

Data Analytics

__/__/__

★ Data Analysis :-

- Data Analysis is a process of obtaining raw data and converting it into information useful for decision-making by users.

★ Data Analytics :-

- It is the science of examining raw data with the purpose of drawing conclusions about that information.

★ Applications of Data Analytics :-

- In commercial & industries, to enable organizations to make more informed business decisions and by scientists.
- By researches, to verify or disapprove scientific models, theories and hypotheses.

★ Analysis vs Analytics :-

- Data analytics is a broader term and includes data analysis as necessary subcomponent.
- Analytics defines the science behind the analysis.
- The science means understanding the cognitive process an analyst uses to understand problems and explore data in meaningful ways.
- Analytics also include data extract, transform, and load; specific tools, techniques, and methods; and how to successfully communicate results.

★ Data Science:-

- Dealing with unstructured and structured data, Data Science is a field that comprises of everything that related to data cleaning, preparation, and analysis.
- Involves in creation of new algorithms.

★ Data Vs. Information

- Data are the facts or details from which information is derived.
- Individual pieces of data are rarely useful alone.
- For data to become information, data needs to be put into context.

<u>Data</u>	<u>Information</u>
• Data means raw facts gathered about someone or something, which is bare and random.	• Facts, concerning a particular event or subject, which are refined by processing is called information.
• It is just text and numbers.	• It is refined data.
• Records and observations.	• Analysis
• Unorganized	• Organized.
• May or may not be useful.	• Always
• No	• Yes
• Does not depend on information	• Without data, information cannot be processed.

Characteristics of data :-

1. Accuracy and Precision
2. Legitimacy and Validity
3. Reliability and Consistency
4. Timeliness and Relevance
5. Completeness and Comprehensiveness
6. Availability and Accessibility
7. Granularity and Uniqueness

1. Accuracy and Precision :- This characteristic refers to the exactness of the data.
2. Legitimacy and Validity :- Requirements governing data set the boundaries of this characteristic.
3. Reliability and Consistency :- Regardless of what source collected the data or where it resides, it cannot contradict a value residing in a different source or collected by a different system.
4. Timeliness and Relevance :- Data collected or too late could misinterpret a situation and drive inaccurate decisions.
5. Completeness and Comprehensiveness :- Incomplete data is as dangerous as inaccurate data.
6. Availability and Accessibility :- This presumes that the data exists and is available for access to be granted.

_ / _ / _

7. Granularity and Uniqueness :- The level of detail at which data is collected is important, because confusion and inaccurate decisions can otherwise occur.

★ Steps involved in Data Analysis :-

1. Data Collection :- Identify good sources.
2. Data Cleaning :- Remove noise.
3. Data Analysis :- Apply algorithms.
4. Interpretation :- Explain results and take actions.

★ Data Munging :-

- Required for improve the quality of gathered data.
- ~~Int~~ Involves cleaning and transformation of messy data available with us.
- Also referred as Data Wrangling.

★ Reasons for ~~the~~ Noise in data :-

→ Lots of potentially incorrect data.

- Faulty instruments
- Human or computer error.
- Transmission errors.

eg. Occupation = " " (missing data)
Salary = " - 10" (an error)

→ Binning :-

- first sort data and partition into (equal-frequency) bins
- then one can smooth by bin means, smooth by bin median, smooth by bin boundaries, etc.

→ Regression :-

- smooth by fitting the data into regression functions.

→ Clustering :-

- detect and remove outliers.

→ Combined computer and human inspection :-

- detect suspicious values and check by human (e.g., deal with possible outliers)

* Data Transformation :-

- Transformation is mapping the data to a new space.

eg :-

- Fourier Transform
- Wavelet Transform

* Data Sampling :-

- obtaining a small sample S to represent the whole data set N .

* Types of Sampling:-

→ Simple Random Sampling:-

- There is an equal probability of selecting any particular item.

→ Stratified Sampling:-

- Partition the data set, and draw samples from each partition (proportionally, i.e., approximating the same percentage of the data).
- Used in conjunction with skewed data.

~~* Types of Sampling:-~~

Stratified Sampling:-

* Sampling without replacement:-

- Once an object is selected, it is removed from the population.

