

Unit-7

Analysis of Variance :-

The analysis of variance is a powerful statistical tool for tests of significance. The test of significance based on t -distⁿ is an adequate procedure only for testing the significance of the difference betⁿ two sample means. In a situation when we have three or more samples to consider at a time an alternative procedure is needed for testing the hypothesis that all the samples are drawn from the same popⁿ. i.e. they have the same mean.

eg :- Five fertilizers are applied to four plots each of wheat and yields of wheat on each of the plot is given. Now we are interested in finding out whether the effect of these fertilizers on the yield is significantly different or in other words, whether the samples have come from the same normal popⁿ.

The answer to this problem is provided by the technique of analysis of variance. Thus basic purpose of the analysis of variance is to test the homogeneity of several means.

The term "Analysis of Variance" was introduced by prof. R. A. Fisher in 1930's to deal with problem in the analysis of agronomical data.

Cochran The : used to justify results relating to the prob. distⁿ of statistics that are used in ANOVA.

Page _____
Date _____

Variation is inherent in nature. The total variation in any set of numerical data is due to a no. of causes which may be classified as

- Assignable causes, and
- Chance causes.

The variation due to assignable causes can be detected and measured whereas the variation due to chance causes is beyond the control of human hand and cannot be traced separately.

Definition :-

According to Prof. R.A. Fisher, Analysis of Variance (ANOVA) is the "separation of variance ^{ascribable} to one group of causes from the variance ascribable to other group.

to credit
or assign
to a cause

The ANOVA consists in the estimation of the amount of variation due to each of the indept. factors (causes) separately and then comparing this estimates due to assignable factors (causes) with the estimates due to chance factor (causes).

Assumption :- For the validity of the F-test in ANOVA, the following assumptions are made :-

- The obserⁿs are indept.
- Parent popⁿ from which observation are taken is normal and
- Various treatment and environmental effects are additive in nature.

Cochran's Theorem :- Let $X_1, X_2, \dots, X_n$ denote a s.s. from normal popⁿ $N(0, \sigma^2)$. Let the sum of the squares of these values are

$$\sum_{i=1}^n X_i^2 = Q_1 + Q_2 + \dots + Q_k$$

Where Q_j is a quadratic form of $X_1, X_2, \dots, X_n$ with rank $(d+1)h_j; j=1, 2, \dots, k$. Then the s.v. $Q_1, Q_2, \dots, Q_k$ are mutually independent and Q_j/k^2 is a χ^2 variate with h_j degrees of freedom.

ANOVA for One-Way Classification :— Let

Suppose that N obsⁿ x_{ij} ($i=1, 2, \dots, k; j=1, 2, \dots, n_i$) of a Random Variable X are grouped on some basis, into k classes of sizes $n_1, n_2, \dots, n_k$ respectively. ($N = \sum_{i=1}^k n_i$) as exhibited below.

				Means	Total
x_{11}	x_{12}	...	x_{1n_1}	$\bar{x}_1$	T_1
x_{21}	x_{22}	...	x_{2n_2}	$\bar{x}_2$	T_2
$\vdots$	$\vdots$		$\vdots$	$\vdots$	$\vdots$
x_{k1}	x_{k2}	...	x_{kn_k}	$\bar{x}_k$	T_k
				$\bar{x}$	T

The total variation in the obsⁿ x_{ij} can be split into the following two components:

- The Variation betⁿ the classes or the variation due to different bases of classification, commonly known as treatments.
- The variations within the classes i.e. the inherent variation of the s.v. within the obsⁿ of a class.

The first type of variation is due to assignable causes which can be detected and controlled by human endeavour and the second type of variation is due to chance causes which are beyond the control of human hand.

The main object of analysis of variance technique is to ~~estimate~~^{examine} if there is significant difference between the class means ~~in view~~ of the due to chance causes i.e. variability within the ~~classes~~^{separate} classes.

eg. In particular let us consider the effect of k different rations on the yield of milk of N cows (of same breed & stock). These N cows are divided into k classes of sizes $n_1, n_2, \dots, n_k$ respectively.

$$N = \sum_{i=1}^k n_i$$

Here the sources of variation are

- (i) Effect of the rations (treatment) t_i $i=1, 2, \dots, k$
- (ii) error ϵ produced by numerous causes.

They are not detected and identified & produce a variation of random nature.

Mathematical Model :- Let the linear model will be

$$x_{ij} = \mu_i + \epsilon_{ij}$$

$$= \mu + (\mu_i - \mu) + \epsilon_{ij}$$

$$= \mu + \alpha_i + \epsilon_{ij} \quad \text{where } (i=1, 2, \dots, k; j=1, 2, \dots, n_i)$$

— (1)

(i) x_{ij} is the yield from the j th cow, ($j=1, 2, \dots, n_i$) fed on the i th rations ($i=1, 2, \dots, k$)

(ii) μ is the general mean effect given by

$$\mu = \frac{\sum_{i=1}^k (n_i \mu_i)}{N} \quad \text{--- (2)}$$

where μ_i is the fixed effect due to the i^{th} ration, i.e. if there were no treatment differences and no chance causes then the yield of each row will be μ .

