

MULTIPLE LINEAR REGRESSION VIEWPOINTS

VOL. 10 NO. 3

***MULTIPLE COMPARISONS
IN
HIGHER DIMENSIONAL
DESIGNS***

JOHN D. WILLIAMS

The University of North Dakota

COPYRIGHT 1980

Multiple Comparisons in Higher Dimensional Designs

TABLE OF CONTENTS

CHAPTER ONE	Introduction	1
CHAPTER TWO	Multiple Comparisons in the Analysis of Covariance	10
CHAPTER THREE	Multiple Comparisons For Treatments X Subjects Designs	20
CHAPTER FOUR	Multiple Comparisons in Two-Way Layouts	27
CHAPTER FIVE	Multiple Comparisons for Disproportionate Two-Way Layouts	47
CHAPTER SIX	Multiple Comparisons in Higher Dimensional Designs	59
CHAPTER SEVEN	Problems Less Amenable to a Direct Regression Solution	65
TABLES		78
REFERENCES		85

CHAPTER I

INTRODUCTION

Several explications of multiple comparisons are available for the usual one-way analysis of variance (ANOVA) situation; Miller's (1966) text might be the most definitive. Most standard texts on statistical analysis typically present one or more methods for the one-way ANOVA situation. A presentation of multiple comparisons in a regression framework was made by Williams (1976) wherein most of the more common multiple comparison procedures were considered.

For research situations more complex than the simple one-way ANOVA, most texts are silent. This silence is understandable; the issues, equations and other considerations become much more complex. Further, if the uncertainties of two or more issues are compounded, the morass of detail can be more than most authors wish to attempt a resolution for general readership. Winer (1971) presents some detail for more complex research situations that might shed light on the reluctance of traditional researchers to address the issue. The computations are often sufficient to ward off all but the most hardy. Perhaps the lack of writing regarding multiple comparisons in more complex situations than one-way ANOVA should serve as a warning, but "fools rush in where angels fear to tread"; lacking wings, we'll rush right in.

Background of the Reader

At this point, it is necessary to give some indication of the expectations regarding readers of this monograph. Some prior exposure to multiple

comparisons would be necessary, though no specific exposure is mandated. Hopefully, the reader is familiar with at least some of these tests: Dunnett's (1955) test for comparing one group to all other groups; Tukey's (1953) tests; Scheffé's (1959) test; Dunn's (1960) test; orthogonal comparisons; the Newman-Keuls test (Newman, 1939; Keuls, 1952); and Duncan's (1955) multiple range test. All of the foregoing tests are shown in a regression format by Williams (1976). For those whose exposure to multiple comparisons is minimal, articles by Ryan (1959, 1962) and Sparks (1963) would be useful.

In that the technique employed in this monograph is multiple linear regression, prior exposure to the use of this technique is helpful. Useful sources on multiple linear regression would include any of the following: Bottenberg and Ward (1963); Kelly, Beggs and McNeil (1969); Mendenhall (1976); Searle (1971); Ward and Jennings (1973); Williams (1976); Kerlinger and Pedhazur (1974); McNeil, Kelly and McNeil (1975); and Cohen and Cohen (1975), among many others.

Error Rates for Multiple Comparisons

In that several different kinds of error rates are used by various authors, it is useful to remind the reader of the more common error rates. Five of these are defined as

Per comparison = $\frac{\text{No. of comparisons incorrectly called significant.}}{\text{total number of comparisons}}$

Per experiment = $\frac{\text{No. of comparisons incorrectly called significant.}}{\text{total number of experiments}}$

Experimentwise = $\frac{\text{No. of experiments containing erroneous statements of significance}}{\text{total number of experiments}}$

Per family = $\frac{\text{No. of comparisons incorrectly called significant,}}{\text{total number of statements in the family}}$ (1.4)

Familywise = $\frac{\text{No. of experiments containing erroneous statements of significance}}{\text{total number of statements in the family}}$ (1.5)

Additional detail on error rates and concerns of a priori and a posteriori tests can be found in Williams (1976, 2-5).

The Use of Binary Coding and Dunnett's Test

In that the technique employed in this monograph is to utilize Dunnett's test with binary coding, a complete example is shown. Because repeated use of the test is employed, it should be recognized that the test is, except when Dunnett's test is being specifically employed, only a quasi-Dunnett's test. That is, Dunnett's test is used as a simple method to arrive at a solution. The test (actually, the table used to judge significance) remains a judgement left to the researcher.

Dunnett's Test for Comparisons of Several Treatment Groups with a Control

Dunnett (1955, 1964) devised a test that would allow the comparisons of several treatment groups with a control group and still retain an experimentwise error rate. This test could also be used whenever an experimenter wished to test a group which might be called the "experimental group" against several existing (but different) groups.

For example, a business educator may have devised a new approach to teaching beginning typewriting. The business educator may find that instead of finding one typical approach to teaching typewriting there may be

- (d.1) several methods being used. Rather than lumping all of the existing methods together and calling them a control group, it would seem more logical to test the new approach against each existing group separately, but in a single experiment. Dunnett's test is appropriate for this situation. So that the various tests can be compared to one another, a single data set is used several times in this monograph. That data set is

TABLE 1.1
DATA FOR DUNNETT'S TEST

Control Group Group One	Group Two	Group Three	Group Four
9	8	13	15
8	7	10	12
6	8	12	10
3	6	11	17
4	6	14	11

$$\bar{Y}_1 = 6.0, \bar{Y}_2 = 7.0, \bar{Y}_3 = 12.0, \bar{Y}_4 = 13.0.$$

Suppose the interest is in comparing the Control Group to Groups Two, Three and Four.

Viewing the problem from a regression viewpoint, it is helpful to define four binary predictors:

$X_1 = 1$ if the score is from a member of the control group (Group One); and 0 otherwise,

$X_2 = 1$ if the score is from a member of Group Two; and 0 otherwise,

$X_3 = 1$ if the score is from a member of Group Three; and 0 otherwise,

$X_4 = 1$ if the score is from a member of Group Four; and 0 otherwise.

A linear model can be written for this situation:

$$Y = b_0 + b_2X_2 + b_3X_3 + b_4X_4 + e_1. \quad (1.6)$$

where

b_0 = the Y-intercept,

b_2 = the regression coefficient for Group Two,

b_3 = the regression coefficient for Group Three,

b_4 = the regression coefficient for Group Four and

e_1 = the error involved in prediction.

It can be noticed that the control group has seemingly been left out. However, if equation 1.6 is solved for the expected value for a member of the control group,

$$E(Y) \hat{=} b_0 + b_2(0) + b_3(0) + b_4(0),$$

$$E(Y) \hat{=} b_0.$$

The expectancy for a member of the control group is by definition $\bar{Y}_1$. Thus, a least squares solution for b_0 is $\bar{Y}_1$, the mean of the control group.

For a member in Group Two, the expected value is

$$E(Y) \hat{=} b_0 + b_2(1) + b_3(0) + b_4(0),$$

$$E(Y) \hat{=} b_0 + b_2,$$

$$E(Y) \hat{=} \bar{Y}_1 + b_2. \quad (1.7)$$

A least squares solution for the expectancy of a given member of Group Two is the mean of Group Two. Thus

$$\bar{Y}_2 = \bar{Y}_1 + b_2, \text{ from equation 1.7, or}$$

$$\bar{Y}_2 - \bar{Y}_1 = b_2. \quad (1.8)$$

Likewise

$$b_3 = \bar{Y}_3 - \bar{Y}_1 \text{ and } b_4 = \bar{Y}_4 - \bar{Y}_1.$$

Equation 1.6 can be rewritten

$$Y = \bar{Y}_1 + (\bar{Y}_2 - \bar{Y}_1)X_2 + (\bar{Y}_3 - \bar{Y}_1)X_3 + (\bar{Y}_4 - \bar{Y}_1)X_4 + e_1. \quad (1.9)$$

Equation 1.9 lists precisely the comparisons of interest for comparing several treatments with a control. Since equation 1.6 (and, therefore, equation 1.9) is the same model as has been given for a one-way analysis of variance (Williams, 1971, 1974a), this approach also yields results identical to the analysis of variance situation. Thus, using equation 1.6, it can be seen that these two useful results can be obtained simultaneously: the usual analysis of variance as one part of the output, and Dunnett's test as the other part.

The information necessary for a regression solution, with equation 1.6 as the linear model, can be conveniently placed in tabular form (see Table

TABLE 1.2

REGRESSION FORMULATION FOR COMPARING SEVERAL
TREATMENTS WITH A CONTROL

Y	X ₁	X ₂	X ₃	X ₄
9	1	0	0	0
8	1	0	0	0
6	1	0	0	0
3	1	0	0	0
4	1	0	0	0
8	0	1	0	0
7	0	1	0	0
8	0	1	0	0
6	0	1	0	0
6	0	1	0	0
13	0	0	1	0
10	0	0	1	0
12	0	0	1	0
11	0	0	1	0
14	0	0	1	0
15	0	0	0	1
12	0	0	0	1
10	0	0	0	1
17	0	0	0	1
11	0	0	0	1

For the data in Table 1.2, a general purpose multiple regression program was used. Table 1.3 contains the printout from that analysis. The criterion variable is given as Y; variable 1 refers to the Control Group, variable 2 to Group Two, variable 3 to Group Three and variable 4 to Group Four. Because variable 1 refers to the Control Group, no information appears in the printout using that variable number. The table of residuals has not been included herein.

Table 1.3 contains the previously mentioned items. It can be recalled that $\bar{Y}_1 = 6.0$, $\bar{Y}_2 = 7.0$, $\bar{Y}_3 = 12.0$, $\bar{Y}_4 = 13.0$. The intercept is 6.0 (within rounding error) and is $\bar{Y}_1$. Also, $b_2 = 1 = \bar{Y}_2 - \bar{Y}_1$, and is in keeping with equation 1.9. Similar statements could be made concerning b_3 and b_4 . The computed t values in Table 1.3 are identically the same values as would result from the use of Dunnett's test. It is only necessary to compare each of these values to Dunnett's table for the test of significance. From Table Id, a computed t value of 3.39 is needed for significance at the .01 level on a two-tailed test. Thus, both Groups Three and Four are significantly higher than the Control Group. It thus can be seen that the computed t values, which are tests of the partial regression weights, should be evaluated in this instance not by the traditional t table, but by use of Dunnett's tables.

TABLE 1.3

OUTPUT OF MULTIPLE REGRESSION PROGRAM

Variable No.	Mean	Standard Deviation	Correlation X vs Y	Regression Coefficient	Standard Error of Regression Coefficient	Computed t Value	Beta
2	0.25000	0.44426	-0.40109	1.00000	1.36014	0.73522	0.12033
3	0.25000	0.44426	0.40109	6.00000	1.36014	4.41130	0.72197
4	0.25000	0.44426	0.56153	7.00000	1.36014	5.14652	0.84230
Dependent Y	9.50000	3.69210					
Intercept		6.00000					
Multiple Correlation		0.84515					
St. Error of Estimate		2.15058					

Analysis of Variance for the Regression

Source of Variation	Degrees of Freedom	Sum of Squares	Mean Squares	F Value
Attributable to Regression	3	185.00027	61.66675	13.33340
Deviation from Regression	16	73.99973	4.62498	
Total	19	259.00000		

Organization of this Monograph

The remaining portions of the monograph address multiple comparison concerns not usually addressed by most available sources. An additional objective is to reduce computations to a minimum. While computer programs could easily be written for each specific situation, all that is necessary is a general usage multiple linear regression program. The print-out shown in Table 1.3 would yield sufficient information for completing any of the situations described in this monograph.

Chapter Two addresses multiple comparisons in the analysis of covariance. The method described is much easier to accomplish than the process described by Winer. Chapter Three considers the treatments X subjects design (repeated measures). Chapters Four and Five consider the two-way ANOVA for the equal cell case and the disproportional case, respectively. The N-way ANOVA is considered in Chapter Six. Finally, Chapter Seven considers situations that do not lend themselves to a simple solution as described here.

CHAPTER II

MULTIPLE COMPARISONS IN THE ANALYSIS OF COVARIANCE

In the analysis of covariance, each separate comparison will have its own standard error of estimate even if equal N occur in each cell. The equation for the standard error of estimate for a comparison in the analysis of covariance is given by Winer (1971, p. 772),

$$s_{\bar{Y}_i \text{adj} - \bar{Y}_j \text{adj}} = \sqrt{MS'_W \left[\frac{1}{n_i} + \frac{1}{n_j} + \frac{(\bar{X}_i - \bar{X}_j)^2}{E_{xx}} \right]} \quad (2.1)$$

where

$\bar{Y}_i \text{adj}$ = the adjusted mean for group i ;

$\bar{Y}_j \text{adj}$ = the adjusted mean for group j ;

MS'_W = the error term in the analysis of covariance;

n_i, n_j = respectively cell frequencies for the i^{th} and j^{th} groups;

$\bar{X}_i, \bar{X}_j$ = respectively the means on the covariate for the i^{th} and j^{th} groups; and

E_{xx} = SS_W for the covariate.

While researchers may feel justifiably ill at ease in attempting to use equation 2.1, the use of regression can eliminate the tedious calculations. Further, more than one covariate can easily be accommodated.

An Example

Table 2.1 is taken from Williams (1974a, p. 104 and p. 109). In Table 2.1, X_1 is a binary variable for membership in group 1, X_2 is a binary variable for membership in group 2 and X_3 is similarly a binary

variable for membership in group 3. Also, X_4 represents a pretest score and X_5 represents a measure of intelligence; the Y value represents a posttest score. Only the pretest is considered as a covariate in this section; both the pretest and intelligence are considered as covariates in the section under multiple covariates.

TABLE 2.1

Data for the Analysis of Covariance

Y	X_1	X_2	X_3	X_4	X_5
35	1	0	0	12	120
27	1	0	0	17	98
32	1	0	0	13	102
29	1	0	0	10	106
27	1	0	0	8	94
38	0	1	0	29	123
25	0	1	0	12	96
36	0	1	0	17	108
35	0	1	0	22	115
31	0	1	0	15	128
27	0	0	1	17	90
35	0	0	1	22	110
19	0	0	1	10	94
17	0	0	1	8	95
32	0	0	1	13	116

Under the assumption of a single regression line on the covariate (the pretest, X_4) an analysis of covariance can be accomplished with two linear models:

$$Y = b_0 + b_1X_1 + b_2X_2 + b_4X_4 + e_1, \quad (2.2)$$

and

$$Y = b_0 + b_4X_4 + e_2. \quad (2.3)$$

In that a large part of the print-out regarding equation 2.2 is useful, the print-out is reproduced in Table 2.2.

TABLE 2.2

Print-Out for Equation 2.2

VARIABLE	MEAN	STAND. DEV.	CORRELATION X VS Y	REG. COEF.	STD. ERROR OF REG. COEF.	COMPUTED T VALUE
4	15.00000	5.85539	0.68903	0.76000	0.22783	3.33582
1	0.33333	0.48795	0.03983	5.52000	2.73396	2.01905
2	0.33333	0.48795	0.39328	3.20000	2.92653	1.09345
DEPENDENT						
Y	29.66666	6.12566				
INTERCEPT		15.36000				
MULTIPLE CORRELATION		0.78714				
STD. ERROR OF ESTIMATE		4.26230				
MULTIPLE CORRELATION SQUARED			0.61959			
ONE MINUS MULTIPLE CORRELATION SQD.			0.38041			

Analysis of Variance for the Regression

SOURCE OF VARIATION	df	SS	MS	F
ATTRIBUTABLE TO REGRESSION	3	325.49292	108.49763	5.97215
DEVIATION FROM REGRESSION	11	199.83984	18.16725	
TOTAL	14	525.33276		

The usual analysis of covariance can be completed by using

$$F = \frac{(R_2^2 - R_3^2)/(g - 1)}{(1 - R_2^2)/dfw_2} = \frac{(.61959 - .47476)/2}{(1 - .61959)/11} = 2.09,$$

which for $df = 2, 11$, $p > .05$.

In equation 2, the X_3 variable has been omitted. Thus $b_1 = \bar{Y}_1\text{adj} - \bar{Y}_3\text{adj}$ and $b_2 = \bar{Y}_2\text{adj} - \bar{Y}_3\text{adj}$. To find the adjusted means, the following equations can be used:

$$\bar{Y}_3\text{adj} = b_0 + b_4\bar{X}_4 = 15.36 + .76(15) = 26.76;$$

$$\bar{Y}_1\text{adj} = b_1 + \bar{Y}_3\text{adj} = 5.52 + 26.76 = 32.28; \text{ and}$$

$$\bar{Y}_2\text{adj} = b_2 + \bar{Y}_3\text{adj} = 3.20 + 26.76 = 29.96.$$

The adjusted values agree with those originally given by Williams (1974a, p. 106), though the method shown here is simplified somewhat.

