

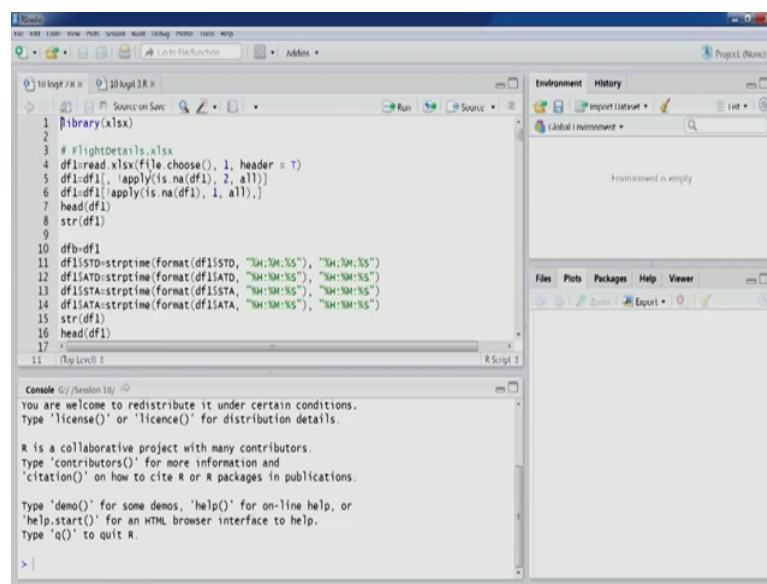
**Business Analytics & Data Mining Modeling Using R**  
**Dr. Gaurav Dixit**  
**Department of Management Studies**  
**Indian Institute of Technology, Roorkee**

**Lecture - 52**  
**Logistic Regression-Part VII**

Welcome to the course Business Analytics and Data Mining Modelling Using R. So, in previous few lectures we have been discussing different aspects of logistic regression so we will continue that discussion in this particular lecture as well.

So, in previous lecture we will be used apply details data set and then promotional offers data set as well. Some of the details regarding modelling exercise we could not cover.

(Refer Slide Time: 00:44)



```
1 library(xlsx)
2
3 # FlightDetails.xlsx
4 df1=read.xlsx(file.choose(), 1, header = T)
5 df1=df1[, !apply(is.na(df1), 2, all)]
6 df1=df1[!apply(is.na(df1), 1, all),]
7 head(df1)
8 str(df1)
9
10 dfb=df1
11 df1$STO=strptime(format(df1$STO, "%H:%M:%S"), "%H:%M:%S")
12 df1$ATO=strptime(format(df1$ATO, "%H:%M:%S"), "%H:%M:%S")
13 df1$STA=strptime(format(df1$STA, "%H:%M:%S"), "%H:%M:%S")
14 df1$ATA=strptime(format(df1$ATA, "%H:%M:%S"), "%H:%M:%S")
15 str(df1)
16 head(df1)
17
18 x
```

Console: C:/Session 10/

you are welcome to redistribute it under certain conditions.  
Type 'license()' or 'licence()' for distribution details.

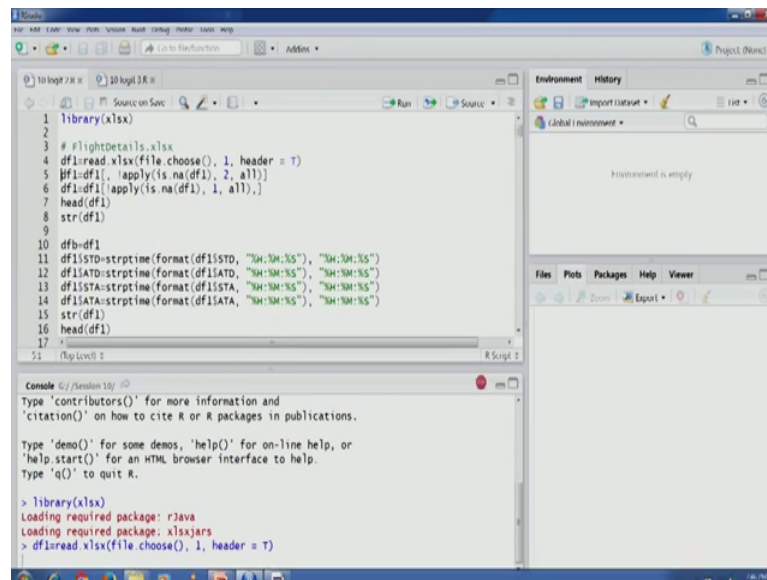
R is a collaborative project with many contributors.  
Type 'contributors()' for more information and  
'citation()' on how to cite R or R packages in publications.

Type 'demo()' for some demos, 'help()' for on-line help, or  
'help.start()' for an HTML browser interface to help.  
Type 'q()' to quit R.

> |

So, we will do that and discuss. So, let us move to let us import this data set. So, we will let us import the library x plus x. So, we are again going to use this particular data set flight details.

(Refer Slide Time: 01:00)



The screenshot shows the RStudio interface. The script editor on the left contains the following R code:

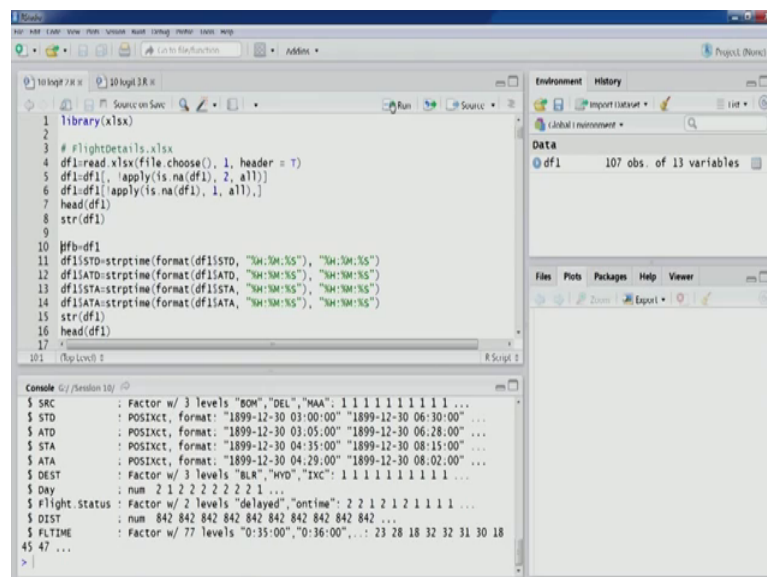
```
1 library(xlsx)
2
3 # flightDetails.xlsx
4 df1=read.xlsx(file.choose(), 1, header = T)
5 df1=df1[, !apply(is.na(df1), 2, all)]
6 df1=df1[!apply(is.na(df1), 1, all),]
7 head(df1)
8 str(df1)
9
10 dfb=df1
11 df1$DST=strptime(format(df1$DST, "%H:%M:%S"), "%H:%M:%S")
12 df1$ATD=strptime(format(df1$ATD, "%H:%M:%S"), "%H:%M:%S")
13 df1$STA=strptime(format(df1$STA, "%H:%M:%S"), "%H:%M:%S")
14 df1$ATA=strptime(format(df1$ATA, "%H:%M:%S"), "%H:%M:%S")
15 str(df1)
16 head(df1)
17
```

The console on the right shows the output of the first few lines of code:

```
> library(xlsx)
Loading required package: rJava
Loading required package: xlsxjars
> df1=read.xlsx(file.choose(), 1, header = T)
```

So, let us import this. Now, let us remove any columns any rows, let us look at these are the observations.

(Refer Slide Time: 01:18)



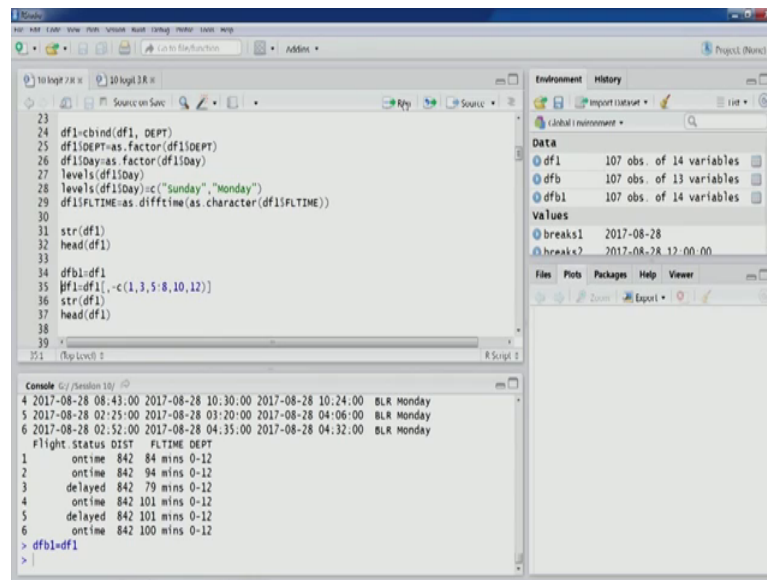
The screenshot shows the RStudio interface with the same script as before. The console now shows the output of the `str(df1)` command:

```
$ SRC      : Factor w/ 3 levels "BOM","DEL","MAA": 1 1 1 1 1 1 1 1 ...
$ STD      : POSIXct, format: "1899-12-30 03:00:00" "1899-12-30 06:30:00" ...
$ ATD      : POSIXct, format: "1899-12-30 03:05:00" "1899-12-30 06:28:00" ...
$ STA      : POSIXct, format: "1899-12-30 04:35:00" "1899-12-30 08:15:00" ...
$ ATA      : POSIXct, format: "1899-12-30 04:29:00" "1899-12-30 08:02:00" ...
$ DEST     : Factor w/ 3 levels "BLR","HYD","INC": 1 1 1 1 1 1 1 1 ...
$ Day      : num 2 1 2 2 2 2 2 2 1 ...
$ Flight.status : Factor w/ 2 levels "delayed","ontime": 2 2 1 2 1 1 1 1 ...
$ DIST     : num 842 842 842 842 842 842 842 842 ...
$ FLTIME   : Factor w/ 77 levels "0:35:00","0:36:00",...: 23 28 18 32 31 30 18
45 47 ...
>
```

The Environment pane on the right shows the data frame `df1` with 107 observations and 13 variables.

Let us the structure for the data frame. So, we will follow some of the steps that we have gone through in previous lectures as well, so we will just go through them.

(Refer Slide Time: 01:34)



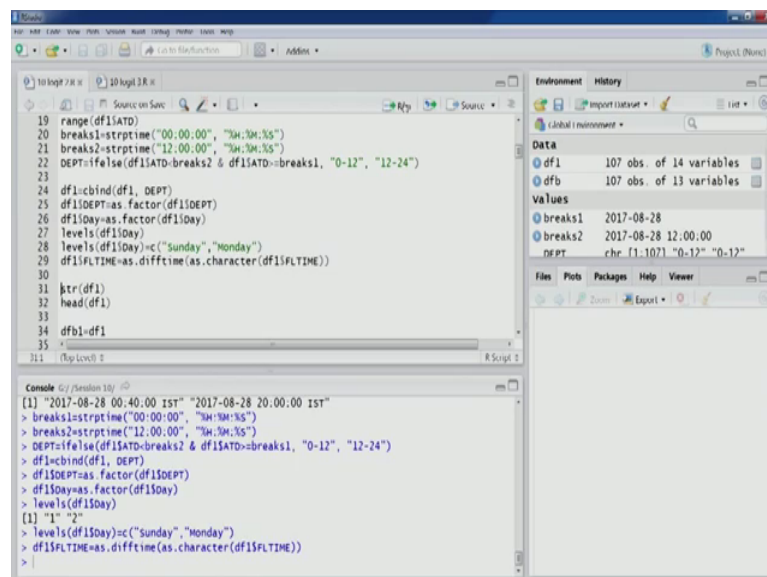
```
23 df1<-cbind(df1, DEPT)
24 df1$DEPT<-as.factor(df1$DEPT)
25 df1$Day<-as.factor(df1$Day)
26 levels(df1$Day)
27 levels(df1$Day)=c("Sunday","Monday")
28 df1$FLTIME<-as.difftime(as.character(df1$FLTIME))
29
30 str(df1)
31 head(df1)
32
33
34 dfb1<-df1
35 dfb1[,-c(1,3,5,8,10,12)]
36 str(dfb1)
37 head(dfb1)
38
39
```

Console

```
4 2017-08-28 08:43:00 2017-08-28 10:30:00 2017-08-28 10:24:00 BLR Monday
5 2017-08-28 02:25:00 2017-08-28 03:20:00 2017-08-28 04:06:00 BLR Monday
6 2017-08-28 02:52:00 2017-08-28 04:35:00 2017-08-28 04:32:00 BLR Monday
Flight Status DIST FLTIME DEPT
1 ontime 842 84 mins 0-12
2 ontime 842 94 mins 0-12
3 delayed 842 79 mins 0-12
4 ontime 842 101 mins 0-12
5 delayed 842 101 mins 0-12
6 ontime 842 100 mins 0-12
> dfb1=df1
>
```

Now, there is this particular exercise in the previous lecture we had used a separate grouping for departure time. So, again I have done certain changes into this grouping also, but this is not very important; however, let us run the model with a new grouping for departure time interval. So, this is the range we already familiar with.

(Refer Slide Time: 02:00)



```
19 range(df1$ATD)
20 breaks1<-strptime("00:00:00", "%H:%M:%S")
21 breaks2<-strptime("12:00:00", "%H:%M:%S")
22 DEPT<-ifelse(df1$ATD>=breaks2 & df1$ATD<=breaks1, "0-12", "12-24")
23
24 df1<-cbind(df1, DEPT)
25 df1$DEPT<-as.factor(df1$DEPT)
26 df1$Day<-as.factor(df1$Day)
27 levels(df1$Day)
28 levels(df1$Day)=c("Sunday","Monday")
29 df1$FLTIME<-as.difftime(as.character(df1$FLTIME))
30
31 str(df1)
32 head(df1)
33
34 dfb1<-df1
35
```

Console

```
[1] "2017-08-28 00:40:00 IST" "2017-08-28 20:00:00 IST"
> breaks1<-strptime("00:00:00", "%H:%M:%S")
> breaks2<-strptime("12:00:00", "%H:%M:%S")
> DEPT<-ifelse(df1$ATD>=breaks2 & df1$ATD<=breaks1, "0-12", "12-24")
> df1<-cbind(df1, DEPT)
> df1$DEPT<-as.factor(df1$DEPT)
> df1$Day<-as.factor(df1$Day)
> levels(df1$Day)
[1] "1" "2"
> levels(df1$Day)=c("Sunday","Monday")
> df1$FLTIME<-as.difftime(as.character(df1$FLTIME))
>
```

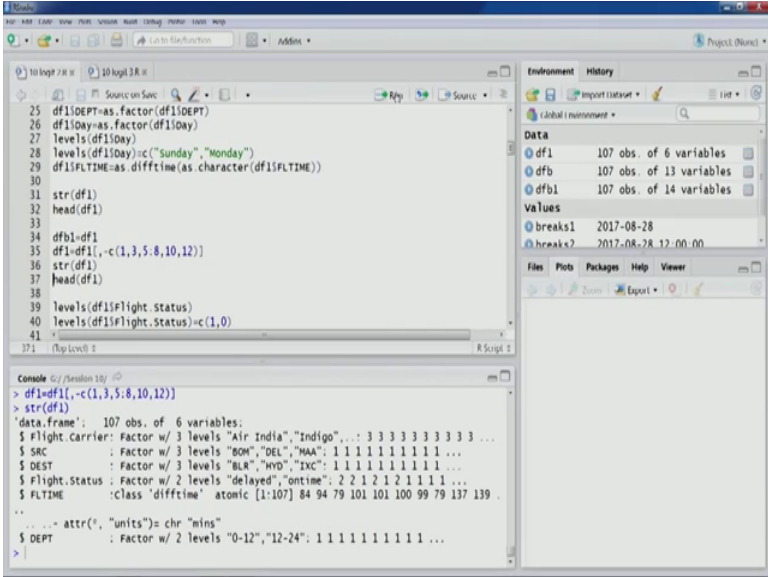
So, break is now 0 and 12, so these are the 2 breaks, 0 hours and 12 hours. Now this is how we are creating department and variable. So, if less than a breaks 2, breaks 2, and

breaks 1, then 0 to 12 so; that means, within if the timing is within these two you know hours 0 hours, and 12 hours.

Then first category that is 0 to 12 otherwise it is going to be the second category that is 12 to 14. So, let us create this variable let us append this to the data frame let us change it to a factor variable day variable as well let us cut the labels, change the labels flight time.

Let us also change it to appropriate format, now this is the structure that we have, now after taking backup we would not like to you know take forward some of the variables so let us get rid of them. Let us look at the structure again now these are the variables, now these are first few values for 6 values you can see everything is ok.

(Refer Slide Time: 03:04)



The screenshot shows the R Studio interface. The script editor on the left contains the following R code:

```
25 df1$DEPT=as.factor(df1$DEPT)
26 df1$Day=as.factor(df1$Day)
27 levels(df1$Day)
28 levels(df1$Day)=c("Sunday","Monday")
29 df1$FLTIME=as.difftime(as.character(df1$FLTIME))
30
31 str(df1)
32 head(df1)
33
34 df1=df1[,c(1,3,5,8,10,12)]
35 str(df1)
36 head(df1)
37
38 levels(df1$flight.status)
39 levels(df1$flight.status)=c(1,0)
40
41
```

The console on the bottom left shows the output of the code:

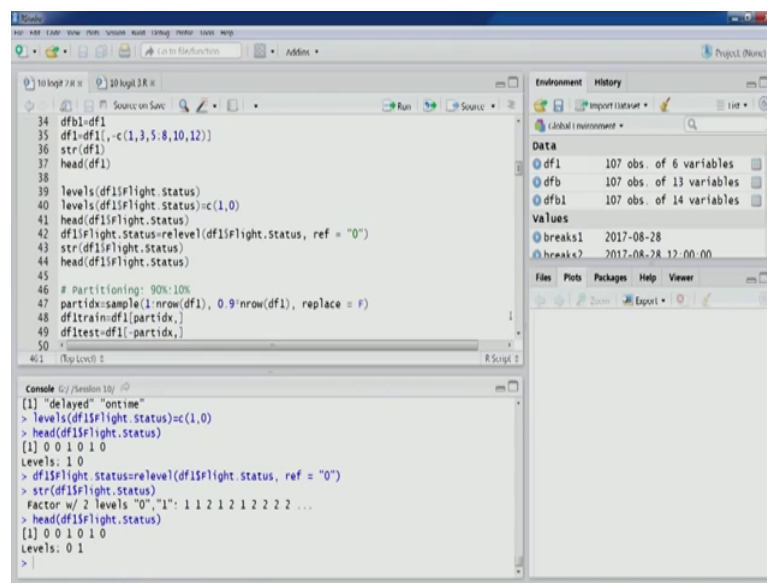
```
> df1=df1[,c(1,3,5,8,10,12)]
> str(df1)
'data.frame': 107 obs. of 6 variables:
 $ flight.carrier: Factor w/ 3 levels "Air India","Indigo",...: 3 3 3 3 3 3 3 3 3 ...
 $ SRC           : Factor w/ 3 levels "BOM","DEL","MAA": 1 1 1 1 1 1 1 1 1 ...
 $ DEST          : Factor w/ 3 levels "BLR","HYD","INC": 1 1 1 1 1 1 1 1 1 ...
 $ flight.status : Factor w/ 2 levels "delayed","ontime": 2 2 1 2 1 2 1 1 1 ...
 $ FLTIME        : class 'difftime' atomic [1:107] 84 94 79 101 101 100 99 79 137 139 ...
 ..- attr(*, "units")= chr "mins"
 $ DEPT          : Factor w/ 2 levels "0-12","12-24": 1 1 1 1 1 1 1 1 1 ...
>
```

The environment pane on the right shows the following objects:

- df1: 107 obs. of 6 variables
- dfb: 107 obs. of 13 variables
- dfb1: 107 obs. of 14 variables
- breaks1: 2017-08-28
- breaks2: 2017-08-28 12:00:00

Now, let us work on the outcome variable like we have been doing in previous lectures. So, let us change it to numeric code, let us change the reference category, now this is what it becomes now this is ok.

