

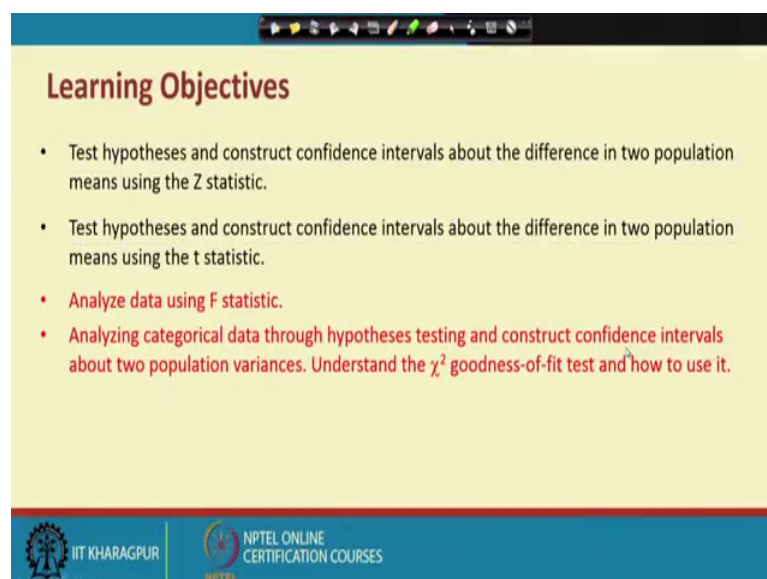
Business Analytics for Management Decision
Prof. Rudra P Pradhan
Vinod Gupta School of Management
Indian Institute of Technology, Kharagpur

Lecture – 25
Inferential Analytics Part – 2 (Contd.)

Hello everybody, this is Rudra Pradhan here, and welcome to BMD lecture series. And today we will continue the concept inferential analytics. And in fact, in the last couple of lectures we have already discussed this issue and it is the case where we are dealing with the business problems relating to multiple sample case. And in the last couple of lectures we have solved problems by using z statistic, t statistics and f statistic.

And today we will continue the similar kind of you know problems by using f statistics and typically we will focus on the concept called as you know ANOVA that is the analysis of variance. And today we will continue the ANOVA by connecting one way structures, two-way structures and three way structure. And finally, it will be continue the concept called as you know MANOV that is multivariate analysis of variance. And the concept is so interesting that if you have a multiple sample case so how we can actually check the kind of within the difference and between the difference, and with the different kind of you know layers or different kind of you know factors, in typical it is a kind of you know multivariate kind of structures.

(Refer Slide Time: 01:42)



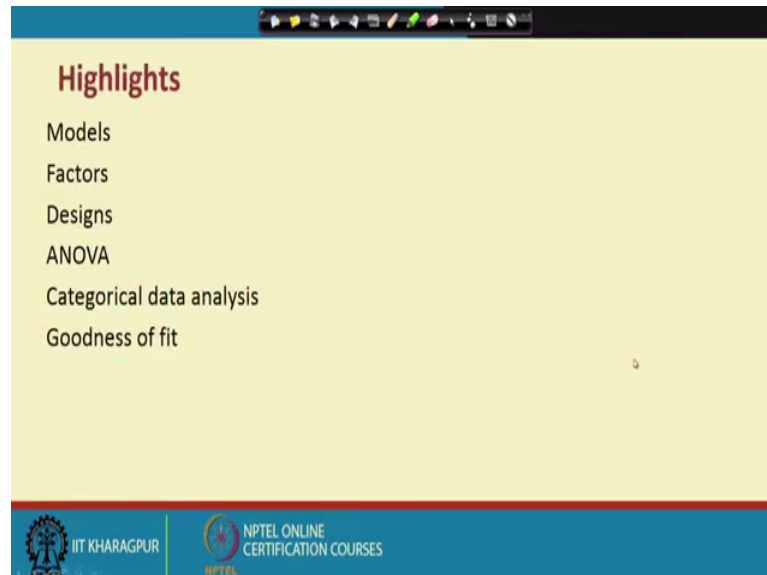
Learning Objectives

- Test hypotheses and construct confidence intervals about the difference in two population means using the Z statistic.
- Test hypotheses and construct confidence intervals about the difference in two population means using the t statistic.
- Analyze data using F statistic.
- Analyzing categorical data through hypotheses testing and construct confidence intervals about two population variances. Understand the χ^2 goodness-of-fit test and how to use it.

IIT KHARAGPUR | NPTEL ONLINE CERTIFICATION COURSES

So, let us see how is this particular concept? And then, we will discuss. So, this is what the kind of you know discussion today.

(Refer Slide Time: 01:46)



The highlights of this discussion or you know will be the ANOVA models. And followed by factors then designs and a we will connect with the concept that is ANOVA. And we will dealing with the problems, categorical data analysis and finally, the goodness of fit.

(Refer Slide Time: 02:08)

The slide is titled "Models, Factors and Designs" in a bold, dark red font. It defines a statistical model as a set of equations and assumptions that capture the essential characteristics of a real-world situation. It then introduces the one-factor ANOVA model with the equation $X_{ij} = \mu_i + \varepsilon_{ij} = \mu + \alpha_i + \varepsilon_{ij}$. Below the equation, it explains that ε_{ij} is the error associated with the j th member of the i th population, and the errors are assumed to be normally distributed with mean 0 and variance σ^2 . The slide has a light yellow background and a blue footer containing the IIT Kharagpur and NPTEL logos.

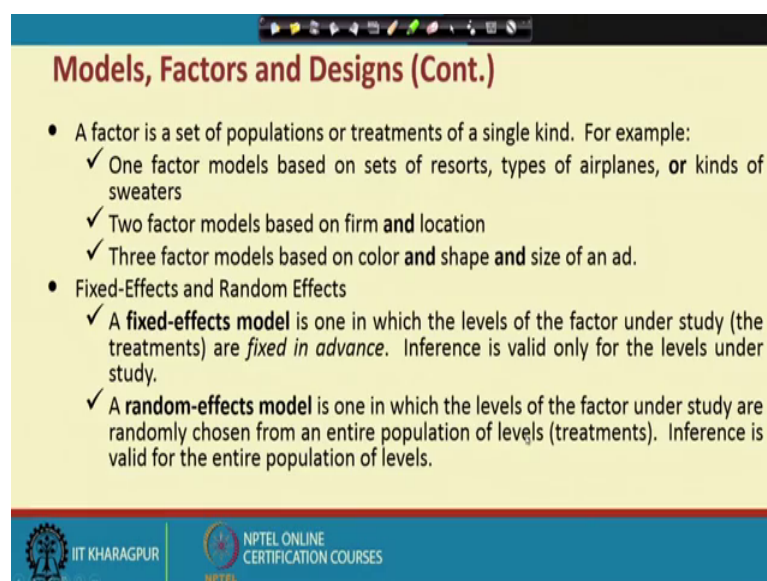
So, in the last couple of lectures, we have already discussed the concept typically the F statistic issue and in a kind of you know structure multiple sample cases. And here so we

will start with the same concept, but with a kind of you know different structure altogether. And let us start with a concept called as you know models. So, you know since it is a kind of multiple sample case or multivariate kind of you know structures, so there is a kind of you know typically you will finding relationships and that relationship is a kind of you know a complex kind of you know structure.

So, obviously, since there are a couple of variables and that to multiple sample case, so you cannot actually simply you know analyze the particular structure and get some kind of you know inference as for the business requirement. So, that is why we need to develop a models and through the models we have to analyze the particular you know problem and then we will get some kind of you know inference.

So, now as we have already discussed typically statistical model is a set of equations and assumption that capture the essential features of the real world situations. And the structure is you know we start with here kind of you know concept called as you know one-factor ANOVA model and the model will be like this you know x_{ij} equal to μ_i plus ϵ_{ij} . And so μ_i is the kind of you know concept called as you know sample means. And then ϵ_{ij} is the error associated with the j th member of the i th populations. And we are assuming that you know these errors are normally distributed with the mean 0 and unit variance that is the sigma square.

(Refer Slide Time: 04:00)



Models, Factors and Designs (Cont.)

- A factor is a set of populations or treatments of a single kind. For example:
 - ✓ One factor models based on sets of resorts, types of airplanes, or kinds of sweaters
 - ✓ Two factor models based on firm and location
 - ✓ Three factor models based on color and shape and size of an ad.
- Fixed-Effects and Random Effects
 - ✓ A **fixed-effects model** is one in which the levels of the factor under study (the treatments) are *fixed in advance*. Inference is valid only for the levels under study.
 - ✓ A **random-effects model** is one in which the levels of the factor under study are randomly chosen from an entire population of levels (treatments). Inference is valid for the entire population of levels.

