

Chapter 224

General Linear Models (GLM) for Fixed Factors

Introduction

This procedure performs analysis of variance (ANOVA) and analysis of covariance (ANCOVA) for factorial models that include fixed factors (effects) and/or covariates. This procedure uses multiple regression techniques to estimate model parameters and compute least squares means. This procedure also provides standard error estimates for least squares means and their differences and calculates many multiple comparison tests (Tukey-Kramer, Dunnett, Bonferroni, Scheffe, Sidak, and Fisher's LSD) with corresponding adjusted P-values and simultaneous confidence intervals. This procedure also provides residuals for checking assumptions.

This procedure cannot be used to analyze models that include nested (e.g., repeated measures ANOVA, split-plot, split-split-plot, etc.) or random factors (e.g., randomized block, etc.). If your model includes nested and/or random factors, use General Linear Models (GLM), Repeated Measures ANOVA, or Mixed Models instead.

If you have only one fixed factor in your model, then you might want to consider using the One-Way Analysis of Variance procedure or the One-Way Analysis of Variance using Regression procedure instead.

If you have only one fixed factor and one covariate in your model, then you might want to consider using the One-Way Analysis of Covariance (ANCOVA) procedure instead.

Kinds of Research Questions

A large amount of research consists of studying the influence of a set of independent variables on a response (dependent) variable. Many experiments are designed to look at the influence of a single independent variable (factor) while holding other factors constant. These experiments are called single-factor experiments and are analyzed with the one-way analysis of variance (ANOVA). A second type of design considers the impact of one factor across several values of other factors. This experimental design is called the factorial design.

The factorial design is popular among researchers because it not only lets you study the individual effects of several factors in a single experiment, but it also lets you study their interaction. Interaction is present when the response variable fails to behave the same at values of one factor when a second factor is varied. Since factors seldom work independently, the study of their interaction becomes very important.

Analysis of covariance (ANCOVA) is another design that may be analyzed using this procedure. ANCOVA is useful when you want to improve precision by removing various extraneous sources of variation from your study.

The Linear Model

We begin with an infinite population of individuals with many measurable characteristics. These individuals are (mentally) separated into two or more treatment populations based on one or more of these characteristics. A random sample of the individuals in each population is drawn. A treatment is applied to everyone in the sample and an outcome is measured. The data so obtained are analyzed using an analysis of variance table that produces an F-test.

A mathematical model may be formulated that underlies each analysis of variance. This model expresses the response variable as the sum of parameters of the population. For example, a linear mathematical model for a two-factor experiment is

$$Y_{ijk} = m + a_i + b_j + (ab)_{ij} + e_{ijk}$$

where $i = 1, 2, \dots, I$; $j = 1, 2, \dots, J$; and $k = 1, 2, \dots, K$. This model expresses the value of the response variable, Y , as the sum of five components:

- m the mean.
- a_i the contribution of the i^{th} level of a factor A.
- b_j the contribution of the j^{th} level of a factor B.
- $(ab)_{ij}$ the combined contribution of the i^{th} level of a factor A and the j^{th} level of a factor B.
- e_{ijk} the contribution of the k^{th} individual. This is often called the "error."

Note that this model is the sum of various constants. This type of model is called a linear model. It becomes the mathematical basis for our discussion of the analysis of variance. Also note that this serves only as an example. Many linear models could be formulated for the two-factor experiment.

Assumptions

The following assumptions are made when using the F-test.

1. The response variable is continuous.
2. The e_{ijk} follow the normal probability distribution with mean equal to zero.
3. The variances of the e_{ijk} are equal for all values of i , j , and k .
4. The individuals are independent.

Normality of Residuals

The residuals are assumed to follow the normal probability distribution with zero mean and constant variance. This can be evaluated using a normal probability plot of the residuals. Also, normality tests are used to evaluate this assumption. The most popular of the five normality tests provided is the Shapiro-Wilk test.

Unfortunately, a breakdown in any of the other assumptions results in a departure from this assumption as well. Hence, you should investigate the other assumptions first, leaving this assumption until last.

Limitations

There are few limitations when using these tests. Sample sizes may range from a few to several hundred. If your data are discrete with at least five unique values, you can assume that you have met the continuous variable assumption. Perhaps the greatest restriction is that your data comes from a random sample of the population. If you do not have a random sample, the F-test will not work.

When missing cells occur in your design, you must take special care to be sure that appropriate interaction terms are removed from the model.

Representing Factor Variables

Categorical factor variables take on only a few unique values. For example, suppose a therapy variable has three possible values: A, B, and C. One question is how to include this variable in the regression model. At first glance, we can convert the letters to numbers by recoding A to 1, B to 2, and C to 3. Now we have numbers. Unfortunately, we will obtain completely different results if we recode A to 2, B to 3, and C to 1. Thus, a direct recode of letters to numbers will not work.

To convert a categorical variable to a form usable in regression analysis, we must create a new set of numeric variables. If a categorical variable has k values, $k - 1$ new binary variables must be generated.

Indicator (Binary) Variables

Indicator (dummy or binary) variables are created as follows. A *reference group* is selected. Usually, the most common value or the control is selected as the reference group. Next, a variable is generated for each of the groups other than the reference group. For example, suppose that C is selected as the reference group. An indicator variable is generated for each of the remaining groups: A and B. The value of the indicator variable is one if the value of the original variable is equal to the value of interest, or zero otherwise. Here is how the original variable T and the two new indicator variables TA and TB look in a short example.

<u>T</u>	<u>TA</u>	<u>TB</u>
A	1	0
A	1	0
B	0	1
B	0	1
C	0	0
C	0	0

The generated variables, TA and TB, would be used as columns in the design matrix, X , in the model.

Representing Interactions of Two or More Categorical Variables

When the interaction between two categorical variables is included in the model, an interaction variable must be generated for each combination of the variables generated for each categorical variable. This can be accomplished automatically in **NCSS** using an appropriate Model statement.

In the following example, the interaction between the categorical variables *T* and *S* are generated. Try to determine the reference value used for variable *S*.

<u>T</u>	<u>TA</u>	<u>TB</u>	<u>S</u>	<u>SD</u>	<u>SE</u>	<u>TASD</u>	<u>TASE</u>	<u>TBSD</u>	<u>TBSE</u>
A	1	0	D	1	0	1	0	0	0
A	1	0	E	0	1	0	1	0	0
B	0	1	F	0	0	0	0	0	0
B	0	1	D	1	0	0	0	1	0
C	0	0	E	0	1	0	0	0	1
C	0	0	F	0	0	0	0	0	0

When the variables, *TASD*, *TASE*, *TBSD*, and *TBSE* are added to the model, they will account for the interaction between *T* and *S*.

Representing Interactions of Numeric and Categorical Variables

When the interaction between a factor variable and a covariate is to be included in the model, all proceeds as above, except that an interaction variable must be generated for each categorical variable. This can be accomplished automatically in **NCSS** using an appropriate Model statement.

In the following example, the interaction between the factor variable *T* and the covariate variable *X* is created.

<u>T</u>	<u>TA</u>	<u>TB</u>	<u>X</u>	<u>XTA</u>	<u>XTB</u>
A	1	0	1.2	1.2	0
A	1	0	1.4	1.4	0
B	0	1	2.3	0	2.3
B	0	1	4.7	0	4.7
C	0	0	3.5	0	0
C	0	0	1.8	0	0

When the variables *XTA* and *XTB* are added to the model, they will account for the interaction between *T* and *X*.

Technical Details

This section presents the technical details of the analysis method (multiple regression) using a mixture of summation and matrix notation.

The Linear Model

The linear model can be written as

$$\mathbf{Y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{e}$$

where $\mathbf{Y}$ is a vector of N responses, $\mathbf{X}$ is an $N \times p$ design matrix, $\boldsymbol{\beta}$ is a vector of p fixed and unknown parameters, and $\mathbf{e}$ is a vector of N unknown, random error values. Define the following vectors and matrices:

$$\mathbf{Y} = \begin{bmatrix} y_1 \\ \vdots \\ y_j \\ \vdots \\ y_N \end{bmatrix}, \mathbf{X} = \begin{bmatrix} 1 & x_{11} & \cdots & x_{1p} \\ \vdots & \vdots & \ddots & \vdots \\ 1 & x_{1j} & \cdots & x_{pj} \\ \vdots & \vdots & \ddots & \vdots \\ 1 & x_{1N} & \cdots & x_{pN} \end{bmatrix}, \mathbf{e} = \begin{bmatrix} e_1 \\ \vdots \\ e_j \\ \vdots \\ e_N \end{bmatrix}, \mathbf{1} = \begin{bmatrix} 1 \\ \vdots \\ 1 \\ \vdots \\ 1 \end{bmatrix}, \boldsymbol{\beta} = \begin{bmatrix} \beta_0 \\ \beta_1 \\ \vdots \\ \beta_p \end{bmatrix}$$

$\mathbf{X}$ is the design matrix that includes covariates and sets of binary variables corresponding to fixed factor variables (and their interactions) included in the model.

Least Squares

Using this notation, the least squares estimates of the model coefficients, $\mathbf{b}$, are found using the equation.

$$\mathbf{b} = (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{Y}$$

The vector of predicted values of the response variable is given by

$$\hat{\mathbf{Y}} = \mathbf{X}\mathbf{b}$$

The residuals are calculated using

$$\mathbf{e} = \mathbf{Y} - \hat{\mathbf{Y}}$$

Estimated Variances

An estimate of the variance of the residuals is computed using

$$s^2 = \frac{\mathbf{e}'\mathbf{e}}{N - p - 1}$$

An estimate of the variance of the model coefficients is calculated using

$$\mathbf{v} \begin{pmatrix} b_0 \\ b_1 \\ \vdots \\ b_p \end{pmatrix} = s^2(\mathbf{X}'\mathbf{X})^{-1}$$

General Linear Models (GLM) for Fixed Factors

An estimate of the variance of the predicted mean of Y at a specific value of X , say X_0 , is given by

$$s_{Y_m|X_0}^2 = s^2(1, X_0)(\mathbf{X}'\mathbf{X})^{-1} \begin{pmatrix} 1 \\ X_0 \end{pmatrix}$$

An estimate of the variance of the predicted value of Y for an individual for a specific value of X , say X_0 , is given by

$$s_{Y_I|X_0}^2 = s^2 + s_{Y_m|X_0}^2$$

Hypothesis Tests of the Intercept and Coefficients

Using these variance estimates and assuming the residuals are normally distributed, hypothesis tests may be constructed using the Student's t distribution with $N - p - 1$ degrees of freedom using

$$t_{b_i} = \frac{b_i - B_i}{s_{b_i}}$$

Usually, the hypothesized value of B_i is zero, but this does not have to be the case.

Confidence Intervals of the Intercept and Coefficients

A $100(1 - \alpha)\%$ confidence interval for the true model coefficient, β_i , is given by

$$b_i \pm (t_{1-\alpha/2, N-p-1})s_{b_i}$$

Confidence Interval of Y for Given X

A $100(1 - \alpha)\%$ confidence interval for the mean of Y at a specific value of X , say X_0 , is given by

$$b'X_0 \pm (t_{1-\alpha/2, N-p-1})s_{Y_m|X_0}$$

A $100(1 - \alpha)\%$ prediction interval for the value of Y for an individual at a specific value of X , say X_0 , is given by

$$b'X_0 \pm (t_{1-\alpha/2, N-p-1})s_{Y_I|X_0}$$

R² (Percent of Variation Explained)

Several measures of the goodness-of-fit of the model to the data have been proposed, but by far the most popular is R^2 . R^2 is the square of the correlation coefficient between Y and $\hat{Y}$. It is the proportion of the variation in Y that is accounted by the variation in the independent variables. R^2 varies between zero (no linear relationship) and one (perfect linear relationship).

R^2 , officially known as the *coefficient of determination*, is defined as the sum of squares due to the linear regression model divided by the adjusted total sum of squares of Y . The formula for R^2 is

$$R^2 = 1 - \left(\frac{\mathbf{e}'\mathbf{e}}{\mathbf{Y}'\mathbf{Y} - \frac{(\mathbf{1}'\mathbf{Y})^2}{\mathbf{1}'\mathbf{1}}} \right)$$

$$= \frac{SS_{Model}}{SS_{Total}}$$

R^2 is probably the most popular measure of how well a model fits the data. R^2 may be defined either as a ratio or a percentage. Since we use the ratio form, its values range from zero to one. A value of R^2 near zero indicates no linear relationship, while a value near one indicates a perfect linear fit. Although popular, R^2 should not be used indiscriminately or interpreted without scatter plot support. Following are some qualifications on its interpretation:

1. *Additional independent variables.* It is possible to increase R^2 by adding more independent variables, but the additional independent variables may cause an increase in the mean square error, an unfavorable situation. This usually happens when the sample size is small.
2. *Range of the independent variables.* R^2 is influenced by the range of the independent variables. R^2 increases as the range of the X 's increases and decreases as the range of the X 's decreases.
3. *Slope magnitudes.* R^2 does not measure the magnitude of the slopes.
4. *Linearity.* R^2 does not measure the appropriateness of a linear model. It measures the strength of the linear component of the model. Suppose the relationship between X and Y was a perfect sphere. Although there is a perfect relationship between the variables, the R^2 value would be zero.
5. *Predictability.* A large R^2 does not necessarily mean high predictability, nor does a low R^2 necessarily mean poor predictability.
6. *Sample size.* R^2 is highly sensitive to the number of observations. The smaller the sample size, the larger its value.

$\bar{R}^2$ (Adjusted R^2)

R^2 varies directly with N , the sample size. In fact, when $N = p$, $R^2 = 1$. Because R^2 is so closely tied to the sample size, an adjusted R^2 value, called $\bar{R}^2$, has been developed. $\bar{R}^2$ was developed to minimize the impact of sample size. The formula for $\bar{R}^2$ is

$$\bar{R}^2 = 1 - \frac{(N - 1)(1 - R^2)}{N - p - 1}$$

Least Squares Means

As opposed to raw or arithmetic means which are simply averages of the grouped raw data values, least squares means are adjusted for the other terms in the model, both the other fixed factors and the covariates. In balanced designs with no covariates, the least squares group means will be equal to the raw group means. In unbalanced designs or when covariates are present, the least squares means usually are different from the raw means.

The least squares means and associated comparisons (i.e., differences) can be calculated using a linear contrast vector, $\mathbf{c}_i$. Means and differences are estimated as

$$\mathbf{c}_i' \mathbf{b},$$

with estimated standard error,

$$SE(\mathbf{c}_i' \mathbf{b}) = s \sqrt{\mathbf{c}_i' (\mathbf{X}' \mathbf{X})^{-1} \mathbf{c}_i}.$$

where s is the square root of the estimated mean square error (MSE) from the model based on v degrees of freedom.

For an ANCOVA model with a fixed factor with 4 levels and a covariate, and if level 4 were the reference value, the components of the contrast vector would take the form

$$c_i = (I, \mu_1, \mu_2, \mu_3, X)$$

where I represents the indicator for the intercept and X is the value of the covariate which the mean or difference is evaluated. The contrast vector used to estimate μ_2 would be

$$c_i = (1, 0, 1, 0, X).$$

The contrast vector used to estimate $\mu_1 - \mu_2$ would be

$$c_i = (0, 1, -1, 0, 0).$$

Confidence intervals for the estimable functions are of the form

$$\mathbf{c}_i' \mathbf{b} \pm c_\alpha SE(\mathbf{c}_i' \mathbf{b}),$$

where c_α is the critical value, usually selected to keep the experimentwise error rate equal to α for the collection of all comparisons (see Multiple Comparisons).

