

USING THE MUTHEN MODEL AS A CHOICE MODEL: MONTE CARLO EXPERIMENTS AND APPLICATION TO PANEL DATA

S. PRASAD KANTAMNENI, *Emporia State University*
CHARLES F. HOFACKER, *Florida State University*

Muthen (1978) developed a method for the factor analysis of dichotomous variables. Using information from the first and second order proportions of the data, the method fits a multiple factor model. Despite successful empirical applications, Monte Carlo experiments regarding the robustness of the Muthen Model are limited. The purpose of this paper is twofold: (1) to test the model in a Monte Carlo context and (2) using panel data, to show the effectiveness of the model for the study of consumer choice. Results show the model to perform well in both simulation and application to panel data.

INTRODUCTION

Muthen (1978) proposed a GLS estimator requiring less computing time, to simplify the estimation for the factor analysis of dichotomous variables. While the method of factor analysis of dichotomous variables has been applied to real world data (e.g., Christofferson 1975; Muthen 1978), the model has not been adequately tested in a simulation context. The only simulation study was done by Potthast (1993) who used the categorical variable methodology developed by Muthen (1984), to conduct confirmatory factor analysis of polychotomous variables using samples of 500 and 1000. One implication of the study was that correct models may seem to fit poorly where samples of less than 1000 observations are used. Monte Carlo work was not done on the model with dichotomous data nor was the model applied to panel data. As the Muthen Model is a factor analytic model, questions such as whether a certain number of variables are required based on the number of factors desired and whether the model is robust to "overfactoring" and "underfactoring", also become important.

The purpose of this research is twofold: (1) conduct Monte Carlo experiments using dichotomous data to show the requirements and limitations of the Muthen Model and (2) test model performance as a choice model, when panel data is used.

Muthen's (1987) LISCOMP (Analysis of Linear Structural Relations using a Comprehensive Measurement Model), a computer program, will be used for the analysis. Muthen's LISCOMP extension is a generalization of the LISREL model and as such allows users of the Muthen Model to conduct causal analysis. In addition to estimating the usual parameters of a factor model, marketers using the Muthen Model will also estimate sample thresholds and sample marginal proportions.

The Muthen Model is similar to the Probit Model (a Random Utility Model) in its distributional assumptions but while the Probit Model requires observations to be of the pick 1/n nature, the Muthen Model will be shown as a "Pick Any" model.

LITERATURE REVIEW

Marketers have used data of 1/n (when the respondent is allowed to pick only one alternative from a set of alternatives, the resulting empirical observation is of the pick 1/n nature), k/n (the respondent is instructed to choose exactly k alternatives) or pick any/n nature, for implementing choice models. Using pick 1/n and pick k/n data make restrictions that otherwise do not exist naturally in all choice situations.

In the "pick any" situation, the number of choices made by a subject as well as the set of alternatives from which they are chosen is unconstrained. Observations thus derived empirically are referred to as "pick any" data (Coombs, 1964). In such situations, a non-chosen alternative is considered as either acceptable or nonacceptable, or acceptable and considered but not chosen.

Empirically "pick any" data sets consist of only zeros (for alternatives not chosen) and ones (for alternatives chosen).

Levine (1979) provides a model and scaling method for "pick any" data. Levine observed that the subject, when asked to pick among alternatives, seemed to do so by picking alternatives that were not too different from his/her concept of an "ideal product." Considering that both subjects and alternatives are representable in a joint space as points, when a group of individuals pick the same brand, this brand is located near the centroid of these individuals. If each individual has picked more than one brand, then that individual should be near the centroid of the brands. Levine's algorithm, used for spatial representation of "pick any" choices, solves for both brands and respondents simultaneously by finding coordinates that satisfy criteria for both centroids.