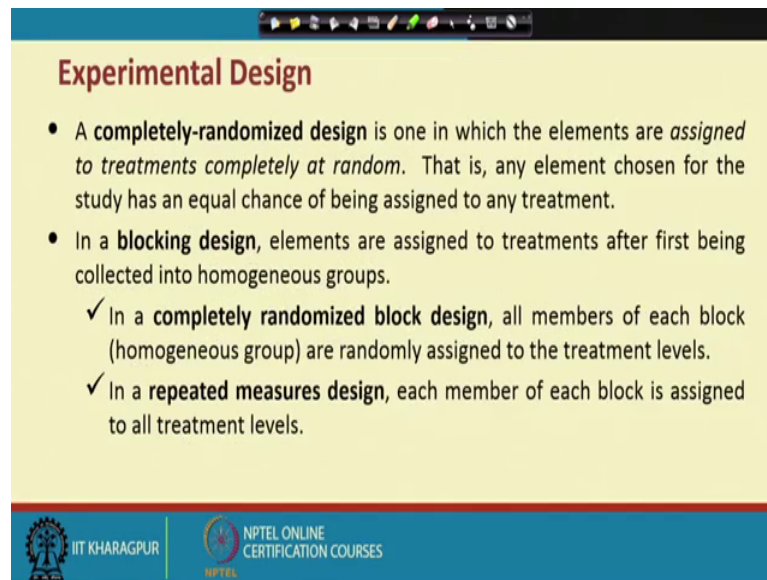
IIT KHARAGPUR | NPTEL ONLINE CERTIFICATION COURSES

And in the other side, a factor is a set of you know populations or you know kind of know treatments of a single kind. For examples, we have actually one-way structure, two-way structure, three-way structure. And in the case of you know one way structures that is called as you know one way typical you know factor models was done you know examples like you know hotels type of you know at aeroplanes and the kinds of you know colors, in a particular product. And two factors models based on business firm and the kind of you know site location; and the three factors models can be like you know color shape and size of a particular you know service or you know kind of you know product. But typically you know analysis will be on the basis of you know the concept called as you know fixed effect models and a random effect models.

In fact, this particular you know ANOVA is well connected with your concept called as you know panel data and a in the last couple of lectures we have already discussed this data issue. And then you know data can be actually connected or you know using a kind of you know business environment or you know business analytics in the form of time series cross sectional and you know penner. So, ANOVA is a kind of you know cluster like you know panel data structure all together. So, we have a separate lectures about the panel data modelling. But in the mean times we will analyze the similar kind of you know problems in a kind of you know structure called as you know of factor models that is one way factor model, two-way factor model and three way factor model depending upon the kind of you know structure of the business problem.

And here is we have actually typically two different you know models which you can connect with these particular you know situations. So, the first one is called as a fixed effect model, and the second one is called as a random effect model. In the fixed effect model the levels of the factor under study are fixed; and in the case of random effect models, the level of factors under the study are randomly chosen from an entire populations and that is the treatments. And the inference is valid for the entire populations for this particular you know study or analysis.

(Refer Slide Time: 06:20)



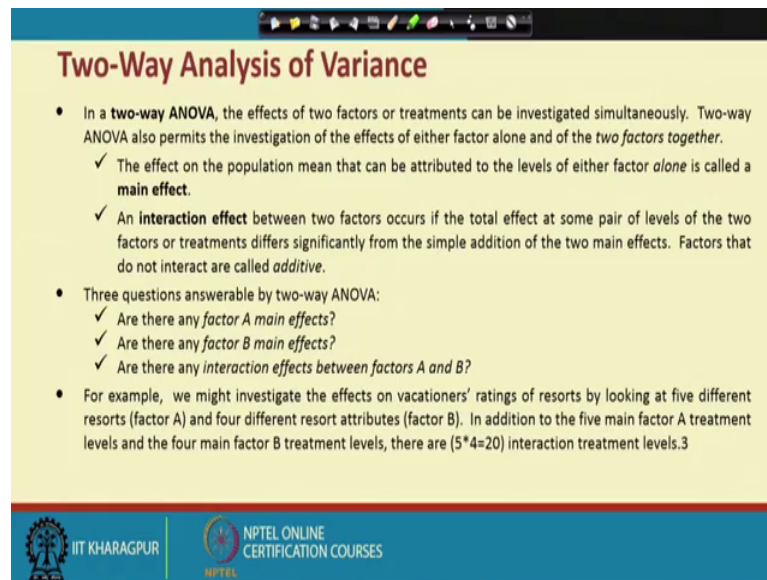
Experimental Design

- A **completely-randomized design** is one in which the elements are *assigned to treatments completely at random*. That is, any element chosen for the study has an equal chance of being assigned to any treatment.
- In a **blocking design**, elements are assigned to treatments after first being collected into homogeneous groups.
 - ✓ In a **completely randomized block design**, all members of each block (homogeneous group) are randomly assigned to the treatment levels.
 - ✓ In a **repeated measures design**, each member of each block is assigned to all treatment levels.

IIT KHARAGPUR | NPTEL ONLINE CERTIFICATION COURSES

And the third component is with respect to experimental design, and basically ANOVA can be connected with the design of experiments that typically called as you know DOE. And here a completely randomized design is one in which the elements are assigned to treatments completely at a random basis, and that is any element chosen for the study has an equal chance of being assigned to any treatment. And in the case of you know block designs elements are assigned to treatments after first being collected into homogeneous groups. So, in a completely randomized block designs, all members of each block that is the homogeneous groups are randomly assigned to the treatment levels. And in the case of you know repeated measures designs, each member of you know each block is assigned to all treatment levels.

(Refer Slide Time: 07:17)



The slide is titled "Two-Way Analysis of Variance" in a bold, dark red font. It contains a bulleted list of concepts and questions related to two-way ANOVA. The background is a light yellow color. At the bottom of the slide, there is a blue banner with the IIT Kharagpur logo and the NPTEL Online Certification Courses logo.

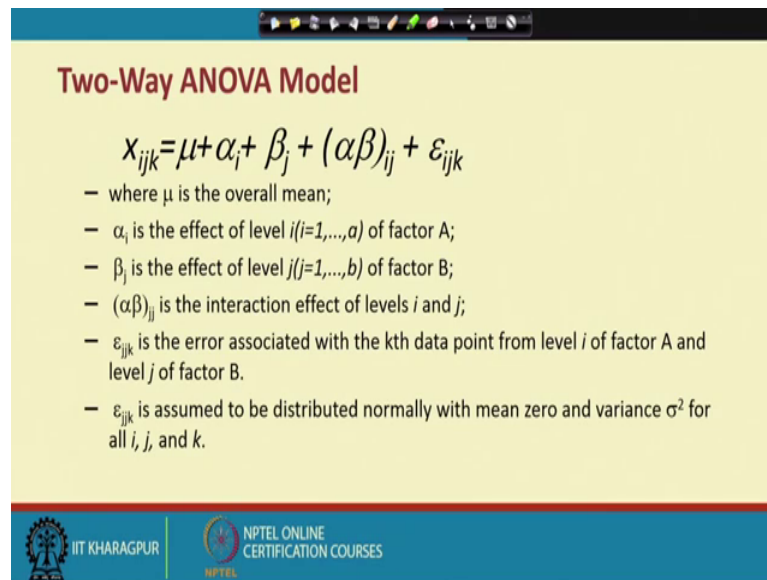
Two-Way Analysis of Variance

- In a **two-way ANOVA**, the effects of two factors or treatments can be investigated simultaneously. Two-way ANOVA also permits the investigation of the effects of either factor alone and of the *two factors together*.
 - ✓ The effect on the population mean that can be attributed to the levels of either factor *alone* is called a **main effect**.
 - ✓ An **interaction effect** between two factors occurs if the total effect at some pair of levels of the two factors or treatments differs significantly from the simple addition of the two main effects. Factors that do not interact are called *additive*.
- Three questions answerable by two-way ANOVA:
 - ✓ Are there any **factor A main effects**?
 - ✓ Are there any **factor B main effects**?
 - ✓ Are there any **interaction effects between factors A and B**?
- For example, we might investigate the effects on vacationers' ratings of resorts by looking at five different resorts (factor A) and four different resort attributes (factor B). In addition to the five main factor A treatment levels and the four main factor B treatment levels, there are $(5 \times 4 = 20)$ interaction treatment levels.³

In the case of you know two analysis of ANOVA typically the effect is with respect to two factors or it is with respect to treatment two treatments, so that means, typically actually we have a concept one factor model, two factor model, three factor model. In the case of one-factor models, so we have a one factor then the kind of you know different sample variation will be there; and then in the two factor models we have actually two different factor through which the investigations can be you can say gently you know investigators. Then you know in the case of three factor models, we have a three factors through which these simultaneous interaction can be investigators.

So, now in the case of you know two factor models, so we have actually a main effects and then the kind of you know interactive effect. So, for instance, if there are two factors A and B then the main effect will be effect A and effect B and the interactive effect will be a you know AB together so that means it is a intersection between A and B. So, in the kind of you know a problem, so multivariate problem is the two factors. So, we defined actually kind of you know a cross correlation or you know kind of you know a concept that need to be connected while you are investing the kind of you know interrelationship among the variables. So, typically in the two ANOVA case, we have a three different effects. So, the main effects that is with respect to two factors and they are you know interactions.

(Refer Slide Time: 08:51)



The slide is titled "Two-Way ANOVA Model" in red text. It features the mathematical model $x_{ijk} = \mu + \alpha_i + \beta_j + (\alpha\beta)_{ij} + \varepsilon_{ijk}$ in black text. Below the equation is a list of six bullet points explaining the terms. The slide has a yellow background and a blue footer bar containing the IIT Kharagpur and NPTEL logos.

