

Complete Solutions

Multiple Regression

Ratiomera Statistics

These complete solutions use the same identifiers and order as the Exercise Sheet. Intermediate values are retained until the stated rounding step, so small differences caused by earlier rounding are acceptable where noted. All settings, values, data, and software outputs are constructed teaching material; they are not empirical findings.

Part I: Theory

A06: Constructing Dummy Indicators and Finding the Reference

T07-A06-V01: Tutorial format

Identify the issue, part (a)

With an intercept, $k - 1 = 2$ indicators are required. The omitted category is represented by the intercept and becomes the comparison baseline.

Reason through the evidence, part (b)

The complete coding is:

Category	D_1 (Video)	D_2 (Interactive)
Text	0	0
Video	1	0
Interactive	0	1

Reason through the evidence, part (c)

Text is the reference because every indicator equals zero in its row. The fitted category values are:

Category	Fitted reasoning score
Text	61.00
Video	64.50
Interactive	67.00

The coefficient on D_1 is 3.50. Therefore, the fitted reasoning score for Video is 3.50 points higher than for Text. The intercept 61.00 is the fitted value for Text.

State the conclusion and its limits, part (d)

For each case, the k category indicators would sum exactly to one, which is already the intercept column. Including all of them with the intercept makes one column an exact combination of the others, so the coefficients are not uniquely identified. Choosing a different reference changes the displayed intercept and category contrasts, but it does not change any category's fitted value.

T07-A06-V02: Study location

Identify the issue, part (a)

With an intercept, $k - 1 = 3$ indicators are required. The omitted category is represented by the intercept and becomes the comparison baseline.

Reason through the evidence, part (b)

The complete coding is:

Category	D_1 (Library)	D_2 (Study room)	D_3 (Outdoors)
Home	0	0	0
Library	1	0	0
Study room	0	1	0
Outdoors	0	0	1

Reason through the evidence, part (c)

Home is the reference because every indicator equals zero in its row. The fitted category values are:

Category	Fitted focus score
Home	54.00
Library	58.00
Study room	56.50
Outdoors	52.50

The coefficient on D_1 is 4.00. Therefore, the fitted focus score for Library is 4.00 points higher than for Home. The intercept 54.00 is the fitted value for Home.

State the conclusion and its limits, part (d)

For each case, the k category indicators would sum exactly to one, which is already the intercept column. Including all of them with the intercept makes one column an exact combination of the others, so the coefficients are not uniquely identified. Choosing a different reference changes the displayed intercept and category contrasts, but it does not change any category's fitted value.

T07-A06-V03: Feedback channel

Identify the issue, part (a)

With an intercept, $k - 1 = 2$ indicators are required. The omitted category is represented by the intercept and becomes the comparison baseline.

Reason through the evidence, part (b)

The complete coding is:

Category	D_1 (Audio)	D_2 (Video)
Written	0	0
Audio	1	0
Video	0	1

Reason through the evidence, part (c)

Written is the reference because every indicator equals zero in its row. The fitted category values are:

Category	Fitted revision score
Written	66.00
Audio	68.00
Video	70.50

The coefficient on D_1 is 2.00. Therefore, the fitted revision score for Audio is 2.00 points higher than for Written. The intercept 66.00 is the fitted value for Written.

State the conclusion and its limits, part (d)

For each case, the k category indicators would sum exactly to one, which is already the intercept column. Including all of them with the intercept makes one column an exact combination of the others, so the coefficients are not uniquely identified. Choosing a different reference changes the displayed intercept and category contrasts, but it does not change any category's fitted value.

T07-A06-V04: Note-taking method

Identify the issue, part (a)

With an intercept, $k - 1 = 3$ indicators are required. The omitted category is represented by the intercept and becomes the comparison baseline.

Reason through the evidence, part (b)

The complete coding is:

Category	D_1 (Tablet)	D_2 (Laptop)	D_3 (Mixed)
Paper	0	0	0
Tablet	1	0	0
Laptop	0	1	0
Mixed	0	0	1

Reason through the evidence, part (c)

Paper is the reference because every indicator equals zero in its row. The fitted category values are:

Category	Fitted recall score
Paper	58.00
Tablet	56.50
Laptop	55.50
Mixed	61.00

The coefficient on D_1 is -1.50 . Therefore, the fitted recall score for Tablet is 1.50 points lower than for Paper. The intercept 58.00 is the fitted value for Paper.

State the conclusion and its limits, part (d)

For each case, the k category indicators would sum exactly to one, which is already the intercept column. Including all of them with the intercept makes one column an exact combination of the others, so the coefficients are not uniquely identified. Choosing a different reference changes the displayed intercept and category contrasts, but it does not change any category's fitted value.

T07-A06-V05: Workshop schedule

Identify the issue, part (a)

With an intercept, $k - 1 = 2$ indicators are required. The omitted category is represented by the intercept and becomes the comparison baseline.

Reason through the evidence, part (b)

The complete coding is:

Category	D_1 (Afternoon)	D_2 (Evening)
Morning	0	0
Afternoon	1	0
Evening	0	1

Reason through the evidence, part (c)

Morning is the reference because every indicator equals zero in its row. The fitted category values are:

Category	Fitted confidence score
Morning	49.00
Afternoon	51.50
Evening	46.00

The coefficient on D_1 is 2.50. Therefore, the fitted confidence score for Afternoon is 2.50 points higher than for Morning. The intercept 49.00 is the fitted value for Morning.

State the conclusion and its limits, part (d)

For each case, the k category indicators would sum exactly to one, which is already the intercept column. Including all of them with the intercept makes one column an exact combination of the others, so the coefficients are not uniquely identified. Choosing a different reference changes the displayed intercept and category contrasts, but it does not change any category's fitted value.

T07-A06-V06: Archive guide

Identify the issue, part (a)

With an intercept, $k - 1 = 3$ indicators are required. The omitted category is represented by the intercept and becomes the comparison baseline.

Reason through the evidence, part (b)

The complete coding is:

Category	D_1 (Map)	D_2 (Mentor)	D_3 (Search tool)
Checklist	0	0	0
Map	1	0	0
Mentor	0	1	0
Search tool	0	0	1

Reason through the evidence, part (c)

Checklist is the reference because every indicator equals zero in its row. The fitted category values are:

Category	Fitted retrieval score
Checklist	63.00
Map	64.50
Mentor	68.00
Search tool	66.00

The coefficient on D_1 is 1.50. Therefore, the fitted retrieval score for Map is 1.50 points higher than for Checklist. The intercept 63.00 is the fitted value for Checklist.

State the conclusion and its limits, part (d)

For each case, the k category indicators would sum exactly to one, which is already the intercept column. Including all of them with the intercept makes one column an exact combination of the others, so the coefficients are not uniquely identified. Choosing a different reference changes the displayed intercept and category contrasts, but it does not change any category's fitted value.

T07-A06-V07: Revision strategy

Identify the issue, part (a)

With an intercept, $k - 1 = 2$ indicators are required. The omitted category is represented by the intercept and becomes the comparison baseline.

Reason through the evidence, part (b)

The complete coding is:

Category	D_1 (Peer review)	D_2 (Instructor review)
Self-review	0	0
Peer review	1	0
Instructor review	0	1

Reason through the evidence, part (c)

Self-review is the reference because every indicator equals zero in its row. The fitted category values are:

Category	Fitted quality score
Self-review	60.00
Peer review	64.00
Instructor review	67.00

The coefficient on D_1 is 4.00. Therefore, the fitted quality score for Peer review is 4.00 points higher than for Self-review. The intercept 60.00 is the fitted value for Self-review.

State the conclusion and its limits, part (d)

For each case, the k category indicators would sum exactly to one, which is already the intercept column. Including all of them with the intercept makes one column an exact combination of the others, so the coefficients are not uniquely identified. Choosing a different reference changes the displayed intercept and category contrasts, but it does not change any category's fitted value.

T07-A06-V08: Museum route

Identify the issue, part (a)

With an intercept, $k - 1 = 4$ indicators are required. The omitted category is represented by the intercept and becomes the comparison baseline.

Reason through the evidence, part (b)

The complete coding is:

Category	D_1 (Thematic)	D_2 (Free choice)	D_3 (Guided)	D_4 (Hybrid)
Chronological	0	0	0	0
Thematic	1	0	0	0
Free choice	0	1	0	0
Guided	0	0	1	0
Hybrid	0	0	0	1

Reason through the evidence, part (c)

Chronological is the reference because every indicator equals zero in its row. The fitted category values are:

Category	Fitted knowledge score
Chronological	57.00
Thematic	60.00
Free choice	56.00
Guided	62.50
Hybrid	61.00

The coefficient on D_1 is 3.00. Therefore, the fitted knowledge score for Thematic is 3.00 points higher than for Chronological. The intercept 57.00 is the fitted value for Chronological.

State the conclusion and its limits, part (d)

For each case, the k category indicators would sum exactly to one, which is already the intercept column. Including all of them with the intercept makes one column an exact combination of the others, so the coefficients are not uniquely identified. Choosing a different reference changes the displayed intercept and category contrasts, but it does not change any category's fitted value.

T07-A06-V09: Study plan

Identify the issue, part (a)

With an intercept, $k - 1 = 2$ indicators are required. The omitted category is represented by the intercept and becomes the comparison baseline.

Reason through the evidence, part (b)

The complete coding is:

Category	D_1 (Twice weekly)	D_2 (Weekly)
Daily	0	0
Twice weekly	1	0
Weekly	0	1

Reason through the evidence, part (c)

Daily is the reference because every indicator equals zero in its row. The fitted category values are:

Category	Fitted retention score
Daily	69.00
Twice weekly	67.00
Weekly	64.00

The coefficient on D_1 is -2.00 . Therefore, the fitted retention score for Twice weekly is 2.00 points lower than for Daily. The intercept 69.00 is the fitted value for Daily.

State the conclusion and its limits, part (d)

For each case, the k category indicators would sum exactly to one, which is already the intercept column. Including all of them with the intercept makes one column an exact combination of the others, so the coefficients are not uniquely identified. Choosing a different reference changes the displayed intercept and category contrasts, but it does not change any category's fitted value.

T07-A06-V10: Task interface

Identify the issue, part (a)

With an intercept, $k - 1 = 3$ indicators are required. The omitted category is represented by the intercept and becomes the comparison baseline.

Reason through the evidence, part (b)

The complete coding is:

Category	D_1 (Board)	D_2 (Calendar)	D_3 (Timeline)
List	0	0	0
Board	1	0	0
Calendar	0	1	0
Timeline	0	0	1

Reason through the evidence, part (c)

List is the reference because every indicator equals zero in its row. The fitted category values are:

Category	Fitted completion score
List	62.00
Board	64.50
Calendar	66.00
Timeline	63.00

The coefficient on D_1 is 2.50. Therefore, the fitted completion score for Board is 2.50 points higher than for List. The intercept 62.00 is the fitted value for List.

State the conclusion and its limits, part (d)

For each case, the k category indicators would sum exactly to one, which is already the intercept column. Including all of them with the intercept makes one column an exact combination of the others, so the coefficients are not uniquely identified. Choosing a different reference changes the displayed intercept and category contrasts, but it does not change any category's fitted value.

Part II: Calculator Practice

A01: Reading a Multiple-Regression Equation and Output

T07-A01-V01: Guided practice and reasoning

Reason before calculating, part (a)

The fitted equation is $\hat{Y} = 38.000 + (2.400)X_1 + (0.310)X_2$. An unstandardized slope uses the original measurement units. A standardized coefficient instead describes the fitted change in outcome standard deviations for a one-standard-deviation increase in a predictor, conditional on the other predictor.

Work through the calculation, part (b)

With prior-preparation score held fixed, a one-unit increase in guided-practice hours is associated with a fitted change of 2.400 points in reasoning score. With guided-practice hours held fixed, a one-unit increase in prior-preparation score is associated with a fitted change of 0.310 points. These are conditional associations, not automatically causal effects.

Work through the calculation, part (c)

For X_1 , $t = 2.400/0.580 = 4.136$ with 77 degrees of freedom, giving $p < 0.0001$; therefore, reject the coefficient null at $\alpha = .05$. For X_2 , $t = 0.310/0.108 = 2.879$, giving $p = 0.0052$; therefore, reject the coefficient null. Each test concerns that one population coefficient conditional on the exact other term in this model.

Interpret and check the result, part (d)

$R^2 = 0.370$ means the fitted two-predictor model accounts for 37.0% of the sample variation in reasoning score. Adjusted $R^2 = 0.354$ applies an in-sample penalty for estimating two slopes; it is not a new-data test. The residual standard error says observed outcomes typically remain spread by roughly 5.60 points around their fitted values, under the model. The standardized slopes, 0.419 and 0.292, differ from the bivariate correlations, 0.550 and 0.480, because each slope separates a predictor's conditional relationship from variation shared with the other predictor.

T07-A01-V02: Archive workflow and retrieval time

Reason before calculating, part (a)

The fitted equation is $\hat{Y} = 70.000 + (-1.750)X_1 + (-0.220)X_2$. An unstandardized slope uses the original measurement units. A standardized coefficient instead describes the fitted change in outcome standard deviations for a one-standard-deviation increase in a predictor, conditional on the other predictor.

Work through the calculation, part (b)

With archive-experience months held fixed, a one-unit increase in checklist-practice sessions is associated with a fitted change of -1.750 minutes in retrieval time. With checklist-practice sessions held fixed, a one-unit increase in archive-experience months is associated with a fitted change of -0.220 minutes. These are conditional associations, not automatically causal effects.

Work through the calculation, part (c)

For X_1 , $t = -1.750/0.467 = -3.747$ with 69 degrees of freedom, giving $p = 0.0004$; therefore, reject the coefficient null at $\alpha = .05$. For X_2 , $t = -0.220/0.093 = -2.366$, giving $p = 0.0208$; therefore, reject the coefficient null. Each test concerns that one population coefficient conditional on the exact other term in this model.

Interpret and check the result, part (d)

$R^2 = 0.316$ means the fitted two-predictor model accounts for 31.6% of the sample variation in retrieval time. Adjusted $R^2 = 0.296$ applies an in-sample penalty for estimating two slopes; it is not a new-data test. The residual standard error says observed outcomes typically remain spread by roughly 4.80 minutes around their fitted values, under the model. The standardized slopes, -0.407 and -0.257, differ from the bivariate correlations, 0.550 and 0.480, because each slope separates a predictor's conditional relationship from variation shared with the other predictor.

−0.510 and −0.420, because each slope separates a predictor’s conditional relationship from variation shared with the other predictor.

T07-A01-V03: Reading routines and comprehension

Reason before calculating, part (a)

The fitted equation is $\hat{Y} = 42.000 + (1.850)X_1 + (0.280)X_2$. An unstandardized slope uses the original measurement units. A standardized coefficient instead describes the fitted change in outcome standard deviations for a one-standard-deviation increase in a predictor, conditional on the other predictor.

Work through the calculation, part (b)

With baseline-vocabulary score held fixed, a one-unit increase in weekly reading hours is associated with a fitted change of 1.850 points in comprehension score. With weekly reading hours held fixed, a one-unit increase in baseline-vocabulary score is associated with a fitted change of 0.280 points. These are conditional associations, not automatically causal effects.

Work through the calculation, part (c)

For X_1 , $t = 1.850/0.443 = 4.179$ with 92 degrees of freedom, giving $p < 0.0001$; therefore, reject the coefficient null at $\alpha = .05$. For X_2 , $t = 0.280/0.084 = 3.340$, giving $p = 0.0012$; therefore, reject the coefficient null. Each test concerns that one population coefficient conditional on the exact other term in this model.

Interpret and check the result, part (d)

$R^2 = 0.322$ means the fitted two-predictor model accounts for 32.2% of the sample variation in comprehension score. Adjusted $R^2 = 0.308$ applies an in-sample penalty for estimating two slopes; it is not a new-data test. The residual standard error says observed outcomes typically remain spread by roughly 5.10 points around their fitted values, under the model. The standardized slopes, 0.383 and 0.306, differ from the bivariate correlations, 0.490 and 0.440, because each slope separates a predictor’s conditional relationship from variation shared with the other predictor.