More importantly, the standard error of the regression coefficients corresponding to X_1 and X_2 are respectively equal to the standard errors from equation 2.1 for comparing $\bar{Y}_1\text{adj}$ to $\bar{Y}_2\text{adj}$ and $\bar{Y}_1\text{adj}$ to $\bar{Y}_3\text{adj}$. Thus, the computed t values given in Table 2.2 are directly usable in whichever multiple comparison procedure the researcher prefers. Were there interest in comparing $\bar{Y}_1\text{adj}$ to $\bar{Y}_2\text{adj}$, a model of the form

$$Y = b_0 + b_1X_1 + b_3X_3 + b_4X_4 + e_1 \quad (2.4)$$

could be used, with focus on the computed t value for the X_1 variable.

Complex Comparisons

Complex comparisons, or contrasts, can be completed in a regression analysis for the analysis of covariance as well. Suppose a contrast of the form

$$\psi_1 = \bar{Y}_3\text{adj} - \frac{1}{2}\bar{Y}_1\text{adj} - \frac{1}{2}\bar{Y}_2\text{adj} \quad (2.5)$$

is contemplated. First, equation 2.2 is reparametrized as

$$Y = b_1X_1 + b_2X_2 + b_3X_3 + b_4X_4 + e_1. \quad (2.6)$$

Then a restriction corresponding to ψ_1 , $b_3 = \frac{1}{2}b_1 + \frac{1}{2}b_2$ is placed on equation 2.6:

$$Y = b_1X_1 + b_2X_2 + (\frac{1}{2}b_1 + \frac{1}{2}b_2)X_3 + b_4X_4 + e_3.$$

Or,

$$Y = b_1(X_1 + \frac{1}{2}X_3) + b_2(X_2 + \frac{1}{2}X_3) + b_4X_4 + e_3. \quad (2.7)$$

Two new variables can be constructed such that $V_1 = 1$ if a member of group 1, $\frac{1}{2}$ if a member of group 3, 0 if a member of group 2; and $V_2 = 1$ if a member of group 2, $\frac{1}{2}$ if a member of group 3, 0 if a member of group 1. Then equation 2.7 can be rewritten as

$$Y = b_1V_1 + b_2V_2 + b_4X_4 + e_3. \quad (2.8)$$

Equation 2.8 (and also equation 2.6) could be processed using a program such as Ward and Jennings' (1973) DATRAN or McNeil et al.'s (1975) LINEAR.

However, equation 2.8 can also be reparametrized back into a form as a unit vector as was done earlier. This can be accomplished by setting either b_1 or b_2 equal to zero. Setting $b_2 = 0$ yields

$$Y = b_0 + b_1V_1 + b_4X_4 + e_3. \quad (2.9)$$

Then $R_9^2 = .50151$

$$\text{To test } \psi_1, t = \frac{R_2^2 - R_9^2}{(1-R_2^2)/11} = \frac{.61959 - .50151}{(1-.61959)/11};$$

$t = 1.85, p > .05.$

Concerns of Homogeneity Regression

To this point, the assumption of a single regression line for the covariate has been made. A test can be made of this assumption; three new variables are defined such that

$$X_6 = X_1 \cdot X_4;$$

$$X_7 = X_1 \cdot X_4; \text{ and}$$

$$X_8 = X_3 \cdot X_4.$$

Then a model can be written as

$$Y = b_0 + b_1X_1 + b_2X_2 + b_6X_6 + b_7X_7 + b_8X_8 + e_4. \quad (2.10)$$

$R^2_{10} = .71825$. To test this for significance,

$$F = \frac{(R^2_{10} - R^2_2)/(g - 1)}{(1 - R^2_{10})/(N - 2g)},$$

$$F = \frac{(.71825 - .61959)/2}{(1 - .71825)/9} \approx 1.58; p > .05.$$

Had the F value been significant, some researchers would prefer to abandon the analysis given earlier; interpretations become increasingly difficult. It would be inappropriate to attempt to use the computed t values for testing b_1 and b_2 in equation 2.10. The "adjusted means" would occur where separate regression lines are used for each group on the covariate. Since the covariance process is occurring separately for each group, differences in the adjusted means would not test any meaningful hypotheses regarding group differences on the criterion score. Table 2.3 should help show why this is so.

TABLE 2.3

Regression Output With Separate Regression Lines for Each Group on the Covariate

VARIABLE	MEAN	STAND. DEV.	CORRELATION X VS Y	REG. COEF.	STD. ERROR OF REG. COEF.	COMPUTED T-VALUE
6	4.00000	6.12955	0.03995	0.02173	0.59792	0.03623
7	6.33333	9.93311	0.50556	0.64045	0.30396	2.10702
8	4.66667	7.46101	-0.16514	1.19841	0.36127	3.31721
1	0.33333	0.48795	0.03983	20.51697	9.14573	2.24344
2	0.33333	0.48795	0.39828	11.60924	8.09397	1.43434
DEPENDENT						
Y	29.66666	6.12567				
INTERCEPT	9.22224					
MULTIPLE CORRELATION	0.84750					
STD. ERROR OF ESTIMATE	4.05527					
MULTIPLE CORRELATION SQUARED			0.71826			
ONE MINUS MULTIPLE CORRELATION SQD			0.28174			

Analysis of Variance for the Regression

SOURCE OF VARIATION	df	SS	MS	F
ATTRIBUTABLE TO REGRESSION	5	377.32568	75.46513	4.58886
DEVIATION FROM REGRESSION	9	148.00757	16.44528	

Does the information given in Table 2.3 suggest that $\bar{Y}_1\text{adj} - \bar{Y}_3\text{adj} = 20.51697$? The answer is a qualified "no". Only under the condition that each group has its separate regression on the covariate, and its separate mean on the covariate would $b_1 = \bar{Y}_1\text{adj} - \bar{Y}_3\text{adj}$. However, that condition is very different than most users of the analysis of covariance would wish to use.

It is clearly quite different from asking, "If the groups were equal (on the covariate) at the beginning of the experiment, how do they compare at the end?" Even if all groups are "adjusted" by using a covariate mean of $\bar{X}_4 = 15$, the difference in the regression coefficients preclude interpreting b_1 as "a treatment difference after covariate adjustment between groups 1 and 3". The analysis of covariance is usually enlisted to test treatment differences in groups whose members were unable to be randomly assigned to a treatment group, so that a statistical control is used. While a test of significance on b_1 can be legitimately done, it does not address questions usually asked by researchers using the analysis of covariance.

If the expected values are found for each group at $X_4 = 15$, then a more useful result can be found:

$$E(Y_1, X_4 = 15) \hat{=} b_0 + b_1 + b_6(15) \hat{=} 9.22224 + 20.51697 + .02173(15) \hat{=} 30.06516.$$

$$\text{Likewise, } E(Y_2, X_4 = 15) \hat{=} b_0 + b_2 + b_7(15) \hat{=} .$$

$$9.22224 + 11.60924 + .64045(15) \hat{=} 30.43823$$

and

$$E(Y_3, X_4 = 15) \hat{=} b_0 + b_8(15) \hat{=} 9.22224 + 1.19841(15) \hat{=} 27.19839.$$

Where testing for the analysis of covariance (locking homogeneity of regression), it is the difference between $E(Y_1, X_4 = 15) \hat{=} 30.06516$ and $E(Y_3, X_4 = 15) \hat{=} 27.19839$ or $30.06516 - 27.19839 \hat{=} 2.86677$ that is likely to be of interest rather than thinking b_1 represents the difference between $\bar{Y}_1 \text{ adj}$ and $\bar{Y}_3 \text{ adj}$.

The specific hypothesis tested in the analysis of covariance is $E(Y_1, X = \bar{X}) = E(Y_2, X = \bar{X}) = \dots = E(Y_g, X = \bar{X})$.

Testing this type of hypothesis is shown in a series of steps in the preface of Ward and Jennings (1973, p. xvii-xviii) test, Introduction to Linear Models. The interested reader is referred to that introduction; in the preface, testing this type of hypothesis is outlined in steps 3 to step 11.

Multiple Covariates

Extensions to more than one covariate can easily be accommodated both for the analysis of covariance and for multiple comparisons. The intelligence score, X_5 , could be used together with the pretest as covariates. Assuming single regression lines for all three groups on the two covariates, the model can be given as

$$Y = b_0 + b_1X_1 + b_2X_2 + b_4X_4 + b_5X_5 + e_5. \quad (2.11)$$

The use of the computed t values for b_1 and b_2 allow a test regarding differences among the adjusted means for comparing groups 1 and 2 with group 3 respectively; $t_1 = 1.95059$ and $t_2 = .38191$. To test the differences between the adjusted means of group 1 to group 2, a model such as

$$Y = b_0 + b_1X_1 + b_3X_3 + b_4X_4 + b_5X_5 + e_5 \quad (2.12)$$

can be used.

Here, $t_1 = 1.33421$; also, $t_3 = -.38191$, reaffirming the t value for the difference between the adjusted means of groups 2 and 3. The sign is changed because the direction of the comparison has changed; for b_2 in equation 11, t_2 addresses $\bar{Y}_2 \text{ adj} - \bar{Y}_3 \text{ adj}$, for b_3 in equation 12, t_3 addresses $\bar{Y}_3 \text{ adj} - \bar{Y}_2 \text{ adj}$.

MULTIPLE COMPARISONS FOR TREATMENTS X SUBJECTS DESIGNS

To consider multiple comparisons for treatments X subjects designs (or repeated measure designs) an example taken from Chapter 7 of Williams (1974a, p. 56) is used. See Table 3.1.

TABLE 3.1

THREE TREATMENT METHODS OF PAIRED-ASSOCIATE LEARNING
WITH EDUCABLE MENTALLY RETARDED SUBJECTS

Subject	Treatment One	Treatment Two	Treatment Three
1	18	27	15
2	17	24	14
3	14	13	12
4	5	8	6
5	11	14	10
6	9	12	8
7	14	16	15
8	12	17	9
9	22	21	16
10	10	18	15

The information in Table 3.1 can be placed in a tabular form suitable for use in a regression format; see Table 3.2.

TABLE 3.2

Illustration of Design Matrix for Treatments X Subjects Designs

Y	X ₁	X ₂	X ₃	X ₄	X ₅	X ₆	X ₇	X ₈	X ₉	X ₁₀	X ₁₁	X ₁₂	X ₁₃	X ₁₄
18	1	0	0	1	0	0	0	0	0	0	0	0	0	60
27	0	1	0	1	0	0	0	0	0	0	0	0	0	60
15	0	0	1	1	0	0	0	0	0	0	0	0	0	60
17	1	0	0	0	1	0	0	0	0	0	0	0	0	55
24	0	1	0	0	1	0	0	0	0	0	0	0	0	55
14	0	0	1	0	1	0	0	0	0	0	0	0	0	55
14	1	0	0	0	0	1	0	0	0	0	0	0	0	39
13	0	1	0	0	0	1	0	0	0	0	0	0	0	39
12	0	0	1	0	0	1	0	0	0	0	0	0	0	39
5	1	0	0	0	0	0	1	0	0	0	0	0	0	19
8	0	1	0	0	0	0	1	0	0	0	0	0	0	19
6	0	0	1	0	0	0	1	0	0	0	0	0	0	19
11	1	0	0	0	0	0	0	1	0	0	0	0	0	35
14	0	1	0	0	0	0	0	1	0	0	0	0	0	35
10	0	0	1	0	0	0	0	1	0	0	0	0	0	35
9	1	0	0	0	0	0	0	0	1	0	0	0	0	29
12	0	1	0	0	0	0	0	0	1	0	0	0	0	29
8	0	0	1	0	0	0	0	0	1	0	0	0	0	29
14	1	0	0	0	0	0	0	0	0	1	0	0	0	45
16	0	1	0	0	0	0	0	0	0	1	0	0	0	45
15	0	0	1	0	0	0	0	0	0	1	0	0	0	45
12	1	0	0	0	0	0	0	0	0	0	1	0	0	38
17	0	1	0	0	0	0	0	0	0	0	1	0	0	38
9	0	0	1	0	0	0	0	0	0	0	1	0	0	38
22	1	0	0	0	0	0	0	0	0	0	0	1	0	59
21	0	1	0	0	0	0	0	0	0	0	0	1	0	59
16	0	0	1	0	0	0	0	0	0	0	0	1	0	59
10	1	0	0	0	0	0	0	0	0	0	0	0	1	43
18	0	1	0	0	0	0	0	0	0	0	0	0	1	43
15	0	0	1	0	0	0	0	0	0	0	0	0	1	43

The values in Table 3.2 are defined as follows:

Y = the criterion score;

X₁ = 1 if the score corresponds to Treatment 1, 0 otherwise;

X₂ = 1 if the score corresponds to Treatment 2, 0 otherwise;

X₃ = 1 if the score corresponds to Treatment 3, 0 otherwise;

X₄ = 1 if the score is obtained from Subject 1, 0 otherwise;

X₅ = 1 if the score is obtained from Subject 2, 0 otherwise;

$X_6 = 1$ if the score is obtained from Subject 3, 0 otherwise;
 $X_7 = 1$ if the score is obtained from Subject 4, 0 otherwise;
 $X_8 = 1$ if the score is obtained from Subject 5, 0 otherwise;
 $X_9 = 1$ if the score is obtained from Subject 6, 0 otherwise;
 $X_{10} = 1$ if the score is obtained from Subject 7, 0 otherwise;
 $X_{11} = 1$ if the score is obtained from Subject 8, 0 otherwise;
 $X_{12} = 1$ if the score is obtained from Subject 9, 0 otherwise;
 $X_{13} = 1$ if the score is obtained from Subject 10, 0 otherwise; and
 X_{14} = the sum of the criterion scores for each subject separately.

The analysis in Williams (1974a) proceeds as follows: three linear models are defined, one for the treatments effect, one for the subjects effect and one for the combined treatments and subjects effects. These models are given as

$$Y = b_0 + b_1X_1 + b_2X_2 + e_1, \quad (3.1)$$

$$Y = b_0 + b_4X_4 + b_5X_5 + b_6X_6 + b_7X_7 + b_8X_8 + b_9X_9 + b_{10}X_{10} + b_{11}X_{11} + b_{12}X_{12} + e_2, \quad (3.2)$$

and

$$Y = b_0 + b_1X_1 + b_2X_2 + b_4X_4 + b_5X_5 + b_6X_6 + b_7X_7 + b_8X_8 + b_9X_9 + b_{10}X_{10} + b_{11}X_{11} + b_{12}X_{12} + e_3. \quad (3.3)$$

The associated R^2 values and sums of squares (SS) for equations 1-3 are $R_1^2 = .1784$; $SS_1 = 136.27$; $R_2^2 = .6823$; $SS_2 = 521.20$; $R_3^2 = .8607$; $SS_3 = 657.47$; $SS_T = 763.86$. A complete summary table is shown in Table 3.3.

TABLE 3.3

SUMMARY TABLE FOR THE TREATMENTS X SUBJECTS DESIGN

Source of Variation	df	SS	MS	F
Treatments	2	136.27	68.13	11.52
Subjects	9	521.20		
Error	18	106.39	5.91	
Total	29	763.86		

An alternative analysis using X_{14} , the sum of criterion scores for each subject, would use the following equations (Williams, 1977a):

$$Y = b_0 + b_1X_1 + b_2X_2 + e_1, \quad (3.1)$$

$$Y = b_0 + b_{14}X_{14} + e_2, \quad (3.4) \text{ and}$$

$$Y = b_0 + b_1X_1 + b_2X_2 + b_{14}X_{14} + e_3. \quad (3.5)$$

The results, in terms of sums of squares and R^2 values, is identical to that already given, with $R_4^2 = R_2^2$, $SS_4 = SS_2$, $R_5^2 = R_3^2$ and $SS_5 = SS_3$. However, care must be taken with the degrees of freedom. Equation 3.4 uses only one predictor; thus, the "apparent" degrees of freedom is one. It must be remembered that the actual $df = N - 1$.

The analysis to this point has been focused on constructing the usual ANOVA summary table rather than being concerned with the possible comparisons of means. Suppose there is interest in testing hypotheses regarding the means: $\mu_1 = \mu_2$; $\mu_1 = \mu_3$ and $\mu_2 = \mu_3$. Testing these hypotheses can be accomplished with two different formulations of the full model, equation 3.3 and

$$Y = b_0 + b_1X_1 + b_3X_3 + b_4X_4 + b_5X_5 + b_6X_6 + b_7X_7 + b_8X_8 + b_9X_9 + b_{10}X_{10} + b_{11}X_{11} + b_{12}X_{12} + e_3. \quad (3.3a)$$

Part of the printout from the use of equation 3.3 is shown in Table 3.4.