Two-Way ANOVA Model

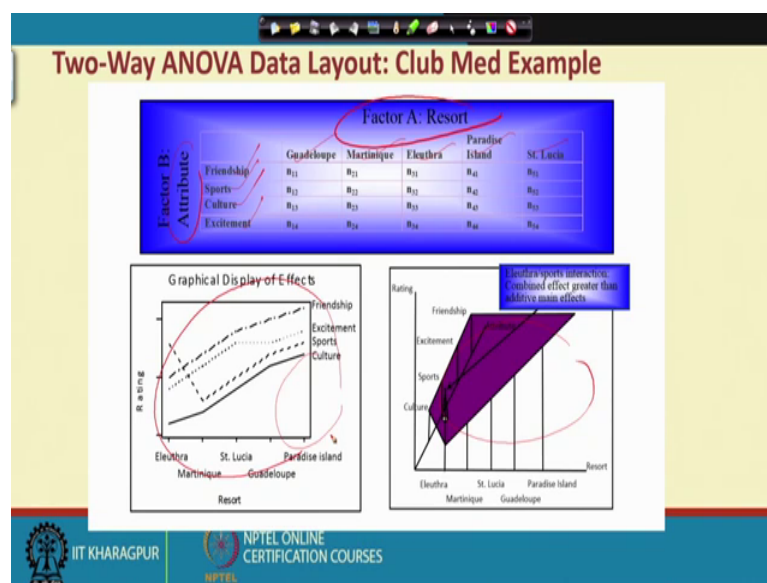
$$x_{ijk} = \mu + \alpha_i + \beta_j + (\alpha\beta)_{ij} + \varepsilon_{ijk}$$

- where μ is the overall mean;
- α_i is the effect of level $i(i=1, \dots, a)$ of factor A;
- β_j is the effect of level $j(j=1, \dots, b)$ of factor B;
- $(\alpha\beta)_{ij}$ is the interaction effect of levels i and j ;
- ε_{ijk} is the error associated with the k th data point from level i of factor A and level j of factor B.
- ε_{ijk} is assumed to be distributed normally with mean zero and variance σ^2 for all i, j , and k .

IIT KHARAGPUR | NPTEL ONLINE CERTIFICATION COURSES

So, let us see how is this particular you know structure. And this is how the two-way ANOVA model and x_{ijk} is the kind of you know a total effect and that is a connected with a μ that is the population mean and subject to α_i plus β_j and plus $\alpha\beta_{ij}$ and ε_{ijk} . So, that means, μ is the overall you know mean and α is the effect of level i for factor A; and β is the effect of level j for factor B; and $\alpha\beta$ is nothing but the interactive effect. And the ε_{ijk} is the kind of you know error associated with this kind of you know link because we are actually structuring with respect to it in two different factors. So, obviously, there is a kind of you know cross correlation and then you need to have a you know interactive effect and along with the you know error terms. And here is there are error terms are you know distributed normally and with the assumption that you know zero mean and again unity variance with respect to i, j and k .

(Refer Slide Time: 10:00)



Typically, two-way ANOVA problem is like this. So, let us take a case you know factor A, so we have an information here the kind of you know let us say this is a kind of you know problem say A resort and then factor B attributes. So, in this particular you know problems we have actually four attributes friendship, sports, culture and excitement, and then we have it you know five different results. And again so we like to check you know how is the kind of you know a rating, a rating of a particular you know hotels and with respect to these you know four attributes.

So, you know it is a kind of you know a hotel business, you will find you know the customers, and they like to enter and they like to leave. And a during these times we usually collect the information and that information will we know will be very helpful to analyze the problem and then you can actually use for some kind of you know improvement or some kind of you know management requirement. So, in this case, so the rating are recorded for different hotels with respect to different attributes. And the usual typical structure of the effects will be like this. So, it is a kind of you know multiple kind of you know plotting and with respect to rating and that too with respect to five different you know resorts. So, the kind of you know interactions, so this is the graphically visualization about this particular problem, and we need actually the kind of you know quantification the total effect or total quantifications how they are you know interacting each other so far as you know the particular problem is concerned. And accordingly we will see how is the particular you know structure then we will discuss.

(Refer Slide Time: 11:48)

Hypothesis Tests a Two-Way ANOVA

- **Factor A main effects test:**
 - ✓ $H_0: \alpha_i = 0$ for all $i=1,2,\dots,a$
 - ✓ H_1 : Not all α_i are 0
- **Factor B main effects test:**
 - ✓ $H_0: \beta_j = 0$ for all $j=1,2,\dots,b$
 - ✓ H_1 : Not all β_j are 0
- **Test for (AB) interactions:**
 - ✓ $H_0: (\alpha\beta)_{ij} = 0$ for all $i=1,2,\dots,a$ and $j=1,2,\dots,b$
 - ✓ H_1 : Not all $(\alpha\beta)_{ij}$ are 0

IIT KHARAGPUR | NPTEL ONLINE CERTIFICATION COURSES

So, what we have already discussed that you know in the two-way process of you know ANOVA, so we have actually three tests. So, the since it is the two factors, so the main two effects are you know factor A and factor B, then finally, the interactive AB, so that is connected with the alpha i beta j and alpha beta ij. So, now accordingly, so the kind of you know a null hypothesis will be alpha equal to 0 for all i again the counterpart alternative hypothesis will be not all alpha ir equal to 0. And similarly we can set for you know bj, so this is for second factors and then finally, there is a interactive effects, so that means, alpha beta j ij equal to 0 means we assume that there is no interactions, but again the alternative will be so at least one interaction is there. So, likewise you know this is how these typical you know structures and that need to be tested with this particular you know concept.

(Refer Slide Time: 12:51)

Hypothesis Tests a Two-Way ANOVA

In a two-way ANOVA:

$$x_{ijk} = \mu + \alpha_i + \beta_j + (\alpha\beta)_{ijk} + \varepsilon_{ijk}$$

- $SST = SSTR + SSE$
- $SST = SSA + SSB + SS(AB) + SSE$

$$SST = SSTR + SSE$$
$$\sum \sum \sum (x - \bar{x})^2 = \sum \sum \sum (\bar{x} - \bar{\bar{x}})^2 + \sum \sum \sum (x - \bar{x})^2$$
$$SSTR = SSA + SSB + SS(AB)$$
$$= \sum \sum \sum (\bar{x}_i - \bar{\bar{x}})^2 + \sum \sum \sum (\bar{x}_j - \bar{\bar{x}})^2 + \sum \sum \sum (\bar{x}_{ij} + \bar{x}_i + \bar{x}_j - \bar{\bar{x}})^2$$

IIT KHARAGPUR | NPTEL ONLINE CERTIFICATION COURSES

So, let us see how is this particular structure. So, it since it is a kind of you know ANOVA that is the analysis of variance. So, we like to check this particular you know effect with respect to variance, because it is a variance statistics. And as a result the typical model is here is this is what the typical models. So, it has a two parts and total SST means total sum of square that is the total variance; and then it is divided into two parts that is the error part and the kind of you know sum square totals the kind of you know factors. So, as a result so the sum square total will be so sum square you know total variance from factor A variance from the factor B that is actually main effect and then this is the interactive effects sum squares through the interactions and then this is the error component.

So, you know the procedure of you know getting this kind of an outcomes like this which we have already discussed in the previous lectures in the case of you know or one-way ANOVA. So, now you know with the help of this data, so you have to find out the variance for a factor A, factor B and for the kind of you know interact interactions and then the error component. And if you add or all together then you will get the sum square totals or else you know yeah you know a sum square totals if you can calculate and SSA and SSB and SSB then by doing some kind of you know simplifications, you can get these sum square errors. But typically so our requirement is a what is SSA, SSB, SS AB and SSE. So, these are the items through which you have to check the effect you know that is the main effect and then the kind of you know interactive effect.

(Refer Slide Time: 14:45)

Two-Way ANOVA

Source of Variation	Sum of Squares	Degrees of Freedom	Mean Square	F Ratio
Factor A	SSA	$a-1$	$MSA = \frac{SSA}{a-1}$	$F = \frac{MSA}{MSE}$
Factor B	SSB	$b-1$	$MSB = \frac{SSB}{b-1}$	$F = \frac{MSB}{MSE}$
Interaction	SS(AB)	$(a-1)(b-1)$	$MS(AB) = \frac{SS(AB)}{(a-1)(b-1)}$	$F = \frac{MS(AB)}{MSE}$
Error	SSE	$ab(n-1)$	$MSE = \frac{SSE}{ab(n-1)}$	
Total	SST	$abn-1$		

A Main Effect Test: $F(a-1, ab(n-1))$ B Main Effect Test: $F(b-1, ab(n-1))$
 (AB) Interaction Effect Test: $F((a-1)(b-1), ab(n-1))$

 IIT KHARAGPUR  NPTEL ONLINE CERTIFICATION COURSES

So, since we have already gone through this ANOVA that is one way ANOVA. In this case it is the similar kind of you know structure and in this case so this is for factor A, this is factor B, and this is interactive effect and then the errors. Since we are dealing with you know two factors and the interactions by default some error will be there in the systems. So, as a result so for factor A so sum will be sum squares for factor A SSA. And for factor B, so it is SSB; and then interactivity is sum squares of you know interactive AB. And for factor a, so the degree of freedom is a minus 1 because we are allowing the factor samples for you know up to a, so that is the size of the samples and by default the degree of freedom will be a minus 1. And similarly for factor B, it will be b minus 1, and the interactive effect will be a minus 1 into b minus 1s.

And the error component accordingly will be a b into n minus 1. And the total for the total that is sum square total, so this is the combined a samples that is a b ns and that to the degree of freedom will be a b n minus 1s. And we typically get to know so what is the mean square of you know factor A, mean square of factor B and mean square of you know interactive. So, this is nothing but actually sum square of factor A you know divided by the degree of freedom and it is followed by factor B and also it is followed by interactive effect.

So, accordingly we have actually three different f statistics for two-way ANOVA, one is for factor A, one is for factor B, and one is for you know interactive effects. So, we like

to know with respect to the total variations which factor is more effective and then whether the interactive effect is statistically significant or not. And in order to justify this, so we have to take and you know live examples.

(Refer Slide Time: 16:45)

Example: Two-Way ANOVA (Location and Artist)

Source of Variation	Sum of Squares	Degrees of Freedom	Mean Square	F Ratio
Location	1824	2	912	8.93
Artist	2230	2	1115	10.93
Interaction	804	4	201	1.93
Error	8262	81	102	
Total	13120	89		

$\alpha = 0.01, F(2,81) = 4.88 \Rightarrow$ Both main effect null hypotheses are rejected.
 $\alpha = 0.05, F(2,81) = 2.48 \Rightarrow$ Interaction effect null hypotheses are not rejected.