* Sampling with replacement:-

- A selected object is not removed from the population.

★ Data Cleaning

Problem Solving

Data % Fat

⋮
⋮
⋮

Q) mean, Median, Mode, Range, Standard Deviation, Variance
Find the above values for age and % fat.

Modes :-

One mode - unimodal

Two modes - bimodal

Three modes - trimodal

more than one mode - multimodal

$$SD = \sqrt{\frac{\sum (x - \bar{x})^2}{n - 1}}$$

$$SD = \sqrt{\text{Variance}}$$

★ Binning :-

Q) Perform SA

23, 23, 27, 27, 39, 41, 47, 49, 50, 52, 54, 54, 56, 57, 58, 60, 61

Perform smoothing operation using

- Smoothing by min means
- Smoothing by min & median
- Smoothing by bin boundaries using a bin depth of 3.

Ans. Divide the elements into bins of depth 3.

(frequency)

Bin 1 :- 23, 23, 27

Bin 2 :- 27, 39, 41

Bin 3 :- 47, 49, 50

Bin 4 :- 52, 54, 54

Bin 5 :- 56, 57, 58

Bin 6 :- 58, 60, 61

• Smoothing by Bin Means :-

Bin 1 :- 24, 24, 24

Bin 2 :- 36, 36, 36

Bin 3 :- 49, 49, 49

Bin 4 :- 53, 53, 53

Bin 5 :- 57, 57, 57

Bin 6 :- 60, 60, 60

{ Add the 3 elements of B1 and $\div$ by 3 }

← 0.5 →

• Smoothing by Bin Medians :-

Bin 1 :- 23, 23, 23

Bin 2 :- 39, 39, 39

Bin 3 :- 49, 49, 49

Bin 4 :- 54, 54, 54

Bin 5 :- 57, 57, 57

Bin 6 :- 60, 60, 60

{ Find the median of the 3 elements of B1 }

← 0.5 →

• Smoothing by Bin Boundaries :-

Bin 1 :- 23, 23, 27

Bin 2 :- 27, 41, 41

Bin 3 :- 47, 50, 50

Bin 4 :- 52, 54, 54

Bin 5 :- 56, 58, 58 or 56, 56, 58

Bin 6 :- 58, 61, 61

{ Find the middle element is close to which min/max value and change it to that value }

→ Equal width Partitioning :-

Width 10

B1 -	1 and 10	= 5, 10	} either do by mean, by Median, by boundary.
B2 -	11 and 20	= 13, 15, 11	
	21 and 30	= 35	
B3 -	31 and 40	7, 8, 35	} 1 values no smooth.
B4 -	41 and 50	= 50	
B5 -	51 and 60	7, 8, 58	
	61 and 70	=	
B6 -	71 and 80	= 72	
	81 and 90	=	
B7 -	91 and 100	= 92	

Q1) Find the missing values. [find by app. means of mean, median or mode]

74

68

56

57

60

?

68

71

53

61

57

68

68

~~50~~ 60

57.

~~74~~

$$74 + 68 + 56 + 57 + 60 + 68 + 71 + 53 + 61 + 57 + 68 + 68$$

$$60 + 57$$

$$14$$

$$\frac{?}{14} \rightarrow (63)$$

$$= 62.71 \approx (63)$$

* Data Sampling :-

91 95 103 115 112 101 102 97 115 99 102 97
100 107 116 107 110 113 95 93

① Simple Random Sampling w/o replacement. (size 5)

randomly :- 102 102 91 110 112 → Sample size = 5

Ways :- $20C_5$

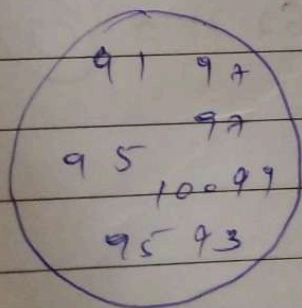
② Simple Random sampling with Replacement