T07-A01-V04: Route rehearsal and navigation time

Reason before calculating, part (a)

The fitted equation is $\hat{Y} = 65.000 + (-2.100)X_1 + (-0.160)X_2$. An unstandardized slope uses the original measurement units. A standardized coefficient instead describes the fitted change in outcome standard deviations for a one-standard-deviation increase in a predictor, conditional on the other predictor.

Work through the calculation, part (b)

With route-familiarity score held fixed, a one-unit increase in route-rehearsal attempts is associated with a fitted change of −2.100 minutes in navigation time. With route-rehearsal attempts held fixed, a one-unit increase in route-familiarity score is associated with a fitted change of −0.160 minutes. These are conditional associations, not automatically causal effects.

Work through the calculation, part (c)

For X_1 , $t = -2.100/0.519 = -4.043$ with 65 degrees of freedom, giving $p = 0.0001$; therefore, reject the coefficient null at $\alpha = .05$. For X_2 , $t = -0.160/0.080 = -1.997$, giving $p = 0.0500$; therefore, do not reject the coefficient null. Each test concerns that one population coefficient conditional on the exact other term in this model.

Interpret and check the result, part (d)

$R^2 = 0.322$ means the fitted two-predictor model accounts for 32.2% of the sample variation in navigation time. Adjusted $R^2 = 0.302$ applies an in-sample penalty for estimating two slopes; it is not a new-data test. The residual standard error says observed outcomes typically remain spread by roughly 6.00 minutes around their fitted values, under the model. The standardized slopes, −0.446 and −0.220, differ from the bivariate correlations, 0.490 and 0.440, because each slope separates a predictor’s conditional relationship from variation shared with the other predictor.

–0.530 and –0.390, because each slope separates a predictor’s conditional relationship from variation shared with the other predictor.

T07-A01-V05: Search practice and catalog accuracy

Reason before calculating, part (a)

The fitted equation is $\hat{Y} = 48.000 + (1.550)X_1 + (0.340)X_2$. An unstandardized slope uses the original measurement units. A standardized coefficient instead describes the fitted change in outcome standard deviations for a one-standard-deviation increase in a predictor, conditional on the other predictor.

Work through the calculation, part (b)

With prior catalog-knowledge score held fixed, a one-unit increase in search-practice sets is associated with a fitted change of 1.550 points in catalog-accuracy score. With search-practice sets held fixed, a one-unit increase in prior catalog-knowledge score is associated with a fitted change of 0.340 points. These are conditional associations, not automatically causal effects.

Work through the calculation, part (c)

For X_1 , $t = 1.550/0.413 = 3.752$ with 107 degrees of freedom, giving $p = 0.0003$; therefore, reject the coefficient null at $\alpha = .05$. For X_2 , $t = 0.340/0.107 = 3.180$, giving $p = 0.0019$; therefore, reject the coefficient null. Each test concerns that one population coefficient conditional on the exact other term in this model.

Interpret and check the result, part (d)

$R^2 = 0.280$ means the fitted two-predictor model accounts for 28.0% of the sample variation in catalog-accuracy score. Adjusted $R^2 = 0.266$ applies an in-sample penalty for estimating two slopes; it is not a new-data test. The residual standard error says observed outcomes typically remain spread by roughly 4.60 points around their fitted values, under the model. The standardized slopes, 0.339 and 0.288, differ from the bivariate correlations, 0.460 and 0.430, because each slope separates a predictor’s conditional relationship from variation shared with the other predictor.

T07-A01-V06: Workshop participation and confidence

Reason before calculating, part (a)

The fitted equation is $\hat{Y} = 30.000 + (2.200)X_1 + (0.450)X_2$. An unstandardized slope uses the original measurement units. A standardized coefficient instead describes the fitted change in outcome standard deviations for a one-standard-deviation increase in a predictor, conditional on the other predictor.

Work through the calculation, part (b)

With baseline-confidence score held fixed, a one-unit increase in workshop sessions is associated with a fitted change of 2.200 points in confidence score. With workshop sessions held fixed, a one-unit increase in baseline-confidence score is associated with a fitted change of 0.450 points. These are conditional associations, not automatically causal effects.

Work through the calculation, part (c)

For X_1 , $t = 2.200/0.546 = 4.027$ with 73 degrees of freedom, giving $p = 0.0001$; therefore, reject the coefficient null at $\alpha = .05$. For X_2 , $t = 0.450/0.125 = 3.590$, giving $p = 0.0006$; therefore, reject the coefficient null. Each test concerns that one population coefficient conditional on the exact other term in this model.

Interpret and check the result, part (d)

$R^2 = 0.363$ means the fitted two-predictor model accounts for 36.3% of the sample variation in confidence score. Adjusted $R^2 = 0.345$ applies an in-sample penalty for estimating two slopes; it is not a new-data test. The residual standard error says observed outcomes typically remain spread by roughly 5.00 points around their fitted values, under the model. The standardized slopes, 0.395 and 0.352, differ from the bivariate correlations, 0.500 and 0.470, because each slope separates a predictor’s conditional relationship from variation shared with the other predictor.

T07-A01-V07: Focus blocks and task accuracy

Reason before calculating, part (a)

The fitted equation is $\hat{Y} = 55.000 + (1.300)X_1 + (1.150)X_2$. An unstandardized slope uses the original measurement units. A standardized coefficient instead describes the fitted change in outcome standard deviations for a one-standard-deviation increase in a predictor, conditional on the other predictor.

Work through the calculation, part (b)

With sleep duration in hours held fixed, a one-unit increase in notification-free blocks is associated with a fitted change of 1.300 points in task-accuracy score. With notification-free blocks held fixed, a one-unit increase in sleep duration in hours is associated with a fitted change of 1.150 points. These are conditional associations, not automatically causal effects.

Work through the calculation, part (c)

For X_1 , $t = 1.300/0.330 = 3.935$ with 117 degrees of freedom, giving $p = 0.0001$; therefore, reject the coefficient null at $\alpha = .05$. For X_2 , $t = 1.150/0.335 = 3.438$, giving $p = 0.0008$; therefore, reject the coefficient null. Each test concerns that one population coefficient conditional on the exact other term in this model.

Interpret and check the result, part (d)

$R^2 = 0.244$ means the fitted two-predictor model accounts for 24.4% of the sample variation in task-accuracy score. Adjusted $R^2 = 0.231$ applies an in-sample penalty for estimating two slopes; it is not a new-data test. The residual standard error says observed outcomes typically remain spread by roughly 4.30 points around their fitted values, under the model. The standardized slopes, 0.329 and 0.288, differ from the bivariate correlations, 0.410 and 0.380, because each slope separates a predictor's conditional relationship from variation shared with the other predictor.

T07-A01-V08: Museum engagement and historical knowledge

Reason before calculating, part (a)

The fitted equation is $\hat{Y} = 40.000 + (2.650)X_1 + (0.370)X_2$. An unstandardized slope uses the original measurement units. A standardized coefficient instead describes the fitted change in outcome standard deviations for a one-standard-deviation increase in a predictor, conditional on the other predictor.

Work through the calculation, part (b)

With prior-history score held fixed, a one-unit increase in museum visits is associated with a fitted change of 2.650 points in historical-knowledge score. With museum visits held fixed, a one-unit increase in prior-history score is associated with a fitted change of 0.370 points. These are conditional associations, not automatically causal effects.

Work through the calculation, part (c)

For X_1 , $t = 2.650/0.619 = 4.283$ with 81 degrees of freedom, giving $p < 0.0001$; therefore, reject the coefficient null at $\alpha = .05$. For X_2 , $t = 0.370/0.118 = 3.144$, giving $p = 0.0023$; therefore, reject the coefficient null. Each test concerns that one population coefficient conditional on the exact other term in this model.

Interpret and check the result, part (d)

$R^2 = 0.350$ means the fitted two-predictor model accounts for 35.0% of the sample variation in historical-knowledge score. Adjusted $R^2 = 0.334$ applies an in-sample penalty for estimating two slopes; it is not a new-data test. The residual standard error says observed outcomes typically remain spread by roughly 5.50 points around their fitted values, under the model. The standardized slopes, 0.411 and 0.302, differ from the bivariate correlations, 0.520 and 0.450, because each slope separates a predictor's conditional relationship from variation shared with the other predictor.

T07-A01-V09: Peer feedback and revision quality

Reason before calculating, part (a)

The fitted equation is $\hat{Y} = 44.000 + (2.100)X_1 + (0.300)X_2$. An unstandardized slope uses the original measurement units. A standardized coefficient instead describes the fitted change in outcome standard deviations for a one-standard-deviation increase in a predictor, conditional on the other predictor.

Work through the calculation, part (b)

With baseline-writing score held fixed, a one-unit increase in peer-feedback rounds is associated with a fitted change of 2.100 points in revision-quality score. With peer-feedback rounds held fixed, a one-unit increase in baseline-writing score is associated with a fitted change of 0.300 points. These are conditional associations, not automatically causal effects.

Work through the calculation, part (c)

For X_1 , $t = 2.100/0.507 = 4.145$ with 89 degrees of freedom, giving $p < 0.0001$; therefore, reject the coefficient null at $\alpha = .05$. For X_2 , $t = 0.300/0.104 = 2.877$, giving $p = 0.0050$; therefore, reject the coefficient null. Each test concerns that one population coefficient conditional on the exact other term in this model.

Interpret and check the result, part (d)

$R^2 = 0.296$ means the fitted two-predictor model accounts for 29.6% of the sample variation in revision-quality score. Adjusted $R^2 = 0.280$ applies an in-sample penalty for estimating two slopes; it is not a new-data test. The residual standard error says observed outcomes typically remain spread by roughly 4.90 points around their fitted values, under the model. The standardized slopes, 0.391 and 0.271, differ from the bivariate correlations, 0.480 and 0.400, because each slope separates a predictor's conditional relationship from variation shared with the other predictor.

T07-A01-V10: Planning sessions and completion time

Reason before calculating, part (a)

The fitted equation is $\hat{Y} = 82.000 + (-1.900)X_1 + (0.850)X_2$. An unstandardized slope uses the original measurement units. A standardized coefficient instead describes the fitted change in outcome standard deviations for a one-standard-deviation increase in a predictor, conditional on the other predictor.

Work through the calculation, part (b)

With task-complexity score held fixed, a one-unit increase in planning sessions is associated with a fitted change of -1.900 minutes in completion time. With planning sessions held fixed, a one-unit increase in task-complexity score is associated with a fitted change of 0.850 minutes. These are conditional associations, not automatically causal effects.

Work through the calculation, part (c)

For X_1 , $t = -1.900/0.384 = -4.954$ with 85 degrees of freedom, giving $p < 0.0001$; therefore, reject the coefficient null at $\alpha = .05$. For X_2 , $t = 0.850/0.185 = 4.590$, giving $p < 0.0001$; therefore, reject the coefficient null. Each test concerns that one population coefficient conditional on the exact other term in this model.

Interpret and check the result, part (d)

$R^2 = 0.361$ means the fitted two-predictor model accounts for 36.1% of the sample variation in completion time. Adjusted $R^2 = 0.346$ applies an in-sample penalty for estimating two slopes; it is not a new-data test. The residual standard error says observed outcomes typically remain spread by roughly 5.70 minutes around their fitted values, under the model. The standardized slopes, -0.430 and 0.398, differ from the bivariate correlations, -0.450 and 0.420, because each slope separates a predictor's conditional relationship from variation shared with the other predictor.

A02: Comparing a Prespecified Nested Model Sequence

T07-A02-V01: Guided practice and reasoning

Reason before calculating, part (a)

Apply $SSE = 1840.0(1 - R^2)$ and subtract consecutive R^2 values.

Work through the calculation, part (b)

Substitute each model's own number of predictors into the adjusted formula:

Model	SSE	Change in R-squared	Adjusted R-squared
M1	1435.20	not a later step	0.2085
M2	1159.20	0.150	0.3512
M3	1122.40	0.020	0.3623

Work through the calculation, part (c)

Ordinary R^2 rises from 0.370 to 0.390 when reflection-session count is added, an increment of 0.020, or 2.0 percentage points of sample variation. Ordinary R^2 cannot fall when a predictor is added to this same-case, same-intercept model. Adjusted R^2 increases from 0.3512 to 0.3623 because it weighs the extra fit against the additional estimated slope. That adjustment is descriptive and in-sample.

Work through the calculation, part (d)

The restricted equation is $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \varepsilon$. The unrestricted equation adds reflection-session count: $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 + \varepsilon$. The null hypothesis is $H_0 : \beta_3 = 0$, conditional on the terms already in M2. The incremental statistic is $F = [(0.390 - 0.370)/1]/[(1 - 0.390)/(70 - 3 - 1)] = 2.1639$ with 1 and 66 degrees of freedom. Its p-value is 0.1460, so the added term does not meet the 5% criterion.

Interpret and check the result, part (e)

M1 is contained in M2, and M2 is contained in M3: setting each newly added coefficient to zero reproduces the preceding model. The outcome, cases, and intercept also stay the same, so the fit changes are comparable as nested steps. The sequence does not randomize predictors, exclude omitted variables, prove a mechanism, or measure prediction on new cases. Those questions require design information and separate validation.

T07-A02-V02: Archive workflow and retrieval time

Reason before calculating, part (a)

Apply $SSE = 1320.0(1 - R^2)$ and subtract consecutive R^2 values.

Work through the calculation, part (b)

Substitute each model's own number of predictors into the adjusted formula:

Model	SSE	Change in R-squared	Adjusted R-squared
M1	950.40	not a later step	0.2708
M2	858.00	0.070	0.3331
M3	856.68	0.001	0.3254

Work through the calculation, part (c)

Ordinary R^2 rises from 0.350 to 0.351 when catalog-familiarity score is added, an increment of 0.001, or 0.1 percentage points of sample variation. Ordinary R^2 cannot fall when a predictor is added to this same-case, same-intercept model. Adjusted R^2 decreases from 0.3331 to 0.3254 because it weighs the extra fit against the additional estimated slope. That adjustment is descriptive and in-sample.

Work through the calculation, part (d)

The restricted equation is $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \varepsilon$. The unrestricted equation adds catalog-familiarity score: $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 + \varepsilon$. The null hypothesis is $H_0 : \beta_3 = 0$, conditional on the terms

already in M2. The incremental statistic is $F = [(0.351 - 0.350)/1]/[(1 - 0.351)/(80 - 3 - 1)] = 0.1171$ with 1 and 76 degrees of freedom. Its p-value is 0.7331, so the added term does not meet the 5% criterion.

Interpret and check the result, part (e)

M1 is contained in M2, and M2 is contained in M3: setting each newly added coefficient to zero reproduces the preceding model. The outcome, cases, and intercept also stay the same, so the fit changes are comparable as nested steps. The sequence does not randomize predictors, exclude omitted variables, prove a mechanism, or measure prediction on new cases. Those questions require design information and separate validation.

T07-A02-V03: Reading routines and comprehension

Reason before calculating, part (a)

Apply $SSE = 1560.0(1 - R^2)$ and subtract consecutive R^2 values.

Work through the calculation, part (b)

Substitute each model's own number of predictors into the adjusted formula:

Model	SSE	Change in R-squared	Adjusted R-squared
M1	1279.20	not a later step	0.1659
M2	1076.40	0.130	0.2858
M3	998.40	0.050	0.3257

Work through the calculation, part (c)

Ordinary R^2 rises from 0.310 to 0.360 when annotation-session count is added, an increment of 0.050, or 5.0 percentage points of sample variation. Ordinary R^2 cannot fall when a predictor is added to this same-case, same-intercept model. Adjusted R^2 increases from 0.2858 to 0.3257 because it weighs the extra fit against the additional estimated slope. That adjustment is descriptive and in-sample.

Work through the calculation, part (d)

The restricted equation is $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \varepsilon$. The unrestricted equation adds annotation-session count: $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 + \varepsilon$. The null hypothesis is $H_0 : \beta_3 = 0$, conditional on the terms already in M2. The incremental statistic is $F = [(0.360 - 0.310)/1]/[(1 - 0.360)/(60 - 3 - 1)] = 4.3750$ with 1 and 56 degrees of freedom. Its p-value is 0.0410, so the added term meets the 5% criterion.

