

Total Quality Management - II
Prof. Raghunandan Sengupta
Department of Industrial and Management Engineering
Indian Institute of Technology, Kanpur

Lecture - 38
Fitting Regression Models - III

Welcome back my dear friends a very good morning, good afternoon, good evening to all of you we are in the 38 lecture for the TQM II course under the NPTEL MOOC series and I am Raghunandan Sengupta from IME department IIT Kanpur. So, if you remember we in the last 2 lectures for the last week, which is the 36th and 37th seventh lecture we did not cover much of the slides.

But we have we had a lot of discussions about regression model, multiple linear regression model or you find out estimate of beta how you find out the sum of the square of errors then how you find out the variance for the betas and the covariance of betas sigma square, what is the degrees of freedom. And then; obviously, I keep kept repeating what are the assumptions and how you do the q q plots, how you do the plots of the errors with respect to x_1 , x_2 , x_3 whichever variable which are there the assumptions being that normality holds true for all the x s for the errors for the y s and expected value of the errors is 0 variances and all is obviously, there and all this things.

Then we slowly went in to the 2 to the power k fractional factor models considering there are k variables they are 2 different levels 0 1 minus 1 plus 1 and so on. So, and we saw that that how the concept of orthogonality can be brought into the picture and we can module in the multiple linear regression models.

And we were discussing that why is the variable and there are beta naught, beta 1, beta 2, beta 3 which means there is x_1 , x_2 , x_3 variables based on which you can proceed and we can find out the errors and the variances of beta naught variances of beta 1, beta 2, beta 3 where the variances of beta 1, beta 2, beta 3 for the last example which were discussing in the 37th lecture all the variances were equal because food for the orthogonality one, orthogonality one.

And I did mention time and again the concept of degrees of freedom you are losing the degrees of freedom as you add more and more variables which are the independent ones which basically affect the value of y . So, continuing the discussion.

(Refer Slide Time: 02:34)

TABLE 10.4
Minitab Output for the Viscosity Regression Model, Example 10.1

Regression Analysis

The regression equation is:
Viscosity = 1566.08 + 7.62 Temp + 8.52 Feed Rate

Predictor	Coef	Std. Dev.	t	P
Constant	1566.08	61.19	25.43	0.000
Temp	7.6213	0.6184	12.32	0.000
Feed Rate	8.585	2.439	3.52	0.004

S = 16.36 R-sq = 92.7% R-sq (adj) = 91.6%

Analysis of Variance

Source	DF	SS	MS	F	P
Regression	2	44157	22079	87.50	0.000
Residual Error	15	3479	232		
Total	17	47636			

Source	DF	Seq. SS
Temp	1	40845
Feed Rate	1	3356

So, if you have the output for the viscosity regression models at the actual viscosity which is y is given by $\hat{y}$ which is 1566 plus β_2 times x_1 plus β_3 times x_2 .

So, this will be β_2 with 7.62 into temperature is x_1 plus 8.52 which is β_3 times x_2 and then if you see the constant. So, which is what is interesting I want to note down and highlight. So, these were the coefficients, the coefficient value is 1 point 566.08 which is here 7.656213 which is here 8.585 which is here. The standard deviations or the standard errors for the constant one; obviously, because they are estimates, estimates would basically have expected value and if there is an expected value they are random and they would be a variance for that. So, variance for all these 3 constant which is β_0 temperatures coefficient which is β_1

And and feed rate coefficient which is β_2 their standard error are given. So, note down for a value of 1566 the standard error deviation is about 61 for 7.62 it is 0.61 for 8.5 it is 2.4. For the temperature difference values which you remember t plus and t minus and the p value level of significance, the r squared value is given as 92.7 adjusted

r square is given as 91.6 adjusted r square basically considers the degrees of freedom in those picture that is all, in a in a very general layman terms I am saying.

So, r square basically gives you the level of prediction we are able to do it, if it is 92.7; that means, we are able to predict about 92 percent of the variations of y with respect to x; obviously, 100 minus 92 would be the effect of white noise or the errors, the adjusted r square when we bring the degrees of freedom into the picture. So, it comes out to be about 91.6; that means, 100 minus 91.6 is the overall errors or the prediction which we are not able to predict using the axis which is basically given by the white noise.

The analysis of variance again we will go back to the general main theme of our discussions you have a table the, the reading in the all the variables of the attributes are there on the first column, then second column you have the sum of squared errors you have the third column in the degrees of freedom you have f value, because you want to test the null hypothesis whether this is true or false and you have the degrees of a level of significance. So, these are the sources, these are the degrees of freedom, these are the sum of squares in a general table this sum of the squares comes here degrees of freedom comes here, you have the mean square error which is division of the sum of squares by the degrees of freedom with a value and the level of significance. So, based on that you can calculate and do the ANOVA table for the regression model also.

