



Sawtooth Software

RESEARCH PAPER SERIES

Situational Choice Experiments for Marketing Research: How to Design, Analyze and Report Them

Keith Chrzan
Sawtooth Software, Inc.

Situational Choice Experiments for Marketing Research: How to Design, Analyze and Report Them

Keith Chrzan, Sawtooth Software

November 2022

Summary: Situational choice experiments (SCE) resemble the more commonly used choice-based conjoint experiments, but (a) they have different experimental design requirements and (b) they ask for different cognitive operations on the part of survey respondents. These differences lead to a statistical model that differs from the conditional multinomial logit typically used in conjoint experiments. After a brief review and taxonomy of choice models and choice experiments, we illustrate the process of executing a SCE, tracing the process from design to data formatting to analysis and reporting. Appendices show how to prepare the data for analysis in different Sawtooth Software packages. In addition to the common example of a SCE fielded to enable pharmaceutical marketers to understand how physicians make therapy decisions, we cover several brief case studies showing how academics and marketers have applied situational choice experiments in practice.

Introduction

Situational choice experiments differ from the choice-based conjoint experiments more commonly used in marketing research. A bit of background will clarify how situational choice experiments differ from other kinds of choice experiments and from other kinds of choice models.

Choice Models

The go-to analysis engine for choice modelers is the conditional multinomial logit (MNL) model (McFadden 1974, Ben-Akiva and Lerman 1985, Train 2003). Imagine that we have a set of attributes that describe products (or services, or, more generally, alternatives) and those products differ from one another in terms of the specific levels they have for the attributes. For example, we might have a set of mobile phones described in terms of attributes such as brand, color, storage and price and a specific set of levels of those attributes for each specific phone, e.g. a black Google Pixel with 128GB of storage and a price of €159. We collect information about decision makers' choices and we model choice among the products as a function of attributes and levels coded to capture the product profiles and with choice among profiles as the dependent variable.

A special case of MNL is the polytomous multinomial logit (P-MNL) model (Theil 1969, Hoffman and Duncan 1988). With P-MNL, the attributes and levels describe not the products, but the chooser, the situation, or the context in which a decision occurs. Respondents choose among the alternatives, but the attributes and levels we have are invariant across those alternatives, because they describe the situation, not the alternatives.

We can execute either of these models with cross-sectional survey data or with data drawn from behavioral databases. Choice modelers know such models as revealed preference (RP) models, because behaviors reveal the decision makers' values. McFadden, in his Nobel lecture describes the classic example of an RP conditional MNL model built from survey data: using data from interviews with

commuters in San Francisco, his team predicted the travel mode choices respondents reported (taking a bus, driving, etc.) as a function of factors such as travel times, wait times and costs for each mode, given each respondent's home and work locations (McFadden 2001, McFadden et al 1977). Guadagni and Little (1983) provide an early example of a RP model based on observed behaviors, using scanner panel data from grocery retailers to predict ground coffee purchases as a function of the brand, package size and prices of the products present in the retail grocers.

Similarly, we can imagine an RP model using P-MNL. In the first model the author worked on we observed choices of pregnancy outcomes (terminate pregnancy, carry baby to term then give it up for adoption, carry baby to term and keep it) among female inmates at a state prison system in the southern USA. We predicted these choices as a function of their situations (facts about their incarceration like length of term, severity of crime, availability of parole plus demographics like age, education, family structure, household income and frequency of attending religious services). Note the invariance of these situational factors across the choice outcomes.

Choice Experiments

If we also apply experimental control, then instead of having an RP model, we have a stated preference (SP) model. In SP models, we show survey respondents hypothetical choice sets of two or more alternatives and we ask which they would select. The most well-known SP models feature multi-profile choice sets analyzed via conditional logit, a combination known as choice-based conjoint experiment or a discrete choice experiment (Louviere and Woodworth 1983, Louviere 1988, Louviere *et al.* 2000). In a choice-based conjoint experiment, we show each respondent a series of a dozen or so questions that look like this:

Figure 1 – Example Choice-Based Conjoint Question

To replace your broken refrigerator, which option would you choose?

(1 of 12)

	Option 1	Option 2	Option 3
Brand	Frigidaire	Kenmore	GE
Color	White	Black	Stainless steel
Freezer configuration	Freezer on top	Side-by-side	Freezer on bottom
Warranty	1 year parts and labor	5 years parts and labor	90 days parts and labor
Price	\$1,189	\$829	\$469
	Select	Select	Select

The experimental design behind the profiles in each question makes the attribute levels independent, so that upon analysis we can quantify the value of each level of each attribute, which values we call

“utilities.” Our data file codes the attributes that comprise the three alternatives and includes the dependent variable (whether the respondent chooses alternative 1, 2 or 3). We then use the conditional MNL model to predict the choice selections and to produce a vector of utilities (model coefficients) for each level of each attribute.

