



EXAMINATION PAPER

Examination Session: May/June	Year: 2025	Exam Code: MATH4287-WE01
---	----------------------	------------------------------------

Title: High-Dimensional Statistics
--

Time:	2 hours	
Additional Material provided:		
Materials Permitted:		
Calculators Permitted:	No	Models Permitted: Use of electronic calculators is forbidden.

Instructions to Candidates:	Answer all questions. Section A is worth 40% and Section B is worth 60%. Within each section, all questions carry equal marks. Write your answer in the white-covered answer booklet with barcodes. Begin your answer to each question on a new page.
-----------------------------	---

Revision:	
------------------	--

SECTION A

- Q1** (a) Consider the clustering problem for a high dimensional data set with $n = 67$ observations and $p = 3000$ variables. Figure 1 shows the **scree plot** from the K -means clustering method with the number of clusters varying from 1 to 15. Based on this scree plot, what number of clusters is most appropriate for the K -means clustering of this data and why?

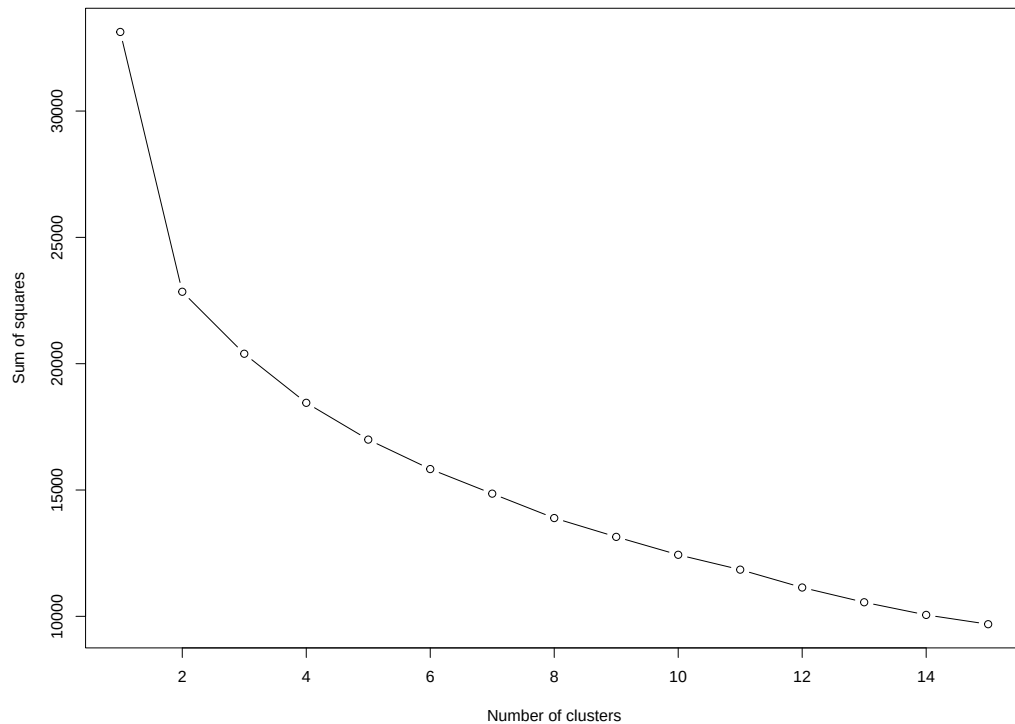


Figure 1: Scree plot for the K -means clustering in part (a) of Q1.

- (b) Now consider hierarchical clustering for the data set in Part (a). Figure 2 shows the dendrogram from hierarchical clustering for this high dimensional data set. Based on this dendrogram, explain what type of hierarchical clustering (agglomerative or divisive) is used here.
- (c) Using the dendrogram shown in Figure 2 on page 4, how many clusters do you suggest for this data? Explain your answer.
- (d) The R output on the next page reports the proportion of variance and the cumulative proportion captured by the principal components for this data. Based on this output, how many principal components capture about 95% of the data variance?
Also, explain briefly why the total variance (100%) is fully captured by only 67 principal components while there are 3000 variables in this data set.