(Refer Slide Time: 05:47)

EXAMPLE 10.3 A 2^3 Factorial Design with a Missing Observation

Consider the 2^3 factorial design with four center points from Example 10.2. Suppose that when this experiment was performed, the run with all variables at the high level (run 8 in Figure 10.5) was missing. This can happen for a variety of reasons; the measurement system can produce a faulty reading, the combination of factor levels may prove infeasible, the experimental unit may be damaged, and so forth. We will fit the main effects model

$$y = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \beta_3 x_3 + \epsilon$$

using the 11 remaining observations. The X matrix and y vector are

$$X = \begin{bmatrix} 1 & -1 & -1 & -1 \\ 1 & 1 & -1 & -1 \\ 1 & -1 & 1 & -1 \\ 1 & 1 & 1 & -1 \\ 1 & -1 & -1 & 1 \\ 1 & 1 & -1 & 1 \\ 1 & -1 & 1 & 1 \\ 1 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 \end{bmatrix} \text{ and } y = \begin{bmatrix} 32 \\ 46 \\ 57 \\ 65 \\ 36 \\ 48 \\ 57 \\ 50 \\ 44 \\ 53 \\ 56 \end{bmatrix}$$

To estimate the model parameters, we form

$$X'X = \begin{bmatrix} 11 & -1 & -1 & -1 \\ -1 & 7 & -1 & -1 \\ -1 & -1 & 7 & -1 \\ -1 & -1 & -1 & 7 \end{bmatrix} \text{ and } X'y = \begin{bmatrix} 544 \\ -23 \\ 17 \\ -59 \end{bmatrix}$$

2 to the power 3 factorial, factorial design with missing observations.

So, you are mixing the observations consider the 2^2 the power 3 factorial design with 4 centre points from same examples suppose that when this example was performed the run with all the variables at high level run were missing. This can happen for a variety reason that measurement errors was there, measurement system can produce a faulty reading, the combination the factor level are not feasible infeasible they were not tempered or tuned as for the experiment and so on and so forth or the unit was damage, we did not get that unit or we miss the unit or the person who was taking the reading was not attentive on all this reasons can be there.

Again the same model, y is equal to β_0 , but β_1 into x_1 plus β_2 into x_2 , plus β_3 into x_3 plus epsilon again is the multiple linear regression model. β_0 , β_1 , β_2 , β_3 are the regression coefficients with respect to β_0 is when is basically the point at which where it is cutting the y axis. β_1 , β_2 , β_3 are the partial regression coefficients or the rate of change of y with respect to x_1 , x_2 , x_3 respectively. Using the 11 remaining observations the x matrix is of this form which is basically n cross the 4, 4 means β_0 , β_1 , β_2 , β_3 .

So, we so if there are x_1 , x_2 , x_3 . So, it will be n plus 4 because the fourth one is basically coming for β_0 , y values again n plus 1 the x , transpose x values is found out here which I am just circling you can use simple excel sheet or a mini tab or a mat lab to find it, then $x^T y$ is given from that you can find out $\hat{\beta}$. So, $\hat{\beta}$ when I am mentioning betas or the x s of the y s are all vectors, matrices, vectors, scalar would only come for the case when you are trying to find out the sigma square.

(Refer Slide Time: 07:58)

Because there is a missing observation, the design is no longer orthogonal. Now

$$\hat{\beta} = (X'X)^{-1}X'y$$

$$= \begin{bmatrix} 9.61538 \times 10^{-2} & 1.92307 \times 10^{-2} \\ 1.92307 \times 10^{-2} & 0.15385 \\ 1.92307 \times 10^{-2} & 2.88462 \times 10^{-2} \\ 1.92307 \times 10^{-2} & 2.88462 \times 10^{-2} \end{bmatrix} \begin{bmatrix} 544 \\ -23 \\ 17 \\ -59 \end{bmatrix}$$

$$= \begin{bmatrix} 51.25 \\ 5.75 \\ 10.75 \\ 1.25 \end{bmatrix}$$

Therefore, the fitted model is

$$\hat{y} = 51.25 + 5.75x_1 + 10.75x_2 + 1.25x_3$$

Compare this model to the one obtained in Example 10.2, where all 12 observations were used. The regression coefficients are very similar. Because the regression coefficients are closely related to the factor effects, our conclusions would not be seriously affected by the missing observation. However, notice that the effect estimates are no longer orthogonal because $X'X$ and its inverse are no longer diagonal. Furthermore the variances of the regression coefficients are larger than they were in the original orthogonal design with no missing data.

So, beta because there is missing observation the design is a no normal orthogonal when you find it out the beta hat values comes out to be 51.25 this is beta naught hat beta 1 hat is 5.75, beta 2 hat is 10.75 beta 3 hat is 1.25 therefore, the fitted model is ah; obviously, they would not be any error term.

So, if you want to find out the nth plus 1 value it will be y hat nth plus 1 which is in the suffix would be 51.25 which is beta naught hat plus 5.75 which is beta 1 hat into x 1. This is the first value is the first x value and its corresponding reading is the nth plus 1, then in plus beta 2 hat which is 10.25 into x 2 and it is corresponding value is the nth plus 1 for the second variable, plus beta 3 hat which is 1.25 multiplied by x 3 and it is corresponding x 3 value being nth plus 1.

So, from that we can find out y hat, then find of the difference between y and y hat for the nth plus 1 nth plus 2 nth plus 3, n is basically depending on at which stage you are which reading you are, find out the errors plot the errors with respect to x 1, x 2, x 3 find out the averages values come out to be 0 yes or no. Then plot the values of errors which should be normally distributed with respect to the standard normal distribution you use the q q plots, plot it along the 45 degrees line check whether they are there. Obviously, they would not be all lying on the 45 degrees line, but you can get an answer for that, compare this model to the obtained earlier the regression coefficients are very similar

because the regression coefficients are closely related to the factors effects. So, what is the effect of the x_1 and y , what is the effect of x_2 and y , what is effect of x_3 and y ?

So, they would be given by β_1 , β_2 , β_3 , but we are using the small sample whatever it is to predict the values using the estimate $\hat{\beta}_1$, $\hat{\beta}_2$, $\hat{\beta}_3$, obviously, we will use $\hat{\beta}$ to find out β . However, notice that the effects estimates are no longer orthogonal because $X^T X$ and it is inverse on no longer orthogonal for the more variances of the regression coefficients are larger than then when it was the original data because you have basically shrunk your sample size to smaller simple size.