(Refer Slide Time: 03:30)



The screenshot shows the RStudio interface. The script editor contains the following R code:

```
34 dfb1=df1
35 df1=df1[,c(1,3,5,8,10,12)]
36 str(df1)
37 head(df1)
38
39 levels(df1$flight.status)
40 levels(df1$flight.status)=c(1,0)
41 head(df1$flight.status)
42 df1$flight.status=relevel(df1$flight.status, ref = "0")
43 str(df1$flight.status)
44 head(df1$flight.status)
45
46 # Partitioning: 90%:10%
47 partidx=sample(1:nrow(df1), 0.9*nrow(df1), replace = F)
48 df1train=df1[partidx,]
49 df1test=df1[-partidx,]
50
```

The console shows the output of the first few lines of code:

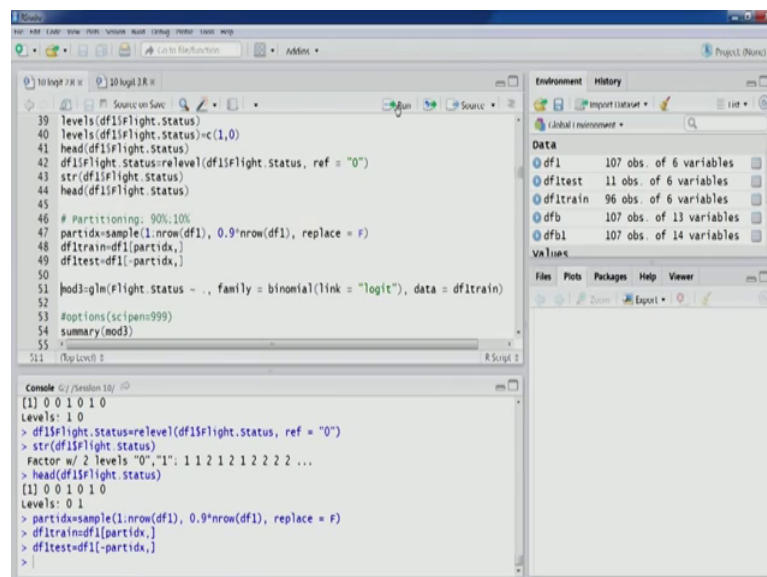
```
[1] "delayed" "ontime"
> levels(df1$flight.status)=c(1,0)
> head(df1$flight.status)
[1] 0 0 1 0 1 0
Levels: 1 0
> df1$flight.status=relevel(df1$flight.status, ref = "0")
> str(df1$flight.status)
Factor w/ 2 levels "0","1": 1 1 2 1 2 1 2 2 2 ...
> head(df1$flight.status)
[1] 0 0 1 0 1 0
Levels: 0 1
>
```

The Environment pane on the right shows the following objects:

| Object  | Class      | Attributes               |
|---------|------------|--------------------------|
| df1     | data.frame | 107 obs. of 6 variables  |
| dfb     | data.frame | 107 obs. of 13 variables |
| dfb1    | data.frame | 107 obs. of 14 variables |
| breaks1 | numeric    | 2017-08-28               |
| breaks2 | numeric    | 2017-08-28 12:00:00      |

Now, we can move ahead let us do our partitioning 2 partitions, training partition, testing partition and 90 percent for training partition and 10 of observation for test partition.

(Refer Slide Time: 03:40)



The screenshot shows the RStudio interface. The script editor contains the following R code:

```
39 levels(df1$flight.status)
40 levels(df1$flight.status)=c(1,0)
41 head(df1$flight.status)
42 df1$flight.status=relevel(df1$flight.status, ref = "0")
43 str(df1$flight.status)
44 head(df1$flight.status)
45
46 # Partitioning: 90%:10%
47 partidx=sample(1:nrow(df1), 0.9*nrow(df1), replace = F)
48 df1train=df1[partidx,]
49 df1test=df1[-partidx,]
50
51 mod3=glm(flight.status ~ ., family = binomial(link = "logit"), data = df1train)
52
53 #options(scipen=999)
54 summary(mod3)
55
```

The console shows the output of the first few lines of code:

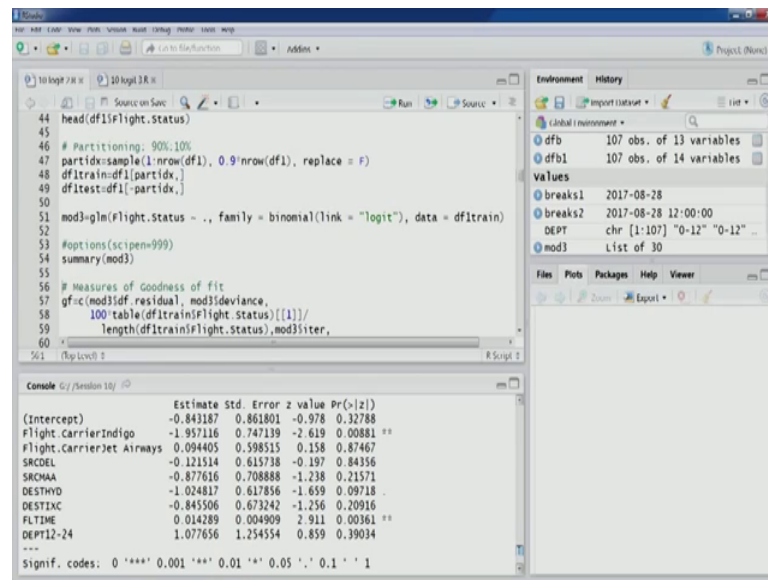
```
[1] 0 0 1 0 1 0
Levels: 1 0
> df1$flight.status=relevel(df1$flight.status, ref = "0")
> str(df1$flight.status)
Factor w/ 2 levels "0","1": 1 1 2 1 2 1 2 2 2 ...
> head(df1$flight.status)
[1] 0 0 1 0 1 0
Levels: 0 1
> partidx=sample(1:nrow(df1), 0.9*nrow(df1), replace = F)
> df1train=df1[partidx,]
> df1test=df1[-partidx,]
>
```

The Environment pane on the right shows the following objects:

| Object   | Class      | Attributes               |
|----------|------------|--------------------------|
| df1      | data.frame | 107 obs. of 6 variables  |
| df1test  | data.frame | 11 obs. of 6 variables   |
| df1train | data.frame | 96 obs. of 6 variables   |
| dfb      | data.frame | 107 obs. of 13 variables |
| dfb1     | data.frame | 107 obs. of 14 variables |

Now, the same function glm that can be used to again model this. So, these are the results.

(Refer Slide Time: 03:53)



```
44 head(df1$Flight.Status)
45
46 # Partitioning: 90%:10%
47 partidx=sample(1:nrow(df1), 0.9*nrow(df1), replace = F)
48 dfitrain=df1[partidx,]
49 dfitest=df1[-partidx,]
50
51 mod3=glm(Flight.Status ~ ., family = binomial(link = "logit"), data = dfitrain)
52
53 #options(scipen=999)
54 summary(mod3)
55
56 # Measures of Goodness of fit
57 gfc(mod3$df.residual, mod3$deviance,
58     100*table(dfitrain$Flight.Status)[[1]]/
59     length(dfitrain$Flight.Status), mod3$iter,
60     ...)
```

Console Output:

|                           | Estimate  | Std. Error | z value | Pr(> z )   |
|---------------------------|-----------|------------|---------|------------|
| (Intercept)               | -0.843187 | 0.861801   | -0.978  | 0.32788    |
| Flight.CarrierIndigo      | -1.957116 | 0.747139   | -2.619  | 0.00881 ** |
| Flight.CarrierJet Airways | 0.094405  | 0.598515   | 0.158   | 0.87467    |
| SRCDEL                    | -0.121514 | 0.615738   | -0.197  | 0.84356    |
| SRCMAA                    | -0.877616 | 0.708888   | -1.238  | 0.21571    |
| DESTHYD                   | -1.024817 | 0.617856   | -1.659  | 0.09718 .  |
| DESTLXC                   | -0.845506 | 0.673242   | -1.256  | 0.20916    |
| FLTIME                    | 0.014289  | 0.004909   | 2.911   | 0.00361 ** |
| DEPT12-24                 | 1.077656  | 1.254554   | 0.859   | 0.39034    |
| ---                       |           |            |         |            |

Signif. codes: 0 '\*\*\*' 0.001 '\*\*' 0.01 '\*' 0.05 '.' 0.1 ' ' 1

Let us look at once again now you would see every time like we have been running this particular model on the same data set and every time. The significance levels have been changing that as I have been explaining smaller data set and therefore, subject to change in terms of as we as the observation that are part of training data set training partition this results will also slightly change and mainly with respect to significance level.

Now, you can see again flight carrier indigo has become significant to star level right and we can also see that destination has also become significant right. And then we can also see flight time has also you know modes you know higher level of significance.

However, more important thing is look at the p varies we can see that this one source for madras as well this is also smaller p values this is anyway significant flight carrier and destination anyway this is significant at ninety percent this is also significant. However, the new grouping that we have created out of department departure time intervals that also not comes out to be significant; however, p value is now smaller.

(Refer Slide Time: 05:14)

```

44 head(df1$Flight.Status)
45
46 # Partitioning: 90%:10%
47 partidx=sample(1:nrow(df1), 0.9*nrow(df1), replace = F)
48 dfltrain=df1[partidx,]
49 dfltest=df1[-partidx,]
50
51 mod3=glm(Flight.Status ~ ., family = binomial(link = "logit"), data = dfltrain)
52
53 #options(scipen=999)
54 summary(mod3)
55
56 # Measures of Goodness of fit
57 gf=c(mod3$residual, mod3$deviance,
58      100*table(dfltrain$Flight.Status)[[1]]/
59      length(dfltrain$Flight.Status), mod3$iter,
60      1-(mod3$deviance/mod3$null.deviance))
61
62 gf=as.data.frame(gf, optional=T)
63 rownames(gf)=c("Residual df", "Std. Dev. Estimate",
64               "% Success in training data", "#Iterations used",
65               "Multiple R-squared")
66
67 gf

```

|           | Flight.CarrierJet Airways | 0.094405 | 0.598515 | 0.158      | 0.87467 |
|-----------|---------------------------|----------|----------|------------|---------|
| SRCDEL    | -0.121514                 | 0.615738 | -0.197   | 0.84356    |         |
| SRCMAA    | -0.877616                 | 0.708888 | -1.238   | 0.21571    |         |
| DESTHYD   | -1.024817                 | 0.617856 | -1.659   | 0.09718    |         |
| DESTTNC   | -0.845506                 | 0.673242 | -1.256   | 0.20916    |         |
| FLTIME    | 0.014289                  | 0.004909 | 2.911    | 0.00361 ** |         |
| DEPT12-24 | 1.077656                  | 1.254554 | 0.859    | 0.39034    |         |

Signif. codes: 0 '\*\*\*' 0.001 '\*\*' 0.01 '\*' 0.05 '.' 0.1 ' ' 1

(Dispersion parameter for binomial family taken to be 1)

So, with this we will discuss the important aspect of logistic regression so that is the majors of goodness of fit.

(Refer Slide Time: 05:19)

```

50
51 mod3=glm(Flight.Status ~ ., family = binomial(link = "logit"), data = dfltrain)
52
53 #options(scipen=999)
54 summary(mod3)
55
56 # Measures of Goodness of fit
57 gf=c(mod3$residual, mod3$deviance,
58      100*table(dfltrain$Flight.Status)[[1]]/
59      length(dfltrain$Flight.Status), mod3$iter,
60      1-(mod3$deviance/mod3$null.deviance))
61
62 gf=as.data.frame(gf, optional=T)
63 rownames(gf)=c("Residual df", "Std. Dev. Estimate",
64               "% Success in training data", "#Iterations used",
65               "Multiple R-squared")
66
67 gf

```

|           | Flight.CarrierJet Airways | 0.094405 | 0.598515 | 0.158      | 0.87467 |
|-----------|---------------------------|----------|----------|------------|---------|
| SRCDEL    | -0.121514                 | 0.615738 | -0.197   | 0.84356    |         |
| SRCMAA    | -0.877616                 | 0.708888 | -1.238   | 0.21571    |         |
| DESTHYD   | -1.024817                 | 0.617856 | -1.659   | 0.09718    |         |
| DESTTNC   | -0.845506                 | 0.673242 | -1.256   | 0.20916    |         |
| FLTIME    | 0.014289                  | 0.004909 | 2.911    | 0.00361 ** |         |
| DEPT12-24 | 1.077656                  | 1.254554 | 0.859    | 0.39034    |         |

Signif. codes: 0 '\*\*\*' 0.001 '\*\*' 0.01 '\*' 0.05 '.' 0.1 ' ' 1

(Dispersion parameter for binomial family taken to be 1)

So, just like the linear regression multiple linear regression logistic regression also is a statistical technique primarily. And therefore, in a statistical modelling the main objective is to fit to data as we have talked about this aspect many times before. So, therefore, in multiple linear regression we have a particular metric called a multiple R square and adjusted R square which are used to assess how good the model is fitting to the data.

So, similarly since logistic regression is also a statistical technique so there also we are required to have matrix for good to measure the goodness of it how well model is fitting to the data. So, what are those matrix so because there is certain key differences in logistic regression and linear regression. So, we will talk about some of those matrix now.

So, you can see in the code that I have created a vector here `gf` first one is `mod3` the model that we have just computed and you can see degree of residual degree of freedom. So, this `df.residual` is one of the value that is returned by and the `glm` function and gives us the residual degree of freedom then we have deviance.

So, this is again the returns the deviance value returned by the `glm` function right then few other things which are mainly for the descriptive purposes. For example, this table result which is for the outcome variable here and then divided by the full number observation that will give us percentage success in training data and then we have iterations.

So, as we talked about the particular estimation technique that is used in that is used in logistic regression is different from multiple linear regression we can look at for more details we can look at here. So, we talked about that `Emily` is used for typically for used for logistic regression; however, you can see that for example, we have been using `glm` function, and within that if we go we look up for the some of the arguments.

(Refer Slide Time: 07:48)

```

50 mod3=glm(Flight.status ~ ., family = binomial(link = "logit"), data = dftrain)
51
52 #options(scipen=999)
53 summary(mod3)
54
55 # Measures of Goodness of fit
56 gf=c(mod3$df.residual, mod3$deviance,
57       100*(table(dftrain$Flight.status)[1])/
58       length(dftrain$Flight.status), mod3$iter,
59       1-(mod3$deviance/mod3$null.deviance))
60 gf=as.data.frame(gf, optional=T)
61 rownames(gf)=c("Residual df", "Std. Dev. Estimate",
62               "% Success in training data", "#iterations used",
63               "Multiple R-squared")
64
65 gf
66
2942 (Top Level)
R Script 2

```

```

Flight.CarrierJet Airways 0.094405 0.598515 0.158 0.87467
SRCDEL -0.121514 0.615738 -0.197 0.84356
SRCHAA -0.877616 0.708888 -1.238 0.21571
DESTIND -1.024817 0.617856 -1.659 0.09718
DESTINC -0.845506 0.673242 -1.256 0.20916
FLTIME 0.014289 0.004909 2.911 0.00361 **
DEPT12-24 1.077656 1.254554 0.859 0.39034
***
Signif. codes: 0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

(Dispersion parameter for binomial family taken to be 1)

```

Environment: Global environment, Project (Blank)

Global Environment

- dfb 107 obs. of 13 variables
- dfb1 107 obs. of 14 variables

Values

- breaks1 2017-08-28
- breaks2 2017-08-28 12:00:00
- DEPT chr [1:107] "0-12" "0-12" ..
- mod3 List of 30

Files Plots Packages Help Viewer

fitting Generalized Linear Models

```

glm1=glm(mod3, family = gaussian, data = wa,
na.action = na.omit, start = NULL, etastart = NULL,
control = list(...), model = TRUE, mc = FALSE, y.trace, contrasts = NULL)
glm.fit(x, y, weights = rep(1, nobs),
start = NULL, classwt = NULL, mu.offset = rep(0, nobs), family = gaussian,
control = list(), intercept = TRUE)
## S3 method for class 'glm'
weights(object, type = c("prior", "working")
Arguments
formula an object of class "formula" (or one

```

Specifically for this purpose the estimation technique purpose will get more detail. So, we will see that glm for dot fit.