Two-sided significance tests for the mean and the difference (against a null value of zero) use the test statistic

$$T_i = \frac{|\mathbf{c}_i' \mathbf{b}|}{SE(\mathbf{c}_i' \mathbf{b})} \geq c_\alpha.$$

Multiple Comparisons

Given that the analysis of variance (ANOVA) test finds a significant difference among treatment means, the next task is to determine which treatments are different. Multiple comparison procedures (MCPs) are methods that pinpoint which treatments are different.

The discussion to follow considers the following experiment. Suppose an experiment studies how two gasoline additives influence the miles per gallon obtained. Three types of gasoline were studied. The first sample received additive W, the second received additive V, and the third did not receive an additive (the control group).

If the F-test from an ANOVA for this experiment is significant, we do not know which of the three possible pairs of groups are different. MCPs can help solve this dilemma.

Multiple Comparison Considerations

Whenever MCPs are to be used, the researcher needs to contemplate the following issues.

Exploration Versus Decision-Making

When conducting exploration (or data snooping), you make several comparisons to discover the underlying factors that influence the response. In this case, you do not have a set of planned comparisons to make. In contrast, in a decision-making mode, you would try to determine which treatment is preferred. In the above example, because you do not know which factors influence gasoline additive performance, you should use the exploration mode to identify those. A decision-making emphasis would choose the gasoline that provides the highest miles per gallon.

Choosing a Comparison Procedure

You should consider two items here. First, will you know before or after experimentation which comparisons are of interest? Second, are you interested in some or all possible comparisons? Your choice of an MCP will depend on how you answer these two questions.

Error Rates

You will need to consider two types of error rates: comparisonwise and experimentwise.

1. Comparisonwise error rate. In this case, you consider each comparison of the means as if it were the only test you were conducting. This is commonly denoted as α . The conceptual unit is the comparison. Other tests that might be conducted are ignored during the calculation of the error rate. If we perform several tests, the probability of a type I error on each test is α .
2. Experimentwise, or familywise, error rate. In this situation, the error rate relates to a group of independent tests. This is the probability of making one or more type I errors in a group of independent comparisons. We will denote this error rate as α_f .

The relationship between these two error rates is:

$$\alpha_f = 1 - (1 - \alpha)^c$$

General Linear Models (GLM) for Fixed Factors

where c is the total number of comparisons in the family. The following table shows these error rates for a few values of c and α . The body of the table consists of the calculated values of α_f .

Calculated Experimentwise Error Rates

α	k				
	2	3	5	10	20
0.20	0.360	0.488	0.672	0.893	0.988
0.10	0.190	0.271	0.410	0.651	0.878
0.05	0.098	0.143	0.226	0.401	0.642
0.02	0.040	0.059	0.096	0.183	0.332
0.01	0.020	0.030	0.049	0.096	0.182

As you can see, the possibility of at least one erroneous result goes up markedly as the number of tests increases. For example, to obtain an α_f of 0.05 with a c of 5, you would need to set α to 0.01.

Multiple Comparison Procedures

The multiple comparison procedures (MCPs) considered here assume that there is independence between treatments or samples, equal variance for each treatment, and normality. In addition, unless stated otherwise, the significance tests are assumed to be two-tailed.

Let $\bar{y}_i$ represent the least squares mean of the i^{th} treatment group, $i = 1, \dots, k$. Let s^2 represent the mean square error for these least squares means based on ν degrees of freedom. As described above, simultaneous confidence intervals are of the form

$$\mathbf{c}_i' \mathbf{b} \pm c_\alpha SE(\mathbf{c}_i' \mathbf{b}),$$

where c_α is the critical value, usually selected to keep the experimentwise error rate equal to α for the collection of all comparisons. Significance tests are of the form

$$T_i = \frac{|\mathbf{c}_i' \mathbf{b}|}{SE(\mathbf{c}_i' \mathbf{b})} \geq c_\alpha.$$

Alpha

This is the α_f , or α , specified for the multiple comparison test. It may be comparisonwise or experimentwise, depending on the test. This alpha usually ranges from 0.01 to 0.10.

All-Pairs Comparison Procedures

For a factor with k levels, when comparing all possible pairs, there are $c = k(k - 1)/2$ comparisons.

Tukey-Kramer

The Tukey-Kramer method (also known as Tukey's HSD (Honest Significant Difference) method) uses the Studentized Range distribution to compute the adjustment to c_α . The Tukey-Kramer method achieves the exact alpha level (and simultaneous confidence level $(1 - \alpha)$) if the group sample sizes are equal and is conservative if the sample sizes are unequal. The Tukey-Kramer test is one of the most powerful all-pairs testing procedures and is widely used.

The Tukey-Kramer adjusted critical value for tests and simultaneous confidence intervals is

$$c_\alpha = \frac{q_{1-\alpha, k, v}}{\sqrt{2}}$$

where $q_{1-\alpha, k, v}$ is the $1 - \alpha$ quantile of the studentized range distribution.

Scheffe

This method controls the overall experimentwise error rate but is generally less powerful than the Tukey-Kramer method. Scheffe's method might be preferred if the group sample sizes are unequal or if the number of comparisons is larger than, say, 20.

Scheffe's adjusted critical value for tests and simultaneous confidence intervals is

$$c_\alpha = \sqrt{(k - 1)F_{1-\alpha, k-1, v}}$$

where $F_{1-\alpha, k-1, v}$ is the $1 - \alpha$ quantile of the F distribution with $k - 1$ numerator and v denominator degrees of freedom.

Bonferroni

This conservative method always controls the overall experimentwise error rate (alpha level), even when tests are not independent. P-values are adjusted by multiplying each individual test p-value by the number of comparisons ($c = k(k - 1)/2$) (if the result is greater than 1, then the adjusted p-value is set to 1). Simultaneous confidence limits are adjusted by simply dividing the overall alpha level by the number of comparisons (α/c) and computing each individual interval at $1 - \alpha/c$. Generally, this MCP is run after the fact to find out which pairs are different.

Bonferroni's adjusted critical value for tests and simultaneous confidence intervals is

$$c_\alpha = T_{1-\frac{\alpha}{2c}, v}$$

where $T_{1-\frac{\alpha}{2c}, v}$ is the $1 - \frac{\alpha}{2c}$ quantile of the T distribution with v degrees of freedom.

General Linear Models (GLM) for Fixed Factors

Sidak

This method is like Bonferroni's method, but is more powerful if the tests are independent.

Sidak's adjusted critical value for tests and simultaneous confidence intervals is

$$c_{\alpha} = T_{(1-\alpha)^{\frac{1}{c}}/2, v}$$

where $T_{(1-\alpha)^{\frac{1}{c}}/2, v}$ is the $(1 - \alpha)^{\frac{1}{c}}/2$ quantile of the T distribution with v degrees of freedom.

Fisher's LSD (No Adjustment)

The only difference between this test and a regular two-sample T-test is that the degrees of freedom here is based on the whole-model error degrees of freedom, not the sample sizes from the individual groups. This method is not recommended since the overall alpha level is not protected.

The unadjusted critical value for tests and simultaneous confidence intervals is

$$c_{\alpha} = T_{1-\frac{\alpha}{2}, v}$$

where $T_{1-\frac{\alpha}{2}, v}$ is the $1 - \frac{\alpha}{2}$ quantile of the T distribution with v degrees of freedom.

Each vs. Reference Value (or Control) Comparison Procedures

For a factor with k levels, when comparing each versus the reference value or control, there are $c = k - 1$ comparisons. Often, it is of interest to find only those treatments that are better (or worse) than the control, so both one- and two-sided versions of these tests are provided.

Dunnett's (*Only available for models without covariates*)

This method uses Dunnett's Range Distribution (either two- or one-sided depending on the test direction) to compute the adjustment (see Hsu(1996)). Dunnett's method controls the overall experimentwise error rate and is the most widely used method for all treatments versus control comparisons.

Dunnett's adjusted critical value for tests and simultaneous confidence intervals is

$$c_{\alpha} = q_{1-\alpha, v}$$

where $q_{1-\alpha, v}$ is the $1 - \alpha$ quantile of Dunnett's Range distribution.

Bonferroni

This conservative method always controls the overall experimentwise error rate (alpha level), even when tests are not independent. P-values are adjusted by multiplying each individual test p-value by the number of comparisons ($c = k - 1$) (if the result is greater than 1, then the adjusted p-value is set to 1).

Simultaneous confidence limits are adjusted by simply dividing the overall alpha level by the number of comparisons (α/c) and computing each individual interval at $1 - \alpha/c$. Generally, this MCP is run after the fact to find out which pairs are different.

Bonferroni's adjusted critical value for tests and simultaneous confidence intervals is

$$c_{\alpha} = T_{1-\frac{\alpha}{2c},v}$$

where $T_{1-\frac{\alpha}{2c},v}$ is the $1 - \frac{\alpha}{2c}$ quantile of the T distribution with v degrees of freedom.

Fisher's LSD (No Adjustment)

The only difference between this test and a regular two-sample T-test is that the degrees of freedom here is based on the whole-model error degrees of freedom, not the sample sizes from the individual groups. This method is not recommended since the overall alpha level is not protected.

The unadjusted critical value for tests and simultaneous confidence intervals is

$$c_{\alpha} = T_{1-\frac{\alpha}{2},v}$$

where $T_{1-\frac{\alpha}{2},v}$ is the $1 - \frac{\alpha}{2}$ quantile of the T distribution with v degrees of freedom.

Recommendations

These recommendations assume that normality and equal variance are valid.

1. Planned all-possible pairs. If you are interested in paired comparisons only and you know this in advance, use either the Bonferroni for pairs or the Tukey-Kramer MCP.
2. Unplanned all-possible pairs. Use Scheffe's MCP.
3. Each versus a control. Use Dunnett's test.

Data Structure

The data must be entered in a format that puts the responses in one column and the values of each of the factors and covariates in the other columns. An example of data that might be analyzed using this procedure is shown next. The data contains a response variable (Test), 2 fixed factors (Dose and Gender), and two covariates (PreTest and Age). The data could be analyzed as a two-way ANOVA by including only the fixed factors in the model. Or the data could be analyzed as an ANCOVA model by including the covariates along with the fixed factors.

2 Factor 2 Covariate Dataset

Test	Dose	Gender	PreTest	Age
64.1	Low	F	52.1	58
69.1	Low	F	56.4	57
69.5	Low	F	53.5	50
81.8	Low	F	72.5	30
82.4	Low	F	65.3	72
86.7	Low	F	73.6	24
87	Low	F	69.2	30
93.8	Low	F	78.5	61
.	.	.	.	.
.	.	.	.	.
.	.	.	.	.
99	High	M	84.6	74
99.5	High	M	76	70
99.5	High	M	77.2	36

Example 1 – ANCOVA Model with Two Fixed Factors and Two Covariates

This section presents an example of how to run an analysis of the data presented above. These data are contained in the 2 Factor 2 Covariate dataset.

In this example, 24 males and 24 females are randomly allocated to three dose groups: low, medium, and high. The age of each subject is recorded. Their response to a certain stimuli is recorded as a pretest. Next, the assigned dose of a certain compound is administered and their response to the stimuli is measured again.

Researchers wish to investigate how the response to the stimuli is affected by the subject's age, gender, dose, and pretest score. This can be done using a two-factor, two-covariate model. They want to allow for the possibility of a difference in variance for males versus females.

This example will run all reports and plots so that they may be documented.

Setup

To run this example, complete the following steps:

1 Open the 2 Factor 2 Covariate example dataset

- From the File menu of the NCSS Data window, select **Open Example Data**.
- Select 2 Factor 2 Covariate and click OK.

2 Specify the General Linear Models (GLM) for Fixed Factors procedure options

- Find and open the **General Linear Models (GLM) for Fixed Factors** procedure using the menus or the Procedure Navigator.
- The settings for this example are listed below and are stored in the **Example 1** settings file. To load these settings to the procedure window, click **Open Example Settings File** in the Help Center or File menu.

Variables, Model Tab

Response (Y).....**Test**
 Fixed Factor(s).....**Dose, Gender**
 Covariate(s).....**PreTest, Age**

Reports Tab

All Reports**Checked**

Plots Tab

All Plots.....**Checked**

Report Options (*in the Toolbar*)

Variable Labels.....**Column Names**

3 Run the procedure

- Click the **Run** button to perform the calculations and generate the output.

Run Summary

Run Summary

Response (Y): Test
 Fixed Factor(s): Dose, Gender
 Covariate(s): PreTest, Age
 Model: PreTest + Age + Dose + Gender

Parameter	Value	Rows	Value
R ²	0.932	Rows Processed	48
Adjusted R ²	0.9239	Rows Filtered Out	0
Coefficient of Variation	0.035	Rows with Y Missing	0
Mean Square Error	8.345328	Rows with X's Missing	0
Square Root of MSE	2.888828	Rows Used in Estimation	48
Average Percent Error	2.797	Completion Status	Normal Completion
Error Degrees of Freedom	42		

This report summarizes the results. It presents the variables used, the model, the number of rows used, and basic summary results.

R²

R^2 , officially known as the *coefficient of determination*, is defined as

$$R^2 = \frac{SS_{Model}}{SS_{Total(Adjusted)}}$$

R^2 is probably the most popular measure of how well a model fits the data. R^2 may be defined either as a ratio or a percentage. Since we use the ratio form, its values range from zero to one. A value of R^2 near zero indicates no linear relationship, while a value near one indicates a perfect linear fit. Although popular, R^2 should not be used indiscriminately or interpreted without scatter plot support. Following are some qualifications on its interpretation:

1. *Additional independent variables.* It is possible to increase R^2 by adding more independent variables, but the additional independent variables may cause an increase in the mean square error, an unfavorable situation. This usually happens when the sample size is small.
2. *Range of the independent variables.* R^2 is influenced by the range of the independent variables. R^2 increases as the range of the X 's increases and decreases as the range of the X 's decreases.
3. *Slope magnitudes.* R^2 does not measure the magnitude of the slopes.
4. *Linearity.* R^2 does not measure the appropriateness of a linear model. It measures the strength of the linear component of the model. Suppose the relationship between X and Y was a perfect sphere. Although there is a perfect relationship between the variables, the R^2 value would be zero.
5. *Predictability.* A large R^2 does not necessarily mean high predictability, nor does a low R^2 necessarily mean poor predictability.

General Linear Models (GLM) for Fixed Factors

6. *Sample size.* R^2 is highly sensitive to the number of observations. The smaller the sample size, the larger its value.

Adjusted R^2

This is an adjusted version of R^2 . The adjustment seeks to remove the distortion due to a small sample size. The formula for adjusted R^2 is

$$\bar{R}^2 = 1 - \frac{(N - 1)(1 - R^2)}{N - p - 1}$$

Coefficient of Variation

The coefficient of variation is a relative measure of dispersion, computed by dividing root mean square error by the mean of the response variable. By itself, it has little value, but it can be useful in comparative studies.

$$CV = \frac{\sqrt{MSE}}{\bar{y}}$$

Average |Percent Error|

This is the average of the absolute percent errors. It is another measure of the goodness of fit of the linear model to the data. It is calculated using the formula

$$AAPE = \frac{100 \sum_{j=1}^N \left| \frac{y_j - \hat{y}_j}{y_j} \right|}{N}$$

Note that when the response variable is zero, its predicted value is used in the denominator.

Descriptive Statistics**Descriptive Statistics**

Variable	Count	Mean	Standard Deviation	Minimum	Maximum
PreTest	48	65.80833	9.542421	50.8	84.6
Age	48	44.14583	16.45883	21	74
(Dose="Medium")	48	0.3333333	0.4763931	0	1
(Dose="High")	48	0.3333333	0.4763931	0	1
(Gender="M")	48	0.5	0.5052912	0	1
Test	48	82.60416	10.47304	63.8	99.5

For each variable, the count, arithmetic mean, standard deviation, minimum, and maximum are computed. This report is particularly useful for checking that the correct variables were selected. Recall that a factor variable with K levels is represented by $K - 1$ binary indicator variables. The reference value is not listed.