(Refer Slide Time: 10:21)

When running a designed experiment, it is sometimes difficult to reach and hold the precise factor levels required by the design. Small discrepancies are not important, but large ones are potentially of more concern. Regression methods are useful in the analysis of a designed experiment where the experimenter has been unable to obtain the required factor levels.

To illustrate, the experiment presented in Table 10.5 shows a variation of the 2^3 design from Example 10.2, where many of the test combinations are not exactly the ones specified in the design. Most of the difficulty seems to have occurred with the temperature variable.

EXAMPLE 10.4 Inaccurate Levels in Design Factors

TABLE 10.5
Experimental Design for Example 10.4

Run	Process Variables			Coded Variables			Yield Y
	Temp (°C)	Pressure (psig)	Time (hr)	X_1	X_2	X_3	
1	125	40	14	-0.75	-0.95	-1.133	32
2	150	40	15	0.00	-1	-1	46
3	125	62	15	-0.95	1.1	-1	57
4	140	60	15	1	1	-1	65
5	148	50	15	-1.10	-1.05	1.14	36
6	160	40	30	1.15	-1	1	44
7	122	40	30	-0.90	1	1	57
8	145	63	30	1.25	1.15	1	68
9	140	60	22.5	0	0	0	50
10	140	40	22.5	0	0	0	44
11	140	60	22.5	0	0	0	53
12	140	40	22.5	0	0	0	56

When running a design experiment is sometimes difficult to reach and hold the precise factor levels required by the design, small description are not important.

So; obviously, they would be, but large once are potentially of more concern to us and we should be careful, regression methods are useful in the analysis of the designed experiments. Where the experimental has been unable to obtain the required factor levels at what factor they are, to illustrate example where very many of the test combinations are not exactly the one specified the in design most of the difficulty seems to have been occurred with the temperature variable.

So, again I am noting down or presenting to you the experimental design for the example where you have the run given in the first column, the temperature which is the x_1 variable at the second, pressure which is x_2 in the third and concentration which is x_3 on the fourth, the coded variables are basically done accordingly such that you can find out the average values and noted down and the y yield which is the dependent variable which is given on the last column.

(Refer Slide Time: 11:29)

We will fit the main effects model

$$y = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \beta_3 x_3 + e$$

The X matrix and y vector are

$$X = \begin{bmatrix} 1 & -0.75 & -0.95 & -1.133 \\ 1 & 0.90 & -1 & -1 \\ 1 & -0.95 & 1.1 & -1 \\ 1 & 1 & 0 & -1 \\ 1 & -1.10 & -1.05 & 1.4 \\ 1 & 1.15 & -1 & 1 \\ 1 & -0.90 & 1 & 1 \\ 1 & 1.25 & 1.15 & 1 \\ 1 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 \end{bmatrix} \quad y = \begin{bmatrix} 32 \\ 46 \\ 57 \\ 65 \\ 36 \\ 48 \\ 57 \\ 68 \\ 50 \\ 44 \\ 53 \\ 56 \end{bmatrix}$$

To estimate the model parameters, we need

$$X'X = \begin{bmatrix} 12 & 0.60 & 0.25 & 0.2670 \\ 0.60 & 8.18 & 0.31 & -0.1403 \\ 0.25 & 0.31 & 8.5375 & -0.3437 \\ 0.2670 & -0.1403 & -0.3437 & 9.2437 \end{bmatrix}$$
$$X'y = \begin{bmatrix} 612 \\ 77.55 \\ 100.7 \\ 19.144 \end{bmatrix}$$

We will fit the main effect model again the same model y is equal to β_0 plus β_1 into x_1 plus β_2 into x_2 , plus β_3 into x_3 plus ϵ and $\hat{y}$ would basically be is equal to $\hat{\beta}_0$, plus $\hat{\beta}_1$ into x_1 plus $\hat{\beta}_2$ into x_2 plus $\hat{\beta}_3$ into x_3 plus; obviously, there is no error.

So, X matrix is given the y matrix is given from this we find out the X' and $X'X$ from that we find out the $\hat{\beta}$ vector.

(Refer Slide Time: 12:03)

Then

$$\hat{\beta} = (X'X)^{-1}X'y$$

$$= \begin{bmatrix} 8.37447 \times 10^{-7} & -6.09871 \times 10^{-7} & -2.33542 \times 10^{-7} & -2.59833 \times 10^{-7} \\ -6.09871 \times 10^{-7} & 0.12289 & -4.20766 \times 10^{-7} & 1.88490 \times 10^{-7} \\ -2.33542 \times 10^{-7} & -4.20766 \times 10^{-7} & 0.11753 & 4.37851 \times 10^{-7} \\ -2.59833 \times 10^{-7} & 1.88490 \times 10^{-7} & 4.37851 \times 10^{-7} & 0.10845 \end{bmatrix} \begin{bmatrix} 612 \\ 77.55 \\ 100.7 \\ 19.144 \end{bmatrix} = \begin{bmatrix} 50.49391 \\ 5.40996 \\ 10.16316 \\ 1.07245 \end{bmatrix}$$

The fitted regression model, with the coefficients reported to two decimal places, is:

$$\hat{y} = 50.49 + 5.41x_1 + 10.16x_2 + 1.07x_3$$

Comparing this to the original model in Example 10.2, where the factor levels were exactly those specified by the design,

we note very little difference. The practical interpretation of the results of this experiment would not be seriously affected by the inability of the experimenter to achieve the desired factor levels exactly.