Interpret and check the result, part (e)

M1 is contained in M2, and M2 is contained in M3: setting each newly added coefficient to zero reproduces the preceding model. The outcome, cases, and intercept also stay the same, so the fit changes are comparable as nested steps. The sequence does not randomize predictors, exclude omitted variables, prove a mechanism, or measure prediction on new cases. Those questions require design information and separate validation.

T07-A02-V04: Route rehearsal and navigation time

Reason before calculating, part (a)

Apply $SSE = 2100.0(1 - R^2)$ and subtract consecutive R^2 values.

Work through the calculation, part (b)

Substitute each model's own number of predictors into the adjusted formula:

Model	SSE	Change in R-squared	Adjusted R-squared
M1	1575.00	not a later step	0.2415
M2	1407.00	0.080	0.3146

Model	SSE	Change in R-squared	Adjusted R-squared
M3	1398.60	0.004	0.3108

Work through the calculation, part (c)

Ordinary R^2 rises from 0.330 to 0.334 when landmark-recall score is added, an increment of 0.004, or 0.4 percentage points of sample variation. Ordinary R^2 cannot fall when a predictor is added to this same-case, same-intercept model. Adjusted R^2 decreases from 0.3146 to 0.3108 because it weighs the extra fit against the additional estimated slope. That adjustment is descriptive and in-sample.

Work through the calculation, part (d)

The restricted equation is $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \varepsilon$. The unrestricted equation adds landmark-recall score: $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 + \varepsilon$. The null hypothesis is $H_0 : \beta_3 = 0$, conditional on the terms already in M2. The incremental statistic is $F = [(0.334 - 0.330)/1]/[(1 - 0.334)/(90 - 3 - 1)] = 0.5165$ with 1 and 86 degrees of freedom. Its p-value is 0.4743, so the added term does not meet the 5% criterion.

Interpret and check the result, part (e)

M1 is contained in M2, and M2 is contained in M3: setting each newly added coefficient to zero reproduces the preceding model. The outcome, cases, and intercept also stay the same, so the fit changes are comparable as nested steps. The sequence does not randomize predictors, exclude omitted variables, prove a mechanism, or measure prediction on new cases. Those questions require design information and separate validation.

T07-A02-V05: Search practice and catalog accuracy

Reason before calculating, part (a)

Apply $SSE = 1750.0(1 - R^2)$ and subtract consecutive R^2 values.

Work through the calculation, part (b)

Substitute each model's own number of predictors into the adjusted formula:

Model	SSE	Change in R-squared	Adjusted R-squared
M1	1225.00	not a later step	0.2929
M2	1032.50	0.110	0.3978
M3	980.00	0.030	0.4225

Work through the calculation, part (c)

Ordinary R^2 rises from 0.410 to 0.440 when query-planning score is added, an increment of 0.030, or 3.0 percentage points of sample variation. Ordinary R^2 cannot fall when a predictor is added to this same-case, same-intercept model. Adjusted R^2 increases from 0.3978 to 0.4225 because it weighs the extra fit against the additional estimated slope. That adjustment is descriptive and in-sample.

Work through the calculation, part (d)

The restricted equation is $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \varepsilon$. The unrestricted equation adds query-planning score: $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 + \varepsilon$. The null hypothesis is $H_0 : \beta_3 = 0$, conditional on the terms already in M2. The incremental statistic is $F = [(0.440 - 0.410)/1]/[(1 - 0.440)/(100 - 3 - 1)] = 5.1429$ with 1 and 96 degrees of freedom. Its p-value is 0.0256, so the added term meets the 5% criterion.

Interpret and check the result, part (e)

M1 is contained in M2, and M2 is contained in M3: setting each newly added coefficient to zero reproduces the preceding model. The outcome, cases, and intercept also stay the same, so the fit changes are comparable as

nested steps. The sequence does not randomize predictors, exclude omitted variables, prove a mechanism, or measure prediction on new cases. Those questions require design information and separate validation.

T07-A02-V06: Workshop participation and confidence

Reason before calculating, part (a)

Apply $SSE = 980.0(1 - R^2)$ and subtract consecutive R^2 values.

Work through the calculation, part (b)

Substitute each model's own number of predictors into the adjusted formula:

Model	SSE	Change in R-squared	Adjusted R-squared
M1	823.20	not a later step	0.1442
M2	695.80	0.130	0.2627
M3	693.84	0.002	0.2504

Work through the calculation, part (c)

Ordinary R^2 rises from 0.290 to 0.292 when reflection-log count is added, an increment of 0.002, or 0.2 percentage points of sample variation. Ordinary R^2 cannot fall when a predictor is added to this same-case, same-intercept model. Adjusted R^2 decreases from 0.2627 to 0.2504 because it weighs the extra fit against the additional estimated slope. That adjustment is descriptive and in-sample.

Work through the calculation, part (d)

The restricted equation is $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \varepsilon$. The unrestricted equation adds reflection-log count: $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 + \varepsilon$. The null hypothesis is $H_0 : \beta_3 = 0$, conditional on the terms already in M2. The incremental statistic is $F = [(0.292 - 0.290)/1]/[(1 - 0.292)/(55 - 3 - 1)] = 0.1441$ with 1 and 51 degrees of freedom. Its p-value is 0.7058, so the added term does not meet the 5% criterion.

Interpret and check the result, part (e)

M1 is contained in M2, and M2 is contained in M3: setting each newly added coefficient to zero reproduces the preceding model. The outcome, cases, and intercept also stay the same, so the fit changes are comparable as nested steps. The sequence does not randomize predictors, exclude omitted variables, prove a mechanism, or measure prediction on new cases. Those questions require design information and separate validation.

T07-A02-V07: Focus blocks and task accuracy

Reason before calculating, part (a)

Apply $SSE = 2280.0(1 - R^2)$ and subtract consecutive R^2 values.

Work through the calculation, part (b)

Substitute each model's own number of predictors into the adjusted formula:

Model	SSE	Change in R-squared	Adjusted R-squared
M1	1801.20	not a later step	0.2033
M2	1504.80	0.130	0.3287
M3	1436.40	0.030	0.3537

Work through the calculation, part (c)

Ordinary R^2 rises from 0.340 to 0.370 when planning-break count is added, an increment of 0.030, or 3.0 percentage points of sample variation. Ordinary R^2 cannot fall when a predictor is added to this same-case,

same-intercept model. Adjusted R^2 increases from 0.3287 to 0.3537 because it weighs the extra fit against the additional estimated slope. That adjustment is descriptive and in-sample.

Work through the calculation, part (d)

The restricted equation is $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \varepsilon$. The unrestricted equation adds planning-break count: $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 + \varepsilon$. The null hypothesis is $H_0 : \beta_3 = 0$, conditional on the terms already in M2. The incremental statistic is $F = [(0.370 - 0.340)/1]/[(1 - 0.370)/(120 - 3 - 1)] = 5.5238$ with 1 and 116 degrees of freedom. Its p-value is 0.0204, so the added term meets the 5% criterion.

Interpret and check the result, part (e)

M1 is contained in M2, and M2 is contained in M3: setting each newly added coefficient to zero reproduces the preceding model. The outcome, cases, and intercept also stay the same, so the fit changes are comparable as nested steps. The sequence does not randomize predictors, exclude omitted variables, prove a mechanism, or measure prediction on new cases. Those questions require design information and separate validation.

T07-A02-V08: Museum engagement and historical knowledge

Reason before calculating, part (a)

Apply $SSE = 1440.0(1 - R^2)$ and subtract consecutive R^2 values.

Work through the calculation, part (b)

Substitute each model's own number of predictors into the adjusted formula:

Model	SSE	Change in R-squared	Adjusted R-squared
M1	1094.40	not a later step	0.2296
M2	979.20	0.080	0.3011
M3	977.76	0.001	0.2923

Work through the calculation, part (c)

Ordinary R^2 rises from 0.320 to 0.321 when exhibit-note count is added, an increment of 0.001, or 0.1 percentage points of sample variation. Ordinary R^2 cannot fall when a predictor is added to this same-case, same-intercept model. Adjusted R^2 decreases from 0.3011 to 0.2923 because it weighs the extra fit against the additional estimated slope. That adjustment is descriptive and in-sample.

Work through the calculation, part (d)

The restricted equation is $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \varepsilon$. The unrestricted equation adds exhibit-note count: $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 + \varepsilon$. The null hypothesis is $H_0 : \beta_3 = 0$, conditional on the terms already in M2. The incremental statistic is $F = [(0.321 - 0.320)/1]/[(1 - 0.321)/(75 - 3 - 1)] = 0.1046$ with 1 and 71 degrees of freedom. Its p-value is 0.7474, so the added term does not meet the 5% criterion.

Interpret and check the result, part (e)

M1 is contained in M2, and M2 is contained in M3: setting each newly added coefficient to zero reproduces the preceding model. The outcome, cases, and intercept also stay the same, so the fit changes are comparable as nested steps. The sequence does not randomize predictors, exclude omitted variables, prove a mechanism, or measure prediction on new cases. Those questions require design information and separate validation.

T07-A02-V09: Peer feedback and revision quality

Reason before calculating, part (a)

Apply $SSE = 1620.0(1 - R^2)$ and subtract consecutive R^2 values.

Work through the calculation, part (b)

Substitute each model's own number of predictors into the adjusted formula:

Model	SSE	Change in R-squared	Adjusted R-squared
M1	1312.20	not a later step	0.1771
M2	1036.80	0.170	0.3394
M3	939.60	0.060	0.3915

Work through the calculation, part (c)

Ordinary R^2 rises from 0.360 to 0.420 when revision-plan score is added, an increment of 0.060, or 6.0 percentage points of sample variation. Ordinary R^2 cannot fall when a predictor is added to this same-case, same-intercept model. Adjusted R^2 increases from 0.3394 to 0.3915 because it weighs the extra fit against the additional estimated slope. That adjustment is descriptive and in-sample.

Work through the calculation, part (d)

The restricted equation is $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \varepsilon$. The unrestricted equation adds revision-plan score: $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 + \varepsilon$. The null hypothesis is $H_0 : \beta_3 = 0$, conditional on the terms already in M2. The incremental statistic is $F = [(0.420 - 0.360)/1]/[(1 - 0.420)/(65 - 3 - 1)] = 6.3103$ with 1 and 61 degrees of freedom. Its p-value is 0.0147, so the added term meets the 5% criterion.

Interpret and check the result, part (e)

M1 is contained in M2, and M2 is contained in M3: setting each newly added coefficient to zero reproduces the preceding model. The outcome, cases, and intercept also stay the same, so the fit changes are comparable as nested steps. The sequence does not randomize predictors, exclude omitted variables, prove a mechanism, or measure prediction on new cases. Those questions require design information and separate validation.

T07-A02-V10: Planning sessions and completion time

Reason before calculating, part (a)

Apply $SSE = 1960.0(1 - R^2)$ and subtract consecutive R^2 values.

Work through the calculation, part (b)

Substitute each model's own number of predictors into the adjusted formula:

Model	SSE	Change in R-squared	Adjusted R-squared
M1	1430.80	not a later step	0.2632
M2	1195.60	0.120	0.3786
M3	1185.80	0.005	0.3779

Work through the calculation, part (c)

Ordinary R^2 rises from 0.390 to 0.395 when progress-check count is added, an increment of 0.005, or 0.5 percentage points of sample variation. Ordinary R^2 cannot fall when a predictor is added to this same-case, same-intercept model. Adjusted R^2 decreases from 0.3786 to 0.3779 because it weighs the extra fit against the additional estimated slope. That adjustment is descriptive and in-sample.

Work through the calculation, part (d)

The restricted equation is $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \varepsilon$. The unrestricted equation adds progress-check count: $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 + \varepsilon$. The null hypothesis is $H_0 : \beta_3 = 0$, conditional on the terms already in M2. The incremental statistic is $F = [(0.395 - 0.390)/1]/[(1 - 0.395)/(110 - 3 - 1)] = 0.8760$ with 1 and 106 degrees of freedom. Its p-value is 0.3514, so the added term does not meet the 5% criterion.

Interpret and check the result, part (e)

M1 is contained in M2, and M2 is contained in M3: setting each newly added coefficient to zero reproduces the preceding model. The outcome, cases, and intercept also stay the same, so the fit changes are comparable as nested steps. The sequence does not randomize predictors, exclude omitted variables, prove a mechanism, or measure prediction on new cases. Those questions require design information and separate validation.

A03: Distinguishing the Global F Test From Coefficient t Tests

T07-A03-V01: Guided practice and reasoning

Reason before calculating, part (a)

The global null is $H_0 : \beta_1 = \beta_2 = \beta_3 = 0$. The statistic is $F = (0.220/3)/[(1 - 0.220)/46] = 4.325$. Because 4.325 is greater than 2.80684, reject the global null at $\alpha = .05$.

Work through the calculation, part (b)

The coefficient calculations are: guided-practice hours: $t = 1.800/0.600 = 3.000$, $p = 0.0043$, so reject the coefficient null; prior-preparation score: $t = 0.220/0.180 = 1.222$, $p = 0.2278$, so do not reject the coefficient null; reflection sessions: $t = 0.120/0.160 = 0.750$, $p = 0.4571$, so do not reject the coefficient null. Thus 1 of the three displayed individual tests rejects at the stated level.

Work through the calculation, part (c)

For predictor X_j , the individual null is $H_0 : \beta_j = 0$ conditional on every other term in this exact model. The global test asks one joint question about all three slopes. Rejecting it says at least one non-intercept population slope differs from zero under the model, but the global statistic does not name a predictor. Not rejecting it likewise is not proof that every population slope equals zero.

Interpret and check the result, part (d)

The two sets of decisions can differ because the global test evaluates the predictors jointly, whereas each t test isolates one conditional coefficient and its uncertainty. Shared predictor variation can make individual standard errors large even when the predictor set has joint explanatory value. Conversely, sampling variation can yield a small individual p-value in a model whose global test is not rejected. A p-value does not measure effect size, practical value, future prediction, or causality.

T07-A03-V02: Archive workflow and retrieval time

Reason before calculating, part (a)

The global null is $H_0 : \beta_1 = \beta_2 = \beta_3 = 0$. The statistic is $F = (0.300/3)/[(1 - 0.300)/56] = 8.000$. Because 8.000 is greater than 2.76943, reject the global null at $\alpha = .05$.

Work through the calculation, part (b)

The coefficient calculations are: checklist-practice sessions: $t = -1.400/0.450 = -3.111$, $p = 0.0029$, so reject the coefficient null; archive-experience months: $t = -0.200/0.160 = -1.250$, $p = 0.2165$, so do not reject the coefficient null; catalog familiarity: $t = 0.300/0.120 = 2.500$, $p = 0.0154$, so reject the coefficient null. Thus 2 of the three displayed individual tests reject at the stated level.

Work through the calculation, part (c)

For predictor X_j , the individual null is $H_0 : \beta_j = 0$ conditional on every other term in this exact model. The global test asks one joint question about all three slopes. Rejecting it says at least one non-intercept population slope differs from zero under the model, but the global statistic does not name a predictor. Not rejecting it likewise is not proof that every population slope equals zero.

Interpret and check the result, part (d)

The two sets of decisions can differ because the global test evaluates the predictors jointly, whereas each t test isolates one conditional coefficient and its uncertainty. Shared predictor variation can make individual standard errors large even when the predictor set has joint explanatory value. Conversely, sampling variation can yield a small individual p-value in a model whose global test is not rejected. A p-value does not measure effect size, practical value, future prediction, or causality.

T07-A03-V03: Reading routines and comprehension

Reason before calculating, part (a)

The global null is $H_0 : \beta_1 = \beta_2 = \beta_3 = 0$. The statistic is $F = (0.100/3)/[(1 - 0.100)/66] = 2.444$. Because 2.444 is not greater than 2.74371, do not reject the global null at $\alpha = .05$.

Work through the calculation, part (b)

The coefficient calculations are: weekly reading hours: $t = 1.100/0.580 = 1.897$, $p = 0.0623$, so do not reject the coefficient null; baseline-vocabulary score: $t = 0.180/0.130 = 1.385$, $p = 0.1708$, so do not reject the coefficient null; annotation sessions: $t = -0.150/0.140 = -1.071$, $p = 0.2879$, so do not reject the coefficient null. Thus 0 of the three displayed individual tests reject at the stated level.