(iii) α_i is the effect of the i^{th} ration (treatment) i.e. given by

$$\alpha_i = \mu - \mu_i \quad (i=1, 2, \dots, k) \quad \text{--- (3)}$$

i.e. the i^{th} treatment (ration) increase or decrease the yield by an amount α_i . On using (2), we get

$$\begin{aligned} \sum_{i=1}^k n_i \alpha_i &= \sum_i n_i (\mu_i - \mu) = \sum n_i \mu_i - \mu \sum n_i \\ &= N \cdot \mu - \mu \cdot N = 0 \end{aligned}$$

(iv) ε_{ij} is the error effect due to chance.

Assumption in the model

- (i) All the obserⁿ x_{ij} are indept.
- (ii) Different effects are additive in nature
- (iii) ε_{ij} are i.i.d. $N(0, \sigma^2)$

under the (ii) assumption, the model (i) becomes

$$E(x_{ij}) = \mu_i = \mu + \alpha_i \quad ; \quad \begin{cases} i=1, 2, \dots, k \\ j=1, 2, \dots, n_i \end{cases}$$

Null hypothesis of one for ANOVA for one-way classification.

Here we want to test the equality of popⁿ means i.e. the homogeneity of different ratios. Hence null hypothesis is given by

$$H_0: \mu_1 = \mu_2 = \dots = \mu_k = \mu$$

which from (3) reduces to

$$H_0: \alpha_1 = \alpha_2 = \dots = \alpha_k = 0$$

Statistical Analysis of the Model :- let us write

$\bar{x}_i$ = mean of the i th class.

$$= \frac{\sum_{j=1}^{n_i} x_{ij}}{n_i} ; (i=1, 2, \dots, k)$$

$$\bar{x}_{..} = \text{overall mean} = \frac{1}{N} \sum_{i=1}^k \sum_{j=1}^{n_i} x_{ij}$$

$$= \frac{1}{N} \sum_{i=1}^k n_i \bar{x}_i$$

The parameters μ and α_i in (1) are estimated by the principle of least squares on minimising the error (residual) sum of squares given by

$$E = \sum_i \sum_j \varepsilon_{ij}^2 = \sum_i \sum_j (x_{ij} - \mu - \alpha_i)^2 \quad \text{--- (1)}$$

The normal eqⁿ for estimating μ and α_i are

$$\frac{\partial E}{\partial \mu} = -2 \sum_i \sum_j (x_{ij} - \mu - \alpha_i) = 0 \quad \text{--- (*)}$$

$$\text{and } \frac{\partial E}{\partial \alpha_i} = -2 \sum_{j=1}^{n_i} (x_{ij} - \mu - \alpha_i) = 0 \quad \text{--- (**)}$$

from (*), we get

$$\sum_i \sum_j x_{ij} - N\mu - \sum_i n_i \alpha_i = 0$$

$$\hat{\mu} = \frac{1}{N} \sum_i \sum_j x_{ij} = \bar{x}_{..}$$

$$[\because \sum_i n_i \alpha_i = 0]$$

from (xx), we get

$$\sum_j x_{ij} - n_i \hat{\mu} - n_i \hat{\alpha} = 0$$

$$\Rightarrow \hat{\alpha} = \frac{1}{n_i} \sum_j x_{ij} - \hat{\mu} = \bar{x}_{i.} - \bar{x}_{..}$$

$$\hat{\alpha} = \bar{x}_{i.} - \bar{x}_{..}$$

Hence substituting the value of μ & α in (i) the model becomes

$$x_{ij} = \bar{x}_{..} + (\bar{x}_{i.} - \bar{x}_{..}) + (x_{ij} - \bar{x}_{i.})$$

We introduce the error component ϵ_{ij} so that both the sides are equal. This is the deviation within the class which is due to randomisation. Transposing $\bar{x}_{..}$ to the left squaring both sides and summing over i & j we get

$$\begin{aligned} \sum_i \sum_j (x_{ij} - \bar{x}_{..})^2 &= \sum_i \sum_j [\bar{x}_{i.} - \bar{x}_{..} + (x_{ij} - \bar{x}_{i.})]^2 \\ &= \sum_i \sum_j (\bar{x}_{i.} - \bar{x}_{..})^2 + \sum_i \sum_j (x_{ij} - \bar{x}_{i.})^2 \\ &\quad + 2 \sum_i (\bar{x}_{i.} - \bar{x}_{..}) \sum_j (x_{ij} - \bar{x}_{i.}) \end{aligned}$$

$$\sum_{i=1}^k \sum_{j=1}^{n_i} (x_{ij} - \bar{x}_{..})^2 = \sum_{i=1}^k \sum_{j=1}^{n_i} (x_{ij} - \bar{x}_{i.})^2 + \sum_i n_i (\bar{x}_{i.} - \bar{x}_{..})^2 + 2 \sum_i \left[(\bar{x}_{i.} - \bar{x}_{..}) \sum_j (x_{ij} - \bar{x}_{i.}) \right]$$