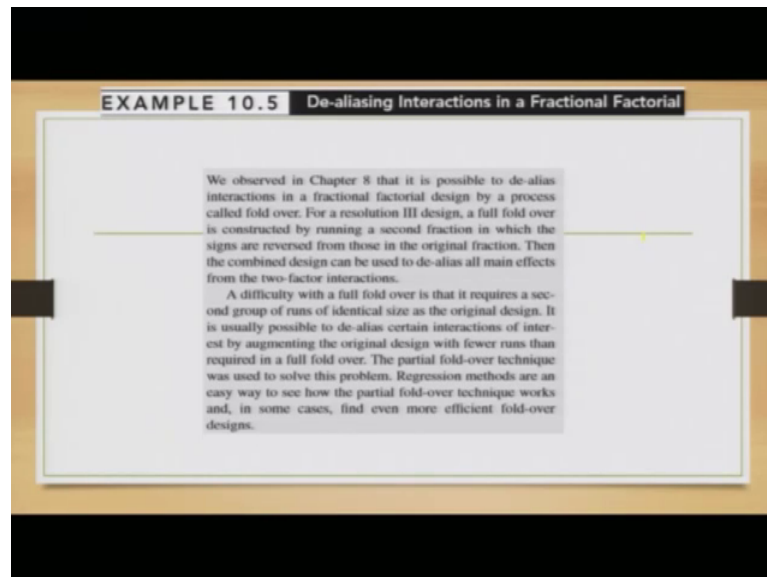
So, vector comes out to be, after all the calculations because the formula or which I am circling always remain the same you do not have to worry ones, actually the basic concept is sum of the squares difference partially differentiate with respect to betas put it to 0 those are the orthogonal equations find out the hats and after that your job is done you proceed with the calculations.

The values of beta naught is given as beta naught hat is 50.49 I am only reading till the second place of decimal beta 1 hat is 5.40, beta 2 hat is 10.16, beta 3 hat is equal to 1.07. The fitted regression model now would be y hat is equal to 50.49 plus 5.40 into x 1 plus 10.16 into x 2, plus 1.07 in to x 3. So, this is what is given, comparing this is with the original model where the factor levels were exactly those specified by the design.

We note a very very little difference, the practical interpretation of the results of the experiment would not be seriously affected by this, hence we are able to utilise this model convincingly. So, even you even if there is inability on the experimental side like missing get was there or the machine was not working to measure rate of the temperature difference were there still there is no problem, because errors are very small errors with respect to the betas and the error terms. And; obviously, we would difilt for all these things I am, I am skipping it for the last example we will draw the q q plots draw the error term on the y axis with be respect to x 1, x 2, x 3 and do all the comparison to find out the mean of the errors is 0.

The variability of the errors should be normal a sigma square hat which want to find out and the normality conditions would also hold true we will also find out the covariance of the beta vector. So, where the principle diagram would be the square of the standard error of the betas and of the diagonal element with the covariance existing between the betas.

(Refer Slide Time: 14:09)



We observed that it is possible to de alias the interaction in the fractional factorial design by a process called folded over.

So, you fold it and try to basically consider in using blocks, block design also for resolution of 3 degrees or 4 degrees a full fold over is constructed by running a second factorial in which the signs are reversed from those in original factors. Then the combined design can be used to de alias all main effects from the 2 factor interactions and basically have a better picture on how the effects are being done. Now, a difficulty was the, with a full fold over is that it requires a second group of runs of identical sizes.

But there that data size may not be applicable because the number of observations which we have for the sample is low, it is usually possible to de alias certain interaction of interest by augmenting the original design with fewer runs than required in a full fold over, the partial fold over technique was used to solve this problem because we would not to do the folds in at one go for considering all of the observations at 1 go, we need to partially break into samples and then do the folds.

The regression methods are an easy way to see how the partial fold over technique works and in some cases find even more efficient fold over design methods.

(Refer Slide Time: 15:25)

To illustrate, suppose that we have run a 2^{4-1}_{III} design. Table 8.3 shows the principal fraction of this design, in which $I = ABCD$. Suppose that after the data from the first eight trials were observed, the largest effects were A , B , C , D (we ignore the three-factor interactions that are aliased with these main effects) and the $AB + CD$ alias chain. The other two alias chains can be ignored, but clearly either AB , CD , or both two-factor interactions are large. To find out which interactions are important, we could, of course, run the alternate fraction, which would require another eight trials. Then all 16 runs could be used to estimate the main effects and the two-factor interactions. An alternative would be to use a partial fold over involving four additional runs.

It is possible to de-alias AB and CD in fewer than four additional trials. Suppose that we wish to fit the model

$$y = \beta_0 + \beta_1x_1 + \beta_2x_2 + \beta_3x_3 + \beta_4x_4 + \beta_{12}x_1x_2 + \beta_{13}x_1x_3 + \beta_{14}x_1x_4 + \beta_{23}x_2x_3 + \beta_{24}x_2x_4 + \beta_{34}x_3x_4 + \epsilon$$

where x_1 , x_2 , x_3 , and x_4 are the coded variables representing A , B , C , and D . Using the design in Table 8.3, the X matrix for this model is

	x_1	x_2	x_3	x_4	x_1x_2	x_1x_3	x_1x_4	x_2x_3	x_2x_4	x_3x_4
1	1	-1	-1	-1	-1	1	1	1	-1	-1
2	1	1	-1	-1	1	-1	-1	-1	1	1
3	1	-1	1	-1	1	1	-1	1	-1	1
4	1	1	1	-1	-1	1	1	-1	1	-1
5	1	-1	-1	1	-1	-1	1	1	-1	1
6	1	1	-1	1	1	-1	-1	-1	1	-1
7	1	-1	1	1	1	1	-1	-1	-1	1
8	1	1	1	1	-1	-1	1	1	1	-1

where we have written the variables above the columns to facilitate understanding. Notice that the x_1x_2 column is identical to the x_3x_4 column (as anticipated, because AB or x_1x_2 is aliased with CD or x_3x_4), implying a linear dependency in the columns of X . Therefore, we cannot estimate both β_{12} and β_{34} in the model. However, suppose that we add a single run $x_1 = -1$, $x_2 = -1$, $x_3 = -1$, and $x_4 = 1$ from the alternate fraction to the original eight runs.