Analysis of Variance

Analysis of Variance

Source	DF	Sum of Squares	Mean Square	F-Ratio	P-Value	Significant at $\alpha = 0.05$?
Model	5	4804.675	960.9351	115.146	0.0000	Yes
PreTest	1	3407.04	3407.04	408.257	0.0000	Yes
Age	1	25.25642	25.25642	3.026	0.0892	No
Dose	2	510.8416	255.4208	30.606	0.0000	Yes
Gender	1	9.977028	9.977028	1.196	0.2805	No
Error	42	350.5038	8.345328			
Total(Adjusted)	47	5155.179	109.6847			

An analysis of variance (ANOVA) table summarizes the information related to the variation in data. PreTest is a significant covariate and Dose is a significant fixed factor. The other two variables, Age and Gender, are not significant and should probably be removed from the model.

Source

This represents a partition of the variation in Y.

DF

The degrees of freedom are the number of dimensions associated with this term. Note that each observation can be interpreted as a dimension in n -dimensional space. The degrees of freedom for the intercept, model, error, and adjusted total are 1, p , $n - p - 1$, and $n - 1$, respectively.

Sum of Squares

These are the sums of squares associated with the corresponding sources of variation. Note that these values are in terms of the dependent variable. The formulas for each are

$$SS_{Model} = \sum (\hat{y}_j - \bar{y})^2$$

$$SS_{Error} = \sum (y_j - \hat{y}_j)^2$$

$$SS_{Total} = \sum (y_j - \bar{y})^2$$

Mean Square

The mean square is the sum of squares divided by the degrees of freedom. This mean square is an estimated variance. For example, the mean square error is the estimated variance of the residuals.

F-Ratio

This is the F -statistic for testing the null hypothesis that all $\beta_j = 0$. This F -statistic has p degrees of freedom for the numerator variance and $n - p - 1$ degrees of freedom for the denominator variance.

General Linear Models (GLM) for Fixed Factors

P-Value

This is the p -value for the above F -test. The p -value is the probability that the test statistic will take on a value at least as extreme as the observed value, if the null hypothesis is true. If the p -value is less than α , say 0.05, the null hypothesis is rejected. If the p -value is greater than α , then the null hypothesis is accepted.

Significant at $\alpha = 0.05$?

This is the decision based on the p -value and the user-entered Tests Alpha value. The default is Tests Alpha = 0.05.

Model Coefficient T-Tests**Model Coefficient T-Tests**

Hypotheses Tested: $H_0: \beta(i) = 0$ vs. $H_1: \beta(i) \neq 0$

Independent Variable	Model Coefficient $b(i)$	Standard Error $Sb(i)$	Hypothesis Tests		
			T-Statistic	P-Value	Reject H_0 at $\alpha = 0.05$?
Intercept	15.43254	3.499119	4.410	0.0001	Yes
PreTest	0.9357576	0.04631231	20.205	0.0000	Yes
Age	0.04904018	0.02818953	1.740	0.0892	No
(Dose="Medium")	3.606826	1.036388	3.480	0.0012	Yes
(Dose="High")	8.151655	1.044059	7.808	0.0000	Yes
(Gender="M")	-0.9868573	0.9025587	-1.093	0.2805	No

This section reports the values and significance tests of the model coefficients.

Independent Variable

The names of the independent variables are listed here. The intercept is the value of the Y intercept.

Note that the name may become very long, especially for interaction terms. These long names may misalign the report. You can force the rest of the items to be printed on the next line by using the Stagger label ... option on the Report Options tab. This should create a better-looking report when the names are extra-long.

Model Coefficient $b(i)$

The coefficients are the least squares estimates of the parameters. The value indicates how much change in Y occurs for a one-unit change in a particular X when the remaining X 's are held constant. These coefficients are often called partial-regression coefficients since the effect of the other X 's is removed. These coefficients are the values of $b_0, b_1, \dots, b_p$.

Standard Error $Sb(i)$

The standard error of the coefficient, s_{b_j} , is the standard deviation of the estimate. It is used in hypothesis tests and confidence limits.

General Linear Models (GLM) for Fixed Factors

T-Statistic

This is the t-test value for testing the hypothesis that $\beta_j = 0$ versus the alternative that $\beta_j \neq 0$ after removing the influence of all other X 's. This t -value has $n - p - 1$ degrees of freedom.

To test for a value other than zero, use the formula below. There is an easier way to test hypothesized values using confidence limits. See the discussion below under Confidence Limits. The formula for the t -test is

$$t_j = \frac{b_j - \beta_j^*}{s_{b_j}}$$

P-Value

This is the p -value for the significance test of the coefficient. The p -value is the probability that this t -statistic will take on a value at least as extreme as the observed value, assuming that the null hypothesis is true (i.e., the coefficient estimate is equal to zero). If the p -value is less than alpha, say 0.05, the null hypothesis of equality is rejected. This p -value is for a two-tail test.

Reject H0 at $\alpha = 0.05$?

This is the decision based on the p -value and the user-entered Tests Alpha value. The default is Tests Alpha = 0.05.

Model Coefficient Confidence Intervals

Model Coefficient Confidence Intervals

Independent Variable	Model Coefficient b(i)	Standard Error Sb(i)	95% Confidence Interval Limits for $\beta(i)$	
			Lower	Upper
Intercept	15.43254	3.499119	8.371028	22.49405
PreTest	0.9357576	0.04631231	0.8422955	1.02922
Age	0.04904018	0.02818953	-0.0078486	0.105929
(Dose="Medium")	3.606826	1.036388	1.51531	5.698341
(Dose="High")	8.151655	1.044059	6.044658	10.25865
(Gender="M")	-0.9868573	0.9025587	-2.808295	0.8345799

Note: The T-Value used to calculate these confidence interval limits was 2.018.

This section reports the values and confidence intervals of the model coefficients.

Independent Variable

The names of the independent variables are listed here. The intercept is the value of the Y intercept.

Note that the name may become very long, especially for interaction terms. These long names may misalign the report. You can force the rest of the items to be printed on the next line by using the Stagger label ... option on the Report Options tab. This should create a better-looking report when the names are extra-long.

Model Coefficient

The coefficients are the least squares estimates of the parameters. The value indicates how much change in Y occurs for a one-unit change in a particular X when the remaining X 's are held constant. These coefficients are often called partial-regression coefficients since the effect of the other X 's is removed. These coefficients are the values of $b_0, b_1, \dots, b_p$.

Standard Error

The standard error of the coefficient, s_{b_j} , is the standard deviation of the estimate. It is used in hypothesis tests and confidence limits.

95% Confidence Interval Limits for $\beta(i)$

These are the lower and upper values of a $100(1 - \alpha)\%$ interval estimate for β_j based on a t -distribution with $n-p-1$ degrees of freedom. This interval estimate assumes that the residuals for the regression model are normally distributed.

These confidence limits may be used for significance testing values of β_j other than zero. If a specific value is not within this interval, it is significantly different from that value. Note that these confidence limits are set up as if you are interested in each regression coefficient separately.

The formulas for the lower and upper confidence limits are:

$$b_j \pm t_{1-\alpha/2, n-p-1} s_{b_j}$$

Note: The T-Value ...

This is the value of $t_{1-\alpha/2, n-p-1}$ used to construct the confidence limits.

Least Squares Means

Least Squares Means

Error Degrees of Freedom (DF): 42

Means Calculated at: PreTest = 65.80833, Age = 44.14583 (Means)

				95% Confidence Interval Limits for the Mean	
Name	Count	Least Squares Mean	Standard Error	Lower	Upper
Intercept					
All	48	82.60416	0.4169664	81.76270	83.44564
Dose					
Low	16	78.68467	0.7313797	77.20869	80.16066
Medium	16	82.29150	0.7288142	80.82069	83.76231
High	16	86.83633	0.7324548	85.35818	88.31448
Gender					
F	24	83.09760	0.6144217	81.85764	84.33755
M	24	82.11074	0.6144217	80.87078	83.35069

General Linear Models (GLM) for Fixed Factors

This section reports the least squares means and associated confidence intervals. In this example, the least squares means for all factor variables are calculated at the means of the covariates, PreTest = 65.80833 and Age = 44.14583. The results are based on $n-p-1 = 42$ degrees of freedom for error.

Name

The names of the fixed factor variables and their levels are listed here. The intercept is the value of the Y intercept.

Note that the name may become very long, especially for interaction terms. These long names may misalign the report. You can force the rest of the items to be printed on the next line by using the Stagger label ... option on the Report Options tab. This should create a better-looking report when the names are extra-long.

Count

This column specifies the number of observations encountered at each level of each fixed factor (or combination in the case of interaction terms).

Least Squares Mean

This is the least squares mean estimate, $\hat{\mu}_j$. The least squares means are adjusted for the other terms in the model, both the other fixed factors and the covariates. In balanced designs with no covariates, the least squares group means will be equal to the raw group means. In unbalanced designs or when covariates are present, the least squares means usually are different from the raw means.

Standard Error

The standard error of the mean, $SE(\hat{\mu}_j)$, is the standard deviation of the estimate. It is used in hypothesis tests and confidence limits.

95% Confidence Interval Limits for the Mean

These are the lower and upper values of a $100(1 - \alpha)\%$ interval estimate for the mean, μ_j , based on a t -distribution with $n-p-1$ degrees of freedom.

The formulas for the lower and upper confidence limits are:

$$\hat{\mu}_j \pm t_{1-\frac{\alpha}{2}, n-p-1} \times SE(\hat{\mu}_j)$$

Least Squares Means with Hypothesis Tests of H0: Mean = 0

Least Squares Means with Hypothesis Tests of H0: Mean = 0

Error Degrees of Freedom (DF): 42

Means Calculated at: PreTest = 65.80833, Age = 44.14583 (Means)

Hypotheses Tested: H0: Mean = 0 vs. H1: Mean $\neq$ 0

		Least Squares Mean	Standard Error	Hypothesis Tests		
Name	Count			T-Statistic	P-Value	Reject H0 at $\alpha = 0.05$?
Intercept						
All	48	82.60416	0.4169664	198.107	0.0000	Yes
Dose						
Low	16	78.68467	0.7313797	107.584	0.0000	Yes
Medium	16	82.29150	0.7288142	112.911	0.0000	Yes
High	16	86.83633	0.7324548	118.555	0.0000	Yes
Gender						
F	24	83.09760	0.6144217	135.245	0.0000	Yes
M	24	82.11074	0.6144217	133.639	0.0000	Yes

This section reports the least squares means and associated hypothesis tests. In this example, the least squares means for all factor variables are calculated at the means of the covariates, PreTest = 65.80833 and Age = 44.14583. The results are based on $n - p - 1 = 42$ degrees of freedom for error.

Name

The names of the fixed factor variables and their levels are listed here. The intercept is the value of the Y intercept.

Note that the name may become very long, especially for interaction terms. These long names may misalign the report. You can force the rest of the items to be printed on the next line by using the Stagger label ... option on the Report Options tab. This should create a better-looking report when the names are extra-long.

Count

This column specifies the number of observations encountered at each level of each fixed factor (or combination in the case of interaction terms).

Least Squares Mean

This is the least squares mean estimate, $\hat{\mu}_j$. The least squares means are adjusted for the other terms in the model, both the other fixed factors and the covariates. In balanced designs with no covariates, the least squares group means will be equal to the raw group means. In unbalanced designs or when covariates are present, the least squares means usually are different from the raw means.

Standard Error

The standard error of the mean, $SE(\hat{\mu}_j)$, is the standard deviation of the estimate. It is used in hypothesis tests and confidence limits.

General Linear Models (GLM) for Fixed Factors

T-Statistic

This is the t -test value for testing the hypothesis that the mean is equal to 0 versus the alternative that it is not equal to 0. This t -value has $n-p-1$ degrees of freedom and is calculated as

$$t_j = \frac{\hat{\mu}_j}{SE(\hat{\mu}_j)}$$

P-Value

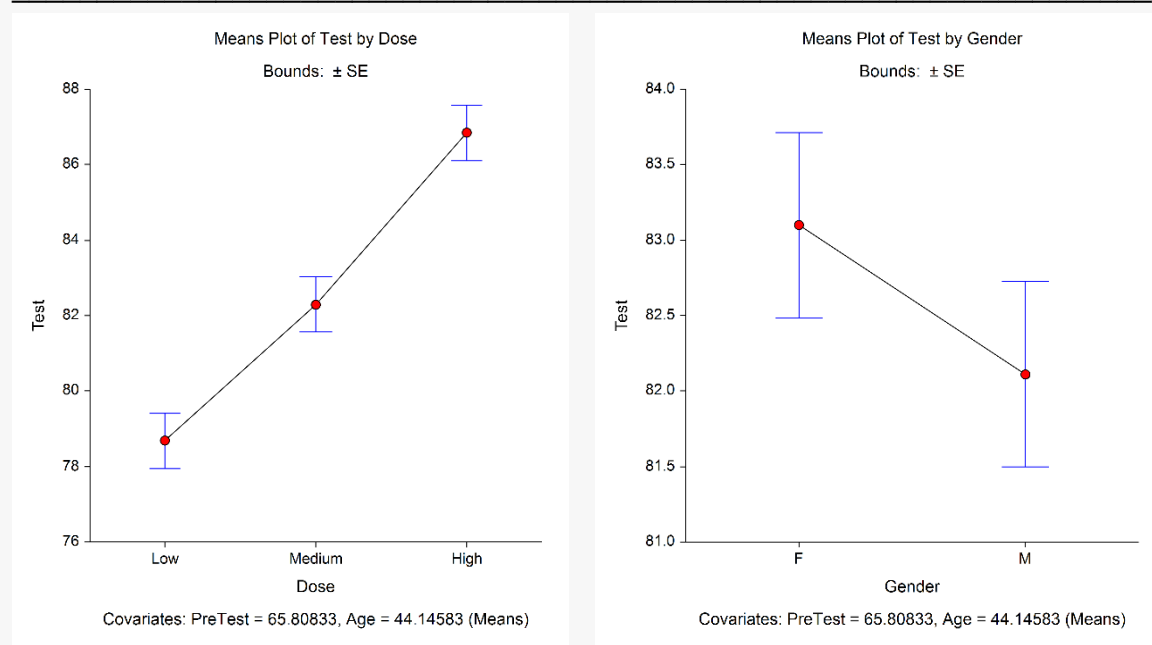
This is the p -value for the significance test of the mean. The p -value is the probability that this t -statistic will take on a value at least as extreme as the observed value, if the null hypothesis is true (i.e., the mean estimate is equal to zero). If the p -value is less than alpha, say 0.05, the null hypothesis of equality is rejected. This p -value is for a two-tail test.

Reject H0 at $\alpha = 0.05$?

This is the decision based on the p -value and the user-entered Tests Alpha value. The default is Tests Alpha = 0.05.

Means Plots

Means Plots



The means plots display the least squares means along with user-selected variability lines, in this case $\pm$ SE.

All-Pairs Comparisons of Least Squares Means

All-Pairs Comparisons of Least Squares Means

Error Degrees of Freedom (DF): 42

Means Calculated at: PreTest = 65.80833, Age = 44.14583 (Means)

Multiple Comparison Type: Tukey-Kramer

Hypotheses Tested: H0: Diff = 0 vs. H1: Diff $\neq$ 0

Comparison	Least Squares Mean Difference	Standard Error	Hypothesis Tests			
			T-Statistic	P-Value		Reject H0 at $\alpha = 0.05$?†
				Unadjusted	Adjusted*	
Dose [3 Comparisons]						
Low - Medium	-3.606826	1.036388	-3.480	0.0012	0.0033	Yes
Low - High	-8.151655	1.044059	-7.808	0.0000	0.0000	Yes
Medium - High	-4.544829	1.038663	-4.376	0.0001	0.0002	Yes
Gender [1 Comparison]						
F - M	0.9868573	0.9025587	1.093	0.2805	0.2805	No

* Adjusted p-values are computed using the number of comparisons and the adjustment type (Tukey-Kramer).

