



JMP035: Durability of Mobile Phone Screen – Part 4

Statistical Modeling: Multivariable Logistic
Regression

Produced by

Kevin Potcner, JMP Global Academic Team
kevin.potcner@jmp.com

Durability of Mobile Phone Screen – Part 4

Statistical Modeling: Multivariable Logistic Regression

Key ideas:

This case study requires building a multivariable logistic regression model to evaluate and compare the durability of different mobile phone screen types in a drop test across various heights.

This case study is an extension to three case studies: JMP032: Durability of Mobile Phone Screen – Part 1, and JMP033: Durability of Mobile Phone Screen – Part 2 and JMP034: Durability of Mobile Phone Screen – Part 3. While it's not necessary to have completed those case studies first, it is recommended that you do so for a more comprehensive description of each stage of the study and subsequent analyses..

Background

The durability of a product is clearly an important quality characteristic for both the end user and the manufacturer. For end users, durability is especially important for mobile phones. Dropping a phone on a hard surface, for example, can result in the screen cracking or even breaking, rendering the phone unusable. To evaluate the durability of these screens, manufacturers subject a sample of screens to a variety of tests to simulate typical wear and tear by a user, such as dropping the phone onto a concrete surface.

In JMP032: Durability of Mobile Phone Screen – Part 1,, material scientists for a screen manufacturer experimented with two new formulations of an aluminosilicate glass (A and B). These two formulations were produced by making a change to a final processing step that uses a specific level of potassium nitrate to strengthen the glass.

A sample of 10 screens of each type was developed for testing. Each screen was installed into the same style of phone. The phones were then dropped in a controlled identical manner from a height of 1 meter onto a concrete surface. A binary variable “Success” (no damage) and “Fail” (screen damage) was recorded.

One of the company's goals was that 97% of the screens manufactured would be able to experience a drop of 1 meter without becoming damaged (i.e., the Population Success Rate).

The analyses that were illustrated in JMP032: Durability of Mobile Phone Screen – Part 1, which were based on 10 phones tested for each Screen Type, failed to generate the statistical evidence needed to demonstrate the desired success rate. The analyses also failed to find any statistically significant difference between the two Screen Types.

The engineers decided to test an additional 40 phones for each of the two Screen Types. Analyzing the results from this expanded test was the objective of the exercises for that case study.

In JMP033: Durability of Mobile Phone Screen – Part 2, the material scientists developed a third Screen Type (C), which required a more expensive process. Tests were done for this Screen Type at 1.0m, as had been done for Screen Types A and B. In an effort to gain a comprehensive understanding of the durability of these screens, the engineers also conducted the drop test at two additional heights (0.5 and 1.5 meters), which resulted in the following data:

	0.5m	1.0m	1.5m
Type A	n=40 S=36 F=4	n=50 S=41 F=9	n=50 S=28 F=22
Type B	n=40 S=40 F=0	n=50 S=48 F=2	n=50 S=40 F=10
Type C	n=40 S=39 F=1	n=50 S=47 F=3	n=50 S=39 F=11

These data were analyzed in case JMP033: Durability of Mobile Phone Screen – Part 2.

To continue developing a deeper understanding of the performance of the Screen Types, the engineers performed additional tests at Drop Heights of 0.25m, 1.25m, 2.0m, 2.5m and 3.0m. In addition, a sample of screens from a leading competitor were obtained and included in the testing, resulting in the following data:

	0.25m	0.5m	1.0m	1.25m	1.5m	2.0m	2.5m	3.0m
Type A	n=25 S=25 F=0	n=40 S=36 F=4	n=50 S=41 F=9	n=50 S=35 F=15	n=50 S=28 F=22	n=40 S=4 F=36	n=40 S=0 F=40	n=25 S=0 F=25
Type B	n=25 S=25 F=0	n=40 S=40 F=0	n=50 S=48 F=2	n=50 S=45 F=5	n=50 S=40 F=10	n=40 S=18 F=22	n=40 S=4 F=36	n=25 S=0 F=25
Type C	n=25 S=25 F=0	n=40 S=39 F=1	n=50 S=47 F=3	n=50 S=45 F=5	n=50 S=39 F=11	n=40 S=16 F=14	n=40 S=3 F=38	n=25 S=0 F=25
Competitor	n=10 S=10 F=0	n=20 S=19 F=1	n=20 S=18 F=2	n=20 S=16 F=4	n=20 S=13 F=7	n=20 S=5 F=15	n=20 S=1 F=19	n=10 S=0 F=10

In case JMP034: Durability of Mobile Phone Screen – Part 3, these data were used to build single-variable logistic regressions models, one for each Screen Type. The logistic regression models provided us with a tool to estimate the Population Success Rate across any Drop Height between 0.25 and 3.0 meters, including those not part of the study.