To illustrate suppose that we have a model of 2 to the power 4 minus 1 and 4 third design model and it shows the that in which we considered the I a matrix or I value as the as the factors ABCD, suppose that after the data from the first 8 trials were observed the largest effects were basically ABCD.

So, we ignore the 3 factor interactions and A plus B, AB plus CD alias chains are ignored, the other 2 aliases can also be ignored, but clearly the AB or AB or both of 2 factors are large hence they cannot be ignored, to find out which interactions are important we could of course, run the on the alternative fraction which would require another 8 trials. Then all these 16 trials could be used to estimate the main effect and the 2 factor interaction and alternative would be to use a partial fold over involving 4 additional runs and we can complete the task accordingly.

So, again now, as we are considering the partial fold model so; obviously, the multiple linear regression would be larger in size and what is that please pay attention to here, it is possible to de alias AB and CD in fewer than 4 additional trails. Hence, we will assume the model as such initially what we had we had y is equal to because it was there was no fold. So, it was y is equal to beta naught plus beta 1, x 1 plus beta 2 x 2 plus beta 3 x 3 plus epsilon.

Now, the model is changing, what is it please pay attention it will be y is equal β_0 plus β_1 what β_1 into x_1 plus β_2 x_2 plus β_3 x_3 this remains as it is. Now, we are basically considering those extra factors as it should be as per the fractional factorial modules considering the partial folds, the fifth term would be x_4 into β_4 plus β_{12} where you are combining the effect of A and B which is β_1 into x_2 and the and th last term would be apart from the error term would be β_{34} into β_{34} which will be x_3 into x_4 plus epsilon.

Where the the factors x_1, x_2, x_3, x_4 are the coded variable represents ABCD and the combination of x_1 x_2 or x_3 x_4 with the combination of AB and CD. Using the design the x x matrix now be given by n remains the same, but the factors which is n cross m k plus one or n cross k plus 2 whatever it is would now depend on the factors.

Which we have consider, what are the factor it is x_1, x_2, x_3, x_4, x_1 and x_2 which is the fifth one and x_3 x_4 which is a sixth one. So; obviously, it would be n cross 7 because the 7 may, seventh means which is the first one would be β_0 . So, there would be $\beta_1, \beta_2, \beta_3, \beta_4, \beta_{12}, \beta_{34}$, where we have written the variables above the columns to facilitate better understanding noticed there is that x_1 x_2 column is identical to x_3 x_4 column ah, because AB or alias CD are coded as x_1 x_2 and x_3 x_4 respectively implying a linear dependence model. Therefore, we cannot estimate both x_1, x_2 and x_3, x_4 in the model, I was suppose that we add additional run and basically consider the x_1 x_2 as different levels.

We can basically have alternate fractions to the original 8 model, 8 runs and basically be able to estimate $\beta_0, \beta_1, \beta_2, \beta_3, \beta_4, \beta_{12}$ and β_{34} .

(Refer Slide Time: 18:58)

The X matrix for the model now becomes

$$X = \begin{bmatrix} 1 & -1 & -1 & -1 & -1 & 1 & 1 \\ 1 & 1 & -1 & -1 & 1 & -1 & -1 \\ 1 & -1 & 1 & -1 & 1 & -1 & -1 \\ 1 & 1 & 1 & -1 & -1 & 1 & 1 \\ 1 & -1 & -1 & 1 & 1 & 1 & 1 \\ 1 & 1 & -1 & 1 & -1 & -1 & -1 \\ 1 & -1 & 1 & 1 & -1 & -1 & -1 \end{bmatrix}$$

Notice that the columns x_1x_2 and x_1x_3 are now no longer identical, and we can fit the model including both the x_1x_2 (AB) and x_1x_3 (CD) interactions. The magnitudes of the regression coefficients will give insight regarding which interactions are important.

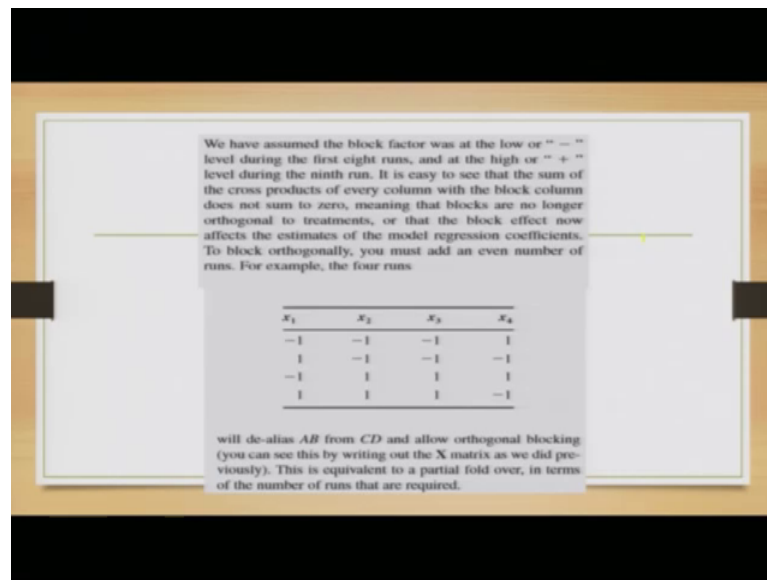
Although adding a single run will de-alias the AB and CD interactions, this approach does have a disadvantage. Suppose that there is a time effect (or a block effect) between the first eight runs and the last run added above. Add a column to the X matrix for blocks, and you obtain the following:

