

On Conjoint Analysis and Quantal Choice Models

I. Conjoint Analysis

Conjoint analysis (or trade-off analysis) has become increasingly used in marketing. Essentially, the techniques attempt to parse the appeals of a product or concept into a set of factors, assign utility values to the separate levels of each factor, and finally, under the assumption of separability, determine the utility value of any product or concept by adding (or multiplying) the utility values of the individual levels of each of the factors embodied in the product or concept. A good expository description of conjoint analysis is given in Green and Srinivasan (1978). For notational convenience, I present a brief description herein.

Formally, if we are considering n factors, with the i th factor ($i = 1, \dots, n$) having m_i levels, then we shall represent a bundle as an n -tuple $(i_1, i_2, \dots, i_n)$, where

i_1 can take on the integer values $1, \dots, m_1$

i_2 can take on the integer values $1, \dots, m_2$

.

.

.

i_n can take on the integer values $1, \dots, m_n$.

There are $M = \prod_{j=1}^n m_j$ possible bundles that can be composed from the $m_1, \dots, m_n$ respective levels of the n factors.

The conjoint analysis model, along with various methods used to elicit trade-off data and to estimate utilities, are described. Comparisons are made between features of this model and that of the quantal choice model.

Sometimes it is convenient to speak of "subbundles," that is, bundles in which the level of one or more factors is unspecified. We can describe subbundles succinctly by the device of introducing for each factor a level 0, denoting the nonexistence of this factor in the subbundle under consideration. The number of bundles and subbundles that can be comprised from the $m_1, \dots, m_n$ respective levels of the n factors is

$$\prod_{j=1}^n (m_j + 1).$$

We denote by $u(i_1, \dots, i_n)$ the utility function of the bundle $(i_1, \dots, i_n)$ and by $v(i, j)$ the utility of the j th level of factor i ($i = 1, \dots, n, j = 0, \dots, m_i$). We say that the utility model is *additive* if $u(i_1, \dots, i_n) = v(1, i_1) + \dots + v(n, i_n)$. We say that the utility model is *multiplicative* if $u(i_1, \dots, i_n) = v(1, i_1) \times v(2, i_2) \times \dots \times v(n, i_n)$. In an additive model, $v(i, 0) = 0$; in a multiplicative model $v(i, 0) = 1$.

The goal of conjoint analysis is to estimate utility values for each level of each factor, such that, upon developing from these (by either addition or multiplication) the utility values of each of the subbundles ranked by the respondent, the ranking of the subbundle utility values matches the respondent's ranking of the subbundles themselves.

There are three main data-elicitation procedures used in conjoint analysis. In the most general "bundle-ranking" approach, one selects from the

$$B = \prod_{j=1}^n (m_j + 1)$$

bundles and subbundles an "appropriate" set of bundles and asks the respondent to rank them. The determination of which are "appropriate" is a topic for discussion in its own right. Many rules have been formulated, such as: (1) do not include a pair of bundles where one clearly dominates the other in utility; (2) make sure each level of each factor is represented equally in the set of bundles to be ranked; (3) make sure that the "design matrix" (to be defined later) is of maximal rank.

In the "trade-off-matrix" approach, one considers the set of $\binom{n}{2}$ pairs of factors, selects some subset of these pairs, and asks the respondent to rank the subbundles comprised of all combinations of levels of each of these pairs of factors (with all other factors set at level 0). The reason this is referred to as a "trade-off-matrix" approach is that one way of presenting all the subbundles comprised of factors i and j is via an $m_i \times m_j$ matrix in which the m_i levels of factor i index the rows, the m_j levels of factor j index the columns, and the (k, l) th matrix element is the subbundle consisting of level k of factor i and level l of

factor j . We may view this task as that of ranking $m_i \times m_j$ subbundles. However, in the trade-off-matrix approach each set of subbundles, defined by the different choices of i and j , is ranked separately and the respondent is not given the task of consolidating his rankings of each of the subbundles presented to him into one overall ranked set.