(Refer Slide Time: 08:00)

The screenshot shows the R Studio environment with a script editor, console, and environment pane. The script defines a glm model for flight status data using the binomial family and logit link. It calculates measures of goodness of fit and displays the results in the console.

```

50 mod3=glm(flight.status ~ ., family = binomial(link = "logit"), data = dfltrain)
51
52 #options(scipen=999)
53 summary(mod3)
54
55 # Measures of Goodness of fit
56 gf=c(mod3$df.residual, mod3$deviance,
57       100*(cable(dfltrain$flight.status)/length(dfltrain$flight.status))^2,
58       1-(mod3$deviance/mod3$null.deviance))
59
60 gf=as.data.frame(gf, optional=T)
61 rownames(gf)=c("Residual df", "Std. Dev. Estimate",
62               "% Success in training data", "Iterations used",
63               "Multiple R-squared")
64
65 gf
66
67
68
69
70
71
72
73
74
75
76
77
78
79
80
81
82
83
84
85
86
87
88
89
90
91
92
93
94
95
96
97
98
99

```

Console output:

```

R flight.carrierjet Airways 0.094405 0.598515 0.158 0.87467
SRDEL -0.121514 0.615738 -0.197 0.84356
SRDMAA -0.877616 0.708888 -1.238 0.21571
DETHYD -1.024817 0.617856 -1.659 0.09718
DETHXC -0.845506 0.673242 -1.256 0.20916
FLTIME 0.014289 0.004909 2.911 0.00361 **
DEPT12-24 1.077656 1.254554 0.859 0.39034
---
Signif. codes: 0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
(Dispersion parameter for binomial family taken to be 1)

```

Environment pane shows the following objects:

- dflb: 107 obs. of 13 variables
- dflb1: 107 obs. of 14 variables
- breaks1: 2017-08-28
- breaks2: 2017-08-28 12:00:00
- DEPT: chr [1:107] "0-12" "0-12" ...
- mod3: List of 30

We will see that glm dot fit method. So, this we can see iteratively re weighted least square is used the particular function that we are using glm iteratively re weighted least square function is used which is quite similar in approach with respect to Emily estimation technique that we talked about.

(Refer Slide Time: 08:41)

## LOGISTIC REGRESSION

- Linear Regression for a categorical outcome variable?
  - Can be done by treating the outcome variable as continuous and coding it numerically
  - However, anomalies will lead to spurious modeling
    - Predictions can take any value, not just dummy values {0,1}
    - Outcome variable or residuals don't follow normal distribution
      - binomial distribution
    - Variance of outcome variable is not constant across all records (violation of homoscedasticity)
      - $np(1-p)$

4

IIT ROORKEE
 NPTEL ONLINE CERTIFICATION COURSE

14

So, MLM techniques, sorry MLM maximum likelihood method that we talked about in our discussion as we can look in these slides as well maximum likelihood method MLM method that we talked about. So, this is quite similar to what this the discussion that we have.

So, in particular R's implementation glm function that we are using it is the iteratively re weighted least square that estimation technique that is used to estimate the coefficients that we have been doing quite similar to a MLM and as we talked about that number of iterations have to be performed to reach to estimate these parameters so that we can get the best model which is the model best model which is fitting the data.

So, number of iterations actually indicate that and then we have one matrix which is quite similar to what we have multiple R square and linear regression right. So, this is actually computed where using the deviance value. So, null deviance and the standard deviation estimate or deviance that we have can look at the returned values here.

So, you can see deviance is the one of the return value so this is one and then we also null deviance is also a return. So, this is null deviance is come with respect to the Knave rule. So, 1 minus this deviance divided by null deviance gives us a value which is quite similar to what we have in multiple linear regression multiple R square.

So, this particular value will be will give us a metric which can be used to understand the goodness of fit follows degradation model and as on its own deviance also can be used it is quite similar to what we have there in as I see some of squares error.

So, this is quite similar deviance is quite similar to that and then we can have one metric as I talked about similar to multiple R square. So, these metrics can be used to assess; the assess the fitness or model goodness of fitness goodness of fit of alloys degradation model.

So, let us compute some of these things so, residual degree of freedom deviance which is similar to SSC in linear regression, then we have this proportions percentage in successive training data. Then we have number of iteration and then we have a multiple R square kind of metric let us compute this.

(Refer Slide Time: 11:34)

The screenshot shows an R script being executed in RStudio. The script defines a generalized linear model (GLM) and summarizes its results. The console output shows the following values:

```

> gf=as.data.frame(gf, optional=T)
> rownames(gf)=c("Residual df", "Std. Dev. Estimate",
+               "% Success in training data", "#Iterations used",
+               "Multiple R-squared")
> gf
Residual df      87.0000000
Std. Dev. Estimate 104.2420821
% Success in training data 57.2916667
#Iterations used      5.0000000
Multiple R-squared   0.2044732

```

Let us create a data frame and row names we have given some and these are the values. So, you can see residual df df is 87 here, so as you can see let us again have a look df 1 training partition we have 96 observations.

(Refer Slide Time: 11:56)

The screenshot shows the same R script as before, but the console output now displays the coefficients of the fitted model:

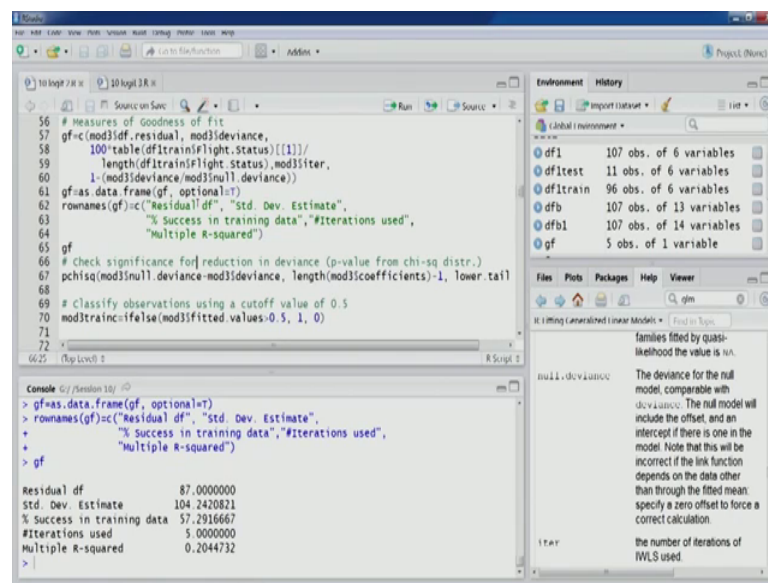
```

Coefficients:
(Intercept)      -0.843187  0.861801  -0.978  0.32788
Flight.CarrierIndigo -1.957116  0.747139  -2.619  0.00881 **
Flight.CarrierJet Airways 0.094405  0.598515  0.158  0.87467
SRCDEL           -0.121514  0.615738  -0.197  0.84356
SRCHMA           -0.877616  0.708888  -1.238  0.21571
DESTHND          -1.024817  0.617856  -1.659  0.09718 .
DESTXNC          -0.845506  0.673242  -1.256  0.20916
FLTIME           0.014289  0.004909  2.911  0.00361 **
DEPT12-24        1.077656  1.254554  0.859  0.39034

```

And if we go back to our summary results right if we go back to our summary results we can see that how many variables we have here 1, 2, 3, 4, 5, 6, 7, 8; 8 8 variables. And we can see that 87 is the residual.

(Refer Slide Time: 12:06)



The screenshot shows the R Studio environment. The script editor contains R code for fitting a generalized linear model (GLM) and calculating measures of goodness of fit. The console shows the output of the `gf` function, and the environment pane lists the objects created.

```
# Measures of Goodness of fit
56 gf=c(mod3$df.residual, mod3$deviance,
57       100*(table(dfitrain$flight.status)[{1}]/
58         length(dfitrain$flight.status))/mod3$iter,
59       1-(mod3$deviance/mod3$null.deviance))
60 gf=as.data.frame(gf, optional=T)
61 rownames(gf)=c("Residual df", "Std. Dev. Estimate",
62               "% Success in training data", "#Iterations used",
63               "Multiple R-squared")
64 gf
65 # Check significance for reduction in deviance (p-value from chi-sq distr.)
66 pchisq(mod3$null.deviance-mod3$deviance, length(mod3$coefficients)-1, lower.tail=F)
67
68 # Classify observations using a cutoff value of 0.5
69 mod3$train=ifelse(mod3$fitted.values>0.5, 1, 0)
70
71
72
```

Console Output:

```
> gf=as.data.frame(gf, optional=T)
> rownames(gf)=c("Residual df", "Std. Dev. Estimate",
+               "% Success in training data", "#Iterations used",
+               "Multiple R-squared")
> gf
      Residual df      Std. Dev. Estimate
104.2420821
% Success in training data 57.2916667
#Iterations used          5.0000000
Multiple R-squared        0.2044732
```

Environment Pane:

| Object   | Class      | Attributes               |
|----------|------------|--------------------------|
| df1      | data.frame | 107 obs. of 6 variables  |
| dfitest  | data.frame | 11 obs. of 6 variables   |
| dfitrain | data.frame | 96 obs. of 6 variables   |
| dfb      | data.frame | 107 obs. of 13 variables |
| dfb1     | data.frame | 107 obs. of 14 variables |
| gf       | data.frame | 5 obs. of 1 variable     |

So, 87 plus 7 that makes it 94 that is  $n$  minus 1, so that is the computation that is how the degrees of freedom have been computed, so this is a correct value here. And then we have deviance value which is also called as standard deviation estimated by some software's some statistical commercial statistical software's.

And then so this is all similar to what we have SSE sum of squares error in multiple linear regression then we have number of iteration that have been used to arrive at the particular model and that we have that is not 3 in this case. Then we have a value similar to multiple are squares we can see 20 percent, of the variability in the outcome variable has been explained by this model. So, all we talk about that this is being computed by 1 minus mod 3 deviance, divided by null deviance.

Now, whether so in terms of on further in terms of deviance the null deviance represents the knave rule value. So, we have to see how much our model has been able to how much reduction in deviance has been done by our model and whether that that is significant or not so that can also be that can also be performed using this chi square test that we can do.

So, we have one function p chi square. So, there we can actually use these two values or we can take a difference of null deviance and deviance so that would be the reduction in deviance from a knave rule and you know how much deduction that our model has done.

And we can look at the number of predictors as degrees of freedom I could be used the number predictors that we use have used could be used degrees of freedom because these are freedom degree of freedom that had been used to reduce the deviance as we talked about from 95 minus 1 available degrees of freedom two we have reduced up to 87 that is residual degree of freedom so 7 predictors have been used.

So, that information can be use to perform chi square test and to find out the significance of whether the reduction has been significant or not. So, the third argument is lowered tail that is specified as false and we can compute this chi square value so we can see that this is a small value. So, therefore, it seems that this redis reduction is reduction deviance is significant which are also clear by the difference between the deviance and null deviance values.

So, we can also compute that we can see this is null deviance; this is null deviance. So, let us look at the value and we can have the deviance. Let us look at the value so you can see.

(Refer Slide Time: 15:17)

The screenshot shows the R Studio environment. The script editor contains the following R code:

```

57 gf=c(mod3$df.residual, mod3$deviance,
58 100*table(dfitrain$flight.status)[[1]]/
59 length(dfitrain$flight.status),mod3$iter,
60 1/(mod3$deviance/mod3$null.deviance))
61 gf=as.data.frame(gf, optional=T)
62 rownames(gf)=c("Residual df", "Std. Dev. Estimate",
63 "% Success in training data", "#Iterations used",
64 "Multiple R-squared")
65 gf
66 # Check significance for reduction in deviance (p-value from chi-sq distr.)
67 pchisq(mod3$null.deviance-mod3$deviance, length(mod3$coefficients)-1, lower.tail=
68 F)
69 # Classify observations using a cutoff value of 0.5
70 mod3train=ifelse(mod3$fitted.values>0.5, 1, 0)
71
72 table("Actual value"=dfitrain$flight.status, "Predicted value"=mod3train)
73

```

The console output shows the following results:

```

Residual df      87.0000000
Std. Dev. estimate 104.2420821
% Success in training data 57.2916667
#Iterations used    5.0000000
Multiple R-squared  0.2044732
> pchisq(mod3$null.deviance-mod3$deviance, length(mod3$coefficients)-1, lower.tail = F)
[1] 0.000767512
> mod3$null.deviance
[1] 131.0353
> mod3$deviance
[1] 104.2421
>

```

The environment pane on the right shows the following objects:

- df1: 107 obs. of 6 variables
- dfitest: 11 obs. of 6 variables
- dfitrain: 96 obs. of 6 variables
- dfb: 107 obs. of 13 variables
- dfb1: 107 obs. of 14 variables
- gf: 5 obs. of 1 variable

The Files pane shows the following files:

- Files: R, Rplots.pdf, Packages, Help, Viewer
- Viewer: R: Fitting Generalized Linear Models

So, the difference is so there is good enough Reduction and deviance and that is why it also came as significant you know difference. So, these are some of the matrix that can actually be used to understand to a measure the goodness of it to asses goodness of it for a model as we talked about I talked about that in a statistical set setting. So, these are some of the matrix which would be more useful.

So, in a statistical modelling we stopped at when we build the model using the training partition. So, typically all the observations are used that are present and then we look at the model with respect to some of these some of these matrix. Now, let us move forward to our next discussion point in logistic regression so that is let us move forward so that is this particular point whether linear regression can be used for a categorical outcome variable right.

So, there are there are some situations where linear regression can be used as a categorical outcome variable which we will discuss later; however, right now we are discussing some of the more important points with respect to overall general applicability of linear regression for a categorical outcome variable.

So, can be done technically it can be applied so we can treat the outcome variable as continuous. So, the categorical outcome variable can be treated as continuous variable. So, we can essentially do the numeric coding and keep it as a numeric variable and technically we can apply we will get the results; however, the results would be meaningful or not that we need to understand.

So, technically it can be applied we can read the categorical outcome variable as continuous variable we can code it numerically so that can be done. However, there are going to be anomalies that would lead to spurious modelling so what could be some of these things.

So, number one predict predictions can take any value not just any values so for example, that binary logistic regression model that we have been performing for on some of the data sets so they are the our outcome variable it is it typically has two classes class 1, and class 0.

And so, the values the remaining variable we will take is to 0 for class 1 and 0 and 1 for class 1; however, when we apply a linear regression model to a categorical outcome variable the prediction can take any values any real value can be taken and not just the dummy values 0 and 1 so that is one challenge.

How do we map some of these predictors values which can which can be any real value to the actual values of the outcome variable 0 and 1. Now, outcome variable or residuals do not follow normal distribution. So, as we have discussed during a linear regression

that this is one of the important assumption that dependent variable that is outcome variable all residuals should follow normal distribution, but that is not the case as we can understand that categorical outcome variable will have just 2 values 0 and 1.

So, it is actually it actually follows a binomial distribution so this particular assumption would also be violated; however, we talked about that because we are in data mining modelling context. So, where as we talked about even if for you know prediction purposes even if this particular you know assumption is violated in terms of prediction it might not much matter much because generally check performance on validation partition and test partitions; however, this case is different.

The deviation from normal distribution is much higher it is actually different distribution binomial distribution so that is one problem. Now, the another assumptions that we talked about in multiple linear regression is homoscedasticity; however, if we apply linear regression to an a to a categorical outcome variable this particular assumption would also be violated. The variance of outcome variable that we that we expect to be constant across all the records that is the homoscedasticity property so we for to apply multiple interrogation we want our outcome variable to follow this to have this property.

So, variance should be constant across all time; however, if we look at the variance for our categorical outcome variable it is going to be this particular value  $n \times p \times (1 - p)$  and as you can see because this is dependent on the value of  $p$ . So, therefore, the variance will change as the value of  $p$  changes.

So, when the value of a probability value is actually close to 0, then the variance would be on the as lower side and when the value of  $p$  is approaches 1 then the variance would be on the higher side. So, therefore, the variance will be will not be constant and it will be it will increase as the probability value you know in case is from 0 to 1.

So, so some of these are some of the problems that we can see in that we can directly understand and why a linear regression and cannot be applied to category outcome variable in a general sense and the problems that we can see here. So, what we will do we will do an exercise in R to understand the same thing to understand this particular aspect.

So, what we will do we will apply a linear regression model on a logistic partition and see the see its applicability and see some of the anomalies or violations that are that could be there. So, for this purpose we are going to use as you can see multiple here the comment is multiple linear regression model for a categorical response. So, promotion offers is the data set that we are going to use for this particular exercise. So, let us import the data set let us edit load.

(Refer Slide Time: 22:16)

```

130 segments(1,0,nrow(df1lft),df1lft$cumActualClass2[nrow(df1lft)], lty = 3)
131 legend(5, 2, inset = 0.005,
132       c("Cumulative Flight Status when sorted using predicted values",
133         "Cumulative Flight Status using average"),
134       lty = c(1,2), bty = "n", cex = 0.7, x.intersp = 0.3, y.intersp = 0.5)
135
136 # Multiple Linear Regression Model for a categorical response
137 # PromotionOffers.xlsx
138 df2<-read.xlsx(file.choose(), 1, header = T)
139 df2<-df2[, !apply(is.na(df2), 2, all)]
140 df2<-df2[!apply(is.na(df2), 2, all),]
141 str(df2)
142
143 dfb2<-df2
144 df2<-df2[, c(1,3,7,9)]
145 str(df2)
146
147
148
149
150
151
152
153
154
155
156
157
158
159
160
161
162
163
164
165
166
167
168
169
170
171
172
173
174
175
176
177
178
179
180
181
182
183
184
185
186
187
188
189
190
191
192
193
194
195
196
197
198
199
200
201
202
203
204
205
206
207
208
209
210
211
212
213
214
215
216
217
218
219
220
221
222
223
224
225
226
227
228
229
230
231
232
233
234
235
236
237
238
239
240
241
242
243
244
245
246
247
248
249
250
251
252
253
254
255
256
257
258
259
260
261
262
263
264
265
266
267
268
269
270
271
272
273
274
275
276
277
278
279
280
281
282
283
284
285
286
287
288
289
290
291
292
293
294
295
296
297
298
299
300
301
302
303
304
305
306
307
308
309
310
311
312
313
314
315
316
317
318
319
320
321
322
323
324
325
326
327
328
329
330
331
332
333
334
335
336
337
338
339
340
341
342
343
344
345
346
347
348
349
350
351
352
353
354
355
356
357
358
359
360
361
362
363
364
365
366
367
368
369
370
371
372
373
374
375
376
377
378
379
380
381
382
383
384
385
386
387
388
389
390
391
392
393
394
395
396
397
398
399
400
401
402
403
404
405
406
407
408
409
410
411
412
413
414
415
416
417
418
419
420
421
422
423
424
425
426
427
428
429
430
431
432
433
434
435
436
437
438
439
440
441
442
443
444
445
446
447
448
449
450
451
452
453
454
455
456
457
458
459
460
461
462
463
464
465
466
467
468
469
470
471
472
473
474
475
476
477
478
479
480
481
482
483
484
485
486
487
488
489
490
491
492
493
494
495
496
497
498
499
500
501
502
503
504
505
506
507
508
509
510
511
512
513
514
515
516
517
518
519
520
521
522
523
524
525
526
527
528
529
530
531
532
533
534
535
536
537
538
539
540
541
542
543
544
545
546
547
548
549
550
551
552
553
554
555
556
557
558
559
560
561
562
563
564
565
566
567
568
569
570
571
572
573
574
575
576
577
578
579
580
581
582
583
584
585
586
587
588
589
590
591
592
593
594
595
596
597
598
599
600
601
602
603
604
605
606
607
608
609
610
611
612
613
614
615
616
617
618
619
620
621
622
623
624
625
626
627
628
629
630
631
632
633
634
635
636
637
638
639
640
641
642
643
644
645
646
647
648
649
650
651
652
653
654
655
656
657
658
659
660
661
662
663
664
665
666
667
668
669
670
671
672
673
674
675
676
677
678
679
680
681
682
683
684
685
686
687
688
689
690
691
692
693
694
695
696
697
698
699
700
701
702
703
704
705
706
707
708
709
710
711
712
713
714
715
716
717
718
719
720
721
722
723
724
725
726
727
728
729
730
731
732
733
734
735
736
737
738
739
740
741
742
743
744
745
746
747
748
749
750
751
752
753
754
755
756
757
758
759
760
761
762
763
764
765
766
767
768
769
770
771
772
773
774
775
776
777
778
779
780
781
782
783
784
785
786
787
788
789
790
791
792
793
794
795
796
797
798
799
800
801
802
803
804
805
806
807
808
809
810
811
812
813
814
815
816
817
818
819
820
821
822
823
824
825
826
827
828
829
830
831
832
833
834
835
836
837
838
839
840
841
842
843
844
845
846
847
848
849
850
851
852
853
854
855
856
857
858
859
860
861
862
863
864
865
866
867
868
869
870
871
872
873
874
875
876
877
878
879
880
881
882
883
884
885
886
887
888
889
890
891
892
893
894
895
896
897
898
899
900
901
902
903
904
905
906
907
908
909
910
911
912
913
914
915
916
917
918
919
920
921
922
923
924
925
926
927
928
929
930
931
932
933
934
935
936
937
938
939
940
941
942
943
944
945
946
947
948
949
950
951
952
953
954
955
956
957
958
959
960
961
962
963
964
965
966
967
968
969
970
971
972
973
974
975
976
977
978
979
980
981
982
983
984
985
986
987
988
989
990
991
992
993
994
995
996
997
998
999
1000