The logistic regression model also provided us with a tool to perform an inverse prediction, specifically estimating the Drop Height that would result in a given Success Rate.

In this case study, we'll build one multivariable logistic regression model created from all of the data together for a single equation to describe the performance of all Screen Types. In addition to having a

single equation that can estimate Success Rates and Drop Heights for any Screen Type, this will also provide us with a tool for conducting formal statistical tests comparing the Success Rates between Screen Types.

We'll finish by comparing the estimates made by this multivariable model to the ones made using separate single-variable models to determine which approach is best to use.

The Task

The primary objectives of this analysis are to:

1. Describe how the performance of each Screen Type changes as the Drop Height changes.
2. Describe any statistical differences in durability between the three Screen Types and determine if that difference is consistent across all eight Drop Heights.
3. Determine how the company's three Screen Types compare to the competitor's.
4. Develop a tool that will estimate the Success Rate for any Screen Type at any Drop Height including those not part of the testing.
5. Determine the estimated Success Rate for each Screen Type at 1.0m Drop Height.
6. Determine the Drop Height for which the Success Rate is estimated to be at least 97% for each Screen Type.
7. Estimate the minimum Drop Height at which the screen is more likely to be damaged versus not damaged.
8. Estimate the greatest Drop Height at which 5% of screen would be damaged versus not damaged?

The Data Drop Test 4.jmp

The data is stored in what's referred to as Outcome/Frequency Table format.

Screen Height	Four screen types (A, B, C, and Comp) Height from which screens were dropped (0.25m, 0.5m, 1.0m, 1.25m, 1.5m, 2.0m, 2.5m, and 3.0m)
Outcome Tested	Two outcomes (Success, Fail)
Count	Number of screens tested
Rate	Number of specimens that resulted in either Success or Fail The percentage of specimens that succeeded (didn't experience damage) or failed (did experience damage) in the test

An additional data set (Drop Test 4_Example) will be used to illustrate the process of building a multivariable logistic regression model to help you build such a model for the actual test results.

JMP Tips

To perform an analysis on a subset of the data in the table, which will be needed throughout these analyses, you'll need to filter the data. It can be done either through the Global Data Filter at the data table level before running an analysis or with the Local Data Filter within the output window from an analysis.

For Global Data Filter, select Rows>Data Filter. To choose a variable to filter, select the variable and click the +. Choose Add to filter by additional variables. Select the values of the variables you wish to include in your analysis and then select the Show and Include boxes. Perform the analysis.

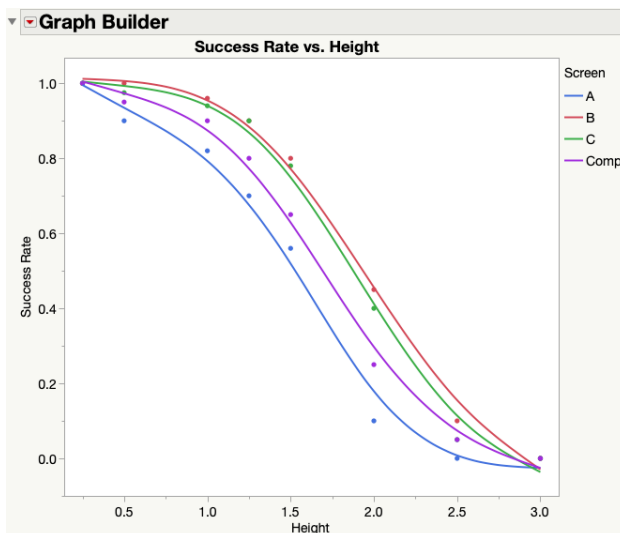
For Local Data Filter, first run the analysis. Then select Local Data Filter under the red triangle at the top of the output. To choose a variable to filter, select the variable and click the +. Choose Add to filter by additional variables. Select the values of the variables by which you wish to subset the data. The numerical and graphical results will change to correspond to the data being included in the analysis. The tools and analysis options available may change based on attributes of the data being included.

Analysis

Graphical

We begin by summarizing the data graphically. Exhibit 1 displays a scatter plot showing points for each of the success rates for the four Screen Types and eight Drop Heights. In addition, fitted curves are added to visualize how the success rate changes across the Drop Heights.

