:
 $S_i^2 = \frac{1}{N} \sum_{t=1}^N (x_i^t - m_i)^2$

$$S_{ij} = \frac{1}{N} \sum_{t=1}^N (x_i^t - m_i)(x_j^t - m_j)$$

sample correlation. $r_{ij} = S_{ij} / S_i S_j$.

Missing Values — some variables may be missing in observations.
imputation — estimate the missing variable

~ mean imputation — substitute the mean of other observations.

imputation by regression — predict these values from existing ones.

but — data may not be missing at random

— if so, need to model why it is missing
(e.g. heights of Americans — Army data)

Multivariate Normal

$$p(\underline{x}) = \frac{1}{(2\pi)^{d/2} |\underline{\Sigma}|^{1/2}} e^{-\frac{1}{2}(\underline{x} - \underline{\mu})^T \underline{\Sigma}^{-1}(\underline{x} - \underline{\mu})}$$

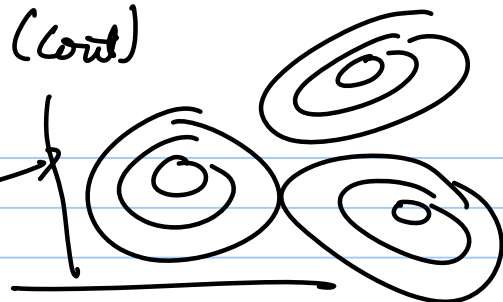
Mahalanobis distance $(\underline{x} - \underline{\mu})^T \underline{\Sigma}^{-1}(\underline{x} - \underline{\mu})$

Level sets $(\underline{x} - \underline{\mu})^T \underline{\Sigma}^{-1}(\underline{x} - \underline{\mu}) = c^2$ are d -dim
hyperellipsoids centered on $\underline{\mu}$, shape & orientation
determined by $\underline{\Sigma}$.

(3)

Multivariate Normal (cont)Bivariate case ($d=2$)

level sets.



$$\text{If } \underline{x} \sim N(\underline{\mu}, \underline{\Sigma})$$

then each dimension of $\underline{x}$ is a univariate normal (converse not necessarily true).

Special case — if the components of $\underline{x}$ are independent then the correlations (covariances) are 0.

$$P(\underline{x}) = \prod_{i=1}^d P(x_i) = \frac{1}{(2\pi)^d \prod_{i=1}^d \sigma_i} e^{-\frac{1}{2} \sum_{i=1}^d \left(\frac{x_i - \mu_i}{\sigma_i} \right)^2}$$

(Note: we can always pick a coordinate rep for which this is true $\rightarrow$ by rotating space to diagonals)

$$\underline{\Sigma} \quad \text{eg.} \quad \underline{\Sigma} \rightarrow \underline{\Phi}^T \underline{\Sigma} \underline{\Phi} \quad \underline{\Phi} \text{ rotation}$$

$\underline{\Phi}$ can be obtained from eigenvectors/eigenvalues of $\underline{\Sigma}$

Importance Property

the projection of a d -dim normal onto a vector $\underline{w}$ is also a Gaussian.

so is the projection onto a subspace $\underline{w}$.

$$\underline{w}^T \underline{x} = w_1 x_1 + \dots + w_d x_d \sim N(\underline{w}^T \underline{\mu}, \underline{w}^T \underline{\Sigma} \underline{w})$$

$$\underline{w}^T \underline{x} \sim N_k(\underline{w}^T \underline{\mu}, \underline{w}^T \underline{\Sigma} \underline{w}) \quad \underline{w} \text{ } d \times k \text{ matrix. rank } k \leq d.$$

(4)

Multivariate Classification

$$p(\underline{x} | C_i) = \frac{1}{(2\pi)^{d/2} |\Sigma_i|^{1/2}} e^{-\frac{1}{2} (\underline{x} - \underline{\mu}_i)^T \Sigma_i^{-1} (\underline{x} - \underline{\mu}_i)}$$

(analytic simplification)

Discriminant function

$$g_i(\underline{x}) = \log p(\underline{x} | C_i) + \log p(C_i)$$

$$g_i(\underline{x}) = -\frac{d}{2} \log 2\pi - \frac{1}{2} \log |\Sigma_i| - \frac{1}{2} (\underline{x} - \underline{\mu}_i)^T \Sigma_i^{-1} (\underline{x} - \underline{\mu}_i) + \log p(C_i)$$

Given training samples $\mathcal{X} = \{ \underline{x}^t, r^t \}$

Estimate the distribution: $\hat{p}(C_i) = \sum_t r_i^t / N$

$$\underline{m}_i = \sum_t r_i^t \underline{x}^t / \sum_t r_i^t$$

$$\underline{S}_i = \sum_t r_i^t (\underline{x}^t - \underline{m}_i)(\underline{x}^t - \underline{m}_i)^T / \sum_t r_i^t$$

Plug into discriminant function:

$$g_i(\underline{x}) = -\frac{1}{2} \log |\underline{S}_i| - \frac{1}{2} (\underline{x} - \underline{m}_i)^T \underline{S}_i^{-1} (\underline{x} - \underline{m}_i) + \log \hat{p}(C_i)$$

Quadratic Discriminant

(dropping constant term)

Can be written as

$$g_i(\underline{x}) = \underline{x}^T \underline{w}_i \underline{x} + \underline{w}_i^T \underline{x} + w_{i0}$$

$$\underline{w}_i = -\frac{1}{2} \underline{S}_i^{-1}, \quad w_{i0} = \underline{S}_i^{-1} \underline{m}_i$$

$$w_{i0} = -\frac{1}{2} \underline{m}_i^T \underline{S}_i^{-1} \underline{m}_i - \frac{1}{2} \log |\underline{S}_i| + \log \hat{p}(C_i)$$

No. parameters: $K \cdot d$ mean + $K \cdot d(d+1)/2$ covariance

(5)

Multivariate Classification (cont)

For large d , samples are small, $\underline{\Sigma}_i$ may be singular and $\underline{\Sigma}^{-1}$ won't exist.