This paper concerns a less common SP model, one we will call a situational choice experiment or “SCE.” In a given SCE question, we have a single experimentally designed profile that describes the situation or context of a decision, and then two or more fixed alternatives from which the respondent can choose. In other words, the experimentally designed profile changes from question to question, but the choice alternatives do not. Across questions, the profiles conform to an experimental design so that upon analysis we can quantify the independent effect of each attribute level on each choice alternative. A single SCE question might look like this:

Figure 2 – Example Situational Choice Experiment Question

Patient 1
81 year old Female BMI: 26.5 Moderate anxiety Former smoker Moderately active

For the patient above, which therapy would you be most likely to prescribe to treat newly diagnosed hepatic sarcoidosis?

☐ Lotomil
☐ Vicodin
☐ Darvon
☐ Opana
☐ Diet and exercise

In the next several questions, the description of the patient changes but the five choice alternatives remain the same. So, the patient in the next question might be a 64 year old inactive female smoker with a BMI of 23 and moderate anxiety, for example. If we are to model the different choice probabilities of the five alternatives, and because the patient description is the same for Lotomil as it is for Darvon as it is for the other alternatives, then the only way for the probabilities to change is if we have a different set of utilities for each of the five choice alternatives. The P-MNL provides just such a matrix of alternative-specific utilities. Whereas conditional MNL produces a single vector of utilities, one for each level of each attribute, and choice probabilities for alternatives differ because the attributes and levels describing the alternatives differ, with P-MNL, we have a matrix of utilities, one vector of utilities for each alternative, and choice probabilities differ across alternatives because the alternatives have different sets of utilities.

To review, a situational choice experiment differs from a choice-based conjoint experiment. Choice-based conjoint features experimentally designed sets of two or more profiles each and analyzes respondents’ choices among profiles as a function of the attributes and levels of the profiles, using conditional MNL. An SCE, on the other hand, features experimentally designed profiles, shown one at a

time, which describe the attributes and levels of the choice situation or context and uses polytomous MNL to generate utilities that predict choices among a set of fixed alternatives.

Few marketers know about SCEs and there does not seem to be a single source reference about them, two limitations this paper may help remedy. The next section walks through the steps involved in conducting an SCE, including a discussion of sample size. The final section describes four commercial marketing examples of SCEs.

Executing a Situational Choice Experiment

Research Design

Before choosing a design strategy, we need to know how many attributes and how many levels per attribute to accommodate. Attributes in a patient study could include the patient's age, sex, details about their condition, concomitant conditions and so on. Levels are the specific values of those attributes shown to respondents, for examples the ages 50, 65 and 80 years old.

SCE's sample size requirements increase with the number of parameters, as discussed below, so more attributes of more levels can increase sample size needs. That said, the stimulus in an SCE question is a single profile, likely requiring less information processing on the part of respondents than would a choice-based conjoint experiment with a similar number of attributes. Most SCEs the author builds have fewer than 10 attributes.

One can construct an SCE design using any software that produces efficient designs for single-profile experiments. For the special case in which all the attributes have the same number of levels (i.e. for a "symmetric" experiment) one can use a traditional orthogonal main effects experimental design plan such as provided by Addelman (1962) or Hahn and Shapiro (1966). The experimental designers currently available in R also make orthogonal main effects designs. Ngene from Choice Metrics, SAS, and Sawtooth Software's Lighthouse Studio all have efficient design algorithms that fit the bill, regardless of the number of attributes or the number of levels per attribute.

In practice we usually use an efficient design program that makes the design in several blocks. Each respondent receives one block of questions and we randomly assign different blocks to different respondents. Using efficient designs and multiple blocks allows us to test for interactions and to include them in our model when significant. Note that efficient designs allow the user to specify how many questions each block contains while the orthogonal design plans tend to come in set sizes (e.g., a design for four 3-level attributes has a nine question orthogonal design while a design for five 3-level attributes requires 18 questions).

Design in hand, we can produce the SCE questionnaire, giving each respondent the questions from a given block in the experimental design.

In addition to asking a single response as in the example above, we might ask for separate responses for different segments, like this:

Figure 3 – Example of a Multi-Response SCE Question

Patient 2

88 year old
Female
BMI: 20.6
Mild anxiety
Currently smokes tobacco
Extremely active

For each of these three types of hepatic sarcoidosis patients, which therapy would you prescribe?