```

Environment

| Object   | Class      | Attributes               |
|----------|------------|--------------------------|
| df1      | data.frame | 107 obs. of 6 variables  |
| df1test  | data.frame | 11 obs. of 6 variables   |
| df1train | data.frame | 96 obs. of 6 variables   |
| dfb      | data.frame | 107 obs. of 13 variables |
| dfb1     | data.frame | 107 obs. of 14 variables |
| gf       | data.frame | 5 obs. of 1 variable     |

Console

```

> pchisq(mod3$null.deviance-mod3$deviance, length(mod3$coefficients)-1, lower.tail = F)
[1] 0.000767512
> mod3$null.deviance
[1] 131.0353
> mod3$deviance
[1] 104.2421
> df2<-read.xlsx(file.choose(), 1, header = T)

```

So, once the observations have been loaded into environment we as we will see in the environment section we will do some of these steps I think it is has been loaded yes. So, df 2 we can see 5000 observation, 9 variables so let us remove any columns, or any rows if there are any let us look at the structure so this is the data set. So, we are already familiar with this.

(Refer Slide Time: 22:41)

```

131 legend(5, 2, inset = 0.005,
132       c("Cumulative Flight Status when sorted using predicted values",
133         "Cumulative Flight Status using average"),
134       lty = c(1,2), bty = "n", cex = 0.7, x.intersp = 0.3, y.intersp = 0.5)
135
136 # Multiple Linear Regression Model for a categorical response
137 # Promoters.xlsx
138 df2<-read.xlsx(file.choose(), 1, header = T)
139 df2<-df2[, !apply(is.na(df2), 2, all)]
140 df2<-df2[!apply(is.na(df2), 2, all),]
141 str(df2)
142
143 dfb2<-df2
144 df2<-df2[, c(1,3,7,9)]
145 str(df2)
146 #df2$Promoter<-as.factor(df2$Promoter)
147
148

```

Console: C:/Rstudio10y

```

> str(df2)
'data.frame': 5000 obs. of 9 variables:
 $ Income      : num 49 35 10 101 45 31 71 23 80 182 ...
 $ Spending    : num 1.6 2.201 0.495 2.73 1 ...
 $ Promoter    : num 0 0 0 0 0 0 0 0 1 ...
 $ Age         : num 25 45 39 35 35 37 53 50 35 34 ...
 $ PIN.Code    : num 110057 110092 110036 110095 110081 ...
 $ Experience  : num 1 19 15 9 8 13 27 24 10 9 ...
 $ Family.Size : num 4 3 1 1 4 4 2 1 3 1 ...
 $ Education   : factor w/ 3 levels "Grad","Hsc","PostGrad": 2 2 1 1 1 1 3 1 3 ...
 $ Online      : num 0 0 0 0 0 1 1 0 1 0 ...

```

Environment: History

- df1 107 obs. of 6 variables
- df1test 11 obs. of 6 variables
- df1train 96 obs. of 6 variables
- df2 5000 obs. of 9 variables
- dfb 107 obs. of 13 variables
- dfb1 107 obs. of 14 variables

Files: Plots: Packages: Help: Viewer

Null deviance

families fitted by quasi-likelihood the value is NA.

The deviance for the null model, comparable with deviance. The null model will include the offset, and an intercept if there is one in the model. Note that this will be incorrect if the link function depends on the data other than through the fitted mean. Specify a zero offset to force a correct calculation.

Iter

the number of iterations of IWLS used.

So, let us go through some of the some of the competitions some of the transformations that we have been doing you know in previous lectures as well take a backup let us get rid of let us select just these 4 variables that is income promotional offer of our outcome variable then we have I think family size and online activity. So, let us select them. So, these are the variable selected for this particular exercise income promotion offer and family size and online.

(Refer Slide Time: 23:12)

```

144 df2<-df2[, c(1,3,7,9)]
145 str(df2)
146 #df2$Promoter<-as.factor(df2$Promoter)
147 #df2$Online<-as.factor(df2$Online)
148 #str(df2)
149
150 # Partitioning: Tr:Te->60:40%
151 partidx=sample(1:nrow(df2), 0.6*nrow(df2), replace = F)
152 df2train=df2[partidx,]
153
154 mod4<-lm(Promoter ~ ., df2train)
155
156 pop=options()
157 options(scipen=999)
158 |
159 mod4sum<-summary(mod4); mod4sum
160
161

```

Console: C:/Rstudio10y

```

> dfb2<-df2
> df2<-df2[, c(1,3,7,9)]
> str(df2)
'data.frame': 5000 obs. of 4 variables:
 $ Income      : num 49 35 10 101 45 31 71 23 80 182 ...
 $ Promoter    : num 0 0 0 0 0 0 0 0 1 ...
 $ Family.Size : num 4 3 1 1 4 4 2 1 3 1 ...
 $ Online      : num 0 0 0 0 0 1 1 0 1 0 ...
> partidx=sample(1:nrow(df2), 0.6*nrow(df2), replace = F)
> df2train=df2[partidx,]
> mod4<-lm(Promoter ~ ., df2train)

```

Environment: History

- gf 5 obs. of 1 variable
- values
- breaks1 2017-08-28
- breaks2 2017-08-28 12:00:00
- DEPT chr [1:107] "0-12" "0-12" ..
- mod3 List of 30
- mod4 Large lm (12 elements, 801 ..

Files: Plots: Packages: Help: Viewer

Null deviance

families fitted by quasi-likelihood the value is NA.

The deviance for the null model, comparable with deviance. The null model will include the offset, and an intercept if there is one in the model. Note that this will be incorrect if the link function depends on the data other than through the fitted mean. Specify a zero offset to force a correct calculation.

Iter

the number of iterations of IWLS used.

So, the variable that we are going to use for our outcome variable is going to be the promotional offer as you can see that we have commented out the lines of code, which we used earlier to convert these numeric variables going to convert numeric variables into the categorical variable. So, promotional offer an online so they are actually categorical variable factor variable, but we are not converting them into factor variable because we are going to apply a linear regression modelling. So, we will give them as numeric and we will apply.

So, a partitioning is the same 60 percent, 40 percent in this case so we can see that let us do partitioning. So, df 2 train you can see 3000 observations, 4 variables. Now the same align function is going to be used now the promotional offer is going to be request against all the predictors that are present in this particular dataset df 2 train. So, let us run this.

(Refer Slide Time: 24:30)

```

148 #str(df2)
149
150 # Partitioning: Tr:Te->60%:40%
151 partidx=sample(1:nrow(df2), 0.6*nrow(df2), replace = F)
152 df2train=df2[partidx,]
153
154 mod4=lm(Promoffer ~ ., df2train)
155
156 #op=options()
157 #options(scipen=999)
158
159 mod4sum=sumary(mod4); mod4sum
160 mod4ava=anova(mod4); mod4ava
161
162 DF=c(mod4sumfstatistic[["r.squared"]],
163      mod4sumfstatistic[["dendf"]],
164      )
165
166 (Top Level) 5
  
```

Console Output:

|  | Min      | 1Q       | Median   | 3Q      | Max     |
|--|----------|----------|----------|---------|---------|
|  | -0.58172 | -0.13744 | -0.03657 | 0.05929 | 0.92044 |

Coefficients:

|             | Estimate   | Std. Error | t value | Pr(> t )   |
|-------------|------------|------------|---------|------------|
| (Intercept) | -0.2388135 | 0.0148069  | -16.129 | <2e-16 *** |
| Income      | 0.0033865  | 0.0001005  | 33.706  | <2e-16 *** |
| Family.size | 0.0372885  | 0.0040443  | 9.220   | <2e-16 *** |
| Online      | -0.0058494 | 0.0093219  | -0.627  | 0.53       |

Signif. codes: 0 '\*\*\*' 0.001 '\*\*' 0.01 '\*' 0.05 '.' 0.1 ' ' 1

Environment:

- breaks1: 2017-08-28
- breaks2: 2017-08-28 12:00:00
- DEPT: chr [1:107] "0-12" "0-12" ..
- mod3: List of 30
- mod4: Large lm (12 elements, 801)
- mod4sum: List of 11

Now, what we will do we will look at the summary table. Let us look at the results. So, as we can see that income is intercepted significant income is significant and family size is significant we can look at the different estimates for example, income quite a small value from this family size is also 0.03, online is not significant it has not only found to significant as we can see.

(Refer Slide Time: 24:55)

The screenshot shows the RStudio interface. The script editor contains the following code:

```

148 #str(df2)
149
150 # Partitioning: Tr:Te->60%:40%
151 partidx=sample(1:nrow(df2), 0.6*nrow(df2), replace = F)
152 df2train=df2[partidx,]
153
154 mod4=lm(Promoffer ~ ., df2train)
155
156 #options()
157 #options(scipen=999)
158
159 mod4sum=summary(mod4); mod4sum
160 mod4ava=anova(mod4); mod4ava
161
162 DF=c(mod4sum$dfstatistic[["numdf"]],
163      mod4sum$dfstatistic[["dendf"]]),
164
165
166
167
168
169
170
171
172
173
174
175
176
177
178
179
180
181
182
183
184
185
186
187
188
189
190
191
192
193
194
195
196
197
198
199
200
201
202
203
204
205
206
207
208
209
210
211
212
213
214
215
216
217
218
219
220
221
222
223
224
225
226
227
228
229
230
231
232
233
234
235
236
237
238
239
240
241
242
243
244
245
246
247
248
249
250
251
252
253
254
255
256
257
258
259
260
261
262
263
264
265
266
267
268
269
270
271
272
273
274
275
276
277
278
279
280
281
282
283
284
285
286
287
288
289
290
291
292
293
294
295
296
297
298
299
300
301
302
303
304
305
306
307
308
309
310
311
312
313
314
315
316
317
318
319
320
321
322
323
324
325
326
327
328
329
330
331
332
333
334
335
336
337
338
339
340
341
342
343
344
345
346
347
348
349
350
351
352
353
354
355
356
357
358
359
360
361
362
363
364
365
366
367
368
369
370
371
372
373
374
375
376
377
378
379
380
381
382
383
384
385
386
387
388
389
390
391
392
393
394
395
396
397
398
399
400
401
402
403
404
405
406
407
408
409
410
411
412
413
414
415
416
417
418
419
420
421
422
423
424
425
426
427
428
429
430
431
432
433
434
435
436
437
438
439
440
441
442
443
444
445
446
447
448
449
450
451
452
453
454
455
456
457
458
459
460
461
462
463
464
465
466
467
468
469
470
471
472
473
474
475
476
477
478
479
480
481
482
483
484
485
486
487
488
489
490
491
492
493
494
495
496
497
498
499
500
501
502
503
504
505
506
507
508
509
510
511
512
513
514
515
516
517
518
519
520
521
522
523
524
525
526
527
528
529
530
531
532
533
534
535
536
537
538
539
540
541
542
543
544
545
546
547
548
549
550
551
552
553
554
555
556
557
558
559
560
561
562
563
564
565
566
567
568
569
570
571
572
573
574
575
576
577
578
579
580
581
582
583
584
585
586
587
588
589
590
591
592
593
594
595
596
597
598
599
600
601
602
603
604
605
606
607
608
609
610
611
612
613
614
615
616
617
618
619
620
621
622
623
624
625
626
627
628
629
630
631
632
633
634
635
636
637
638
639
640
641
642
643
644
645
646
647
648
649
650
651
652
653
654
655
656
657
658
659
660
661
662
663
664
665
666
667
668
669
670
671
672
673
674
675
676
677
678
679
680
681
682
683
684
685
686
687
688
689
690
691
692
693
694
695
696
697
698
699
700
701
702
703
704
705
706
707
708
709
710
711
712
713
714
715
716
717
718
719
720
721
722
723
724
725
726
727
728
729
730
731
732
733
734
735
736
737
738
739
740
741
742
743
744
745
746
747
748
749
750
751
752
753
754
755
756
757
758
759
760
761
762
763
764
765
766
767
768
769
770
771
772
773
774
775
776
777
778
779
780
781
782
783
784
785
786
787
788
789
790
791
792
793
794
795
796
797
798
799
800
801
802
803
804
805
806
807
808
809
810
811
812
813
814
815
816
817
818
819
820
821
822
823
824
825
826
827
828
829
830
831
832
833
834
835
836
837
838
839
840
841
842
843
844
845
846
847
848
849
850
851
852
853
854
855
856
857
858
859
860
861
862
863
864
865
866
867
868
869
870
871
872
873
874
875
876
877
878
879
880
881
882
883
884
885
886
887
888
889
890
891
892
893
894
895
896
897
898
899
900
901
902
903
904
905
906
907
908
909
910
911
912
913
914
915
916
917
918
919
920
921
922
923
924
925
926
927
928
929
930
931
932
933
934
935
936
937
938
939
940
941
942
943
944
945
946
947
948
949
950
951
952
953
954
955
956
957
958
959
960
961
962
963
964
965
966
967
968
969
970
971
972
973
974
975
976
977
978
979
980
981
982
983
984
985
986
987
988
989
990
991
992
993
994
995
996
997
998
999
1000

```

The console output shows the following results:

```

lmfit: 0.2388135 0.0148069 -16.129 <2e-16 ***
Income 0.0033865 0.0001005 33.706 <2e-16 ***
Family.size 0.0372885 0.0040443 9.220 <2e-16 ***
Online -0.0058494 0.0093219 -0.627 0.53
---
Signif. codes: 0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

Residual standard error: 0.251 on 2996 degrees of freedom
Multiple R-squared: 0.2774, Adjusted R-squared: 0.2767
F-statistic: 383.4 on 3 and 2996 DF, p-value: < 2.2e-16

```

The Environment pane shows the following objects:

- breaks1: 2017-08-28
- breaks2: 2017-08-28 12:00:00
- DEPT: chr [1:107] "0-12" "0-12" ..
- mod3: List of 30
- mod4: Large lm (12 elements, 801)
- mod4sum: List of 11

And we can have a look at the other values we can see adjusted R square R square and multiple R square. So, we can see if we you know 27 percent multiple R square value we look at the p well it is quite as small. So, the model is significant ah; however, as we talked about certain problems as we have discussed could be there certain anomalies could be there.

(Refer Slide Time: 25:21)

The screenshot shows the RStudio interface. The script editor contains the following code:

```

158 mod4sum=summary(mod4); mod4sum
159 mod4ava=anova(mod4); mod4ava
160
161
162 DF=c(mod4sum$dfstatistic[["numdf"]],
163      mod4sum$dfstatistic[["dendf"]]),
164
165
166
167
168
169
170
171
172
173
174
175
176
177
178
179
180
181
182
183
184
185
186
187
188
189
190
191
192
193
194
195
196
197
198
199
200
201
202
203
204
205
206
207
208
209
210
211
212
213
214
215
216
217
218
219
220
221
222
223
224
225
226
227
228
229
230
231
232
233
234
235
236
237
238
239
240
241
242
243
244
245
246
247
248
249
250
251
252
253
254
255
256
257
258
259
260
261
262
263
264
265
266
267
268
269
270
271
272
273
274
275
276
277
278
279
280
281
282
283
284
285
286
287
288
289
290
291
292
293
294
295
296
297
298
299
300
301
302
303
304
305
306
307
308
309
310
311
312
313
314
315
316
317
318
319
320
321
322
323
324
325
326
327
328
329
330
331
332
333
334
335
336
337
338
339
340
341
342
343
344
345
346
347
348
349
350
351
352
353
354
355
356
357
358
359
360
361
362
363
364
365
366
367
368
369
370
371
372
373
374
375
376
377
378
379
380
381
382
383
384
385
386
387
388
389
390
391
392
393
394
395
396
397
398
399
400
401
402
403
404
405
406
407
408
409
410
411
412
413
414
415
416
417
418
419
420
421
422
423
424
425
426
427
428
429
430
431
432
433
434
435
436
437
438
439
440
441
442
443
444
445
446
447
448
449
450
451
452
453
454
455
456
457
458
459
460
461
462
463
464
465
466
467
468
469
470
471
472
473
474
475
476
477
478
479
480
481
482
483
484
485
486
487
488
489
490
491
492
493
494
495
496
497
498
499
500
501
502
503
504
505
506
507
508
509
510
511
512
513
514
515
516
517
518
519
520
521
522
523
524
525
526
527
528
529
530
531
532
533
534
535
536
537
538
539
540
541
542
543
544
545
546
547
548
549
550
551
552
553
554
555
556
557
558
559
560
561
562
563
564
565
566
567
568
569
570
571
572
573
574
575
576
577
578
579
580
581
582
583
584
585
586
587
588
589
590
591
592
593
594
595
596
597
598
599
600
601
602
603
604
605
606
607
608
609
610
611
612
613
614
615
616
617
618
619
620
621
622
623
624
625
626
627
628
629
630
631
632
633
634
635
636
637
638
639
640
641
642
643
644
645
646
647
648
649
650
651
652
653
654
655
656
657
658
659
660
661
662
663
664
665
666
667
668
669
670
671
672
673
674
675
676
677
678
679
680
681
682
683
684
685
686
687
688
689
690
691
692
693
694
695
696
697
698
699
700
701
702
703
704
705
706
707
708
709
710
711
712
713
714
715
716
717
718
719
720
721
722
723
724
725
726
727
728
729
730
731
732
733
734
735
736
737
738
739
740
741
742
743
744
745
746
747
748
749
750
751
752
753
754
755
756
757
758
759
760
761
762
763
764
765
766
767
768
769
770
771
772
773
774
775
776
777
778
779
780
781
782
783
784
785
786
787
788
789
790
791
792
793
794
795
796
797
798
799
800
801
802
803
804
805
806
807
808
809
810
811
812
813
814
815
816
817
818
819
820
821
822
823
824
825
826
827
828
829
830
831
832
833
834
835
836
837
838
839
840
841
842
843
844
845
846
847
848
849
850
851
852
853
854
855
856
857
858
859
860
861
862
863
864
865
866
867
868
869
870
871
872
873
874
875
876
877
878
879
880
881
882
883
884
885
886
887
888
889
890
891
892
893
894
895
896
897
898
899
900
901
902
903
904
905
906
907
908
909
910
911
912
913
914
915
916
917
918
919
920
921
922
923
924
925
926
927
928
929
930
931
932
933
934
935
936
937
938
939
940
941
942
943
944
945
946
947
948
949
950
951
952
953
954
955
956
957
958
959
960
961
962
963
964
965
966
967
968
969
970
971
972
973
974
975
976
977
978
979
980
981
982
983
984
985
986
987
988
989
990
991
992
993
994
995
996
997
998
999
1000

```

The console output shows the following results:

```

lmfit: 0.2388135 0.0148069 -16.129 <2e-16 ***
Income 0.0033865 0.0001005 33.706 <2e-16 ***
Family.size 0.0372885 0.0040443 9.220 <2e-16 ***
Online -0.0058494 0.0093219 -0.627 0.53
---
Signif. codes: 0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

Residual standard error: 0.251 on 2996 degrees of freedom
Multiple R-squared: 0.2774, Adjusted R-squared: 0.2767
F-statistic: 383.4 on 3 and 2996 DF, p-value: < 2.2e-16

```

The Environment pane shows the following objects:

- breaks1: 2017-08-28
- breaks2: 2017-08-28 12:00:00
- DEPT: chr [1:107] "0-12" "0-12" ..
- mod3: List of 30
- mod4: Large lm (12 elements, 801)
- mod4sum: List of 11

So, what we will do? We will look at look at computing some of those things to check whether those anomalies are there in this particular case. So, what we will do we are

running and over to extract some of the parameters here. So, we can see so sum of square errors mean of square errors and F value of a statistic and the probability values for these predictors income family size is given there in the ANOVA table.