TABLE 3.4

OUTPUT OF FULL MODEL FOR TREATMENTS X SUBJECTS DESIGN

Variable No.	Mean	Standard Deviation	Correlation X VS Y	Regression Coefficient	Std. Error Of Reg. Coef.	Computed T Value	Beta
1	0.33333	0.47946	-0.12145	1.19998	1.08732	1.10362	0.11210
2	0.33333	0.47946	0.41105	4.99997	1.08732	4.59846	0.46710
4	0.10000	0.30513	0.39195	5.66663	1.98515	2.85451	0.33690
5	0.10000	0.30513	0.28185	4.00001	1.98515	2.01496	0.23781
6	0.10000	0.30513	-0.07046	-1.33331	1.98515	-0.67164	-0.07927
7	0.10000	0.30513	-0.51085	-7.99992	1.98515	-4.12987	-0.47562
8	0.10000	0.30513	-0.15854	-2.66665	1.98515	-1.34329	-0.15854
9	0.10000	0.30513	-0.29066	-4.66664	1.98515	-2.35077	-0.27745
10	0.10000	0.30513	0.06166	0.66668	1.98515	0.33583	0.03964
11	0.10000	0.30513	-0.09248	-1.66665	1.98515	-0.83956	-0.09909
12	0.10000	0.30513	0.36993	5.33332	1.98515	2.68661	0.31708
Dependent Y	14.06667	5.13226					
INTERCEPT		12.26667					
MULTIPLE CORRELATION		0.92774					
STD. ERROR OF ESTIMATE		2.43131					
MULTIPLE CORRELATION SQUARED			0.86070				
ONE MINUS MULTIPLE CORRELATION SQD.			0.13930				

ANALYSIS OF VARIANCE FOR THE REGRESSION

Source of Variation	Degrees Of Freedom	Sum Of Squares	Mean Squares	F Value
Attributable To Regression	11	657.46021	59.76910	10.11102
Deviation From Regression	18	106.40308	5.91128	
Total	29	763.86328		

From Table 3.4, it can be seen that $t_1 = 1.10362$ and $t_2 = 4.59846$; these t values are respectively the tests regarding comparing $\bar{Y}_1$ to $\bar{Y}_3$ and $\bar{Y}_2$ to $\bar{Y}_3$, taking into account that the subjects serve as their own controls. A similar printout could be generated using a model corresponding to equation 3.3a. Values from this printout show $t_1 = -3.49484$ $t_3 = -4.59847$; these t values correspond to comparing $\bar{Y}_1$ to $\bar{Y}_2$ and $\bar{Y}_3$ to $\bar{Y}_2$. Also, the corresponding means are $\bar{Y}_1 = 13.20$, $\bar{Y}_2 = 17.00$ and $\bar{Y}_3 = 12.00$. Because the comparisons are being done on an a posteriori basis, and because all possible simple comparisons are being evaluated, Tukey's (H.S.D.) test is an appropriate multiple comparison procedure to evaluate the differences in means. Tables IIa and IIb utilize the computed t values directly for the Studentized range statistic; for $\alpha = .01$, $t = 3.326$ and for $\alpha = .05$, $t = 2.553$. Two t values are significant; comparing $\bar{Y}_1$ to $\bar{Y}_2$ ($t = 3.49484$, disregarding the sign) and comparing $\bar{Y}_2$ to $\bar{Y}_3$ ($t = 4.59846$), both exceed 3.326 ($p < .01$). To complete these comparisons using Tukey's test, $v = df_w = 18$ and $r = 3$.

Using the Shortcut Method

The solution just given in the last section presumed that each subject (except one) is separately coded using a binary coding scheme. Clearly, if the number of subjects is at all large, the coding procedure described in Williams (1977a) and using equations 3.1, 3.4, and 3.5 might be preferable. However, one difficulty with using this shortcut procedure is that the standard error of the regression coefficients for X_1 and X_2 are too small due to the degrees of freedom,

as generated by the computer program, not being accurate for deviation from regression. These t values could be adjusted by multiplying by an appropriate constant. The appropriate constant is: $c = \frac{MS_{w3.5}}{MS_{w3.3}}$ (3.6)

where $MS_{w3.5}$ is the mean square within (or deviation from regression) for equation 3.5 and $MS_{w3.3}$ is the mean square within for equation 3.3. The $MS_{w3.5}$ is 4.09225 and $MS_{w3.3}$ is 5.91125. Thus, $c = .83203$. The

values generated by equation 3.5 for t_1 and t_2 (comparing $\bar{Y}_1$ to $\bar{Y}_3$ and $\bar{Y}_2$ to $\bar{Y}_3$) are $t_1 = 1.32641$ and $t_2 = 5.52678$. Multiplying t_1 and t_2 by c yields corrected $t_1 = 1.10361$ and corrected $t_2 = 4.59845$, within rounding error of the values found earlier. Of course, $MS_{w3.3}$ would not be available were the researcher using the shortcut method.

However, $MS_{w3.3} = \frac{SS_{w3.5}}{N-S-g+1}$ (3.7) where N is the total number of scores,

S is the number of subjects and g is the number of groups. The denominator in equation 3.7 can also be found as $(S-1)(g-1)$.

MULTIPLE COMPARISONS IN TWO-WAY LAYOUTS

Before considering specific questions in a multiple comparison framework given two-way designs, it is useful to consider several unique concerns of such analyses. First, some rectification needs to be done regarding decisions on whether to employ a family approach, or to use the experiment with all of its dimensions as the unit for deciding upon the number of comparisons. That is, are results for rows to be interpreted without considering the results for columns and/or interactions? Similar questions could be posed for columns and interactions as well. On the one hand, statistical tradition would be on the side of using familywise error rates where the row comparisons, column comparisons and comparisons involving interactions would be interpreted with their own error rates. On the other hand, if the row, column and interaction effects are tested independently at, say, the .05 level, then the overall error rate is $1 - (1-.05)^3$. The logic of multiple comparisons would suggest that all comparisons in the experiment be taken into account in reporting probabilities.

In the long run, the researcher has to take the responsibility of deciding whether to use a family approach or to use an approach that uses the experiment as the unit for comparisons. Whatever decision is made, communication should be made to readers so that they understand the ground rules used by the researcher. Having at least brought up some of the unique concerns of a two-way layout, we can now turn our attention to

an actual problem.

An Example for A Two-Way Layout

The following example is taken from Williams (1974a)

Suppose a researcher wishes to measure the effects of three different approaches to teaching arithmetic combined with four different methods of assigning homework. Suppose there are 24 students available, so the students are randomly assigned to one of the 12 (3x4) treatment combinations. When the experiment is concluded, a standardized arithmetic test consisting of 20 items with the items in a multiple choice format is administered. The three different methods of teaching arithmetic are 1) teaching arithmetic with the problems very similar to the final test; 2) teaching arithmetic with half the problems similar to the final test, and the other half dissimilar to the final test; and 3) teaching arithmetic with the problems dissimilar to the final test. The four different methods of homework are 1) problems in a multiple-choice format; 2) homework is of a project nature; 3) problems from the book in which the student supplies his own answer, and 4) no homework. The data for this experiment are in Table 4.1.

TABLE 4.1

DATA FOR METHOD OF TEACHING X TYPE OF HOMEWORK

Type of Homework	<u>Method of Teaching</u>		
	Very Similar Material to Final	Half Similar, Half Dissimilar Material to Final	Dissimilar Material to Final
Multiple Choice Items	18 10	10 6	12 10
Project	5 3	2 2	9 3
Student Supplies Own Answer	7 3	6 2	4 2
No Homework	2 0	7 5	10 6

When proportionality is present, several alternative approaches will yield the same end products in the summary table. A simple to use approach is to successively analyze the data into three different one-way layouts. One layout would consider each treatment combination as a separate group. Thus, there would be $(4)(3) = 12$ groups. An analysis of variance could be generated by the following model (Model 1):

$$Y = b_0 + b_1X_1 + b_2X_2 + b_3X_3 + b_4X_4 + b_5X_5 + b_6X_6 + b_7X_7 + b_8X_8 + b_9X_9 + b_{10}X_{10} + b_{11}X_{11} + e_1, \quad (4.1)$$

where

Y = the criterion variable,

b_0 = the Y -intercept (in this model, b_0 will equal the mean of Group Twelve),

$X_1 = 1$ if a score is from a member of Group 1; 0 otherwise,

$X_2 - X_{12} = 1$ if a score is from a member of the group corresponding to the subscript number; 0 otherwise,

$b_1 - b_{11}$ = the regression coefficients for $X_1 - X_{11}$, and

e_1 = the error in prediction with this model.

A second model (Model 2) would include only information regarding the rows, disregarding the particular column:

$$Y = b_0 + b_{13}X_{13} + b_{14}X_{14} + b_{15}X_{15} + e_2, \quad (4.2)$$

where

b_0 = the Y -intercept for this model (in this Model, b_0 will equal the mean of the fourth row, no homework),

$X_{13} = 1$ if the score is from a member in row 1 (Multiple Choice Items); 0 otherwise;

$X_{14} = 1$ if the score is from a member in row 2 (Project);
0 otherwise;

$X_{15} = 1$ if the score is from a member in row 3 (Student
Supplies Own Answer); 0 otherwise;

b_{13} to b_{15} = the regression coefficients corresponding to
 X_{13} to X_{15} and

e_2 = the error in prediction for this model

A third model (Model 3) would include only information regarding the
columns, disregarding the particular row:

$$Y = b_0 + b_{17}X_{17} + b_{18}X_{18} + e_3, \quad (4.3)$$

where

b_0 = the Y-intercept for this model (in this case, b_0 will equal
the mean of the third column, Dissimilar Material to the Final),

$X_{17} = 1$ if the score is from a member in column 1 (Very Similar
Material to the Final),

$X_{18} = 1$ if the score is from a member in column 2 (Half Similar,
Half Dissimilar Material to the Final),

b_{17}, b_{18} = are regression coefficients for X_{17}, X_{18} , and

e_3 = the error in prediction for this model.

A fourth model (Model 4), while not necessary when the data is proportional,
is useful in understanding the process:

$$Y = b_0 + b_{13}X_{13} + b_{14}X_{14} + b_{15}X_{15} + b_{17}X_{17} + b_{18}X_{18} + e_4, \quad (4.4)$$

where

b_0 = the Y-intercept for this model.

X_{13} to X_{18} = as defined in Models 2 and 3 (excluding X_{16})

b_{13} to b_{18} = are regression coefficients for X_{13} to X_{18} for this model (excluding b_{16}), and

e_4 = the error in prediction with this model.

Note that X_{12} , X_{16} and X_{19} , corresponding respectively to the score is from

X_{12} = 1 if the score is from a member of the 12th cell; 0 otherwise,

X_{16} = 1 if the score is from a member of row 4; 0 otherwise and

X_{19} = 1 if the score is from a member of column 3; 0 otherwise

are not used in the formulation of the analysis; in that they may be useful for multiple comparisons, they should be included. All of the information regarding a regression solution is given in Table 4.2.

Using the four models (Models 1-4), the portion of the summary table for regression can be used directly. The portions of interest, taken directly from the printout, are reported in Table 4.3.

TABLE 4.3
VALUES FOUND FROM THE REGRESSION ANALYSES

	df	SS	R ²
Model 1			
Attributable to Regression	11	312	.77228
Deviation from Regression	12	92	.22772
Total Sum of Squares	23	404	
Model 2 (Rows)			
Attributable to Regression	3	204	.50495
Model 3 (Columns)			
Attributable to Regression	2	16	.03960
Model 4 (Rows & Columns)			
Attributable to Regression	5	220	.54455

The data from Table 4.3 can be put into a summary table. There are four sources of variation of interest in a two-way analysis of variance: rows, columns, interaction (among rows and columns) and the within group variation. The sum of squares for rows can be found directly from Model 2 as the value for attributable to regression and is 204. The sum of squares for columns can be found directly from Model 3 as the value for attributable to regression and is 16. The sum of squares within is equal to the deviation from regression for Model 1 and is 92. The interaction can be found as the difference in sum of squares between the attributable to regression for Model 1 and the attributable to regression for Model 4: $312 - 220 = 92$.

Also, the degrees of freedom are necessary for each source of variation. The degrees of freedom for rows is the number of rows minus one ($r - 1$); the degrees of freedom for columns is the number of columns minus one ($c - 1$);

the degrees of freedom for interaction is the product of the degrees of freedom for rows and columns $[(r - 1)(c - 1)]$; the total degrees of freedom are the total number of subjects minus one $(N - 1)$.

All of the information for the data in Table 4.3 can be put into a summary table. See Table 4.4.

TABLE 4.4

Summary Table For Data For Method
Of Teaching X Type of Homework

Source of Variation	df	SS	MS	F
Rows	3	204.00	68.00	8.87
Columns	2	16.00	8.00	1.04
Interaction	6	92.00	15.33	2.00
Within	<u>12</u>	<u>92.00</u>	7.67	
Total	23	404.00		

Traditionally, each F value is interpreted as essentially a separate experiment; the wording used is independent. For this model, the row and column effects are independent, as well as both effects being independent of interaction. This is not the same as saying the F values are independent; the F values are correlated due to sharing a common denominator. The numerator sums of squares are independent. The probability of a Type I error for the experiment is not α , but $1 - (1 - \alpha)^3$. For those who prefer to use a family approach and interpret rows separately from columns, etc., they might repeat the error rate as α for a given effect, but not for the experiment. In any event, the data reported here are unequivocal in their interpretation. Clearly there is a significant row effect, and no other.

An Alternative Formulation for Model 1

The linear model given for Model 1 was a model that ignored row and column information, and viewed the problem as one that contained twelve groups. An alternative formulation (Model 1a) is:

$$Y = b_0 + b_{13}X_{13} + b_{14}X_{14} + b_{15}X_{15} + b_{17}X_{17} + b_{18}X_{18} + b_{20}X_{20} + b_{21}X_{21} + b_{22}X_{22} + b_{23}X_{23} + b_{24}X_{24} + b_{25}X_{25} + e_1, \quad (4.5)$$

where

b_0 = the Y-intercept (and equal to the mean of group 12),

X_{13} to X_{18} = as defined previously,

X_{20} = the product of X_{13} times X_{17} ($X_{13} \cdot X_{17}$),

$X_{21} = X_{13} \cdot X_{18}$,

$X_{22} = X_{14} \cdot X_{17}$,

$X_{23} = X_{14} \cdot X_{18}$,

$X_{24} = X_{15} \cdot X_{17}$,

$X_{25} = X_{15} \cdot X_{18}$,

b_{13} to b_{25} = are regression coefficients corresponding to

X_{13} to X_{25} (excluding X_{16} and X_{19}) and

e_1 = the error in prediction with this model.

Both Model 1 and Model 1a can be referred to as full models. It can be seen that X_{13} to X_{15} impart row information, X_{17} to X_{18} impart column information, and X_{20} to X_{25} report the product of row and column information. Thus, interaction is the difference between the full model and the model that does not include the product of row and column information.

Also, reflection on X_{20} to X_{25} would show that $X_{20} = X_1$, $X_{21} = X_2$, $X_{22} = X_4$, $X_{23} = X_5$, $X_{24} = X_7$ and $X_{25} = X_8$; thus, X_1 , X_2 , X_4 , X_5 , X_7 and X_8

can represent X_{20} to X_{25} . See Table 4.5.

Multiple Comparisons for Rows

Most often, solutions to the two-way ANOVA design follow some form similar to equation 4.5 (Model 1a) where the row, column and interaction effects are shown in the same model. The first model (Model 1, equation 4.1) is more useful as a full model if any multiple comparisons are contemplated. Before proceeding, an intermediate model, a reparameterization of Model 1 can be given as

$$Y = b_1X_1 + b_2X_2 + b_3X_3 + b_4X_4 + b_5X_5 + b_6X_6 + b_7X_7 + b_8X_8 + b_9X_9 + b_{10}X_{10} + b_{11}X_{11} + b_{12}X_{12} + e_1 \quad (4.6) \quad (\text{Model 1b}).$$

The difference between Models 1 and 1b is that Model 1 contains a unit vector but not $b_{12}X_{12}$, whereas Model 1b contains $b_{12}X_{12}$ but not a unit vector. Model 1b serves as a good conceptual starting point that, depending upon the program available, may have to be reparameterized.