† Rejection decisions are based on adjusted p-values.

This section reports the least squares mean differences and associated multiple comparison hypothesis tests. In this example, the least squares means for all factor variables are calculated at the means of the covariates, PreTest = 65.80833 and Age = 44.14583. The results are based on $n-p-1 = 42$ degrees of freedom for error.

You should only consider tests for factors that were found significant in the ANOVA table. In this example, only Dose was significant according to the ANOVA F-test so you should ignore the results for Gender. All Dose level means (Low, Medium, and High) are found here to be significantly different from one another based on the Tukey-Kramer-adjusted p -values, which are based on 3 comparisons.

Comparison

The names of the fixed factor variables and the number of comparisons within each factor, along with the individual comparisons are listed here. The multiple comparison adjustment within each factor is based on the selected Multiple Comparison Adjustment Type and the number of simultaneous comparisons.

Note that the name may become very long, especially for interaction terms. These long names may misalign the report. You can force the rest of the items to be printed on the next line by using the Stagger label ... option on the Report Options tab. This should create a better-looking report when the names are extra-long.

Least Squares Mean Difference

This is the least squares mean difference estimate for group i minus group j , $\hat{\mu}_i - \hat{\mu}_j$.

Standard Error

The least squares mean difference estimate, $SE(\hat{\mu}_i - \hat{\mu}_j)$, is the standard deviation of the estimate. It is used in hypothesis tests and confidence limits.

General Linear Models (GLM) for Fixed Factors

T-Statistic

This is the t-test value for testing the hypothesis that the mean difference is equal to 0 versus the alternative that it is not equal to 0. This t -value has $n-p-1$ degrees of freedom and is calculated as

$$t_j = \frac{\hat{\mu}_i - \hat{\mu}_j}{SE(\hat{\mu}_i - \hat{\mu}_j)}$$

Unadjusted P-Value

This is the unadjusted p -value for the significance test of the mean difference and assumes no multiple comparisons. This p -value is valid if you are only interested in one of the comparisons. You can use this p -value to make your own adjustment (e.g., Bonferroni) if needed.

Adjusted P-Value

This is the adjusted p -value for the significance test of the mean difference. Adjusted p -values are computed using the number of comparisons within each fixed factor and the multiple comparison adjustment type, which in this example is Tukey-Kramer.

Reject H0 at $\alpha = 0.05$?

This is the decision based on the adjusted p -value and the user-entered Tests Alpha value. The default is Tests Alpha = 0.05.

Simultaneous Confidence Intervals for All-Pairs Comparisons of Least Squares Means

Simultaneous Confidence Intervals for All-Pairs Comparisons of Least Squares Means

Error Degrees of Freedom (DF): 42

Means Calculated at: PreTest = 65.80833, Age = 44.14583 (Means)

Multiple Comparison Type: Tukey-Kramer

Comparison	Least Squares Mean Difference	Standard Error	95% Simultaneous Confidence Interval Limits*	
			Lower	Upper
Dose [3 Comparisons]				
Low - Medium	-3.606826	1.036388	-6.124794	-1.088858
Low - High	-8.151655	1.044059	-10.68826	-5.615049
Medium - High	-4.544829	1.038663	-7.068325	-2.021334
Gender [1 Comparison]				
F - M	0.9868573	0.9025587	-0.8345799	2.808295

* Confidence interval limits are adjusted based on the number of comparisons and the adjustment type (Tukey-Kramer).

This section reports the least squares mean differences and associated multiple comparison simultaneous confidence intervals. In this example, the least squares means for all factor variables are calculated at the means of the covariates, PreTest = 65.80833 and Age = 44.14583. The results are based on $n - p - 1 = 42$ degrees of freedom for error.

General Linear Models (GLM) for Fixed Factors

Since the simultaneous confidence intervals are adjusted for the multiplicity of tests, there is a one-to-one correspondence between the intervals and the hypothesis tests--- all differences for which the $100(1 - \alpha)\%$ simultaneous confidence interval does not include zero will be significant at level α . As you can see, the 95% simultaneous confidence intervals for all Dose level mean differences do not include 0, so the corresponding tests (Low, Medium, and High) are found here to be significantly different from one another.

Comparison

The names of the fixed factor variables and the number of comparisons within each factor, along with the individual comparisons are listed here. The multiple comparison adjustment within each factor is based on the selected Multiple Comparison Adjustment Type and the number of simultaneous comparisons.

Note that the name may become very long, especially for interaction terms. These long names may misalign the report. You can force the rest of the items to be printed on the next line by using the Stagger label ... option on the Report Options tab. This should create a better-looking report when the names are extra-long.

Least Squares Mean Difference

This is the least squares mean difference estimate for group i minus group j , $\hat{\mu}_i - \hat{\mu}_j$.

Standard Error

The least squares mean difference estimate, $SE(\hat{\mu}_i - \hat{\mu}_j)$, is the standard deviation of the estimate. It is used in hypothesis tests and confidence limits.

95% Simultaneous Confidence Interval Limits

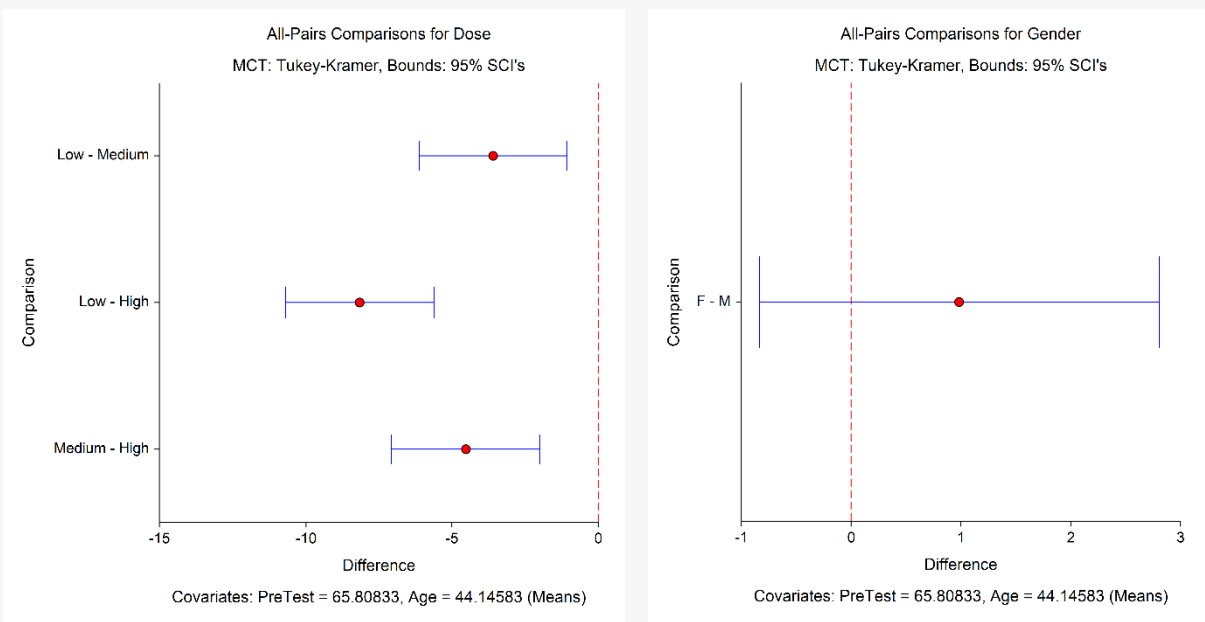
These are the lower and upper values of a $100(1 - \alpha)\%$ simultaneous confidence interval estimates for the mean difference, $\mu_i - \mu_j$. The formulas for the lower and upper confidence limits are:

$$\hat{\mu}_i - \hat{\mu}_j \pm c_\alpha \times SE(\hat{\mu}_i - \hat{\mu}_j)$$

where c_α is the adjusted critical value, usually selected to keep the experimentwise error rate equal to α for the collection of all comparisons.

All-Pairs Comparisons Plots

All-Pairs Comparisons Plots



These multiple comparison plots display the mean differences along with 95% simultaneous confidence intervals. All comparisons for which the interval does not contain zero are significant.

Each vs. Reference Value Comparisons of Least Squares Means

Each vs. Reference Value Comparisons of Least Squares Means

Error Degrees of Freedom (DF): 42
Means Calculated at: PreTest = 65.80833, Age = 44.14583 (Means)
Multiple Comparison Type: Bonferroni
Hypotheses Tested: H0: Diff = 0 vs. H1: Diff $\neq$ 0

Comparison	Least Squares Mean Difference	Standard Error	Hypothesis Tests			
			T-Statistic	P-Value		Reject H0 at $\alpha = 0.05$?†
				Unadjusted	Adjusted*	
Dose [2 Comparisons]						
Medium - Low	3.606826	1.036388	3.480	0.0012	0.0024	Yes
High - Low	8.151655	1.044059	7.808	0.0000	0.0000	Yes
Gender [1 Comparison]						
M - F	-0.9868573	0.9025587	-1.093	0.2805	0.2805	No

* Adjusted p-values are computed using the number of comparisons and the adjustment type (Bonferroni).
† Rejection decisions are based on adjusted p-values.

This section reports the least squares mean differences and associated multiple comparison hypothesis tests for each value versus the reference. In this example, the least squares means for all factor variables

General Linear Models (GLM) for Fixed Factors

are calculated at the means of the covariates, PreTest = 65.80833 and Age = 44.14583. The results are based on $n - p - 1 = 42$ degrees of freedom for error.

You should only consider tests for factors that were found significant in the ANOVA table. In this example, only Dose was significant according to the ANOVA F-test so you should ignore the results for Gender. Both Dose means (Medium and High) are found here to be significantly different from the reference Dose (Low) based on the Bonferroni-adjusted p -values, which are based on 2 comparisons.

Often, Dunnett's test is used to compare individual groups to the control, but when there are covariates in the model, only the Bonferroni adjustment is available in NCSS.

Simultaneous Confidence Intervals for Each vs. Reference Value Comparisons of Least Squares Means

Simultaneous Confidence Intervals for Each vs. Reference Value Comparisons of Least Squares Means

Error Degrees of Freedom (DF): 42

Means Calculated at: PreTest = 65.80833, Age = 44.14583 (Means)

Multiple Comparison Type: Bonferroni

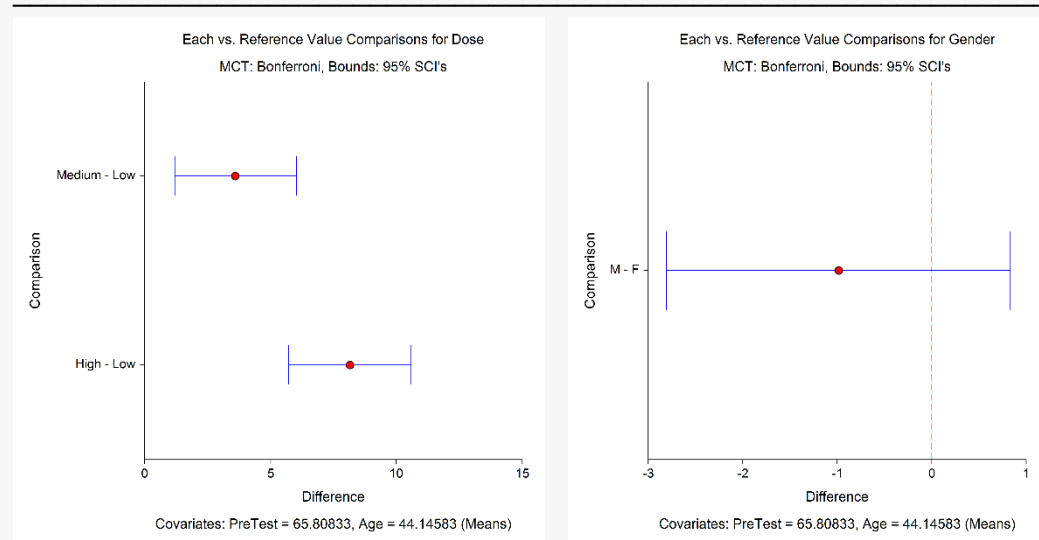
Comparison	Least Squares Mean Difference	Standard Error	95% Simultaneous Confidence Interval Limits*	
			Lower	Upper
Dose [2 Comparisons]				
Medium - Low	3.606826	1.036388	1.197618	6.016034
High - Low	8.151655	1.044059	5.724614	10.5787
Gender [1 Comparison]				
M - F	-0.9868573	0.9025587	-2.808295	0.8345799

* Confidence interval limits are adjusted based on the number of comparisons and the adjustment type (Bonferroni).

This section reports the least squares mean differences and associated multiple comparison simultaneous confidence intervals. In this example, the least squares means for all factor variables are calculated at the means of the covariates, PreTest = 65.80833 and Age = 44.14583. The results are based on $n - p - 1 = 42$ degrees of freedom for error.

Each vs. Reference Value Comparisons Plots

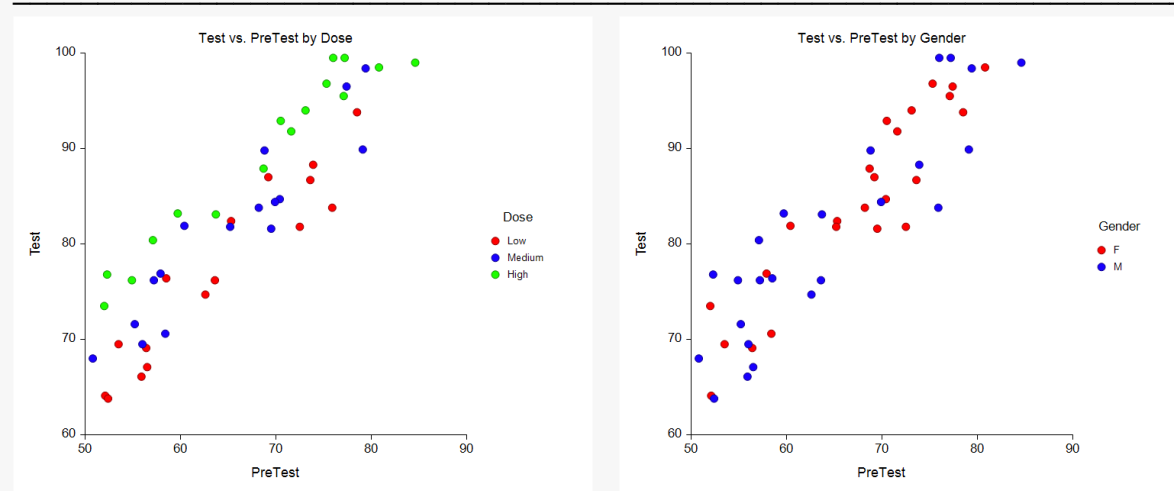
Each vs. Reference Value Comparisons Plots



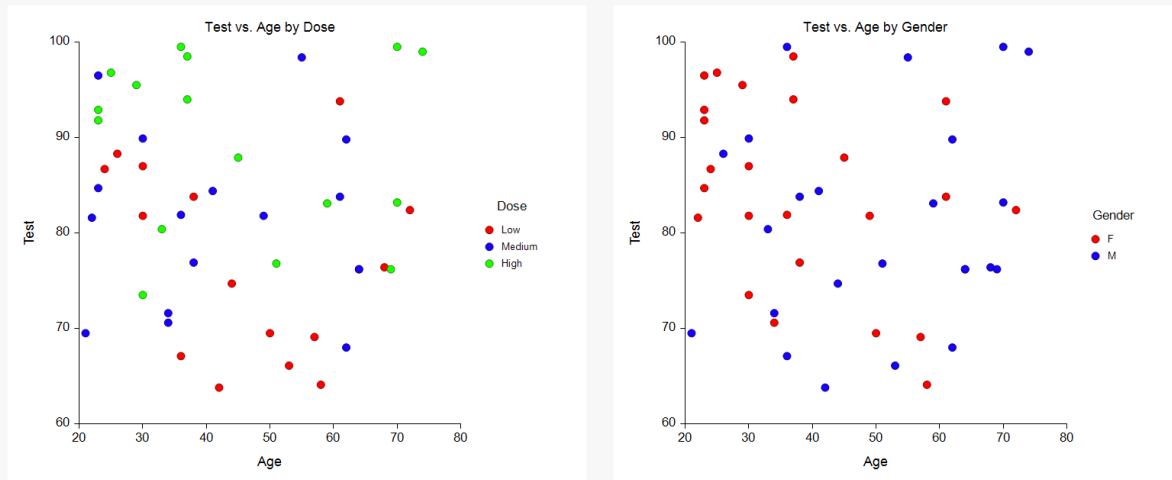
These multiple comparison plots display the mean differences along with 95% simultaneous confidence intervals. All comparisons for which the interval does not contain zero are significant.