	Lotomil	Vicodin	Darvon	Opana	Diet and exercise
Mild symptoms	☐	☐	☐	☐	☐
Moderate symptoms	☐	☐	☐	☐	☐
Severe symptoms	☐	☐	☐	☐	☐

In some cases, we might want respondents to allocate their last 10 patients to treatments, or to ask about the percentage of patients to which they would recommend each therapy:

Figure 4 – Example Allocation SCE Question

Patient 7

88 year old
Male
BMI: 26.5
Mild anxiety
Currently vapes
Moderately active

To what percent of your hepatic sarcoidosis patients like the one described above would you prescribe each of the following therapies?

Lotomil

Vicodin

Darvon

Opana

Diet and exercise

Total

Data Structure

After collecting survey respondents' choices for each question, we append them to the experimental design matrix and we have our data file ready for analysis. For example, for our hepatic sarcoidosis

experiment our data from a single respondent might look like this, assuming we wanted to treat each of the independent variables as a categorical predictor:

Table 1 – SCE Data File, All Categorical Predictors

<u>Question</u>	<u>Age</u>	<u>Sex</u>	<u>BMI</u>	<u>Anxiety</u>	<u>Smoking</u>	<u>Activity</u>	<u>Choice</u>
1	3	1	3	3	4	1	2
2	1	2	1	2	2	2	4
3	5	2	2	1	3	3	3
4	4	1	2	2	1	2	3
5	2	1	3	1	2	3	4
6	2	2	1	3	1	1	1
7	4	2	3	2	3	1	2
8	5	1	1	1	4	2	1
9	1	2	2	3	4	3	2
10	3	1	1	2	3	3	1
11	3	2	3	1	1	2	2
12	5	1	2	3	2	1	2

In the analysis software we would specify our predictors to be categorical (or “factors” in R); the software will recode these variables appropriately, usually treating a k-level variable as k-1 dummy variables. Alternatively, we can treat age and BMI not as categorical variables but as continuous variables for which we want to estimate a single slope coefficient rather than a coefficient for each dummy variable. In this case, our data from the first respondent might look like this:

Table 2 – SCE Data File, Some Continuous Predictors

<u>Question</u>	<u>Age</u>	<u>Sex</u>	<u>BMI</u>	<u>Anxiety</u>	<u>Smoking</u>	<u>Activity</u>	<u>Choice</u>
1	62	1	31.0	3	4	1	2
2	25	2	20.6	2	2	2	4
3	88	2	26.5	1	3	3	3
4	81	1	26.5	2	1	2	3
5	43	1	31.0	1	2	3	4
6	43	2	20.6	3	1	1	1
7	81	2	31.0	2	3	1	2
8	88	1	20.6	1	4	2	1
9	25	2	26.5	3	4	3	2
10	62	1	20.6	2	3	3	1
11	62	2	31.0	1	1	2	2
12	88	1	26.5	3	2	1	2

To create our analysis data file we simply concatenate the 12 rows of data we get from each respondent; in this example where each respondent makes choices for each of 12 patients, our analysis data file would contain 12 times as many rows as we have respondents.

Modeling and Reporting

One can find canned P-MNL routines in general purpose statistical software like SAS, SPSS or SYSTAT. You can also access such programs through specialty choice modeling packages like Nlogit from Econometric Software, the mlogit and nnet packages in R or in the MBC program for logit modeling from Sawtooth Software. Appendix 1 contains a screenshot showing the settings in the MBC software that produce a P-MNL model.

Because P-MNL is a special case of conditional logit, you can also trick conditional logit software into running P-MNL analysis, a fact that comes in handy when you collect constant sum or allocation data as in Figure 4 above. Appendix 2 shows how to format SCE data for analysis in Sawtooth Software's conditional MNL programs, CBC/LC and CBC/HB.

