

Exercise Sheet: Practical Data Collection

Data, Data Storage, Data Collection

Exercise 1 - Designing a Sampling Algorithm

You are tasked with collecting data on student study habits at your university. The target population is all 10,000 students, but you only have resources to survey 500. The university has 4 faculties with the following enrollments:

Faculty	Number of Students
Science	3000
Humanities	2000
Engineering	4000
Business	1000

(1) Systematic Sampling:

You are given a list of all 10,000 students.

- Write an algorithm to select every k -th student from the list to achieve a sample of 500. Shift the index of the first taken by a random value. The input should be the list of students *studentList*, and the size of the desired sample *sampleSize* (here 500), and the output should be a systematic sample of the associated size.
- What potential bias could this method introduce? How would you mitigate it?

(2) Stratified sampling

Stratified sampling is a technique where the population is divided into subgroups called *strata*, which share a common characteristic (here, the faculty). Instead of drawing a simple random sample from the entire population, we sample from each stratum proportionally to its size in the population.

This approach ensures that important subgroups are represented in the sample, reducing the risk that some groups are over- or underrepresented. It is particularly useful when the variable of interest may differ systematically between strata. In our example, study habits might vary across faculties, so sampling proportionally ensures that the collected data better reflects the diversity of the full student population.

- Calculate the number of students to sample from each faculty to ensure proportional representation in a sample of 500.
- Write an algorithm to generate a stratified random sample of 500 students. The input should be the list of students grouped by faculties as a list of pairs *faculties* = $[(facultyName, studentList)]$, and the size of the desired sample *sampleSize* (here 500), and the output should be a stratified sample of the associated size. You can reuse the previous algorithm and call it as a function *systematicSampling(studentList, strataSampleSize)*.
- Why is this approach better than simple random sampling for ensuring representation across faculties?

(3) Non-Response Bias

- (a) Suppose only 60% of the sampled students respond. Propose a method to adjust your analysis to account for non-response bias.

Exercise 2 - Data Collection and the Curse of Dimensionality

In this exercise, we explore the challenges of high-dimensional data collection and propose strategies to address them. The issue is known as the “curse of dimensionality”: as the number of dimensions (variables) increases, the data becomes increasingly sparse, making it difficult to:

- Find meaningful patterns or relationships.
- Generalize results to new data (risk of overfitting).
- Apply traditional statistical or machine learning methods effectively.

We consider the following scenario: A research team wants to collect data to predict student academic performance based on a wide range of factors. They propose measuring 50 variables, including:

- Demographic information (e.g., age, gender, socioeconomic status)
- Behavioral data (e.g., library visits, extracurricular activities)
- Digital footprints (e.g., online study platform usage, social media activity)
- Psychometric scores (e.g., stress levels, motivation)

The team plans to collect this data for 1,000 students.

- (1) Explain why collecting 50 variables for 1,000 students might lead to the “curse of dimensionality.” Use a simple calculation to illustrate the sparsity of the data. How could this affect the reliability of any conclusions drawn from the data?

Hint: what happens if every variable is binary (True or False)?

Therefore, the goal is to reduce the dimensionality of this dataset while preserving the pairwise distances between students as much as possible. Indeed, when analyzing high-dimensional data, the relationships between data points (e.g., students) are often captured by how “close” or “far” they are from each other in the original space. For instance, two students with similar study habits, socioeconomic backgrounds, and academic performance will be “close” in the 50-dimensional space, while two students with very different profiles will be “far” apart. If we reduce the dimensionality (e.g., from 50 to 10 variables), we want to ensure that these relationships are preserved.

The *Johnson-Lindenstrauss Lemma* states that a set of points in a high-dimensional space can be embedded into a lower-dimensional space in such a way that distances between the points are nearly preserved. Specifically, given $0 < \epsilon < 1$, an integer n , a set X of n points in $\mathbb{R}^D$, and an integer $d > \frac{8 \ln n}{\epsilon^2}$, there exists a (linear) map $f : \mathbb{R}^D \rightarrow \mathbb{R}^d$ such that for all pairs of points u, v in X ,

$$(1 - \epsilon)\|u - v\|^2 \leq \|f(u) - f(v)\|^2 \leq (1 + \epsilon)\|u - v\|^2$$

where $\|\cdot\|$ is the Euclidean norms (or L_2 norm) in the appropriate space.

- (2) Suppose you want to preserve the pairwise distances between students with an error margin of $\epsilon = 0.3$. Calculate the minimum number of dimensions d required to embed the dataset according to the Johnson-Lindenstrauss Lemma. Use the approximation $d \approx \frac{8 \ln n}{\epsilon^2}$, you can also approximate $\ln n$ as $\frac{2 \log n}{3}$.
- (3) If you reduce the dataset to $d = 10$ dimensions, what is the maximum error margin ϵ you can guarantee for the 1,000 students? Use the same approximation.

Assume you have implemented a dimensionality reduction technique and reduced the dataset to 10 dimensions. You now want to perform a k -nearest neighbors (k -NN) search on this reduced dataset. Let us recall how the k -NN algorithm works. Given a dataset and a query point, the algorithm identifies the k closest points (neighbors) to the query in the dataset, based on a distance metric (usually Euclidean distance). The label or value of the query point is then determined by the majority vote (for classification) or the average (for regression) of its k neighbors.

- (4) How might the error margin ϵ affect the results of your k -NN search? Discuss the implications for the accuracy of your search.
- (5) Discuss one trade-off between collecting more variables and collecting more samples (e.g., more students). How would you balance this trade-off in practice?