Response vs. Covariate by Factor Scatter Plots

Response vs. Covariate by Factor Scatter Plots



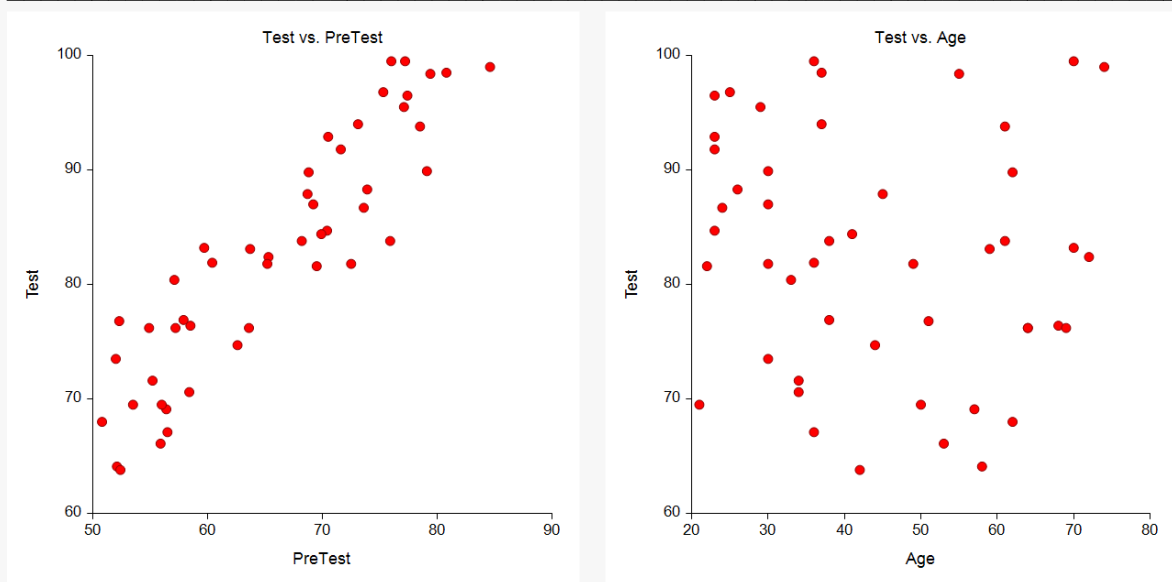
General Linear Models (GLM) for Fixed Factors



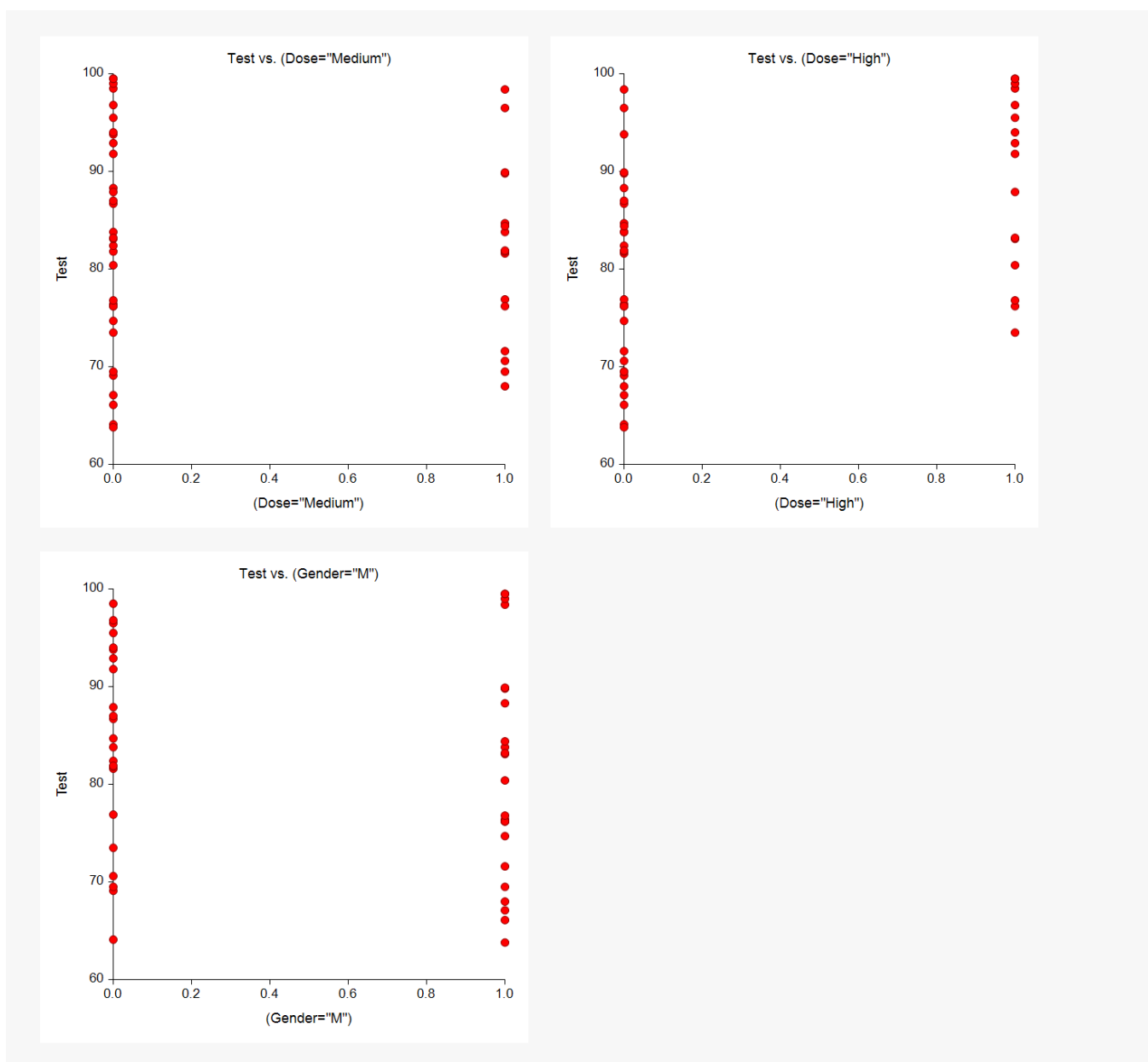
The plots in this section present the response variable plotted against each covariate by each factor variable. Use these plots to look for outliers, curvilinear relationships, or other anomalies.

Response (Y) vs. X Scatter Plots

Response (Y) vs. X Scatter Plots



General Linear Models (GLM) for Fixed Factors



The first plots in this section present the response variable plotted against the values in each column of the design matrix, X . Use these plots to look for outliers, curvilinear relationships, or other anomalies.

Residual Normality Assumption Tests

Residual Normality Assumption Tests

Normality Test	Test Statistic	P-Value	Reject Residual Normality at $\alpha = 0.2$?
Shapiro-Wilk	0.978	0.5032	No
Anderson-Darling	0.435	0.2996	No
D'Agostino Skewness	-0.057	0.9543	No
D'Agostino Kurtosis	-1.302	0.1928	Yes
D'Agostino Omnibus (Skewness and Kurtosis)	1.700	0.4275	No

This report gives the results of applying several normality tests to the residuals. The Shapiro-Wilk test is probably the most popular, so it is given first. These tests are discussed in detail in the Normality Tests section of the Descriptive Statistics procedure.

Graphic Residual Analysis

The residuals can be graphically analyzed in numerous ways. You should examine all of the basic residual graphs: the histogram, the density trace, the normal probability plot, the scatter plot of the residuals versus the predicted value of the dependent variable, and the scatter plot of the residuals versus each of the independent variables.

For the basic scatter plots of residuals versus either the predicted values of Y or the independent variables, Hoaglin (1983) explains that there are several patterns to look for. You should note that these patterns are very difficult, if not impossible, to recognize for small data sets.

Point Cloud

A point cloud, basically in the shape of a rectangle or a horizontal band, would indicate no relationship between the residuals and the variable plotted against them. This is the preferred condition.

Wedge

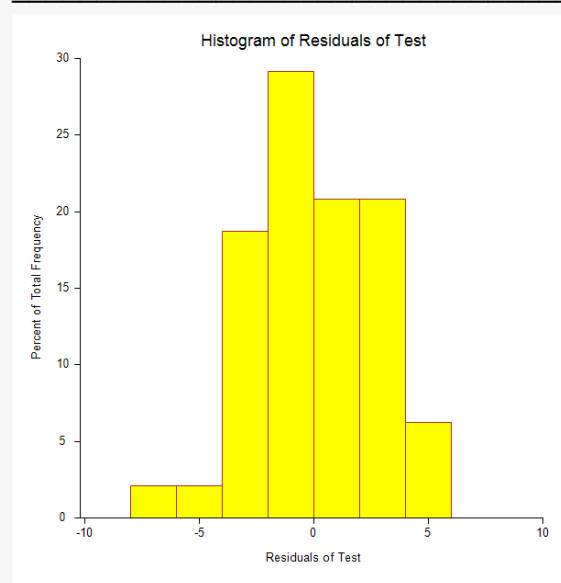
An increasing or decreasing wedge would be evidence that there is increasing or decreasing (non-constant) variation. A transformation of Y may correct the problem.

Bowtie

This is similar to the wedge above in that the residual plot shows a decreasing wedge in one direction while simultaneously having an increasing wedge in the other direction. A transformation of Y may correct the problem.

Histogram of Residuals

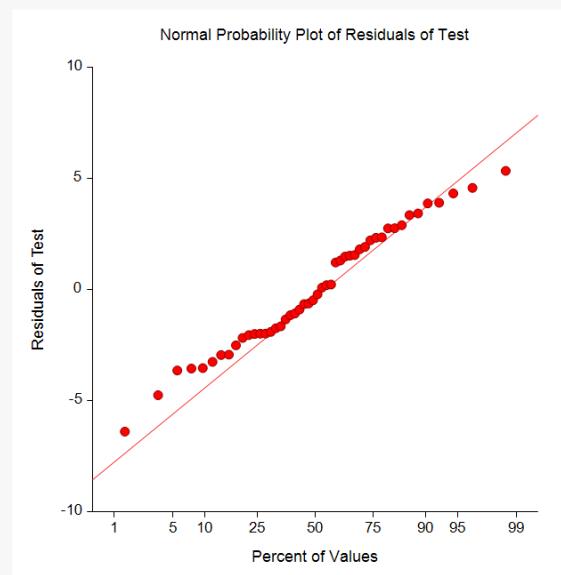
Residual Analysis Plots



The purpose of the histogram and density trace of the residuals is to evaluate whether they are normally distributed. Unless you have a large sample size, it is best not to rely on the histogram for visually evaluating normality of the residuals. The better choice would be the normal probability plot.

Probability Plot of Residuals

Residual Analysis Plots



If the residuals are normally distributed, the data points of the normal probability plot will fall along a straight line through the origin with a slope of 1.0. Major deviations from this ideal picture reflect

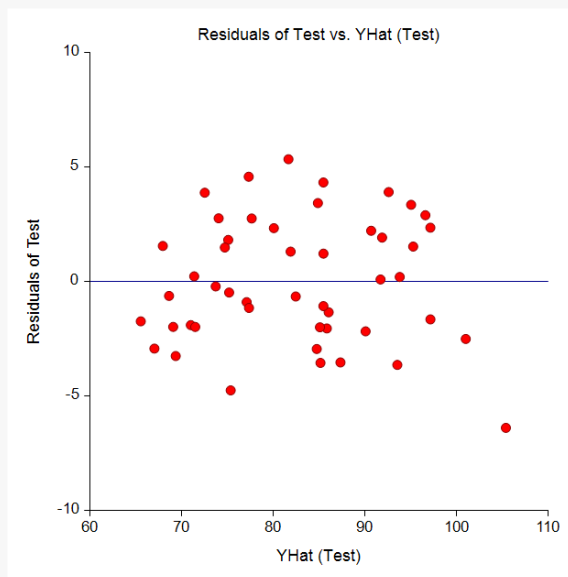
General Linear Models (GLM) for Fixed Factors

departures from normality. Stragglers at either end of the normal probability plot indicate outliers, curvature at both ends of the plot indicates long or short distributional tails, convex or concave curvature indicates a lack of symmetry, and gaps or plateaus or segmentation in the normal probability plot may require a closer examination of the data or model. Of course, use of this graphic tool with very small sample sizes is not recommended.

If the residuals are not normally distributed, then the t-tests on regression coefficients, the F-tests, and any interval estimates are not valid. This is a critical assumption to check.