So, let us take this examples and then we will see how is the kind of you know structure. So, this is the kind of you know with a particular samples and that with the two kind of you know structures location and artist, and which is you know with kind of you know data analysis. So, we have a SSA here and that to SSA here and SSB here, and then this interactive effect, and then the SSE that is sum square errors.

And if you add all these three, so you will get actually the total effect that is called as you know sum square totals. And the corresponding degree of freedom is 2, 3 and then 4, and that is for error 81 and the total samples for this particular you know data is nothing but 89 and that means, 90 and then this degree of freedom will be nine 89 for the SST. And accordingly, so this will be mean square for factor A, this is mean square for factor B, and this is mean square for interactive.

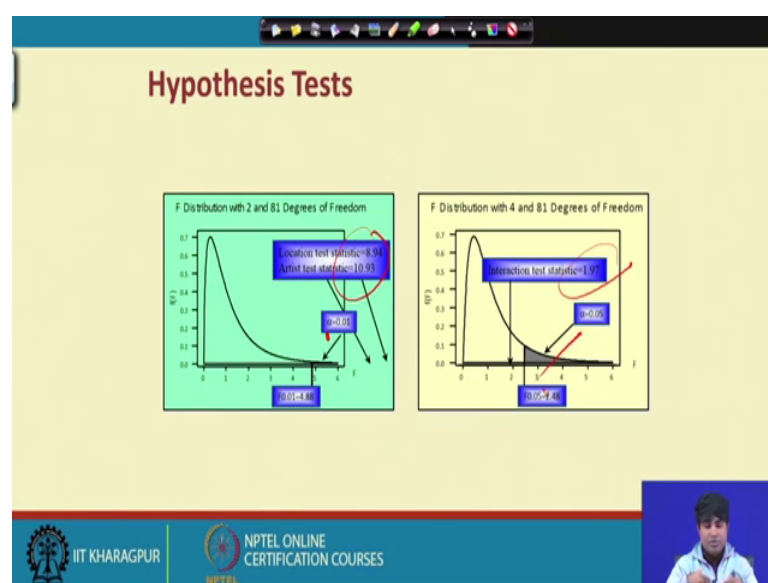
And then corresponding to this, this is the D calculated for factor A, F calculated for factor B, and F calculated for interactions. And then so corresponding to this you know degree of freedom and then you will go to the F critical. And check you know whether you know this F calculated will overtake the F critical depending upon the kind of you know alpha which we have actually fixed here at a 1 percent and 5 percent so that means,

typically actually even if you do not fix actually alpha. So, first you have the calculated value depending upon the sample information and the problem.

And then you have actually a critical value depending upon the kind of you know structure which you can actually fix or you know allow. For instance, you can fix alpha equal to 1 percent, you can fix alpha equal to 5 percent, you can fix alpha equal to 10 percent, so accordingly so you can get you know three critical values. So, now this calculated value can be compared with a 1 percent, 5 percent and 10 percent; typically if it is actually starts really significant at 1 percent that is the highest weightage you can give into these kind of you know difference followed by 5 percent and you know 10 percent.

So, in this case let us see what is happening. So, the critical value will be 4.88 and then it is actually 2.48. So, in the both the cases actually and the calculated value is over taking the tabulated values. So, as a result you are rejecting the true null hypothesis, but in the case of you know interactive effect, it is actually coming you know less that is less than two critical value. So, we cannot actually reject in this case so that means, in this particular process, so the main effects are you know statistically significant while the interactive effect is not statistically significant, so that is how the kind of you know structure.

(Refer Slide Time: 19:35)



(Refer Slide Time: 20:28)

And this is actually the kind of you know excel template while you analyzing two-way ANOVA. And I will discuss these problems through you know excel sheet.

(Refer Slide Time: 20:39)

Extension of ANOVA to Three Factors

Source of Variation	Sum of Squares	Degrees of Freedom	Mean Square	F Ratio
Factor A	SSA	$a-1$	$MSA = \frac{SSA}{a-1}$	$F = \frac{MSA}{MSE}$
Factor B	SSB	$b-1$	$MSB = \frac{SSB}{b-1}$	$F = \frac{MSB}{MSE}$
Factor C	SSC	$c-1$	$MSC = \frac{SSC}{c-1}$	$F = \frac{MSC}{MSE}$
Interaction (AB)	SS(AB)	$(a-1)(b-1)$	$MS(AB) = \frac{SS(AB)}{(a-1)(b-1)}$	$F = \frac{MS(AB)}{MSE}$
Interaction (AC)	SS(AC)	$(a-1)(c-1)$	$MS(AC) = \frac{SS(AC)}{(a-1)(c-1)}$	$F = \frac{MS(AC)}{MSE}$
Interaction (BC)	SS(BC)	$(b-1)(c-1)$	$MS(BC) = \frac{SS(BC)}{(b-1)(c-1)}$	$F = \frac{MS(BC)}{MSE}$
Interaction (ABC)	SS(ABC)	$(a-1)(b-1)(c-1)$	$MS(ABC) = \frac{SS(ABC)}{(a-1)(b-1)(c-1)}$	$F = \frac{MS(ABC)}{MSE}$
Error	SSE	$abc(n-1)$	$MSE = \frac{SSE}{abc(n-1)}$	
Total	SST	$abcn-1$		

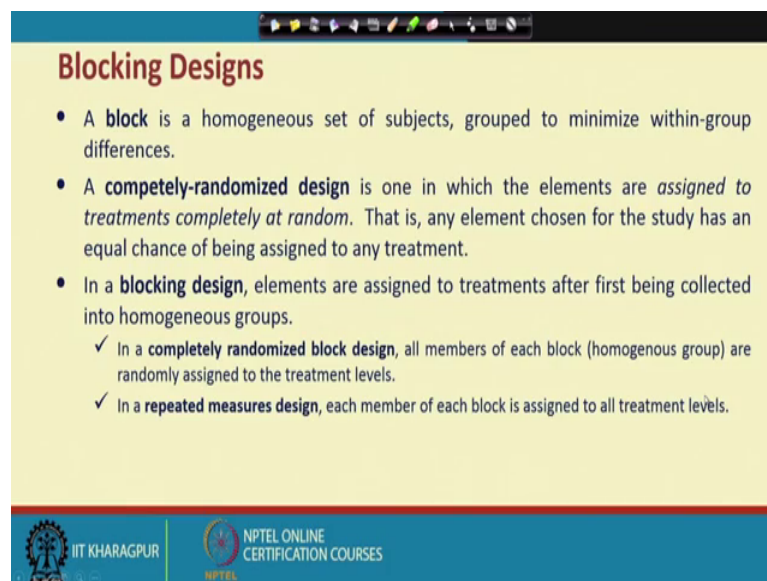
Let me first complete this one then I will go to this particular excel view. And in this case you know the a and this is actually three factor model and this is extension of you know again ANOVA. And in this case we have actually three factors. And the three factors is actually a, b and c. And in the case of you know two factors we have a main effects a and b and the interactive effect will be ab, but when there are actually three factors or four factors and five factors like this, so the interactive effect will be you know continue in a kind of you know complex process. And a typically for these three factors models. So, we have a three main effects that is SSA, SSB and SSC that is sum squared of you know factor A sum square of factor B and sum square of you know factor C. Then by default so there are you know three interactive effect, so between AB, AC and BC and then finally, the combined interactive effect which is nothing but actually ABC.

So, as a result sum square AB, sum square AC, sum square BC and sum square ABC and finally, when you are dealing with more factors and by default there will be error component and that will be represented by sum square errors. In fact, some square error was also there in the case of you know two-way ANOVA. And in that case, we have actually to two main effect and one interactive effect. Here is we have actually three main effects ABC and then three interactive effect AB, AC, BC and then combined effect that is the interactive between A Band C. So, the corresponding the degree of freedom will be a minus 1 b minus 1 c minus 1. And like the previous one, it is a minus 1 into b

minus 1 a minus into c minus 1 b 1 minus 1 to c minus 1 and then a minus 1 b minus 1 c minus 1 that is for you know joint effect of you know abc.

And finally, this is for you know error sum squares, and this is for you know sum square totals. And accordingly you can find out mean square you know for factor A, factor B and factor C and this is what you know the interactive effect between AB interactive between AC and interactivity to BC and this is the interactive effect combined ABC. And accordingly we have actually in that case of you know three factor models, so these are the three main factors and then these are the three actually you know kind of you know interactive factors. So, all together, so we have actually seven different test statistic which you would like to report in the case of you know three factor models. So, while in the case of you know two factor models we have a two main effects and then one is the interactive effect. So, let us see here is, so connecting to this, so this is actually three factor model situation.

(Refer Slide Time: 23:25)



Blocking Designs

- A **block** is a homogeneous set of subjects, grouped to minimize within-group differences.
- A **completely-randomized design** is one in which the elements are *assigned to treatments completely at random*. That is, any element chosen for the study has an equal chance of being assigned to any treatment.
- In a **blocking design**, elements are assigned to treatments after first being collected into homogeneous groups.
 - ✓ In a **completely randomized block design**, all members of each block (homogenous group) are randomly assigned to the treatment levels.
 - ✓ In a **repeated measures design**, each member of each block is assigned to all treatment levels.