$$X = \begin{bmatrix} 1 & -1 & -1 & -1 & -1 & 1 & 1 & 1 \\ 1 & 1 & -1 & -1 & 1 & -1 & -1 & 1 \\ 1 & -1 & 1 & -1 & 1 & -1 & -1 & 1 \\ 1 & 1 & 1 & -1 & -1 & 1 & 1 & 1 \\ 1 & -1 & -1 & 1 & 1 & 1 & 1 & 1 \\ 1 & 1 & -1 & 1 & -1 & -1 & -1 & 1 \\ 1 & -1 & 1 & 1 & -1 & -1 & -1 & 1 \\ 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \end{bmatrix}$$

Once you have the other matrix it is again of size n cross, I am repeating the word n in order to me basically denote the number of observations, n cross 7. Notice that the column x 1, x 2, x 3 and x 4 are no longer, are now no longer identical and you can fit the model increasing and find to try to find out the estimates of beta 1 2 hat and beta 3 4 hat 3 4 basically this is a suffixes, the magnitude of the regression coefficients will give inside regarding which interactions an important.

So, once we do, do that we can basically do the same methodology we are trying to basically find out the beta hats, try to find out the errors, try to fit in the model to find out the average of the errors, try to fit and find out the distribution of the errors and everything and would basically remain the same. And; obviously, we will come into the main table which is for the ANOVA over where the reading on the variables of the attributes, I am using the attributes because it can be attributes and variables also depending on what type of model you are considering any of the sum of the squares you have the degrees of freedom you have the f values, the mean square value and values the f value and basically you have the level of significance.

(Refer Slide Time: 20:09)



We have assumed the block factor was at the low or " - " level during the first eight runs, and at the high or " + " level during the ninth run. It is easy to see that the sum of the cross products of every column with the block column does not sum to zero, meaning that blocks are no longer orthogonal to treatments, or that the block effect now affects the estimates of the model regression coefficients. To block orthogonally, you must add an even number of runs. For example, the four runs

x_1	x_2	x_3	x_4
-1	-1	-1	1
1	-1	-1	-1
-1	1	1	1
1	1	1	-1

will de-alias AB from CD and allow orthogonal blocking (you can see this by writing out the X matrix as we did previously). This is equivalent to a partial fold over, in terms of the number of runs that are required.

We have assume the block factors were at the level of minus in the first 8 runs then at the high level of during the ninth run it is easy to see that the sum of the cross product of every column with the block columns does not sum to 0, meaning the block are longer orthogonal to treatment or that the block effect.

Now, effects the estimate the model regression coefficients, to block orthogonality we must add an even number of runs for example, we have extra 4 runs where it will give us because there are what, there are x_1 , x_2 and x_3 , x_4 combined which is AB and CD . So, hence we basically add up more number of combinations for that, such that we can basically do the calculations and try to basically differentiate the effects in much better fashion.

So, with the analysis for AB which is x_1 , x_2 and CD which is x_3 , x_4 and for allow orthogonal blocking this is equivalent to a partial fold over in terms of the number of runs and that are required in order to estimate.

(Refer Slide Time: 21:08)

Hypothesis Testing in Multiple Regression: Test for Significance of Regression

The test for significance of regression is a test to determine whether a linear relationship exists between the response variable y and a subset of the regressor variables $x_1, x_2, \dots, x_k$. The appropriate hypotheses are

$$H_0: \beta_1 = \beta_2 = \dots = \beta_k = 0 \quad (10.20)$$
$$H_1: \beta_j \neq 0 \quad \text{for at least one } j$$

Rejection of H_0 in Equation 10.20 implies that at least one of the regressor variables $x_1, x_2, \dots, x_k$ contributes significantly to the model. The test procedure involves an analysis of variance partitioning of the total sum of squares SS_T into a sum of squares due to the model (or to regression) and a sum of squares due to residual (or error), say

$$SS_T = SS_R + SS_E \quad (10.21)$$

Now, you want to do that hypothesis testing, again we are going to come back to the ANOVA table. So, the test of significance for regression is a test to determine whether the linear relationship exists between the response variables y and a subset of the regressor variables $x_1, x_2, x_3, \dots, x_k$. Now, the actual hypothesis would be stated which is in H_0 case we will consider β_1 to β_k all to be 0 and in H_1 which is the alternative 1 one cases we will consider that for at least one of these j, j is 1 to k they are not 0.

So, in this case 0 means that the effects of them are 0; that means, we need you to the partial differentiation the rate of change of y with respect to β_1, x_1, x_2, x_3 whatever is 0; that means, they are not being affected. So, rejection of H_0 in equation in this equation implies that are at least one of the regression variables x_1 to x_3 contribute significance so; that means, if you are rejecting H_0 been basically b tech taking cognizance of the fact that H_0 is true.