Exhibit 1 Scatterplot with Smoother



To create this graph from the data table, the data will need to be filtered so that it only uses the rows corresponding to the success outcome. To do so, use the Global Data Filter. Choose Rows>Data Filter. Filter the data by the variable Outcome. See JMP Tips on page 4 for instructions.

To create, Graph>Graph Builder. Select Success Rate as the Y Variable, Height as the Freq Variable, and Screen as the Overlay Variable.

This graph provides a nice visualization of the change in success rate across the full range of Drop Heights, as well as potential differences between the Screen Types. For example, it appears that the success rates for Screen Types B and C are very similar, having the best performance. Screen Type A appears to perform the worst, with the success rate for the competitor's screen being somewhere in between. The exercises will have you perform statistical tests to decide if the performance between the Screen Types is considered statistically different.

Building a multiple variable logistic regression model

It is often best to build one statistical model that incorporates all of the data instead of separate individual models for each of the four Screen Types. Each separate model is built using only the data from that Screen Type. A multivariable logistic regression model will be built from all of the data. The structure of a multivariable logistic regression model for our study is:

$$\text{Prob(Event)} = \frac{1}{1 + e^{-(\beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_{1,2} X_1 X_2)}}$$

where X_1 is a categorical variable representing the four Screen Types, X_2 is the Drop Height, and $X_1 X_2$ is the interaction between the two.

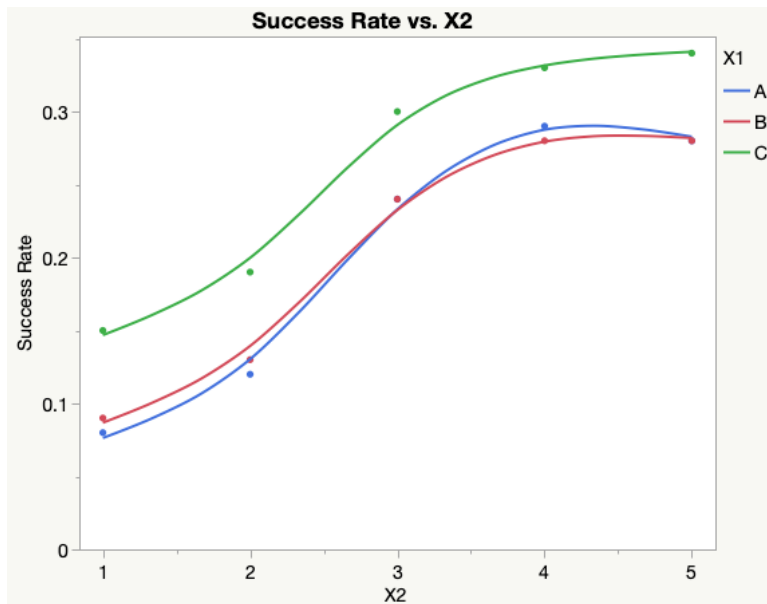
Technical Note: The model above is written in a simplified version. The terms involving the categorical variable X_1 actually require a set of variables and parameters to represent the number of categories. For simplicity's sake, that form is not shown.

In a statistical model building process, we seek to determine which of the X 's are useful in describing changes in the outcome variable and should be included in the model and which ones are not and should not be included.

In the exercises, you will be asked to build the multivariable logistic regression model for the actual results of the drop test. To illustrate that model-building process, a different fabricated data set contained in the file Drop Test 4_Example will be used. This data is of a similar structure to the drop test data. There are two variables: X_1 , a categorical variable with four different categories (A, B, C, and D); and X_2 , a continuous variable with numerical values (1, 2, 3, 4, and 5).

We begin with a plot of the success rate, as shown in Exhibit 2. It's clear from this graph that the success rate increases as X_2 increases and that category A and B have very similar results, with C having a higher success rate across the full range of X_2 .

Exhibit 2 Scatter Plot with Smoother



To create this graph from the data table, the data will need to be filtered so that it only uses the rows corresponding to the success outcome, which is best done with the Global Data Filter. Choose Rows>Data Filter. Filter the data by the variable Outcome. See JMP Tips on page 4 for instructions.