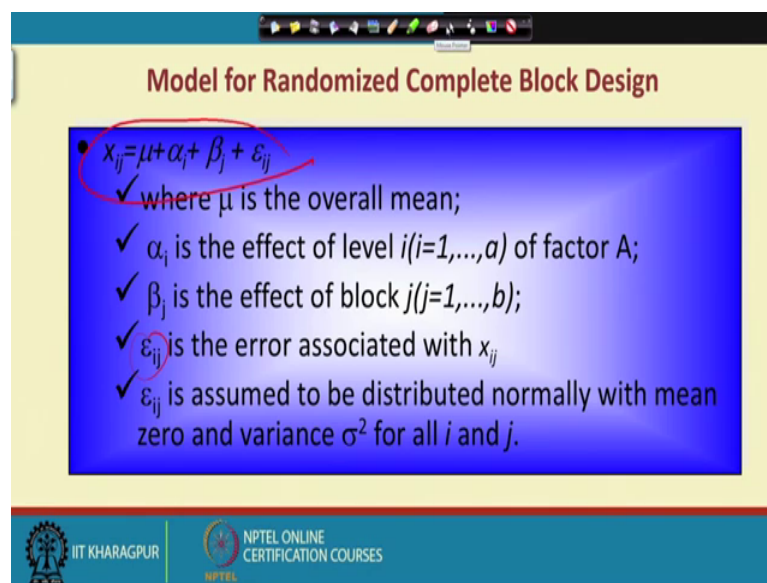
IIT KHARAGPUR | NPTEL ONLINE CERTIFICATION COURSES

And accordingly we can connect with you know four factor model five factor model. So, when we add one after another layers then the degree of complexity and the kind of you know effect will be more and more. So, which we have already seen from the moment from one to two, and two to three like this. So, when will you add one after another factor, so then the combined effect or the kind of you know interactive effect will be complex and then and then it will be very in fact, it will be very interesting and also it

will be through complex process we have to investigate and then we can analyze the process. In fact, for you know more factors, and you know more kind of you know typical samples, so manually it is not possible through software we have actually you can deal with the kind of you know an analysis then finally, we will check the kind of you know variance.

And in the case of you know block designs, so let us see what is actually block. A block is a homogeneous a set of subjects a group to minimize within group differences and. In fact, we have what we have already immense on the completely randomized design in one which the elements are you know assigned to treatments completely at the random. And we in fact, we have already mentioned all these things completely random designs and block designs.

(Refer Slide Time: 24:44)



Model for Randomized Complete Block Design

$$x_{ij} = \mu + \alpha_i + \beta_j + \varepsilon_{ij}$$

- ✓ where μ is the overall mean;
- ✓ α_i is the effect of level $i (i=1, \dots, a)$ of factor A;
- ✓ β_j is the effect of block $j (j=1, \dots, b)$;
- ✓ ε_{ij} is the error associated with x_{ij}
- ✓ ε_{ij} is assumed to be distributed normally with mean zero and variance σ^2 for all i and j .

IIT KHARAGPUR | NPTEL ONLINE CERTIFICATION COURSES

So, let me take the kind of you know a model part of this you know block designs. And like the previous you know factor models, you know a block design model usually represented like this. So, the main effect which is actually $\mu + \alpha_i + \beta_j$ you know summation you know α_i and then and so it is followed by the α component and that is the effect of level 1. And β_j the it is effect of you know block j , and μ is the overall mean. And like the kind of you know factor models, so a this is the error term. And again we are assuming that you know the error term will be normally distributed

with a mean zero and unit variance right and that is the homogeneous variance. So, let us see how is the particular you know cases.

(Refer Slide Time: 25:35)

ANOVA Table for Blocking Designs: Example

Source of Variation	Sum of Squares	Degress of Freedom	Mean Square	F Ratio
Blocks	SSBL	$n - 1$	$(MSBL = SSBL/(n-1))$	$F = MSBL/MSR$
Treatments	SSTR	$r - 1$	$(MSTR = SSTR/(r-1))$	$F = MSTR/MSR$
Error	SSE	$(n-1)(r-1)$	$(MSE = SSE/((n-1)(r-1)))$	
Total	SST	$nr - 1$		

Source of Variation	Sum of Squares	df	Mean Square	F Ratio
Blocks	2750	39	70.51	0.69
Treatments	2640	2	1320	12.93
Error	7960	78	102.05	
Total	13350	119		

$\alpha = 0.01, F(2, 78) = 4.88$

IIT KHARAGPUR | NPTEL ONLINE CERTIFICATION COURSES

This is how the typical problems and it is it like you know or kind of typically ANOVA structures, but here the items are you know just like you know different, because it is a kind of you know block design problems. So, as a result so compared to the previous one, so the first effect will be blocks then treatments, so this will be block, statements, and the error component, and then finally total So, again similarly sum square you know for blocks and sum square you know for errors it, it means and then sum square you know for errors. And if you combine all these then you will get you know total sum of squares and corresponding to this is actually block wise is nothing but actually the kind of you know column variations and this is the row variations.

So, typically so the degree of freedom will be n minus 1 and r minus 1s and then finally, so for you know errors it will be n minus 1 into r minus 1. And like the ANOVA so we will find out the mean square for you know blocks mean square for treatments and then mean square for you know errors. And in this case we are interested about the a block effect and the treatment effect and accordingly we have a two f statistics. And with a particular you know sample problems, so we have calculated here actually sum of square for the block, sum of square for the treatments and sum square of the error. The first hand check is actually these three should actually you know if close to us total sum square and

the corresponding the degree of freedom is reported here. And the simplification you define mean square for a blocks and mean square for treatment.

And then this is what the calculated value and again fixing alpha and then reporting the critical value for f and you can check the value. So, since actually this value is lesser to this particular value. So, we are not in a position to a reject the null hypothesis, but in this case we are actually rejecting the term, because it is over taking the kind of you know critical value.

(Refer Slide Time: 27:41)

Template for the Randomized Complete Block Design

	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z	AA	
1	Randomized Complete Block Design																											
2																												
3	Mean	10.714	13	15.429																								
4		A	B	C		Mean																						
5	1	12	16	15		14.333																						
6	2	11	10	14		11.667																						
7	3	10	11	17		12.667																						
8	4	12	15	19		15.333																						
9	5	11	14	17		14																						
10	6	10	12	13		11.667																						
11	7	9	13	13		11.667																						
12																												
13																												
14																												
15																												
16																												

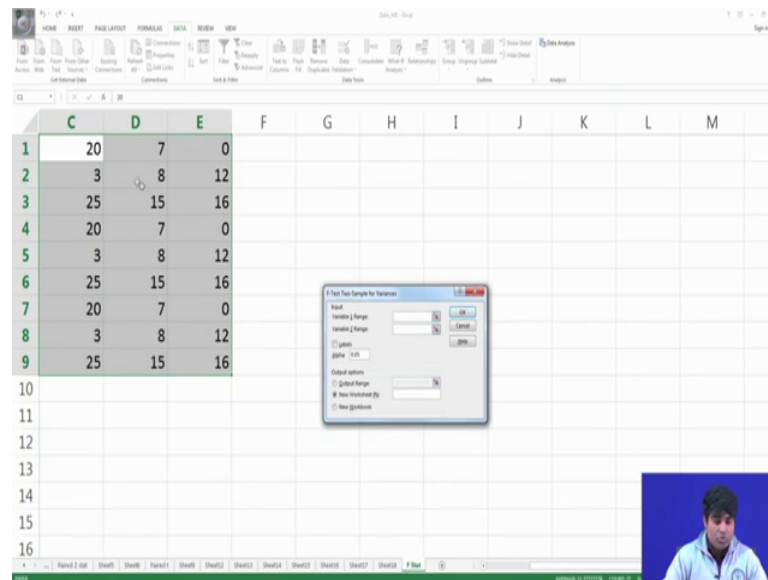
Source	SS	df	MS	F	F _{critical}	p-value
Treatment	40.952	6	6.8254	3.1273	2.996	0.0439
Block	77.81	2	38.905	17.825	3.895	0.0003
Error	26.19	12	2.1825			
Total	144.95	20				

α
5%
Reject
Reject

NPTEL ONLINE
CERTIFICATION COURSES

So, this is actually again excel template how you have to do the kind of you know an analysis. And in fact so before we go to this particular process, let me highlight the kind of you know examples how you can go to the analysis and solve the problems.

(Refer Slide Time: 27:56)



So, this is actually spreadsheet and we have actually a couple of samples here and then we like to test through oh you know simple you know start with you know simple f then we will go for you know one-way ANOVA, two-way ANOVA, and in three-way ANOVA like this. And then since actually this is so go to the data analysis package and in fact, you will find here is you will find for he has these are the tests. So, we need to test first you know f statistics. So, this is actually the f statistic for two sample variance let us say. So, this is with respect to the previous class discussion.

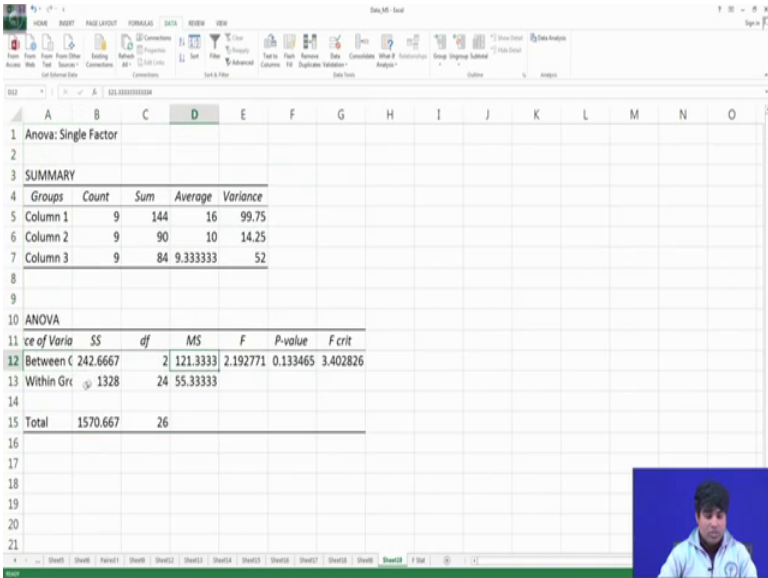
So, when we will go for you know means it is actually multi sample case. So, the multi sample case means at least you must have actually two different you know samples and then we must have some sample points to test the you know kind of you know variance. So, now accordingly you have to specify here the first sample case, and then you have to specify the second sample case, and like you know t and z so you have to just you know put. So, you will finally, this is how the f distribution this is the f_k you know distribution values and this is f critical.