A third method of eliciting information from which to determine utilities is that of "paired comparisons of subbundles." In this method some number of subbundles from among the $B(B - 1)/2$ such pairs is selected, presented to respondent one pair at a time, and the respondent selects that subbundle which he prefers. (When we restrict the subbundles to be full bundles, the maximum number of pairs reduces to $M(M - 1)/2$. Clearly, if there is monotonicity in the utility of the levels within a factor, this monotonicity should carry over into the utility of the subbundles, and so by dominance the maximum number of pairs can be reduced substantially.)

Suppose we were somehow able to measure the utilities for each of the bundles or subbundles presented to the respondent via any of the above data-elicitation techniques. (In the trade-off-matrix approach, each matrix would elicit $m_i \times m_j$ utility values, where i and j index the factors defining the trade-off matrix. In the bundles-ranking approach, each bundle presented would elicit a utility value. In the paired-comparison-of-subbundles approach, each pair of subbundles would elicit two utility values.) In general, let N denote the number of utilities which the method of data elicitation would produce, if only we were able to measure them.

Let $M = \sum_{i=1}^n m_i$ and V be the M -vector of utilities for each of the factor levels. Let A be the $N \times M$ "design matrix," where the rows correspond to the various bundles or subbundles presented to the respondent in the course of the data-elicitation procedure, and where all elements of the i th row of A are 0 except for those columns corresponding to factor levels present in subbundle i , in which case $a_{ij} = 1$. Let Y be the N -vector of utilities which could have been generated by the data-elicitation procedure, if only we could observe them. Then $Y = AV$, and a "best" V could be determined using the method of least squares.

Unfortunately, we cannot observe Y , but merely the vector $r(Y)$ —a vector of the ranks of the elements of Y (in the bundles-ranking procedure), a vector of the ranks of elements of Y associated with each of the matrices (in the matrix-ranking approach), or pairs of subbundles (in the paired-comparisons-of-subbundles approach). Our problem is to determine V such that $r(AV)$ is "as close as possible" to $r(V)$.

The two most widely used methods of analyzing these data are MONANOVA and MONOTONE regression. The procedure of MONANOVA (Kruskal 1965) seeks to find a monotone (ascending) transformation of the ranks so that we can carry out a main-effects-

only analysis of variance of the transformed values. The main-effects-only analysis-of-variance model is one wherein each cell of the analysis-of-variance data table can be expressed as the summation of each of the main effects defining that cell. The main effects are used as the measures of the "utilities" or values of each of the levels of the variables. The criterion for the success of the transformation is in the ability of the summation of main effects to produce a set of cell values whose ranks match the original input ranks. A measure of this is embodied in Kruskal's measure of "stress," the ratio of the mean-square difference between the fitted and transformed values of the total variance.

The thrust of the monotone-regression (Johnson 1975) procedure is to estimate the analysis-of-variance main effects, as described above, via a regression-like procedure. Monotone regression is used to estimate utilities such that the ranks of the predicted dependent variables match best the ranks of the (unobserved) dependent variables. Here, too, a measure of "stress" or "badness of fit" is used to determine how well the solution (i.e., the set of utilities) reproduces the stated preferences for the objects' features.

II. Quantal Choice Analysis

McFadden describes quantal choice analysis fully in this issue. A most concise description of quantal choice analysis is given in McFadden (1976), which we shall quote here: "The prototypical study of quantal choice considers individual subjects placed in one of N possible experimental choice settings. Each choice setting requires one of J possible responses from the subject. Choice setting n is characterized in general by a vector c_n of measured characteristics of the subject and vectors $x_{1n}, \dots, x_{Jn}$ of measured attributes of the alternatives. In R_n repetitions of the presentation of the choice setting n to a subject, S_{jn} responses of alternative j are observed."

In effect, the S_{jn} are the number of times out of R_n that alternative j is selected in choice setting n . "The response of a subject in choice setting n can be depicted as a drawing from a multinomial distribution with probabilities $(P_{1n}, \dots, P_{Jn})$; we term these the selection probabilities, and note they are "observable" in the sense that they are estimated consistently by the observed relative frequencies $(S_{1n}/R_n, \dots, S_{Jn}/R_n)$ as the number of repetitions approaches infinity."