Though software packages format their results differently, each produces a set of model coefficients (utilities), one per level in the dependent variable. One of the vectors, the last on the right, represents the reference level of the dependent variable and all coefficients in that vector are zero. For example, the utilities resulting from our hepatic sarcoidosis example might look like this:

Table 3 – SCE Utilities (disguised example)

	<u>Lotomil</u>	<u>Vicodin</u>	<u>Darvon</u>	<u>Opana</u>	<u>Diet and exercise</u>
Constant	-1.162	0.304	1.661	0.111	0.000
62 year old	-1.205	0.554	-1.314	-1.227	0.000
81 year old	0.712	-0.386	0.132	-1.122	0.000
88 year old	0.000	0.000	0.000	0.000	0.000
Female	-0.762	-0.933	-0.948	1.498	0.000
Male	0.000	0.000	0.000	0.000	0.000
BMI: 20.6	0.804	-0.997	1.634	0.176	0.000
BMI: 26.5	-2.330	-1.188	-1.159	0.303	0.000
BMI: 31.0	0.000	0.000	0.000	0.000	0.000
Mild anxiety	-1.091	-1.373	0.977	0.903	0.000
Moderate anxiety	0.616	0.128	-0.943	0.263	0.000
Severe anxiety	0.000	0.000	0.000	0.000	0.000
Non-smoker	-0.006	-0.632	-1.661	-0.548	0.000
Former smoker	0.451	0.450	0.913	0.949	0.000
Currently vapes	0.023	-0.261	-0.674	0.043	0.000
Currently smokes	0.000	0.000	0.000	0.000	0.000
Inactive	-0.079	-0.153	0.459	1.006	0.000
Moderately active	0.312	-0.643	0.142	-1.215	0.000
Extremely active	0.000	0.000	0.000	0.000	0.000

You'll notice that each column has an alternative-specific constant (the top row of utilities) which measures the utility attached to that alternative apart from the utility captured by the various levels of the various attributes.

Note that in this model, we treated each of the predictor variables as categorical, and we used dummy coding of those categorical variables, so that each has a reference level row set to zero. In many studies we can instead model quantitative attributes as linear functions. Of course, the statistical software will also produce standard errors and model fit statistics. These allow us to calculate the p-value of each of our utilities and to test alternative model formulations (e.g., whether the categorical or linear coding best captures the effect of quantitative variables, whether interactions significantly improve the model, etc.). We can also generate the utilities at the respondent level using mixed logit or hierarchical Bayesian (HB) MNL, again a topic covered in the Appendix.

Some academic modelers prefer to express the results as odds, so they exponentiate the utilities. With most marketing audiences, however, taking a number they do not understand, transforming it in a way they understand even less to produce numbers without an intuitive meaning is hardly a winning communication strategy.

Simulations

Clients find SCE results easier to understand when we deliver the model results to clients in Excel simulators. Using the familiar logit choice rule, the same equation that calibrates the utilities from the choice response data in the first place, we can estimate shares for each level of the dependent variable for any profile specified in terms of the attributes and levels. As a result, marketers need not even see the utilities: using drop-down menus for each attribute, the user can create a profile and then get share estimates for each level of the dependent variable:

Figure 5 – Screenshot of Simulator

The screenshot displays a web-based simulator interface. At the top, there is a section titled "Patient Profile" with a list of attributes and their corresponding values: Age (81 year old), Gender (Female), BMI (26.5), Anxiety (Moderate anxiety), Smoking status (Currently vapes), and Activity Level (Extremely active). The Activity Level is shown as a dropdown menu. Below this, there is a section titled "Preference Shares" which lists five categories with their respective percentages: Lotomil (1.28%), Vicodin (2.26%), Darvon (3.38%), Opana (69.77%), and Diet and exercise (23.31%).

Patient Profile	
Age	81 year old
Gender	Female
BMI	BMI: 26.5
Anxiety	Moderate anxiety
Smoking status	Currently vapes
Activity Level	Extremely active

Preference Shares	
Lotomil	1.28%
Vicodin	2.26%
Darvon	3.38%
Opana	69.77%
Diet and exercise	23.31%

Sensitivity analysis showing how shares change according to changes in the profile can directly inform marketing decisions, with no need for marketers to interpret the utilities themselves.

Other Modeling Options

While modelers usually prefer mixed logit (which produces a set of utilities for each individual respondent) for their choice-based conjoint questions, the same does not necessarily hold for SCEs. Compared to a choice-based conjoint experiment, where observations of multiple product profiles in each choice set inform a single vector of utilities, in an SCE respondents see one profile per question and the model outputs multiple vectors of utilities. As a result, the sparse data matrix for SCE modeling may lack the amount of information needed to support robust respondent-level utility estimation.

Sample Size Considerations

Peduzzi *et al.* (1996) recommend that sample size for a logit model should be at least 10 times the number of parameters in the model divided by the choice probability for an alternative:

$$n \geq \frac{10k}{p}$$

where

- n is the sample size
- k is the number of non-zero parameters (utilities) to be estimated by the model
- p is the probability of an alternative chosen

A SCE with six 4-level attributes and five choice alternatives, will have a constant and $6 \times 3 = 18$ parameters per utility function, and four non-zero utility functions, for a total of 76 parameters. With five choice alternatives, the average probability of choice is 20%, which suggests a number of observations of at least

$$n \geq \frac{10(76)}{0.20} \text{ or } n = 3,800$$

If we ask 10 SCE questions of 380 respondents, we can achieve our minimum sample size target of 3,800 observations.