And then finally, this is f critical and this is what the f calculated on the basis of the samples and these are all mean reporting of these two variables and variance of these two variables, these are all counts and that is sample observation. And f reporter is this much, and the degree of freedom is this much that is actually an n minus 1. And then finally, with comparison since this f calculator is you are taking they are critical. So, it is

statistically significant. So, this is what simple f statistic through f statistic you have to analyze this problem that means there is a significance difference between these two samples and then accordingly the management decision will be a taken into consideration.

Now, again you go to the data analysis same structures. We will analyze this problem through ANOVA and you see here is another single factor models, which you have already highlighted then you know you just actually highlight again, so the entire sample is just a highlight and you know then you give the kind of ok.

(Refer Slide Time: 30:24)



The screenshot shows the 'Data Analysis' tool in Excel with 'ANOVA: Single Factor' selected. The SUMMARY table is as follows:

Groups	Count	Sum	Average	Variance
Column 1	9	144	16	99.75
Column 2	9	90	10	14.25
Column 3	9	84	9.333333	52

The ANOVA table is as follows:

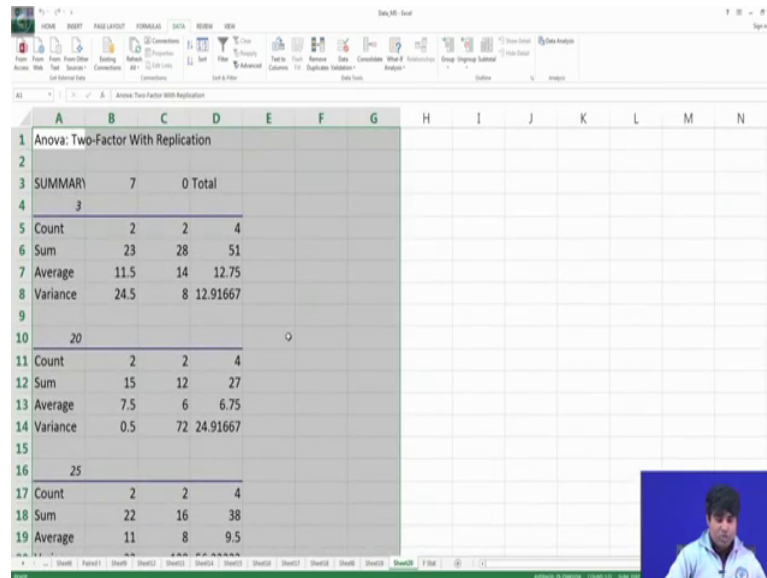
Source of Variation	SS	df	MS	F	P-value	F crit
Between Groups	242.6667	2	121.3333	2.192771	0.139465	3.402826
Within Groups	1328	24	55.33333			
Total	1570.667	26				

So, this is how the usual representation of this is what actually the output. And we are actually having between the difference and within the difference and the ratio between the given f value and this is what the probability which you have P and that to this is critical and this is what the calculated. And since the f critical is over taking the f calculated as a result, so this is not actually statistically significant even at actually 10 percent. And these are actually the kind of you know summary statistics for this particular you know kind of you know problem.

So, this is the actually the kind of you know output side for single factor model. And again you go to the spreadsheet and then you allow the data analysis again. So, again you allow one ANOVA two factor models. So, it is extension again. So, in this case again you

specify the range and allow that you know there are you know two since two factors. So, you allow the two factor sampling and then again put ok.

(Refer Slide Time: 31:28)

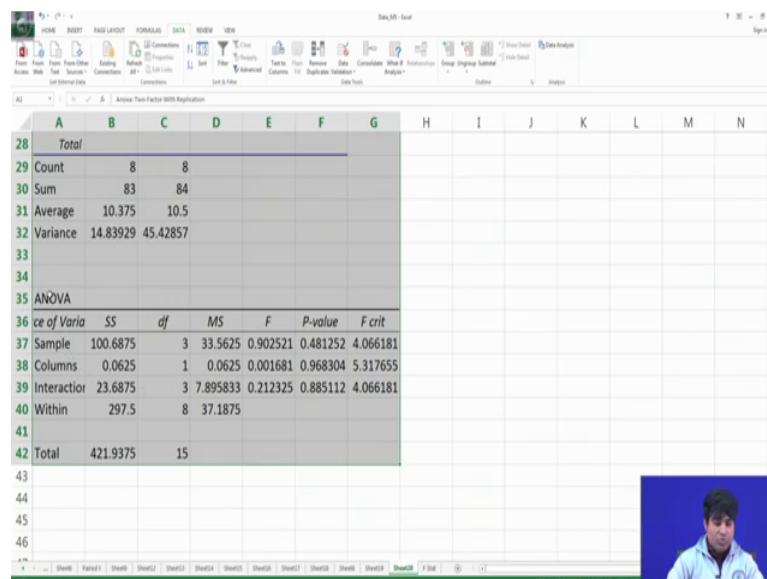


The screenshot shows an Excel spreadsheet with the title 'Anova: Two-Factor With Replication'. The data is organized in columns A through G. The summary table is as follows:

	A	B	C	D	E	F	G
1	Anova: Two-Factor With Replication						
2							
3	SUMMARY	7	0	Total			
4	3						
5	Count	2	2	4			
6	Sum	23	28	51			
7	Average	11.5	14	12.75			
8	Variance	24.5	8	12.91667			
9							
10	20						
11	Count	2	2	4			
12	Sum	15	12	27			
13	Average	7.5	6	6.75			
14	Variance	0.5	72	24.91667			
15							
16	25						
17	Count	2	2	4			
18	Sum	22	16	38			
19	Average	11	8	9.5			

So, this is how the analysis of actually a two factor kind of you know situations, you will find plenty of you know ANOVA so that means, you is compared to the previous one. So, in the previous one, you really find here plenty of results.

(Refer Slide Time: 31:45)



The screenshot shows an Excel spreadsheet with the title 'Anova: Two-Factor With Replication'. The data is organized in columns A through G. The ANOVA table is as follows:

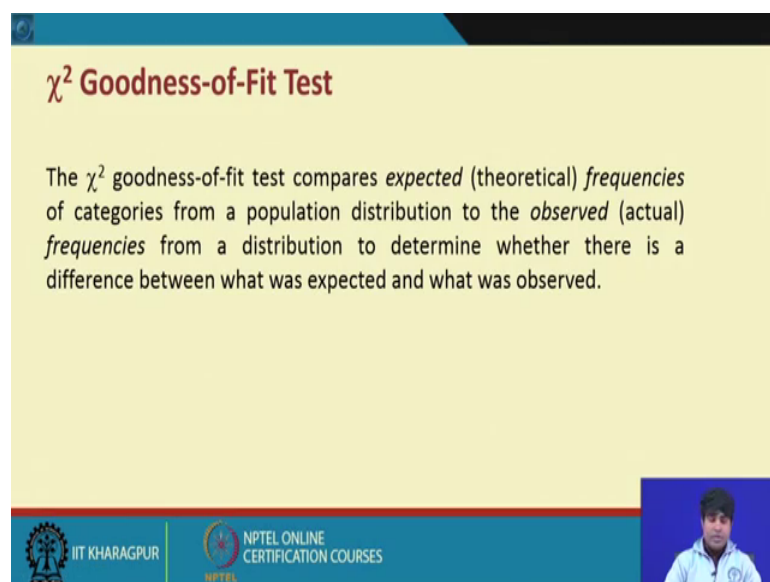
	A	B	C	D	E	F	G
28	Total						
29	Count	8	8				
30	Sum	83	84				
31	Average	10.375	10.5				
32	Variance	14.83929	45.42857				
33							
34							
35	ANOVA						
36	Source of Variation	SS	df	MS	F	P-value	F crit
37	Sample	100.6875	3	33.5625	0.902521	0.481252	4.066181
38	Columns	0.0625	1	0.0625	0.001681	0.968304	5.317655
39	Interaction	23.6875	3	7.895833	0.212325	0.885112	4.066181
40	Within	297.5	8	37.1875			
41							
42	Total	421.9375	15				

And then finally you come to this you know a final outcome. And here this is the sample the factor A and this is the factor B and this is what the interactive way; that means, this

is the SSA, this is SSB, and this is SSAB. And then it is a kind of you know a error component, and then this is means sum squares of A, and sum square of B, and this is the interactive mean sum square and accordingly. So, these are all critical values f critical values corresponding to a particular alpha levels. And this is actually f calculated. And compared to these critical, so we are not in a position to reject for the sample one that is the factor one case and in the case of factor A, B also it is also not in a position to reject. And again in the interactive, you are also not in a position to reject.

So, that means, in this particular in a sample and problem, so we do not find any significant difference. In fact, the problem is actually this may be with respect to in a small sample case, because this is actually multiple sample case means you must have actually more number of data points to validate this particular you know a problems or the kind of you know issues. And then you can get some kind of you know better inference. So, likewise you can extend this particular model for three factor models then you can go for you know completely you know four factor models and then MANOVA like the kind of you know situations right. So, typically what I like to say that you know f is a kind of you know interesting factor through which you have to discuss the particular you know problem.