103 103 103 103 103 → sample size = 5

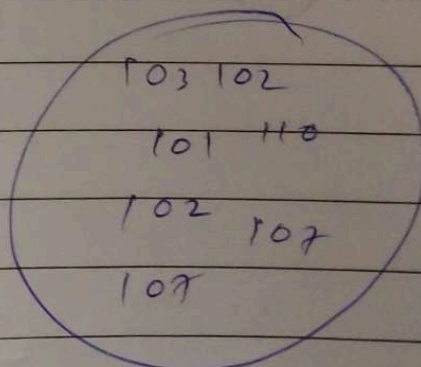
$|P| = 20 \rightarrow 20 \times 20 \times 20 \times 20 \times 20$

③ stratified sampling → heterogeneous sampling

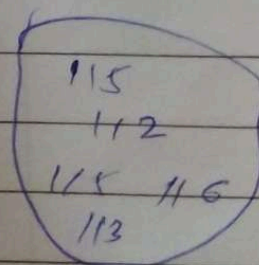
Score ≤ 100



Score ≤ 110



Score > 110



2

1

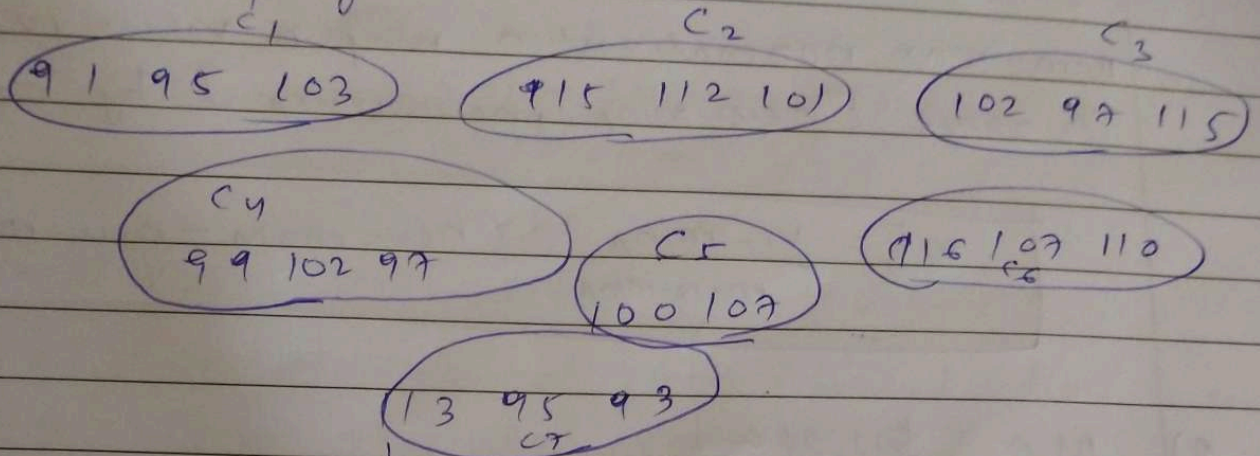
2

5

11
selective sample of size - 5 from stratified sample.

91 97 103 115 112

④ cluster sampling.



heterogeneous sampling.

Out of 7 clusters any 5 can be chosen
↓
clusters.

→ sampling $C_1 C_2 C_3 C_4 C_5$

//_

* Min-max Normalization:-

Min_A - minimum value of an attribute

Max_A - maximum value of an attribute, A .

Min-max normalization maps a value, v_i of A to v_i' in range $[new_min_A, new_max_A]$ by computing

$$v_i' = \frac{v_i - min_A}{max_A - min_A} (new_max_A - new_min_A) + new_min_A$$

Q) $Min = \$12,000$

$Max = \$98,000$

Map income to the range $[0.0, 1.0]$

By min-max normalization, a value of \$73,000 for income is transformed to:-

$$\frac{73,000 - 12,000}{98,000 - 12,000} \cdot (1.0 - 0.0) + 0 = 0.715$$

* Z-score Normalization:- (Zero mean)

~~Zero~~

the values of an attribute, A are normalized based on the mean and standard deviation of A .

$$v_i' = \frac{v_i - \bar{A}}{\sigma_A}$$

$\bar{A}$ = Mean

σ_A = S.D.

Q) Mean = \$54,000
SD = \$16,000
Z score, 73,600.

$$\begin{aligned} V_i' &= \frac{V_i - \bar{A}}{s_A} \\ &= \frac{73,600 - 54,000}{16,000} \\ &= 1.225 \end{aligned}$$

* Decimal Scale Normalization :-

The number of decimal points moved depends on the maximum absolute value of A.