Another sample size rule of thumb comes from thinking about the sampling error around simulation shares. We know that a sample size of 100 produces margins of error of 0.098 for percentages and that halving that margin of error requires quadrupling the sample size. The sample size of 380 above would produce shares with margins of error of 0.05. For small experiments, the simulator margin of error will drive sample size more than will the Peduzzi *et al.* rule of thumb, but the latter will have more influence as the number of attributes, levels and (especially) choice alternatives increases.

For example, an experiment with four 2-level attributes and three choice alternatives would suggest a minimum of 300 observations:

$$n \geq \frac{10(10)}{0.333} \text{ or } n = 300$$

Asking each of 30 respondents 10 SCE questions would get you the minimum number of observations from the Peduzzi *et al.* formula, but the sampling error around shares for a sample of size 30 would be an excessive +/- 0.18, or 18 percentage points.

Power Analysis

Some academic users of SCEs need to write grant proposals and to express their sample size decisions in terms of the precision of their estimates and the power for finding significant differences. For these cases one can estimate the size of the standard errors by creating a data file with as many copies of the design blocks as needed for a given sample size, then populating the choices with random numbers between 1 and the number of choice alternatives (in Excel one can use the =RANDBETWEEN() function to create these random responses). Modeling this data file with a P-MNL analysis program will produce standard errors for each model parameter. One can then use a shortcut method to estimate standard errors for any other sample size. For example, say you generated a random data set with 100 respondents answering eight questions each (or 800 observations in total). To estimate the standard errors for any other number of observations, just take the square root of the ratio of the original 800 observations to the new number of observations. For a given parameter with a standard error of 0.05, the new sample size of 1,600 observations would yield standard errors of $\sqrt{800/1600}$ or 0.707 times as large (i.e., we expect the standard error of 0.05 to shrink to 0.035 if the number of observations doubles from 800 to 1,600). We can do similar calculations on power analyses for different sample sizes.

Case Studies

Therapy Choice Experiments

The patient type experiments such as the hepatic sarcoidosis example above constitute the archetypal commercial application of SCE. These experiments tend to have small numbers of attributes, typically four to eight, and also small sample sizes: the populations from which we draw physician-respondents are small to begin with and like many B2B respondents, they expect expensive honoraria for their participation in survey research.

Antibiotic Decision-Making for Nursing Home Patients

As reported by Kistler *et al.* (2020), in order to determine which nursing home resident characteristics most influence clinicians' prescribing decisions for antibiotics for suspected urinary tract infections (UTI) in nursing home patients. The SCE that included 10 patient and diagnostic characteristics and it revealed that existing guidelines for diagnosing a UTI requiring antibiotics were considered less important than some diagnostic test results not currently included in clinical guidelines.

Patient Segmentation

The data structure for SCE, a choice among two or more options on the basis of variables describing a single profile, makes SCE appropriate for machine learning methods, particularly decision trees like CART. Decision trees branch the total sample into successively smaller segments, highly discriminating with respect to the dependent variable, therapy choice. Moreover, tree models generate these

branches using the variables that define the profiles. For example, a client wanted to understand how physicians group patients for differential treatment. A CART decision tree analysis identified five segments of patients differing with respect to three situational variables (age, disease progression history and insurance coverage). The pharmaceutical marketers created messaging plans for the two segments of patients for whom their therapeutic offering fit best.

The ability to predict, combined with the visualization available from a decision tree analysis might suggest using trees for the choice modeling itself. In a recent empirical comparison of nine SCE studies, the P-MNL model had better cross-validated predictive validity than CART trees in six of the studies while CART trees outperformed P-MNL in three of the studies (Chrzan and Retzer 2019). Moreover, when large sample sizes combine with the absence of correlations among the attributes (owing to the experimental design) the trees often get complex, with many branches and leaves. This complexity can negate the benefit of visualization, as a tree with 40 branches will be difficult for researchers to communicate and for marketers to digest.