Suppose there is an interest in comparing row one to row four. A comparison of interest might be:

$$\frac{n_1\bar{Y}_1 + n_2\bar{Y}_2 + n_3\bar{Y}_3}{n_1 + n_2 + n_3} = \frac{n_{10}\bar{Y}_{10} + n_{11}\bar{Y}_{11} + n_{12}\bar{Y}_{12}}{n_{10} + n_{11} + n_{12}} \quad (4.7)$$

In that each $n_i = 2$, equation 4.7 reduces to

$$\bar{Y}_1 + \bar{Y}_2 + \bar{Y}_3 = \bar{Y}_{10} + \bar{Y}_{11} + \bar{Y}_{12}.$$

In terms of the regression coefficients, the comparison can be stated as $b_1 + b_2 + b_3 = b_{10} + b_{11} + b_{12}$. (4.8)

Transforming equation 4.8 by isolating b_1 ,

$$b_1 = b_{10} + b_{11} + b_{12} - b_2 - b_3. \quad (4.9)$$

REGRESSION OUTPUT FOR FULL MODEL WITH ROW, COLUMN AND INTERACTION EFFECTS

Variable No.	Mean	Standard Deviation	Correlation X vs Y	Regression Coefficient	Standard Error of Regression Coefficient	Computed t Value	Beta
13	0.25000	0.44233	0.70360	2.99993	2.76885	1.08346	0.31661
14	0.25000	0.44233	-0.28144	-1.99999	2.76884	-0.72232	-0.21108
15	0.25000	0.44233	-0.28144	-5.00002	2.76884	-1.80582	-0.52770
17	0.33333	0.48154	-0.00000	-7.00007	2.76885	-2.52815	-0.80429
18	0.33333	0.48154	-0.17235	-1.99997	2.76884	-0.72231	-0.22979
1	0.08333	0.28233	0.58791	10.00017	3.91578	2.55381	0.67365
2	0.08333	0.28233	0.14698	-0.99997	3.91576	-0.25537	-0.06736
4	0.08333	0.28233	-0.14698	5.00005	3.91576	1.27690	0.33682
5	0.08333	0.28233	-0.29395	-2.00003	3.91574	-0.51077	-0.13473
7	0.08333	0.28233	-0.07349	9.00010	3.91576	2.29843	0.60629
8	0.08333	0.28233	-0.14698	2.99998	3.92574	0.76613	0.20209

Dependent Y 6.00000 4.19108

Intercept 8.00000

Multiple Correlation 0.87879

Std. Error of Estimate 2.76886

Multiple Correlation Squared 0.77228

One Minus Multiple Correlation Sqd 0.22772

ANALYSIS OF VARIANCE FOR THE REGRESSION

Source of Variation	Degrees of Freedom	Sum of Squares	Mean Squares	F Value
Attributable to Regression	11	311.99951	28.36359	3.69965
Deviation from Regression	12	91.99878	7.66656	
Total	23	403.99829		

Equation 4.9 is a restriction on Model 1b (equation 4.6); the right hand side of equation 4.9 can be substituted for b_1 in Model 1b:

$$Y = (b_{10} + b_{11} + b_{12} - b_2 - b_3)X_1 + b_2X_2 + b_3X_3 + b_4X_4 + b_5X_5 + b_6X_6 + b_7X_7 + b_8X_8 + b_9X_9 + b_{10}X_{10} + b_{11}X_{11} + b_{12}X_{12} + e_5;$$

$$Y = b_2(X_2 - X_1) + b_3(X_3 - X_1) + b_4X_4 + b_5X_5 + b_6X_6 + b_7X_7 + b_8X_8 + b_9X_9 + b_{10}(X_{10} + X_1) + b_{11}(X_{11} + X_1) + b_{12}(X_{12} + X_1) + e_5. \quad (4.10)$$

Five new variables can be defined:

$V_1 = 1$ if from cell 2, -1 if from cell 1, 0 otherwise;

$V_2 = 1$ if from cell 3, -1 if from cell 1, 0 otherwise;

$V_3 = 1$ if from cell 10 or cell 1, 0 otherwise;

$V_4 = 1$ if from cell 11 or cell 1, 0 otherwise; and

$V_5 = 1$ if from cell 12 or cell 1, 0 otherwise.

Equation 4.10 can be transformed to

$$Y = b_2V_1 + b_3V_2 + b_4X_4 + b_5X_5 + b_6X_6 + b_7X_7 + b_8X_8 + b_9X_9 + b_{10}V_3 + b_{11}V_4 + b_{12}V_5 + e_5. \quad (4.11)$$

Equation 4.11 can be used directly by programs that do not require a unit vector. If the available programs require the use of the unit vector (that is, a constant term is included) then any one of the variables could be excluded (say V_5) by setting one of the regression coefficients equal to zero. If $b_{12} = 0$, then equation 4.11 is reparameterized to be

$$Y = b_0 + b_2V_1 + b_3V_2 + b_4X_4 + b_5X_5 + b_6X_6 + b_7X_7 + b_8X_8 + b_9X_9 + b_{10}V_3 + b_{11}V_4 + e_5. \quad (4.12)$$

Using equation 4.12, the following result was obtained:

$$R_{12}^2 = .50496.$$

Now, the following equation is necessary:

$$F = \frac{(R_F^2 - R_R^2) / k}{(1 - R_F^2) / df_w} \quad (4.13)$$

where

R_F^2 = the R^2 value for the full model,

R_R^2 = the R^2 value for the restricted model,

k = the number of restrictions, and

df_w = the degrees of freedom for the within term in the full model,

For the comparison of row 1 to row 4,

$R_F^2 = .77228$, $k = 1$, $R_R^2 = .50496$ and $df_w = 12$.

Then

$$F = \frac{(.77228 - .50496) / 1}{(1 - .77228) / 12} = 14.08682;$$

$\sqrt{F} = t = 3.75324$. This t value could be compared to an appropriate table of critical values for multiple comparison procedures. This brings us back to the original point made in this chapter; the unique concerns of two-way ANOVA's as it relates to multiple comparisons need to be addressed.

Choosing An Appropriate Multiple

Comparison Procedure for Two-Way ANOVA's

Ideally, a researcher will have chosen a priori comparisons (or contrasts). Perhaps the most likely a priori procedure to be of use is Dunn's (1961) test. If a limited number of comparisons are chosen on a reasonable basis, then the two-way ANOVA design can be handled without any new complications.

The more difficult situation arises if a posteriori contrasts are envisioned. For example, the t value found in the previous section might be

tested by Scheffe's test; the t value could be directly compared to $S = \sqrt{(k-1)F}$ where k is the number of groups (12) and F is the tabled F value for $k-1$ and $N-k$ degrees of freedom; $S = \sqrt{11(2.72)}$ at the .05 level; $S = 5.47$. This outcome may seem bewildering; what seems to be a likely significant outcome, testing row 1, with a mean of 11, against row 5, with a mean of 5, fails to achieve significance. Clearly, some other strategy needs to be used for the two-way situation.

Reference can be made back to the idea of a family of comparisons; there are three likely kinds of families; the family of comparisons for rows, the family of comparisons for columns, and the family of comparisons for interactions. By apportioning the error rate to different portions of the experiment, the total experiment can remain as the basis for interpretations. For example, the researcher might report the overall error rate as .05, apportioning .02 to each of the main effects and .01 to the interactions. Then the critical values would be

- a) for rows, $S_r = \sqrt{(r-1) .02 F_{r-1, df_w}}$;
- b) for columns, $S_c = \sqrt{(c-1) .02 F_{c-1, df_w}}$; and
- c) for interactions $S_i = \sqrt{(k-1) .01 F_{k-1, df_w}}$.

These values are respectively $S_r = 3.86$; $S_c = 3.38$; and $S_i = 6.81$. Even using the first critical value, the previously found t value ($t = 3.75324$) fails to achieve significance. Had a slightly different allocation been made (say .025 for both row and columns), eliminating comparisons of cell means, then $S_r = 3.66$ and $S_c = 3.19$. In that case, the previously found t value would be significant. If row 2 and row 3 are compared to row 1 in a manner similar to that described here for comparing row 1 to row 4,

then in both cases, $t = 4.37883$, both showing significance by the earlier critical value of 3.86.

Multiple Comparisons for Columns

Testing for differences among columns can proceed in a manner similar to that shown for rows, columns 2 and 3 could be compared. The corresponding comparison (or contrast) is:

$$\frac{n_2\bar{Y}_2 + n_5\bar{Y}_5 + n_8\bar{Y}_8 + n_{11}\bar{Y}_{11}}{n_2 + n_5 + n_8 + n_{11}} = \frac{n_3\bar{Y}_3 + n_6\bar{Y}_6 + n_9\bar{Y}_9 + n_{12}\bar{Y}_{12}}{n_3 + n_6 + n_9 + n_{12}} \quad (4.14)$$

Since all of the n 's are equal, equation 4.14 reduces to

$$\bar{Y}_2 + \bar{Y}_5 + \bar{Y}_8 + \bar{Y}_{11} = \bar{Y}_3 + \bar{Y}_6 + \bar{Y}_9 + \bar{Y}_{12}.$$

In terms of the regression coefficients, the comparison can be stated

as

$$b_2 + b_5 + b_8 + b_{11} = b_3 + b_6 + b_9 + b_{12}.$$

If b_2 is isolated,

$$b_2 = b_3 + b_6 + b_9 + b_{12} - b_5 - b_8 - b_{11}. \quad (4.15)$$

Equation 4.15 is a restriction on the full model (equation 4.6) corresponding to the null hypothesis in equation 4.14, comparing columns 2 and 3.

Substituting equation 4.15 for b_2 in equation 4.6 yields

$$\begin{aligned} Y &= b_1X_1 + (b_3 + b_6 + b_9 + b_{12} - b_5 - b_8 - b_{11})X_2 + b_3X_3 + b_4X_4 + \\ &\quad b_5X_5 + b_6X_6 + b_7X_7 + b_8X_8 + b_9X_9 + b_{10}X_{10} + b_{11}X_{11} + b_{12}X_{12} + e_6; \\ Y &= b_1X_1 + b_3(X_3+X_2) + b_4X_4 + b_5(X_5-X_2) + b_6(X_6+X_2) + b_7X_7 + b_8(X_8-X_2) + \\ &\quad b_9(X_9+X_2) + b_{10}X_{10} + b_{11}(X_{11}-X_2) + b_{12}(X_{12}+X_2) + e_6. \end{aligned} \quad (4.16)$$

Seven new variables can be defined

$$V_6 = 1 \text{ if cell 3 or cell 2, 0 otherwise;}$$

$$V_7 = 1 \text{ if cell 5, -1 if cell 2, 0 otherwise;}$$

$$V_8 = 1 \text{ if cell 6 or cell 2, 0 otherwise;}$$

$$V_9 = 1 \text{ if cell 8, -1 if cell 2, 0 otherwise;}$$

$$V_{10} = 1 \text{ if cell 9 or cell 2, 0 otherwise;}$$

$$V_{11} = 1 \text{ if cell 11, -1 if cell 2, 0 otherwise; and}$$

$$V_{12} = 1 \text{ if cell 12 or cell 2, 0 otherwise.}$$

Equation 4.16 can be transformed to

$$Y = b_1X_1 + b_3V_6 + b_4X_4 + b_5V_7 + b_6V_8 + b_7X_7 + b_8V_9 + b_9V_{10} + b_{10}X_{10} + b_{11}V_{11} + b_{12}V_{12} + e_6.$$

Reparameterizing by setting $b_{12} = 0$ yields

$$Y = b_0 + b_1X_1 + b_3V_6 + b_4X_4 + b_5V_7 + b_6V_8 + b_7X_7 + b_8V_9 + b_9V_{10} + b_{10}X_{10} + b_{11}V_{11} + e_6. \quad (4.17)$$

Using equation 4.17 as the restricted model and equation 4.6 as the full model yields

$$F = \frac{(.77228 - .73267)/1}{(1 - .77228)/12} = 2.08730.$$

$t = \sqrt{F} = 1.44475$. This t value is of course non-significant. Other column comparisons could be achieved in a similar manner.

Comparisons of Cell Means

Comparisons regarding cell means can be fairly simply conceived by viewing the analysis as a one-way layout; comparisons can be achieved into different approaches (but which yield identical t values). The more

important issue is: Are the number of comparisons of experimental interest, decided on prior to the analysis, sufficiently small to justify the use of Dunn's test? If so, then a much more powerful test can be performed.

Ideally, the number of cell comparisons can be integrated into the row and column comparisons so that the interpretation can be made in an unambiguous manner. However, the error rate can be partitioned into three distinct units and a separate portion for comparisons among cell means can be used.

If any comparison (including complex contrasts) is envisioned, the use of full and restricted models can be followed as outlined for both rows and columns earlier in this chapter. For example, suppose there is interest in comparing $\bar{Y}_1$ to $\bar{Y}_{12}$. The restriction to accomplish this comparison is

$b_1 = b_{12}$ when using Model 1b (Equation 4.6) as the full model. Then

$$Y = b_{12}X_1 + b_2X_2 + b_3X_3 + b_4X_4 + b_5X_5 + b_6X_6 + b_7X_7 + b_8X_8 \\ + b_9X_9 + b_{10}X_{10} + b_{11}X_{11} + b_{12}X_{12} + e_7,$$

or

$$Y = b_{12}(X_1 + X_{12}) + b_2X_2 + b_3X_3 + b_4X_4 + b_5X_5 + b_6X_6 + b_7X_7 \\ + b_8X_8 + b_9X_9 + b_{10}X_{10} + b_{11}X_{11} + e_7;$$

$$Y = b_{12}V_{13} + b_2X_2 + b_3X_3 + b_4X_4 + b_5X_5 + b_6X_6 + b_7X_7 + b_8X_8 \\ + b_9X_9 + b_{10}X_{10} + b_{11}X_{11} + e_7$$

where

$$V_{13} = 1 \text{ if cell 1 or cell 12, 0 otherwise } (V_{13} \text{ happens to equal } V_5).$$

Then

$$F = \frac{(R_F^2 - R_R^2) / 1}{(1 - R_F^2) / df_W} = \frac{(.77228 - .68317) / 1}{(1 - .77228) / 12}$$

$F = 4.69576$; therefore $t = \sqrt{F} = 2.16697$.

In the specific case wherein two cell means are being compared, a simplified process will yield the same t value; if either b_1 or b_{12} is set equal to zero and if a unit vector is available (as is true of most general purpose multiple linear regression computer programs), then the remaining b value is tested by the computed t value as a test of the partial regression coefficient. As an example, $b_{12} = 0$, so that a model identical to equation 4.6 is employed. See Table 4.6.

REGRESSION SUMMARY FOR FULL MODEL USING ALL VARIABLES

Variable No.	Mean	Standard Deviation	Correlation X vs Y	Regression Coefficient	Standard Error of Regression Coefficient	Computed t Value	Beta
1	0.08333	0.28233	0.58791	6.00002	2.76889	2.16693	0.40419
2	0.08333	0.28233	0.14698	0.00001	2.76889	0.00000	0.00000
3	0.08333	0.28233	0.36744	3.00002	2.76889	1.08347	0.20209
4	0.08333	0.28233	-0.14698	-4.00000	2.76889	-1.44462	-0.26946
5	0.08333	0.28233	-0.29395	-5.99999	2.76888	-2.16693	-0.40418
6	0.08333	0.28233	0.00000	-1.99998	2.76888	-0.72231	-0.13473
7	0.08333	0.28233	-0.07349	-2.99998	2.76888	-1.08346	-0.20209
8	0.08333	0.28233	-0.14698	-3.99998	2.76888	-1.44462	-0.26946
9	0.08333	0.28233	-0.22047	-4.99999	2.76888	-1.80578	-0.33682
10	0.08333	0.28233	-0.36744	-6.99999	2.76888	-2.52810	-0.47155
11	0.08333	0.28233	0.00000	-1.99998	2.76888	-0.72231	-0.13473
Dependent Y	6.00000	4.19108					
Intercept		7.99999					
Multiple Correlation		0.87879					
Std. Error of Estimate		2.76886					
Multiple Correlation Squared			0.77228				
One Minus Multiple Correlation Sqd			0.22772				

45

ANALYSIS OF VARIANCE FOR THE REGRESSION

Source of Variation	Degrees of Freedom	Sum of Squares	Mean squares	F Value
Attributable to Regression	11	311.99927	28.36356	3.69963
Deviation From Regression	12	91.99902	7.66658	
Total	23	403.99829		

As can be seen from Table 4.6, the computed t value for X_1 (variable 1) is 2.16693, within rounding error of the value found earlier through the use of the full and restricted models; this computed t value can be directly compared to the appropriate table, depending upon the multiple comparison test used. The remaining t values are available in case there were interest in testing any of cells 2 through 11 to cell 12.