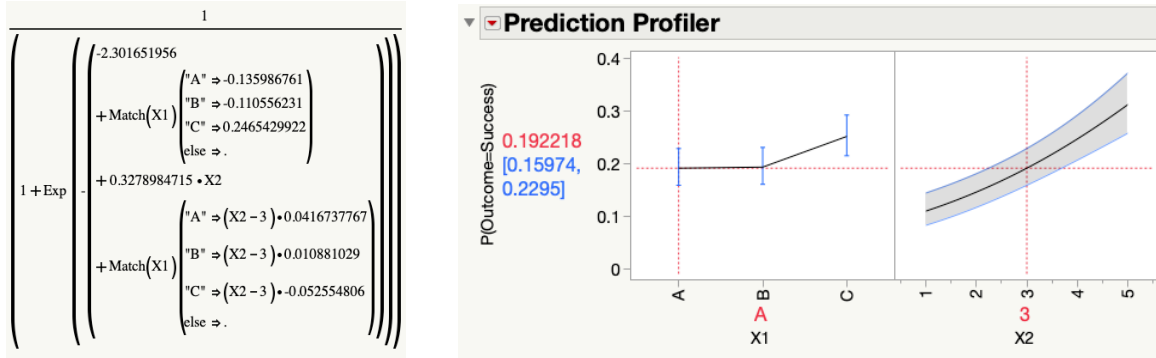
To create, Graph>Graph Builder. Select Success Rate as the Y Variable, X2 as the Freq Variable, and X1 as the Overlay Variable.

To build the multivariable logistic regression model, we will begin by fitting the full model of the form:

$$\text{Prob}(\text{Event}) = \frac{1}{1 + e^{-(\beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_{1,2} X_1 X_2)}}$$

Exhibit 3 displays both the numerical formula and a geometric visualization for the fitted model.

Exhibit 3 Fitted Logistic Regression Model



For this analysis, the full data in Outcome/Frequency format needs to be used. If data filter from Graph Builder is on, return to including all the data in the data table by selecting Rows>Clear Row State.

To create, Analyze>Fit Model. Select Outcome as the Y Variable. Add X1 and X2 into the Model Effects field. Add the interaction effect by selecting X1 and X2 in the variable list and selecting Cross. Choose Generalized Linear Model under Personality, Binomial as the Distribution, and Logit as the Link Function.

Select Profiler under the top red triangle. Select Save Columns>Prediction Formula under the top red triangle. Select the column in which the prediction formula was saved. Right-click and select Formula to display the equation for the fitted model.

To evaluate if this is a good model for describing the data, we will examine important elements of the output. Exhibit 4 displays some of the logistic regression output.

The hypothesis for the whole model test can be written as:

$$H_0: \beta_1 = 0 ; \beta_2 = 0 ; \beta_{1,2} = 0$$

$$H_A: \beta_1 , \beta_2 \text{ and/or } \beta_{1,2} \neq 0$$

The p-value for the whole model test is <0.0001, which is a highly significant result. It is as expected since the graph shown in Exhibit 1 made it clear that the success rate significantly changes across X₂ and that category C for X₁ looks to be statistically different than A and B.

Exhibit 4 Statistical Tests for Logistic Regression Model

▼ Whole Model Test				
Model	-LogLikelihood	ChiSquare	DF	Prob>ChiSq
Difference	29.4789	58.9578	5	<.0001*
Full	765.914717			
Reduced	795.393617			
Goodness Of Fit				
Statistic	ChiSquare	DF	Prob>ChiSq	
Pearson	1486.578	1494	0.5493	
Deviance	1531.829	1494	0.2424	
AICc				
1543.8857				
▼ Effect Summary				
Source	LogWorth			PValue
X2	12.319			0.00000
X1	1.627			0.02360
X1*X2	0.163			0.68700
Remove Add Edit ☐ FDR				

For this analysis, the full data in Outcome/Frequency format needs to be used. If data filter from Graph Builder is on, return to including all the data the in the data table by selecting Rows>Clear Row State.

To create, Analyze>Fit Model. Select Outcome as the Y Variable. Add X1 and X2 into the Model Effects field. Add the interaction effect by selecting X1and X2 in the variable list and selecting Cross. Choose Generalized Linear Model under Personality, Binomial as the Distribution, and Logit as the Link Function.

The effect summary table and graph show individual tests. The test that should be examined first is the one with the most complicated model term. In this instance, it would be the interaction term X_1X_2 .

$$H_0: \beta_{1,2} = 0$$

$$H_A: \beta_{1,2} \neq 0$$

The p-value for this test is 0.687, which is an insignificant result. It indicates that the interaction term is not helpful in describing the data and should not be included in the model. The interpretation of not needing an interaction term in this model is that the change in the success rate across X_2 is not statistically different for the four categories of X_1 . That is, the same estimate of the coefficient for X_2 can be used for all categories of X_1 .