So, what we are going to do is we are going to compute these particular values in a in a format that can be used for tabular presentation later on. So, mod 4 so on; we have a F statistics value. So, that is written as part of the summary function of models.

(Refer Slide Time: 26:19)

```

158 mod4sum<-summary(mod4); mod4sum
159 mod4ava<-anova(mod4); mod4ava
160
161 DF<-c(mod4sum$dfstatistic[["numdf"]],
162       mod4sum$dfstatistic[["dendf"]],
163       mod4sum$dfstatistic[["numdf"]]+mod4sum$dfstatistic[["dendf"]])
164
165 SS<-c(sum(head(mod4ava[, "Sum Sq"],-1)),
166      mod4ava[["Residuals"], "Sum Sq"],
167      sum(mod4ava[, "Sum Sq"]))
168
169 MS<-c(mean(head(mod4ava[, "Mean Sq"],-1)),
170      mod4ava[["Residuals"], "Mean Sq"], "")
171
172 FSTATISTIC<-mod4sum$FSTATISTIC[["value"]] == ""
173
174
175
176
177
178
179
180
181
182
183
184
185
186
187
188
189
190
191
192
193
194
195
196
197
198
199
200
201
202
203
204
205
206
207
208
209
210
211
212
213
214
215
216
217
218
219
220
221
222
223
224
225
226
227
228
229
230
231
232
233
234
235
236
237
238
239
240
241
242
243
244
245
246
247
248
249
250
251
252
253
254
255
256
257
258
259
260
261
262
263
264
265
266
267
268
269
270
271
272
273
274
275
276
277
278
279
280
281
282
283
284
285
286
287
288
289
290
291
292
293
294
295
296
297
298
299
300
301
302
303
304
305
306
307
308
309
310
311
312
313
314
315
316
317
318
319
320
321
322
323
324
325
326
327
328
329
330
331
332
333
334
335
336
337
338
339
340
341
342
343
344
345
346
347
348
349
350
351
352
353
354
355
356
357
358
359
360
361
362
363
364
365
366
367
368
369
370
371
372
373
374
375
376
377
378
379
380
381
382
383
384
385
386
387
388
389
390
391
392
393
394
395
396
397
398
399
400
401
402
403
404
405
406
407
408
409
410
411
412
413
414
415
416
417
418
419
420
421
422
423
424
425
426
427
428
429
430
431
432
433
434
435
436
437
438
439
440
441
442
443
444
445
446
447
448
449
450
451
452
453
454
455
456
457
458
459
460
461
462
463
464
465
466
467
468
469
470
471
472
473
474
475
476
477
478
479
480
481
482
483
484
485
486
487
488
489
490
491
492
493
494
495
496
497
498
499
500
501
502
503
504
505
506
507
508
509
510
511
512
513
514
515
516
517
518
519
520
521
522
523
524
525
526
527
528
529
530
531
532
533
534
535
536
537
538
539
540
541
542
543
544
545
546
547
548
549
550
551
552
553
554
555
556
557
558
559
560
561
562
563
564
565
566
567
568
569
570
571
572
573
574
575
576
577
578
579
580
581
582
583
584
585
586
587
588
589
590
591
592
593
594
595
596
597
598
599
600
601
602
603
604
605
606
607
608
609
610
611
612
613
614
615
616
617
618
619
620
621
622
623
624
625
626
627
628
629
630
631
632
633
634
635
636
637
638
639
640
641
642
643
644
645
646
647
648
649
650
651
652
653
654
655
656
657
658
659
660
661
662
663
664
665
666
667
668
669
670
671
672
673
674
675
676
677
678
679
680
681
682
683
684
685
686
687
688
689
690
691
692
693
694
695
696
697
698
699
700
701
702
703
704
705
706
707
708
709
710
711
712
713
714
715
716
717
718
719
720
721
722
723
724
725
726
727
728
729
730
731
732
733
734
735
736
737
738
739
740
741
742
743
744
745
746
747
748
749
750
751
752
753
754
755
756
757
758
759
760
761
762
763
764
765
766
767
768
769
770
771
772
773
774
775
776
777
778
779
780
781
782
783
784
785
786
787
788
789
790
791
792
793
794
795
796
797
798
799
800
801
802
803
804
805
806
807
808
809
810
811
812
813
814
815
816
817
818
819
820
821
822
823
824
825
826
827
828
829
830
831
832
833
834
835
836
837
838
839
840
841
842
843
844
845
846
847
848
849
850
851
852
853
854
855
856
857
858
859
860
861
862
863
864
865
866
867
868
869
870
871
872
873
874
875
876
877
878
879
880
881
882
883
884
885
886
887
888
889
890
891
892
893
894
895
896
897
898
899
900
901
902
903
904
905
906
907
908
909
910
911
912
913
914
915
916
917
918
919
920
921
922
923
924
925
926
927
928
929
930
931
932
933
934
935
936
937
938
939
940
941
942
943
944
945
946
947
948
949
950
951
952
953
954
955
956
957
958
959
960
961
962
963
964
965
966
967
968
969
970
971
972
973
974
975
976
977
978
979
980
981
982
983
984
985
986
987
988
989
990
991
992
993
994
995
996
997
998
999
1000

```

```

Response: Promoffer
Income      1  67.073  67.073 1064.8362 <2e-16 ***
Family size  1   5.345   5.345  84.8536 <2e-16 ***
Online      1   0.025   0.025   0.3938 0.5304
Residuals 2996 188.716   0.063
---
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

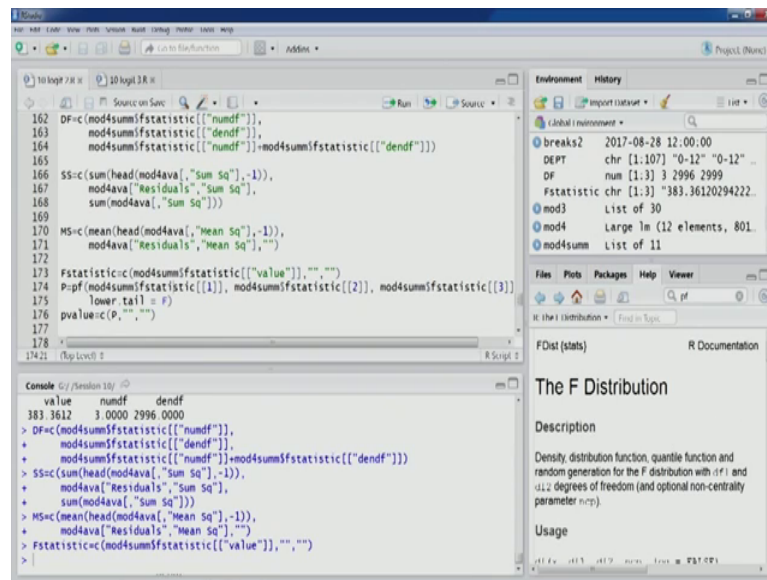
> mod4sum$FSTATISTIC
      value numdf dendf
383.3612   3.0000 2996.0000

```

So, when we apply summery on the model object so you can see this is nothing, but F statistic value for the model and then we have degree of freedom and here. So, we can see the degree of freedom residuals degree of freedom and the residuals degree of freedom and the regression degree of freedom here so that is going to be stored in this so in this data frame.

First we have the regression degree of freedom, then we have the residual degree of freedom and then we have the total so this data frame is about degree of freedom as is sustained by its named DF. Then we will compute some other square error so first we can see that in the ANOVA table, from the ANOVA table we are trying to extract out this these values.

(Refer Slide Time: 27:21)



So, first some other square for the you know regressors regression and then for the residuals, and then total. So, this would be recorded in this particular variable essence. And then we look at the mean or mean of square errors, so that is also being extracted from the ANOVA table results, you can see this particular column mean square and these values are being instructed.

So, first 3 values so this head function as you can see now this role of this head function is quite different and both these computations we have used head function.

So, you can see the second argument is minus 1 so what it actually does is it gives us all the values except the last value in the vector. So, for example, mean square or some other square. So, there are 4 values and these 2 columns, and these 2 vectors.

So, except the last values that is corresponding to residuals the first 3 values are going to be written; that means, last value will be left out and the remaining and remaining  $n$  minus 1 values are going to be written. So, that is what we want. So, that will give us the corresponding sum of squares or mean sum of squares are for the deviation.

So, first that then residuals so let us compute this. Then let us also extract the F statistics, from the this particular vector that we have already seen. So, let us do that then what we are trying to do we are trying to compute the corresponding probability value.

So, probability value corresponding to the F for test so that that is how it is being committed pf is the function that can be used for more detail you can go into the help section and find out more information about pf. So, you can see this is F distribution.

(Refer Slide Time: 29:09)

The screenshot shows the R Studio environment. The script editor contains the following code:

```

167 mod4ava["Residuals", "Sum Sq",
168 sum(mod4ava[, "Sum Sq"])]
169
170 MS=c(mean(head(mod4ava[, "Mean Sq"], -1)),
171 mod4ava["Residuals", "Mean Sq"], "")
172
173 Fstatistic=c(mod4sumsfstatistic[["value"]], "", "")
174 P=pf(mod4sumsfstatistic[[1]], mod4sumsfstatistic[[2]], mod4sumsfstatistic[[3]],
175 lower.tail = F)
176 pvalue=c(P, "", "")
177
178 anovadf=data.frame(Df, SS, MS, Fstatistic, pvalue)
179 rownames(anovadf)=c("Regression", "error", "Total")
180 anovadf
181
182 # Prediction for a new observation:
183
184
185

```

The console shows the execution of the script, including the calculation of the F-statistic and the resulting data frame:

```

> mod4sumsfstatistic[["numdf"]]=mod4sumsfstatistic[["dendf"]]
> SS=c(sum(head(mod4ava[, "Sum Sq"], -1)),
+ mod4ava["Residuals", "Sum Sq"],
+ sum(mod4ava[, "Sum Sq"])]
> MS=c(mean(head(mod4ava[, "Mean Sq"], -1)),
+ mod4ava["Residuals", "Mean Sq"], "")
> Fstatistic=c(mod4sumsfstatistic[["value"]], "", "")
> P=pf(mod4sumsfstatistic[[1]], mod4sumsfstatistic[[2]], mod4sumsfstatistic[[3]],
+ lower.tail = F)
> pvalue=c(P, "", "")
> anovadf=data.frame(Df, SS, MS, Fstatistic, pvalue)
>

```

The environment pane on the right shows the following objects:

- anovadf: 3 obs. of 5 variables
- df1: 107 obs. of 6 variables
- df1test: 11 obs. of 6 variables
- df1train: 96 obs. of 6 variables
- df2: 5000 obs. of 4 variables
- df2train: 3000 obs. of 4 variables

The help pane on the right shows the documentation for the `pf` function, including its usage and arguments.

So, within F distribution we have F function which can be used to compute the corresponding p value. So, what we need is F statistics and the first argument then we need degree of freedom as second argument that is responding to a regression here and the second is then last one is residuals degree of freedom third argument and then we have a specified lower tail as false.

So, we will get the p value corresponding p value. So, let us record it in this format once this is done we can create this table data frame. And let us assign a names for this and let us have a look at this particular table. So, we can see that now once computed.

(Refer Slide Time: 29:50)

The screenshot shows the RStudio interface. The main editor contains R code for fitting an ANOVA model and predicting values. The console shows the execution of the code, including the calculation of F-statistics and p-values. The environment pane on the right shows the data frame 'anovadf' with 3 observations and 5 variables.

```

170 MS=c(mean(head(mod4ava[, "Mean Sq"], -1)),
171       mod4ava["Residuals", "Mean Sq"], "")
172
173 Fstatistic=c(mod4sumsfstatistic[["value"]], "", "")
174 P=pf(mod4sumsfstatistic[[1]], mod4sumsfstatistic[[2]], mod4sumsfstatistic[[3]])
175 lower.tail = F
176 pvalue=c(P, "", "")
177
178 anovadf=data.frame(Df, SS, MS, Fstatistic, pvalue)
179 rownames(anovadf)=c("Regression", "Error", "Total")
180 anovadf
181
182 # Prediction for a new observation:
183 # Annual income of A's lakhs with two family members
184 # who is not active online
185 predict(mod4, data.frame(income=5, family.size=2, online=0))
186
187
188

```

```

> Fstatistic=c(mod4sumsfstatistic[["value"]], "", "")
> P=pf(mod4sumsfstatistic[[1]], mod4sumsfstatistic[[2]], mod4sumsfstatistic[[3]])
+ lower.tail = F
> pvalue=c(P, "", "")
> anovadf=data.frame(Df, SS, MS, Fstatistic, pvalue)
> rownames(anovadf)=c("Regression", "Error", "Total")
> anovadf
      Df      SS      MS      Fstatistic      pvalue
Regression  3 72.44317 24.1477243151055 383.361202942226 1.00014894386617e-210
Error      2996 188.71649 0.062989483885631
Total      2999 261.15967

```

**Data**

| Variable | Observations             |
|----------|--------------------------|
| anovadf  | 3 obs. of 5 variables    |
| df1      | 107 obs. of 6 variables  |
| dfitest  | 11 obs. of 6 variables   |
| dfitrain | 96 obs. of 6 variables   |
| df2      | 5000 obs. of 4 variables |
| df2train | 3000 obs. of 4 variables |

**Usage**

```

df(n, df1, df2, ncp, log = FALSE)
pf(q, d1, d2, ncp, lower.tail = TRUE, 1)
qf(p, df1, df2, ncp, lower.tail = TRUE, 1)
z.f(n, d1, d2, ncp)

```

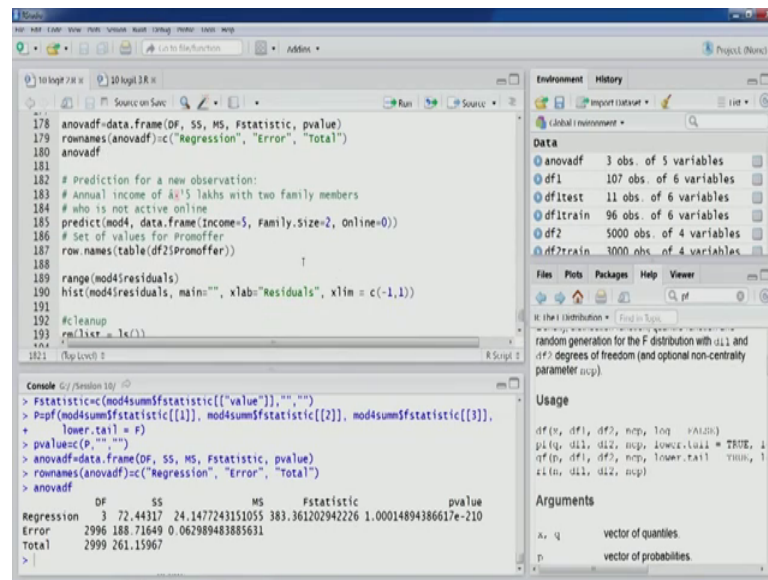
**Arguments**

- `q`: vector of quantiles.
- `p`: vector of probabilities.

So, first column we have degrees of freedom. So, for regression there are 3 and then because as remember that we have used just 4 variable, so one being outcome variables, so 3 predictors. So, therefore, degree of freedom degree of freedom is here is 3 for regression then residuals it is the remaining that is n minus 1 is 2999, total number of observation and is 3000.

So, residual degree of freedom 2996 sum of a square for regression and for residual is also there. So, you can see that residual sum of square is much higher so that also is we can we saw that the variance was also on the lower side right earlier we saw that the variance that we had computed was on the lower side so that is also indicated here, mean square of errors is also there and then we have F statis F statistic and then we have p value is small value.