Suppose there was an interest in testing $\bar{Y}_3$ to $\bar{Y}_6$. This could be accomplished by either of the following models:

$$Y = b_0 + b_1 X_1 + b_2 X_2 + b_3 X_3 + b_4 X_4 + b_5 X_5 + b_7 X_7 + b_8 X_8 + b_9 X_9 + b_{10} X_{10} + b_{11} X_{11} + b_{12} X_{12} + e_1 \quad (4.18) \quad \text{or}$$

$$Y = b_0 + b_1 X_1 + b_2 X_2 + b_4 X_4 + b_5 X_5 + b_6 X_6 + b_7 X_7 + b_8 X_8 + b_{10} X_{10} + b_{11} X_{11} + b_{12} X_{12} + e_1. \quad (4.19)$$

In equation 4.18, $b_6 X_6$ has been omitted, and hence, $b_0 = \bar{Y}_6$ and $b_3 = \bar{Y}_3 - \bar{Y}_6$; the test of the regression coefficient (the computed t value) is the test of significance for the difference between $\bar{Y}_3$ and $\bar{Y}_6$. Equation 4.19 simply reverses the roles of variables 3 and 6 respectively for equations 4.18 and 4.19. Therefore, $b_3 = -b_6$ and $t_3 = -t_6$ from equation 4.18 to 4.19. Of course, the use of full and restricted models would have yielded the same result.

Clearly, the concerns of multiple comparisons when addressed to the two-way layout are complex. Most questions of substance require the use of full and restricted models, which require the researcher to state comparisons in terms of the regression coefficients. Comparisons among cell means can be accomplished either through the use of full and restricted models, or through the use of the unit vector together with the omission of one of the groups of interest as a predictor.

MULTIPLE COMPARISONS FOR DISPROPORTIONATE TWO-WAY LAYOUTS

Many readers are aware of some of the complexities involved with two-way disproportionate analyses of variance. Several different points of view have been made; Jennings (1967) proposed a solution through the use of full and restricted models. Williams (1972) showed that his unadjusted main effects solution agreed with Jennings. Overall and Spiegel (1969) described three solutions: 1) a fitting constants solution (also described by Anderson and Bancroft, 1952); 2) a hierarchical solution (Cohen, 1968); and 3) a full rank model solution, more recently called the "standard" solution for Timm and Carlson (1975) and apparently accepted as such by Overall, Spiegel and Cohen (1975) despite Overall and Spiegel's previously showing that the full rank model failed to provide an additive solution given that proportionate, but unequal cell frequencies occurred. Speed and Hocking (1976) described various solutions in terms of the hypotheses tested. Williams (1977b) showed that the solution proposed by Jennings and arithmetically identical to the unadjusted main effects solution was also properly a full rank model solution.

Often the researcher who is more concerned with getting any solution, rather than worry about what might seem to be a complex argument among scholars, is content to use a solution available through the local computing facility. Given the difficulties involved with interpreting a disproportionate two-way analysis of variance, and given the complexities involved with multiple comparisons in a two-way layout described in Chapter Four, the reader might fairly ask, "How much

more complex are multiple comparisons for the disproportionate case?" Luckily, the answer is that, for a given set of hypotheses (or comparisons) multiple comparisons for the disproportionate case are not too much more complex than they are for the equal cell frequency case.

An Example

Table 5.1 contains data for the two-way disproportionate case.

	Column 1	Column 2	Column 3
Row 1	8 8 6 6 4 4	1 1 1 1	6 2 2 6
Row 2	10 10	7 7 5 5 4 4 4 4 3 3	10 10 9 9 7 7 5 5 4 4

The systematic nature of the data within any cell is by no means coincidental; the data first reported in Williams (1972) which had an $N = 18$ has been doubled to an $N = 36$ by recording each score twice.

First, several variables can be defined:

X_1 = 1 if in cell 1 (row 1, column 1), 0 otherwise;

X_2 = 1 if in cell 2 (row 1, column 2), 0 otherwise;

X_3 = 1 if in cell 3 (row 1, column 3), 0 otherwise;

X_4 = 1 if in cell 4 (row 2, column 1), 0 otherwise;

X_5 = 1 if in cell 5 (row 2, column 2), 0 otherwise;

X_6 = 1 if in cell 6 (row 2, column 3), 0 otherwise;

X_7 = 1 if in row 1, 0 if in row 2;

X_8 = 1 if in column 1, 0 otherwise;

X_9 = 1 if in column 2, 0 otherwise;

X_{10} = $X_7 \cdot X_8$; and

X_{11} = $X_7 \cdot X_9$.

Two different full models could be given:

$$Y = b_0 + b_1X_1 + b_2X_2 + b_3X_3 + b_4X_4 + b_5X_5 + e_1; \quad (5.1)$$

or

$$Y = b_0 + b_7X_7 + b_8X_8 + b_9X_9 + b_{10}X_{10} + b_{11}X_{11} + e_1. \quad (5.2)$$

A row model could be given as

$$Y = b_0 + b_7X_7 + e_2. \quad (5.3)$$

The column model is

$$Y = b_0 + b_8X_8 + b_9X_9 + e_3. \quad (5.4)$$

The combined row and column model is

$$Y = b_0 + b_7X_7 + b_8X_8 + b_9X_9 + e_4. \quad (5.5)$$

If there is interest in reproducing the sum of squares, they could be found as

$$SS_{\text{ROWS}} = SS_T (R_{5.3}^2); \quad (5.6)$$

$$SS_{\text{COLUMNS}} = SS_T (R_{5.4}^2); \quad (5.7)$$

$$SS_{\text{ROWS} \times \text{COLUMNS}} = SS_T (R_{5.2}^2 - R_{5.5}^2); \quad (5.8)$$

and

$$SS_{\text{WITHIN}} = SS (1 - R_{5.2}^2). \quad (5.9)$$

Equations 5.6 to 5.9 will produce a solution that has been called the unadjusted main effects solution that is identical for each effect to the solution proposed by Jennings (see Williams, 1977b). If the full-rank model solution preferred by Timm and Carlson (1975) is desired, several changes would be necessary; these changes are also described in Williams (1977b).

The interest here, however, is in multiple comparisons; the whole issue of which analysis of variance solution should be used might be avoided and the comparisons of interest pursued. However, Table 5.2 contains results from the use of equations 5.1 to 5.5.

TABLE 5.2

SUMMARY INFORMATION FROM THE USE OF EQUATIONS 5.1-5.5

$$SS_T = 264.$$

Full Model(s) (Both Equations 5.1 and 5.2)

$$R_{5.2}^2 = .61212 \quad SS_{\text{FULL}} = 161.60.$$

$$\text{Rows: } R_{5.3}^2 = .15427 \quad SS_{\text{ROWS}} = 40.73.$$

$$\text{Columns: } R_{5.4}^2 = .28355 \quad SS_{\text{COLUMNS}} = 74.86.$$

$$\text{Rows, Columns: } R_{5.5}^2 = .60797; \quad SS_{\text{ROWS} \times \text{COLUMNS}} = 1.10.$$

$$\text{Within: } SS_{\text{WITHIN}} = 102.40.$$

Testing for Comparisons Between Rows

Since there are only two rows, clearly the only comparison will involve row 1 and row 2. A likely comparison of interest is

$$\frac{n_1 \bar{Y}_1 + n_2 \bar{Y}_2 + n_3 \bar{Y}_3}{n_1 + n_2 + n_3} = \frac{n_4 \bar{Y}_4 + n_5 \bar{Y}_5 + n_6 \bar{Y}_6}{n_4 + n_5 + n_6} \quad (5.10)$$

In terms of the regression coefficients, and assuming a reparameterized full model, eliminating the constant,

$$Y = b_1 X_1 + b_2 X_2 + b_3 X_3 + b_4 X_4 + b_5 X_5 + b_6 X_6 + e_1, \quad (5.11)$$

equation 5.10 can be stated as

$$\frac{n_1 b_1 + n_2 b_2 + n_3 b_3}{n_1 + n_2 + n_3} = \frac{n_4 b_4 + n_5 b_5 + n_6 b_6}{n_4 + n_5 + n_6} \quad (5.12)$$

Solving equation 5.12 for b_1 yields

$$b_1 = \frac{(n_4 b_4 + n_5 b_5 + n_6 b_6)(n_1 + n_2 + n_3) - n_2 b_2 - n_3 b_3}{n_1(n_4 + n_5 + n_6)} \quad (5.13)$$

Substituting this value back into equation 5.11 yields

$$Y = \left[\frac{n_4 b_4 + n_5 b_5 + n_6 b_6}{n_1} \frac{(n_1 + n_2 + n_3)}{(n_4 + n_5 + n_6)} - \frac{n_2 b_2}{n_1} - \frac{n_3 b_3}{n_1} \right] X_1 + b_2 X_2 + b_3 X_3 + b_4 X_4 + b_5 X_5 + b_6 X_6 + e_5.$$

After some algebraic unravelling, the previous equation yields

$$Y = b_2 \left(X_2 - \frac{n_2}{n_1} X_1 \right) + b_3 \left(X_3 - \frac{n_3}{n_1} X_1 \right) + b_4 \left[X_4 + \frac{n_4}{n_1} \frac{(n_1 + n_2 + n_3)}{(n_4 + n_5 + n_6)} X_1 \right] + b_5 \left[X_5 + \frac{n_5}{n_1} \frac{(n_1 + n_2 + n_3)}{(n_4 + n_5 + n_6)} X_1 \right] + b_6 \left[X_6 + \frac{n_6}{n_1} \frac{(n_1 + n_2 + n_3)}{(n_4 + n_5 + n_6)} X_1 \right] + e_5 \quad (5.14).$$

Since $n_1 = 6$, $n_2 = 4$, $n_3 = 4$, $n_4 = 2$, $n_5 = 10$ and $n_6 = 10$, five new

variables can be defined to correspond to the variables in equation 5.14:

$V_1 = 1$ if from cell 2, $-.6667$ if from cell 1, 0 otherwise;

$V_2 = 1$ if from cell 3, $-.6667$ if from cell 1, 0 otherwise;

$V_3 = 1$ if from cell 4, .2121 if from cell 1, 0 otherwise;

$V_4 = 1$ if from cell 5, 1.0606 if from cell 1, 0 otherwise; and

$V_5 = 1$ if from cell 6, 1.0606 if from cell 1, 0 otherwise.

An equation can be formed with the new variables:

$$Y = b_1V_1 + b_2V_2 + b_3V_3 + b_4V_4 + b_5V_5 + e_5.$$

Any one of b_1 through b_5 can be set equal to zero, and a reparameterization will be made such that the constant is reintroduced. Letting $b_5 = 0$,

$$Y = b_0 + b_1V_1 + b_2V_2 + b_3V_3 + b_4V_4 + e_5. \quad (5.15)$$

The use of equation 5.15 yields $R^2_{5.15} = .45785$, so that

$$t = \sqrt{F} = \sqrt{\frac{(.61212 - .45785)/1}{(.38788)/24}} = \sqrt{\frac{.15427}{.01616}} = 9.54641,$$

$t = 3.08973$. These results are identical to the usual F test of the row main effect; given that there are only two rows, this was to be expected (though perhaps not after all that manipulation!).

A Different Hypothesis

Suppose that, in equation 5.12, no concern is placed upon the varying cell frequencies, and a hypothesis is formed

$$\frac{Y_1 + Y_2 + Y_3}{3} = \frac{Y_4 + Y_5 + Y_6}{3}. \quad (5.16)$$

The corresponding equation, in relation to the regression coefficients, is

$$b_1 + b_2 + b_3 = b_4 + b_5 + b_6 \quad \text{or} \quad b_1 = b_4 + b_5 + b_6 - b_2 - b_3.$$

Substituting this value for b_1 in equation 5.11 yields

$$Y = (b_4 + b_5 + b_6 - b_2 - b_3)X_1 + b_2X_2 + b_3X_3 + b_4X_4 + b_5X_5 + b_6X_6 + e_6;$$

$$Y = b_2(X_2 - X_1) + b_3(X_3 - X_1) + b_4(X_4 + X_1) + b_5(X_5 + X_1) + b_6(X_6 + X_1) + e_6. \quad (5.17)$$

After defining new variables to correspond to each of the portions in equation 5.17 in a manner similar to the previous solution,

$$F = \frac{(.61212 - .30070)/1}{(.38788)/30} = 26.12584; t = 5.11134.$$

Curiously, this result is identical to the row effect for using the full rank model solution described by Timm and Carlson (1975). Here is the curiosity; if a restriction is placed on the full model taking into cognizance the unequal cell frequencies the outcome is in agreement with the unadjusted main effects solution. However, if the difference in cell frequencies is disregarded, the solution is identical to the full rank model solution as described by Timm and Carlson. One could argue that the solution of Timm and Carlson, commonly referred to as the full rank model solution, is actually the solution that truly leaves the data "unadjusted" for unequal cell frequencies. If presented with the question, "Which hypothesis, 5.10 or 5.16, more fairly addresses the difference between row 1 and row 2, given that the cell frequencies are disproportionate?", it is my guess that most researchers would prefer 5.10 (i.e., taking into account the disproportionate cell frequencies). However, if researchers are aware of the exact hypothesis being tested, they may opt for whatever solution best answers their specific research questions.

Comparisons Among Columns

Comparisons among columns can be accomplished in a manner similar to that for rows; suppose there is an interest in comparing columns 2, separately to columns 1 and 3. The hypothesis for comparing column 1 to column 2 is given by

$$\frac{n_1\bar{Y}_1 + n_4\bar{Y}_4}{n_1 + n_4} = \frac{n_2\bar{Y}_2 + n_5\bar{Y}_5}{n_2 + n_5}. \quad (5.18)$$

In terms of the regression coefficients, the hypothesis can be stated as

$$\frac{n_1 b_1 + n_4 b_4}{n_1 + n_4} = \frac{n_2 b_2 + n_5 b_5}{n_2 + n_5}$$

Solving for b_1 ,

$$b_1 = \frac{(n_2 b_2 + n_5 b_5)}{n_1} \frac{(n_1 + n_4)}{(n_2 + n_5)} - \frac{n_4}{n_1} b_4$$

Substituting this value for b_1 into equation 5.11 yields

$$Y = \left[\frac{(n_2 b_2 + n_5 b_5)}{n_1} \frac{(n_1 + n_4)}{(n_2 + n_5)} - \frac{n_4}{n_1} b_4 \right] X_1 + b_2 X_2 + b_3 X_3 + b_4 X_4 + b_5 X_5 + b_6 X_6 + e_7.$$

After some algebraic manipulations,

$$Y = b_2 \left[X_2 + \frac{n_2 (n_1 + n_4)}{n_1 (n_2 + n_5)} X_1 \right] + b_3 X_3 + b_4 \left[X_4 - \frac{n_4}{n_1} X_1 \right] + b_5 \left[X_5 + \frac{n_5 (n_1 + n_4)}{n_1 (n_2 + n_5)} X_1 \right] + b_6 X_6 + e_7. \quad (5.19)$$

Three new variables can be defined

$V_7 = 1$ if X_2 , $8/21$ if X_1 , 0 otherwise;

$V_8 = 1$ if X_4 , $-1/3$ if X_1 , 0 otherwise; and

$V_9 = 1$ if X_5 , $20/21$ if X_1 , 0 otherwise.

An equation can be formed with the new variables:

$$Y = b_7 V_7 + b_8 V_8 + b_9 V_9 + b_3 X_3 + b_6 X_6 + e_7.$$

Any one of the b 's in the previous equation can be set equal to zero;

letting $b_6 = 0$,

$$Y = b_0 + b_7 V_7 + b_8 V_8 + b_9 V_9 + b_3 X_3 + e_7.$$

This yields $F = \frac{(.61212 - .38544)/1}{.38788/24} = 14.02584$; $t = 3.74511$.

In a similar manner, testing column 2 against column 3 yields

$$F = \frac{(.61212 - .45433)/1}{.38788/30} = 12.21285; t = 3.49469. \text{ The testing of sig-}$$

nificance is postponed until after the next section.

Individual Comparison of Cell Means

Individual comparisons of cell means can be accomplished, as far as the computations are concerned, by using any $g-1$ of the variables indicating cell membership (X_1 through X_6). The "left out group" will be compared to all remaining groups. The process can be continued until all comparisons of interest are found. One such set of comparisons, made with cell 6 compared to all other cells, using equation 5.1, is shown in Table 5.3.

