

***Exploring Management 6th edition Schermerhorn
Solutions Manual***

SKU: 9781119395805-SOLUTIONS

A Guide to Modern Econometrics, 5th edition

Answers to selected exercises - Chapter 2

Exercise 2.1

- a. See page 7-9.
- b. Assumption (A1) and (A2), see page 15.
- c. See page 25. We also require assumptions (A3) and (A4) to make the confidence interval approximately valid. We need (A3)-(A4) to make sure that $\text{se}(b_2)$ is the correct standard error. In small samples, the confidence interval is exactly valid if (A5), normality of the error terms, is imposed too.

- d. The hypothesis that $\beta_3 = 1$ can be tested by means of a t -test. The test statistic is

$$t = \frac{b_3 - 1}{\text{se}(b_3)},$$

which – under the null hypothesis – has an approximate standard normal distribution (assumptions (A1)-(A4)). At the 95% confidence level, we reject the null of $|t| > 1.96$.

- e. The hypothesis that $\beta_2 + \beta_3 = 0$ can also be tested by means of a t -test. The test statistic is

$$t = \frac{b_2 + b_3}{\text{se}(b_2 + b_3)}.$$

See Subsection 2.5.3 for an explanation on how to obtain the standard error. (Most software will do this automatically.)

- f. The joint hypothesis that $\beta_2 = \beta_3 = 0$ can be tested by means of an F -test. The test statistic is most easily obtained from the R^2 and we use (2.61) (or – more generally – (2.59)). We compare the test statistic with the critical values from an F distribution with 2 (the number of restrictions) and $N - 3$ degrees of freedom.
- g. For consistency we need assumptions (A6) and (A7). The first of these is often referred to as a “regularity condition”. See page 34.
- h. This is an example of exact multicollinearity. Estimation will break down. Regression software will either give an error message (“singular matrix”) or automatically drop either x_{i2} or x_{i3} from the model.

i. Write

$$\begin{aligned}y_i &= \beta_1^* + \beta_2^*x_{i2}^* + \beta_3^*x_{i3} + \varepsilon_i \\&= \beta_1^* + \beta_2^*(2x_{i2} - 2) + \beta_3^*x_{i3} + \varepsilon_i \\&= [\beta_1^* - 2\beta_2^*] + 2\beta_2^*x_{i2} + \beta_3^*x_{i3} + \varepsilon_i.\end{aligned}$$

From this it follows that in the new model, $\beta_2^* = 0.5\beta_2$, i.e. the coefficient for x_{i2}^* will be half the coefficient for x_{i2} . The intercept term β_1^* will be equal to $\beta_1 + \beta_2$. The coefficient for x_{i3} is unaffected, while the R^2 s of the two models are identical. These two models are statistically equivalent. In projection terminology, they provide the linear projection upon the same space.

j. Write

$$\begin{aligned}y_i &= \beta_1^* + \beta_2^*u_i + \beta_3^*x_{i3} + \varepsilon_i \\&= \beta_1^* + \beta_2^*(x_{i2} - x_{i3}) + \beta_3^*x_{i3} + \varepsilon_i \\&= \beta_1^* + \beta_2^*x_{i2} + (\beta_3^* - \beta_2^*)x_{i3} + \varepsilon_i.\end{aligned}$$

In the new model, we find that $\beta_3^* = \beta_2 + \beta_3$. The other coefficients and the R^2 s are not affected. This model is also statistically equivalent to the original model. The coefficient for x_{i3} now has a different interpretation, because the ceteris paribus condition has changed. See the discussion in Section 3.1 of Chapter 3.

Exercise 2.4

- a. This statement is false. The phrase “linear” in the acronym BLUE refers to the fact that the OLS estimator is a linear function of y (see p. 17). The model is called linear because it is a linear function of the unknown parameters. It is possible to estimate a linear model by a nonlinear estimation method. For example, feasible generalized least squares (as discussed in Chapter 4) is typically nonlinear (in y), even though the model of interest is linear in β .
- b. This statement is false. What is meant here – of course – is whether we need all four Gauss-Markov conditions to argue that a t -test is valid. It is not needed to assume that all regressors are independent of all error terms (assumption (A2)). All that is required is that regressors and error terms are independent for the same observation (assumption (A8) on page 36). This is important (particularly in time series applications) because it allows the inclusion of, for example, a lagged dependent variable in the model, without affecting the validity of standard t - and F -tests. Recall that the assumption of normal error terms is not needed for the t -test either. It would be needed to argue that the t -statistic has an exact t

distribution (in small samples). However, even without normality the t distribution or a standard normal distribution are approximately valid. The approximation is more accurate for large samples.