(Refer Slide Time: 31:00)



The screenshot shows the RStudio environment. The script editor contains R code for ANOVA, prediction, and residual analysis. The console shows the output of the ANOVA and prediction functions. The Environment pane on the right lists the objects created in the script.

```
178 anovadf=data.frame(Df, SS, MS, Fstatistic, pvalue)
179 rownames(anovadf)=c("Regression", "Error", "Total")
180 anovadf
181
182 # Prediction for a new observation:
183 # Annual income of 5 lakhs with two family members
184 # who is not active online
185 predict(mod4, data.frame(Income=5, Family.Size=2, online=0))
186 # Set of values for Promoffer
187 row.names(table(df2$Promoffer))
188
189 range(mod4$residuals)
190 hist(mod4$residuals, main="", xlab="Residuals", xlim = c(-1,1))
191
192 # cleanup
193 rm(list = ls())
194
```

```
> Fstatistic=c(mod4sumsfstatistic[["value"]], "", "")
> P=pf(mod4sumsfstatistic[[1]], mod4sumsfstatistic[[2]], mod4sumsfstatistic[[3]],
+ lower.tail = F)
> pvalue=c(P, "", "")
> anovadf=data.frame(Df, SS, MS, Fstatistic, pvalue)
> rownames(anovadf)=c("Regression", "Error", "Total")
> anovadf
      Df      SS      MS      Fstatistic      pvalue
Regression 3 72.44317 24.1477243151055 383.361202942226 1.00014894386617e-210
Error      2996 188.71649 0.062989483885631
Total      2999 261.15967
```

**Environment**

- Global Environment
- Data
  - anovadf 3 obs. of 5 variables
  - df1 107 obs. of 6 variables
  - dfitest 11 obs. of 6 variables
  - dfitrain 96 obs. of 6 variables
  - df2 5000 obs. of 4 variables
  - df2train 3000 obs. of 4 variables
- Files
- Plots
- Packages
- Help
- Viewer

**Usage**

```
df(r, df1, df2, ncp, log = FALSE)
pf(q, d1, d2, ncp, lower.tail = TRUE, 1)
qf(q, df1, df2, ncp, lower.tail = TRUE, 1)
z.f(q, d1, d2, ncp)
```

**Arguments**

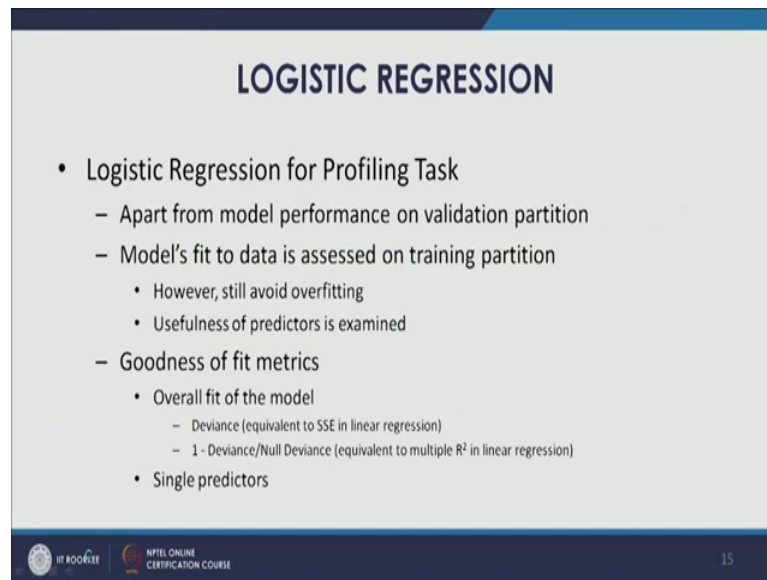
- A, q vector of quantiles.
- p vector of probabilities.

So, this gives us some information about the model and now what we will do is we will use this model to the score of a particular observation. So, let us do a prediction for a new observation let us say this is our new observation is annual income of rupees 5 lakhs with 2 family members who is not active online.

So, this is information about the particular customers customer whether is going to accept the promotional offer or not. So, annual income is 5 lakh family size is 2 and not acting online. So, you can use the predict function first argument is as usual the model object mod 4 then in a data frame we are trying to pass on the values operate predictors for example, income 5.

So, it should be in the same unit as was used for the modelling exercise in the training partition. So, you can see 5 and family size is 2 and then online is because this particular customer is not active so 0. So, once we do this we will be able to predict this particular observation. So, you can see this comes out to be a negative value. Now that was one of the first point that we discussed in the slide.

(Refer Slide Time: 32:11)



**LOGISTIC REGRESSION**

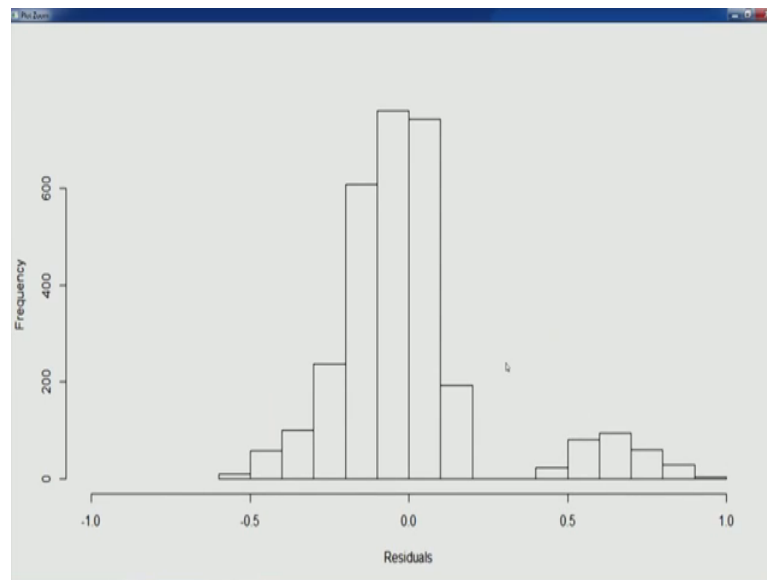
- Logistic Regression for Profiling Task
  - Apart from model performance on validation partition
  - Model's fit to data is assessed on training partition
    - However, still avoid overfitting
    - Usefulness of predictors is examined
  - Goodness of fit metrics
    - Overall fit of the model
      - Deviance (equivalent to SSE in linear regression)
      - $1 - \text{Deviance}/\text{Null Deviance}$  (equivalent to multiple  $R^2$  in linear regression)
    - Single predictors

IIIT ROOKEE | NPTEL ONLINE CERTIFICATION COURSE | 15

Let us go back you can see here that the predictions can take any value not just dummy values. So, we can see that a negative value has been taken here and so that is in one anomaly that we can clearly see. And set of values for them was no offer that we already know 0 and 1, and the value is that comes out to be predicted value comes out to be minus 0.14 now.

So, one difficulty is how do we do our classification in this case? Now, let us look at the residuals so second anomaly that I discussed that outcome variable of residuals do not follow normal distribution. So, let us look at that let us plot a histogram and find out whether this is being followed.

(Refer Slide Time: 33:07)



You can see residuals when we plot residuals when we get histogram you can see this is clearly not being followed normal distribution is not being followed one grouping here and other grouping where this is of this lower values, lower frequency and this is higher frequency. So, this is clearly looks binomial or different groups kind of thing. So, this definitely not following so normal distribution and distortions due to real binomial distribution can be seen here.

So, the typically what exercises and the discussion that we have been doing was pertaining 2 classification tasks and as we talked about that this is a statistical technique logistic regression being a statistical technique is also applied in so used in statistical modelling. So, the kind of task that we generally do in a statistical modelling are quite similar to what we can call profiling tasks.

So, as we talked about in the you know starting lecture of logistic regression and that it is about understanding similarities and differences between two groups. So, logistic regression can also be used in can also use to understand what are the variables which you know which can bring out some of the similarities or differences between group, so let us discuss that aspect as well.

So, in profiling tasks when we talk about profiling tasks the situation is slightly different. So, in the classification task we typically build our model and look at the performance of that classifier that particular model using the classification matrix using the overall

accuracy or overall error matrix and you know some deviations when we have a class of interest.

So, those are the things that we typically do we also typically look at left chart especially when we have a class of interest to in left chart and decide charge to see whether it is still despite you know higher error whether it still the model is useful in class of interest with respect to knave you know knave rule or a or a average case.

So, some of those things that we do in classification tasks; however, in profiling tasks what we do in classification we follow that. So, apart from model performance on validation partition we also asses models fit to data on training partition right because as I taught what in a statistical technique typically the whole sample is used; however, since we are using training partition for model building.

So, the models fit to data is assessed on training partition; however, model performance is assessed on validation partition for profiling tasks and models fit to data assessed on training partition. So, some of these things we talked about when we talk about goodness of fit measures we talked about the deviance, we talked about 1 minus deviance divided by null deviance that is equivalent of multiple R square some of those things you can see here as well.

So, models fit to data is assessed on training partition ah; however, we still focus on avoiding over fitting because as we talked about when we have matrix which look for and when we do modelling to achieve the you know goodness of fit then it can it can lead to over fitting. So, it still we would like to avoid over fitting and still be able to you know check the performance still be able to have good classification performance as well.

So, usefulness of usefulness of predictors is also examined in this particular case profiling so because it is about understanding similarities and differences between 2 groups. So, therefore, which predictor is more helpful in terms of bringing out those differences those similarities so when we build our model. We also look at the significance levels of some of these predictors we look at which predictors are significant which are not significant whether the not significant predictors can be dropped from the model.

And this has also should be looked at from the perspective of model performance because as we have taught about data mining model we would like to keep the insignificant variables also in the order if they provide some practical importance in terms of scoring new observation.

How and logistic statistical modelling we just drop the insignificant variables because we are just interested in understanding the phenomena, understanding the underlying relationship. So, profiling is quite similar to you know that that approach; however, because we are doing data mining modelling. So, we have to balance between these two, we have to balance between performance and also the main profiling tasks.

So, we have to really see whether the predictors can be dropped just like in statistical training or whether they have to be kept in the model because they also provide some practical significance for scoring new observation so that balance has to be achieved.

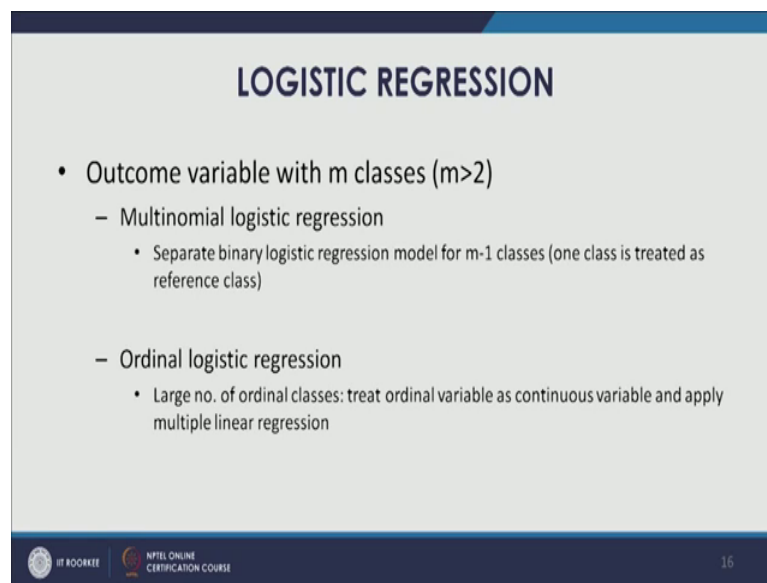
So, we have to avoid over fitting, we have to look for using usefulness or predictors in both the context data mining context prediction point of view and also statistical context in terms of understanding you know finding understanding variables which differentiate the groups, which bring out similarity or differences between groups.

So, this kind of exercise is done in profiling tasks and goodness of fit matrix that through an exercise in R that we have already understood over phila. So, we look to understand first we look to understand the overall fit of the model, so if the overall fit of the model is good only then we go ahead and look at the individual variables.

So, first step typically is in profiling on a statistical modelling we look at the overall fit of the model. So, in this particular case logistic regression as we talked about the deviance is the metric we taught one previous lecture that could be used and we also said that this is equal to SSE in linear regression and  $1 - \text{deviance} / \text{null devian}$  is equivalent to multiple R square in linear regression.

So, these are two matrix that could be used to assess the overall fit of the model and then once this is done then we look at the single predictors. We look at whether they are significant or not as I talked about and whether we can strike a balance in terms of prediction performance, classification performance, versus the profiling that is and also statistical modelling context.

(Refer Slide Time: 40:04)



**LOGISTIC REGRESSION**

- Outcome variable with  $m$  classes ( $m > 2$ )
  - Multinomial logistic regression
    - Separate binary logistic regression model for  $m-1$  classes (one class is treated as reference class)
  - Ordinal logistic regression
    - Large no. of ordinal classes; treat ordinal variable as continuous variable and apply multiple linear regression

16

So, with this we move to our next discussion that is about so till now the exercises that we have been doing they were mainly focused on binary classification, and binary logistic regression model we just had 2 classes class 0, and class 1. Can logistic regression we extended to a scenario where we are dealing with more than 2 classes where we are dealing with you know  $m$  classes. So, yes it is possible so we will discuss some of those things.

So, first one is multinomial logistic regression so multinomial logistic regression. So, the categories the classes that we have they are you know so the categorical variable is nominal, so, in that case we can apply multinomial logistic regression. So, what happens in multinomial logistic regression first out of those  $m$  classes that we have we have to select one we have to pick one as the reference category, and for the remaining  $m$  minus 1 classes we create separate binary logistic regression model.

So, for each of the  $m$  minus 1 classes apart from the reference category class. So, we will have  $n$  minus 1 classes so for each of the  $n$  minus 1 classes will create separate binary logistic regression model; that means, for a class 1 we will have the scenario where the observation probability of belonging to class 1 and probability of long not belonging to class 1 so that kind of binary scenario we will have and that for each of the  $n$  minus 1 classes.

So, we will be dealing with  $m - 1$  binary logistic regression equations and so using that we can compute all those probabilities values with the help of the predictors and then the remaining reference class of that probability for that can always be computed by the  $m - 1$  probabilities values for the  $m - 1$  classes.

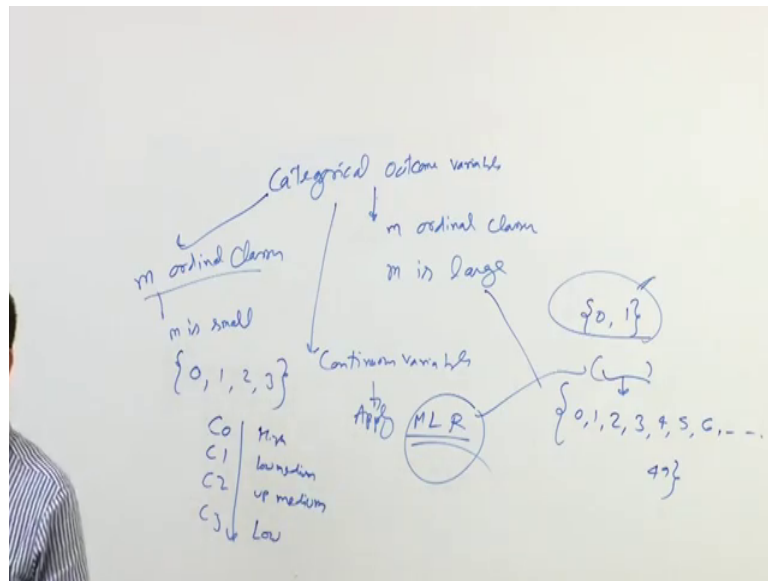
So, we can just subtract that from one for you know and then we can get for the probability value for the different class. And once we have the probabilities values for all  $m$  classes then we can apply our most probable class method routine where the class having highest probability value would be assigned to the new observation.

So, this is how we can go about applying logistic regression to an outcome variable with  $m$  classes. So, this is called this is a multinomial logistic regression and this is applicable mainly to nominal categorical nominal variable right, so the categorical variable having nominal classes.

The second scenario second scenario is about when we have a categorical variable with ordinal classes we have an ordinal variable. So, in those cases we can apply ordinal logistic regression. So, so again within this ordinal logistic regression we can have 2 scenarios. So, as we have understood in some of the initial lectures and supplementary lectures that ordinal variables they have order among different labels is also important that is also meaningful.

So, less than or equal to or greater than or equal to operations they are also applicable in this case, so the first scenario is large number of ordinal classes. So, if our outcome variable which is categorical variable with ordinal classes if that variable have is having large number of ordinal classes, then one solution is treat that ordinal variable as continuous variable and apply multiple linear regression right.

(Refer Slide Time: 44:10)



So, when we have a categorical variable with ordinal classes so we have a categorical variable. Categorical outcome variable when  $m$  ordinal classes, and  $m$  is large then I as I talked about we can treat this particular variable as a continuous variable and apply multiple linear regression.

So, multiple linear regression can be applied and reason is one reason one justification for this is as we talked about earlier that in binary situation we have just two values for a categorical variable and the predicted values up using MLR can range anywhere any real value so, that was the one main problem here.

But when we have  $m$  is large; that means, set of values could be you know many more. So, it could be you know if there are 50 groups let us so in this fashion we can go on up to this. So, the number of values that can be taken by this particular this particular ordinal variable are many more.

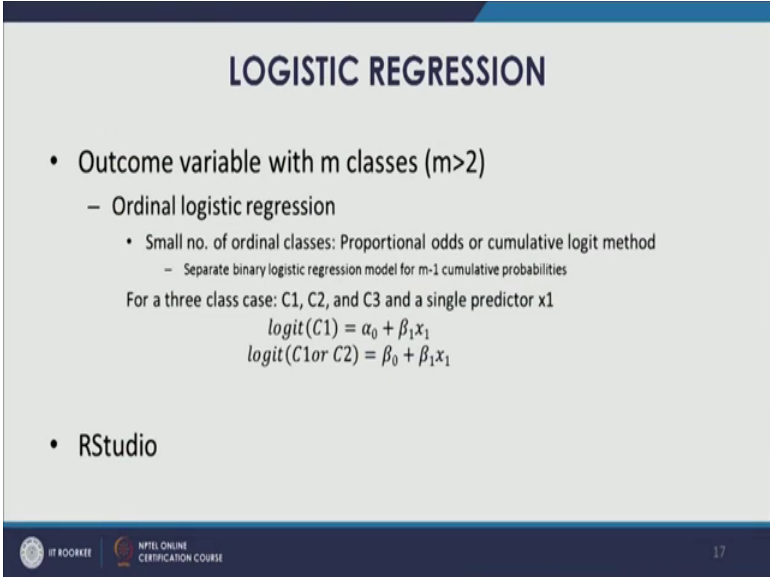
So, therefore, the predicted values this can be easily mapped to some of these values could be close to some of these values, and probably multiple linear regression can be still applied so this is one way. When we have a ordinal variable with many number of classes with large number of classes, when  $m$  is large then is still you know instead of logistic regression we can apply multiple linear regression so that is first scenario.

The second one is whatever we have a small number of ordinal classes. So, if  $m$  is in this case  $m$  ordinal classes that we have if this is a small  $m$  is small, then probably we will we will run into the same problem like for binary classification. So, similar problem would be there we will have only few values 0, 1, 2, 3 let us say these many.