The underlying assumption is that there exists a (random) utility function u defined on the vectors $x_{1n}, \dots, x_{Jn}$, $n = 1, \dots, N$, and a function also of c_n , such that the probability that j is selected is given by $P_{jn} = \text{prob} [u(x_{jn}, c_n) > u(x_{in}, c_n) \text{ for all } i \neq j]$. Quantal choice models are a set of models describing this probability mechanism.

III. Congruences

Given these descriptions of conjoint analysis and quantal choice analysis, let us find the points of congruence. By virtue of (1) setting the experiment as the selection of the best of J alternatives; (2) introducing c_n , the subject characteristics associated with experiment n ; (3) allowing the x_{in} to be *measured* attributes of the alternatives rather than merely the 0 or 1 values of the columns of the design matrix; and (4) allowing the utility function to be random, one sees that the quantal choice models are in some ways more general and in some ways more restrictive than that of conjoint analysis. Let us discuss each of these points in turn.

In both the bundles-ranking and trade-off-matrix approach, the respondent is given a more complex task than that of the quantal choice model, namely, the selection of only the most preferred subbundle or pair of factor-level combinations. In conjoint analysis he is asked to rank them all. The ranking task may be decomposed into a number of paired comparisons, and we can reformulate these approaches in such a way that, for each trade-off matrix or set of subbundles to be ranked, N is the equivalent number of paired comparisons. The paired-comparisons-of-bundles approach is closest to the quantal choice paradigm, in that the bundles being compared correspond to the $J = 2$ choices to be considered by the respondent. The N is in this case the number of paired-comparison sets presented to the respondent. So in what follows let us assume that $J = 2$.

In general, the factor levels are nominally scaled. Occasionally they are ordinarily scaled. Rarely is there an interval or ratio scale available so that the conjoint analyst can exploit the x_{1n} and x_{2n} vectors associated with the two alternatives in the paired comparison. Neither does he have the luxury in general of a vector c_n of measured characteristics of the respondent in setting n . We are so far left with the view that conjoint analysis deals with the quantal choice setting in which $J = 2$; the x_{1n} and x_{2n} are 0, 1 vectors; and no c_n is available.

We begin to see major distinctions when we contrast how the two approaches view individual respondents and aggregates of responses. In the main, conjoint analysis deals with each respondent—his preferences and utilities—separately. Each person is taken as he comes, and the problem of interpersonal comparison of utility never arises.¹ (The only aggregation across respondents is performed in a post-conjoint analysis “market share simulation,” in which, after each respondent’s

1. The MONANOVA method, though, does allow one to make interpersonal utility comparisons by making the assumption that each respondent’s factor-level utilities are but a random error away from the overall population factor-level utility. MONANOVA can be used to estimate the population factor-level utilities.

factor-level utilities are derived, one calculates from these the utilities of each of a new set of bundles and determines the bundle whose utility is highest for each respondent. The resulting frequency distribution across the set of bundles is an estimate of the market share of each of those bundles, assuming that the respondents are a random sample, that each member of the population is a utility maximizer, and that the additivity assumption holds.)

In quantal choice methods, by contrast, one deals with aggregate data, the observed selection probabilities, and hopes to derive from them the underlying utilities. It is assumed here that variations in utilities are a function of the respondent's attributes c_n and random error and are not fundamentally irreconcilable utilities. In other words, the utility function in quantal choice analysis is a commonly held function, not a separate function for each individual.

This use of the selection probability in economic research stems from the underlying notion of the random utility model. In this model, it is assumed that a respondent selects a utility function "at random" from some set of utility functions, so that we have intrapersonal random variation in utility functions as well as interpersonal random (and nonrandom) variations in utility functions in the extreme form of this model. We can see from this description why one requires replication (both within and between respondents); why one models a decision as not being firm but due to some chance mechanism, the selection probability; and why one tries to parametrize or otherwise structure both the set of utility functions and the process by which a utility is selected "at random."