The next step is to examine the tests for the individual terms X_1 and X_2 . Exhibit 5 shows the effect summary table and graph with the interaction term removed. Notice that both X_1 and X_2 are statistically significant, with X_1 (the continuous variable) being the most significant.

Exhibit 5 Generalized Linear Model

Effect Summary		
Source	LogWorth	PValue
X2	12.171	0.00000
X1	1.482	0.03294
Remove Add Edit Undo ☐ FDR		

To create, select the term $X1 \times X2$ in the Effect Summary Table, and select Remove.

The hypothesis for the categorical variable X_1 can be written as:

$$H_0: \beta_1 = 0$$

$$H_A: \beta_1 \neq 0$$

The significant result (p-value of 0.03294) indicates that there is a difference in the success rate between the categories of X_1 . The results of the test, however, do not inform us which ones are different. The graph from Exhibit 1 does give us an indication of where those differences might be. To more formally reach a conclusion about potential differences, how to use the model to perform statistical tests via contrasts is illustrated next.

Contrasts

Exhibit 6 shows the results of all possible contrast comparisons (A vs. B, A vs. C, and B vs. C).

The p-value for A vs. B is 0.9363, which is a highly insignificant result that indicates there is no evidence suggesting the success rates for category A and B are different. The p-value for A vs. C is 0.0220 and B vs. C is 0.0270. Both are statistically significant, indicating that category C is different than A and B.

Exhibit 6 Contrast Comparisons

Contrast		Contrast		Contrast	
Test Detail		Test Detail		Test Detail	
Level		Level		Level	
X1[A]	1.0000	X1[A]	1.0000	X1[A]	0.0000
X1[B]	-1.0000	X1[B]	0.0000	X1[B]	1.0000
X1[C]	0.0000	X1[C]	-1.0000	X1[C]	-1.0000
Value	-0.0128	Value	-0.3506	Value	-0.3378
Std Error	0.1598	Std Error	0.1536	Std Error	0.1533
ChiSquare	0.0064	ChiSquare	5.2470	ChiSquare	4.8881
Prob>ChiSq	0.9363	Prob>ChiSq	0.0220	Prob>ChiSq	0.0270
-LogLikelihood	766.2933	-LogLikelihood	768.9136	-LogLikelihood	768.7342
-LogLikelihood	766.2933	-LogLikelihood	768.9136	-LogLikelihood	768.7342
DF	1.0000	DF	1.0000	DF	1.0000
L-R ChiSquare	0.0064	L-R ChiSquare	5.2470	L-R ChiSquare	4.8881
Prob>ChiSq	0.9363	Prob>ChiSq	0.0220	Prob>ChiSq	0.0270

To create, select Contrast under the red triangle at the top-left of the output. Choose X_1 as the variable. Select 1 for A and -1 for B. Repeat for the other two comparisons (A vs. C and B vs. C).

The Prediction Profiler shown in Exhibit 3 provides a visualization of the model and the significant difference in the success rate between category C and both A and B, as well as the insignificant difference in success rates between A and B.

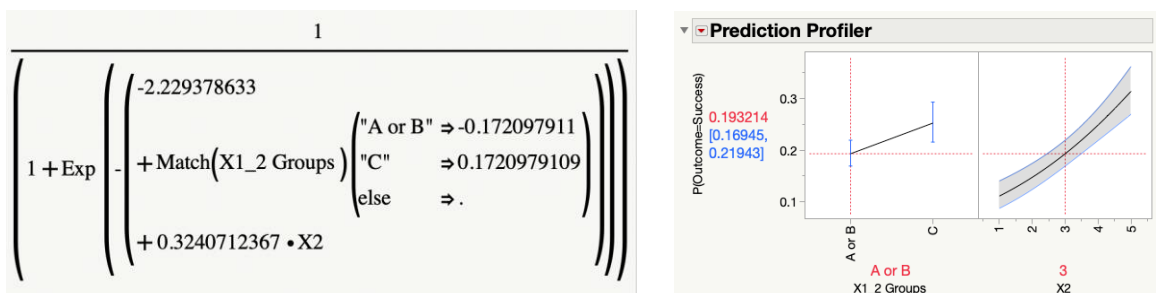
A common approach for building a statistical model is to reduce the model even further so as not to include components that are not significant. It can be done by recoding the variable X_1 so that it only has two values ("A or B" and "C"). Exercise caution when using this method so as not to remove features in the data that are worth retaining in the model. Many practitioners will use a lower confidence level to determine which terms to keep in a model. For example, a p-value of 0.14, though not significant at the 95% confidence level, is significant at the 85% confidence level. In such cases, it may be best to keep that term in a final model. The p-value for A vs. B in our data is 0.9363, thus highly insignificant. It would be very reasonable to recode the data to only two categories ("A or B" and "C") and refit the model.