(Refer Slide Time: 33:27)



χ^2 Goodness-of-Fit Test

The χ^2 goodness-of-fit test compares *expected* (theoretical) *frequencies* of categories from a population distribution to the *observed* (actual) *frequencies* from a distribution to determine whether there is a difference between what was expected and what was observed.

The slide features a blue header with a logo, a yellow main content area, and a blue footer. The footer contains the IIT Kharagpur logo and the text 'NPTEL ONLINE CERTIFICATION COURSES'. A small video inset of a speaker is visible in the bottom right corner.

And the last but not the least in this particular you know multi sample case is the testing through chi square observations. And in fact, chi square we have already discussed with

one sample case. And here is we are dealing with the chi square with one sample case and taking you to the clue actually for multiple sample case and typically for categorical variables, because variable means information may be numeric, but the variables may be kind of you know qualitative in nature. So, when a problem is having actually qualitative variables and then how you have to dealing the situations and then we will get the inference. So, chi square is the best statistic through which you can actually analyze the particular situations.

(Refer Slide Time: 34:10)

χ^2 Goodness-of-Fit Test

$$\chi^2 = \sum \frac{(f_o - f_e)^2}{f_e}$$
$$df = k - 1 - c$$

where: f_o = frequency of observed values
 f_e = frequency of expected values
 k = number of categories
 c = number of parameters estimated from the sample data

IIT KHARAGPUR | NPTEL ONLINE CERTIFICATION COURSES

So, chi square is a goodness of fit measures and this is how the test structure through actually you have to test the procedure. It is more or less same like you know z test, t test and f test. In fact, a chi square one sample case we have already tested. So, it is actually sum total of you know the difference between observed frequency and expected frequency. And for all by the you know you will get the kind of you know calculated and then followed by the kind of you know critical and then you check whether there is a significant difference between observed frequency and the expected frequency.

(Refer Slide Time: 34:43)

Milk Sales Data for Demonstration Problem	
Month	Gallons
January	1,610
February	1,585
March	1,649
April	1,590
May	1,540
June	1,397
July	1,410
August	1,350
September	1,495
October	1,564
November	1,602
December	1,655
	18,447

For examples, let us take example like this is you know milk sales data over the you know last twelve months of a year and we like to check whether there is a significant difference about these samples with respect to sample means. So, the total of these sales is nothing but 18,447, and as a result these are all individual sales.

(Refer Slide Time: 35:07)

Hypotheses and Decision Rules for Demonstration Problem

H_0 : The monthly milk figures for milk sales are uniformly distributed

H_a : The monthly milk figures for milk sales are not uniformly distributed

$\alpha = .01$
 $df = k - 1 - c$
 $= 12 - 1 - 0$
 $= 11$
 $\chi^2_{.01, 11} = 24.725$

If $\chi^2_{cal} > 24.725$, reject H_0 .
If $\chi^2_{cal} \leq 24.725$, do not reject H_0 .

And while doing the analysis, so this is how the null hypothesis really specify and then and the type of you know the structure which you have fixed for testing. So, chi square calculated is this much, and then we will see how is the kind of you know difference.

(Refer Slide Time: 35:24)

Calculations for Demonstration :

Month	f_o	f_e	$(f_o - f_e)^2/f_e$
January	1,610	1,537.25	3.44
February	1,585	1,537.25	1.48
March	1,649	1,537.25	8.12
April	1,590	1,537.25	1.81
May	1,540	1,537.25	0.00
June	1,397	1,537.25	12.80
July	1,410	1,537.25	10.53
August	1,350	1,537.25	22.81
September	1,495	1,537.25	1.16
October	1,564	1,537.25	0.47
November	1,602	1,537.25	2.73
December	1,655	1,537.25	9.02
	18,447	18,447.00	74.38

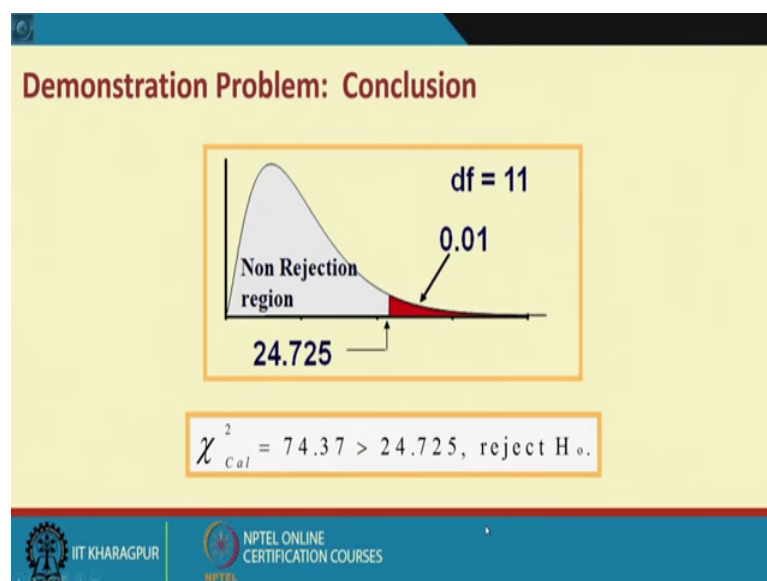
$$f_e = \frac{18447}{12}$$

$$= 1537.25$$

$$\chi^2_{cal} = 74.37$$

So, the procedure of you know calculation is like this. So, now, having the sample information this ones and then what will you do. So, you have to calculate the kind of you know mean and then you have to find out the difference between the individual items within a sample mean. And then with the help of you know this particular you know different squaring of these difference and followed by expected frequency, and once you have it this sum that is what is actually total chi square variance. And that need to be tested through critical and that the critical is here.

(Refer Slide Time: 36:03)



And f critical is here coming actually 24.725 corresponding to alpha at actually 1 percent. And depending upon the degree of freedom that is 11 that means, the sample size is a 12. So, the calculated value is coming 74.37. So, as a results. So, we are in a position to research so; that means, there is a significant difference between the individual items corresponding to sample you know I mean or mean of this particular you know series.

(Refer Slide Time: 36:29)

Using a χ^2 Goodness-of-Fit Test to Test a Population Proportion

$\alpha = .05$ $df = k - 1 - c$ $= 2 - 1 - 0$ $= 1$ $\chi^2_{.05,1} = 3.841$	$H_0: P = .08$ $H_a: P \neq .08$
	If $\chi^2_{\text{Cal}} > 3.841$, reject H_0 . If $\chi^2_{\text{Cal}} \leq 3.841$, do not reject H_0 .

 IIT KHARAGPUR
  NPTEL ONLINE CERTIFICATION COURSES

So, likewise actually you can go for you know testing for population proportions, it is more on the same.

(Refer Slide Time: 36:36)

Using a χ^2 Goodness-of-Fit Test to Test a Population Proportion: Calculations

	f_o	f_e
Defects	33	16
Nondefects	167	184
n	200	200

Defects $f_e = n \cdot P$
 $f_e = (200)(.08)$
 $= 16$

Nondefects $f_e = n \cdot (1 - P)$
 $f_e = (200)(.92)$
 $= 184$

$$\chi^2 = \sum \frac{(f_o - f_e)^2}{f_e}$$
$$= \frac{(33-16)^2}{16} + \frac{(167-184)^2}{184}$$
$$= 18.0625 + 1.5707$$
$$= 19.6332$$

IIT KHARAGPUR | NPTEL ONLINE CERTIFICATION COURSES

And here the problem of proper population proportion will be like this. So, this is the observed frequency then you have a expected frequency. And it is more or less same against, so you have to find out the difference. And the squaring of the difference need to be tested again compared with the you know kind of the critical value. And on the basis of this data, so we have actually calculated chi square where region 19.64 and then again that need to be tested and again which is actually in a position to reject. So, this is typically actually one sample case you know these two problems are with respect to one sample case.

(Refer Slide Time: 37:12)

χ^2 Test of Independence

Used to analyze the frequencies of two variables with multiple categories to determine whether the two variables are independent.

Qualitative Variables
Nominal Data

IIT KHARAGPUR | NPTEL ONLINE CERTIFICATION COURSES

But I like to highlight you know the concept with you know qualitative variables and the nominal data structure.

(Refer Slide Time: 37:20)

χ^2 Test of Independence: Investment Example

- In which *region of the country* do you reside?
A. Northeast B. Midwest C. South D. West
- Which *type of financial investment* are you most likely to make today?
E. Stocks F. Bonds G. Treasury bills

Contingency Table

		Type of financial Investment			
		E	F	G	
Geographic Region	A			O_{11}	n_A
	B				n_B
	C				n_C
	D				n_D
		n_E	n_F	n_G	N

And let me stake you the problems. So, usually it is the sexually chi square test for you know test of independence and the examples will be let us say investment and then you know find kind of you know how it is varying with respect to different you know locations. It is a kind of you know multiple sample case like you know two factor model which you have actually discussed. So, here all the similar kind of you know structures. And the same problems can be analyzed with you know chi square test statistics. Let us have a look here is how is the kind of you know structure (Refer Time: 37:52) are there.

(Refer Slide Time: 37:52)

χ^2 Test of Independence: Investment Example

If A and F are independent,
 $P(A \cap F) = P(A) \cdot P(F)$

$$P(A) = \frac{n_A}{N} \quad P(F) = \frac{n_F}{N}$$

$$P(A \cap F) = \frac{n_A}{N} \cdot \frac{n_F}{N}$$

$$e_{AF} = N \cdot P(A \cap F)$$

$$= N \left(\frac{n_A}{N} \cdot \frac{n_F}{N} \right)$$

$$= \frac{n_A \cdot n_F}{N}$$

		Type of Financial Investment			
		E	F	G	
Geographic Region	A		e_{12}		n_A
	B				n_B
	C				n_C
	D				n_D
		n_E	n_F	n_G	N

NPTEL ONLINE
CERTIFICATION COURSES

And if we add up all these column wise sums, and row by sums then we will find actually combined you know samples and then it will be tested. So, the typical idea actually like this all right. So, this is actually the kind of you know structure. And then what do you like to do here, you have to find out the again the difference between the actual of positions and the expected positions. And the expected position that means, actually in this case it is actually three into four matrix and that means, typically each entry is the kind of you know actual kind of you know observation or actual kind of you know effects. And then we will find out the expected effect and procedure to calculate the expected effect is like this.