Importance of components:

	PC1	PC2	PC3	PC4	PC5	PC6	PC7
Standard deviation	14.4311	7.9997	5.84184	4.51285	4.28174	3.95003	3.53042
Proportion of Variance	0.4149	0.1275	0.06798	0.04057	0.03652	0.03108	0.02483
Cumulative Proportion	0.4149	0.5424	0.61034	0.65091	0.68743	0.71852	0.74335
	PC8	PC9	PC10	PC11	PC12	PC13	PC14
Standard deviation	3.3459	3.1447	3.11112	2.95285	2.61072	2.41062	2.27749
Proportion of Variance	0.0223	0.0197	0.01928	0.01737	0.01358	0.01158	0.01033
Cumulative Proportion	0.7657	0.7853	0.80463	0.82200	0.83558	0.84715	0.85749
	PC15	PC16	PC17	PC18	PC19	PC20	PC21
Standard deviation	2.08230	2.02676	1.99309	1.87063	1.82158	1.76700	1.66751
Proportion of Variance	0.00864	0.00818	0.00791	0.00697	0.00661	0.00622	0.00554
Cumulative Proportion	0.86612	0.87431	0.88222	0.88919	0.89580	0.90202	0.90756
	PC22	PC23	PC24	PC25	PC26	PC27	PC28
Standard deviation	1.61992	1.59664	1.53104	1.46407	1.3993	1.34242	1.30520
Proportion of Variance	0.00523	0.00508	0.00467	0.00427	0.0039	0.00359	0.00339
Cumulative Proportion	0.91279	0.91787	0.92254	0.92681	0.9307	0.93430	0.93769
	PC29	PC30	PC31	PC32	PC33	PC34	PC35
Standard deviation	1.29655	1.26377	1.21520	1.1848	1.10976	1.10597	1.07783
Proportion of Variance	0.00335	0.00318	0.00294	0.0028	0.00245	0.00244	0.00231
Cumulative Proportion	0.94104	0.94422	0.94716	0.9500	0.95241	0.95485	0.95716
	PC36	PC37	PC38	PC39	PC40	PC41	PC42
Standard deviation	1.06867	1.05511	1.01839	0.99151	0.97300	0.96405	0.9501
Proportion of Variance	0.00228	0.00222	0.00207	0.00196	0.00189	0.00185	0.0018
Cumulative Proportion	0.95944	0.96166	0.96372	0.96568	0.96757	0.96942	0.9712
	PC43	PC44	PC45	PC46	PC47	PC48	PC49
Standard deviation	0.94500	0.93754	0.91271	0.88823	0.85382	0.83616	0.82573
Proportion of Variance	0.00178	0.00175	0.00166	0.00157	0.00145	0.00139	0.00136
Cumulative Proportion	0.97300	0.97475	0.97641	0.97798	0.97943	0.98082	0.98218
	PC50	PC51	PC52	PC53	PC54	PC55	PC56
Standard deviation	0.81650	0.80496	0.78883	0.78213	0.77324	0.76232	0.75980
Proportion of Variance	0.00133	0.00129	0.00124	0.00122	0.00119	0.00116	0.00115
Cumulative Proportion	0.98351	0.98480	0.98604	0.98726	0.98845	0.98961	0.99076
	PC57	PC58	PC59	PC60	PC61	PC62	PC63
Standard deviation	0.73549	0.72290	0.71643	0.7082	0.70021	0.69464	0.65413
Proportion of Variance	0.00108	0.00104	0.00102	0.0010	0.00098	0.00096	0.00085
Cumulative Proportion	0.99184	0.99288	0.99390	0.9949	0.99587	0.99684	0.99769
	PC64	PC65	PC66	PC67			
Standard deviation	0.65224	0.62365	0.58815	1.992e-14			
Proportion of Variance	0.00085	0.00077	0.00069	0.000e+00			
Cumulative Proportion	0.99854	0.99931	1.00000	1.000e+00			