- c. This statement is true. The result holds by virtue of the first order conditions of OLS (see page 8-9). It tells us that the residuals cannot be explained by the regressors. This is obvious, because any part of y_i that could be explained by the regressors ends up in the fitted value $\hat{y}_i$ and thus not in the residuals. One conclusion from this is that it does not make sense to test whether the error terms are uncorrelated with the regressors (as implied by assumptions (A2) or (A7)) by regressing the residuals upon the regressors.
- d. This statement is false. Actually, the formulation of the statement is incorrect. We never test hypotheses about the *estimators* that we use. Estimators are random variables and could take many different values depending upon the sample that we have. Instead, we always test hypotheses about the true but unknown coefficients. The statement should therefore be formulated as: The hypothesis that a (beta) coefficient is equal to zero can be tested by means of a t -test.
- e. This statement is false. Asymptotic theory tells us something about the distribution of the estimators we use, not about the error terms themselves. Clearly, if the error terms have a certain distribution, for example uniform over the interval $[-1, 1]$ (to make sure that they have zero mean), this does not change if we get more observations. Our estimators, which are typically (weighted) averages of these error terms (apart from a nonrandom component), do get different distributions if the number of observations increases. For most estimators in econometrics we find that they are asymptotically normally distributed (see Subsection 2.6.2), which follows from a central limit theorem. This means that in large samples, these estimators have approximately a normal distribution.
- f. This statement is not true. The statement is not necessarily false, but in general we have to be careful to accept a given hypothesis. The best way to formulate our decision is as follows: “If the absolute t -value of a coefficient is smaller than 1.96, we *cannot reject* the null hypothesis that the coefficient is zero, with 95% confidence.” While it is possible to not reject a number of different mutually exclusive hypotheses, it is somewhat awkward to accept a number of mutually exclusive hypotheses (see Subsection 2.5.7). For example, perhaps we cannot reject that some coefficient β_k is 0 and we can also not reject that it is 1, but we cannot accept that β_k is both 0 and 1 at the same time. Sometimes, standard errors are so high (or the power of a test is so low), that it is hard to reject a wide range of alternative null hypotheses. On the basis of the results of a t -test, we may decide to accept the null hypothesis. For example, when we test whether a variable can be omitted from the model, we typically

accept that it has a coefficient of zero if the null hypothesis cannot be rejected. However, one has to keep in mind that this does not mean that the null hypothesis is actually true. We may simply do not have sufficient evidence to reject it.

- g. This statement is false. The two occurrences of “best” refer to two completely different things. The fact that OLS provides the *best* linear approximation is an algebraic result and *best* has the algebraic interpretation of making the sum of squared differences between the observed values and the fitted values as small as possible. This holds for the sample values that we observe, irrespective of whatever is assumed about a statistical model and its properties. The fact that OLS gives *best* linear unbiased estimators (under assumptions (A1)-(A4) listed in Subsection 2.3.1), is a statistical result and means that the OLS estimators have the smallest variance (within the class of linear unbiased estimators) (see page 17). This statistical result refers to the sampling distribution of the OLS estimator. Whether or not it is true depends upon the assumptions we are willing to make. If the Gauss-Markov conditions are not satisfied, the OLS estimator is often not the *best* linear unbiased estimator, while it would still produce the *best* linear approximation of y_i . Also note that the assumptions we make determine what we mean by the true value of β (for example, a causal parameter). In Section 2.1 there is no true beta we are trying to estimate.
- h. This statement is false. Actually, it works the other way around. If a variable is significant at the 5% level, it is also significant at the 10% level. This is most easily explained on the basis of the p -value (see page 31). If the p -value is smaller than 0.05 (5%) we say that a variable is significant at the 5% level. Clearly, if p is smaller than 0.05 it is certainly smaller than 0.10 (10%).
- i. This statement is not true. A p -value is informative if we wish to make a conclusion about whether or not to reject the null hypothesis. It tells us nothing about the economic significance of the effect, and is also not helpful if we are interested in a different null hypothesis than the one under test (e.g. that a parameter is equal to one rather than zero). A confidence interval provides both information about the size of the estimated effect, as well as its precision. Nevertheless, p -values are more frequently reported than confidence intervals. See Subsection 2.5.7 and the references therein for more information.
- j. This statement is not true. It is correct that removing outliers often leads to lower standard errors, but there is no standard way to deal with outliers and just selecting those observations in the sample that give you the “nicest?? results is a questionable research practice. It is advisable to investigate how influential outliers are upon your estimation results. What ever you do, be explicit and transparant. See Subsection 2.9.1.

- k. This statement is false. The p -value corresponds to the probability of finding a given test value, or a more extreme one, if the null hypothesis is true. This is very different from the probability that the null hypothesis is true, given the data we observe, although you would not be the first one to confuse these two. There is no direct or obvious link between the two probabilities. See Subsection 2.5.7 and the references therein for more information.

© 2017, John Wiley and Sons