Green and DeSarbo (1981) developed an external unfolding algorithm for handling "pick any" data. The method involves using PREFMAP-2 (an algorithm where a matrix derived from ratings of I respondents on J stimuli yields both a stimulus space and a set of ideal points for respondents) to develop a joint space of stimuli and attributes of stimuli, with the list of attributes prespecified. Once the joint space for each individual is derived, the Levine's pick k/n approach can be used to compute an individual's ideal point. The ideal points will be positioned at the centroid of the brands chosen in the existing perceptual space. DeSarbo and Hoffman (1987) introduced a multidimensional unfolding threshold model for pick any/ n data. The model is stochastic, represents products and respondents as points, and can handle all types of binary choice data. According to the multidimensional unfolding threshold model, an unobservable "disutility" variable D_{ij} is defined, which if less than or equal to d_i - individual i 's threshold value - a choice will be made for product j . The utility function thus used posits choice as a function of an unobservable variable and an error function. The resulting space, in terms of ideal point representation, can be used to understand the competitive market structure and also for product positioning. Where the ideal point model is unfit to represent respondents for whom 'more is better', a vector model can be substituted (DeSarbo & Hoffman, 1987, p. 52).

Elrod (1988) developed the Choice Map Model to simultaneously infer brand positions and the distribution of consumer preferences. Choice Map uses real world data (panel data) that typically represent pick any situations. Choice Map is a stochastic brand choice model. While most brand choice models do not con-

sider the attributes of brands, Choice Map captures information on both, time variant and invariant, attributes of the brands.

Pick Any models are both static (e.g. Levine, 1979) and dynamic (e.g. Elrod, 1988). These models only deal with discrete choice situations. The choice rules for these models are either deterministic (e.g. Levine, 1979) or stochastic (e.g. DeSarbo & Hoffman, 1987).

The Muthen Model

The Muthen Model is a factor analytic procedure that uses dichotomized variables to extract the underlying dimensions of choice. Considering a household's purchase history of a specific category of products, the model attempts to extract the latent or unobservable dimensions that seem to explain household buying behavior. The model is particularly useful when panel data is considered. The extracted factors (dimensions) can represent price, quantity, flavors, colors, packaging, and other marketing elements that help in marketing the product. The Muthen Model assumes a latent variable, a so called response strength, to underlie each dichotomous response. Based on how the entire sample of subjects responds to an item, the model estimates a threshold value (a cutoff based on the strength of belief). Let us consider an example to understand the concept of thresholds. The Cabbage Patch Kids were first introduced to the market at a price of \$9.99 that resulted in rather weak sales. On the advice of consultants, the price was raised to \$39.99 to result in phenomenal sales growth. It was understood that the initial price may have signaled a lower quality to the consumers. Consumers seemed to have used an imaginary cutoff for the price at which they either purchased or did not purchase the dolls.

The Muthen Model is dynamic but only to the extent that it looks at more than one purchase occasion at the end of a certain period. The model is also stochastic in that choice is represented as a result of the unobservable factors and a stochastic component. The model is limited to handling only discrete choice situations. While the DeSarbo and Hoffman (1987) model and the Muthen Model involve the threshold concept, the former deals with a geometric representation (MDS) whereas the latter is based on a factor model.

MODEL DEVELOPMENT AND METHODOLOGY

The discussion in this section is based on the work presented by Christoffersson (1975) and Muthen

(1978). Consider that the vector y (of order $m \times 1$), which represents the utilities of the various brands, allows the usual K - factor model,

$$y = \mu + \Lambda\eta + \delta.$$

where Λ is a $m \times k$ matrix of factor loadings, η is the k by 1 vector of factors, δ is the vector of residuals and μ denotes the expectation for y .

The covariance matrix of y is given by:

$$E[(y - \mu)(y - \mu)'] = \Sigma,$$

and that

$$\Sigma = \Lambda\Psi\Lambda' + \theta,$$

where Ψ is the covariance matrix of factors and θ is the covariance matrix of δ .

Suppose that the observed variable y^* is such that its i th component y_i^* equals one if the y_i is greater than or equal to a number h_i called the threshold, and otherwise zero. That is,

$$y_i^* = \begin{cases} 1 & \text{if } y_i \geq h_i \\ 0 & \text{if } y_i < h_i \end{cases}$$

Here y_i^* denotes the utility for brand i . If it is equal to or greater than h_i , the subject chooses (otherwise rejects) the product.