So, if H_0 is true; that means, one of the betas are definitely not equal to 0, which means that there is a linear relationship between y and that particular x . The test procedure involves an analysis of variance exactly as ANOVA table where we find out the total number of squares into the sum of the squares due to the model and due to the sum of the errors total sum is equal to sum of the variables plus sum of the errors. So, this is the, this

is the main thing which you have been discussing and from there we will decide we do the ANOVA table.

(Refer Slide Time: 22:48)

Now if the null hypothesis $H_0: \beta_1 = \beta_2 = \dots = \beta_k = 0$ is true, then SS_R/σ^2 is distributed as χ^2_k , where the number of degrees of freedom for χ^2 is equal to the number of regressor variables in the model. Also, we can show that SS_R/σ^2 is distributed as χ^2_{n-k-1} and that SS_R and SS_E are independent. The test procedure for $H_0: \beta_1 = \beta_2 = \dots = \beta_k = 0$ is to compute

$$F_0 = \frac{SS_R/k}{SS_E/(n-k-1)} = \frac{MS_R}{MS_E} \quad (10.22)$$

and to reject H_0 if F_0 exceeds $F_{\alpha, k, n-k-1}$. Alternatively, we could use the P -value approach to hypothesis testing and, reject H_0 if the P -value for the statistic F_0 is less than α . The test is usually summarized in an analysis of variance table such as Table 10.6.

TABLE 10.6
Analysis of Variance for Significance of Regression in Multiple Regression

Source of Variation	Sum of Squares	Degrees of Freedom	Mean Square	F_0
Regression	SS_R	k	MS_R	MS_R/MS_E
Error or residual	SS_E	$n - k - 1$	MS_E	
Total	SS_T	$n - 1$		

Now, if the full a null hypothesis which is H_0 is 2 then SS_R which is for the regression models divided by sigma squared is distributed to chi square, because chi square is basically distribution will consider with respect of sigma square and the standard error distribution, which we are considering where the number of degrees of freedom of chi square is equal to the number of degrees variables in the model.

Also we know that the sum of the squares of the errors by sigma square is distributed, again chi square with the certain change in the degrees of freedom and that sum of the squares of the errors and sum of the squares of the total effects are independent. That is what is very important to note down, considering the fact that we are considering normal days, normal distribution to be true for all the x s errors and y , this is very important to be noted by all of you for basically listening to this lecture and who are taking this course.

The test procedure would be basically for H_0 would be to find out if F statistics is a sum of the squares of the product of the case that by k , with k is the degree of freedom depending on the number of variables divided by sum of the square of errors is divided by it is degree of freedom. That will give me the mean square of r by divided by mean square of the errors and we will reject if it H_0 , F_0 under H_0 exceeds

that value F_{α} , α is the level of significance which you have alternatively we could use the p values to approach to test the hypothesis.

And reject F_{naught} with p value of the statistics F_{naught} is less than or if it is greater than α we accept that α can be 1 percent, can be 5 percent, can be 10 percent so on and so forth. The test statistic is usually summarized again in same ANOVA table. So, what we have is this here the regression errors in the residual are given in the first column, sum of the squares is given, degrees of freedom is given divide SS R by k is given, then the mean squared errors SS e by n minus k minus 1 will give me the mean squared errors based on that you can find out the SS F statistics and then we can comment with whether it is right or wrong depending on the α value which is the level of significance.

(Refer Slide Time: 25:08)

the regression sum of squares is

$$SS_R = \hat{\beta}'X'y - \frac{\left(\sum_{i=1}^n y_i\right)^2}{n} \quad (10.23)$$

and the error sum of squares is

$$SS_E = y'y - \hat{\beta}'X'y \quad (10.24)$$

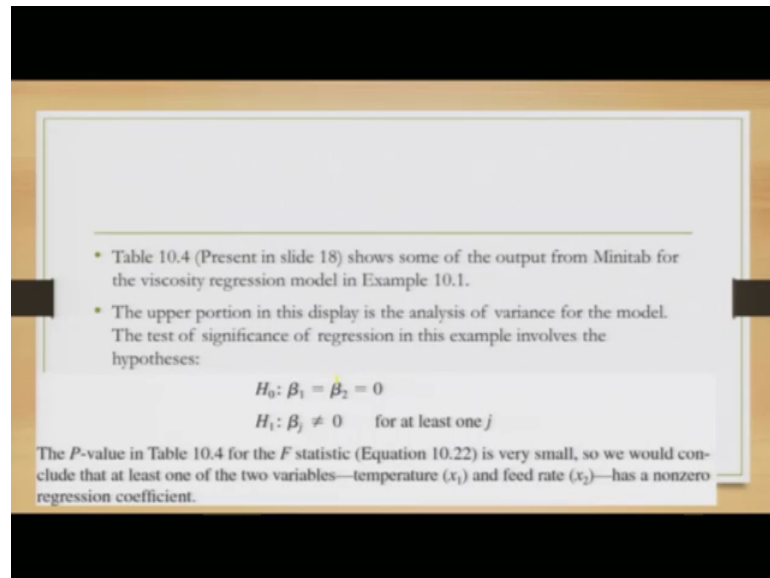
and the total sum of squares is

$$SS_T = y'y - \frac{\left(\sum_{i=1}^n y_i\right)^2}{n} \quad (10.25)$$

The regressions on the squares by the formula is given by beta hat transpose. So, beta hat what basically of size k plus 1 plus 1. So, it transpose would be just the opposite, multiplied by x transpose y minus the overall effect of y s, y s are the actual value divided by this.

So, this is the basically the best fit line which you have for the actual value and the sum of the squares would be given by, this is the total error minus this would give me the error sum of the square of the error. So, total value would be given by addition of them and you can find out the calculations calculate the values accordingly.