Durable Acquisition Decision

In a study of high-priced industrial durables, a client wanted to know which product profiles drive preference among brands, a question that a choice-based conjoint experiment can answer very nicely. The client also wanted to know, however, whether the industrial customers would opt to lease or to buy the products they chose. For this, we created a choice-based conjoint experiment in which respondents faced a choice among three experimentally designed product profiles. After choosing the profile they most preferred, respondents answered a second question about that most preferred profile: given what they know about the market, the products and prices available and their budgets, would they (a) buy the product they selected, (b) lease the product they selected or (c) neither buy nor lease the product. We built both models into an Excel-based simulator and the client was able to see, for any product specified and in any competitive set, how many respondents preferred it more than the other products, and how many of those would lease the product, buy it, or go without it.

Retirement Hybrid Experiment

A financial institution specializing in retirement savings accounts wanted to be able to forecast the proportion of its savers who would choose to retire and start drawing down their retirement accounts, as a function of changing economic conditions. The SCE featured attributes and levels that described the economic conditions (interest rates on investments, growth in home prices, inflation, recent and forecasted economic growth). In addition, the financial institution had information on their savers stored in a database: the saver's age, the amount of savings with the institution and with other institutions, the saver's income, the saver's credit history and so on. This additional information didn't conform to an experimental design but we included it in the model because it would seem silly to respondents if we asked a 64 year old married male with \$750,000 in retirement savings to imagine his retirement choice if he were a 58 year old single woman with \$1,200,000 in retirement savings. In this hybrid experiment, we had experimental control over the variables we designed into the SCE, but we also had the database variables to provide even more context to the respondents' reactions to the SCE questions.

Identifying Incremental Spend

A fashion retailer wanted to understand if offering a 'Buy Now Pay Later' credit offering would increase the average spend across its customer base. While a standard conjoint would tell us which of the offerings would be the most appealing, this client wanted to understand if the new mechanism would drive incremental spend – or purchases which would not have occurred without a credit option in place. The SCE presented a range of differing spend scenarios alongside an accompanying credit option. Within each scenario respondents were asked if they would buy with cash or the credit option, or not buy at all. In order to identify spend incrementality, those who chose the credit option were asked if they would have still purchased the item if the credit option was no longer available. This study helped the retailer to identify the optimal credit option, as well as the optimal spend threshold to offer it at.

Maximizing Contract Renewal Rates

A telecom provider sought to maximize the amount of revenue it could generate from customers reaching the point of contract renewal. Typically, a customer's price increases sharply once the initial contract expires, leading to a large level of churn as customers leave for cheaper competitive offers. A SCE presented customers with a series of hypothetical contract renewal letters, highlighting their current spend, planned contract price increase and a series of potential loyalty benefits. Within each scenario customers could choose to accept the offer, exit their contract or contact the company for a better rate. This model allowed the company to set the standard price increase at a level which minimized churn while maximizing the monthly spend of those who signed up to a new contract.

.....

Thanks to my colleague Dean Tindall for the last two case studies.

References

- Addelman S. and O. Kempthorne (1961). *Orthogonal main effect plans*. Report for the United States Airforce Aeronautical Research Laboratory. Contract no. AF 33 (619)-5599, Project no. 7071, Task no. 70418. Wright Patterson AFB, OH: ARL.
- Addelman S. (1962) "Orthogonal main-effect plans for asymmetrical factorial experiments," *Technometrics* 4(1): 21–46.
- Ben-Akiva M. and S.R. Lerman (1985) *Discrete Choice Analysis: Theory and Application to Travel Demand*. Cambridge, MA: MIT Press.
- Chrzan, K. and J. Retzer (2019) "Trees, forests and situational choice experiments," in *2019 Sawtooth Software Conference Proceedings*. Provo: Sawtooth Software, pp. 121-129.
- Guadagni, P.M. and J.D.C. Little (1983) "A logit model of brand choice calibrated on scanner data," *Marketing Science* 2(3): 203–238.
- Hahn G. and S. Shapiro (1966) *A catalogue and computer program for the design and analysis of orthogonal symmetric and asymmetric fractional factorial designs*. Report for Research Development Center. Report no. 66-C-165. Schenectady, NY: General Electric Corporation.
- Hoffman, S.D. and G.J. Duncan (1988) "Multinomial and conditional logit discrete-choice models in demography," *Demography* 25: 415-427.
- Hosmer, D.W. and S. Lemeshow (1989) *Applied Logistic Regression*. New York: Wiley.
- Kistler, C. E., A.S. Beeber, S. Zimmerman, K. Ward, C.E. Farel, K. Chrzan, C.J. Wretman, M.H. Boynton, M. Pignone, P.D. Sloane (2020). *Nursing Home Clinicians' Decision to Prescribe Antibiotics for a Suspected Urinary Tract Infection: Findings From a Discrete Choice Experiment*. *Journal of the American Medical Directors Association*. doi:10.1016/j.jamda.2019.12.004
- Louviere, J.J. and G.G. Woodworth (1983). "Design and analysis of simulated consumer choice or allocation experiments: an approach based on aggregate data," *Journal of Marketing Research* 20: 350-367.
- Louviere, J.J. 1988. *Analyzing decision making*. Sage University Paper series on Quantitative Applications in the Social Sciences, 07-067. Beverly Hills: Sage.
- Louviere, J.J., D.A. Hensher and J.D. Swait (2000) *Stated Choice Models: Analysis and Application*. Cambridge: Cambridge University.
- McFadden, D. (1974) "Conditional logit analysis of qualitative choice behavior," in: P. Zarembka (ed.) *Frontiers in Econometrics*. New York: Academic Press, pp. 105–142.
- McFadden, D., A. Talvitie and Associates (1977) Report for Urban Travel Demand Forecasting Project, Phase 1 Final Report Series, Volume V. Report no. UCB-ITS-SR-77-9. Berkeley: University of California.