The model assumes that the h_i , the threshold values for each brand, are constant across the population of consumers. Given this, the distribution of y^* is of the following type

$$g(y_i^* = (1, 0, 1, \dots, 0)) = \int_{h_1}^{\infty} \int_{-\infty}^{h_2} \int_{h_3}^{\infty} \dots \int_{-\infty}^{h_M} f(x) dx,$$

and we further assume

$$f(x) = \frac{1}{\Sigma^{1/2} (2\pi)^{M/2}} e^{-x'\Sigma^{-1}x/2}$$

Note that $f(x)$, in the above equation, is the m dimensional normal density function which is also assumed by the probit model. The distribution of y^* is such that, a maximum likelihood estimation of the parameters would involve a good deal of complexity. An alternative estimation procedure has the prerequisite that such

estimation be based on marginal distributions of single and pairs of items. The marginal proportions P_i^* and P_{ij}^* allow us the following model:

$$P_i = P_i^* + \epsilon_i$$

$$P_{ij} = P_{ij}^* + \epsilon_{ij}$$

or in vector form

$$P = P^* + \epsilon$$

where ϵ has expectation zero and covariance Σ_ϵ . (Christofferson, 1975, shows what Σ_ϵ looks like). Generalized least squares estimation (GLS) can be used to estimate the parameters in the factor model. The GLS technique will minimize,

$$(P^* - P)' S^{-1}_\epsilon (P^* - P),$$

where S_ϵ estimates Σ_ϵ . Muthen's (1978) GLS estimator requires a transformation 'F' such that

$$p = f(\theta) + \epsilon,$$

where $\theta = (\theta_1', \theta_2')$ is an m -dimensional vector with $\theta_1 = h$, and θ_2 denotes the vector of elements below the diagonal of Σ . θ_1 is the vector of population thresholds and θ_2 is the vector of population tetrachoric correlations.

MONTE CARLO EXPERIMENT

Boomsma's (1985) robustness study on maximum likelihood estimation with LISREL addressed problems of nonconvergence, improper solutions and choice of starting values. The results of the Boomsma study are considered relevant here because Muthen's LISCOMP extension is a generalization of the LISREL approach to dichotomous data. Based on covariance structures of multivariate normal distributions, Boomsma generated independent samples ranging from size 25 to 400. Monte Carlo results suggested that the sample size should be at least 100 for LISREL users. The choice of ideal starting values did not in any way influence the Monte Carlo results. The results also show nonconvergence to be more likely if the number of variables to number of factors ratio (NV/NF ratio) decreases: 10 variable two-factor models converge more often than five variable two-factor models. A solution to the problem of nonconvergence required that the sample size be large i.e., $N > 100$.

Sharma, Durvasula and Dillon (1989) conducted Monte Carlo experiments on the behavior of alternative covariance structure estimators. A factorial design with three levels each of kurtosis, skewness and sample sizes was used in the experiments. The results show that in the case of non-normal samples, GLS performed poorly. As the Muthen model has been shown to work with dichotomous data in previous studies, we did not consider testing the model with regard to distributional assumptions.

In order to conduct a Monte Carlo experiment for the Muthen Model, the following steps are necessary. First, consider a set of variables, for example $m = 10$. Then produce Λ , Ψ , and Θ matrices. This gives us:

$$\Sigma = \Lambda\Psi\Lambda' + \Theta.$$

Second, generate samples of normal deviates with $\mu =$ zero and Σ as specified above. Third, dichotomize the normal deviates assuming that $h_i = 0.5$, $i = 1, 2, \dots, 10$ (Note that 0.5 was arbitrarily set). Finally, use the data to fit a factor model. In the second step, a sample of normal deviates are generated to allow the Muthen Model's assumption that the latent variables are multivariate normally distributed. Any robustness testing of the model for non-normal distributions is hence ignored.

Based on the work of Boomsma (1985), only samples that converge in the Monte Carlo experiment will be considered for analysis. The minimum sample size required for the model to work will be tested. In this regard, whether the model works with sample sizes less than 100 will also be tested (Boomsma recommends that N be greater than 100). Sample sizes ranging from 100 to 500 will be used to see how well the model behaves. The root mean square errors will be used to examine the consistency of the model across different sample sizes. As the Muthen Model extracts the factors that underlie the data, both "underfactoring" and "overfactoring" will be tested in the analysis. This is done by extracting a number of factors different from that represented by the covariance matrix. Finally, whether a desired number of factors require a minimum number of variables (NV/NF ratio) will also be tested.

Muthen's LISCOMP (1987) program will be used to conduct the Monte Carlo experiment. Data in the form of 0s and 1s will be generated. Each sample size, for example 250, will be generated 30 times (replications). Exploratory factor analysis will be used to determine the underlying factor structures. The objective, given the data generation with each replication, is to note the

number of times the correct model is rejected. Significance levels of 0.05 and 0.01 will be used to reject the model. An arbitrary threshold value of 0.5 will be considered for all items. The covariance structure that represents the number of factors of interest will also be used as input.