(Refer Slide Time: 25:53)



- Table 10.4 (Present in slide 18) shows some of the output from Minitab for the viscosity regression model in Example 10.1.
- The upper portion in this display is the analysis of variance for the model. The test of significance of regression in this example involves the hypotheses:
$$H_0: \beta_1 = \beta_2 = 0$$
$$H_1: \beta_j \neq 0 \quad \text{for at least one } j$$

The P -value in Table 10.4 for the F statistic (Equation 10.22) is very small, so we would conclude that at least one of the two variables—temperature (x_1) and feed rate (x_2)—has a nonzero regression coefficient.

So, in the in the this following slide it will shows the sum of the output mini tab for the viscosity regression model for the upper portion of the displays is the analysis and variance for the model.

The test of significance of the regression in this example involves the hypothesis, that beta 1 is equal to beta 2 is given 0 and either and in the alternative one would be at least one of them is not 0, the p value of the table for the f statistic very small. So, you would conclude that at least one of the variables which is temperature or feed rate have a non 0 value hence the linearity relationship holds.

(Refer Slide Time: 26:31)

R^2

Table 10.4 also reports the coefficient of multiple determination R^2 , where

$$R^2 = \frac{SS_R}{SS_T} = 1 - \frac{SS_E}{SS_T} \quad (10.26)$$

Just as in designed experiments, R^2 is a measure of the amount of reduction in the variability of y obtained by using the regressor variables $x_1, x_2, \dots, x_k$ in the model. However, as we have noted previously, a large value of R^2 does not necessarily imply that the regression model is a good one. Adding a variable to the model will always increase R^2 , regardless of whether the additional variable is statistically significant or not. Thus, it is possible for models that have large values of R^2 to yield poor predictions of new observations or estimates of the mean response.

The table also gives a R square value, which is that adjusted R square and adjusted R square, R square would be the sum of the squares which are we are able to predict from the variables divided by sum of the total of the total errors. So, that would also be equal to 1 minus, because the ratio is always one for the total was the total minus the sum of the square by the error by the total.

Just in design experiment R square is major amount of deduction the variability of y obtained by using the regression variables x_1, x_2, x_3 to x_k ; however, as we have noted previously a large value R square does not necessary imply the regression model is a good one. This is interesting and listen to it adding a variable in the model will always increase r square; obviously, because we are adding more and more variables we are able to think that we are able to think more and more. Regardless of whether the additional variable in statistically significant or not, thus is it is possible for models that have large R values to yield poor prediction of the new observation or estimates of the mean values as we proceed for the n th plan predictions to n th predictions and so on and so forth.

So, basically we should take the minimum number of x s try to predict the maximum, basically have a value of R square where it will be able to give us the predicted errors as least as possible for the for the n th plus 1 n th plus 2, n th plus 3 data.

(Refer Slide Time: 28:00)

Because R^2 always increases as we add terms to the model, some regression model builders prefer to use an **adjusted R^2 statistic** defined as

$$R^2_{adj} = 1 - \frac{SS_E(n-p)}{SS_T(n-1)} = 1 - \left(\frac{n-1}{n-p}\right)(1 - R^2) \quad (10.27)$$

In general, the adjusted R^2 statistic will not always increase as variables are added to the model. In fact, if unnecessary terms are added, the value of R^2_{adj} will often decrease.

Because R square always increases as we had terms in the model some regression models builders prefer to use adjusted R square which I mentioned is basically by dividing by the degrees of freedom. So and adjusted r square would be given 1 minus so, initially it was sum of the square of the errors divided by same squares of the total.

Now, it will basically be n minus 1 by n minus p p is basically k plus 1 depending of beta not is there or not multiplied by R 1 minus R square in general adjusted R square statically R not always increases variables are adjusted to the model. In fact, if necessary terms are added, the value of R square adjusted would often decrease and basically give us an actual value of the prediction level of the model.

(Refer Slide Time: 28:41)

Tests on Individual Regression Coefficients and Groups of Coefficients

Adding a variable to the regression model always causes the sum of squares for regression to increase and the error sum of squares to decrease. We must decide whether the increase in the regression sum of squares is sufficient to warrant using the additional variable in the model. Furthermore, adding an unimportant variable to the model can actually increase the mean square error, thereby decreasing the usefulness of the model.

The hypotheses for testing the significance of any individual regression coefficient, say β_j , are

$$H_0: \beta_j = 0$$
$$H_1: \beta_j \neq 0$$

So, now we will basically start of the test and individual regression coefficient and the group of coefficients. So, adding a variable regression model always causes the sum of squares for regression module. So, increase in the errors of the sum of of of the squares to decrease.

So, obviously, we want as x s are increasing R square would basically be give us the maximum output, but; obviously, the question remains that to what level are we able to predict for the n th plus 1 n th plus 2 data points that is point number 1. Whether the errors as per the assumptions have normally distributed with the 0 mean and certain variance whether the non on normality assumption holds and the covariance of the beta which you want to find out are really small or not.

Because if they are very large; obviously, apart from un biasedness the consistent property would not hold, we must decide whether the increase in the regression sum of the squares is sufficient to warrant using additional variable in the model. So, we will try to basically find out the hypothesis that other than testing the significance of any of them collecting as a group, we will try to take them as an individual and try to predict as maximum as possible.

With this I will end the 38 lecture and try to continue with this combinations of the hypothesis testing such that we take individually the beta and try to predict with H naught is true or H a is true and basically pass judgement accordingly.

Thank you very much and have a nice day.