TABLE 5.3

REGRESSION OUTPUT COMPARING CELL 6 TO ALL OTHER CELLS

Variable No.	Mean	Standard Deviation	Correlation X VS Y	Regression Coefficient	Std. Error Of Reg. Coef.	Computed T Value	Beta
1	0.16667	0.37796	0.11010	-1.00000	0.95406	-1.04815	-0.13762
2	0.11111	0.31873	-0.56575	-5.99999	1.09301	-5.48942	-0.69631
3	0.11111	0.31873	-0.17408	-2.99999	1.09301	-2.74471	-0.34815
4	0.05556	0.23231	0.41796	2.99999	1.43108	2.09631	0.25376
5	0.27778	0.45426	-0.16794	-2.39999	0.82624	-2.90473	-0.39696
Dependent Y	5.33333	2.74643					
INTERCEPT		6.99999					
MULTIPLE CORRELATION		0.78238					
STD. ERROR OF ESTIMATE		1.84752					
MULTIPLE CORRELATION SQUARED			0.61212				
ONE MINUS MULTIPLE CORRELATION SQD			0.38788				

ANALYSIS OF VARIANCE FOR THE REGRESSION

Source of Variation	Degrees Of Freedom	Sum Of Squares	Mean Squares	F Value
Attributable To Regression	5	161.59964	32.31992	9.46869
Deviation From Regression	30	102.40036	3.41335	
Total	35	264.00000		

The use of the "Computed t value" column allows testing $\bar{Y}_6$ against any of the other cells of interest. Comparing $\bar{Y}_3$ to $\bar{Y}_6$ yields $t = 2.74471$. If some other comparison were of interest, say $\bar{Y}_2$ to $\bar{Y}_5$, then either b_2 or b_5 can be set equal to 0 with all other regression coefficients b_1 through b_6 non-zero.

Tests of Significance

The concerns for tests of significance regarding multiple comparisons for two-way analysis of variance are the same whether the data is proportional or disproportional. From an experimental viewpoint, one reasonable approach is to construct a priori a small number of hypotheses likely to be of interest. If the number of tests is held down, then Dunn's test can be a most profitable approach.

Suppose that the hypotheses of interest are those described in this chapter:

Row 1 to Row 2, $t = 3.45549$;

Column 1 to Column 2, $t = 4.16587$;

Column 2 to Column 3, $t = 3.49469$; and

cell 3 to cell 6, $t = 2.74471$ (disregarding the sign).

Using Dunn's test, m , the number of comparisons, is 4, and v , the degrees of freedom for the within term, is 30. With $\alpha = .05$, the critical t ratio is 2.66 so that all four comparisons would be significant at the .05 level. Had the number of comparisons been increased to 6, then the comparison of cell 3 to cell 6 would become nonsignificant. Had the number of comparison been increased to 10, then the comparison of row 1 to row 2 would become nonsignificant. While the table in appendix does not show critical values for $10 < m < 20$, the comparison of columns 2 to column 3

would be nonsignificant if there were 12 or more comparisons. Thus, it can be seen that there is some virtue in limiting the comparisons to those of direct experimental interest.

MULTIPLE COMPARISONS IN HIGHER DIMENSIONAL DESIGNS

Clearly, N-way Designs have complexities of their own; the considerations of multiple comparisons heightens these considerations. On the other hand, the issues regarding multiple comparisons reflect the same concerns that were evident with the two-way design. Should the researcher employ a family error rate? Or should an error rate be used that encompasses the entire experiment? The new complication is that as the dimension of the design increases, the number of families are $2^N - 1$. The problem of deciding upon whether to use a family error rate (and of course reporting the type of error rate used to the reader) is compounded by two competing forces; on the one hand, if an overall experiment error is used, significance becomes difficult to obtain; on the other hand, the use of a family error rate, by increasing the number of families, increases the likelihood of finding a significant outcome.

Part of the problem, and a hint at a useful solution, is that researchers, particularly those who are more inexperienced, often run statistical tests indiscriminately without any overall schema other than the proverbial pragmatist's, "Let's run it through and see what we've got". Clearly, such an approach is inappropriate for hypothesis testing, though it might be useful for hypothesis generation. One useful approach is to construct, on an a priori basis, a limited number of hypotheses that bear directly on the *raison d'etre* for the experiment. If additional, essentially untested, results are to be reported as interesting rela-

tionships that might bear future investigation, then such results should be labeled for what they really are: conjectural possibilities that just might lead to some interesting hypotheses testing on a future data set.

An Example - A Three-Way Layout

Suppose a $2 \times 3 \times 2$ design is used; Figure 6.1 shows the structure for this design.

		C_1		
		B_1	B_2	B_3
A_1	$A_1B_1C_1$ Cell 1	$A_1B_2C_1$ Cell 2	$A_1B_3C_1$ Cell 3	
	$A_2B_1C_1$ Cell 4	$A_2B_2C_1$ Cell 5	$A_2B_3C_1$ Cell 6	
		C_2		
		B_1	B_2	B_3
A_1	$A_1B_1C_2$ Cell 7	$A_1B_2C_2$ Cell 8	$A_1B_3C_2$ Cell 9	
	$A_2B_1C_2$ Cell 10	$A_2B_2C_2$ Cell 11	$A_2B_3C_2$ Cell 12	

Figure 6.1 Structure for a $2 \times 3 \times 2$ Design

Rather than focus on the usual three-way solution (but see Winer, 1971; and Lindquist, 1953; also, a four-factor solution is shown in a regression format by Williams, 1974b) interest is placed upon possible multiple comparisons (more correctly, contrasts). To make the problem somewhat more pragmatic, suppose that the criterion scores represent salaries, the

A category represents males (A_1) and females (A_2); the B category represents respectively assistant professor (B_1), associate professor (B_2), and professor (B_3); and the C category represents a measure of locals (C_1) and non-locals (C_2) where a local is considered to be a person who has received any of their academic degrees from an in-state institution. While many comparisons might conceivably be of interest, suppose the following are to be tested:

- 1) male locals compared to female locals

$$\text{Restriction: } \frac{n_1b_1 + n_2b_2 + n_3b_3}{n_1 + n_2 + n_3} = \frac{n_4b_4 + n_5b_5 + n_6b_6}{n_4 + n_5 + n_6});$$

- 2) male locals compared to male non-locals

$$\text{Restriction: } \frac{n_1b_1 + n_2b_2 + n_3b_3}{n_1 + n_2 + n_3} = \frac{n_7b_7 + n_8b_8 + n_9b_9}{n_7 + n_8 + n_9});$$

- 3) female professors compared to male professors

$$\text{Restriction: } \frac{n_6b_6 + n_{12}b_{12}}{n_6 + n_{12}} = \frac{n_3b_3 + n_9b_9}{n_3 + n_9});$$

- 4) female professors compared to male associate professors

$$\text{Restriction: } \frac{n_6b_6 + n_{12}b_{12}}{n_6 + n_{12}} = \frac{n_2b_2 + n_8b_8}{n_2 + n_8}); \text{ and}$$

- 5) male local professors compared to non-local professors

$$\text{Restriction: } b_3 = \frac{n_9b_9 + n_{12}b_{12}}{n_9 + n_{12}}).$$

If these are the precise comparisons of interest, a most effective test of significance that retains a per experiment error rate is Dunn's (1961) test. In that only 5 comparisons are involved, Dunn's test would be quite sensitive to any differences that exist. Note also that data such as would be likely to occur at any college or university would be

highly disproportionate. While this disproportionality reduces the efficiency of the tests (i.e., somewhat larger samples are necessary than if all cell frequencies are equal) it does not invalidate the tests in any way. In fact, the issue of disproportionality and choosing an appropriate solution is not particularly relevant if the interest is focused on the five research questions given earlier rather than focusing on the main effects and interactions.

Higher Dimensional Designs

Four-way and higher dimensional designs can of course be accommodated to a regression format. If the questions of a major research interest can be limited to a relative few then no major problem is encountered by increasing the dimensionality. If on the other hand, "traditional" analyses of variances are thought necessary before conducting any further tests, then a considerable amount of power is lost.

While it might seem argumentative to do so, it seems fair to ask each researcher to state as clearly as possible precisely what is intended in a given research application. Seldom, it would seem, would a researcher be interested in the specific hypotheses involved in a testing a main effect. If the research question can be stated and written in relationship to the regression coefficients such as was done for the example earlier in this chapter, then a more relevant data analysis can take place. Care must be taken to use an appropriate test of significance; some researchers are prone to run each test against the overall F distribution; such an approach is equivalent to using a per hypothesis error rate.

Missing Cells

Insofar as making multiple comparisons are concerned, missing cells do not present any major computational difficulties. Obviously, if a cell is missing, the corresponding b's can neither be included in any formulation of the full model, nor in any restrictions. Consider that cells 1 and 6 are missing for Figure 6.1.

The full model can be stated either as

$$Y = b_2X_2 + b_3X_3 + b_4X_4 + b_5X_5 + b_7X_7 + b_8X_8 + b_9X_9 + b_{10}X_{10} + b_{11}X_{11} + b_{12}X_{12} + e_1. \quad (6.1)$$

or

$$Y = b_0 + b_2X_2 + b_3X_3 + b_4X_4 + b_5X_5 + b_7X_7 + b_8X_8 + b_9X_9 + b_{10}X_{10} + b_{11}X_{11} + e_1. \quad (6.2)$$

In that either b_1 or b_6 (or both) are involved in restrictions 1 through 4 given earlier, only restriction 5 is unchanged.

Restrictions 1 through 4 become:

$$1'. \quad \text{Restriction: } \frac{n_2b_2 + n_3b_3}{n_2 + n_3} = \frac{n_4b_4 + n_5b_5}{n_4 + n_5};$$

$$2'. \quad \text{Restriction: } \frac{n_2b_2 + n_3b_3}{n_2 + n_3} = \frac{n_7b_7 + n_8b_8 + n_9b_9}{n_7 + n_8 + n_9};$$

$$3'. \quad \text{Restriction: } b_{12} = \frac{n_3b_3 + n_9b_9}{n_3 + n_9}; \text{ and}$$

$$4'. \quad \text{Restriction: } b_{12} = \frac{n_2b_2 + n_8b_8}{n_2 + n_8}.$$

Interpretation of disproportionate cell frequency problems are sometimes considerably more complex, particularly with missing cells. Why are the cells disproportionate? Answering that question is much more an art than a science. Certainly the use of data will be helpful, but the interpre-

tation will undoubtedly rest on the insights (and perhaps prejudices) of the investigator.

PROBLEMS LESS AMENABLE TO A DIRECT REGRESSION SOLUTION

While fixed effects designs or random effects designs yield solutions directly for multiple comparisons, when the design becomes mixed with both fixed and random effects resulting in two or more error terms, solutions to multiple comparison problems will usually prove to be evasive (at least to this writer). The difficulty is that isolating the correct error term, together with predictor variables such that useful hypotheses can be tested, can be either difficult or seemingly impossible. In these circumstances, the more traditional computational procedures might prove to be more effective. To show the difficulty with the mixed design, an example is given.

An Example

The following problem is taken from Williams (1974a).

A researcher may have an interest in the differential effect of two or more methods of instruction over time; thus, measures can be taken at specified intervals on the several instructional methods. From the point of view of the experiment, a repeated measures design can be conceptualized as a treatments X subjects design repeated for each instructional method.*

To make the example more specific, suppose a researcher is interested in investigating the differences among three approaches to a human relations experience.**

*This design is called a Type I design by Lindquist (1953).

**By human relations experience is meant the meeting of a group of people that has variously been called the T-Group (training group), the encounter group, or some similar name.

The three different approaches toward the human relations groups selected are (1) structured sessions in which the group participates toward concrete problem solving, (2) an unstructured group, where the group decides upon its own goals, and (3) a group designed to allow the individual to focus on his personal problems with the interest being to help solve these problems. Five groups with 7-9 individuals in each group are assigned to each of the three human relation group situations; i.e., there are five separate groups for each treatment situation. Each group is to have a two hour session once a week for four weeks.

While there are several things that might be of interest to measure, the researcher is interested specifically in the amount of aggression exhibited in the group setting.

Videotapes are made of all sessions, and a group of five experts independently judge the amount of aggression expressed during the sessions on a continuum from 0 to 10, where 0 represents no aggression and 10 represents an extreme amount of aggression. The measurements are made with the group as the unit of analysis. The score to be used is the mean of the five ratings. Results are as follows:

TABLE 7.1

GROUP SCORES FROM THREE HUMAN RELATION GROUP METHODS FOR FIVE SESSIONS
(ARTIFICIAL DATA)

Methods 1	(Structured Groups)			
Group	Session 1	Session 2	Session 3	Session 4
1	3.2	3.4	3.2	2.8
2	4.6	4.0	3.8	3.4
3	5.0	3.8	5.0	3.2
4	2.0	2.0	2.4	1.6
5	3.6	3.2	3.4	3.0

Method 2 (Unstructured Groups)

Group	Session 1	Session 2	Session 3	Session 4
6	6.2	5.8	6.8	5.0
7	3.6	3.8	7.2	5.4
8	4.0	6.8	7.8	6.0
9	5.0	5.8	6.0	5.0
10	4.8	5.0	6.4	5.8

Method 3 (Personal Problems)

Group	Session 1	Session 2	Session 3	Session 4
11	7.4	7.6	6.8	5.2
12	6.4	6.4	5.6	4.0
13	7.0	6.6	6.6	6.0
14	5.8	7.4	5.0	4.8
15	6.4	5.2	4.0	3.6

To analyze the data in Table 7.1, it is first useful to define several variables:

Y = the criterion variable,

$X_1 - X_{15}$ are binary variables that identify each group,

X_{16} = 1 if the score is from a group in the structured treatments; 0 otherwise;

X_{17} = 1 if the score is from a group in the unstructured treatment; 0 otherwise,

X_{18} = 1 if the score is from a group in the problems treatment.

X_{19} = 1 if the score is from Session 1; 0 otherwise,

X_{20} = 1 if the score is from Session 2; 0 otherwise,

X_{21} = 1 if the score is from Session 3; 0 otherwise,

X_{22} = 1 if the score is from Session 4; 0 otherwise,

$X_{23} = X_{16} \cdot X_{19}$,

$X_{24} = X_{16} \cdot X_{20}$,

$X_{25} = X_{16} \cdot X_{21}$,

$X_{26} = X_{17} \cdot X_{19}$,

$X_{27} = X_{17} \cdot X_{20}$, and

$X_{28} = X_{17} \cdot X_{21}$.

Tables similar to those found in previous chapters could be constructed for the preceding information. In the interest of economy of space, it is omitted here. The reader is reminded that each score (rather than each group) is the unit of analysis, thus, there are 60 scores for the data in Table 7.1. When preparing the data cards for a computer analysis, 60 data cards would be made.

To analyze the data in Table 7.1, it is useful to consider two separate analyses; one analysis can be a treatments X subjects design, temporarily disregarding the three different kinds of groups; then, it is useful to conceptualize the data in a two-way analysis of variance, disregarding for the time being that a given group has been measured several times.

The linear models that are useful for conceptualizing the data in Table 7.1 as a treatments X subjects design are as follows:

$$Y = b_0 + b_1X_1 + b_2X_2 + \dots + b_{14}X_{14} + e_1 \quad \text{(for the subjects (groups) effect)} \quad (7.1);$$

$$Y = b_0 + b_{19}X_{19} + b_{20}X_{21} + b_{21}X_{21} + e_2 \quad \text{(for the trend effect)} \quad (7.2);$$

and

$$Y = b_0 + b_1X_1 + b_2X_2 + \dots + b_{14}X_{14} + b_{19}X_{19} + b_{20}X_{20} + b_{21}X_{21} + e_3. \quad (7.3)$$

When these linear models are used, the following results can be found:

from equation 7.1, $SS_S = 104.14$;

from equation 7.2, $SS_{TREND} = 8.63$; and

from equation 7.3, $SS_{ERROR_1} = 32.52$; also, $SS_T = 145.29$.

While the preceding information would be sufficient for a treatments X subjects design, it should be recalled that in this formulation, the type of human relation group was disregarded.

Actually, the treatments effect is "nested," i.e., totally contained in the variation among subjects. Before pursuing this "nesting" further at this point, it is first useful to complete the analysis for the two-way formulation.

The following four linear models are sufficient:

$$Y = b_0 + b_{16}X_{16} + b_{17}X_{17} + e_4, \quad (\text{for the treatments effect}) \quad (7.4)$$

$$Y = b_0 + b_{19}X_{19} + b_{20}X_{20} + b_{21}X_{21} + e_2 \quad (\text{for the trend effect}) \quad (7.2)$$

$$Y = b_0 + b_{16}X_{16} + b_{17}X_{17} + b_{19}X_{19} + b_{20}X_{20} + b_{21}X_{21} + e_5 \quad (7.5); \text{ and}$$

$$Y = b_0 + b_{16}X_{16} + b_{17}X_{17} + b_{19}X_{19} + b_{20}X_{20} + b_{21}X_{21} + b_{23}X_{23} + \dots + b_{28}X_{28} + e_6 \quad (\text{Full Model}) \quad (7.6)$$

When these linear models are used, the following results can be found:

from equation 7.4, $SS_{\text{METH}} = 78.87$;

from equation 7.2, $SS_{\text{TREND}} = 8.63$;

from equation 7.6, $SS_{\text{ERROR}} = 39.71$.