A value of V_i of A is normalized to V_i' by computing :-

$$V_i' = \frac{V_i}{10^j}$$

where j is the smallest integer such that $\max(|V_i|)$

Q) 'A' range from -986 to 917. The max. absolute value of A is 986. divide each value by 1000 (i.e. $j=3$) (i.e. $j=3$) so that -986 normalized to -0.986 and 917 normalized to 0.917.

1) Use to normalize :-

200, 300, 400, 600, 1000

(a) min-max normalization by setting $\min = 0$ and $\max = 1$

$$\text{New-min} = 0$$

$$\text{New-max} = 1$$

$$\min A = 200$$

$$\max A = 1000$$

$$V_i' = \frac{V_i - \min A}{\max A - \min A} (\text{New-max} - \text{new-min}) + \text{new-min}$$

$$V_i' = \frac{200 - 200}{1000 - 200} \times (1 - 0) + 0$$

(b) Z-score normalization

$$V_i' = \frac{V_i - \bar{A}}{\sigma_A}$$

$$\bar{A} = 500$$

$$\sigma_A = 316.22$$

$$V_i = 200$$

$$V_i' = \frac{200 - 500}{316.22} = -0.99$$

(c) decimal scaling :-

$$V_i' = \frac{V_i}{10^i}$$

$$V_i' = \frac{200}{1000} = 0.2$$

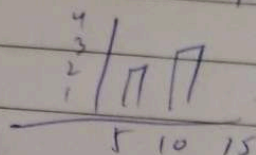
* Histograms :-

Buckets are of 3 types :-

1. Singleton
2. Equal width
2. Equal frequency.

1, 1, 5, 5

The no. has come the no. of times $\rightarrow$ histogram



* Data Wrangling in Python :-

- Data wrangling is the process of gathering, collecting, and transforming raw data into another format for better understanding, decision-making, accessing, and analysis in less time. Data wrangling is also known as Data Munging.

$\rightarrow$ Data Wrangling in Python deals with the below functionalities :-

1. Data exploration.
2. Dealing with missing values.
3. Reshaping data.
4. Filtering data.
5. ~~etc~~

* Data Scraping:-

- Data scraping refers to the process of extracting data for

Q) $x \leftarrow c(1, 2)$
 $y \leftarrow c(3, NA)$
 $x + y$

O/P [1] 4 NA

Q) $x \leftarrow 2:4$
 $y \leftarrow 6:10$
 $x + y$

O/P [1] 8 10 12 11 13

Q) Vector
 $x \leftarrow c(4, 3, 5, 6, 11, 19)$
 $x[x < 5] + x[x < 6]$

[1] 8 6 9

on $x[x < 5] + x[x < 6]$:

longer object length is not a multiple of shorter object length.

Q) repeat-each <- ~~rep~~ rep(c(1,2,3), each=3)
repeat-each

ans - [1] 1 1 1 2 2 2 3 3 3

Q) x <- rep(c(1,2,3), times=3)
x

[1] 1 2 3 1 2 3 1 2 3

Q) x <- rep(c(1,2,2), times=c(5,2,1))
x

[1] 1 1 1 1 1 2 2 2

Q) v <- seq(from=1, to=3.3, by=0.4)
v

v[c(-1,-2)]

[1] 1.0 1.4 1.8 2.2 2.6 3.0

[1] 1.8 2.2 2.6 3.0

Q) b <- vector(2, 3.5, 4, 5.5)

b

6/1- Error in vector(2, 3.5, 4, 5.5)
unused arguments (4, 5.5)

Q) sum(0:9+1) - sum(0:(9+1))

[1] 0

Varname can start
with ~~var~~ letter or $_$ / $_$ / $_$

Q1) `v <- sample(-10:10, 10)`

`print(v)`

`print(max(v) - min(v))` $(-10 + 9) (10 - (-9)) = 19$

`print(mean(v))`

• `print(mean(v, trim = 0.2))`

`print(median(v))`

O/P `[1] -9 -4 5 -5 10 10 -2 7 3 0 -8` [any 10
random
nos b/w the
range]