Work through the calculation, part (c)

For predictor X_j , the individual null is $H_0 : \beta_j = 0$ conditional on every other term in this exact model. The global test asks one joint question about all three slopes. Rejecting it says at least one non-intercept population slope differs from zero under the model, but the global statistic does not name a predictor. Not rejecting it likewise is not proof that every population slope equals zero.

Interpret and check the result, part (d)

The two sets of decisions can differ because the global test evaluates the predictors jointly, whereas each t test isolates one conditional coefficient and its uncertainty. Shared predictor variation can make individual standard errors large even when the predictor set has joint explanatory value. Conversely, sampling variation can yield a small individual p-value in a model whose global test is not rejected. A p-value does not measure effect size, practical value, future prediction, or causality.

T07-A03-V04: Route rehearsal and navigation time

Reason before calculating, part (a)

The global null is $H_0 : \beta_1 = \beta_2 = \beta_3 = 0$. The statistic is $F = (0.250/3)/[(1 - 0.250)/76] = 8.444$. Because 8.444 is greater than 2.72494, reject the global null at $\alpha = .05$.

Work through the calculation, part (b)

The coefficient calculations are: route-rehearsal attempts: $t = -1.800/0.550 = -3.273$, $p = 0.0016$, so reject the coefficient null; route-familiarity score: $t = -0.120/0.100 = -1.200$, $p = 0.2339$, so do not reject the coefficient null; landmark recall: $t = 0.280/0.110 = 2.545$, $p = 0.0129$, so reject the coefficient null. Thus 2 of the three displayed individual tests reject at the stated level.

Work through the calculation, part (c)

For predictor X_j , the individual null is $H_0 : \beta_j = 0$ conditional on every other term in this exact model. The global test asks one joint question about all three slopes. Rejecting it says at least one non-intercept population slope differs from zero under the model, but the global statistic does not name a predictor. Not rejecting it likewise is not proof that every population slope equals zero.

Interpret and check the result, part (d)

The two sets of decisions can differ because the global test evaluates the predictors jointly, whereas each t test isolates one conditional coefficient and its uncertainty. Shared predictor variation can make individual standard errors large even when the predictor set has joint explanatory value. Conversely, sampling variation can yield

a small individual p-value in a model whose global test is not rejected. A p-value does not measure effect size, practical value, future prediction, or causality.

T07-A03-V05: Search practice and catalog accuracy

Reason before calculating, part (a)

The global null is $H_0 : \beta_1 = \beta_2 = \beta_3 = 0$. The statistic is $F = (0.080/3)/[(1 - 0.080)/86] = 2.493$. Because 2.493 is not greater than 2.71065, do not reject the global null at $\alpha = .05$.

Work through the calculation, part (b)

The coefficient calculations are: search-practice sets: $t = 1.000/0.570 = 1.754$, $p = 0.0829$, so do not reject the coefficient null; prior catalog-knowledge score: $t = 0.150/0.120 = 1.250$, $p = 0.2147$, so do not reject the coefficient null; query planning: $t = 0.180/0.140 = 1.286$, $p = 0.2020$, so do not reject the coefficient null. Thus 0 of the three displayed individual tests reject at the stated level.

Work through the calculation, part (c)

For predictor X_j , the individual null is $H_0 : \beta_j = 0$ conditional on every other term in this exact model. The global test asks one joint question about all three slopes. Rejecting it says at least one non-intercept population slope differs from zero under the model, but the global statistic does not name a predictor. Not rejecting it likewise is not proof that every population slope equals zero.

Interpret and check the result, part (d)

The two sets of decisions can differ because the global test evaluates the predictors jointly, whereas each t test isolates one conditional coefficient and its uncertainty. Shared predictor variation can make individual standard errors large even when the predictor set has joint explanatory value. Conversely, sampling variation can yield a small individual p-value in a model whose global test is not rejected. A p-value does not measure effect size, practical value, future prediction, or causality.

T07-A03-V06: Workshop participation and confidence

Reason before calculating, part (a)

The global null is $H_0 : \beta_1 = \beta_2 = \beta_3 = 0$. The statistic is $F = (0.350/3)/[(1 - 0.350)/96] = 17.231$. Because 17.231 is greater than 2.69939, reject the global null at $\alpha = .05$.

Work through the calculation, part (b)

The coefficient calculations are: workshop sessions: $t = 2.100/0.500 = 4.200$, $p < 0.0001$, so reject the coefficient null; baseline-confidence score: $t = 0.380/0.140 = 2.714$, $p = 0.0079$, so reject the coefficient null; reflection logs: $t = -0.100/0.130 = -0.769$, $p = 0.4436$, so do not reject the coefficient null. Thus 2 of the three displayed individual tests reject at the stated level.

Work through the calculation, part (c)

For predictor X_j , the individual null is $H_0 : \beta_j = 0$ conditional on every other term in this exact model. The global test asks one joint question about all three slopes. Rejecting it says at least one non-intercept population slope differs from zero under the model, but the global statistic does not name a predictor. Not rejecting it likewise is not proof that every population slope equals zero.

Interpret and check the result, part (d)

The two sets of decisions can differ because the global test evaluates the predictors jointly, whereas each t test isolates one conditional coefficient and its uncertainty. Shared predictor variation can make individual standard errors large even when the predictor set has joint explanatory value. Conversely, sampling variation can yield a small individual p-value in a model whose global test is not rejected. A p-value does not measure effect size, practical value, future prediction, or causality.

T07-A03-V07: Focus blocks and task accuracy

Reason before calculating, part (a)

The global null is $H_0 : \beta_1 = \beta_2 = \beta_3 = 0$. The statistic is $F = (0.200/3)/[(1 - 0.200)/106] = 8.833$. Because 8.833 is greater than 2.69030, reject the global null at $\alpha = .05$.

Work through the calculation, part (b)

The coefficient calculations are: notification-free blocks: $t = 1.300/0.400 = 3.250$, $p = 0.0015$, so reject the coefficient null; sleep duration in hours: $t = 0.120/0.110 = 1.091$, $p = 0.2778$, so do not reject the coefficient null; planning breaks: $t = 0.250/0.150 = 1.667$, $p = 0.0985$, so do not reject the coefficient null. Thus 1 of the three displayed individual tests rejects at the stated level.

Work through the calculation, part (c)

For predictor X_j , the individual null is $H_0 : \beta_j = 0$ conditional on every other term in this exact model. The global test asks one joint question about all three slopes. Rejecting it says at least one non-intercept population slope differs from zero under the model, but the global statistic does not name a predictor. Not rejecting it likewise is not proof that every population slope equals zero.

Interpret and check the result, part (d)

The two sets of decisions can differ because the global test evaluates the predictors jointly, whereas each t test isolates one conditional coefficient and its uncertainty. Shared predictor variation can make individual standard errors large even when the predictor set has joint explanatory value. Conversely, sampling variation can yield a small individual p-value in a model whose global test is not rejected. A p-value does not measure effect size, practical value, future prediction, or causality.

T07-A03-V08: Museum engagement and historical knowledge

Reason before calculating, part (a)

The global null is $H_0 : \beta_1 = \beta_2 = \beta_3 = 0$. The statistic is $F = (0.280/3)/[(1 - 0.280)/116] = 15.037$. Because 15.037 is greater than 2.68281, reject the global null at $\alpha = .05$.

Work through the calculation, part (b)

The coefficient calculations are: museum visits: $t = 2.000/0.480 = 4.167$, $p < 0.0001$, so reject the coefficient null; prior-history score: $t = 0.310/0.130 = 2.385$, $p = 0.0187$, so reject the coefficient null; exhibit notes: $t = 0.080/0.120 = 0.667$, $p = 0.5063$, so do not reject the coefficient null. Thus 2 of the three displayed individual tests reject at the stated level.

Work through the calculation, part (c)

For predictor X_j , the individual null is $H_0 : \beta_j = 0$ conditional on every other term in this exact model. The global test asks one joint question about all three slopes. Rejecting it says at least one non-intercept population slope differs from zero under the model, but the global statistic does not name a predictor. Not rejecting it likewise is not proof that every population slope equals zero.

Interpret and check the result, part (d)

The two sets of decisions can differ because the global test evaluates the predictors jointly, whereas each t test isolates one conditional coefficient and its uncertainty. Shared predictor variation can make individual standard errors large even when the predictor set has joint explanatory value. Conversely, sampling variation can yield a small individual p-value in a model whose global test is not rejected. A p-value does not measure effect size, practical value, future prediction, or causality.

T07-A03-V09: Peer feedback and revision quality

Reason before calculating, part (a)

The global null is $H_0 : \beta_1 = \beta_2 = \beta_3 = 0$. The statistic is $F = (0.160/3)/[(1 - 0.160)/71] = 4.508$. Because 4.508 is greater than 2.73365, reject the global null at $\alpha = .05$.

Work through the calculation, part (b)

The coefficient calculations are: peer-feedback rounds: $t = 1.200/0.520 = 2.308$, $p = 0.0239$, so reject the coefficient null; baseline-writing score: $t = 0.190/0.150 = 1.267$, $p = 0.2094$, so do not reject the coefficient null; revision planning: $t = -0.090/0.130 = -0.692$, $p = 0.4910$, so do not reject the coefficient null. Thus 1 of the three displayed individual tests rejects at the stated level.

Work through the calculation, part (c)

For predictor X_j , the individual null is $H_0 : \beta_j = 0$ conditional on every other term in this exact model. The global test asks one joint question about all three slopes. Rejecting it says at least one non-intercept population slope differs from zero under the model, but the global statistic does not name a predictor. Not rejecting it likewise is not proof that every population slope equals zero.

Interpret and check the result, part (d)

The two sets of decisions can differ because the global test evaluates the predictors jointly, whereas each t test isolates one conditional coefficient and its uncertainty. Shared predictor variation can make individual standard errors large even when the predictor set has joint explanatory value. Conversely, sampling variation can yield a small individual p-value in a model whose global test is not rejected. A p-value does not measure effect size, practical value, future prediction, or causality.

T07-A03-V10: Planning sessions and completion time

Reason before calculating, part (a)

The global null is $H_0 : \beta_1 = \beta_2 = \beta_3 = 0$. The statistic is $F = (0.240/3)/[(1 - 0.240)/61] = 6.421$. Because 6.421 is greater than 2.75548, reject the global null at $\alpha = .05$.

Work through the calculation, part (b)

The coefficient calculations are: planning sessions: $t = -1.600/0.500 = -3.200$, $p = 0.0022$, so reject the coefficient null; task-complexity score: $t = 0.420/0.170 = 2.471$, $p = 0.0163$, so reject the coefficient null; progress checks: $t = 0.160/0.140 = 1.143$, $p = 0.2576$, so do not reject the coefficient null. Thus 2 of the three displayed individual tests reject at the stated level.

Work through the calculation, part (c)

For predictor X_j , the individual null is $H_0 : \beta_j = 0$ conditional on every other term in this exact model. The global test asks one joint question about all three slopes. Rejecting it says at least one non-intercept population slope differs from zero under the model, but the global statistic does not name a predictor. Not rejecting it likewise is not proof that every population slope equals zero.

Interpret and check the result, part (d)

The two sets of decisions can differ because the global test evaluates the predictors jointly, whereas each t test isolates one conditional coefficient and its uncertainty. Shared predictor variation can make individual standard errors large even when the predictor set has joint explanatory value. Conversely, sampling variation can yield a small individual p-value in a model whose global test is not rejected. A p-value does not measure effect size, practical value, future prediction, or causality.

A04: Semipartial Correlation and Incremental R-Squared

T07-A04-V01: Guided practice and reasoning

Reason before calculating, part (a)

Each candidate is residualized against the current predictors, but the outcome remains in its original form. That one-sided residualization defines a semipartial correlation. A partial correlation would residualize both the candidate and the outcome against the current predictor set.

Work through the calculation, part (b)

Squaring each semipartial correlation gives the one-predictor increment:

Candidate	Semipartial r	Increment in R-squared	New R-squared
reflection sessions	0.240	0.0576	0.3576
study-partner meetings	0.100	0.0100	0.3100
planning checks	-0.180	0.0324	0.3324

Work through the calculation, part (c)

The largest squared semipartial correlation is 0.0576, for reflection sessions. A forward rule based only on the displayed candidates would add that predictor first, raising the sample R^2 from 0.300 to 0.3576.

Interpret and check the result, part (d)

The step ranks these three candidates by the additional sample variation each explains after the current predictors. Squaring removes the sign, so the sign of r_{sp} still matters for the direction of association even though it does not affect ΔR^2 . The ranking is conditional on the present model, candidates, and sample. After another predictor enters, shared variation changes what remains in every other candidate. Selection does not prove truth, causal effect, substantive importance, or performance on new data.

T07-A04-V02: Archive workflow and retrieval time

Reason before calculating, part (a)

Each candidate is residualized against the current predictors, but the outcome remains in its original form. That one-sided residualization defines a semipartial correlation. A partial correlation would residualize both the candidate and the outcome against the current predictor set.

Work through the calculation, part (b)

Squaring each semipartial correlation gives the one-predictor increment:

Candidate	Semipartial r	Increment in R-squared	New R-squared
catalog familiarity	-0.120	0.0144	0.2744
desk-map use	-0.270	0.0729	0.3329
mentor consultations	0.080	0.0064	0.2664

Work through the calculation, part (c)

The largest squared semipartial correlation is 0.0729, for desk-map use. A forward rule based only on the displayed candidates would add that predictor first, raising the sample R^2 from 0.260 to 0.3329.

Interpret and check the result, part (d)

The step ranks these three candidates by the additional sample variation each explains after the current predictors. Squaring removes the sign, so the sign of r_{sp} still matters for the direction of association even though it does not affect ΔR^2 . The ranking is conditional on the present model, candidates, and sample. After another predictor enters, shared variation changes what remains in every other candidate. Selection does not prove truth, causal effect, substantive importance, or performance on new data.

T07-A04-V03: Reading routines and comprehension

Reason before calculating, part (a)

Each candidate is residualized against the current predictors, but the outcome remains in its original form. That one-sided residualization defines a semipartial correlation. A partial correlation would residualize both the candidate and the outcome against the current predictor set.

Work through the calculation, part (b)

Squaring each semipartial correlation gives the one-predictor increment:

Candidate	Semipartial r	Increment in R-squared	New R-squared
annotation sessions	0.150	0.0225	0.3625
discussion posts	0.310	0.0961	0.4361
quiet-reading blocks	0.200	0.0400	0.3800

Work through the calculation, part (c)

The largest squared semipartial correlation is 0.0961, for discussion posts. A forward rule based only on the displayed candidates would add that predictor first, raising the sample R^2 from 0.340 to 0.4361.

Interpret and check the result, part (d)

The step ranks these three candidates by the additional sample variation each explains after the current predictors. Squaring removes the sign, so the sign of r_{sp} still matters for the direction of association even though it does not affect ΔR^2 . The ranking is conditional on the present model, candidates, and sample. After another predictor enters, shared variation changes what remains in every other candidate. Selection does not prove truth, causal effect, substantive importance, or performance on new data.

T07-A04-V04: Route rehearsal and navigation time

Reason before calculating, part (a)

Each candidate is residualized against the current predictors, but the outcome remains in its original form. That one-sided residualization defines a semipartial correlation. A partial correlation would residualize both the candidate and the outcome against the current predictor set.

Work through the calculation, part (b)

Squaring each semipartial correlation gives the one-predictor increment:

Candidate	Semipartial r	Increment in R-squared	New R-squared
landmark recall	-0.280	0.0784	0.3684
map checks	-0.140	0.0196	0.3096
route previews	0.190	0.0361	0.3261

Work through the calculation, part (c)

The largest squared semipartial correlation is 0.0784, for landmark recall. A forward rule based only on the displayed candidates would add that predictor first, raising the sample R^2 from 0.290 to 0.3684.