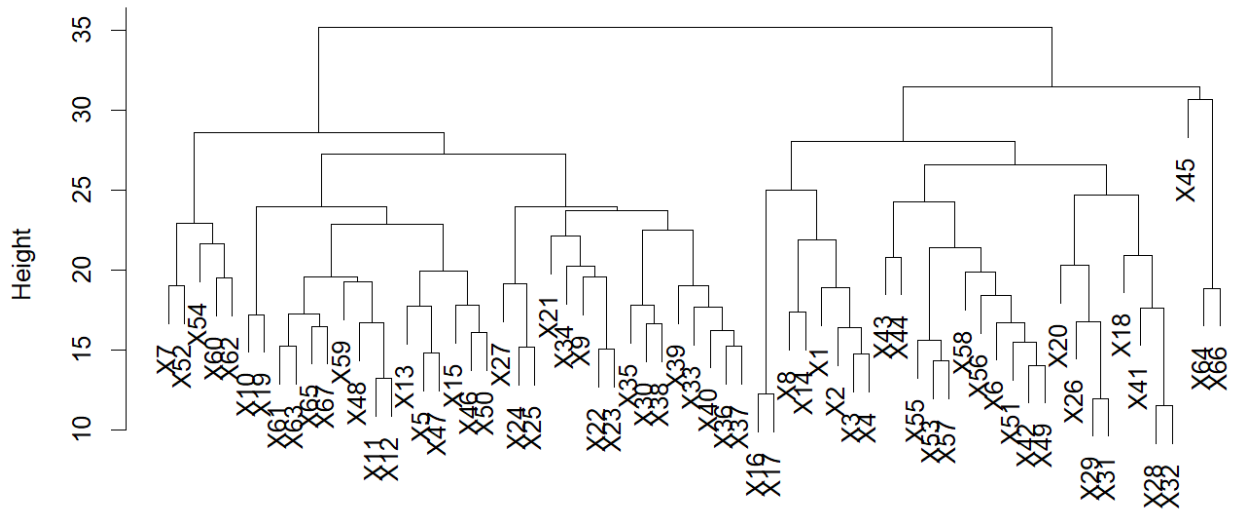


Figure 2: Dendrogram from the hierarchical clustering in part (b) of Q1.

Q2 Consider the linear regression model $\mathbf{Y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon}$. The elastic net estimator of $\boldsymbol{\beta}$ is defined as follows

$$\hat{\boldsymbol{\beta}}_{\text{EN}}(\lambda_1, \lambda_2) = \arg \min_{\boldsymbol{\beta} \in \mathbb{R}^p} \left\{ \|\mathbf{Y} - \mathbf{X}\boldsymbol{\beta}\|_2^2 + \lambda_1 \|\boldsymbol{\beta}\|_1 + \lambda_2 \|\boldsymbol{\beta}\|_2^2 \right\},$$

where $\lambda_1, \lambda_2 \geq 0$ are two separate regularisation parameters.

- (a) Prove that the elastic net estimator with a fixed λ_2 can be computed using a lasso problem.

Hint: Make an appropriate enlargement of the design matrix $\mathbf{X}$ using some additional matrix and proceed from there.

- (b) Consider a simple case where the columns of $\mathbf{X}$ are orthonormal, that is, we have $\mathbf{X}^T \mathbf{X} = \mathbf{I}_p$ with $\mathbf{I}_p$ being the $p \times p$ identity matrix, meaning that the columns of $\mathbf{X}$ are orthogonal with norm 1. Show that the elastic net estimator in this simple case can be obtained in closed form as follows

$$[\hat{\boldsymbol{\beta}}_{\text{EN}}(\lambda_1, \lambda_2)]_j = \frac{\text{sign}([\mathbf{X}^T \mathbf{Y}]_j) (|[\mathbf{X}^T \mathbf{Y}]_j| - \frac{\lambda_1}{2})_+}{1 + \lambda_2}, \quad j = 1, \dots, p,$$

where $\text{sign}(x) = 1(x > 0) - 1(x < 0)$ for $x \neq 0$, and also $(x)_+ = \max(x, 0)$.