Residuals vs. Yhat (Predicted) Plot

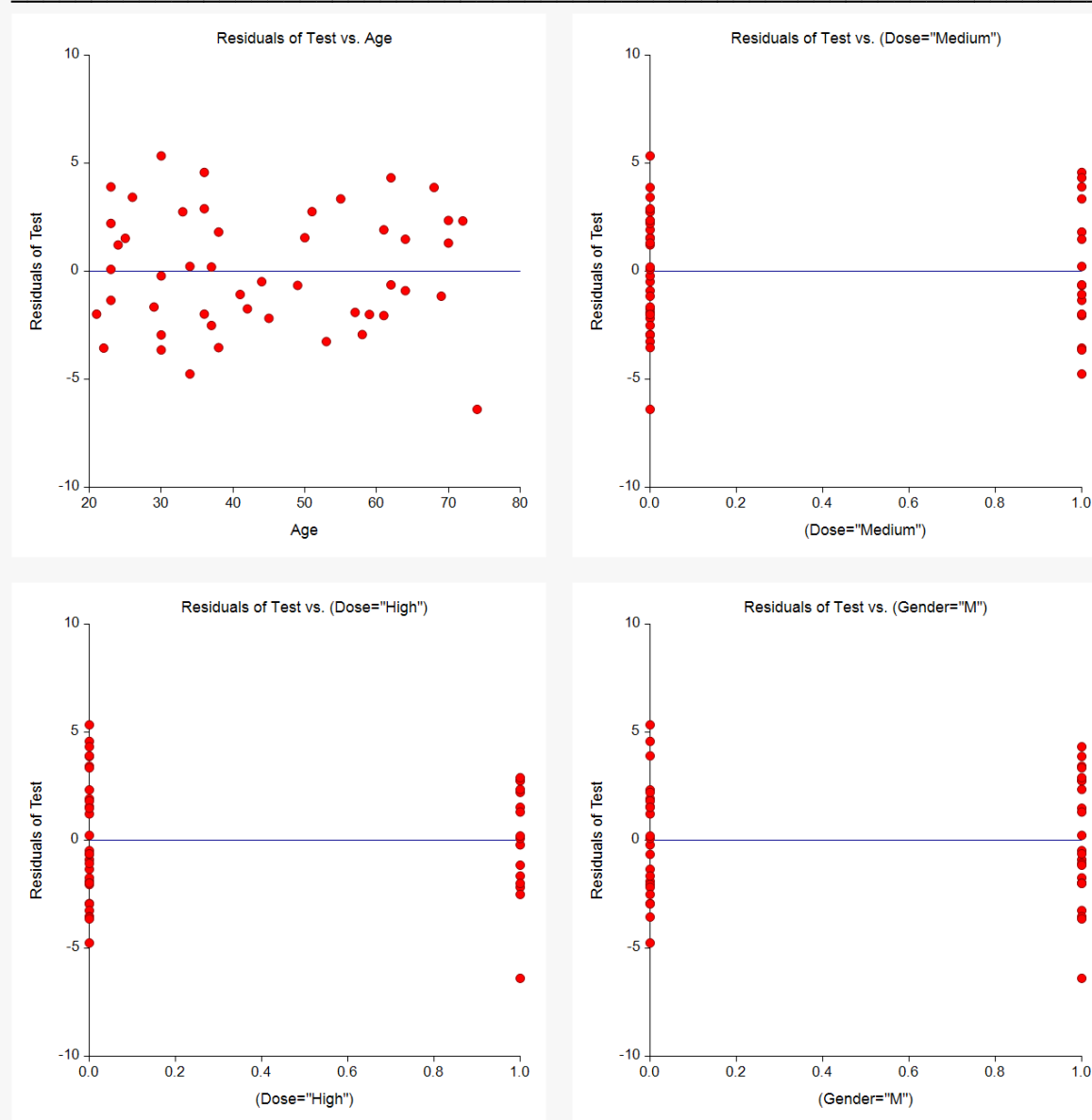
Residual Analysis Plots



This plot should always be examined. The preferred pattern to look for is a point cloud or a horizontal band. A wedge or bowtie pattern is an indicator of non-constant variance, a violation of a critical assumption. The sloping or curved band signifies inadequate specification of the model. The sloping band with increasing or decreasing variability suggests non-constant variance and inadequate specification of the model.

Residuals vs. X Plots

Residual Analysis Plots



These are scatter plots of the residuals versus each independent variable. Again, the preferred pattern is a rectangular shape or point cloud. Any other nonrandom pattern may require a redefining of the model.

Residuals List

Residuals List

Row	Test		Residual	Absolute Percent Error	Sqrt(MSE) Without This Row
	Actual	Predicted			
1	64.1	67.02984	-2.9298360	4.571	2.881651
2	69.1	71.00455	-1.9045540	2.756	2.906575
3	69.5	67.94758	1.5524250	2.234	2.912372
4	81.8	84.74616	-2.9461650	3.602	2.883133
5	82.4	80.06840	2.3316010	2.830	2.896264
6	86.7	85.48125	1.2187430	1.406	2.916780
7	87.0	81.65816	5.3418350	6.140	2.788943
8	93.8	91.88096	1.9190440	2.046	2.905113
9	63.8	65.53906	-1.7390630	2.726	2.909403
10	66.1	69.35365	-3.2536570	4.922	2.874725
11	67.1	69.08143	-1.9814280	2.953	2.905100
12	74.7	75.18187	-0.4818707	0.645	2.922780
13	76.2	77.09843	-0.8984319	1.179	2.920119
14	76.4	72.52223	3.8777710	5.076	2.852854
15	83.8	87.33321	-3.5332050	4.216	2.862403
16	88.3	84.87321	3.4267920	3.881	2.863982
17	70.6	75.35497	-4.7549700	6.735	2.816864
18	76.9	75.08325	1.8167480	2.362	2.908450
19	81.6	85.15340	-3.5533970	4.355	2.864755
20	81.8	82.45373	-0.6537243	0.799	2.921861
.	.	.	.	.	.
.	.	.	.	.	.
.	.	.	.	.	.

This section reports on the sample residuals, or *e*'s.

Actual

This is the actual value of *Y*.

Predicted

The predicted value of *Y* using the values of the independent variables given on this row.

Residual

This is the error in the predicted value. It is equal to the *Actual* minus the *Predicted*.

Absolute Percent Error

This is percentage that the absolute value of the *Residual* is of the *Actual* value. Scrutinize rows with the large percent errors.

Sqrt(MSE) Without This Row

This is the value of the square root of the mean square error that is obtained if this row is deleted. A perusal of this statistic for all observations will highlight observations that have an inflationary impact on mean square error and could be outliers.

Predicted Values with Confidence Intervals for the Means

Predicted Values with Confidence Intervals for the Means

Row	Test		Standard Error of Predicted	95% Confidence Interval Limits for the Mean	
	Actual	Predicted		Lower	Upper
1	64.1	67.02984	1.1011910	64.80754	69.25213
2	69.1	71.00455	1.0065210	68.97331	73.03580
3	69.5	67.94758	1.0096140	65.91009	69.98506
4	81.8	84.74616	0.9339331	82.86141	86.63092
5	82.4	80.06840	1.2050600	77.63649	82.50031
6	86.7	85.48125	1.0096200	83.44376	87.51875
7	87.0	81.65816	0.8993402	79.84322	83.47311
8	93.8	91.88096	1.2166660	89.42562	94.33629
9	63.8	65.53906	1.0190780	63.48248	67.59565
10	66.1	69.35365	0.8832454	67.57120	71.13612
11	67.1	69.08143	1.0157530	67.03156	71.13130
12	74.7	75.18187	0.8710524	73.42402	76.93973
13	76.2	77.09843	0.8955054	75.29123	78.90563
14	76.4	72.52223	0.9391295	70.62699	74.41747
15	83.8	87.33321	1.0946160	85.12418	89.54223
16	88.3	84.87321	1.2028300	82.44580	87.30061
17	70.6	75.35497	0.9229798	73.49232	77.21762
18	76.9	75.08325	0.9291070	73.20824	76.95827
19	81.6	85.15340	0.9119249	83.31306	86.99374
20	81.8	82.45373	0.9209040	80.59527	84.31219
.	.	.	.	.	.
.	.	.	.	.	.
.	.	.	.	.	.

Confidence intervals for the mean response of Y given specific levels for the factor and covariate variables are provided here. It is important to note that violations of any assumptions will invalidate these interval estimates.

Actual

This is the actual value of Y .

Predicted

The predicted value of Y . It is predicted using the values of the factor and covariate variables for this row. If the input data had all factor and covariate values but no value for Y , the predicted value is still provided.

Standard Error of Predicted

This is the standard error of the mean response for the specified values of the factor and covariate variables. Note that this value is not constant for all variable values. In fact, it is a minimum at the average value of each factor and covariate variable.

95% Confidence Interval Limits for the Mean

These are the lower and upper limits of a 95% confidence interval estimate of the mean of Y for this observation.

Predicted Values with Prediction Intervals for Individuals

Predicted Values with Prediction Intervals for Individuals

Row	Test		Standard Error of Predicted	95% Prediction Interval Limits for an Individual	
	Actual	Predicted		Lower	Upper
1	64.1	67.02984	3.091593	60.79075	73.26892
2	69.1	71.00455	3.059152	64.83093	77.17817
3	69.5	67.94758	3.060171	61.77190	74.12325
4	81.8	84.74616	3.036043	78.61918	90.87315
5	82.4	80.06840	3.130095	73.75161	86.38519
6	86.7	85.48125	3.060173	79.30558	91.65694
7	87.0	81.65816	3.025581	75.55229	87.76404
8	93.8	91.88096	3.134582	85.55511	98.20680
9	63.8	65.53906	3.063307	59.35706	71.72107
10	66.1	69.35365	3.020836	63.25736	75.44995
11	67.1	69.08143	3.062202	62.90165	75.26120
12	74.7	75.18187	3.017293	69.09273	81.27102
13	76.2	77.09843	3.024443	70.99486	83.20200
14	76.4	72.52223	3.037646	66.39201	78.65244
15	83.8	87.33321	3.089257	81.09883	93.56758
16	88.3	84.87321	3.129237	78.55815	91.18826
17	70.6	75.35497	3.032692	69.23475	81.47519
18	76.9	75.08325	3.034562	68.95926	81.20724
19	81.6	85.15340	3.029346	79.03993	91.26686
20	81.8	82.45373	3.032061	76.33478	88.57267
.	.	.	.	.	.
.	.	.	.	.	.
.	.	.	.	.	.

A prediction interval for the individual response of Y given specific values of the factor and covariate variables is provided here for each row.

Actual

This is the actual value of Y .

Predicted

The predicted value of Y . It is predicted using the values of the factor and covariate variables for this row. If the input data had all factor and covariate values but no value for Y , the predicted value is still provided.

Standard Error of Predicted

This is the standard error of the mean response for the specified values of the factor and covariate variables. Note that this value is not constant for all variable values. In fact, it is a minimum at the average value of each factor and covariate variable.

95% Prediction Interval Limits for the Mean

These are the lower and upper limits of a 95% prediction interval of the individual value of Y for this observation.

Example 2 – ANCOVA Interaction Model with One Fixed Factor and One Covariate (with One-Sided Multiple Comparison Tests and Simultaneous Confidence Intervals)

In Example 1 we analyzed a two-factor, two-covariate model using the 2 Factor 2 Covariate dataset and found that the fixed factor Dose and the covariate PreTest were significant predictors of the response, Test. The fixed factor Gender and covariate Age were not significant. Let's now remove those two terms from the model. We'll test the equal slopes assumption of ANCOVA by testing the Dose*PreTest interaction. Let's also calculate the means at when the covariate, PreTest, is equal to 70, 75, and 80.

Setup

To run this example, complete the following steps:

1 Open the 2 Factor 2 Covariate example dataset

- From the File menu of the NCSS Data window, select **Open Example Data**.
- Select 2 Factor 2 Covariate and click OK.

2 Specify the General Linear Models (GLM) for Fixed Factors procedure options

- Find and open the **General Linear Models (GLM) for Fixed Factors** procedure using the menus or the Procedure Navigator.
- The settings for this example are listed below and are stored in the **Example 2a** settings file. To load these settings to the procedure window, click **Open Example Settings File** in the Help Center or File menu.

Variables, Model Tab

Response (Y).....**Test**
 Fixed Factor(s).....**Dose**
 Covariate(s).....**PreTest**
 Calculate Fixed Factor Means at**70 75 80**
 Terms**Full Model**
 Include Covariate Interaction Terms**Checked**

Reports Tab

Run Summary.....**Checked**
 ANOVA Table**Checked**
 All Other Reports**Unchecked**

Plots Tab

All Plots.....**Unchecked**

General Linear Models (GLM) for Fixed Factors

Report Options (*in the Toolbar*)

Variable Labels.....Column Names

3 Run the procedure

- Click the **Run** button to perform the calculations and generate the output.

Output**Run Summary**

Response (Y): Test
 Fixed Factor(s): Dose
 Covariate(s): PreTest
 Model: PreTest + Dose + PreTest*Dose

Parameter	Value	Rows	Value
R ²	0.9301	Rows Processed	48
Adjusted R ²	0.9218	Rows Filtered Out	0
Coefficient of Variation	0.0355	Rows with Y Missing	0
Mean Square Error	8.578494	Rows with X's Missing	0
Square Root of MSE	2.928907	Rows Used in Estimation	48
Average Percent Error	2.777	Completion Status	Normal Completion
Error Degrees of Freedom	42		

Analysis of Variance

Source	DF	Sum of Squares	Mean Square	F-Ratio	P-Value	Significant at $\alpha = 0.05$?
Model	5	4794.882	958.9765	111.788	0.0000	Yes
PreTest	1	3527.334	3527.334	411.183	0.0000	Yes
Dose	2	54.85988	27.42994	3.198	0.0510	No
PreTest*Dose	2	17.7435	8.871752	1.034	0.3644	No
Error	42	360.2968	8.578494			
Total(Adjusted)	47	5155.179	109.6847			

The p-value of 0.3644 for the PreTest*Dose interaction term indicates that we cannot reject the assumption of equal slopes among groups. Let's go ahead and remove the interaction term from the model and re-run, now outputting the least squares means and one-sided Each vs. Reference Value multiple comparisons.

General Linear Models (GLM) for Fixed Factors

4 Adjust the model and select additional reports and plots

- The settings for this example are listed below and are stored in the **Example 2b** settings file. To load these settings to the procedure window, click **Open Example Settings File** in the Help Center or File menu.

Variables, Model Tab

Terms**1-Way**

Reports Tab

Least Squares Means**Checked**Compare Each vs. Reference Value.....**Checked**

Plots Tab

Means Plots**Checked**Multiple Comparisons Plots**Checked**Report Options (*in the Toolbar*)Variable Labels.....**Column Names****5 Run the procedure**

- Click the **Run** button to perform the calculations and generate the output.

Output

Run Summary

Response (Y): Test
 Fixed Factor(s): Dose
 Covariate(s): PreTest
 Model: PreTest + Dose

Parameter	Value	Rows	Value
R ²	0.9267	Rows Processed	48
Adjusted R ²	0.9217	Rows Filtered Out	0
Coefficient of Variation	0.0355	Rows with Y Missing	0
Mean Square Error	8.591824	Rows with X's Missing	0
Square Root of MSE	2.931181	Rows Used in Estimation	48
Average Percent Error	2.905	Completion Status	Normal Completion
Error Degrees of Freedom	44		

General Linear Models (GLM) for Fixed Factors

Analysis of Variance

Source	DF	Sum of Squares	Mean Square	F-Ratio	P-Value	Significant at $\alpha = 0.05$?
Model	3	4777.139	1592.38	185.337	0	Yes
PreTest	1	3530.287	3530.287	410.889	0	Yes
Dose	2	503.197	251.5985	29.283	0	Yes
Error	44	378.0403	8.591824			
Total(Adjusted)	47	5155.179	109.6847			

The ANOVA table indicates that both PreTest (covariate) and Dose (fixed factor) are highly significant, so we'll look at the least squares means and multiple comparison tests with PreTest evaluated at 3 different values: 70, 75, and 80.