`[1] 19`

`[1] 0.5`

`[1]`

`[1] 1.5`

Q2) `v <- 1:10`

`sum(v[-1])`

$(\cancel{1} + 10 + (1 + 2 + \dots + 10) - 1)$

`length(v[-5:-8])`

`v[seq(-2, -8, -2)]`

• `v[c(-2, -4, -6, -8)]`

`[1] 54`

`[1] 6`

`[1] 1 3 5 7 9 10`

`[1] 1 3 5 7 9 10`

Q1) $v_1 \leftarrow -5:5$
 $v_2 \leftarrow v_1 [c(\text{length}(v_1):1)]$
 $\text{sum}(v_2 - v_1)$

O/P [1] 0

Q2) $v \leftarrow \text{sample}(-5:5, 5)$
 v
 $v_1 \leftarrow \text{sort}(v)$
 v_1
 $v_1[\text{length}(v_1) - v_1[1]]$

O/P [1] -4 1 0 -5 -1

[1] -5 -4 -3 -2 -1

[1] 6

Q3) $v \leftarrow \text{seq}(-5, 5, 2)$

v

$\text{append}(v, 10)$

$\text{append}(v, 10, 3)$

[Add 10 at index position 3]

O/P [1] -5 -3 -1 1 3 5

[1] -5 -3 -1 1 3 5 10

[1] -5 -3 -1 10 1 3 5

Q) $v \leftarrow c(1, 2, 3, 4, 5, 6)$
 $v \leftarrow v \times 2 \rightarrow \text{remainder}$
 $\text{sum}(v)$

O/P [1] 3

Q) $x \leftarrow c(1, 1, 4, 5, 4, 6)$
 $\text{unique}(x)$ [once ~~no~~ value]

O/P [1] 1, 4, 5, 6

Q) $x \leftarrow c(1, 1, 4, 5, 4, 6)$
 $\text{duplicated}(x)$ [if repeated $\rightarrow$ True, otherwise false]
 $x[\text{duplicated}(x)]$
 $x[!\text{duplicated}(x)]$

O/P [1] FALSE TRUE FALSE FALSE TRUE FALSE

[1] 1 4

[1] 1 4 5 6

Q) $m \leftarrow \text{matrix}(\text{seq}(\text{from} = 10, \text{to} = 18, \text{by} = 2), \text{nrow} = 3, \text{ncol} = 3)$

$\text{sum}(m[1,])$

	[1,]	[2,]	[3,]
[1,]	10	16	12
[2,]	12	18	14
[3,]	14	10	16
[1]	38	[Add column 1]	

Matrix :-

Q) $x \leftarrow 11:13$
 $y \leftarrow 10:12$ → rows
 $\dim(t(rbind(x, y)))$
 dimension → (row, column)

O/P [1] 3 2

Q) $a \leftarrow \text{matrix}(1:2, \text{nrow} = 2, \text{ncol} = 2)$
 $\text{sum}(a[1,] + a[,1])$

O/P [1] 5

Q) $\text{mat1} \leftarrow \text{matrix}(1:3, 3, 3, \text{byrow} = \text{T})$ → True - Row wise
 $\text{mat2} \leftarrow \text{matrix}(1:3, 3, 3)$ False - Column wise
 $\text{print}(\text{sum}(\text{mat1} - \text{mat2}))$

O/P [1] 0

Q) $x \leftarrow \text{matrix}(1:7, \text{nrow} = 3)$
 x

1 → 7 can isn't a multiple of 3.

Q) $\dim(t(\text{matrix}(1:10)))$

[1] 1 10
 col. Row

Q) $m \leftarrow \text{matrix}(1:6, \text{byrow} = \text{TRUE})$

$m[,1] + m[1,]$
↓ ↓
 1st column 1st row

[1] 2 3 4 5 6 7

Q) m1 <- matrix(1:6, byrow = TRUE)
 m2 <- matrix(1:6, byrow = FALSE)
 m1 - m2

	[,1]
[1,]	0
[2,]	0
[3,]	0
[4,]	0
[5,]	0
[6,]	0

Q) a <- c('abc', 'def', 'ghi', 'jkl')
 b <- c(11, 23, 355, 426)
 mat <- matrix(a, 2, 2)
 mat