The sum of squares attributed to regression for the full model (equation 7.6) is 105.58. The sum of squares attributed to regression for equation 7.5 is 87.50. The difference between these two values is equal to the interaction. Thus, $SS_{\text{METH} \times \text{TREND}} = 105.58 - 87.50 = 18.08$.

A summary table that would contain the foregoing information would appear as follows:

TABLE 7.2

SUMMARY TABLE FOR THE HUMAN RELATION GROUPS DATA IN TABLE 7.1

Source of Variation	df	SS	MS	F
Among Subjects				
Method	2	78.87	39.44	18.69**
Error (a)	12	25.27	2.11	
Total Among Subjects	14	104.14		
Within Subjects				
trend	3	8.63	2.88	7.20**
meth x trend	6	18.08	3.01	7.52**
error (b)	36	14.44	.40	
Total Within Subjects	45	41.15		
Total	59	145.29		

**Significant at .01 level

The interest, then, is in finding ways to isolate error (a) and error (b). Isolation of these two error terms is effected through the use of equation 7.1. Because of the complexity of the solution, it is useful to reproduce part of the regression output relating to equation 7.1; see Table 7.3.

REGRESSION OUTPUT FOR THE SUBJECTS EFFECT

Variable No.	Mean	Standard Deviation	Correlation X vs Y	Regression Coefficient	Standard Error of Regression Coefficient	Computed t value	Beta
1	0.06667	0.25155	-0.30801	-1.64995	0.67618	-2.44011	-0.26449
2	0.06667	0.25155	-0.17061	-0.84996	0.67618	-1.25700	-1.13625
3	0.06667	0.25155	-0.11908	-0.54996	0.67618	-0.81334	-0.08816
4	0.06667	0.25155	-0.50552	-2.79995	0.67618	-4.14084	-0.44883
5	0.06667	0.25155	-0.28224	-1.49994	0.67618	-2.21827	-0.24044
6	0.06667	0.25155	0.17289	1.15002	0.67618	1.70076	0.18435
7	0.06667	0.25155	0.09730	0.67618	0.67618	0.29583	0.03207
8	0.06667	0.25155	0.20724	1.35003	0.67618	1.99655	0.21641
9	0.06667	0.25155	0.08702	0.65002	0.67618	0.96132	0.10420
10	0.06667	0.25155	0.09561	0.70004	0.67618	1.03528	0.11222
11	0.06667	0.25155	0.31030	1.95003	0.67618	2.88388	0.31259
12	0.06667	0.25155	0.11278	0.80003	0.67618	1.18316	0.12825
13	0.06667	0.25155	0.27595	1.75003	0.67618	2.58810	0.28053
14	0.06667	0.25155	0.13855	0.95003	0.67618	1.40499	0.15229
Dependent Y	4.94332	1.56923					
Intercept		4.79995					
Multiple Correlation		0.84662					
St. Error of Estimate		0.95627					
Multiple Correlation Squared			0.71676				
One Minus Multiple Correlation Sqd.			0.28324				

ANALYSIS OF VARIANCE FOR THE REGRESSION

Source of Variation	Degrees of Freedom	Sum of Squares	Mean Squares	F Value
Attributable to Regression	14	104.13638	7.43831	8.13413
Deviation from Regression	45	41.15059	0.91446	
Total	59	145.28697		

The isolation of these error terms should be accompanied by isolating the sources of variance. That is, two separate analyses are indicated; the analysis among subjects and the analysis within subjects. This in turn can be accomplished by constructing two new criteria Y^* and Y^{**} such that $Y^* = \hat{Y}$, where the $\hat{Y}$ are the predicted scores found from using the regression coefficients in Table 7.3; also, $Y^{**} = Y - \hat{Y}$, or the error term associated with equation 7.1.

In particular,

$$\hat{Y} = 4.79995 - 1.64995X_1 - .84996X_2 - .54996X_3 + \dots + .95003X_{14}.$$

Several interesting things occur with using Y^* and Y^{**} (clearly, $Y^* + Y^{**} = Y$); the SS_T for $Y^* = 104.13701$ (within rounding of the attributable to regression in Table 7.3); and, the SS_T for $Y^{**} = 41.14972$ (within rounding of the deviation from regression in Table 7.3). Also, an equation can be formed such that $Y^* = b_0 + b_{16}X_{16} + b_{17}X_{17} + e_7$. (7.7) For equation 7.7, $SS_{\text{ATTRIBUTABLE}} = 78.87$ and $SS_{\text{DEVIATION}} = 25.27$; this is precisely the same results respectively as the Among Subjects variation for Method and Error (a) given in Table 7.2. Thus, hypotheses regarding the methods effect can be tested using Y^* and making restrictions on $Y^* = b_{16}X_{16} + b_{17}X_{17} + b_{18}X_{18} + e_7$, (and reparameterizing). For example, suppose that we wish to test for differences between groups one and three ($H_0 : \mu_{16} = \mu_{18}$).

$$\text{Then } Y^* = b_{16}X_{16} + b_{17}X_{17} + b_{18}X_{18} + e_8,$$

$$Y^* = b_{16}(X_{16} + X_{18}) + b_{17}X_{17} + e_8.$$

Let $V_1 = 1$ if a member of either group one or group three; then

$$Y^* = b_{16}V_1 + b_{17}X_{17} + e_8. \quad (7.8)$$

Equation 7.8 can be reparameterized to $Y^* = b_0 + b_{17}X_{17} + e_8$. (7.9)

The R^2 associated with equation 7.7 is .75735 and the R^2 associated with equation 7.9 is .128035. Then $F = \frac{.75735 - .128040}{(1 - .75735)/12} = 31.12161$;
 $\sqrt{F} = t = 5.57867$.

Without using regression, the same result can be found by

$$t = \frac{\bar{Y}_1 - \bar{Y}_3}{\sqrt{MS_{(a)} \left(\frac{1}{n_1} + \frac{1}{n_3} \right)}}$$

$$t = \frac{3.33 - 5.89}{\sqrt{2.11 \left(\frac{1}{20} + \frac{1}{20} \right)}} = -5.57129.$$

The latter answer has only two decimal point accuracy; had $MS_{(a)}$ been taken to five decimal points, $MS_{(a)} = 2.10583$, then $t = 5.57868$, very close to the earlier given answer.

Finding Models for Within Subjects

The second constructed criterion, Y^{**} , can be used (with extreme caution) for tests of trend differences and for comparing one cell mean to another.

Before testing for these effects, it is first useful to construct a full set of cell variables. Actually, several cell variables have already been used in the analysis:

- $X_{23} = X_{16} \cdot X_{19}$ (structural treatment, session 1, or cell 1),
- $X_{24} = X_{16} \cdot X_{20}$ (structured treatment, session 2, or cell 2),
- $X_{25} = X_{16} \cdot X_{21}$ (structured treatment, session 3, or cell 3),
- $X_{29} = X_{16} \cdot X_{22}$ (structured treatment, session 4, or cell 4),
- $X_{26} = X_{17} \cdot X_{19}$ (unstructured treatment, session 1, or cell 5),
- $X_{27} = X_{17} \cdot X_{20}$ (unstructured treatment, session 2, or cell 6),
- $X_{28} = X_{17} \cdot X_{21}$ (unstructured treatment, session 3, or cell 7),

$X_{30} = X_{17} \cdot X_{22}$ (unstructured treatment, session 4, or cell 8),

$X_{31} = X_{18} \cdot X_{19}$ (problems treatment, session 1, or cell 9),

$X_{32} = X_{18} \cdot X_{20}$ (problems treatment, session 2, or cell 10),

$X_{33} = X_{18} \cdot X_{21}$ (problems treatment, session 3, or cell 11) and

$X_{34} = X_{18} \cdot X_{22}$ (problems treatment, session 4, or cell 12).

Now, suppose there is interest in testing among the sessions for group differences, in particular, differences among session 1 and session 3 group means. Session 1 data is contained in cells 1,5 and 9 (X_{23} , X_{26} and X_{31}), session 3 data is contained in cells 3,7 and 11 (X_{24} , X_{28} and X_{33}).

Now, a full model should be constructed that contains the cells of interest.

One such full model is as follows:

$$Y^{**} = b_0 + b_1X_1 + b_2X_2 + \dots + b_{14}X_{14} + b_{23}X_{23} + b_{24}X_{24} + b_{25}X_{25} + b_{26}X_{26} + b_{27}X_{27} + b_{28}X_{28} + b_{31}X_{31} + b_{32}X_{32} + b_{33}X_{33} + e_9. \quad (7.10)$$

Care must be taken to make sure equation 7.10 is actually a true full model; SSE_8 should equal $SS_{ERROR(b)}$ (14.44). A restriction reflecting the hypothesis of interest can be written using the regression coefficients: $b_{23} + b_{26} + b_{31} = b_{25} + b_{28} + b_{33}$. In term of b_{23} this yields $b_{23} = b_{25} + b_{28} + b_{33} - b_{26} - b_{31}$. Placing this restriction on equation 7.10 yields:

$$Y^{**} = b_0 + b_1X_1 + \dots + b_{14}X_{14} + (b_{25} + b_{28} + b_{33} - b_{26} - b_{31})X_{23} + b_{24}X_{24} + b_{25}X_{25} + b_{26}X_{26} + b_{27}X_{27} + b_{28}X_{28} + b_{29}X_{29} + b_{30}X_{30} + b_{31}X_{31} + b_{32}X_{32} + b_{33}X_{33} + e_9.$$

$$Y^{**} = b_0 + b_1X_1 + \dots + b_{14}X_{14} + b_{24}X_{24} + b_{25}(X_{25} + X_{23}) + b_{26}(X_{26} - X_{23}) + b_{27}X_{27} + b_{28}(X_{28} + X_{23}) + b_{29}X_{29} + b_{30}X_{30} + b_{31}(X_{31} - X_{23}) + b_{32}X_{32} + b_{33}(X_{33} + X_{23}) + e_{10}. \quad (7.11)$$

Let $V_2 = 1$ if a member of X_{23} or X_{21} (cells 1 or 3), 0 otherwise,

$V_3 = 1$ if a member of X_{26} (cell 5), -1 if a member of X_{23} (cell 1),
0 otherwise,

$V_4 = 1$ if a member of X_{28} or X_{23} (cells 7 or 1), 0 otherwise,

$V_5 = 1$ if a member of X_{31} (cell 9), -1 if a member of X_{23} (cell 1),
0 otherwise and

$V_6 = 1$ if a member of X_{31} or X_{23} (cells 10 or 1), 0 otherwise.

Then equation 7.11 becomes:

$$Y^{**} = b_0 + b_1X_1 + \dots + b_{14}X_{14} + b_{24}X_{24} + b_{25}V_2 + b_{26}V_3 + b_{27}X_{27} + \\ b_{28}V_4 + b_{29}X_{29} + b_{30}X_{30} + b_{31}V_5 + b_{32}X_{32} + b_{33}V_6 + e_{10}. \quad (7.12)$$

Equation 7.10 yields $R_{10}^2 = .64899$; equation 7.12 yields $R_{12}^2 = .62874$.

Then

$$F = \frac{(.64889 - .62874)/1}{(1 - .64889)/36} = 2.06603;$$

$$t = \sqrt{F} = 1.43737.$$

On the other hand, this same hypothesis could be much more easily tested by

$$t = \frac{\bar{Y}_{.1} - \bar{Y}_{.3}}{\sqrt{MS(b) \left(\frac{1}{n_{.1}} + \frac{1}{n_{.3}} \right)}} = \frac{5.00 - 5.33}{\sqrt{.40 \left(\frac{1}{15} + \frac{1}{15} \right)}} = -1.43478.$$

The last calculation is only accurate to two decimal points.

Testing for Cell Differences

If there is interest in testing for differences among cells (say cell 1 and cell 2), then these tests can be fairly easily accommodated.

In that the full model in equation 7.10 contains both X_{23} and X_{24} (cells 1 and 2), a restriction can be made regarding these two cells:

$$b_{23} = b_{24}.$$

Then

$$Y^{**} = b_0 + b_1 X_1 + \dots + b_{14} X_{14} + b_{24} (X_{23} + X_{24}) + b_{25} X_{25} + \dots + b_{33} X_{33} + e_{11}. \quad (7.13)$$

Let $V_7 = 1$ if a member of either X_{23} or X_{24} (cells 1 and 2), 0 otherwise.

Equation 7.13 becomes:

$$Y^{**} = b_0 + b_1 X_1 + \dots + b_{14} X_{14} + b_{24} V_7 + b_{25} X_{25} + \dots + b_{33} X_{33} + e_{11}. \quad (7.14)$$

Equation 7.14 yields $R_{14}^2 = .63927$; previously, $R_{10}^2 = .64889$. Therefore,

$$F = \frac{(.64889 - .63927)/1}{(1 - .64889)/36} = .98636,$$

$$t = \sqrt{F} = .99316.$$

This cell mean difference could also be tested by

$$t = \frac{Y_{11} - Y_{12}}{\sqrt{MS(b) \left(\frac{1}{n_{11}} + \frac{1}{n_{12}} \right)}} = \frac{3.68 - 3.28}{\sqrt{(.40) \left(\frac{1}{5} + \frac{1}{5} \right)}} = 1.00,$$

which is close to the previously given value of $t = .99316$.

Practical Limitations of the Regression Approach

One of the selling points regarding the use of multiple linear regression as a general analysis technique, including solving problems more commonly done by the analysis of variance, is that not only is the process more easily conceptualized, but also the computations are usually greatly simplified. This present chapter should serve as ample proof that such is not always the case. Clearly, the multiple comparisons performed here could more easily be accomplished by the usual analysis variance techniques. Perhaps this is but another way of saying that limiting yourself in the possible analysis techniques may be in the long

run too limiting; the same point was made much earlier and much better in "Don't put all your eggs in one basket". While multiple linear regression is an excellent tool for both conceptualizing and calculating many tests, the traditional techniques can occasionally prove to be easier to use, even for a hard-nosed regression aficionado.

Using Appropriate Tables

Perhaps not enough was said regarding using proper tables to evaluate an experiment. In general, the total number (and type) of contrasts run should always be considered when choosing a multiple comparison technique (and hence, choosing a table to evaluate significance). Every time a set of tests are run, the number (and type) of tests should be considered when reporting levels of significance. Those who would contend that "standard practise" would allow them to use such techniques as Duncan's tests or orthogonal comparisons (which do not usually retain an experimentwise error rate) should remember that "standard practise" is the defense that physicians use in court when the patient dies.

TABLES

Tables Ia, Ib, Ic, and Id are reproduced from C. W. Dunnett, A multiple comparison procedure for comparing several treatments with a control. Journal of the American Statistical Association, 1955, 50, 1096-122, and C. W. Dunnett, New Tables for multiple comparisons with a control. Biometrics, 1964, 20, 482-491, with permission of the author.

Tables IIa and IIb were calculated by the present writer by transforming the values in Harter (1960) by division by $\sqrt{2}$.

Tables IIIa and IIIb were calculated by the present writer by transforming values in Duncan (1955) by division by $\sqrt{2}$.

Tables IVa and IVb are reproduced from O. J. Dunn, Multiple comparisons among means. Journal of the American Statistical Association, 1961, 56: 52-64 by permission of the author.

In Tables Ia, Ib, Ic, Id, IIa, IIb, IIIa, IIIb, IVa and IVb, v refers to the degrees of freedom within (df_w).

In Tables IIa, IIb, IIIa, IIIb, r refers to the number of means in the range.

In Tables Ia, Ib, Ic, Id, k refers to the number of groups compared to the control (excluding the control).

In Tables IVa and IVb, m refers to the total number of a priori contrasts.