So, again small number of classes small number of ordinal classes, so we will have I will run into same problem. So, so we what we do is we use a different version of logistic regression called proportional odds or cumulative logit method as indicated in the slide, so small number of ordinal classes so we would like to use proportional odds, or cumulative cumulative, or logit method.

So, what we do here in this particular method is we create separate binary logistic regression model for  $m$  minus 1 cumulative probabilities, so, we talked about that when we so, when we discussed multinomial logistic regression; So, for all  $m$  minus 1 classes will have separate binary logistic regression model.

(Refer Slide Time: 47:50)



**LOGISTIC REGRESSION**

- Outcome variable with  $m$  classes ( $m > 2$ )
  - Ordinal logistic regression
    - Small no. of ordinal classes: Proportional odds or cumulative logit method
      - Separate binary logistic regression model for  $m-1$  cumulative probabilities
- For a three class case:  $C_1$ ,  $C_2$ , and  $C_3$  and a single predictor  $x_1$ 
$$\text{logit}(C_1) = \alpha_0 + \beta_1 x_1$$
$$\text{logit}(C_1 \text{ or } C_2) = \beta_0 + \beta_1 x_1$$
- RStudio

IT ROOKIE NPTEL ONLINE CERTIFICATION COURSE 17

However, if you see that here will have separate binary logistic regression model not for  $m$  minus 1 classes not for presence of that class or absence of that class, but  $m$  minus 1 cumulative probabilities.

So, let us understand what we mean by that so let us take an example for a 3 class case as written in this slide for a 3 class case  $C_1$ ,  $C_2$ , and  $C_3$ . Let us say these are the 3 classes

ordinal classes  $C_1$ ,  $C_2$ ,  $C_3$  and a single predictor  $x_1$  that is being used. So, our logit equations could be something like this logit for  $C_1$  it could be  $\alpha_0 + \beta_1 x_1$  and logit for  $C_2$  so; that means, from  $C_1$  first logit equation is just for you know observation belonging to  $C_1$ . The second is observation belonging to  $C_2$  so that gives us the cumulative sense. So, as we talked about that ordinal the order is important so that that means, you know different classes can be compared.

So, therefore,  $C_1$ , or  $C_2$  is a you know is a meaningful here in the sense that if we look at the right part of the equation  $\beta_0 + \beta_1 x_1$  you can see that  $\beta_1$  is same and both these equation  $x_1$  so you can see because the this is the comparison can be done.

So, the coefficient so the intercept  $R$  only difference because the comparison so one is when we talk about ordinal, so one particular class this ordering this ordering is you know this is meaningful. So, when the ordering of classes is meaningful; that means, one can be you know said you know less than one particular class the or higher than one particular class just like the categories that we might have high, you know low medium, upper medium, medium and then low.

So, this kind of this kind of classes we might have all we might have you know a strongly agree to strongly disagree. So, those kind of ordinal classes where the order is meaningful, where the order is meaningful then in those regression we can have something like this cumulative probabilities values, and beta coefficient is can be used the same beta coefficient can be used for both these logistic regression.

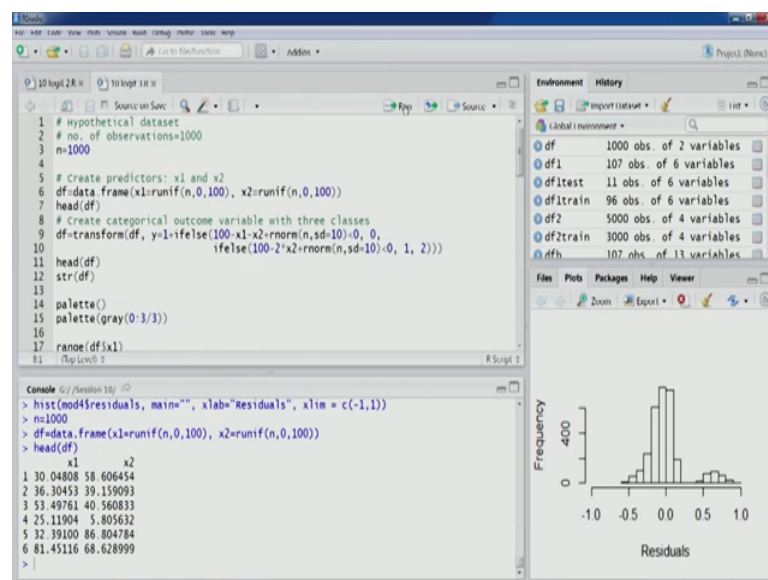
And from this we can compute the cumulative probabilities values and once these cumulative probabilities value values for two of the classes for these classes have been computed the actual probabilities value for  $C_1$ ,  $C_2$ ,  $C_3$  can be derived from using the probability value that formulation that we have so, it is from logit, we can derive the probabilities values.

So, once these cumulative so from there once we have the probabilities values for all these classes then we can again apply the most probable class method and also assign the class based on the based on the probability value and again this so in this fashion in this fashion we can apply ordinary logistic regression to a scenario with fewer number of ordinal classes.

So, what we will do we will go through R Studio and do an exercise for this when we have classes more than m is greater than 2, so a logistic regression modelling for classes greater than 2. So, what we will do we will create a hypothetical data set here, so number of observations are 100 in this case, so as you can see and it is 1000.

So, let us create this now what we are going to do is we are going to create a data frame having that data frame where we have 2 variables x 1 and x 2, x 1 as you can see we are using run if for x 1 and x 2 both so the observations. So, the wave values would lie between 0 and 100 and n number of values would be created that is 1000 values would be created lying between 0 and 100.

(Refer Slide Time: 52:23)



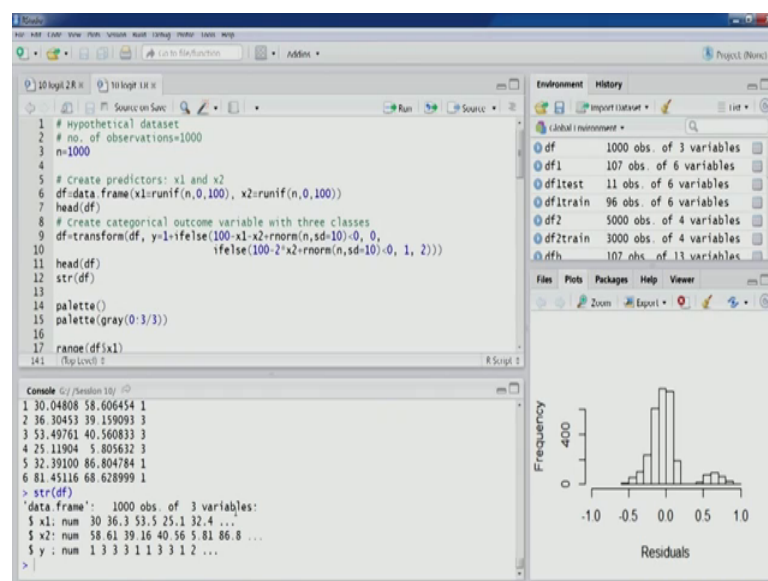
So, let us create this data frame, let us look at some of the observation. For 6 observations you can see all the values they are lying for both the variables x 1 and x 2; they are between 0 and 100.

Now, what we will do we will create a categorical outcome variable with 3 classes. So, this particular data set that we are trying to create we are going to use it for both the scenarios multinomial scenario, and ordinal scenario right, so this is just for illustration purpose so we are not specifying whether the variable is ordinal or not s, we are just going to use it for both the scenarios.

So, what we are going to do is we are using transform function to create our outcome variable, categorical outcome variables. So, you can see that we are trying to compute y as 1 plus if else and if the value is less than 0, this particular value 100 minus x 1 minus x 2 plus again a certain value is being taken from normal distribution no standard deviation 10000 values.

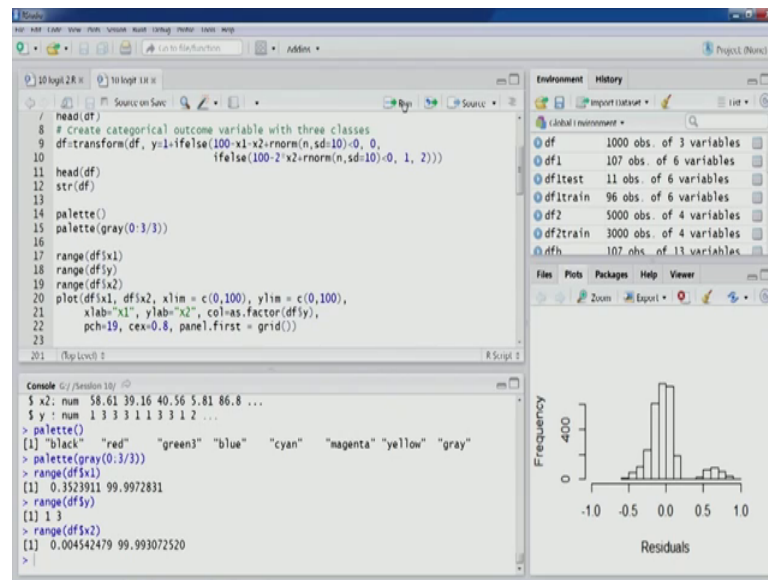
So, if it is less than 0, then you know we are using this information and that is information based on x 1, x 2, and certain additional computation we are trying to assign it a class 0 or then the second if that is not true then second computation so it will get a class 1 or class 2. So, let us compute this once this is done let us look at the observations you can see another variable y that is categorical variable has been created having you can see 1 and 3 2 values are being taken. Let us look at the structure of this particular data frame.

(Refer Slide Time: 54:00)



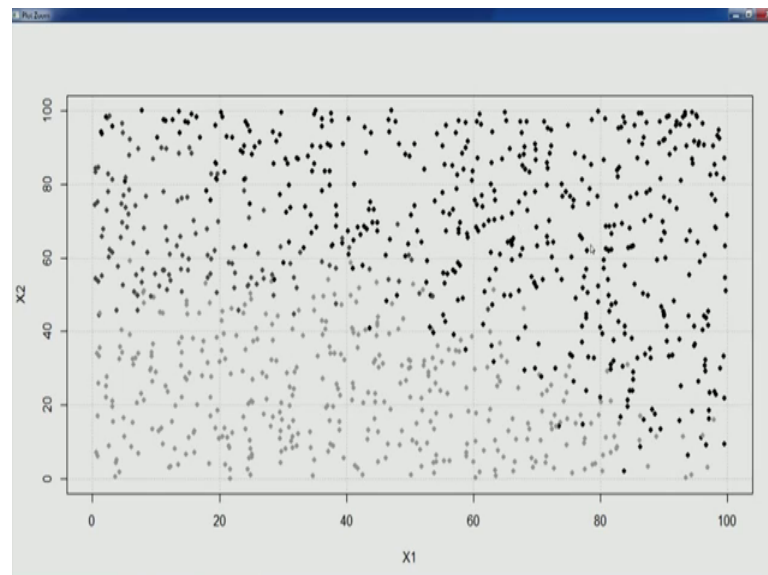
So, these are 3 variables x 1, x 2 values between 0 and 100, and y is taken value 3 values 1, 2, and 3. So, let us plot this particular data set. So, this is our default palette so; however, I would like to use this palette agree 3 sets. So, in this fashion we can have sets of no number of sets that we require. So, let us look at the range of x 1, y, and x 2 which is already.

(Refer Slide Time: 54:25)



Because we have just now created these variables studies actually clearly understood as well so limits 0 to 100, and 0 to 100 and then colouring is with using this particular factor y. So, you can see as dot factor we are using this particular variable and let us create this plot.

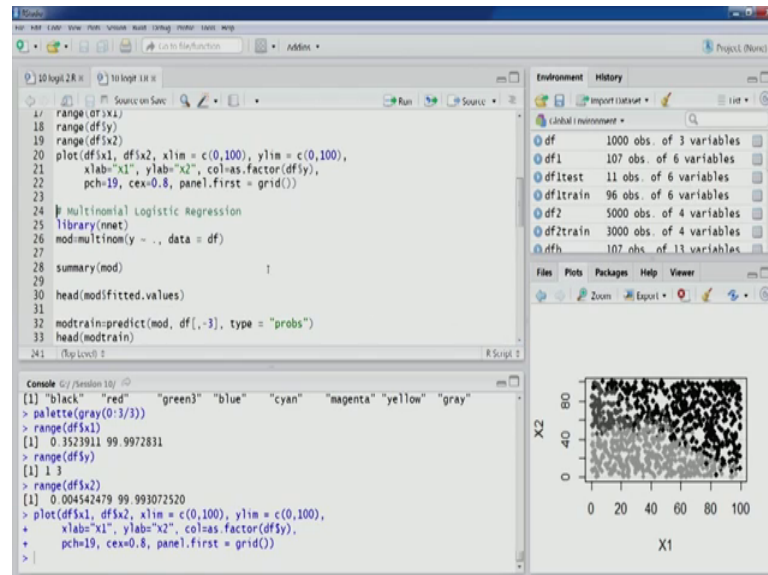
(Refer Slide Time: 54:58)



So, you can see so this is our plotting. So, a one group is here, the second group is having some medium level gray set here and the third group is here, this is lighter gray code

colour. So, these are the values that we are going to use x 1, x 2 this is plot between x 1, and x 2 and the outcome variable has been color coded, so, 3 categories.

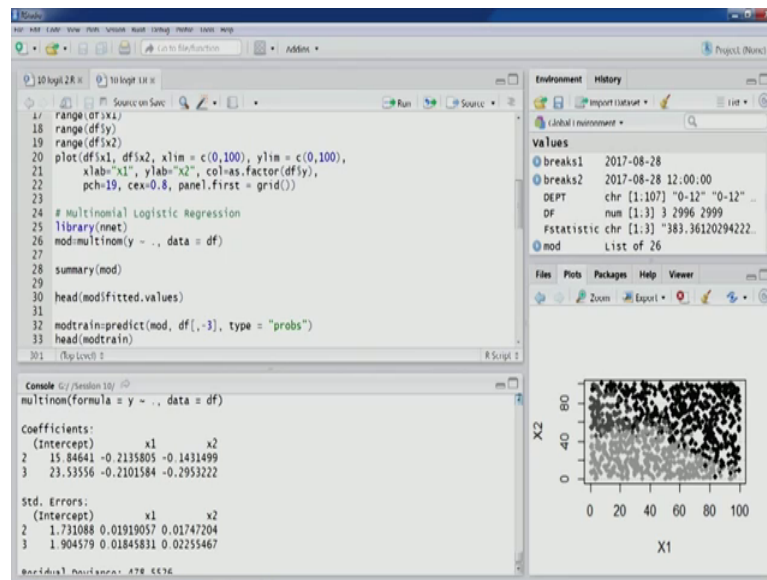
(Refer Slide Time: 55:25)



Now, what we will do is we will use the multinomial logistic regression. So, for this we need this package the library, and net package this package actually for a neural network, but it provides us offers us this function which can be used for multinomial logistic regression.

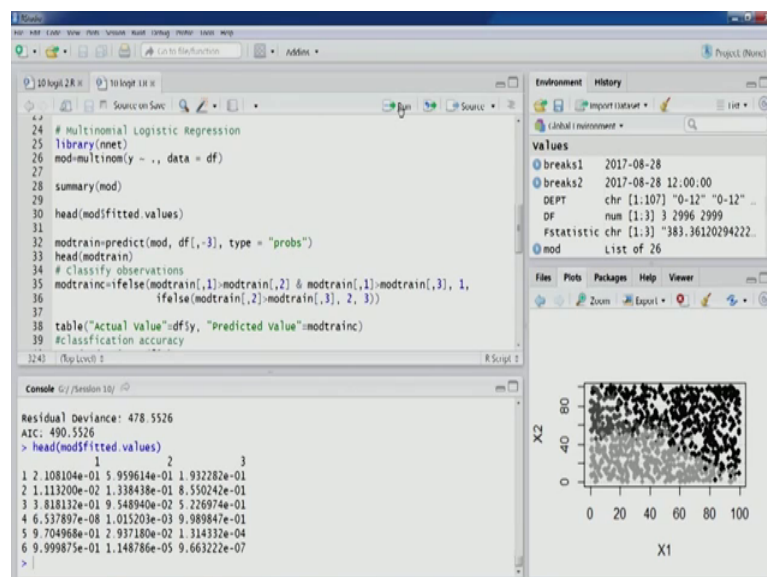
So, a multi norm is the function so what we are now going to do is the outcome variable y we are going to regress it with the remaining variables that is predictors in this particular data. So, all the observations we are going to use here, so let us run, this let us look at the summary.

(Refer Slide Time: 56:00)



So, we can see coefficients x 1 and x 2 right. We can also see the error values residual deviance and the AIC values, are also indicated here. Now, fitted values also we can look at first 6 observation here.

(Refer Slide Time: 56:19)



So, we will get these fitted values which are nothing, but estimated probabilities values for these 3 classes 1, 2, and 3 so what this is what we have. Now, if we had another partition validation or test partition you can use predict function to the score those

partitions and have probabilities values, estimated probability values, for those new observation.

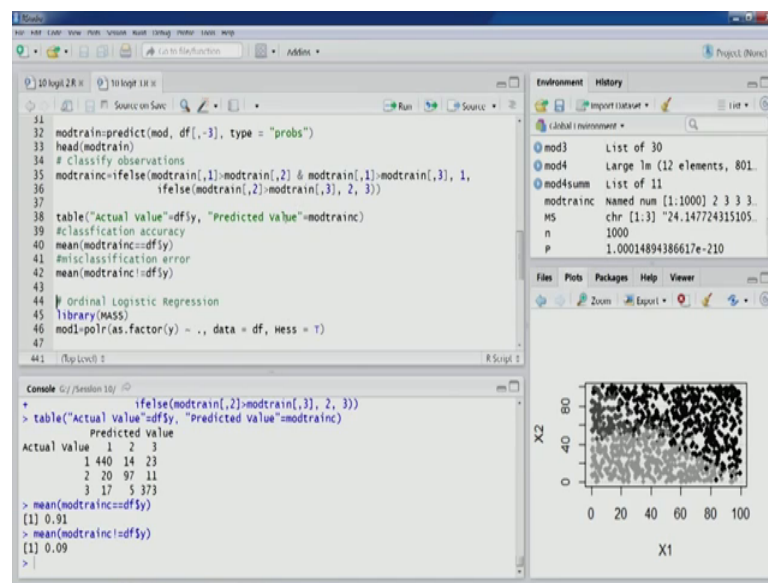
However for the demonstration purpose, we are applying predict function on the training partition that is the full data set itself. So, you can see another argument that we can see is type which is probs which is proper probabilities values. So, let us run this and you would see we will have the same probabilities values which were fitted. So, we can see here that same values you can see first row it is same, same values. So, the fitted values and all we have got the same one using predict function.