	[,1]	[,2]
[1,]	"abc"	"ghi"
[2,]	"def"	"jkl"

Q) x <- matrix(c(1:7), 3, 2)
 y <- matrix(c(1:7), 3, byrow = T)
 x + y

	[,1]	[,2]	[,3]
[1,]	2	6	10
[2,]	6	10	7
[3,]	10	7	4

logical - TRUE, FALSE

integer = 36L

Complex = 5+2i

numeric = -24, 36

Char = " " " "

Data Types :-

Q1) `x <- c("a", "b", "TRUE")`
`as.logical(x)`

[1] NA NA TRUE

Q1) `x <- charToRaw("a"), charToRaw("b")`
`as.logical(x)`

[1] TRUE TRUE

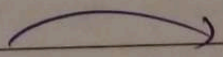
Q1) `a <- c(48, 54.5, 33i)`
`class(a)`

[1] "Complex"

[1] 48+0i 54.5+0i 0+33i

Q1) `as.logical(sum(as.numeric(F) + as.numeric(T)))`

[1] TRUE

Q1) `dim(numeric(6))`  0 0 0 0 0 0

[1] NULL

Q1) `x <- c(10L, 10, 10i, 10')`
`class(x)`

[1] "10" "10" "0+10i" "10"

[1] "character"

Q) print (sum (complex (500)))

[1] 0+0i

Q) print (complex (2))

[1] 0+0i 0+0i

Q) c <- c (TRUE, 24, FALSE, -34, 0)

c <- as.character (c)

c

[1] "1" "24" "0" "-34" "0"

Q) c <- c (TRUE, FALSE)

c <- as.character (c)

c

d <- c (TRUE, 24, FALSE, -34, 0)

d <- as.character (d)

d

[1] "TRUE" "FALSE"

When single T=T, F=F.

[1] "1" "24" "0" "-34" "0" [when with no,
T → 1, F → 0]

Q) cd <- c ("TRUE", "FALSE", "FALSE")

class (cd)

[1] character

AND		
A	B	Y
0	0	0
0	1	0
1	0	0
1	1	1

OR		
A	B	Y
0	0	0
0	1	1
1	0	1
1	1	1

d) in logical - 1 → True, 0 → False.

Q1) $x \leftarrow C(0, 1, 0, 1, 0, 1)$
 $y \leftarrow C(1, 0, 1, 0, 1, 0)$

print (sum(as.numeric(x & y)) + as.numeric(x & y))

[1] 0

Q) $x \leftarrow C(10, 20, 30, 40)$

$y \leftarrow C(11, 21, 31, 41)$

$x \& y$

$x \& y$

$x \& y$

$x \& y$

[1] TRUE

[1] TRUE TRUE TRUE TRUE

[1] TRUE

[1] TRUE TRUE TRUE TRUE

Q) $x \leftarrow C(1, 0, 1, 0)$

$y \leftarrow C(01, 11, 01, 11)$

$x \& y + x \& y$

[1] FALSE TRUE FALSE TRUE

array (2, 3, 4)
 ↓ ↓ ↓
 row col no. of matrices

1/1

Array :

a) this array <- c (1:18)

multiarray <- array (this array, dim = c (4, 3, 2))

multiarray [2, 3, 1] → 1st matrix

multiarray [2, 3, 2] → 2nd matrix

[1] 10

[2] 4

q) this array <- c (1:18)

multiarray <- array (this array, dim = c (4, 3, 2))

multiarray [, 1]

↓
print matrix 1.

multiarray [, 1] → col. 1 + col. 2
 ↓ ↓
 m1 m2

q) c (2, 3, 2, 2)

↓ ↓ ↓
 row col no. of matrices
 follow up

q) c (2, 3, 2, 2)

multiarray [2, 3, 1,]

List

Q1) $n \leftarrow c(2, 3, 5)$
 $s \leftarrow c("aa", "bb", "cc", "dd", "ee")$
 $b \leftarrow c(TRUE, FALSE, TRUE, FALSE, FALSE)$
 $lst \leftarrow list(n, s, b)$
 lst

$[[1]]$

$[1] 2 3 5$

$[[2]]$

$[1] "aa" "bb" "cc" "dd" "ee"$

$[[3]]$

$[1] TRUE FALSE TRUE FALSE FALSE$

And in list index position starts with 1.