Interpret and check the result, part (d)

The step ranks these three candidates by the additional sample variation each explains after the current predictors. Squaring removes the sign, so the sign of r_{sp} still matters for the direction of association even though it does not affect ΔR^2 . The ranking is conditional on the present model, candidates, and sample. After another

predictor enters, shared variation changes what remains in every other candidate. Selection does not prove truth, causal effect, substantive importance, or performance on new data.

T07-A04-V05: Search practice and catalog accuracy

Reason before calculating, part (a)

Each candidate is residualized against the current predictors, but the outcome remains in its original form. That one-sided residualization defines a semipartial correlation. A partial correlation would residualize both the candidate and the outcome against the current predictor set.

Work through the calculation, part (b)

Squaring each semipartial correlation gives the one-predictor increment:

Candidate	Semipartial r	Increment in R-squared	New R-squared
query planning	0.110	0.0121	0.3821
keyword rehearsals	0.220	0.0484	0.4184
catalog hints used	0.290	0.0841	0.4541

Work through the calculation, part (c)

The largest squared semipartial correlation is 0.0841, for catalog hints used. A forward rule based only on the displayed candidates would add that predictor first, raising the sample R^2 from 0.370 to 0.4541.

Interpret and check the result, part (d)

The step ranks these three candidates by the additional sample variation each explains after the current predictors. Squaring removes the sign, so the sign of r_{sp} still matters for the direction of association even though it does not affect ΔR^2 . The ranking is conditional on the present model, candidates, and sample. After another predictor enters, shared variation changes what remains in every other candidate. Selection does not prove truth, causal effect, substantive importance, or performance on new data.

T07-A04-V06: Workshop participation and confidence

Reason before calculating, part (a)

Each candidate is residualized against the current predictors, but the outcome remains in its original form. That one-sided residualization defines a semipartial correlation. A partial correlation would residualize both the candidate and the outcome against the current predictor set.

Work through the calculation, part (b)

Squaring each semipartial correlation gives the one-predictor increment:

Candidate	Semipartial r	Increment in R-squared	New R-squared
reflection logs	0.260	0.0676	0.3876
peer meetings	0.170	0.0289	0.3489
practice demonstrations	-0.090	0.0081	0.3281

Work through the calculation, part (c)

The largest squared semipartial correlation is 0.0676, for reflection logs. A forward rule based only on the displayed candidates would add that predictor first, raising the sample R^2 from 0.320 to 0.3876.

Interpret and check the result, part (d)

The step ranks these three candidates by the additional sample variation each explains after the current predictors. Squaring removes the sign, so the sign of r_{sp} still matters for the direction of association even though

it does not affect ΔR^2 . The ranking is conditional on the present model, candidates, and sample. After another predictor enters, shared variation changes what remains in every other candidate. Selection does not prove truth, causal effect, substantive importance, or performance on new data.

T07-A04-V07: Focus blocks and task accuracy

Reason before calculating, part (a)

Each candidate is residualized against the current predictors, but the outcome remains in its original form. That one-sided residualization defines a semipartial correlation. A partial correlation would residualize both the candidate and the outcome against the current predictor set.

Work through the calculation, part (b)

Squaring each semipartial correlation gives the one-predictor increment:

Candidate	Semipartial r	Increment in R-squared	New R-squared
planning breaks	0.130	0.0169	0.2669
screen-free intervals	0.210	0.0441	0.2941
task previews	0.070	0.0049	0.2549

Work through the calculation, part (c)

The largest squared semipartial correlation is 0.0441, for screen-free intervals. A forward rule based only on the displayed candidates would add that predictor first, raising the sample R^2 from 0.250 to 0.2941.

Interpret and check the result, part (d)

The step ranks these three candidates by the additional sample variation each explains after the current predictors. Squaring removes the sign, so the sign of r_{sp} still matters for the direction of association even though it does not affect ΔR^2 . The ranking is conditional on the present model, candidates, and sample. After another predictor enters, shared variation changes what remains in every other candidate. Selection does not prove truth, causal effect, substantive importance, or performance on new data.

T07-A04-V08: Museum engagement and historical knowledge

Reason before calculating, part (a)

Each candidate is residualized against the current predictors, but the outcome remains in its original form. That one-sided residualization defines a semipartial correlation. A partial correlation would residualize both the candidate and the outcome against the current predictor set.

Work through the calculation, part (b)

Squaring each semipartial correlation gives the one-predictor increment:

Candidate	Semipartial r	Increment in R-squared	New R-squared
exhibit notes	0.180	0.0324	0.3424
guided-tour stops	0.120	0.0144	0.3244
follow-up readings	0.250	0.0625	0.3725

Work through the calculation, part (c)

The largest squared semipartial correlation is 0.0625, for follow-up readings. A forward rule based only on the displayed candidates would add that predictor first, raising the sample R^2 from 0.310 to 0.3725.

Interpret and check the result, part (d)

The step ranks these three candidates by the additional sample variation each explains after the current predictors. Squaring removes the sign, so the sign of r_{sp} still matters for the direction of association even though it does not affect ΔR^2 . The ranking is conditional on the present model, candidates, and sample. After another predictor enters, shared variation changes what remains in every other candidate. Selection does not prove truth, causal effect, substantive importance, or performance on new data.

T07-A04-V09: Peer feedback and revision quality

Reason before calculating, part (a)

Each candidate is residualized against the current predictors, but the outcome remains in its original form. That one-sided residualization defines a semipartial correlation. A partial correlation would residualize both the candidate and the outcome against the current predictor set.

Work through the calculation, part (b)

Squaring each semipartial correlation gives the one-predictor increment:

Candidate	Semipartial r	Increment in R-squared	New R-squared
revision planning	0.090	0.0081	0.3681
peer comments used	0.280	0.0784	0.4384
editing passes	0.160	0.0256	0.3856

Work through the calculation, part (c)

The largest squared semipartial correlation is 0.0784, for peer comments used. A forward rule based only on the displayed candidates would add that predictor first, raising the sample R^2 from 0.360 to 0.4384.

Interpret and check the result, part (d)

The step ranks these three candidates by the additional sample variation each explains after the current predictors. Squaring removes the sign, so the sign of r_{sp} still matters for the direction of association even though it does not affect ΔR^2 . The ranking is conditional on the present model, candidates, and sample. After another predictor enters, shared variation changes what remains in every other candidate. Selection does not prove truth, causal effect, substantive importance, or performance on new data.

T07-A04-V10: Planning sessions and completion time

Reason before calculating, part (a)

Each candidate is residualized against the current predictors, but the outcome remains in its original form. That one-sided residualization defines a semipartial correlation. A partial correlation would residualize both the candidate and the outcome against the current predictor set.

Work through the calculation, part (b)

Squaring each semipartial correlation gives the one-predictor increment:

Candidate	Semipartial r	Increment in R-squared	New R-squared
progress checks	-0.230	0.0529	0.3329
calendar reminders	-0.110	0.0121	0.2921
task previews	0.200	0.0400	0.3200

Work through the calculation, part (c)

The largest squared semipartial correlation is 0.0529, for progress checks. A forward rule based only on the displayed candidates would add that predictor first, raising the sample R^2 from 0.280 to 0.3329.

Interpret and check the result, part (d)

The step ranks these three candidates by the additional sample variation each explains after the current predictors. Squaring removes the sign, so the sign of r_{sp} still matters for the direction of association even though it does not affect ΔR^2 . The ranking is conditional on the present model, candidates, and sample. After another predictor enters, shared variation changes what remains in every other candidate. Selection does not prove truth, causal effect, substantive importance, or performance on new data.

A05: Comparing Prespecified Candidate Models With AIC

T07-A05-V01: Guided practice and reasoning

Reason before calculating, part (a)

For example, M1 gives $-2(-155.0) + 2(3) = 316.00$. Applying the same rule to all four models gives:

Model	AIC	Delta AIC
M1	316.00	21.00
M2	300.00	5.00
M3	295.00	0.00
M4	295.80	0.80

Work through the calculation, part (b)

Step 1 selects M2 because 300.00 is lower than the other displayed Step 1 values and lower than M1's 316.00. Step 2 selects M3 because its AIC is lower than the current M2 value. No later product term is selected on this forward path.

Work through the calculation, part (c)

The selected path coordinates are (0, 316.00), (1, 300.00), (2, 295.00). Plot step on the horizontal axis and AIC on the vertical axis, connect only consecutive selected models, and stop where the rule stops. The downward movements show improvements in the relative fit-complexity balance along this particular path.

Work through the calculation, part (d)

The final selected formula is $\text{reasoning score} \sim \text{guided-practice hours} + \text{prior-preparation score} + \text{reflection-session count}$. Its terms describe conditional fitted associations for this outcome and these cases. They do not by themselves identify causes.

Interpret and check the result, part (e)

A forward path recomputes the choice after each selected term, so an addition that looks useful at one stage can become redundant at another. The path can also stop before reaching the globally lowest AIC among combinations it never made reachable. AIC rewards fit but adds a complexity penalty. It does not establish that a selected model is the data-generating truth or that its predictions will generalize. New-data performance requires separate validation, and AIC values from different outcomes or case sets are not one comparable candidate family.

T07-A05-V02: Archive workflow and retrieval time

Reason before calculating, part (a)

For example, M1 gives $-2(-142.0) + 2(3) = 290.00$. Applying the same rule to all four models gives:

Model	AIC	Delta AIC
M1	290.00	14.40

Model	AIC	Delta AIC
M2	276.00	0.40
M3	276.80	1.20
M4	275.60	0.00

Work through the calculation, part (b)

Step 1 selects M2 because 276.00 is lower than the other displayed Step 1 values and lower than M1's 290.00. Step 2 stops because neither addition has an AIC below the current M2 value. No later product term is selected on this forward path.

Work through the calculation, part (c)

The selected path coordinates are (0, 290.00), (1, 276.00). Plot step on the horizontal axis and AIC on the vertical axis, connect only consecutive selected models, and stop where the rule stops. The downward movements show improvements in the relative fit-complexity balance along this particular path.

Work through the calculation, part (d)

The final selected formula is `retrieval time ~ checklist-practice sessions + archive-experience months`. Its terms describe conditional fitted associations for this outcome and these cases. They do not by themselves identify causes.

Interpret and check the result, part (e)

A forward path recomputes the choice after each selected term, so an addition that looks useful at one stage can become redundant at another. The path can also stop before reaching the globally lowest AIC among combinations it never made reachable. AIC rewards fit but adds a complexity penalty. It does not establish that a selected model is the data-generating truth or that its predictions will generalize. New-data performance requires separate validation, and AIC values from different outcomes or case sets are not one comparable candidate family.

T07-A05-V03: Reading routines and comprehension

Reason before calculating, part (a)

For example, M1 gives $-2(-180.0) + 2(3) = 366.00$. Applying the same rule to all four models gives:

Model	AIC	Delta AIC
M1	366.00	24.00
M2	348.00	6.00
M3	342.00	0.00
M4	343.00	1.00

Work through the calculation, part (b)

Step 1 selects M2 because 348.00 is lower than the other displayed Step 1 values and lower than M1's 366.00. Step 2 selects M3 because its AIC is lower than the current M2 value. No later product term is selected on this forward path.

Work through the calculation, part (c)

The selected path coordinates are (0, 366.00), (1, 348.00), (2, 342.00). Plot step on the horizontal axis and AIC on the vertical axis, connect only consecutive selected models, and stop where the rule stops. The downward movements show improvements in the relative fit-complexity balance along this particular path.

Work through the calculation, part (d)

The final selected formula is comprehension score $\sim$ weekly reading hours + baseline-vocabulary score + annotation-session count. Its terms describe conditional fitted associations for this outcome and these cases. They do not by themselves identify causes.

Interpret and check the result, part (e)

A forward path recomputes the choice after each selected term, so an addition that looks useful at one stage can become redundant at another. The path can also stop before reaching the globally lowest AIC among combinations it never made reachable. AIC rewards fit but adds a complexity penalty. It does not establish that a selected model is the data-generating truth or that its predictions will generalize. New-data performance requires separate validation, and AIC values from different outcomes or case sets are not one comparable candidate family.

T07-A05-V04: Route rehearsal and navigation time

Reason before calculating, part (a)

For example, M1 gives $-2(-130.0) + 2(3) = 266.00$. Applying the same rule to all four models gives:

Model	AIC	Delta AIC
M1	266.00	6.00
M2	260.00	0.00
M3	261.00	1.00
M4	262.40	2.40

Work through the calculation, part (b)

Step 1 selects M2 because 260.00 is lower than the other displayed Step 1 values and lower than M1's 266.00. Step 2 stops because neither addition has an AIC below the current M2 value. No later product term is selected on this forward path.

Work through the calculation, part (c)

The selected path coordinates are (0, 266.00), (1, 260.00). Plot step on the horizontal axis and AIC on the vertical axis, connect only consecutive selected models, and stop where the rule stops. The downward movements show improvements in the relative fit-complexity balance along this particular path.

Work through the calculation, part (d)

The final selected formula is navigation time $\sim$ route-rehearsal attempts + route-familiarity score. Its terms describe conditional fitted associations for this outcome and these cases. They do not by themselves identify causes.

Interpret and check the result, part (e)

A forward path recomputes the choice after each selected term, so an addition that looks useful at one stage can become redundant at another. The path can also stop before reaching the globally lowest AIC among combinations it never made reachable. AIC rewards fit but adds a complexity penalty. It does not establish that a selected model is the data-generating truth or that its predictions will generalize. New-data performance requires separate validation, and AIC values from different outcomes or case sets are not one comparable candidate family.

T07-A05-V05: Search practice and catalog accuracy

Reason before calculating, part (a)

For example, M1 gives $-2(-200.0) + 2(3) = 406.00$. Applying the same rule to all four models gives:

Model	AIC	Delta AIC
M1	406.00	34.00
M2	384.00	12.00
M3	376.00	4.00
M4	372.00	0.00

Work through the calculation, part (b)

Step 1 selects M2 because 384.00 is lower than the other displayed Step 1 values and lower than M1's 406.00. Step 2 selects M3 because its AIC is lower than the current M2 value. Step 3 then selects M4 because its AIC is below M3.

Work through the calculation, part (c)

The selected path coordinates are (0, 406.00), (1, 384.00), (2, 376.00), (3, 372.00). Plot step on the horizontal axis and AIC on the vertical axis, connect only consecutive selected models, and stop where the rule stops. The downward movements show improvements in the relative fit-complexity balance along this particular path.

Work through the calculation, part (d)

The final selected formula is catalog-accuracy score ~ search-practice sets + prior catalog-knowledge score + query-planning score + a prespecified product term. Its terms describe conditional fitted associations for this outcome and these cases. They do not by themselves identify causes.

Interpret and check the result, part (e)

A forward path recomputes the choice after each selected term, so an addition that looks useful at one stage can become redundant at another. The path can also stop before reaching the globally lowest AIC among combinations it never made reachable. AIC rewards fit but adds a complexity penalty. It does not establish that a selected model is the data-generating truth or that its predictions will generalize. New-data performance requires separate validation, and AIC values from different outcomes or case sets are not one comparable candidate family.

T07-A05-V06: Workshop participation and confidence

Reason before calculating, part (a)

For example, M1 gives $-2(-165.0) + 2(3) = 336.00$. Applying the same rule to all four models gives:

Model	AIC	Delta AIC
M1	336.00	14.00
M2	322.00	0.00
M3	322.80	0.80
M4	323.60	1.60

Work through the calculation, part (b)

Step 1 selects M2 because 322.00 is lower than the other displayed Step 1 values and lower than M1's 336.00. Step 2 stops because neither addition has an AIC below the current M2 value. No later product term is selected on this forward path.

Work through the calculation, part (c)

The selected path coordinates are (0, 336.00), (1, 322.00). Plot step on the horizontal axis and AIC on the vertical axis, connect only consecutive selected models, and stop where the rule stops. The downward movements show improvements in the relative fit-complexity balance along this particular path.

Work through the calculation, part (d)

The final selected formula is `confidence score ~ workshop sessions + baseline-confidence score`. Its terms describe conditional fitted associations for this outcome and these cases. They do not by themselves identify causes.