(Refer Slide Time: 38:40)

χ^2 Test of Independence: Formulas

Expected
Frequencies

$$e_{ij} = \frac{n_{i.}n_{.j}}{N}$$

where: i = the row
 j = the column
 $n_{i.}$ = the total of row i
 $n_{.j}$ = the total of column j
 N = the total of all frequencies

Calculated χ^2
(Observed χ^2)

$$\chi^2 = \sum \sum \frac{(f_o - f_e)^2}{f_e}$$

where: $df = (r - 1)(c - 1)$
 r = the number of rows
 c = the number of columns

So, it is the actually row sum into column sums divided by the combined sum and that is the expected point. So, that is why $n_{i.}$ into $n_{.j}$ means it is i th columns and that is the i th total and the $n_{.j}$ means it is a j th column and that to j th totals. So, the multiplication of that two components divided by combined means will give you the position of that you know and that to expected frequency of that particular samples. Likewise for every samples, you can find out the expected frequency and then finally, you have to find out the chi square value which is the difference between the observed frequency and expected frequency and that needs to be again tested.

(Refer Slide Time: 39:28)

χ^2 Test of Independence: Gasoline Preference Versus Income Category

$\alpha = .01$
 $df = (r - 1)(c - 1)$
 $= (4 - 1)(3 - 1)$
 $= 6$
 $\chi^2_{.01, 6} = 16.812$

$r = 4$ $c = 3$

Income	Type of Gasoline		
	Regular	Premium	Extra Premium
Less than \$30,000			
\$30,000 to \$49,999			
\$50,000 to \$99,000			
At least \$100,000			

If $\chi^2_{cal} > 16.812$, reject H_0 .
 If $\chi^2_{cal} \leq 16.812$, do not reject H_0 .

H_0 : Type of gasoline is independent of income
 H_a : Type of gasoline is not independent of income

And this is the sample problems and that is with respect to income category and the type of you know gasoline. And there are three different types regular, premium and extra premium and then subject to income variations. There are four income variations. And this is how the develop the types of gasoline is independent of you know income that is the kind of you know problems which you need to investigate with these examples. And the corresponding alternative hypothesis that you know the gasoline is not in dependent of you know income.

(Refer Slide Time: 40:00)

Gasoline Preference Versus Income Category: Expected Frequencies

$$e_{ij} = \frac{(n_{i.})(n_{.j})}{N}$$

$$e_{11} = \frac{(107)(238)}{385} = 66.15$$

$$e_{12} = \frac{(107)(88)}{385} = 24.46$$

$$e_{13} = \frac{(107)(59)}{385} = 16.40$$

Income	Type of Gasoline			
	Regular	Premium	Extra Premium	
Less than \$30,000	(66.15) 85	(24.46) 16	(16.40) 6	107
\$30,000 to \$49,999	(87.78) 102	(32.46) 27	(21.76) 13	142
\$50,000 to \$99,000	(45.13) 36	(16.69) 22	(11.19) 15	73
At least \$100,000	(38.95) 15	(14.40) 23	(9.65) 25	63
	238	88	59	385

So, accordingly the sample points are like this. So, corresponding to this income and you know types of gasoline, these are you know observed frequency and accordingly with this you know some quarter or row sums and column sums for a particular positions, you have to find out the expected frequency. And then accordingly so this is how the procedure through which actually the all the bracketed figures are actually the expected frequency. And then so you will find here is for observed frequency there is a expected frequency and then you will find the difference and that is the square of the difference will you know need to be tested. So, whether they are statistically different and corresponding to expected frequency.

(Refer Slide Time: 40:45)

Gasoline Preference Versus Income Category: χ^2 Calculation

$$\chi^2 = \sum \sum \frac{(f_o - f_e)^2}{f_e}$$

$$= \frac{(88-66.15)^2}{66.15} + \frac{(16-24.46)^2}{24.46} + \frac{(6-16.40)^2}{16.40} + \frac{(102-87.78)^2}{87.78} + \frac{(27-32.46)^2}{32.46} + \frac{(13-21.76)^2}{21.76} + \frac{(36-45.13)^2}{45.13} + \frac{(22-16.69)^2}{16.69} + \frac{(15-11.19)^2}{11.19} + \frac{(15-38.95)^2}{38.95} + \frac{(23-14.40)^2}{14.40} + \frac{(25-9.65)^2}{9.65}$$

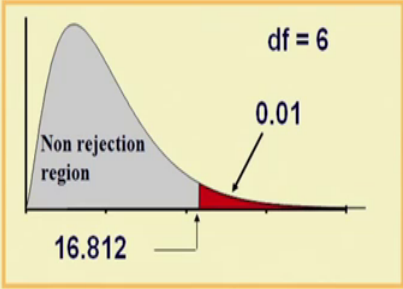
$$= 70.78$$

IIT KHARAGPUR | NPTEL ONLINE CERTIFICATION COURSES

And the since it is actually variant statistics. So, the procedures same procedure we are following in this particular case and accordingly. So, after simplifications you are getting the total chi square is equal to 70.78.

(Refer Slide Time: 41:03)

Gasoline Preference Versus Income Category: Conclusion



$df = 6$
 0.01
Non rejection region
 16.812

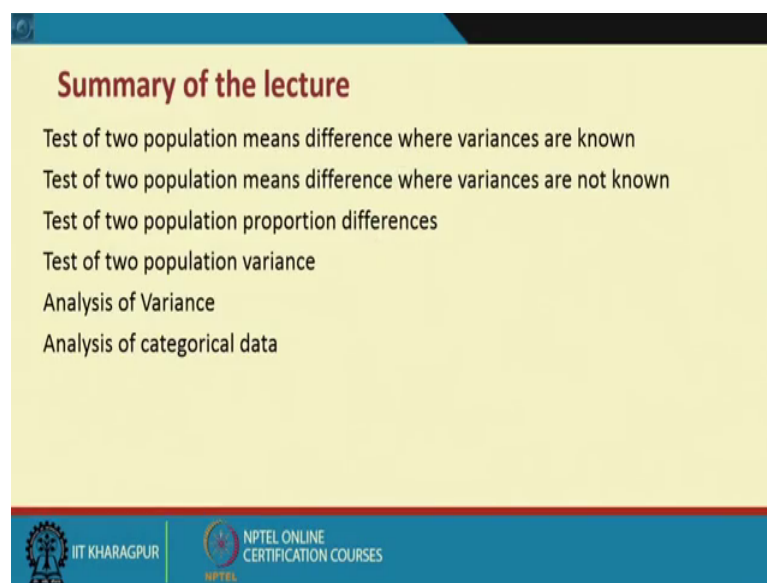
$\chi^2_{Cal} = 70.78 > 16.812, \text{ reject } H_0.$

IIT KHARAGPUR | NPTEL ONLINE CERTIFICATION COURSES

And in the case of you know critical, the critical value χ^2 is showing here 16.812 and that to four degree of freedom six that is $r - 1$ into $c - 1$ and that is $3 - 1$ into $2 - 1$ minus. So, that is coming actually it is $4 - 1$ into $3 - 1$ that is coming actually 6 and at 1 percent you know you know alpha levels, so you will find critical value is χ^2

16.818. So, here the calculated is 70.78 and the critical is 16.812. So, as a result so what is happening, so we are in a position to reject the true null hypothesis so that means, the hypothesis you have said here is that you know the types of gasoline is independent of income so that means they have the kind of you know you know dependencies. So, since we are you know rejecting so that means, the conclusion will be or the inference will be like that you know there is a kind of you know interdependency between these gasoline which subject to the income variations. So, likewise we have actually discussed couple of problems.

(Refer Slide Time: 42:10)



So, in this particular you know unit that is inferential analytics two, so we have actually solved some of the business problems, relating to multiple sample case and we have checked how is the kind of you know inference by using the t statistic, z statistic, chi square statistic and f statistic. And typically the variation in the kind of you know analysis of variance one factor model, two factor model, three factor models and the kind of you know ANOVA. So, that means, actually in a kind of you know a real life scenario or the kind of you know dynamic kind of you know business environment, you will find a some of the problems you know can be investigated with the multiple sample case. And again still find you know lots of you know interactive or interrelationship among these variables or the kind of you know samples. And the inferential analytics will give you the kind of you know exposures, how you have to understand these problems, and how you could you have to pick up a particular you know test statistic. And what is the kind of

you know optimum sampling, or you know the kind of you know or the kind of you know sample size which you can fix, and then how much or what is the kind of you know confident confidence interval we can build among these particular samples.

And these are the things we are actually you know supposed to know from this particular you know in inferential analytics too. So, you have to see what are the business problems. And after knowing all these techniques and the kind of you know concept the kind of you know connections, so I am very serious you know this kind of problems can be investigated very nicely or easily and then you can get some kind of you know inference as per the kind of you know business requirement. So, that means, technically you can you may be in a position to come with a some kind of you know management decision corresponding to a particular you know problems where we are dealing with a problem with multiple sample case. That means, typically a multivariate situations depending upon the kind of you know number of samples or you know number of variables, you can actually apply these techniques. And then solve some kind of you know business problems. With this, we will stop here.

And thank you very much, Have a nice day.