Exhibit 7 displays the equation and the Prediction Profiler for a model of the form:

$$\text{Prob(Event)} = \frac{1}{1 + e^{-(\beta_0 + \beta_1 X_1 + \beta_2 X_2)}}$$

where X_1 only has 2 Groups ("A or B" and "C")

Exhibit 7 Equation and Prediction Profiler



To create, recode the variable X_1 by highlighting the column X_1 and choosing Cols>Recode. Recode the value for Screen Types A and B to "A or B" and leave C as is. Name this new variable " X_1_2 Groups".

Analyze>Fit Model. Select Outcome as the Y Variable. Add X_1_2 Groups and X_2 into the Model Effects field. Choose Generalized Linear Model under Personality, Binomial as the Distribution, and Logit as the Link Function.

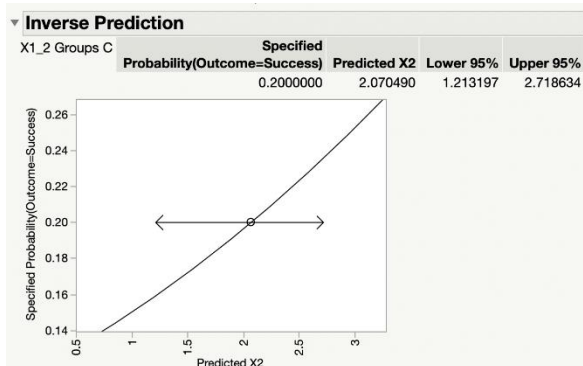
Select Profiler under the top red triangle. Select Save Columns>Prediction Formula under the top red triangle. Select the column in which the prediction formula was saved. Right-click and select Formula to display the equation for the fitted model.

Using final model to make predictions

The Prediction Profiler can be used to make predictions. The estimated success rate for X_1 ="A or B" and X_2 =3 is shown to be 0.193, with a confidence interval of [0.169, 0.219]. The estimated success rate for category C at the same value for X_2 (not shown) is 0.253, with a 95% confidence interval of [0.216, 0.293].

Another useful type of prediction is to estimate the value for the continuous variable X_2 that would achieve a given success rate, called an inverse prediction. To illustrate, we choose a success rate of 0.20. Exhibit 8 shows the results of making an inverse prediction to achieve a 0.20 success rate when X_1 =C. Here we see that a success rate of 0.20 is estimated to occur when X_2 =2.07, with a 95% confidence interval of [1.21, 2.72].

Exhibit 8 Inverse Prediction



To create, choose Inverse Prediction under the red triangle at the upper-left of the output. Type in 0.20 for the Probability(Outcome=Success) field and choose Group C for X_1 . Note: the above graph was zoomed in to the area of interest.

Odds ratio

Another useful way to numerically summarize the results of a logistic regression model is through the odds ratio.

The odds for an outcome are the probability of the outcome occurring divided by the probability of it not occurring; it is written mathematically as:

$$\frac{p_A}{(1 - p_A)}$$

In this case, p_A would be the probability of success for Scenario A and $(1 - p_A)$ is the probability a failure. For example, if the probability of success is 0.90, then the probability of failure is 0.10 and the odds would be $0.90/0.10 = 9$. This is interpreted as the probability of a success being 9 times more likely than the probability of a failure for Scenario A.

The odds ratio is the ratio of two different odds. If, for example, the probability of success for Scenario B is 0.80, the odds of success for Scenario B is $0.80/0.20 = 4$. The odds ratio between the two scenarios would be $9/4 = 2.25$. This is interpreted as the odds of success for Scenario A being 2.25 the size of the odds of success for Scenario B.

An odds ratio equal to 1 would occur when the odds of success for the two scenarios are the same. In other words, when the success rates are the same.

Exhibit 9 shows the odds ratio for this model.

Exhibit 9 Odds Ratios

Odds Ratios

For Outcome odds of Success versus Fail

Unit Odds Ratios

Per unit change in regressor

Term	Odds Ratio	Lower 95%	Upper 95%	Reciprocal
X2	1.382746	1.262899	1.513966	0.7231987

Range Odds Ratios

Per change in regressor over entire range

Term	Odds Ratio	Lower 95%	Upper 95%	Reciprocal
X2	3.65569	2.543752	5.253685	0.2735461

Odds Ratios for X1_2 Groups

Level1	/Level2	Odds Ratio	Prob>Chisq	Lower 95%	Upper 95%
C	A or B	1.4108549	0.0086*	1.0914335	1.8237588
A or B	C	0.7087901	0.0086*	0.5483181	0.9162262

The Odds Ratio are available under the Nominal Logistic Regression Personality in the Fit Model Platform.