Least Squares Means

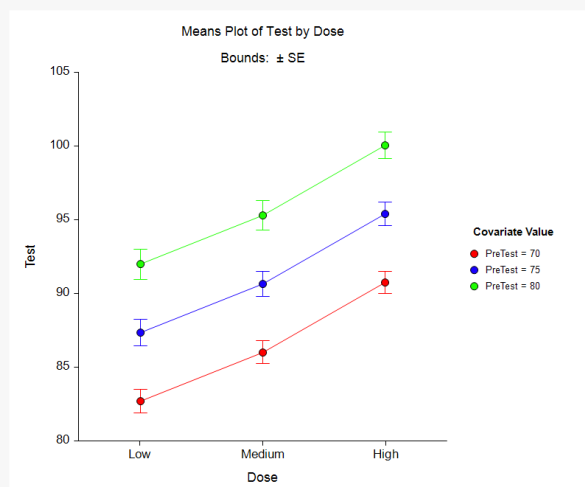
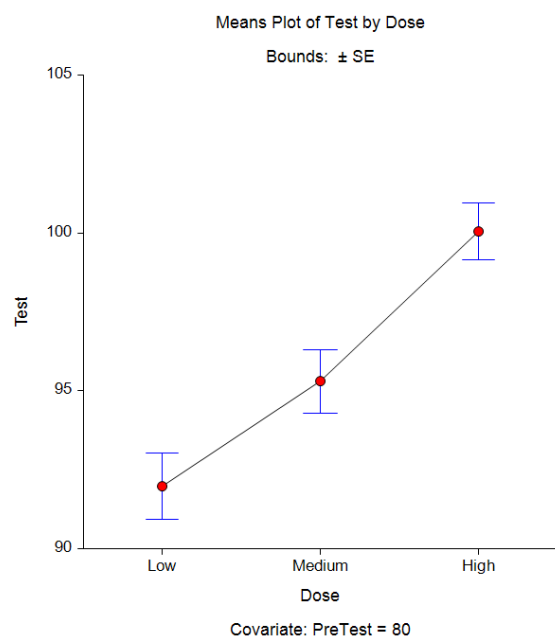
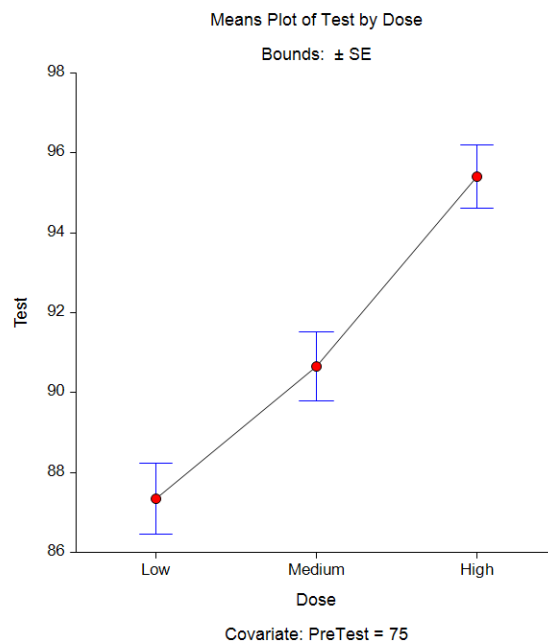
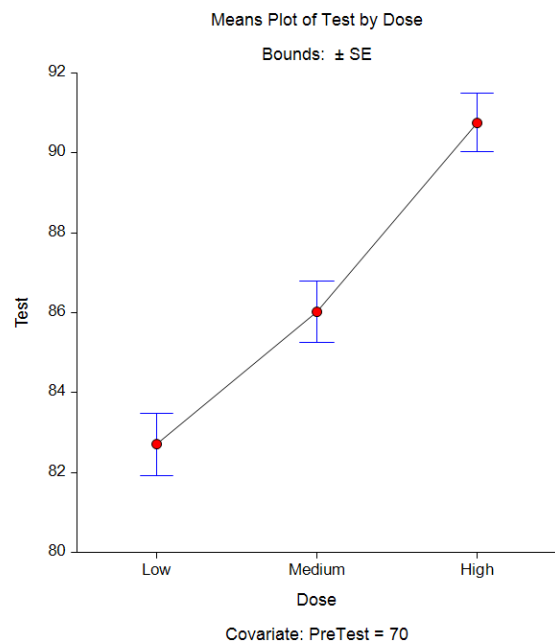
Error Degrees of Freedom (DF): 44

Means Calculated at: Covariate Values 1 to 3 (See below)

				95% Confidence Interval Limits for the Mean	
		Least Squares Mean	Standard Error		
Name	Count			Lower	Upper
Covariate Value 1: PreTest = 70					
Intercept					
All	48	86.49380	0.4645610	85.55753	87.43005
Dose					
Low	16	82.70145	0.7862548	81.11685	84.28604
Medium	16	86.01933	0.7645398	84.47850	87.56016
High	16	90.76061	0.7363901	89.27651	92.24471
Covariate Value 2: PreTest = 75					
Intercept					
All	48	91.13351	0.5966998	89.93094	92.33608
Dose					
Low	16	87.34116	0.8950093	85.53738	89.14494
Medium	16	90.65904	0.8583225	88.92921	92.38888
High	16	95.40032	0.7924203	93.80331	96.99734
Covariate Value 3: PreTest = 80					
Intercept					
All	48	95.77322	0.7752848	94.21074	97.33571
Dose					
Low	16	91.98087	1.0433930	89.87805	94.08369
Medium	16	95.29876	0.9968430	93.28976	97.30776
High	16	100.04000	0.9046390	98.21686	101.86320

General Linear Models (GLM) for Fixed Factors

Means Plots



The least squares means are adjusted based on the model and depend on the covariate values used. The means may change each time a new model is run.

General Linear Models (GLM) for Fixed Factors

Each vs. Reference Value Comparisons of Least Squares Means

Error Degrees of Freedom (DF): 44
 Means Calculated at: Covariate Values 1 to 3 (See below)
 Multiple Comparison Type: Bonferroni
 Hypotheses Tested: H0: Diff $\leq$ 0 vs. H1: Diff $>$ 0

Comparison	Least Squares Mean Difference	Standard Error	Hypothesis Tests			
			T-Statistic	P-Value		Reject H0 at $\alpha = 0.05$?†
				Unadjusted	Adjusted*	
Covariate Value 1: PreTest = 70						
Dose [2 Comparisons]						
Medium - Low	3.317883	1.038489	3.195	0.0013	0.0026	Yes
High - Low	8.059164	1.057851	7.618	0.0000	0.0000	Yes
Covariate Value 2: PreTest = 75						
Dose [2 Comparisons]						
Medium - Low	3.317883	1.038489	3.195	0.0013	0.0026	Yes
High - Low	8.059164	1.057851	7.618	0.0000	0.0000	Yes
Covariate Value 3: PreTest = 80						
Dose [2 Comparisons]						
Medium - Low	3.317883	1.038489	3.195	0.0013	0.0026	Yes
High - Low	8.059164	1.057851	7.618	0.0000	0.0000	Yes

* Adjusted p-values are computed using the number of comparisons and the adjustment type (Bonferroni).

† Rejection decisions are based on adjusted p-values.

General Linear Models (GLM) for Fixed Factors

Simultaneous Confidence Intervals for Each vs. Reference Value Comparisons of Least Squares Means

Error Degrees of Freedom (DF): 44

Means Calculated at: Covariate Values 1 to 3 (See below)

Multiple Comparison Type: Bonferroni

Comparison	Least Squares Mean Difference	Standard Error	95% Simultaneous Confidence Interval Limits*	
			Lower	Upper

Covariate Value 1: PreTest = 70**Dose [2 Comparisons]**

Medium - Low	3.317883	1.038489	1.224945	Infinity
High - Low	8.059164	1.057851	5.927207	Infinity

Covariate Value 2: PreTest = 75**Dose [2 Comparisons]**

Medium - Low	3.317883	1.038489	1.224945	Infinity
High - Low	8.059164	1.057851	5.927207	Infinity

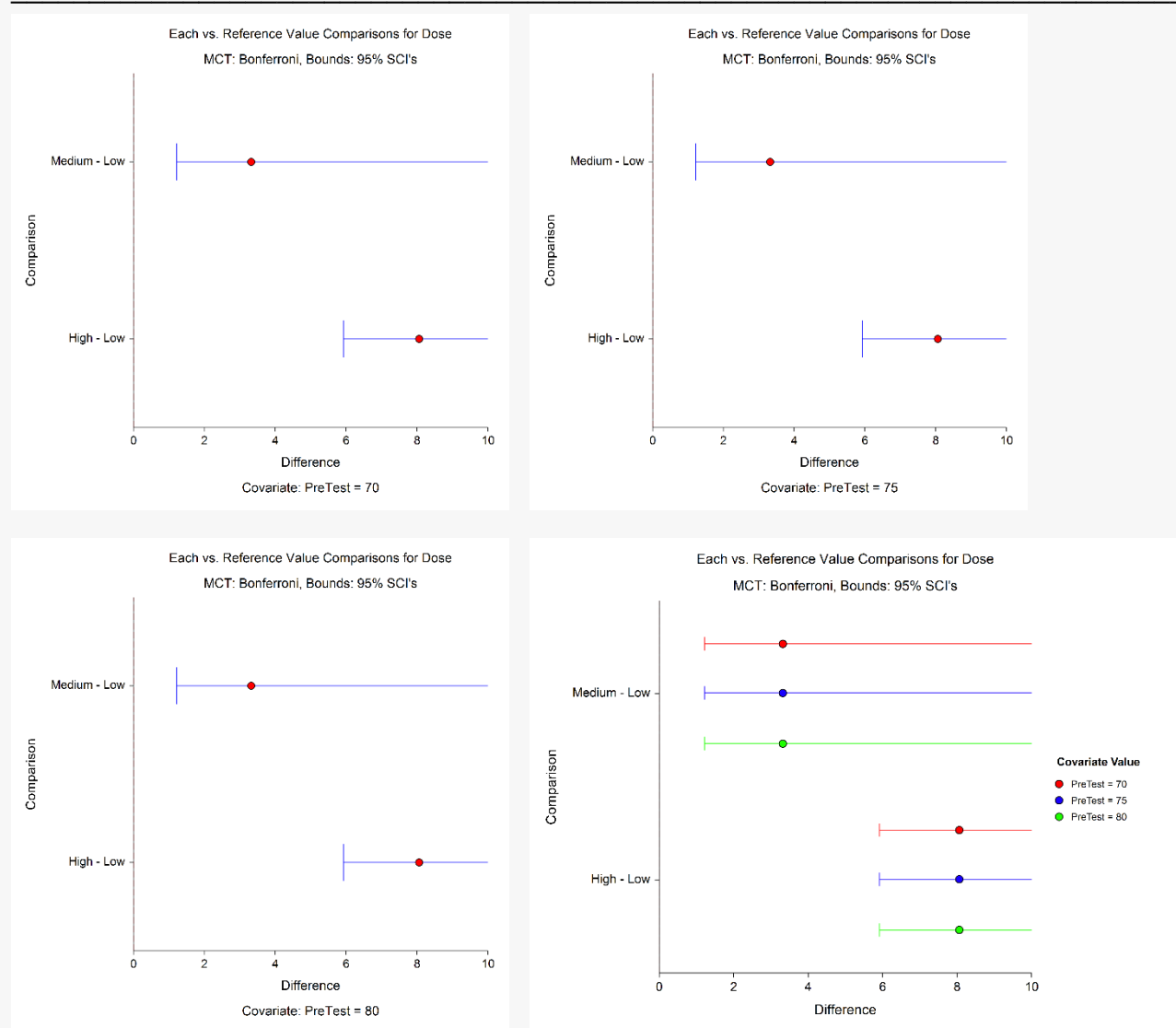
Covariate Value 3: PreTest = 80**Dose [2 Comparisons]**

Medium - Low	3.317883	1.038489	1.224945	Infinity
High - Low	8.059164	1.057851	5.927207	Infinity

* Confidence interval limits are adjusted based on the number of comparisons and the adjustment type (Bonferroni).

General Linear Models (GLM) for Fixed Factors

Each vs. Reference Value Comparisons Plots



The one-sided multiple comparison tests and simultaneous confidence intervals both indicate that both Medium and High doses are significantly different from the Low dose. Notice that the one-sided simultaneous confidence intervals have only a lower bound. We should point out that the test results are the same, regardless of the covariate value used--- this is because there is no interaction in the model. If the model included the covariate-by-factor interaction, then the test results would be different for the various covariate values.

Example 3 – Two-Way ANOVA Model

In Example 1 we analyzed a two-factor, two-covariate model using the 2 Factor 2 Covariate dataset and found that the fixed factor Dose and the covariate PreTest were significant predictors of the response, Test. Let's see now what the results would be if we were to ignore the covariates and run the two-way main effects ANOVA model instead, which allows us to perform Dunnett's test versus the control.

Setup

To run this example, complete the following steps:

1 Open the 2 Factor 2 Covariate example dataset

- From the File menu of the NCSS Data window, select **Open Example Data**.
- Select 2 Factor 2 Covariate and click OK.

2 Specify the General Linear Models (GLM) for Fixed Factors procedure options

- Find and open the **General Linear Models (GLM) for Fixed Factors** procedure using the menus or the Procedure Navigator.
- The settings for this example are listed below and are stored in the **Example 3** settings file. To load these settings to the procedure window, click **Open Example Settings File** in the Help Center or File menu.

Variables, Model Tab

Response (Y).....**Test**
 Fixed Factor(s).....**Dose, Gender**

Reports Tab

Run Summary.....**Checked**
 ANOVA Table**Checked**
 Least Squares Means.....**Checked**
 Compare Each vs. Reference Value.....**Checked**
 All Other Reports**Unchecked**

Plots Tab

Multiple Comparisons Plots**Checked**

Report Options (*in the Toolbar*)

Variable Labels.....**Column Names**

3 Run the procedure

- Click the **Run** button to perform the calculations and generate the output.

General Linear Models (GLM) for Fixed Factors

Output

Run Summary

Response (Y): Test
 Fixed Factor(s): Dose, Gender
 Covariate(s): [None]
 Model: Dose + Gender

Parameter	Value	Rows	Value
R ²	0.2685	Rows Processed	48
Adjusted R ²	0.2186	Rows Filtered Out	0
Coefficient of Variation	0.1121	Rows with Y Missing	0
Mean Square Error	85.70373	Rows with X's Missing	0
Square Root of MSE	9.257631	Rows Used in Estimation	48
Average Percent Error	9.261	Completion Status	Normal Completion
Error Degrees of Freedom	44		

Analysis of Variance

Source	DF	Sum of Squares	Mean Square	F-Ratio	P-Value	Significant at $\alpha = 0.05$?
Model	3	1384.215	461.405	5.384	0.0030	Yes
Dose	2	1246.852	623.4258	7.274	0.0019	Yes
Gender	1	137.3633	137.3633	1.603	0.2122	No
Error	44	3770.964	85.70373			
Total(Adjusted)	47	5155.179	109.6847			

Least Squares Means

Error Degrees of Freedom (DF): 44

Name	Count	Least Squares Mean	Standard Error	95% Confidence Interval Limits for the Mean	
				Lower	Upper
Intercept					
All	48	82.60416	1.336224	79.91119	85.29715
Dose					
Low	16	76.92500	2.314408	72.26062	81.58939
Medium	16	81.60000	2.314408	76.93562	86.26438
High	16	89.28750	2.314408	84.62312	93.95188
Gender					
F	24	84.29583	1.889706	80.48738	88.10429
M	24	80.91250	1.889706	77.10405	84.72095

General Linear Models (GLM) for Fixed Factors

Each vs. Reference Value Comparisons of Least Squares Means

Error Degrees of Freedom (DF): 44
 Multiple Comparison Type: Dunnett
 Hypotheses Tested: $H_0: \text{Diff} = 0$ vs. $H_1: \text{Diff} \neq 0$

Comparison	Least Squares Mean Difference	Standard Error	Hypothesis Tests			
			T-Statistic	P-Value		Reject H0 at $\alpha = 0.05$?†
				Unadjusted	Adjusted*	
Dose [2 Comparisons]						
Medium - Low	4.675	3.273067	1.428	0.1603	0.2709	No
High - Low	12.3625	3.273067	3.777	0.0005	0.0009	Yes
Gender [1 Comparison]						
M - F	-3.383333	2.672448	-1.266	0.2122	0.2122	No

* Adjusted p-values are computed using the number of comparisons and the adjustment type (Dunnett).

† Rejection decisions are based on adjusted p-values.