Now, let us classify these observations. So, this is how we can classify if the probability value for 1 is greater than the probability value for class 2 and also greater than the probability value over class 3, then of course, this observation has to be classified to class 1.

Otherwise we will again compare the probability value for class 2, with probability value for class 3, if again it is greater and I know class 3 value divided then the class 2 is assigned, otherwise class 3.

So, in this fashion we can assign all the observation into appropriate classes so this is implementation of most probable class method. So, once this is done you can see mod train see for all 1000 variables have been created and observations have been assigned to go to appropriate classes.

(Refer Slide Time: 58:13)



Now, let us look at the classification matrix. So, we can see actual value, predicted values, now till now the classification matrices that we have been observing that we have been creating they had 2 values 0 and 1. And we had 2 by 2 classification matrix now this time we are seeing 3 by 3 classification matrix.

So, 3 possibilities for actual values 1, 2, and 3 and the predicted values and we have the corresponding numbers here. So, you can see that again the diagonal element; that means, these 3 elements they represent the values records which have been correctly classified, and the remaining observation off diagonal elements you know they have the records, they have the counts of record, which have been incorrectly classified right.

So, from this we can compute the classification accuracy that is 91 percent in this case, and error that is the remaining 9 percent. So, in this fashion we can apply multinomial logistic regression to a particular data set. Now, let us move to next part that is ordinal logistic regression.

So, in this product case ordinal logistic regression as I talked about 2 scenarios, so 1 where m is large; So, in the in that scenario we can apply multiple regression one exercise that we had done in previous lecture, we had applied multiple linear regression, but that was for the class which had just two you know for a variable category which just had 2 classes 0 and 1.

In the same fashion for the first scenario the multiple linear regression, can be applied. So, there is not much much different in terms of applicability. So, well what we will do we will do an exercise for this scenario where we have to apply this ordinary ordinal logistic regression the cumulative probabilities method, cumulative logit method so for this we need this package mass.

(Refer Slide Time: 60:33)

```

36 ifelse(modtrain[,2]>modtrain[,3], 2, 3))
37
38 table("Actual value"=df$y, "Predicted value"=modtrainc)
39 #classification accuracy
40 mean(modtrainc==df$y)
41 #misclassification error
42 mean(modtrainc!=df$y)
43
44 # ordinal Logistic Regression
45 library(MASS)
46 mod1=polr(as.factor(y) ~ ., data = df, Hess = T)
47
48 summary(mod1)
49
50 head(mod1$fitted.values)
51
52 mod1$residuals(mod1, df[, 1:3], xname = "xname")
53 #Top Level 3

```

Environment: global environment

- mod3: List of 30
- mod4: Large lm (12 elements, 801)
- mod4sum: List of 11
- modtrainc: Named num [1:1000] 2 3 3 3
- MS: chr [1:3] "24.147724315105"
- n: 1000
- P: 1.00014894386617e-210

Console: G:/Lesson 10/

```

> table("Actual value"=df$y, "Predicted value"=modtrainc)
      Predicted value
Actual value 1 2 3
1 440 14 23
2 20 97 11
3 17 5 373
> mean(modtrainc==df$y)
[1] 0.91
> mean(modtrainc!=df$y)
[1] 0.09
> library(MASS)
>

```

polr (MASS) R Documentation

### Ordered Logistic or Probit Regression

Description

Fits a logistic or probit regression model to an ordered factor response. The default logistic case is proportional odds logistic regression, after which the function is named.

Usage

So, let us load this particular library and we have polr is the function for this. So, this is actually for profit however, it can also be used for it can also be used for the cumulative logit method. So, we will just see in the in the in the help section. So, you can see polr this is ordered logistic, or probit regression right.

So, what we are interested in ordered logistic regression, so so here ordered logistic regression. So, you can see in the method argument the first one is draw logistic this is actually ordered logistic method. This is also called proportional odds logistic regression which we have discussed right.

So, let us go and build this model so polr is the function. So, first as you can see y variable I have converted into a factor variable and then the request against all other variables that which are predictors data is full, then we need this argument as well as which is mainly if we want to apply for summary function later on the model object, so, this is also true.

So, now let us apply summary, so we will get the results we can see the coefficient values here x 1, and x 2, the value and the error and the T values are also there. And we have the residual deviance, and AIC value as well for this particular model and we are interested in finding the fitted values so that is also returned by the model. So, let us look at some of these values.

(Refer Slide Time: 62:00)

The screenshot shows the R Studio environment. The script editor contains the following R code:

```

39 #classification accuracy
40 mean(modtrain==df$y)
41 #misclassification error
42 mean(modtrain!=df$y)
43
44 # Ordinal Logistic Regression
45 library(MASS)
46 mod1=polr(as.factor(y) ~ ., data = df, Hess = T)
47
48 summary(mod1)
49
50 head(mod1$fitted.values)
51
52 modtrain=predict(mod1, df[,3], type = "probs")
53 head(modtrain)
54 # classify observations
55 modtrain=ifelse(modtrain[,1]>modtrain[,2] & modtrain[,1]>modtrain[,3], 1,
56
57
58
59
60
61
62
63
64
65
66
67
68
69
70
71
72
73
74
75
76
77
78
79
80
81
82
83
84
85
86
87
88
89
90
91
92
93
94
95
96
97
98
99
100
101
102
103
104
105
106
107
108
109
110
111
112
113
114
115
116
117
118
119
120
121
122
123
124
125
126
127
128
129
130
131
132
133
134
135
136
137
138
139
140
141
142
143
144
145
146
147
148
149
150
151
152
153
154
155
156
157
158
159
160
161
162
163
164
165
166
167
168
169
170
171
172
173
174
175
176
177
178
179
180
181
182
183
184
185
186
187
188
189
190
191
192
193
194
195
196
197
198
199
200
201
202
203
204
205
206
207
208
209
210
211
212
213
214
215
216
217
218
219
220
221
222
223
224
225
226
227
228
229
230
231
232
233
234
235
236
237
238
239
240
241
242
243
244
245
246
247
248
249
250
251
252
253
254
255
256
257
258
259
260
261
262
263
264
265
266
267
268
269
270
271
272
273
274
275
276
277
278
279
280
281
282
283
284
285
286
287
288
289
290
291
292
293
294
295
296
297
298
299
300
301
302
303
304
305
306
307
308
309
310
311
312
313
314
315
316
317
318
319
320
321
322
323
324
325
326
327
328
329
330
331
332
333
334
335
336
337
338
339
340
341
342
343
344
345
346
347
348
349
350
351
352
353
354
355
356
357
358
359
360
361
362
363
364
365
366
367
368
369
370
371
372
373
374
375
376
377
378
379
380
381
382
383
384
385
386
387
388
389
390
391
392
393
394
395
396
397
398
399
400
401
402
403
404
405
406
407
408
409
410
411
412
413
414
415
416
417
418
419
420
421
422
423
424
425
426
427
428
429
430
431
432
433
434
435
436
437
438
439
440
441
442
443
444
445
446
447
448
449
450
451
452
453
454
455
456
457
458
459
460
461
462
463
464
465
466
467
468
469
470
471
472
473
474
475
476
477
478
479
480
481
482
483
484
485
486
487
488
489
490
491
492
493
494
495
496
497
498
499
500
501
502
503
504
505
506
507
508
509
510
511
512
513
514
515
516
517
518
519
520
521
522
523
524
525
526
527
528
529
530
531
532
533
534
535
536
537
538
539
540
541
542
543
544
545
546
547
548
549
550
551
552
553
554
555
556
557
558
559
560
561
562
563
564
565
566
567
568
569
570
571
572
573
574
575
576
577
578
579
580
581
582
583
584
585
586
587
588
589
590
591
592
593
594
595
596
597
598
599
600
601
602
603
604
605
606
607
608
609
610
611
612
613
614
615
616
617
618
619
620
621
622
623
624
625
626
627
628
629
630
631
632
633
634
635
636
637
638
639
640
641
642
643
644
645
646
647
648
649
650
651
652
653
654
655
656
657
658
659
660
661
662
663
664
665
666
667
668
669
670
671
672
673
674
675
676
677
678
679
680
681
682
683
684
685
686
687
688
689
690
691
692
693
694
695
696
697
698
699
700
701
702
703
704
705
706
707
708
709
710
711
712
713
714
715
716
717
718
719
720
721
722
723
724
725
726
727
728
729
730
731
732
733
734
735
736
737
738
739
740
741
742
743
744
745
746
747
748
749
750
751
752
753
754
755
756
757
758
759
760
761
762
763
764
765
766
767
768
769
770
771
772
773
774
775
776
777
778
779
780
781
782
783
784
785
786
787
788
789
790
791
792
793
794
795
796
797
798
799
800
801
802
803
804
805
806
807
808
809
810
811
812
813
814
815
816
817
818
819
820
821
822
823
824
825
826
827
828
829
830
831
832
833
834
835
836
837
838
839
840
841
842
843
844
845
846
847
848
849
850
851
852
853
854
855
856
857
858
859
860
861
862
863
864
865
866
867
868
869
870
871
872
873
874
875
876
877
878
879
880
881
882
883
884
885
886
887
888
889
890
891
892
893
894
895
896
897
898
899
900
901
902
903
904
905
906
907
908
909
910
911
912
913
914
915
916
917
918
919
920
921
922
923
924
925
926
927
928
929
930
931
932
933
934
935
936
937
938
939
940
941
942
943
944
945
946
947
948
949
950
951
952
953
954
955
956
957
958
959
960
961
962
963
964
965
966
967
968
969
970
971
972
973
974
975
976
977
978
979
980
981
982
983
984
985
986
987
988
989
990
991
992
993
994
995
996
997
998
999
1000

```

The console window shows the following output:

```

Residual Deviance: 785.2117
AIC: 793.2117
> head(mod1$fitted.values)
      1      2      3
1 0.3188775524 0.4415837004 0.2395387471
2 0.0515084371 0.2176331067 0.7308584562
3 0.2585359098 0.442402184 0.2972238718
4 0.0001625156 0.0009384905 0.9988989939
5 0.9697389188 0.0256803680 0.0045807132
6 0.9964164603 0.0030534632 0.0005300764

```

The help window for the 'polr' function is also visible, showing the description: 'Fits a logistic or probit regression model to an ordered factor response. The default logistic case is proportional odds logistic regression, after which the function is named.'

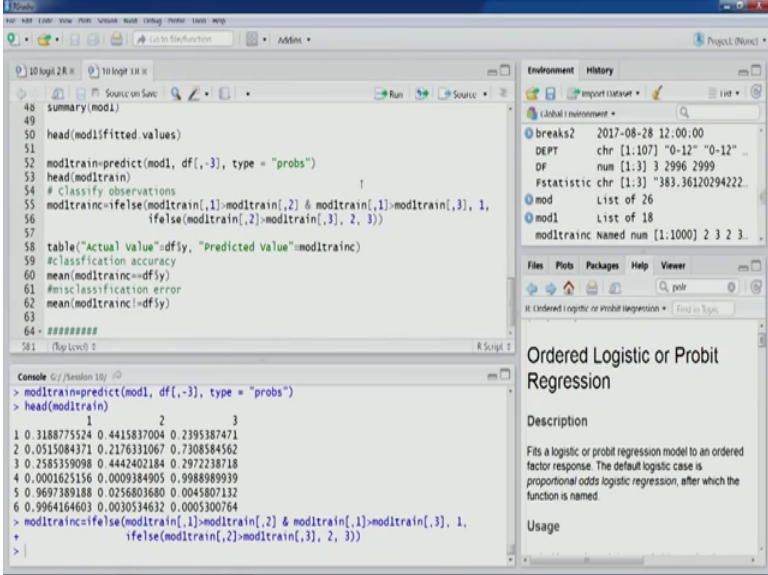
So, you can see for each of these classes class 1, 2, and 3, so these values are actually probability estimated probabilities values. Now using these values we can again apply the most probable class method so first we need to compute the first we need to compute and do the assignment as per the most probable class method like we did in previous exercise.

So, as I talked about predict function can again be used to score new data. So, in this case we are scoring the training partition itself again. So, we expect to get the same values as you can see last row, you can see here, here also the same values are there. So, it is scoring off for the training partition itself. So, we will get the same observation.

Now, what we will do classify observation. So, as we did in previous exercise mod 1 train for class 1 and probability value for class 1 greater than value class 2, and greater than class 3. Then assign it to class 1 otherwise again we look on more comparison the probability value for class 2 is greater than probability value class 3, then assign it to

class 2, otherwise class 3. So, in this way we will have the appropriate classification scores.

(Refer Slide Time: 63:17)



The screenshot shows the R Studio environment. The script editor contains the following code:

```
summary(mod1)
48
49 head(mod1$fitted.values)
50
51 mod1train=predict(mod1, df[,3], type = "probs")
52
53 head(mod1train)
54
55 # Classify observations
56 mod1trainc=ifelse(mod1train[,1]>mod1train[,2] & mod1train[,1]>mod1train[,3], 1,
57 ifelse(mod1train[,2]>mod1train[,3], 2, 3))
58
59 table("Actual Value"=df$y, "Predicted Value"=mod1trainc)
60
61 #classification accuracy
62 mean(mod1trainc==df$y)
63
64 #misclassification error
65 mean(mod1trainc!=df$y)
66
67 #####
68
69 581 | Flip Level: 3
```

The console output shows the results of the predictions:

```
> mod1train=predict(mod1, df[,3], type = "probs")
> head(mod1train)
      1      2      3
1 0.3188775524 0.4415837004 0.2395387471
2 0.0515084371 0.2176331067 0.7308584562
3 0.2585359098 0.4442402184 0.2972238718
4 0.0001625156 0.0009384905 0.9988989939
5 0.9697389188 0.0256803680 0.0045807132
6 0.9964164603 0.0030534632 0.0005300764
> mod1trainc=ifelse(mod1train[,1]>mod1train[,2] & mod1train[,1]>mod1train[,3], 1,
+ ifelse(mod1train[,2]>mod1train[,3], 2, 3))
>
```

The Environment pane on the right shows the following objects:

- breaks2: 2017-08-28 12:00:00
- DEPT: chr [1:107] "0-12" "0-12" ...
- DF: num [1:3] 3 2996 2999
- FStatistic: chr [1:3] "383.36120294222"
- mod: List of 26
- mod1: List of 18
- mod1trainc: Named num [1:1000] 2 3 2 3 ...

The Files pane shows the following files:

- Ordered Logistic or Probit Regression

The Description pane shows the following text:

Description

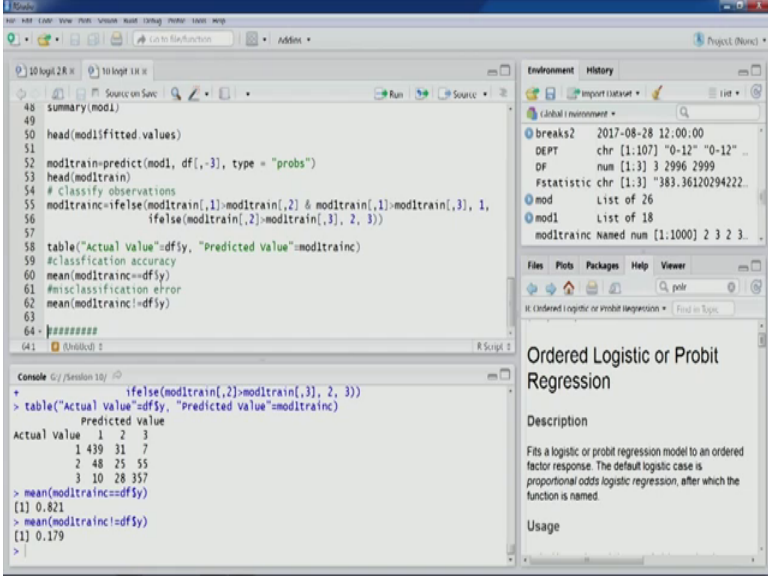
Fits a logistic or probit regression model to an ordered factor response. The default logistic case is proportional odds logistic regression, after which the function is named.

Usage

Now, let us generate the classification matrix here. so you can see 3 by 3 matrix we have 3 classes actual values 3 possibilities, predicted values, 3 predicted classes, so again diagonal elements they represent the correct classification values and off diagonal elements they represent the incorrect classifications.

So, what we can do is so let us compute a classification accuracy, and you can see eighty two percent and the remaining is error.

(Refer Slide Time: 63:49)



The screenshot shows the R Studio environment. The script editor contains the following code:

```
summary(mod1)
49
50 head(mod1$fitted.values)
51
52 mod1train=predict(mod1, df[,3], type = "probs")
53 head(mod1train)
54 # Classify observations
55 mod1trainc=ifelse(mod1train[,1]>mod1train[,2] & mod1train[,1]>mod1train[,3], 1,
56 ifelse(mod1train[,2]>mod1train[,3], 2, 3))
57
58 table("Actual Value"=df$y, "Predicted Value"=mod1trainc)
59
60 #classification accuracy
61 mean(mod1trainc==df$y)
62 #misclassification error
63 mean(mod1trainc!=df$y)
64
65 #####
66 > ifelse(mod1train[,2]>mod1train[,3], 2, 3))
> table("Actual Value"=df$y, "Predicted Value"=mod1trainc)
      Predicted value
Actual Value  1  2  3
1  439  31  7
2  48  25  55
3  10  28  357
> mean(mod1trainc==df$y)
[1] 0.821
> mean(mod1trainc!=df$y)
[1] 0.179
> |
```

The console output shows the results of the classification:

```
> ifelse(mod1train[,2]>mod1train[,3], 2, 3))
> table("Actual Value"=df$y, "Predicted Value"=mod1trainc)
      Predicted value
Actual Value  1  2  3
1  439  31  7
2  48  25  55
3  10  28  357
> mean(mod1trainc==df$y)
[1] 0.821
> mean(mod1trainc!=df$y)
[1] 0.179
> |
```

The Environment pane on the right shows the following objects:

- breaks2: 2017-08-28 12:00:00
- DEPT: chr [1:107] "0-12" "0-12" ..
- DF: num [1:3] 3 2996 2999
- Fstatistic: chr [1:3] "383.36120294222..
- mod: List of 26
- mod1: List of 18
- mod1trainc: Named num [1:1000] 2 3 2 3 ..

The Files pane shows the following files:

- Ordered Logistic or Probit Regression

The Description pane shows the following text:

Description

Fits a logistic or probit regression model to an ordered factor response. The default logistic case is proportional odds logistic regression, after which the function is named.

Usage

So, with this we have completed our discussion on logistic regression and so today we have been also able to cover the scenarios where more than 2 classes are present in our categorical variable. What happens when the classes are nominal and how we can apply logistic regression when classes are ordinal, so we have seen that we have also done an exercise in R.

So, we stop here and we will continue our discussions in next structure for a new technique.

Thank you.