TABLE Ia
PERCENTAGE OF DUNNETT'S TEST
(05 LEVEL)

One-tailed									
$v \backslash k$	1	2	3	4	5	6	7	8	9
5	2.02	2.44	2.68	2.85	2.96	3.08	3.16	3.24	3.30
6	1.94	2.34	2.56	2.71	2.83	2.92	3.00	3.07	3.12
7	1.89	2.27	2.48	2.62	2.73	2.82	2.89	2.95	3.01
8	1.86	2.22	2.42	2.55	2.66	2.74	2.81	2.87	2.92
9	1.83	2.18	2.37	2.50	2.60	2.68	2.75	2.81	2.86
10	1.81	2.15	2.34	2.47	2.56	2.64	2.70	2.76	2.81
11	1.80	2.13	2.31	2.44	2.53	2.60	2.67	2.72	2.77
12	1.78	2.11	2.29	2.41	2.50	2.58	2.64	2.69	2.74
13	1.77	2.09	2.27	2.39	2.48	2.55	2.61	2.66	2.71
14	1.76	2.08	2.25	2.37	2.46	2.53	2.59	2.64	2.69
15	1.75	2.07	2.24	2.36	2.44	2.51	2.57	2.62	2.67
16	1.75	2.06	2.23	2.34	2.43	2.50	2.56	2.61	2.65
17	1.74	2.05	2.22	2.33	2.42	2.49	2.54	2.59	2.64
18	1.73	2.04	2.21	2.32	2.41	2.48	2.53	2.58	2.62
19	1.73	2.03	2.20	2.31	2.40	2.47	2.52	2.57	2.61
20	1.72	2.03	2.19	2.30	2.39	2.46	2.51	2.56	2.60
30	1.70	1.99	2.15	2.25	2.33	2.40	2.45	2.50	2.54
40	1.68	1.97	2.13	2.23	2.31	2.37	2.42	2.47	2.51
60	1.67	1.95	2.10	2.21	2.28	2.35	2.39	2.44	2.48
120	1.66	1.93	2.08	2.18	2.26	2.32	2.37	2.41	2.45
∞	1.64	1.92	2.06	2.16	2.23	2.29	2.34	2.38	2.42

TABLE 1b
 PERCENTAGE POINTS OF DUNNETT'S TEST
 (.01 LEVEL)

One-tailed									
k	1	2	3	4	5	6	7	8	9
	3.37	3.90	4.21	4.43	4.60	4.73	4.85	4.94	5.03
	3.14	3.61	3.88	4.07	4.21	4.33	4.43	4.51	4.59
	3.00	3.42	3.66	3.83	3.96	4.07	4.15	4.23	4.30
	2.90	3.29	3.51	3.67	3.79	3.88	3.96	4.03	4.09
	2.82	3.19	3.40	3.55	3.66	3.75	3.83	3.89	3.94
	2.76	3.11	3.31	3.45	3.56	3.64	3.71	3.78	3.83
	2.72	3.06	3.25	3.38	3.48	3.56	3.63	3.69	3.74
	2.68	3.01	3.19	3.32	3.42	3.50	3.56	3.62	3.67
	2.65	2.97	3.15	3.27	3.37	3.44	3.51	3.56	3.61
	2.62	2.94	3.11	3.23	3.32	3.40	3.46	3.51	3.56
	2.60	2.91	3.08	3.20	3.29	3.36	3.42	3.47	3.52
	2.58	2.88	3.05	3.17	3.26	3.33	3.39	3.44	3.48
	2.57	2.86	3.03	3.14	3.23	3.30	3.36	3.41	3.45
	2.55	2.84	3.01	3.12	3.21	3.27	3.33	3.38	3.42
	2.54	2.83	2.99	3.10	3.18	3.25	3.31	3.36	3.40
	2.53	2.81	2.97	3.08	2.17	3.23	3.29	3.34	3.38
	2.46	2.72	2.87	2.97	3.05	3.11	3.16	3.21	3.24
	2.42	2.68	2.82	2.92	2.99	3.05	3.10	3.14	3.18
	2.39	2.64	2.78	2.87	2.94	3.00	3.04	3.08	3.12
	2.36	2.60	2.73	2.82	2.89	2.94	2.99	3.03	3.06
	2.33	2.56	2.68	2.77	2.84	2.89	2.93	2.97	3.00

TABLE IIb
 PERCENTAGE POINTS OF THE STUDENTIZED RANGE
 REPORTED AT t VALUES (.01 LEVEL)

$v \backslash r$	2	3	4	5	6	7	8	9	10	20
5	4.032	4.933	5.518	5.955	6.302	6.591	6.837	7.051	7.241	8.456
6	3.707	4.476	4.973	5.343	5.638	5.882	6.099	6.271	6.433	7.453
7	3.499	4.185	4.627	4.953	5.213	5.430	5.614	5.774	5.917	6.821
8	3.356	3.985	4.387	4.685	4.921	5.117	5.285	5.431	5.560	6.383
9	3.259	3.838	4.212	4.489	4.708	4.890	5.044	5.180	5.300	6.062
10	3.169	3.726	4.079	4.339	4.545	4.716	4.861	4.953	5.100	5.817
11	3.106	3.639	3.975	4.221	4.417	4.579	4.718	4.838	4.944	5.623
12	3.055	3.568	3.891	4.127	4.314	4.470	4.601	4.716	4.818	5.467
13	3.012	3.510	3.821	4.050	4.229	4.378	4.474	4.616	4.714	5.337
14	2.977	3.461	3.763	3.982	4.158	4.303	4.425	4.532	4.627	5.229
15	2.947	3.420	3.714	3.929	4.098	4.238	4.357	4.461	4.553	5.136
16	2.921	3.384	3.671	3.881	4.046	4.183	4.299	4.400	4.489	5.057
17	2.898	3.353	3.635	3.840	4.002	4.134	4.248	4.347	4.434	4.987
18	2.879	3.326	3.602	3.804	3.962	4.093	4.203	4.300	4.385	4.927
19	2.861	3.302	3.574	3.772	3.927	4.055	4.164	4.258	4.342	4.873
20	2.845	3.280	3.548	3.743	3.896	4.022	4.129	4.221	4.304	4.825
30	2.750	3.150	3.393	3.569	3.707	3.819	3.915	3.997	4.070	4.530
40	2.705	3.088	3.321	3.487	3.616	3.723	3.813	3.891	3.959	4.390
60	2.660	3.028	3.249	3.407	3.529	3.630	3.714	3.787	3.852	4.253
120	2.618	2.970	3.180	3.330	3.445	3.539	3.619	3.687	3.747	4.120
∞	2.576	2.913	3.113	3.255	3.364	3.452	3.526	3.591	3.647	3.992

TABLE IIIa

PERCENTAGE POINTS OF THE STUDENTIZED RANGE
REPORTED AS t VALUES (.05 LEVEL)

$v \backslash r$	2	3	4	5	6	7	8	9	10	20
5	2.570	2.651	2.685	2.697	2.697	2.697	2.697	2.697	2.697	2.697
6	2.447	2.536	2.580	2.602	2.612	2.614	2.614	2.614	2.614	2.614
7	2.365	2.459	2.509	2.537	2.553	2.561	2.564	2.564	2.564	2.564
8	2.306	2.403	2.457	2.490	2.510	2.522	2.528	2.531	2.531	2.531
9	2.262	2.361	2.361	2.418	2.454	2.476	2.491	2.500	2.508	2.508
10	2.228	2.329	2.387	2.425	2.450	2.469	2.478	2.486	2.490	2.493
11	2.201	2.302	2.363	2.402	2.429	2.448	2.461	2.470	2.476	2.482
12	2.179	2.280	2.343	2.383	2.411	2.432	2.446	2.456	2.462	2.474
13	2.160	2.263	2.326	2.367	2.396	2.418	2.434	2.445	2.453	2.468
14	2.145	2.247	2.311	2.354	2.384	2.406	2.423	2.435	2.444	2.464
15	2.131	2.234	2.298	2.342	2.373	2.396	2.413	2.427	2.437	2.461
16	2.120	2.223	2.287	2.332	2.364	2.387	2.406	2.420	2.430	2.459
17	2.110	2.213	2.278	2.323	2.355	2.380	2.399	2.413	2.425	2.458
18	2.101	2.205	2.270	2.315	2.345	2.373	2.392	2.408	2.419	2.456
19	2.093	2.197	2.262	2.308	2.341	2.367	2.386	2.402	2.412	2.456
20	2.086	2.190	2.256	2.302	2.336	2.361	2.382	2.398	2.411	2.456
30	2.042	2.146	2.214	2.262	2.298	2.326	2.349	2.368	2.384	2.454
40	2.021	2.126	2.193	2.242	2.280	2.309	2.333	2.353	2.370	2.453
50	2.000	2.104	2.173	2.222	2.261	2.292	2.317	2.338	2.357	2.452
60	1.980	2.084	2.153	2.203	2.243	2.275	2.301	2.324	2.343	2.451
70	1.960	2.063	2.133	2.184	2.225	2.258	2.285	2.309	2.329	2.451

REFERENCES

- Anderson, R. L. and Bancroft, T. A. Statistical theory in research. New York: McGraw-Hill, 1952.
- Bottenberg, R. A., and Ward, J. H. Applied multiple linear regression. Lackland Air Force Base, Texas: Personnel Research Laboratory PRL-TDR-63-6, 1963.
- Cohen, J. Multiple regression as a general data-analytic system. Psychological Bulletin, 1968, 70, 526-443.
- Cohen, J. and Cohen, P. Applied multiple regression/correlation analysis for the behavioral sciences. New York: Lawrence Erlbaum Associates, 1975.
- Duncan, D. B. Multiple range and multiple F tests. Biometrics, 1955, 11, 1-42.
- Dunn, O. J. Multiple comparisons among means. Journal of the American Statistical Association, 1961, 56, 52-64.
- Dunnett, C. W. A multiple comparison procedure for comparing several treatments with a control. Journal of the American Statistical Association, 1955, 50, 1096-1121.
- Dunnett, C. W. New tables for multiple comparisons with a control. Biometrics, 1964, 20, 482-491.
- Harter, H. L. Tables of range and studentized range. Annals of Mathematical Statistics, 1960, 31, 1122-1147.
- Jennings, E. Fixed effects analysis of variance by regression analysis. Multivariate Behavioral Research, 1967, 2, 95-108.
- Kelly, F. J., Beggs, D. W., and McNeil, K. A. Multiple regression approach. Carbondale, Ill.: Southern Illinois University Press, 1969.
- Kerlinger, F. N. and Pedhazur, E. J. Multiple regression in behavioral research. New York: Holt, Rinehart and Winston, 1974.
- Keuls, M. The use of studentized range in connection with an analysis of variance, Euphytica, 1952, 1, 112-122.
- Lindquist, E. F. Design and analysis of experiments in psychology and education. Boston: Houghton Mifflin Company, 1953.

- McNeil, K. A., Kelly, F. J. and McNeil, J. T. Testing research hypotheses using multiple linear regression. Carbondale, Ill: Southern Illinois University Press, 1975.
- Mendenhall, W. The design and analysis of experiments. Belmont, Calif.: Wadsworth, 1968.
- Miller, R. G. Simultaneous statistical inference. New York: McGraw-Hill, 1966.
- Newman, D. The distribution of range in samples from a normal population, expressed in terms of an independent estimate of standard deviation. Biometrika, 1939, 31, 20-30
- Overall, J. E. and Spiegel, D. H. Concerning least squares analysis of experimental data. Psychological Bulletin, 1969, 72, 311-322.
- Overall, J. E. Spiegel, D. H. and Cohen, J. Equivalence of orthogonal and nonorthogonal analysis of variance. Psychological Bulletin, 1975, 82, 182-186.
- Ryan, T. A. Multiple comparisons in psychological research. Psychological Bulletin, 1959, 56, 26-47.
- Ryan, T. A. The experiment as the unit for computing rates of error. Psychological Bulletin, 1962, 59, 301-305.
- Scheffe, H. The analysis of variance. New York: John Wiley & Sons, 1959.
- Sparks, J. N. Expository notes on the problem of making multiple comparisons in a completely randomized design. Journal of Experimental Education, 1963, 31, 343-349.
- Searle, S. R. Linear models. New York: Wiley, 1971.
- Speed, F. A., and Hocking, R. R. The use of R ()- Notation with unbalanced data. The American Statistician, 1976, 30, 30-33
- Tukey, J. W. The problem of multiple comparisons. Dittoed, Princeton University, 1953.
- Timm, N. H. and Carlson, J. E. Analysis of variance through full rank models. Multivariate Behavioral Research: Monograph, 1975, 75-1.
- Ward, J. H. and Jennings, E. Introduction to linear models. Englewood Cliffs, N. J.: Prentice-Hall, 1973.

Williams, J. D. Two way fixed effects analysis of variance with disproportionate cell frequencies. Multivariate Behavioral Research, 1972, 7, 67-83.

Williams, J. D. Regression analysis in educational research. New York: MSS Information Corporation, 1974a.

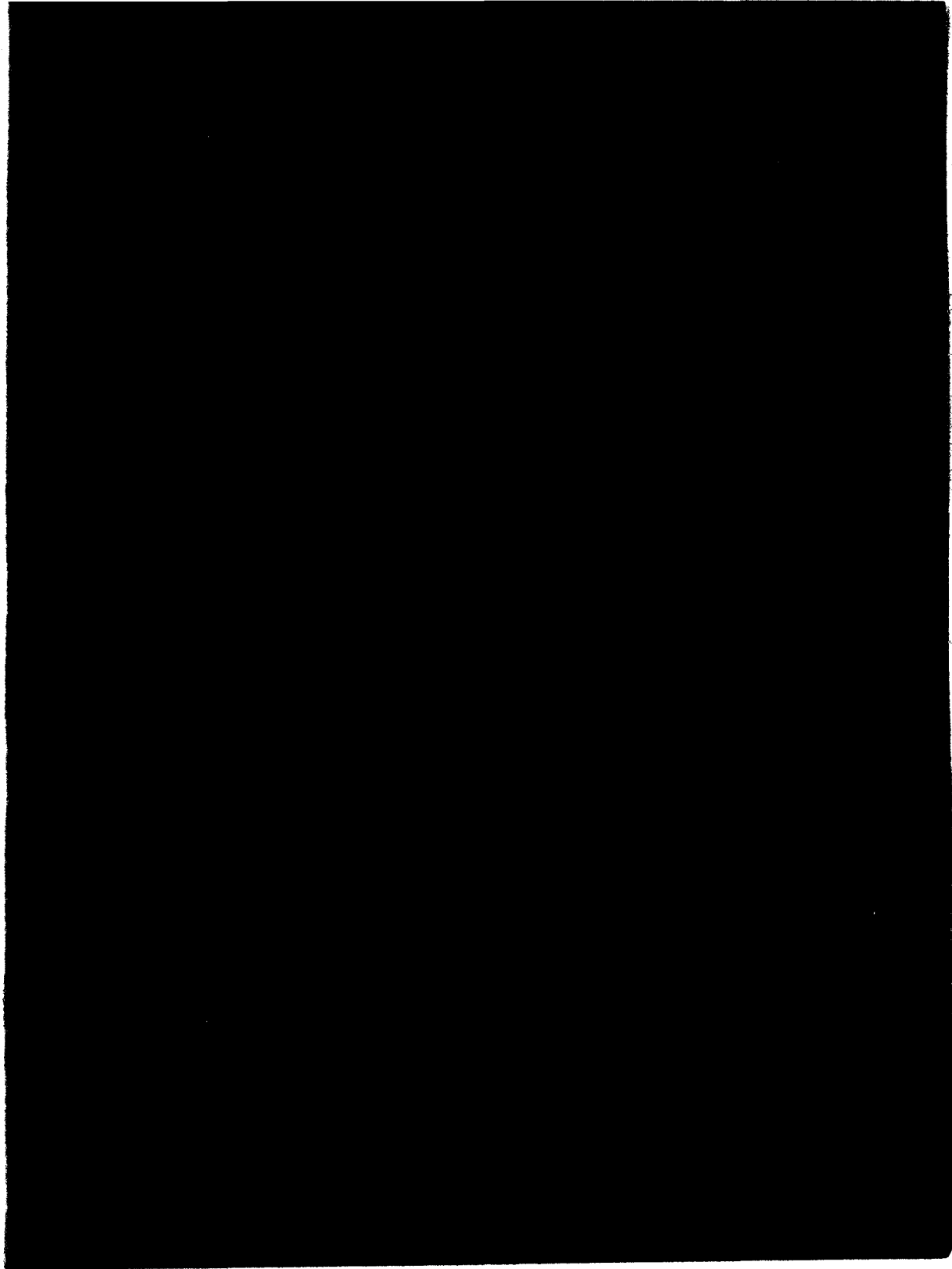
Williams, J. D. A simplified regression formulation of Tukey's test. Journal of Experimental Education, 1974b, 42, No.4, 80-82.

Williams, J. D. Multiple comparisons by multiple linear regression. Multiple Linear Regression Viewpoints Monograph Series, 2, 1976, 7, No. 1

Williams, J. D. Full rank and non-full rank models with contrast and binary coding systems for two-way disproportionate all frequency analyses. Multiple Linear Regression Viewpoints, 1977a, 8, No.1, 1-18.

Williams, J. D. A note on coding the subjects effect in treatments X subjects designs. Multiple Linear Regression Viewpoints, 1977b, 8, No. 1, 32-35.

Winer, B. J. Statistical principles in experimental design. 2nd ed. New York: McGraw-Hill, 1971.



146K0GE1
ROGERS, BRUCE G.
DEPT OF ED PSYCH
UNIV OF NO ICWA
CEDAR FALLS, IOWA 50613