McFadden, D. (2001) "Economic Choices," *American Economic Review* 91(3): 351-378.

Peduzzi P., J. Concato, E. Kemper, T.R. Holford, and A.R. Feinstein (1996) "A simulation study of the number of events per variable in logistic regression analysis," *Journal of Clinical Epidemiology* 49(12): 1373-1379.

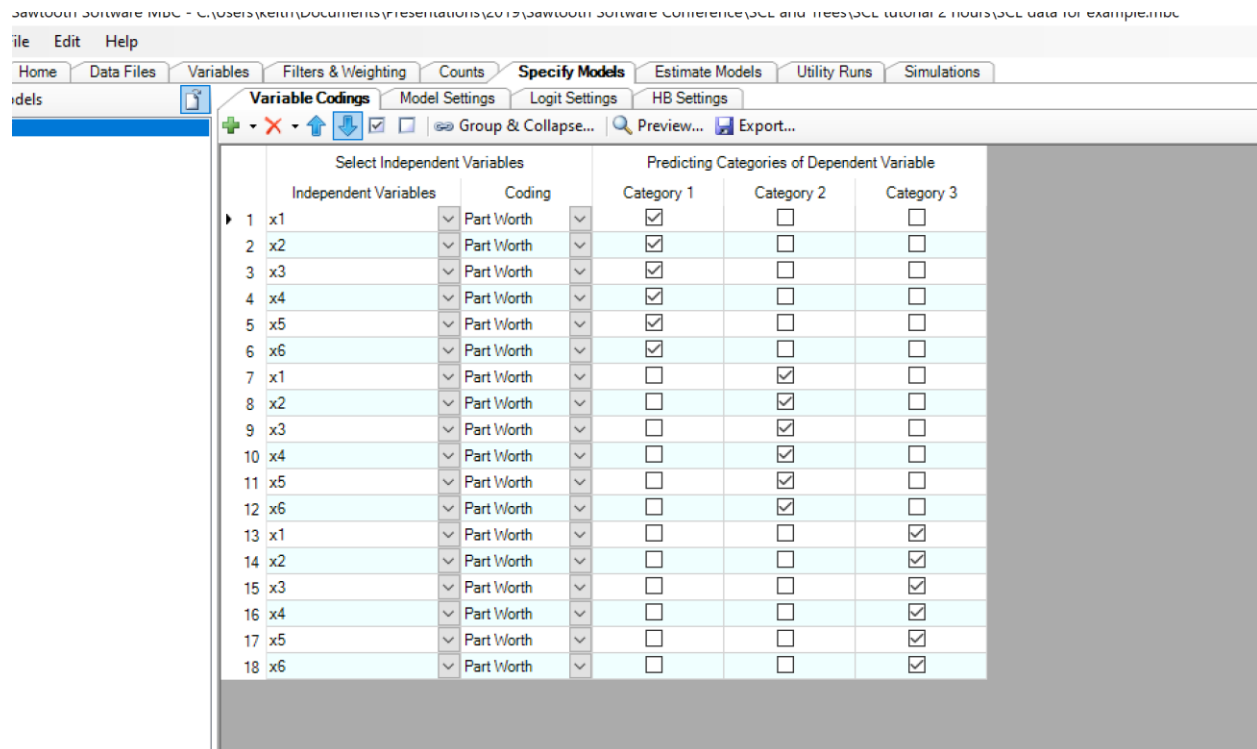
Theil, H. (1969) "A multinomial extension of the linear logit model," *International Economic Review*, 10(3): 251-259.

Train, K. (2003) *Discrete Choice Methods with Simulation*. Cambridge: Cambridge University Press.