May need to decrease dimensionality (next lecture) or pool data and estimate common covariance matrix: $\underline{\Sigma} = \sum_i \hat{P}(C_i) \underline{\Sigma}_i$.

discriminant reduces to

$$g_i(\underline{x}) = -\frac{1}{2} (\underline{x} - \underline{m}_i)^T \underline{\Sigma}^{-1} (\underline{x} - \underline{m}_i) + \log \hat{P}(C_i)$$

Further simplification

assuming off-diagonal elements covariances are zero.

naïve Bayes classifiers

$p(x_j | C_i)$ univariate Gaussian

$\underline{\Sigma}$ is diagonal (so is $\underline{\Sigma}^{-1}$)

$$g_i(\underline{x}) = -\frac{1}{2} \sum_{j=1}^d \left(\frac{x_j - m_{ij}}{s_j} \right)^2 + \log \hat{P}(C_i)$$

m_{ij} is mean of j^{th} component of i^{th} class.

$\rightarrow$ reduces complexity of $\underline{\Sigma}$ from $O(d^2)$ to $O(d)$.

Extreme simplification: assume all variances are the same. Then Mahalanobis distance = Euclidean distance.

$$g_i(\underline{x}) = -\frac{1}{2} \sum_{j=1}^d (x_j - m_{ij})^2 + \log \hat{P}(C_i)$$

If priors are equal - $g_i(\underline{x}) = -\|\underline{x} - \underline{m}_i\|^2$
nearest mean classifier. $\rightarrow$ prototype/template matching

$$g_i(\underline{x}) = \underline{w}_i^T \underline{x} + w_{i0}, \quad \underline{w}_i = \underline{m}_i, \quad w_{i0} = -\frac{1}{2} \|\underline{m}_i\|^2$$

6) Tuning Complexity

As before, for polynomial regression, we can choose between models with different complexity.

→ Full Covariance $d + \frac{1}{2}d(d+1)$ parameters

→ Indep-Same Variances $d + 1$

Same procedures as before:

(e.g. cross-validation, regularization, etc.)

Discrete Features:

Discrete attributes - color $\in \{\text{red, blue, green, ...}\}$
pixel $\in \{\text{on, off}\}$

Suppose x_j are binary (Bernoulli)

$$p_{ij} = p(x_j = 1 | C_i)$$

$$p(x | C_i) = \prod_{j=1}^d p_{ij}^{x_j} (1 - p_{ij})^{(1-x_j)}$$

Discriminant

$$g_i(x) = \log p(x | C_i) + \log P(C_i)$$

$$= \sum_j \{x_j \log p_{ij} + (1-x_j) \log (1-p_{ij})\} + \log P(C_i)$$

Estimator $\hat{p}_{ij} = \frac{\sum_t x_j^t r_i^t}{\sum_t r_i^t}$ linear in x

General Case: multinomial x_j chosen from $(v_1, \dots, v_{n_j})$

$p_{ijk} = \text{prob}(x_j \text{ from class } C_i \text{ takes value } v_k)$

$$p_{ijk} = p(z_{jk} = 1 | C_i) = p(x_j = v_k | C_i)$$

(7)

Discrete Features (cont)

$$z_{jk}^t = \begin{cases} 1, & \text{if } x_j^t = v_k \\ 0, & \text{otherwise} \end{cases}$$

If attributes are indep. n_j

$$p(x|C) = \prod_{j=1}^d \prod_{k=1}^{n_j} p_{ijk}$$

discriminant $g_i(x) = \sum_j \sum_k z_{jk} \log p_{ijk} + \log \hat{p}(C_i)$

Max like estimate for p_{ijk} is $\hat{p}_{ijk} = \sum_t z_{jk}^t r_{it}^+ / \sum_t r_{it}^+$

Multivariate Regression

$$r^t = g(\underline{x}^t | \omega_0, \dots, \omega_d) + \epsilon = \omega_0 + \omega_1 x_1^t + \dots + \omega_d x_d^t + \epsilon.$$

least squares:

$$E(\omega_0, \dots, \omega_d | X) = \frac{1}{2} \sum_t (r^t - \omega_0 - \omega_1 x_1^t - \dots - \omega_d x_d^t)^2.$$

$$\underline{X} = \begin{bmatrix} 1 & x_1^1 & \dots & x_d^1 \\ \vdots & \vdots & & \vdots \\ 1 & x_1^N & \dots & x_d^N \end{bmatrix}, \quad \underline{\omega} = \begin{bmatrix} \omega_0 \\ \vdots \\ \omega_d \end{bmatrix}, \quad \underline{r} = \begin{bmatrix} r_1 \\ r_2 \\ \vdots \\ r_N \end{bmatrix}$$

minimize wrt. ω 's

$$\underline{X}^T \underline{X} \underline{\omega} = \underline{X}^T \underline{r}.$$

$$\hat{\underline{\omega}} = (\underline{X}^T \underline{X})^{-1} \underline{X}^T \underline{r}.$$