To create, Analyze>Fit Model. Select Outcome as the Y Variable. Add X1_2 Groups and X2 into the Model Effects field. Choose Nominal Logistic under Personality. Choose Odds Ratio under the red triangle in the upper-left of the output.

For the continuous variable X_2 , the unit odds ratio is 1.38, with a 95% confidence interval of [1.26, 1.51]. Thus, we are 95% confident that the odds of success increase by a factor of 1.26 to 1.51 for each value increase in X_2 .

To illustrate, by using the Prediction Profiler we can find that the estimated success rate at $X_1 = C$ and $X_2 = 2$ is 0.1964. Thus, the odds of success are $0.1964/(1 - 0.1964) = 0.2444$. The odds of success at $X_1 = C$ and $X_2 = 3$ are $0.2525/(1 - 0.2525) = 0.3378$. The odds ratio describing a one unit change in X_2 is therefore $0.3378/0.2444 = 1.38$ as shown. It would be the same value in comparing the odds between any one unit change in X_2 (e.g., calculating the odds ratio for $X_1 = "A \text{ or } B"$ and $X_2 = 4$ and $X_1 = "A \text{ or } B"$ and $X_2 = 5$ would also be 1.38).

The range odds ratio quantifies the change in odds from the lowest level to the highest level of X_2 . Here we see the range odds ratio is 3.66. Thus we would estimate that the odds of success when $X_2=5$ are 3.66 times larger than the odds of success when $X_2=1$.

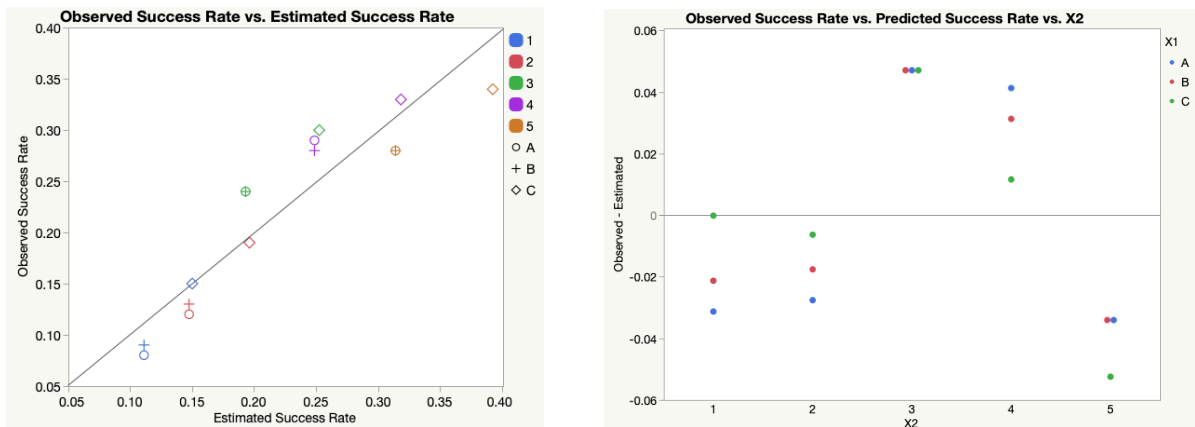
Also provided is the odds ratios for changes in the categorical variable X_1 . Here we see that the odds of success when $X_1 = C$ are estimated to be 1.41 times the size of when $X_1 = "A \text{ or } B"$. This difference in the odds is the same regardless of the value of X_2 .

Confidence intervals are also provided, quantifying the amount of uncertainty in these estimates.

Evaluating model performance

Exhibit 10 illustrates two common ways to visualize prediction error in a logistic regression model. These graphs show how large or small the prediction error is. Points close to the reference lines represent when the estimated success rate from the logistic regression model is close to the observed success rate from the actual data. Points above the line represent when the observed success rate is larger than the estimated success rate; points below the line represent when the estimated is larger than the observed. For example, the model predicts the success rate to be lower than the observed success rate when $X_2 = 3$ or 4 for all categories of X_1 , with $X_2 = 3$ having the largest error of the two. The model predicts the success rate to be higher than the observed success rate for most of the other cases, with $X_2 = 5$ having the largest error. For one case ($X_1 = C$ and $X_2 = 1$), the prediction error is essentially 0. These differences between observed and estimated success rates reflect the error that exists in all statistical models.