Appendix 1

MBC Settings for P-MNL

Assume we use 6 categorical variables to predict a 4-category dependent variable. We set up the variables in the MBC software so that each of the 6 variables predicts each of the first 3 levels of the DV (again, the 4th level is the reference level, set to all zeros so that the model may be identified).



Finally, in the Model Settings tab, indicate that you want the model to contain a constant – this will produce the constant (the first row in the utility table above), which quantifies the appeal of each alternative not explained by the attributes and levels in the experimental design.

Appendix 2

Formatting SCE Data for Analysis in Sawtooth Software’s Conditional MNL Programs

For this example, assume we have three predictors, X1-X3, and a dependent variable (DV) with 3 levels. If each respondent sees 6 choice questions, then a data set for 1,000 respondents will look something like this:

ID	Set	X1	X2	X3	DV
1	1	5	3	2	2
1	2	1	2	4	3
1	3	3	1	3	2
1	4	4	3	1	1
1	5	3	3	3	1
1	6	5	3	4	2
2	1	4	2	3	1
2	2	1	1	2	2
...					
1000	6	2	2	2	1

If you have the MBC software, you could import the data as is and start running your analysis. If, however, you want to trick Sawtooth Software’s CBC/LC software (to run a total-sample P-MNL model) or CBC/HB software (to run respondent-level P-MNL models) then you’ll need to reformat the data file using the following steps.

First, expand that original data file as follows

- Convert each row of the original file into 3 rows in the new data file
- Add a column to contain a variable called “Concept” which tells the software how many choice alternatives there are and which alternative corresponds to each row (in this case each choice question has 3 alternatives, hence 3 rows in the expanded data file)
- Add a column that’s a duplicate of the Concept column, but label it ASC – this column codes the alternative-specific constant (Sawtooth Software’s CBC/HB and CBC/LC use effects coding rather than dummy coding, which will zero-center the utilities for the ASC so that all 4 columns may have non-zero values).
- Expand X1-X3 so that there are two columns for each: Columns AX1-AX3 and BX1-BX3. The AX variables are just copies of X1-X3 that apply ONLY to the first choice alternative while the BX variables are copies of X1-X3 that apply ONLY to the second choice alternative. Notice that AX1-BX3 are coded all zeros for the third alternative (again, it’s the reference alternative, coded all to zero to allow the model to be identified).
- What results is a data file whose first 6 rows look like this (yellow highlight shows how the data transforms for the first respondent’s first choice question:

Original						New Data Set										
ID	Set	X1	X2	X3	DV	ID	Set	Concept	ASC	AX1	AX2	AX3	BX1	BX2	BX3	DV
1	1	5	3	2	2	1	1	1	1	5	3	2	0	0	0	0
1	2	1	2	4	3	1	1	2	2	0	0	0	5	3	2	1
1	3	3	1	3	2	1	1	3	3	0	0	0	0	0	0	0
1	4	4	3	1	1	1	2	1	1	1	2	4	0	0	0	0
1	5	3	3	3	1	1	2	2	2	0	0	0	1	2	4	0
1	6	5	3	4	2	1	2	3	3	0	0	0	0	0	0	1
2	1	4	2	3	1											
2	2	1	1	2	2											
...																
1000	6	2	2	2	1											

Next, notice that we've added a column for the dependent variable, coded as 1 for the chosen alternative and as 0 for the non-chosen alternatives.

Original						New Data Set										
ID	Set	X1	X2	X3	DV	ID	Set	Concept	ASC	AX1	AX2	AX3	BX1	BX2	BX3	DV
1	1	5	3	2	2	1	1	1	1	5	3	2	0	0	0	0
1	2	1	2	4	3	1	1	2	2	0	0	0	5	3	2	1
1	3	3	1	3	2	1	1	3	3	0	0	0	0	0	0	0
1	4	4	3	1	1	1	2	1	1	1	2	4	0	0	0	0
1	5	3	3	3	1	1	2	2	2	0	0	0	1	2	4	0
1	6	5	3	4	2	1	2	3	3	0	0	0	0	0	0	1
2	1	4	2	3	1											
2	2	1	1	2	2											
...																
1000	6	2	2	2	1											

That's it. Do this for all 6 choice sets and for all 1,000 respondents and you have a data file ready to be imported into CBC/HB or CBC/LC for analysis. Though we designed the software to run conditional MNL, because P-MNL is a special case of conditional MNL, reformatting the data in this way (essentially so that all effects are alternative-specific) enables the software to run the P-MNL model. How cool is that?