$lst[[2]] \rightarrow$ second list

$lst[[2]][4] \rightarrow$ 4th element of second list.

Q1) $x \leftarrow c(5, 6)$

$y \leftarrow matrix(1:4, nrow=2, ncol=2)$

$l \leftarrow list(x, y)$

$l[[2]][1,] \rightarrow$ 1st row of 2nd list.

$[1] 1 3$

$[[2]] [1,2]$

$\rightarrow$ 2nd matrix, 1st row 2nd column

$$10f(8) = 18$$

$$10f - (8) \quad 8x = 20$$

$$10f \quad x = \frac{20}{8} = \frac{5}{2}$$

$$10f \quad x = 30$$

```

q) f <- function(x)
{
  if (x == 2)
  {
    return(2)
  }
  else
  {
    print('f')
    f(x-1)
  }
}

```

print(f(6))

[1] "+"

[2] "+"

[3] "+"

[4] "+"

[5] 2

```

q) f <- function(x)
{
  if (x > 0)
  {
    return(x + f(x-2))
  }
  else
  {
    return(0)
  }
}

```

print(f(10))

0) 1 1 1 1 1
 2 2 2 2
 3 3 3
 4 4
 5

```
stars = NULL
k = 5
for (i in 1:5)
{
  for (j in 1:k)
  {
    stars = c(stars, i)
  }
  print(stars)
  k = k - 1
  stars = NULL
}
```

0) 1
 2 2
 3 3 3
 4 4 4 4
 5 5 5 5 5

```
stars = NULL
for (i in 1:5) {
  for (j in 1:i) {
    stars = c(stars, i)
  }
  print(stars)
  stars = NULL
}
```


Q) 1
 1 2
 1 2 3
 1 2 3 4
 1 2 3 4 5

```
stars = NULL
k < 1
for (i in 1:5) {
  for (j in 1:k) {
    stars = c(stars, j)
  }
  print(stars)
  k = k + 1
  stars = NULL
}
```

Q) 1 2 3 4 5
 1 2 3 4
 1 2 3
 1 2
 1

```
stars = NULL
k < 5
for (i in 1:5) {
  for (j in 1:k) {
    stars = c(stars, j)
  }
  print(stars)
  k = k - 1
  stars = NULL
}
```


q) 1 0 0 0 0
 0 1 0 0 0
 0 0 1 0 0
 0 0 0 1 0
 0 0 0 0 1

```
stars = NULL
for (i in 1:5) {
  for (j in 1:5) {
    if (i == j)
      stars = c(stars, 1)
    else
      stars = c(stars, 0)
  }
}
print(stars)
stars = NULL
}
```

q) 1 0 0 0 0
 1 1 0 0 0
 1 1 1 0 0
 1 1 1 1 0
 1 1 1 1 1

if (i < j)

q) 1 1 1 1 1
 0 1 1 1 1
 0 0 1 1 1
 0 0 0 1 1
 0 0 0 0 1

if (i <= j)

Q1)

```

0
1 1
0 0 0
1 1 1 1
0 0 0 0 0

```

stars = NULL

for (i in 1:5) {

 for (j in 1:i) {

 if ((i % 2) == 0) {

 stars = C(stars, 1) }

 else

 {

 stars = C(stars, 0)

 }

 }

 print(stars)

 stars = NULL

}

Data Frame :-

```

1) x <- c('one', 'two', 'three', 'four')
   y <- c(1, 2, 3, 4)
   z <- c(4.5, 8.9, 3.3, 4.5)
   df <- data.frame(x, y, z)
   df

```

	x	y	z
1	one	1	4.5
2	two	2	8.9
3	three	3	3.3
4	four	4	4.5

class = "data.frame"

df[1]

one
two
three
four

df['x']

one
two
three
four

df\$x

one
one
one
one

two

three

four

class df\$x → "character"

df[, c(2, 3)]

↓ ↓
print col.2, col.3

df[c(2, 3)]