Hint: For this question and throughout the paper you may use the following formulae on matrix differentiation. Let $\mathbf{x}$ be a vector, and further suppose that matrix $\mathbf{A}$ and vectors $\mathbf{a}$ and $\mathbf{b}$ are all not functions of $\mathbf{x}$. The following matrix differentiation results hold:

$$\begin{aligned} \frac{\partial \mathbf{a}}{\partial \mathbf{x}} &= \mathbf{0} \\ \frac{\partial \mathbf{x}}{\partial \mathbf{x}} &= \mathbf{I} \\ \frac{\partial \mathbf{Ax}}{\partial \mathbf{x}} &= \mathbf{A} \\ \frac{\partial \mathbf{x}^T \mathbf{A}}{\partial \mathbf{x}} &= \mathbf{A}^T \\ \frac{\partial \mathbf{x}^T \mathbf{Ax}}{\partial \mathbf{x}} &= \mathbf{x}^T (\mathbf{A} + \mathbf{A}^T) \quad \text{the row representation} \\ \frac{\partial \mathbf{x}^T \mathbf{Ax}}{\partial \mathbf{x}} &= (\mathbf{A} + \mathbf{A}^T) \mathbf{x} \quad \text{the column representation} \\ \frac{\partial \mathbf{a}^T \mathbf{xb}}{\partial \mathbf{x}} &= \mathbf{ab}^T \\ \frac{\partial \mathbf{a}^T \mathbf{x}^T \mathbf{b}}{\partial \mathbf{x}} &= \mathbf{ba}^T \\ \frac{\partial \mathbf{a}^T \mathbf{x}^T \mathbf{xb}}{\partial \mathbf{x}} &= \mathbf{x}^T (\mathbf{ab}^T + \mathbf{ba}^T) \quad \text{the row representation} \\ \frac{\partial \mathbf{a}^T \mathbf{xx}^T \mathbf{b}}{\partial \mathbf{x}} &= (\mathbf{ab}^T + \mathbf{ba}^T) \mathbf{x} \quad \text{the column representation.} \end{aligned}$$

SECTION B

- Q3** (a) Recall the linear regression model $\mathbf{Y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon}$, here with a scaled matrix $\mathbf{X}$ (i.e., all columns of $\mathbf{X}$ have mean 0 and variance 1). Assume the random errors $\boldsymbol{\varepsilon} = (\varepsilon_1, \dots, \varepsilon_n)$ are Gaussian and all independent with mean 0 and variance σ^2 . Let $\mathcal{F} := \left\{ \frac{2}{n} \|\boldsymbol{\varepsilon}^T \mathbf{X}\|_\infty \leq \lambda_0 \right\}$, where $\|\cdot\|_\infty$ denotes the L_∞ -norm of a vector (e.g., $\|\mathbf{a}\|_\infty = \max\{|a_1|, \dots, |a_p|\}$). Prove that, on $\mathcal{F}$ with $\lambda \geq 2\lambda_0$, we have

$$\frac{2}{n} \|\mathbf{X}(\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}^0)\|_2^2 + \lambda \|\hat{\boldsymbol{\beta}}_{S_0^c}\|_1 \leq 3\lambda \|\hat{\boldsymbol{\beta}}_{S_0} - \boldsymbol{\beta}_{S_0}^0\|_1^1,$$

where $\hat{\boldsymbol{\beta}}$ is the lasso estimator, $\boldsymbol{\beta}^0$ is the vector of unknown true parameter values, and S_0 denotes the active set.