Chi-square tests for significance both at the 5% and the 1% levels will be used to determine the number of factors to be retained in the Monte Carlo experiment. The root mean square error and its behavior as the sample size varies will be examined.

ICE CREAM DATA

To apply the model to a "real world" situation, panel data collected by Information Resources Incorporated (IRI) in its Chicago market area, consisting of household purchases in the category of frozen novelties and packaged ice creams, will be used. The IRI data were acquired in the time period between February 1, 1988, and January 29, 1989. The data basically consisted of purchase histories of households, each of which used its Shoppers Hotline card at least 3 times every 8 weeks and made at least one category purchase during the 52 week period. 300 households were randomly selected from all those satisfying these requirements.

From the IRI data, only information describing the brand will be considered. More or less such information provides the flavor of ice cream purchased on a particular choice occasion by each household. Examination of the data set showed a variety of ice creams. To keep the items at an optimal level for the analysis, all flavors of ice creams were grouped into 10 items (based on ad hoc reasoning), each item representing a flavor or a combination. These are ordered as follows: 1) Vanilla, 2) Chocolate, 3) Strawberry, 4) Neapolitan, 5) Butter Pecan, 6) Fruit and Nuts, 7) Heavenly Hash/Assorted/etc., 8) Danish toffee/Candy/French Toffee/Heath Candy, 9) Mint Flavors/Spearmint/Mint, and 10) all others.

As the ice cream data only indicates whether a flavor was purchased on an occasion or not, the focus will be to understand the latent variables that may underlie such choices. Each household's purchases are represented by the vector ' y^* ' of order 1 by 10 (as there are 10 flavors). Based on how the households have chosen a particular flavor on occasion, the sample thresholds h_i are estimated for all the 10 flavors. The model assumes that such threshold values are constant across the population of households. Suppose that the utility or "preference" ' y ', for each flavor is such that it is equal

to or greater than the estimated threshold value for that flavor, then the response is coded 1 and otherwise 0. The model assumes that the 'y's' are multivariate normally distributed.

As we are fitting a factor model, we assume that:

$$y^* = \Lambda\eta + \delta.$$

The Λ matrix represents the factor loadings. The k by 1 vector η represents the factors. Ψ is the covariance matrix of the factors that shows the correlation between the factors. Finally, to estimate the parameters, the Muthen Model uses the sample marginal proportions P_{i*} and P_{ij*} . These numbers represent the proportion of households that bought a particular flavor on an occasion and the proportion of households, for example, who bought flavor 1 also bought flavor 2, respectively. These proportions show marketers the popularity of flavors.

The marketing manager using the Muthen Model and having access only to limited information, such as the ice cream data, can to some extent analyze the company product against a competitor's. To the extent the factors are interpretable, marketers can benefit from such an analysis.

RESULTS

Monte Carlo Work

Boomsma (1985) recommends sample sizes of at least 100 for LISREL users. However, in order to see if LISCOMP and the Muthen Model can handle sample sizes less than 100, the first simulation was run using a sample size of 50.

To run a Monte Carlo simulation, LISCOMP requires input considerations in a certain manner. To begin, the threshold value was arbitrarily set at 0.5 for all items. A lambda matrix was developed for both one-factor and two-factor models. The respective covariance matrices were also specified for the one-factor and two-factor models. LISCOMP was then used to generate dichotomous data in sample sizes of 50. Assuming 10 variables, LISCOMP was used to produce Ψ and Θ matrices. This gives,

$$\Sigma = \Lambda\Psi\Lambda' + \Theta.$$

Next a sample of 50 normal deviates with $\eta = 0$ and Σ as specified above were generated. The normal deviates were dichotomized assuming that all 10

threshold values were 0.5. The data were then used to fit a factor model. The simulation had 30 replications each for the one-factor and two-factor models. GLS estimation with the full weight matrix approach was used. Exploratory factor analysis was performed on the data. Thus each of the 30 replications included the transformation into dichotomous data, the generation of normal deviates, specification of the covariance matrix based on whether the model had one factor or two factors, specification of the threshold values, GLS estimation with the full weight matrix and exploratory factor analysis for extracting one and two factors. In addition to this, as 10 variables were considered for the IRI data i.e., 10 flavors of ice creams, 10 equal threshold values of 0.5 (set arbitrarily), were considered for the simulations.