In conjoint analysis, we can create pseudoselection probabilities, as Johnson (1976), does, as follows. Let u_{n1} and u_{n2} be the utilities that a respondent associates with bundles 1 and 2 of the n th pair presented to him. Assume that these utilities do not determine the respondent's selection (i.e., 1 if $u_{n1} > u_{n2}$, 2 if $u_{n1} < u_{n2}$) but rather determine a probability

$$p_{n1} = \frac{u_{n1}}{u_{n1} + u_{n2}}$$

of selecting bundle 1. (This is very much in the spirit of the Luce model, as described by McFadden [1976]). Johnson then goes on to assume multiplicative utilities and uses as a criterion that

$$\prod_{n=1}^N p_{n1}$$

be maximized. But, leaving details such as this aside, the point to be made here is that in conjoint analysis the "selection probability" is respondent based and not a population parameter.

The work of Block and Marschak (1960) is the only preference-choice model in economics that is free from assumptions about additional data or functional relations between selection probabilities and other variables. Yet that work, too, does not catch the spirit of conjoint analysis as it is used in marketing. For that work also hinges on the underlying notion of the random utility model, but, rather than structure the model, it adopts a nonparametric approach, as follows.

Since the utility functions are unobservable, instead of positing a probability distribution on utility functions, let us create equivalence classes of utility functions, wherein all utility functions leading to identical rankings of all subbundles are in the same class. We now deal with a probability distribution across the equivalence classes, that is, over all possible sets of rankings. Given selection probabilities based on paired comparisons of subbundles, Block and Marschak determine conditions under which these selection probabilities are consistent with any probability distribution across the set of ranked subbundles. But since this work still deals with aggregate selection probabilities, the only use it can have in conjoint analysis is to help us check whether the aggregate of data collected via conjoint analysis is consistent with any probability distribution defined on the set of ranked subbundles. Since there certainly should be such a distribution, detection that aggregate data based on conjoint analysis are inconsistent with such a distribution should lead us to closer inspection of our respondent-by-respondent data.

But consistency checks closer to the spirit of conjoint analysis do exist. First of all, one can calculate Kendall's tau as a measure of concordance between the observed and imputed rankings. (Indeed, one criterion for utility estimation, akin to Manski's score maximization in quantal choice [1975], is to maximize tau.) But we typically do not get values of tau near 1 in conjoint analysis. This is usually attributed to a breakdown in our assumption of a separable utility function. (One could have alternatively checked this assumption by applying some of the results of Krantz and Tversky [1971] to the ranked data.) And so we might add interaction terms and reestimate the utility function, all the time trying to achieve tau equal to 1 for each respondent.

Perhaps the conjoint analysts are making life too hard for themselves. With a random-utility-model viewpoint they might brush off inconsistencies not as misspecification of the utility function but as a manifestation of different, inconsistent randomly drawn utility functions at each choice point. They would then tend only to explain gross trends or predilections in decisions and not each person's specific decision in each choice presented to him.

References

- Block, H., and Marschak, J. 1960. Random orderings and stochastic theories of responses. In I. Olkin (ed.), *Contributions to Probability and Statistics*. Stanford, Calif.: Stanford University Press.
- Green, P. E., and Srinivasan, V. 1978. Conjoint analysis in consumer research: issues and outlook. *Journal of Consumer Research* 5 (September): 102-23.
- Johnson, R. M. 1975. A simple method of pairwise monotone regression. *Psychometrika* 40 (June): 163-68.
- Johnson, R. M. 1976. In B. B. Anderson (ed.), *Advances in Consumer Research*. Vol. 3. Proceedings of Association for Consumer Research, Sixth Annual Conference. Association for Consumer Research.
- Krantz, D. H., and Tversky, A. 1971. Conjoint measurement analysis of composition rules in psychology. *Psychological Review* 78 (March): 151-69.
- Kruskal, J. B. 1965. Analysis of factorial experiments by estimating monotone transformations of the data. *Journal of the Royal Statistical Society*, ser. B, 27:151-69.
- McFadden, D. 1976. Quantal choice analysis: a survey. *Annals of Economic and Social Measurement* 5 (Fall): 363-90.
- Manski, C. F. 1975. Maximum score estimation of the stochastic utility model of choice. *Journal of Econometrics* 3 (August): 205-28.

Copyright of Journal of Business is the property of University of Chicago Press and its content may not be copied or emailed to multiple sites or posted to a listserv without the copyright holder's express written permission. However, users may print, download, or email articles for individual use.