Interpret and check the result, part (e)

A forward path recomputes the choice after each selected term, so an addition that looks useful at one stage can become redundant at another. The path can also stop before reaching the globally lowest AIC among combinations it never made reachable. AIC rewards fit but adds a complexity penalty. It does not establish that a selected model is the data-generating truth or that its predictions will generalize. New-data performance requires separate validation, and AIC values from different outcomes or case sets are not one comparable candidate family.

T07-A05-V07: Focus blocks and task accuracy

Reason before calculating, part (a)

For example, M1 gives $-2(-175.0) + 2(3) = 356.00$. Applying the same rule to all four models gives:

Model	AIC	Delta AIC
M1	356.00	22.00
M2	340.00	6.00
M3	334.00	0.00
M4	334.40	0.40

Work through the calculation, part (b)

Step 1 selects M2 because 340.00 is lower than the other displayed Step 1 values and lower than M1's 356.00. Step 2 selects M3 because its AIC is lower than the current M2 value. No later product term is selected on this forward path.

Work through the calculation, part (c)

The selected path coordinates are (0, 356.00), (1, 340.00), (2, 334.00). Plot step on the horizontal axis and AIC on the vertical axis, connect only consecutive selected models, and stop where the rule stops. The downward movements show improvements in the relative fit-complexity balance along this particular path.

Work through the calculation, part (d)

The final selected formula is `task-accuracy score ~ notification-free blocks + sleep duration in hours + planning-break count`. Its terms describe conditional fitted associations for this outcome and these cases. They do not by themselves identify causes.

Interpret and check the result, part (e)

A forward path recomputes the choice after each selected term, so an addition that looks useful at one stage can become redundant at another. The path can also stop before reaching the globally lowest AIC among combinations it never made reachable. AIC rewards fit but adds a complexity penalty. It does not establish that a selected model is the data-generating truth or that its predictions will generalize. New-data performance requires separate validation, and AIC values from different outcomes or case sets are not one comparable candidate family.

T07-A05-V08: Museum engagement and historical knowledge

Reason before calculating, part (a)

For example, M1 gives $-2(-145.0) + 2(3) = 296.00$. Applying the same rule to all four models gives:

Model	AIC	Delta AIC
M1	296.00	11.20
M2	288.00	3.20
M3	286.00	1.20
M4	284.80	0.00

Work through the calculation, part (b)

Step 1 selects M2 because 288.00 is lower than the other displayed Step 1 values and lower than M1's 296.00. Step 2 selects M3 because its AIC is lower than the current M2 value. Step 3 then selects M4 because its AIC is below M3.

Work through the calculation, part (c)

The selected path coordinates are (0, 296.00), (1, 288.00), (2, 286.00), (3, 284.80). Plot step on the horizontal axis and AIC on the vertical axis, connect only consecutive selected models, and stop where the rule stops. The downward movements show improvements in the relative fit-complexity balance along this particular path.

Work through the calculation, part (d)

The final selected formula is `historical-knowledge score ~ museum visits + prior-history score + exhibit-note count + a prespecified product term`. Its terms describe conditional fitted associations for this outcome and these cases. They do not by themselves identify causes.

Interpret and check the result, part (e)

A forward path recomputes the choice after each selected term, so an addition that looks useful at one stage can become redundant at another. The path can also stop before reaching the globally lowest AIC among combinations it never made reachable. AIC rewards fit but adds a complexity penalty. It does not establish that a selected model is the data-generating truth or that its predictions will generalize. New-data performance requires separate validation, and AIC values from different outcomes or case sets are not one comparable candidate family.

T07-A05-V09: Peer feedback and revision quality

Reason before calculating, part (a)

For example, M1 gives $-2(-190.0) + 2(3) = 386.00$. Applying the same rule to all four models gives:

Model	AIC	Delta AIC
M1	386.00	16.00
M2	370.00	0.00
M3	370.60	0.60
M4	371.80	1.80

Work through the calculation, part (b)

Step 1 selects M2 because 370.00 is lower than the other displayed Step 1 values and lower than M1's 386.00. Step 2 stops because neither addition has an AIC below the current M2 value. No later product term is selected on this forward path.

Work through the calculation, part (c)

The selected path coordinates are (0, 386.00), (1, 370.00). Plot step on the horizontal axis and AIC on the vertical axis, connect only consecutive selected models, and stop where the rule stops. The downward movements show improvements in the relative fit-complexity balance along this particular path.

Work through the calculation, part (d)

The final selected formula is revision-quality score ~ peer-feedback rounds + baseline-writing score. Its terms describe conditional fitted associations for this outcome and these cases. They do not by themselves identify causes.

Interpret and check the result, part (e)

A forward path recomputes the choice after each selected term, so an addition that looks useful at one stage can become redundant at another. The path can also stop before reaching the globally lowest AIC among combinations it never made reachable. AIC rewards fit but adds a complexity penalty. It does not establish that a selected model is the data-generating truth or that its predictions will generalize. New-data performance requires separate validation, and AIC values from different outcomes or case sets are not one comparable candidate family.

T07-A05-V10: Planning sessions and completion time

Reason before calculating, part (a)

For example, M1 gives $-2(-158.0) + 2(3) = 322.00$. Applying the same rule to all four models gives:

Model	AIC	Delta AIC
M1	322.00	22.00
M2	306.00	6.00
M3	300.00	0.00
M4	300.80	0.80

Work through the calculation, part (b)

Step 1 selects M2 because 306.00 is lower than the other displayed Step 1 values and lower than M1's 322.00. Step 2 selects M3 because its AIC is lower than the current M2 value. No later product term is selected on this forward path.

Work through the calculation, part (c)

The selected path coordinates are (0, 322.00), (1, 306.00), (2, 300.00). Plot step on the horizontal axis and AIC on the vertical axis, connect only consecutive selected models, and stop where the rule stops. The downward movements show improvements in the relative fit-complexity balance along this particular path.

Work through the calculation, part (d)

The final selected formula is completion time ~ planning sessions + task-complexity score + progress-check count. Its terms describe conditional fitted associations for this outcome and these cases. They do not by themselves identify causes.

Interpret and check the result, part (e)

A forward path recomputes the choice after each selected term, so an addition that looks useful at one stage can become redundant at another. The path can also stop before reaching the globally lowest AIC among combinations it never made reachable. AIC rewards fit but adds a complexity penalty. It does not establish that a selected model is the data-generating truth or that its predictions will generalize. New-data performance requires separate validation, and AIC values from different outcomes or case sets are not one comparable candidate family.

A07: Interpreting an Additive Group Model

T07-A07-V01: Tutorial support and reasoning

Reason before calculating, part (a)

For Self-guided, set $G = 0$: $\hat{Y} = 42.00 + (3.00)X$. For Tutored, set $G = 1$: $\hat{Y} = 47.00 + (3.00)X$. The intercept 42.00 is the fitted reasoning score for Self-guided when practice hours equals zero. It may be mathematically necessary but substantively unhelpful if zero lies outside the meaningful range.

Work through the calculation, part (b)

Within either group, a one-unit increase in practice hours is associated with a fitted change of 3.00 units in reasoning score. At the same value of practice hours, Tutored is fitted 5.00 units higher than Self-guided. “At the same value” expresses the model’s conditional comparison, not an intervention.

Work through the calculation, part (c)

Substitution gives:

Group	X	Fitted reasoning score
Self-guided	2.0	48.00
Self-guided	6.0	60.00
Tutored	2.0	53.00
Tutored	6.0	65.00

Interpret and check the result, part (d)

Both equations have slope 3.00, so equal horizontal changes produce equal fitted vertical changes. Their intercepts differ by 5.00, and subtracting the two fitted values at either displayed X gives that same constant gap. The model contains no XG product term, so it imposes parallel fitted lines. The gap is an adjusted association; without suitable design and assumptions, it does not prove that changing group membership would change the outcome.

T07-A07-V02: Archive experience and retrieval

Reason before calculating, part (a)

For New staff, set $G = 0$: $\hat{Y} = 36.00 + (-1.80)X$. For Experienced staff, set $G = 1$: $\hat{Y} = 32.00 + (-1.80)X$. The intercept 36.00 is the fitted retrieval time for New staff when practice sessions equals zero. It may be mathematically necessary but substantively unhelpful if zero lies outside the meaningful range.

Work through the calculation, part (b)

Within either group, a one-unit increase in practice sessions is associated with a fitted change of -1.80 units in retrieval time. At the same value of practice sessions, Experienced staff is fitted 4.00 units lower than New staff. “At the same value” expresses the model’s conditional comparison, not an intervention.

Work through the calculation, part (c)

Substitution gives:

Group	X	Fitted retrieval time
New staff	1.0	34.20
New staff	5.0	27.00
Experienced staff	1.0	30.20
Experienced staff	5.0	23.00

Interpret and check the result, part (d)

Both equations have slope -1.80 , so equal horizontal changes produce equal fitted vertical changes. Their intercepts differ by -4.00 , and subtracting the two fitted values at either displayed X gives that same constant gap. The model contains no XG product term, so it imposes parallel fitted lines. The gap is an adjusted association; without suitable design and assumptions, it does not prove that changing group membership would change the outcome.

T07-A07-V03: Reading format and comprehension

Reason before calculating, part (a)

For Print, set $G = 0$: $\hat{Y} = 51.00 + (2.20)X$. For Digital, set $G = 1$: $\hat{Y} = 48.50 + (2.20)X$. The intercept 51.00 is the fitted comprehension score for Print when reading hours equals zero. It may be mathematically necessary but substantively unhelpful if zero lies outside the meaningful range.

Work through the calculation, part (b)

Within either group, a one-unit increase in reading hours is associated with a fitted change of 2.20 units in comprehension score. At the same value of reading hours, Digital is fitted 2.50 units lower than Print. “At the same value” expresses the model’s conditional comparison, not an intervention.

Work through the calculation, part (c)

Substitution gives:

Group	X	Fitted comprehension score
Print	2.0	55.40
Print	7.0	66.40
Digital	2.0	52.90
Digital	7.0	63.90

Interpret and check the result, part (d)

Both equations have slope 2.20 , so equal horizontal changes produce equal fitted vertical changes. Their intercepts differ by -2.50 , and subtracting the two fitted values at either displayed X gives that same constant gap. The model contains no XG product term, so it imposes parallel fitted lines. The gap is an adjusted association; without suitable design and assumptions, it does not prove that changing group membership would change the outcome.

T07-A07-V04: Route aid and navigation

Reason before calculating, part (a)

For Paper map, set $G = 0$: $\hat{Y} = 44.00 + (-2.00)X$. For App map, set $G = 1$: $\hat{Y} = 41.00 + (-2.00)X$. The intercept 44.00 is the fitted navigation time for Paper map when rehearsal attempts equals zero. It may be mathematically necessary but substantively unhelpful if zero lies outside the meaningful range.

Work through the calculation, part (b)

Within either group, a one-unit increase in rehearsal attempts is associated with a fitted change of -2.00 units in navigation time. At the same value of rehearsal attempts, App map is fitted 3.00 units lower than Paper map. “At the same value” expresses the model’s conditional comparison, not an intervention.

Work through the calculation, part (c)

Substitution gives:

Group	X	Fitted navigation time
Paper map	1.0	42.00
Paper map	4.0	36.00
App map	1.0	39.00
App map	4.0	33.00

Interpret and check the result, part (d)

Both equations have slope -2.00 , so equal horizontal changes produce equal fitted vertical changes. Their intercepts differ by -3.00 , and subtracting the two fitted values at either displayed X gives that same constant gap. The model contains no XG product term, so it imposes parallel fitted lines. The gap is an adjusted association; without suitable design and assumptions, it does not prove that changing group membership would change the outcome.

T07-A07-V05: Search guide and accuracy

Reason before calculating, part (a)

For No guide, set $G = 0$: $\hat{Y} = 55.00 + (2.50)X$. For Checklist, set $G = 1$: $\hat{Y} = 59.00 + (2.50)X$. The intercept 55.00 is the fitted accuracy score for No guide when practice sets equals zero. It may be mathematically necessary but substantively unhelpful if zero lies outside the meaningful range.

Work through the calculation, part (b)

Within either group, a one-unit increase in practice sets is associated with a fitted change of 2.50 units in accuracy score. At the same value of practice sets, Checklist is fitted 4.00 units higher than No guide. “At the same value” expresses the model’s conditional comparison, not an intervention.

Work through the calculation, part (c)

Substitution gives:

Group	X	Fitted accuracy score
No guide	0.0	55.00
No guide	4.0	65.00
Checklist	0.0	59.00
Checklist	4.0	69.00

Interpret and check the result, part (d)

Both equations have slope 2.50 , so equal horizontal changes produce equal fitted vertical changes. Their intercepts differ by 4.00 , and subtracting the two fitted values at either displayed X gives that same constant gap. The model contains no XG product term, so it imposes parallel fitted lines. The gap is an adjusted association; without suitable design and assumptions, it does not prove that changing group membership would change the outcome.

T07-A07-V06: Workshop mode and confidence

Reason before calculating, part (a)

For Online, set $G = 0$: $\hat{Y} = 38.00 + (3.20)X$. For In person, set $G = 1$: $\hat{Y} = 41.50 + (3.20)X$. The intercept 38.00 is the fitted confidence score for Online when sessions attended equals zero. It may be mathematically necessary but substantively unhelpful if zero lies outside the meaningful range.

Work through the calculation, part (b)

Within either group, a one-unit increase in sessions attended is associated with a fitted change of 3.20 units in confidence score. At the same value of sessions attended, In person is fitted 3.50 units higher than Online. “At the same value” expresses the model’s conditional comparison, not an intervention.

Work through the calculation, part (c)

Substitution gives:

Group	X	Fitted confidence score
Online	1.0	41.20
Online	5.0	54.00
In person	1.0	44.70
In person	5.0	57.50

Interpret and check the result, part (d)

Both equations have slope 3.20, so equal horizontal changes produce equal fitted vertical changes. Their intercepts differ by 3.50, and subtracting the two fitted values at either displayed X gives that same constant gap. The model contains no XG product term, so it imposes parallel fitted lines. The gap is an adjusted association; without suitable design and assumptions, it does not prove that changing group membership would change the outcome.

T07-A07-V07: Focus setting and accuracy

Reason before calculating, part (a)

For Shared room, set $G = 0$: $\hat{Y} = 60.00 + (1.70)X$. For Quiet room, set $G = 1$: $\hat{Y} = 64.50 + (1.70)X$. The intercept 60.00 is the fitted task-accuracy score for Shared room when focus blocks equals zero. It may be mathematically necessary but substantively unhelpful if zero lies outside the meaningful range.

Work through the calculation, part (b)

Within either group, a one-unit increase in focus blocks is associated with a fitted change of 1.70 units in task-accuracy score. At the same value of focus blocks, Quiet room is fitted 4.50 units higher than Shared room. “At the same value” expresses the model’s conditional comparison, not an intervention.

Work through the calculation, part (c)

Substitution gives:

Group	X	Fitted task-accuracy score
Shared room	2.0	63.40
Shared room	8.0	73.60
Quiet room	2.0	67.90
Quiet room	8.0	78.10

Interpret and check the result, part (d)

Both equations have slope 1.70, so equal horizontal changes produce equal fitted vertical changes. Their intercepts differ by 4.50, and subtracting the two fitted values at either displayed X gives that same constant gap. The model contains no XG product term, so it imposes parallel fitted lines. The gap is an adjusted association; without suitable design and assumptions, it does not prove that changing group membership would change the outcome.

T07-A07-V08: Museum guide and knowledge

Reason before calculating, part (a)

For Self-guided, set $G = 0$: $\hat{Y} = 47.00 + (4.00)X$. For Guided, set $G = 1$: $\hat{Y} = 53.00 + (4.00)X$. The intercept 47.00 is the fitted knowledge score for Self-guided when visits equals zero. It may be mathematically necessary but substantively unhelpful if zero lies outside the meaningful range.

Work through the calculation, part (b)

Within either group, a one-unit increase in visits is associated with a fitted change of 4.00 units in knowledge score. At the same value of visits, Guided is fitted 6.00 units higher than Self-guided. “At the same value” expresses the model’s conditional comparison, not an intervention.