But $\sum_{j=1}^{n_i} (x_{ij} - \bar{x}_{i.}) = 0$ [algebraic sum of the deviates of the ratios from their mean is zero]

$$\therefore \sum_{i=1}^k \sum_{j=1}^{n_i} (x_{ij} - \bar{x}_{..})^2 = \sum_{i=1}^k \sum_{j=1}^{n_i} (x_{ij} - \bar{x}_{i.})^2 + \sum_i n_i (\bar{x}_{i.} - \bar{x}_{..})^2$$

$S_T^2 = \sum_{i=1}^k \sum_{j=1}^{n_i} (x_{ij} - \bar{x}_{..})^2$ is known as total sum of squares (T.S.S.)

$S_E^2 = \sum_{i=1}^k \sum_{j=1}^{n_i} (x_{ij} - \bar{x}_{i.})^2$ is called within sum of squares or error sum of squares (S.S.E.) and

$S_t^2 = \sum_i n_i (\bar{x}_{i.} - \bar{x}_{..})^2$ is called S.S. due to treatments (S.S.T.)

Then T.S.S. = S.S.E. + S.S.T.

Degrees of freedom for various S.S. :-

S_T^2 , the total S.S. will carry $(N-1)$ d.f. (one d.f. lost because of the linear constraint)

$$\sum_{i=1}^k \sum_{j=1}^{n_i} (x_{ij} - \bar{x}_{..}) = 0$$

similarly the treatment sum of squares S_t^2 will have $(k-1)$ d.f.

& S_E^2 the error s.s. will have $(N-k)$ d.f.

Hence we see that the

$$N-1 = (N-k) + (k-1)$$

Mean Sum of Squares (M.S.S.)

Defⁿ H_0

The sum of squares divided by its degrees of freedom gives the corresponding variance or the mean sum of squares (M.S.S.). Thus

$$\frac{S_t^2}{(k-1)} = \frac{S.S.T.}{(k-1)} = S_t^2 \text{ is M.S.S. due to treatment}$$

$$\& \frac{S_E^2}{N-k} = \frac{S.S.E.}{(N-k)} = S_E^2 \text{ is M.S.S. due to error.}$$

Hence the test statistic for H_0 provided by the variance ratio

$$F = \frac{S_t^2}{S_E^2} \quad (\text{---XXX---})$$

If H_0 is true the F - should take the value 1, otherwise it should be greater than unity. In order to find out if an observed value of F is significantly greater than unity. we have

to obtain the sampling distⁿ of the F-statistic defined as

$$F = \frac{S_t^2}{S_E^2}$$

Under H_0 by Cochran's thm., $\frac{S_t^2}{\sigma_e^2} \sim \chi^2_{k-1}$ & $\frac{S_E^2}{\sigma_e^2} \sim \chi^2_{N-k}$ are independently distributed as χ^2 variates with $(k-1)$ and $(N-k)$ d.f. respectively.

Hence the statistic

$$F = \left[\frac{S_t^2}{\sigma_e^2} \cdot \frac{1}{k-1} \right] \div \left[\frac{S_E^2}{\sigma_e^2} \cdot \frac{1}{N-k} \right] = \frac{S_t^2}{S_E^2}$$

follows Snedecor's F (central) distⁿ with $[k-1, N-k]$ d.f.

Thus if an observed value of F obtained from xxx is greater than tabulated value of F for $(k-1, N-k)$ d.f. at specified level of significance (usually 5% or 1%) then H_0 is refuted at that level otherwise H_0 may be retained.

The above statistical analysis is very elegantly presented in the following table

ANOVA Table for ONE-WAY CLASSIFIED DATA

Sources of Variation	Sum of Squares	d.f.	Mean Sum of Squares	Variance Ratio
Treatment (Treatments)	S_t^2	$k-1$	$\Delta_t^2 = \frac{S_t^2}{k-1}$	$\frac{\Delta_t^2}{\Delta_E^2} = F_{k-1, N-k}$
Error	S_E^2	$N-k$	$\Delta_E^2 = \frac{S_E^2}{N-k}$	
Total	S_T^2	$N-1$		