The simulations resulted in only one replication converging for the one-factor model and none for the two-factor model. For both models the weight matrix was found not to be positive definite. This precluded any tests for "overfactoring" and "underfactoring". In summary, a sample of 50 observations is not sufficient for the model to work.

The subsequent runs were considered with samples of 100 or more observations. Conditions for the subsequent runs also included conducting "overfactoring" and "underfactoring", change in the sample size, specifically to those of 100, 125, 150, 200, 250, 300, 350, 400, 450 and 500, and use of a covariance matrix for either a one-factor or a two-factor model. As expected "underfactoring" was only possible when Σ stems from two factors and for the two-factor model, the chi-square test for model fit to improve from extracting one factor to extracting two factors. "Underfactoring" was, of course, never done for the one-factor model.

On the other hand, "overfactoring" did not always produce appropriate parameter estimates. Using the sample size of 100, even though LISCOMP was required to generate 30 replications, Heywood cases (outliers with a variance of 20) precluded any simulation. This occurred, obviously, due to "overfactoring" or forcing a solution for more factors than the data actually represent. "Overfactoring" was also done for sample sizes of 125 and 150. As these runs were also unsuccessful, "overfactoring" was ignored in further simulations. The results of the Monte Carlo simulation are presented in graphical form.

The simulations for the one-factor model are shown in Figures 1, 2 and 3.

FIGURE 1
REJECTION RATE FOR THE ONE-FACTOR
MODEL WITH ALPHA = 0.05

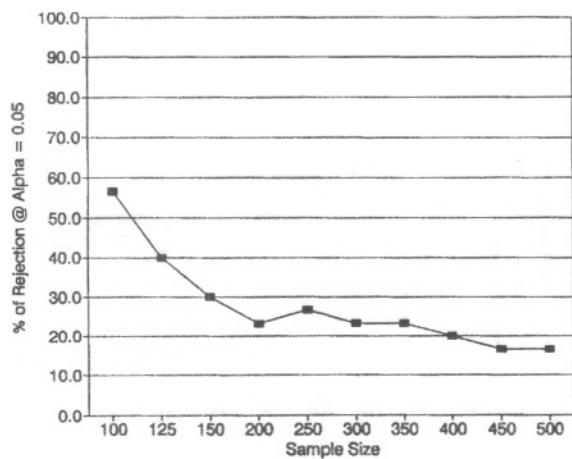
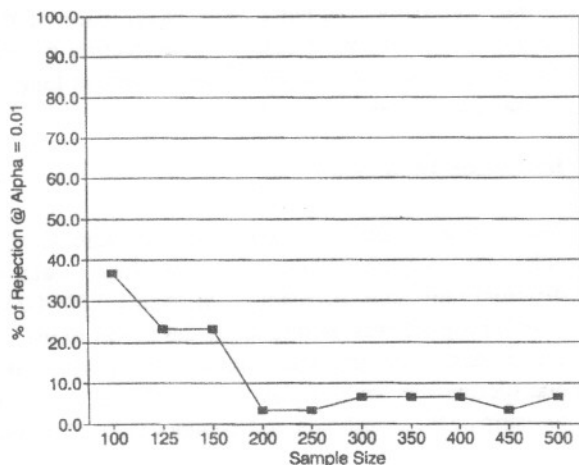


Figure 1 shows the percentage of replications, using a significance level of 0.05, in which the model is rejected. Notice the general drop in rejection rate as the sample size goes up from 100 to 500. This shows that the model behaves consistently across different sample sizes. However the curve does not seem to hit the 5% rejection rate as it is supposed to do. In order to test whether the model behaves better with a higher sample size, another simulation with 1000 observations was run. This brought down the rejection rate to 13.3% an improvement from the 16.66% for the simulation with 500 observations. Whether any increase in sample size would produce better results was not further tested.

FIGURE 2
REJECTION RATE FOR THE ONE-FACTOR
MODEL WITH ALPHA = 0.02



The only difference for Figure 2 is that a significance level of 0.01 was used. Notice from the figure that the rejection rate falls to as low as 3.3% with a sample of 200 observations. Again the curve does not hit the required rejection rate of 1%. Here also a sample of 1000 observations was considered. The results showed the rejection rate to actually increase to 10%. Based on Figure 2, a sample size of 250 observations is recommended for the one-factor model with 10 variables.