Simultaneous Confidence Intervals for Each vs. Reference Value Comparisons of Least Squares Means

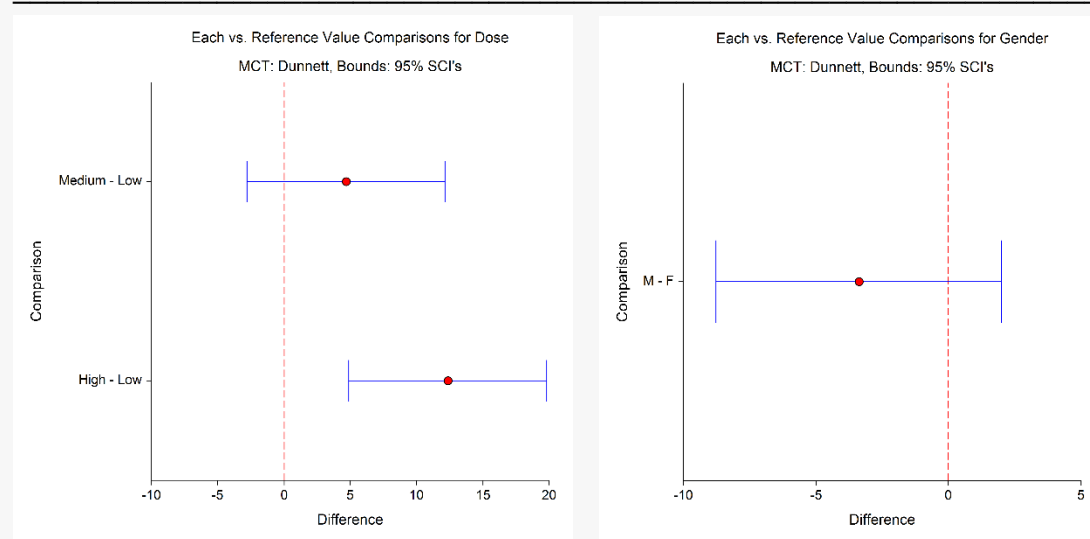
Error Degrees of Freedom (DF): 44
 Multiple Comparison Type: Dunnett

Comparison	Least Squares Mean Difference	Standard Error	95% Simultaneous Confidence Interval Limits*	
			Lower	Upper
Dose [2 Comparisons]				
Medium - Low	4.675	3.273067	-2.799391	12.14939
High - Low	12.3625	3.273067	4.88811	19.83689
Gender [1 Comparison]				
M - F	-3.383333	2.672448	-8.769299	2.002631

* Confidence interval limits are adjusted based on the number of comparisons and the adjustment type (Dunnett).

General Linear Models (GLM) for Fixed Factors

Each vs. Reference Value Comparisons Plots



The ANOVA table indicates that only Dose is significant. Of the comparisons against the reference level, Low, only the High vs. Low comparison is significant using Dunnett's test. Note that when we included the covariates in the model in Example 1, both High vs. Low and Medium vs. Low were deemed significantly different. This illustrates the increase in power that may be obtained when including covariates in the model that account for some of the variability in the response and reduce the overall error.

Example 4 – One-Way ANOVA Model

This example will demonstrate how to analyze a One-Way ANOVA model using the Corn Yield 2 dataset, which is setup for input into this procedure. The dataset contains yield values for three different corn varieties.

Note: This is the same data (i.e., from Corn Yield) that were analyzed in Example 1 of the One-Way Analysis of Variance procedure, except that the data have been stacked column by column.

Corn Yield 2 Dataset

Corn	Yield
A	452
A	874
A	554
.	.
.	.
.	.
B	546
B	547
B	774
.	.
.	.
.	.
C	785
C	458
C	886
.	.
.	.
.	.

Setup

To run this example, complete the following steps:

1 Open the Corn Yield 2 example dataset

- From the File menu of the NCSS Data window, select **Open Example Data**.
- Select **Corn Yield 2** and click **OK**.

2 Specify the General Linear Models (GLM) for Fixed Factors procedure options

- Find and open the **General Linear Models (GLM) for Fixed Factors** procedure using the menus or the Procedure Navigator.
- The settings for this example are listed below and are stored in the **Example 4** settings file. To load these settings to the procedure window, click **Open Example Settings File** in the Help Center or File menu.

General Linear Models (GLM) for Fixed Factors

Variables, Model Tab

Response (Y).....Yield

Fixed Factor(s).....Corn

Reports Tab

Run Summary.....Checked

ANOVA Table.....Checked

Least Squares Means.....Checked

Compare All Pairs.....Checked

All Other Reports.....Unchecked

3 Run the procedure

- Click the **Run** button to perform the calculations and generate the output.

Output**Run Summary**

Response (Y): Yield
 Fixed Factor(s): Corn
 Covariate(s): [None]
 Model: Corn

Parameter	Value	Rows	Value
R ²	0.272	Rows Processed	43
Adjusted R ²	0.2356	Rows Filtered Out	0
Coefficient of Variation	0.2202	Rows with Y Missing	0
Mean Square Error	17964.36	Rows with X's Missing	0
Square Root of MSE	134.0312	Rows Used in Estimation	43
Average Percent Error	19.222	Completion Status	Normal Completion
Error Degrees of Freedom	40		

Analysis of Variance

Source	DF	Sum of Squares	Mean Square	F-Ratio	P-Value	Significant at $\alpha = 0.05$?
Model	2	268532.4	134266.2	7.474	0.0017	Yes
Corn	2	268532.4	134266.2	7.474	0.0017	Yes
Error	40	718574.3	17964.36			
Total(Adjusted)	42	987106.6	23502.54			

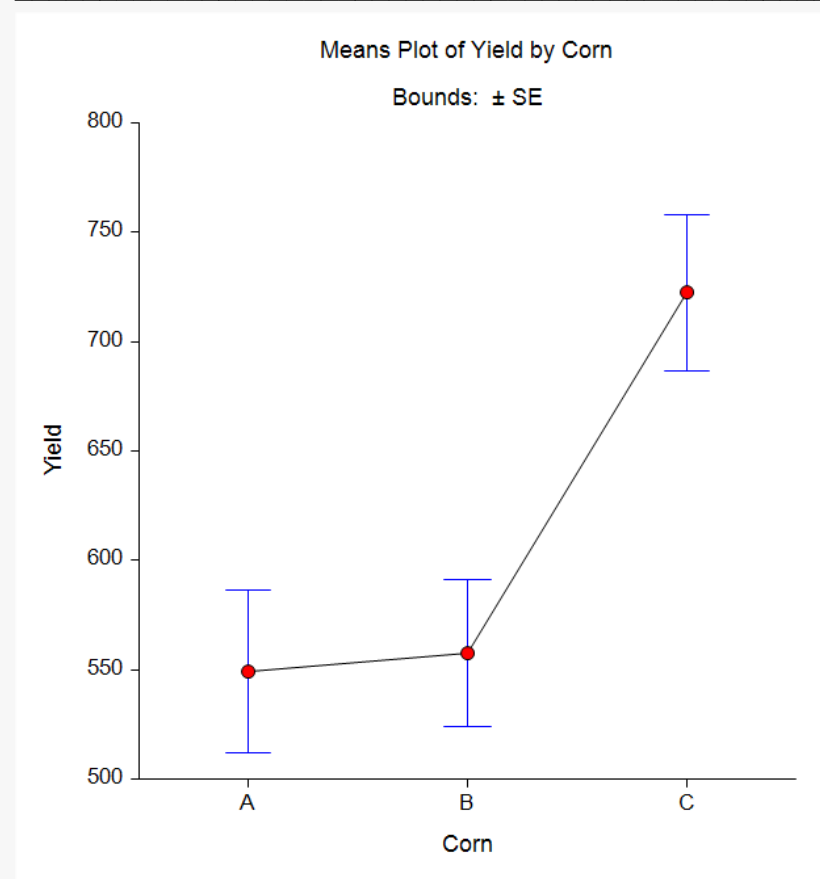
General Linear Models (GLM) for Fixed Factors

Least Squares Means

Error Degrees of Freedom (DF): 40

				95% Confidence Interval Limits for the Mean	
Name	Count	Least Squares Mean	Standard Error	Lower	Upper
Intercept					
All	43	609.7473	20.51507	568.2847	651.2098
Corn					
A	13	549.3846	37.17356	474.2541	624.5152
B	16	557.5000	33.50779	489.7782	625.2218
C	14	722.3571	35.82134	649.9595	794.7548

Means Plots



General Linear Models (GLM) for Fixed Factors

All-Pairs Comparisons of Least Squares Means

Error Degrees of Freedom (DF): 40

Multiple Comparison Type: Tukey-Kramer

Hypotheses Tested: H0: Diff = 0 vs. H1: Diff $\neq$ 0

		Hypothesis Tests				
Comparison	Least Squares Mean Difference	Standard Error	T-Statistic	P-Value		Reject H0 at $\alpha = 0.05$?†
				Unadjusted	Adjusted*	
Corn [3 Comparisons]						
A - B	-8.115385	50.04644	-0.162	0.8720	0.9659	No
A - C	-172.9725	51.62405	-3.351	0.0018	0.0049	Yes
B - C	-164.8571	49.05039	-3.361	0.0017	0.0048	Yes

* Adjusted p-values are computed using the number of comparisons and the adjustment type (Tukey-Kramer).

† Rejection decisions are based on adjusted p-values.

Simultaneous Confidence Intervals for All-Pairs Comparisons of Least Squares Means

Error Degrees of Freedom (DF): 40

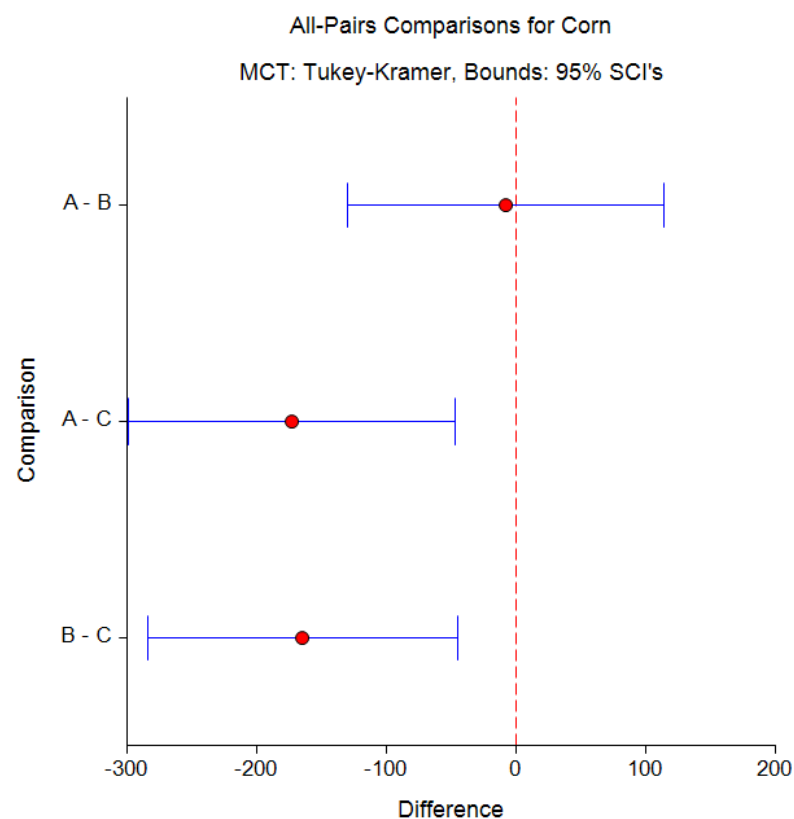
Multiple Comparison Type: Tukey-Kramer

Comparison	Least Squares Mean Difference	Standard Error	95% Simultaneous Confidence Interval Limits*	
			Lower	Upper
Corn [3 Comparisons]				
A - B	-8.115385	50.04644	-129.9271	113.6964
A - C	-172.9725	51.62405	-298.6241	-47.32092
B - C	-164.8571	49.05039	-284.2445	-45.46975

* Confidence interval limits are adjusted based on the number of comparisons and the adjustment type (Tukey-Kramer).

General Linear Models (GLM) for Fixed Factors

All-Pairs Comparisons Plots



The ANOVA table indicates that not all the means are equal. The Tukey-Kramer multiple comparison tests indicate that both varieties A and B are significantly different from variety C, but not significantly different from each other.

Example 5 – Checking the Parallel Slopes Assumption in ANCOVA

Example 3 of the Multiple Regression procedure documentation and Example 4 of the General Linear Models (GLM) procedure documentation discuss how to check the parallel slopes assumption in ANCOVA. This example will show how to do this very quickly using this General Linear Models (GLM) for Fixed Factors procedure.

The analysis of covariance uses features from both analysis of variance and multiple regression. The usual one-way classification model in analysis of variance is

$$Y_{ij} = \mu_i + e_{1ij}$$

where Y_{ij} is the j^{th} observation in the i^{th} group, μ_i represents the true mean of the i^{th} group, and e_{1ij} are the residuals or errors in the above model (usually assumed to be normally distributed). Suppose you have measured a second variable with values X_{ij} that is linearly related to Y . Further suppose that the slope of the relationship between Y and X is constant from group to group. You could then write the analysis of covariance model

$$Y_{ij} = \mu_i + \beta(X_{ij} - \bar{X}_{..}) + e_{2ij}$$

where $\bar{X}_{..}$ represents the overall mean of X . If X and Y are closely related, you would expect that the errors, e_{2ij} , would be much smaller than the errors, e_{1ij} , giving you more precise results.

The classical analysis of covariance is useful for many reasons, but it does have the (highly) restrictive assumption that the slope is constant over all the groups. This assumption is often violated, which limits the technique's usefulness. You will want to study more about this technique in statistical texts before you use it.

The ANCOVA dataset contains three variables: State, Age, and IQ. The researcher wants to test for IQ differences across the three states while controlling for each subjects age. An analysis of covariance should include a preliminary test of the assumption that the slopes between age and IQ are equal across the three states. Without parallel slopes, differences among mean state IQ's depend on age.

General Linear Models (GLM) for Fixed Factors

ANCOVA Dataset

State	Age	IQ
Iowa	12	100
Iowa	13	102
Iowa	12	97
.	.	.
.	.	.
.	.	.
Utah	14	104
Utah	11	105
Utah	12	106
.	.	.
.	.	.
.	.	.
Texas	15	105
Texas	16	106
Texas	12	103
.	.	.
.	.	.
.	.	.

Setup

To run this example, complete the following steps:

1 Open the ANCOVA example dataset

- From the File menu of the NCSS Data window, select **Open Example Data**.
- Select **ANCOVA** and click **OK**.

2 Specify the General Linear Models (GLM) for Fixed Factors procedure options

- Find and open the **General Linear Models (GLM) for Fixed Factors** procedure using the menus or the Procedure Navigator.
- The settings for this example are listed below and are stored in the **Example 5** settings file. To load these settings to the procedure window, click **Open Example Settings File** in the Help Center or File menu.

Variables, Model Tab

Response (Y).....**IQ**
 Fixed Factor(s).....**State**
 Covariate(s)**Age**
 Calculate Fixed Factor Means at**Covariate Means**
 Terms**Full Model**
 Include Covariate Interaction Terms**Checked**

General Linear Models (GLM) for Fixed Factors

Reports Tab

ANOVA Table**Checked**
 All Other Reports**Unchecked**

Plots Tab

All Plots.....**Unchecked**

3 Run the procedure

- Click the **Run** button to perform the calculations and generate the output.

Output**Analysis of Variance**

Source	DF	Sum of Squares	Mean Square	F-Ratio	P-Value	Significant at $\alpha = 0.05$?
Model	5	80.15984	16.03197	1.547	0.2128	No
Age	1	9.740934	9.740934	0.940	0.3419	No
State	2	46.57466	23.28733	2.248	0.1274	No
Age*State	2	38.72052	19.36026	1.869	0.1761	No
Error	24	248.6402	10.36001			
Total(Adjusted)	29	328.8	11.33793			

The F-Value for the Age*State interaction term is 1.869. This matches the result that was obtained in Multiple Regression Example 3 and by hand calculations in General Linear Models (GLM) Example 4. Since the p-value of 0.1761 is not significant, we cannot reject the assumption that the three slopes are equal.