↓ ↓
col.2, col.3

q1) `x <- c('one', 'two', 'three', 'four')`

`y <- c(1, 2, 3, 4)`

`z <- c(4.5, 8.9, 3.3, 4.5)`

`df <- data.frame(x, y, z)`

`df[c(1, 2)]`

1st and 2nd values of df.

	x	y	z
1	one	1	4.5
2	two	2	8.9

q1) `x <- c`

`y <- c`

`z <- c`

`df <-`

`df`

`df[c(seq(1, nrow(df), 2)), ]`

q1) `x`

`y`

`z`

`df`

`rownames(df) <- c('first', 'second', 'third', 'four')`

`df`

	x	y	z
first			
second			
third			
four			

Q1 `df1 = data.frame ( student Id = c (1:4), subject = c ("OS", "DBMS", "Maths", "Java"))`

`df2 = data.frame ( student Id = c (5:7), subject = c ("C++", "Python", "RNR"))`

`df3 <- rbind (df1, df2)`

`df3`

	Student Id	subject
1	1	OS
2	2	DBMS
3	3	Maths
4	4	Java
5	5	C++
6	6	Python
7	7	RNR

Q1 `df <- data.frame ( student Id = c (1:4), subject = c ("OS", "DBMS", "Maths", "Java"))`

`newrow <- c (5, "C")`

`df [nrow(df)+1] <- newrow`

`df`

1	OS
2	DBMS
3	Maths
4	Java
5	C

VVI Q1) ~~df~~ <- data.frame (StudentId =)
~~df~~ \$ GPA <- c (8.5, 9.2, 8.9, 9.3)
~~df~~

St. Id	subject	GPA
1	OS	8.5
2	DBMS	9.2
3	Maths	8.9
4	Java	9.3

or ~~df~~ ['GPA'] <- c (8.9, 9.2, 8.9, 9.3)

or GPA <- c (8.9, 9.2, 9.3, 8.9)
~~df~~ <- cbind (~~df~~, GPA)

Q2) ~~df~~ <- data.frame (StudentId = c (1:4),
 subject = c ("OS", "DBMS", "Maths", "Java"),
 GPA <- c (8.5, 9.2, 8.9, 9.3)
~~df~~ <- cbind (~~df~~ [1], GPA, ~~df~~ [2])

~~df~~ 2 <- ~~df~~

StudentId GPA Subject

Q1) Program to find the HCF of two numbers.

```
hcf <- function(x, y) {
```

```
  if (x > y)
```

```
  {
```

```
    smaller = y
```

```
  }
```

```
  else
```

```
  {
```

```
    smaller = x
```

```
  }
```

```
  for (i in 1:smaller)
```

```
  {
```

```
    if ((x % i == 0) && (y % i == 0))
```

```
    {
```

```
      hcf = i
```

```
    }
```

```
  }
```

```
  return(hcf)
```

```
}
```

```
n1 = as.integer(readline(prompt = "Enter first no"))
```

```
n2 = as.integer(readline(prompt = "Enter 2nd no"))
```

```
print(paste("The hcf of", n1, "and", n2, "is",  
            hcf(n1, n2)))
```


Q1) Palindrome

```
is_palindrome <- function(num) {
```

```
  num_str <- as.character(num)
```

```
  reversed_num_str <- rev(strsplit(num_str, "")[[1]])
```

```
  if (identical(num_str, reversed_num_str)) {
```

```
    return(TRUE)
```

```
  }
```

```
  else
```

```
  {
```

```
    return(FALSE)
```

```
  }
```

```
}
```

Q2) Perfect number:

```
is_perfect <- function(num)
```

```
{
```

```
  if (num <= 0)
```

```
  {
```

```
    return(FALSE)
```

```
  }
```

```
  divisors <- numeric(0)
```

```
  for (i in 1:(num/2)) {
```

```
    if (num % i == 0) {
```

```
      divisors <- c(divisors, i)
```

```
    }
```

```
}
```



```

if (sum(divisors) == num)
{
    return(TRUE)
}
else
{
    return(FALSE)
}
}

```

Q1) Fibonacci

```

fibonacci <- function(n)
{
    fib <- c(0, 1)

```

```

    for (i in 3:n) {
        next_term = fib[i-1] + fib[i-2]
        fib[i] <- c(fib, next_term)
    }
    return(fib[1:n])

```

```

}

```