Work through the calculation, part (c)

Substitution gives:

Group	X	Fitted knowledge score
Self-guided	0.0	47.00
Self-guided	3.0	59.00
Guided	0.0	53.00
Guided	3.0	65.00

Interpret and check the result, part (d)

Both equations have slope 4.00, so equal horizontal changes produce equal fitted vertical changes. Their intercepts differ by 6.00, and subtracting the two fitted values at either displayed X gives that same constant gap. The model contains no XG product term, so it imposes parallel fitted lines. The gap is an adjusted association; without suitable design and assumptions, it does not prove that changing group membership would change the outcome.

T07-A07-V09: Feedback mode and revision

Reason before calculating, part (a)

For Written, set $G = 0$: $\hat{Y} = 52.00 + (3.50)X$. For Conversation, set $G = 1$: $\hat{Y} = 54.00 + (3.50)X$. The intercept 52.00 is the fitted revision score for Written when feedback rounds equals zero. It may be mathematically necessary but substantively unhelpful if zero lies outside the meaningful range.

Work through the calculation, part (b)

Within either group, a one-unit increase in feedback rounds is associated with a fitted change of 3.50 units in revision score. At the same value of feedback rounds, Conversation is fitted 2.00 units higher than Written. “At the same value” expresses the model’s conditional comparison, not an intervention.

Work through the calculation, part (c)

Substitution gives:

Group	X	Fitted revision score
Written	1.0	55.50
Written	4.0	66.00
Conversation	1.0	57.50
Conversation	4.0	68.00

Interpret and check the result, part (d)

Both equations have slope 3.50, so equal horizontal changes produce equal fitted vertical changes. Their intercepts differ by 2.00, and subtracting the two fitted values at either displayed X gives that same constant gap. The model contains no XG product term, so it imposes parallel fitted lines. The gap is an adjusted association; without suitable design and assumptions, it does not prove that changing group membership would change the outcome.

T07-A07-V10: Planning format and completion

Reason before calculating, part (a)

For Paper, set $G = 0$: $\hat{Y} = 70.00 + (-2.40)X$. For Digital, set $G = 1$: $\hat{Y} = 66.50 + (-2.40)X$. The intercept 70.00 is the fitted completion time for Paper when planning sessions equals zero. It may be mathematically necessary but substantively unhelpful if zero lies outside the meaningful range.

Work through the calculation, part (b)

Within either group, a one-unit increase in planning sessions is associated with a fitted change of -2.40 units in completion time. At the same value of planning sessions, Digital is fitted 3.50 units lower than Paper. “At the same value” expresses the model’s conditional comparison, not an intervention.

Work through the calculation, part (c)

Substitution gives:

Group	X	Fitted completion time
Paper	1.0	67.60
Paper	6.0	55.60
Digital	1.0	64.10
Digital	6.0	52.10

Interpret and check the result, part (d)

Both equations have slope -2.40 , so equal horizontal changes produce equal fitted vertical changes. Their intercepts differ by -3.50 , and subtracting the two fitted values at either displayed X gives that same constant gap. The model contains no XG product term, so it imposes parallel fitted lines. The gap is an adjusted association; without suitable design and assumptions, it does not prove that changing group membership would change the outcome.

A08: Releveling Without Changing Fitted Relationships

T07-A08-V01: Practice format releveling

Reason before calculating, part (a)

The new reference is the old $G = 1$ group, so its old intercept becomes the new intercept: $b'_0 = 40.00 + (4.50) = 44.50$. The common slope remains $b'_1 = 2.80$. The contrast reverses direction, so $b'_2 = -(4.50) = -4.50$.

Work through the calculation, part (b)

For Partnered, $H = 0$, giving $\hat{Y} = 44.50 + (2.80)X$. For Independent, $H = 1$, giving $\hat{Y} = 44.50 + (2.80)X + (-4.50) = 40.00 + (2.80)X$. At the same X , Independent is fitted 4.50 units lower than Partnered.

Work through the calculation, part (c)

Both codings give:

Group	X	Fit from old coding	Fit from new coding
Independent	1.0	42.80	42.80

Group	X	Fit from old coding	Fit from new coding
Independent	5.0	54.00	54.00
Partnered	1.0	47.30	47.30
Partnered	5.0	58.50	58.50

Interpret and check the result, part (d)

Every row has identical fitted values under the two codings. Releveling changes which group is represented by the intercept and reverses the displayed group contrast, but it describes the same two lines. Because each case keeps the same fitted value, subtracting that fit from its observed outcome also leaves every residual unchanged. Reference choice changes representation, not model fit or the underlying fitted relationships.

T07-A08-V02: Archive role releveling

Reason before calculating, part (a)

The new reference is the old $G = 1$ group, so its old intercept becomes the new intercept: $b'_0 = 35.00 + (-5.00) = 30.00$. The common slope remains $b'_1 = -1.60$. The contrast reverses direction, so $b'_2 = -(-5.00) = 5.00$.

Work through the calculation, part (b)

For Coordinator, $H = 0$, giving $\hat{Y} = 30.00 + (-1.60)X$. For Assistant, $H = 1$, giving $\hat{Y} = 30.00 + (-1.60)X + (5.00) = 35.00 + (-1.60)X$. At the same X , Assistant is fitted 5.00 units higher than Coordinator.

Work through the calculation, part (c)

Both codings give:

Group	X	Fit from old coding	Fit from new coding
Assistant	0.0	35.00	35.00
Assistant	4.0	28.60	28.60
Coordinator	0.0	30.00	30.00
Coordinator	4.0	23.60	23.60

Interpret and check the result, part (d)

Every row has identical fitted values under the two codings. Releveling changes which group is represented by the intercept and reverses the displayed group contrast, but it describes the same two lines. Because each case keeps the same fitted value, subtracting that fit from its observed outcome also leaves every residual unchanged. Reference choice changes representation, not model fit or the underlying fitted relationships.

T07-A08-V03: Reading medium releveling

Reason before calculating, part (a)

The new reference is the old $G = 1$ group, so its old intercept becomes the new intercept: $b'_0 = 50.00 + (-3.00) = 47.00$. The common slope remains $b'_1 = 2.00$. The contrast reverses direction, so $b'_2 = -(-3.00) = 3.00$.

Work through the calculation, part (b)

For Audio, $H = 0$, giving $\hat{Y} = 47.00 + (2.00)X$. For Print, $H = 1$, giving $\hat{Y} = 47.00 + (2.00)X + (3.00) = 50.00 + (2.00)X$. At the same X , Print is fitted 3.00 units higher than Audio.

Work through the calculation, part (c)

Both codings give:

Group	X	Fit from old coding	Fit from new coding
Print	2.0	54.00	54.00
Print	6.0	62.00	62.00
Audio	2.0	51.00	51.00
Audio	6.0	59.00	59.00

Interpret and check the result, part (d)

Every row has identical fitted values under the two codings. Releveling changes which group is represented by the intercept and reverses the displayed group contrast, but it describes the same two lines. Because each case keeps the same fitted value, subtracting that fit from its observed outcome also leaves every residual unchanged. Reference choice changes representation, not model fit or the underlying fitted relationships.

T07-A08-V04: Navigation display releveling

Reason before calculating, part (a)

The new reference is the old $G = 1$ group, so its old intercept becomes the new intercept: $b'_0 = 46.00 + (-4.00) = 42.00$. The common slope remains $b'_1 = -2.20$. The contrast reverses direction, so $b'_2 = -(-4.00) = 4.00$.

Work through the calculation, part (b)

For Interactive, $H = 0$, giving $\hat{Y} = 42.00 + (-2.20)X$. For Static, $H = 1$, giving $\hat{Y} = 42.00 + (-2.20)X + (4.00) = 46.00 + (-2.20)X$. At the same X , Static is fitted 4.00 units higher than Interactive.

Work through the calculation, part (c)

Both codings give:

Group	X	Fit from old coding	Fit from new coding
Static	1.0	43.80	43.80
Static	5.0	35.00	35.00
Interactive	1.0	39.80	39.80
Interactive	5.0	31.00	31.00

Interpret and check the result, part (d)

Every row has identical fitted values under the two codings. Releveling changes which group is represented by the intercept and reverses the displayed group contrast, but it describes the same two lines. Because each case keeps the same fitted value, subtracting that fit from its observed outcome also leaves every residual unchanged. Reference choice changes representation, not model fit or the underlying fitted relationships.

T07-A08-V05: Catalog aid releveling

Reason before calculating, part (a)

The new reference is the old $G = 1$ group, so its old intercept becomes the new intercept: $b'_0 = 53.00 + (3.00) = 56.00$. The common slope remains $b'_1 = 2.60$. The contrast reverses direction, so $b'_2 = -(3.00) = -3.00$.

Work through the calculation, part (b)

For Search bar, $H = 0$, giving $\hat{Y} = 56.00 + (2.60)X$. For Index, $H = 1$, giving $\hat{Y} = 56.00 + (2.60)X + (-3.00) = 53.00 + (2.60)X$. At the same X , Index is fitted 3.00 units lower than Search bar.

Work through the calculation, part (c)

Both codings give:

Group	X	Fit from old coding	Fit from new coding
Index	0.0	53.00	53.00
Index	3.0	60.80	60.80
Search bar	0.0	56.00	56.00
Search bar	3.0	63.80	63.80

Interpret and check the result, part (d)

Every row has identical fitted values under the two codings. Releveling changes which group is represented by the intercept and reverses the displayed group contrast, but it describes the same two lines. Because each case keeps the same fitted value, subtracting that fit from its observed outcome also leaves every residual unchanged. Reference choice changes representation, not model fit or the underlying fitted relationships.

T07-A08-V06: Workshop setting releveling

Reason before calculating, part (a)

The new reference is the old $G = 1$ group, so its old intercept becomes the new intercept: $b'_0 = 37.00 + (5.00) = 42.00$. The common slope remains $b'_1 = 3.00$. The contrast reverses direction, so $b'_2 = -(5.00) = -5.00$.

Work through the calculation, part (b)

For Classroom, $H = 0$, giving $\hat{Y} = 42.00 + (3.00)X$. For Remote, $H = 1$, giving $\hat{Y} = 42.00 + (3.00)X + (-5.00) = 37.00 + (3.00)X$. At the same X , Remote is fitted 5.00 units lower than Classroom.

Work through the calculation, part (c)

Both codings give:

Group	X	Fit from old coding	Fit from new coding
Remote	1.0	40.00	40.00
Remote	4.0	49.00	49.00
Classroom	1.0	45.00	45.00
Classroom	4.0	54.00	54.00

Interpret and check the result, part (d)

Every row has identical fitted values under the two codings. Releveling changes which group is represented by the intercept and reverses the displayed group contrast, but it describes the same two lines. Because each case keeps the same fitted value, subtracting that fit from its observed outcome also leaves every residual unchanged. Reference choice changes representation, not model fit or the underlying fitted relationships.

T07-A08-V07: Focus room releveling

Reason before calculating, part (a)

The new reference is the old $G = 1$ group, so its old intercept becomes the new intercept: $b'_0 = 59.00 + (4.00) = 63.00$. The common slope remains $b'_1 = 1.80$. The contrast reverses direction, so $b'_2 = -(4.00) = -4.00$.

Work through the calculation, part (b)

For Private room, $H = 0$, giving $\hat{Y} = 63.00 + (1.80)X$. For Open room, $H = 1$, giving $\hat{Y} = 63.00 + (1.80)X + (-4.00) = 59.00 + (1.80)X$. At the same X , Open room is fitted 4.00 units lower than Private room.

Work through the calculation, part (c)

Both codings give:

Group	X	Fit from old coding	Fit from new coding
Open room	2.0	62.60	62.60
Open room	7.0	71.60	71.60
Private room	2.0	66.60	66.60
Private room	7.0	75.60	75.60

Interpret and check the result, part (d)

Every row has identical fitted values under the two codings. Releveling changes which group is represented by the intercept and reverses the displayed group contrast, but it describes the same two lines. Because each case keeps the same fitted value, subtracting that fit from its observed outcome also leaves every residual unchanged. Reference choice changes representation, not model fit or the underlying fitted relationships.

T07-A08-V08: Museum route releveling

Reason before calculating, part (a)

The new reference is the old $G = 1$ group, so its old intercept becomes the new intercept: $b'_0 = 45.00 + (6.50) = 51.50$. The common slope remains $b'_1 = 4.20$. The contrast reverses direction, so $b'_2 = -(6.50) = -6.50$.

Work through the calculation, part (b)

For Curated route, $H = 0$, giving $\hat{Y} = 51.50 + (4.20)X$. For Free route, $H = 1$, giving $\hat{Y} = 51.50 + (4.20)X + (-6.50) = 45.00 + (4.20)X$. At the same X , Free route is fitted 6.50 units lower than Curated route.

Work through the calculation, part (c)

Both codings give:

Group	X	Fit from old coding	Fit from new coding
Free route	0.0	45.00	45.00
Free route	3.0	57.60	57.60
Curated route	0.0	51.50	51.50
Curated route	3.0	64.10	64.10

Interpret and check the result, part (d)

Every row has identical fitted values under the two codings. Releveling changes which group is represented by the intercept and reverses the displayed group contrast, but it describes the same two lines. Because each case keeps the same fitted value, subtracting that fit from its observed outcome also leaves every residual unchanged. Reference choice changes representation, not model fit or the underlying fitted relationships.

T07-A08-V09: Revision meeting releveling

Reason before calculating, part (a)

The new reference is the old $G = 1$ group, so its old intercept becomes the new intercept: $b'_0 = 51.00 + (2.50) = 53.50$. The common slope remains $b'_1 = 3.40$. The contrast reverses direction, so $b'_2 = -(2.50) = -2.50$.

Work through the calculation, part (b)

For Live, $H = 0$, giving $\hat{Y} = 53.50 + (3.40)X$. For Asynchronous, $H = 1$, giving $\hat{Y} = 53.50 + (3.40)X + (-2.50) = 51.00 + (3.40)X$. At the same X , Asynchronous is fitted 2.50 units lower than Live.

Work through the calculation, part (c)

Both codings give:

Group	X	Fit from old coding	Fit from new coding
Asynchronous	1.0	54.40	54.40
Asynchronous	5.0	68.00	68.00
Live	1.0	56.90	56.90
Live	5.0	70.50	70.50

Interpret and check the result, part (d)

Every row has identical fitted values under the two codings. Releveling changes which group is represented by the intercept and reverses the displayed group contrast, but it describes the same two lines. Because each case keeps the same fitted value, subtracting that fit from its observed outcome also leaves every residual unchanged. Reference choice changes representation, not model fit or the underlying fitted relationships.

T07-A08-V10: Planning tool releveling

Reason before calculating, part (a)

The new reference is the old $G = 1$ group, so its old intercept becomes the new intercept: $b'_0 = 72.00 + (-4.00) = 68.00$. The common slope remains $b'_1 = -2.50$. The contrast reverses direction, so $b'_2 = -(-4.00) = 4.00$.

Work through the calculation, part (b)

For Calendar, $H = 0$, giving $\hat{Y} = 68.00 + (-2.50)X$. For Notebook, $H = 1$, giving $\hat{Y} = 68.00 + (-2.50)X + (4.00) = 72.00 + (-2.50)X$. At the same X , Notebook is fitted 4.00 units higher than Calendar.

Work through the calculation, part (c)

Both codings give:

Group	X	Fit from old coding	Fit from new coding
Notebook	1.0	69.50	69.50
Notebook	6.0	57.00	57.00
Calendar	1.0	65.50	65.50
Calendar	6.0	53.00	53.00

Interpret and check the result, part (d)

Every row has identical fitted values under the two codings. Releveling changes which group is represented by the intercept and reverses the displayed group contrast, but it describes the same two lines. Because each case keeps the same fitted value, subtracting that fit from its observed outcome also leaves every residual unchanged. Reference choice changes representation, not model fit or the underlying fitted relationships.

A09: Interpreting a Group-by-Quantitative-Predictor Interaction

T07-A09-V01: Practice hours by tutorial support

Reason before calculating, part (a)

When $G = 0$, the product XG is zero for every X . When $G = 1$, $XG = X$.

Work through the calculation, part (b)

Substitution gives Self-guided: $\hat{Y} = 40.00 + (2.00)X$, with slope 2.00. For Tutored: $\hat{Y} = 44.00 + (3.20)X$, with slope $b_1 + b_3 = 2.00 + (1.20) = 3.20$.