Exhibit 10 Visualizing Prediction Error



To create the graph on the left, change the modeling type for variable X2 to Ordinal. Right-click on the column X2. Choose Column Info. Change Modeling Type to Ordinal. Use Graph>Graph Builder to plot the Observed Success Rate vs. Estimated Success Rate. Place X1 in the Overlay role and X2 in the Color role. A line was added manually via Tools>Line to show when Observed = Predicted.

For the graph on right, save the estimated success rates by selecting Save Columns>Predicted Values under the red triangle in the upper-left of the output. Create a new column for the difference between the observed success rate and the estimated success rate by selecting Cols>New Column. Choose Column Properties>Formula. Create the formula Success Rate – Pred Outcome. Use Graph>Graph Builder to plot the Observed – Estimated values vs. the values of X2. Place X1 in the Color role.

The percent difference between observed vs. estimated could also be used to how large the difference is relative to the observed value.

Summary

Statistical insights

We illustrated building a multivariable logistic regression model using an example data set that has a similar structure to the drop test data (i.e., one categorical variable and one continuous variable). This multivariable model incorporates all of the data instead of building separate models – one for each level of the categorical variable.

We illustrated how to statistically determine which variables should be included in a model and if the categorical variable can be simplified. This model is a tool that gives a single equation to estimate success rates for any values of the two variables. We also illustrated inverse predictions, which estimate the value of a continuous variable that can achieve a given success rate.

Implications

In the exercises, you will build a single multivariable logistic regression model for the drop test data. You'll use that model to describe the performance of each Screen Type across the range of Drop Heights, estimate success rates, make inverse predictions, compare results, and summarize the findings.

Exercises

Use the data in file [Drop Test 3.jmp](#) data set to answer the following questions:

1. a) Build a single multivariable logistic regression model that describes the data.

Hint: Begin by creating a graph that displays the success rate by Drop Height for all the Screen Types. Next, fit a full logistic regression model to all the data using the Fit Model>Generalized Linear Model platform. The model should contain Screen Type, Drop Height, and the Screen Type by Drop Height interaction.

b) Evaluate a statistical test to determine if the interaction term is useful in the model. If needed, reduce the model so that it contains only statistically significant terms.

c) Perform a set of statistical tests comparing all possible pairs of Screen Types (A vs. B, A vs. C, A vs. Competitor, B vs. C, B vs. Competitor, and C vs. Competitor). Is it necessary to conduct these comparisons separately for each Drop Height or can it be done across all Drop Heights together? If so, why?

d) If there's statistical justification, reduce the model further by combining Screen Types that are not statistically different.

2. Provide the equation and a visual representation of your final recommended statistical model.

3. Estimate and provide interpretation of the odds ratios.

4. a) Provide a point estimate and a 95% confidence interval estimate of the Success Rate at a 1.0m Drop Height for each Screen Type. Is there statistical evidence to suggest that one or more of the Screen Types reach the desired goal of having a 97% Success Rate at the 1.0m Drop Height? Create a table and a visualization that display all the confidence intervals as a means to present all the results together.

b) If case JMP034: Durability of Mobile Phone Screen – Part 3 has been completed, compare the precision of the estimates of the Success Rates using a separate logistic regression model for each Screen Type versus the one multivariable model built here. Which estimates are more precise? What might be a reason for this?

5. a) Provide the point estimate and a 95% confidence interval estimate of the Drop Height that would result in a 97% Success Rate for each Screen Type. Is there statistical evidence to suggest that one or more of the Screen Types reach the desired goal of having a 97% Success Rate at the 1.0m Drop Height? Create a table and a visualization that display all the confidence intervals as a means to present all the results together.

b) If case JMP034: Durability of Mobile Phone Screen – Part 3 has been completed, compare the precision of the inverse predictions using a separate logistic regression model for each Screen Type versus the one multivariable model built here. Which estimates are more precise? What might be a reason for this?

6. What is the estimated minimum Drop Height that each Screen Type is more likely to be damaged versus not damaged?
7. What is the estimated greatest Drop Height that up to 5% of screen would be damaged versus not damaged?
8. Create a graphical display that shows the error in the model's predictions.
9. Provide an executive summary of your conclusions by choosing two or three visualizations and up to 10 brief sentences summarizing the results of the analyses. Include recommendations for next steps. Do not use any statistical terminology.