Hint: Use the basic inequality for $\hat{\boldsymbol{\beta}}$, that is,

$$\frac{1}{n} \|\mathbf{X}(\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}^0)\|_2^2 + \lambda \|\hat{\boldsymbol{\beta}}\|_1 \leq \frac{2}{n} \boldsymbol{\varepsilon}^T \mathbf{X}(\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}^0) + \lambda \|\boldsymbol{\beta}^0\|_1^1,$$

and the Holder's inequality for two vectors $\mathbf{u}$ and $\mathbf{v}$ which is

$$\mathbf{u}^T \mathbf{v} \leq \|\mathbf{u}\|_q^1 \|\mathbf{v}\|_r^1, \quad \frac{1}{q} + \frac{1}{r} = 1.$$

- (b) Suppose that the compatibility condition holds for S_0 , that is, for some constant $\phi_0 > 0$ we can write

$$\|\hat{\boldsymbol{\beta}}_{S_0} - \boldsymbol{\beta}_{S_0}^0\|_2^2 \leq \|\mathbf{X}(\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}^0)\|_2^2 / (n\phi_0^2).$$

Prove that, on $\mathcal{F}$ with $\lambda \geq 2\lambda_0$, we have

$$\frac{1}{n} \|\mathbf{X}(\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}^0)\|_2^2 + \lambda \|\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}^0\|_1^1 \leq 4\lambda^2 s_0 / \phi_0^2,$$

where $s_0 = |S_0|$ is the sparsity index which is the cardinality of active set.

- (c) Assume $\lambda = 4\sigma \sqrt{\frac{t^2 + 2\log(p)}{n}}$, with a known σ (e.g., by prior knowledge or an appropriate estimate), for some constant $t > 0$. What would the result in part (b) claim about the optimality of lasso estimator $\hat{\boldsymbol{\beta}}$ with such choice of λ ? Justify your claim mathematically.

Q4 (a) Consider the sparse PCA problem

$$\min_{\mathbf{v}, \boldsymbol{\theta}} \sum_{i=1}^n \|\mathbf{X}_i - \boldsymbol{\theta} \mathbf{v}^T \mathbf{X}_i\|_2^2 + \lambda \|\mathbf{v}\|_2^2 + \lambda_1 \|\mathbf{v}\|_1, \quad \text{subject to } \boldsymbol{\theta}^T \boldsymbol{\theta} = 1.$$

For the case when both regularisation parameters λ and λ_1 are zero and $n > p$, prove that $\mathbf{v} = \boldsymbol{\theta}$ and is the leading principal component direction.

- (b) Write down the sparse PCA problem for the general case of d components with $\mathbf{v}_1, \dots, \mathbf{v}_d$. Also, briefly explain how to carry out the computation for this general sparse PCA problem.
- (c) We know that in K -means clustering, the (squared) Euclidean distance

$$d(\mathbf{X}_i, \mathbf{X}_j) = \|\mathbf{X}_i - \mathbf{X}_j\|_2^2 = \sum_{l=1}^p (\mathbf{X}_{il} - \mathbf{X}_{jl})^2$$

is often used as the dissimilarity measure for clustering the observations. Now suppose that we instead use the weighted Euclidean distance

$$d_W(\mathbf{X}_i, \mathbf{X}_j) = \frac{\sum_{l=1}^p w_l (\mathbf{X}_{il} - \mathbf{X}_{jl})^2}{\sum_{l=1}^p w_l},$$

where the w_l are some non-negative weights for clustering. Show that the weighted Euclidean distance satisfies

$$d_W(\mathbf{X}_i, \mathbf{X}_j) = d(\mathbf{Z}_i, \mathbf{Z}_j) = \sum_{l=1}^p (\mathbf{Z}_{il} - \mathbf{Z}_{jl})^2,$$

where

$$\mathbf{Z}_{il} = \mathbf{X}_{il} \left(\frac{w_l}{\sum_{l=1}^p w_l} \right)^{1/2}.$$

What is the implication of this result when applying the K -means clustering method in this case? Explain your answer.