Work through the calculation, part (c)

The product terms and fitted coordinates are:

Group	G	X	XG	Fitted reasoning score
Self-guided	0	1.0	0.0	42.00
Self-guided	0	5.0	0.0	50.00
Tutored	1	1.0	1.0	47.20
Tutored	1	5.0	5.0	60.00

Work through the calculation, part (d)

Put practice hours on the horizontal axis and fitted reasoning score on the vertical axis. For Self-guided, connect its two table coordinates. For Tutored, connect its two coordinates in a second labeled line. Draw vertical segments between the lines at $X = 1.0$ and $X = 5.0$ and label their lengths 5.20 and 10.00. The nonparallel slopes make the changing gap visible.

Interpret and check the result, part (e)

$b_1 = 2.00$ is the practice hours slope in the reference group. $b_2 = 4.00$ is the fitted Tutored minus Self-guided difference specifically at $X = 0$. It remains interpretable there, although zero may not be substantively central. $b_3 = 1.20$ is the difference between the two group slopes. Consequently, the fitted group gap is $b_2 + b_3X$: it equals 5.20 at $X = 1.0$ and 10.00 at $X = 5.0$. The interaction describes how a conditional association differs by group. It does not establish that group or X causes the outcome.

T07-A09-V02: Practice sessions by archive role

Reason before calculating, part (a)

When $G = 0$, the product XG is zero for every X . When $G = 1$, $XG = X$.

Work through the calculation, part (b)

Substitution gives New staff: $\hat{Y} = 38.00 + (-1.20)X$, with slope -1.20 . For Experienced staff: $\hat{Y} = 35.00 + (-2.00)X$, with slope $b_1 + b_3 = -1.20 + (-0.80) = -2.00$.

Work through the calculation, part (c)

The product terms and fitted coordinates are:

Group	G	X	XG	Fitted retrieval time
New staff	0	0.0	0.0	38.00
New staff	0	4.0	0.0	33.20
Experienced staff	1	0.0	0.0	35.00
Experienced staff	1	4.0	4.0	27.00

Work through the calculation, part (d)

Put practice sessions on the horizontal axis and fitted retrieval time on the vertical axis. For New staff, connect its two table coordinates. For Experienced staff, connect its two coordinates in a second labeled line. Draw vertical

segments between the lines at $X = 0.0$ and $X = 4.0$ and label their lengths -3.00 and -6.20 . The nonparallel slopes make the changing gap visible.

Interpret and check the result, part (e)

$b_1 = -1.20$ is the practice sessions slope in the reference group. $b_2 = -3.00$ is the fitted Experienced staff minus New staff difference specifically at $X = 0$. It remains interpretable there, although zero may not be substantively central. $b_3 = -0.80$ is the difference between the two group slopes. Consequently, the fitted group gap is $b_2 + b_3X$: it equals -3.00 at $X = 0.0$ and -6.20 at $X = 4.0$. The interaction describes how a conditional association differs by group. It does not establish that group or X causes the outcome.

T07-A09-V03: Reading hours by medium

Reason before calculating, part (a)

When $G = 0$, the product XG is zero for every X . When $G = 1$, $XG = X$.

Work through the calculation, part (b)

Substitution gives Print: $\hat{Y} = 49.00 + (2.60)X$, with slope 2.60. For Audio: $\hat{Y} = 51.00 + (1.60)X$, with slope $b_1 + b_3 = 2.60 + (-1.00) = 1.60$.

Work through the calculation, part (c)

The product terms and fitted coordinates are:

Group	G	X	XG	Fitted comprehension score
Print	0	2.0	0.0	54.20
Print	0	6.0	0.0	64.60
Audio	1	2.0	2.0	54.20
Audio	1	6.0	6.0	60.60

Work through the calculation, part (d)

Put reading hours on the horizontal axis and fitted comprehension score on the vertical axis. For Print, connect its two table coordinates. For Audio, connect its two coordinates in a second labeled line. Draw vertical segments between the lines at $X = 2.0$ and $X = 6.0$ and label their lengths 0.00 and -4.00 . The nonparallel slopes make the changing gap visible.

Interpret and check the result, part (e)

$b_1 = 2.60$ is the reading hours slope in the reference group. $b_2 = 2.00$ is the fitted Audio minus Print difference specifically at $X = 0$. It remains interpretable there, although zero may not be substantively central. $b_3 = -1.00$ is the difference between the two group slopes. Consequently, the fitted group gap is $b_2 + b_3X$: it equals 0.00 at $X = 2.0$ and -4.00 at $X = 6.0$. The interaction describes how a conditional association differs by group. It does not establish that group or X causes the outcome.

T07-A09-V04: Rehearsal by navigation display

Reason before calculating, part (a)

When $G = 0$, the product XG is zero for every X . When $G = 1$, $XG = X$.

Work through the calculation, part (b)

Substitution gives Static: $\hat{Y} = 48.00 + (-1.50)X$, with slope -1.50 . For Interactive: $\hat{Y} = 46.00 + (-2.40)X$, with slope $b_1 + b_3 = -1.50 + (-0.90) = -2.40$.

Work through the calculation, part (c)

The product terms and fitted coordinates are:

Group	G	X	XG	Fitted navigation time
Static	0	1.0	0.0	46.50
Static	0	5.0	0.0	40.50
Interactive	1	1.0	1.0	43.60
Interactive	1	5.0	5.0	34.00

Work through the calculation, part (d)

Put rehearsal attempts on the horizontal axis and fitted navigation time on the vertical axis. For Static, connect its two table coordinates. For Interactive, connect its two coordinates in a second labeled line. Draw vertical segments between the lines at $X = 1.0$ and $X = 5.0$ and label their lengths -2.90 and -6.50 . The nonparallel slopes make the changing gap visible.

Interpret and check the result, part (e)

$b_1 = -1.50$ is the rehearsal attempts slope in the reference group. $b_2 = -2.00$ is the fitted Interactive minus Static difference specifically at $X = 0$. It remains interpretable there, although zero may not be substantively central. $b_3 = -0.90$ is the difference between the two group slopes. Consequently, the fitted group gap is $b_2 + b_3X$: it equals -2.90 at $X = 1.0$ and -6.50 at $X = 5.0$. The interaction describes how a conditional association differs by group. It does not establish that group or X causes the outcome.

T07-A09-V05: Practice sets by catalog aid

Reason before calculating, part (a)

When $G = 0$, the product XG is zero for every X . When $G = 1$, $XG = X$.

Work through the calculation, part (b)

Substitution gives Index: $\hat{Y} = 52.00 + (2.00)X$, with slope 2.00. For Search bar: $\hat{Y} = 55.00 + (2.70)X$, with slope $b_1 + b_3 = 2.00 + (0.70) = 2.70$.

Work through the calculation, part (c)

The product terms and fitted coordinates are:

Group	G	X	XG	Fitted accuracy score
Index	0	0.0	0.0	52.00
Index	0	4.0	0.0	60.00
Search bar	1	0.0	0.0	55.00
Search bar	1	4.0	4.0	65.80

Work through the calculation, part (d)

Put practice sets on the horizontal axis and fitted accuracy score on the vertical axis. For Index, connect its two table coordinates. For Search bar, connect its two coordinates in a second labeled line. Draw vertical segments between the lines at $X = 0.0$ and $X = 4.0$ and label their lengths 3.00 and 5.80. The nonparallel slopes make the changing gap visible.

Interpret and check the result, part (e)

$b_1 = 2.00$ is the practice sets slope in the reference group. $b_2 = 3.00$ is the fitted Search bar minus Index difference specifically at $X = 0$. It remains interpretable there, although zero may not be substantively central. $b_3 = 0.70$ is the difference between the two group slopes. Consequently, the fitted group gap is $b_2 + b_3X$: it

equals 3.00 at $X = 0.0$ and 5.80 at $X = 4.0$. The interaction describes how a conditional association differs by group. It does not establish that group or X causes the outcome.

T07-A09-V06: Sessions by workshop setting

Reason before calculating, part (a)

When $G = 0$, the product XG is zero for every X . When $G = 1$, $XG = X$.

Work through the calculation, part (b)

Substitution gives Remote: $\hat{Y} = 36.00 + (2.40)X$, with slope 2.40. For Classroom: $\hat{Y} = 41.00 + (3.20)X$, with slope $b_1 + b_3 = 2.40 + (0.80) = 3.20$.

Work through the calculation, part (c)

The product terms and fitted coordinates are:

Group	G	X	XG	Fitted confidence score
Remote	0	1.0	0.0	38.40
Remote	0	5.0	0.0	48.00
Classroom	1	1.0	1.0	44.20
Classroom	1	5.0	5.0	57.00

Work through the calculation, part (d)

Put sessions on the horizontal axis and fitted confidence score on the vertical axis. For Remote, connect its two table coordinates. For Classroom, connect its two coordinates in a second labeled line. Draw vertical segments between the lines at $X = 1.0$ and $X = 5.0$ and label their lengths 5.80 and 9.00. The nonparallel slopes make the changing gap visible.

Interpret and check the result, part (e)

$b_1 = 2.40$ is the sessions slope in the reference group. $b_2 = 5.00$ is the fitted Classroom minus Remote difference specifically at $X = 0$. It remains interpretable there, although zero may not be substantively central. $b_3 = 0.80$ is the difference between the two group slopes. Consequently, the fitted group gap is $b_2 + b_3X$: it equals 5.80 at $X = 1.0$ and 9.00 at $X = 5.0$. The interaction describes how a conditional association differs by group. It does not establish that group or X causes the outcome.

T07-A09-V07: Focus blocks by room type

Reason before calculating, part (a)

When $G = 0$, the product XG is zero for every X . When $G = 1$, $XG = X$.

Work through the calculation, part (b)

Substitution gives Open room: $\hat{Y} = 58.00 + (2.10)X$, with slope 2.10. For Private room: $\hat{Y} = 62.00 + (1.50)X$, with slope $b_1 + b_3 = 2.10 + (-0.60) = 1.50$.

Work through the calculation, part (c)

The product terms and fitted coordinates are:

Group	G	X	XG	Fitted task-accuracy score
Open room	0	2.0	0.0	62.20
Open room	0	7.0	0.0	72.70
Private room	1	2.0	2.0	65.00

Group	G	X	XG	Fitted task-accuracy score
Private room	1	7.0	7.0	72.50

Work through the calculation, part (d)

Put focus blocks on the horizontal axis and fitted task-accuracy score on the vertical axis. For Open room, connect its two table coordinates. For Private room, connect its two coordinates in a second labeled line. Draw vertical segments between the lines at $X = 2.0$ and $X = 7.0$ and label their lengths 2.80 and -0.20 . The nonparallel slopes make the changing gap visible.

Interpret and check the result, part (e)

$b_1 = 2.10$ is the focus blocks slope in the reference group. $b_2 = 4.00$ is the fitted Private room minus Open room difference specifically at $X = 0$. It remains interpretable there, although zero may not be substantively central. $b_3 = -0.60$ is the difference between the two group slopes. Consequently, the fitted group gap is $b_2 + b_3X$: it equals 2.80 at $X = 2.0$ and -0.20 at $X = 7.0$. The interaction describes how a conditional association differs by group. It does not establish that group or X causes the outcome.

T07-A09-V08: Visits by museum route

Reason before calculating, part (a)

When $G = 0$, the product XG is zero for every X . When $G = 1$, $XG = X$.

Work through the calculation, part (b)

Substitution gives Free route: $\hat{Y} = 44.00 + (3.50)X$, with slope 3.50. For Curated route: $\hat{Y} = 47.00 + (5.00)X$, with slope $b_1 + b_3 = 3.50 + (1.50) = 5.00$.

Work through the calculation, part (c)

The product terms and fitted coordinates are:

Group	G	X	XG	Fitted knowledge score
Free route	0	0.0	0.0	44.00
Free route	0	3.0	0.0	54.50
Curated route	1	0.0	0.0	47.00
Curated route	1	3.0	3.0	62.00

Work through the calculation, part (d)

Put visits on the horizontal axis and fitted knowledge score on the vertical axis. For Free route, connect its two table coordinates. For Curated route, connect its two coordinates in a second labeled line. Draw vertical segments between the lines at $X = 0.0$ and $X = 3.0$ and label their lengths 3.00 and 7.50. The nonparallel slopes make the changing gap visible.

Interpret and check the result, part (e)

$b_1 = 3.50$ is the visits slope in the reference group. $b_2 = 3.00$ is the fitted Curated route minus Free route difference specifically at $X = 0$. It remains interpretable there, although zero may not be substantively central. $b_3 = 1.50$ is the difference between the two group slopes. Consequently, the fitted group gap is $b_2 + b_3X$: it equals 3.00 at $X = 0.0$ and 7.50 at $X = 3.0$. The interaction describes how a conditional association differs by group. It does not establish that group or X causes the outcome.

T07-A09-V09: Feedback rounds by meeting mode

Reason before calculating, part (a)

When $G = 0$, the product XG is zero for every X . When $G = 1$, $XG = X$.

Work through the calculation, part (b)

Substitution gives Asynchronous: $\hat{Y} = 50.00 + (2.80)X$, with slope 2.80. For Live: $\hat{Y} = 54.00 + (2.30)X$, with slope $b_1 + b_3 = 2.80 + (-0.50) = 2.30$.

Work through the calculation, part (c)

The product terms and fitted coordinates are:

Group	G	X	XG	Fitted revision score
Asynchronous	0	1.0	0.0	52.80
Asynchronous	0	5.0	0.0	64.00
Live	1	1.0	1.0	56.30
Live	1	5.0	5.0	65.50

Work through the calculation, part (d)

Put feedback rounds on the horizontal axis and fitted revision score on the vertical axis. For Asynchronous, connect its two table coordinates. For Live, connect its two coordinates in a second labeled line. Draw vertical segments between the lines at $X = 1.0$ and $X = 5.0$ and label their lengths 3.50 and 1.50. The nonparallel slopes make the changing gap visible.

Interpret and check the result, part (e)

$b_1 = 2.80$ is the feedback rounds slope in the reference group. $b_2 = 4.00$ is the fitted Live minus Asynchronous difference specifically at $X = 0$. It remains interpretable there, although zero may not be substantively central. $b_3 = -0.50$ is the difference between the two group slopes. Consequently, the fitted group gap is $b_2 + b_3X$: it equals 3.50 at $X = 1.0$ and 1.50 at $X = 5.0$. The interaction describes how a conditional association differs by group. It does not establish that group or X causes the outcome.

T07-A09-V10: Planning by tool type

Reason before calculating, part (a)

When $G = 0$, the product XG is zero for every X . When $G = 1$, $XG = X$.

Work through the calculation, part (b)

Substitution gives Notebook: $\hat{Y} = 74.00 + (-1.80)X$, with slope -1.80. For Calendar: $\hat{Y} = 72.00 + (-2.70)X$, with slope $b_1 + b_3 = -1.80 + (-0.90) = -2.70$.

Work through the calculation, part (c)

The product terms and fitted coordinates are:

Group	G	X	XG	Fitted completion time
Notebook	0	1.0	0.0	72.20
Notebook	0	6.0	0.0	63.20
Calendar	1	1.0	1.0	69.30
Calendar	1	6.0	6.0	55.80

Work through the calculation, part (d)

Put planning sessions on the horizontal axis and fitted completion time on the vertical axis. For Notebook, connect its two table coordinates. For Calendar, connect its two coordinates in a second labeled line. Draw

vertical segments between the lines at $X = 1.0$ and $X = 6.0$ and label their lengths -2.90 and -7.40 . The nonparallel slopes make the changing gap visible.

Interpret and check the result, part (e)

$b_1 = -1.80$ is the planning sessions slope in the reference group. $b_2 = -2.00$ is the fitted Calendar minus Notebook difference specifically at $X = 0$. It remains interpretable there, although zero may not be substantively central. $b_3 = -0.90$ is the difference between the two group slopes. Consequently, the fitted group gap is $b_2 + b_3X$: it equals -2.90 at $X = 1.0$ and -7.40 at $X = 6.0$. The interaction describes how a conditional association differs by group. It does not establish that group or X causes the outcome.