FIGURE 3
ROOT MEAN SQUARE ERRORS FOR
THE ONE-FACTOR MODEL

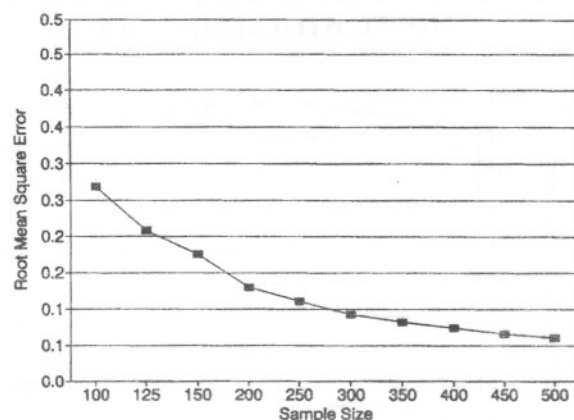


Figure 3 shows the behavior of root mean square error as the sample size increases. The reduction in root mean square error shows consistency of the model. Both Muthen (1978) and Christofferson (1975) do not prove consistency for their estimators.

Figures 4, 5, and 6 show the simulation results for a two-factor model. It is obvious from Figures 4 and 5 that "underfactoring" was possible with all sample sizes but extracting two factors always decreased the rejection rate. Note that a sample of 50 observations considered earlier for a two-factor model resulted in nonconvergence. It is recommended that a minimum sample of 100 observations be used for a two-factor model to insure "underfactoring".

Figure 4 shows the rejection rate using the significance level of 0.05 while Figure 5 shows the same but at the 0.01 level. The proportion of Type I errors is more than triple that of the expected proportions. A simulation with 1000 observations was run for the two-factor model. The rejection rate was 16.6% with a 0.05

significance level but remained the same at 3.3% for the 0.01 significance level.

Figure 6 shows the behavior of the root mean square error as the sample size increases, for both one-factor and two-factor models. The model seems to behave consistently across various sample sizes except for 500 observations where there is a slight increase in the root mean square error. It is not clear as to what might have caused this increase. Moreover, the Type I error is too high at these sample sizes for this model.

FIGURE 4
REJECTION RATE FOR THE TWO-FACTOR
MODEL WITH ALPHA = 0.05

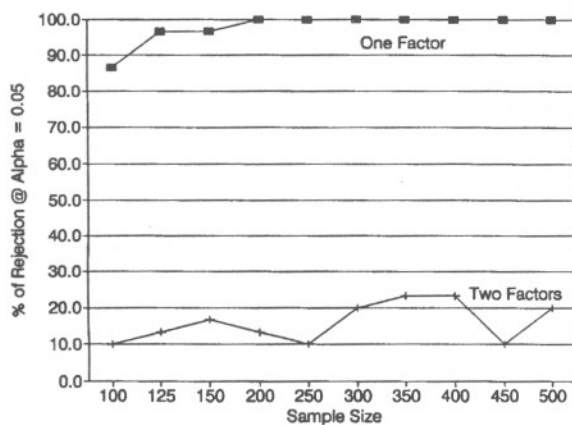


FIGURE 5
REJECTION RATE FOR THE TWO-FACTOR
MODEL WITH ALPHA = 0.01

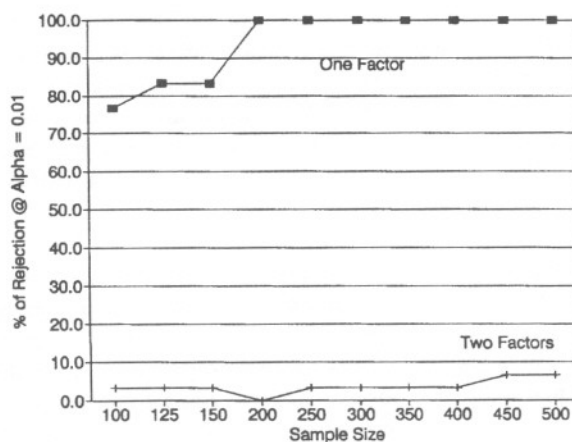
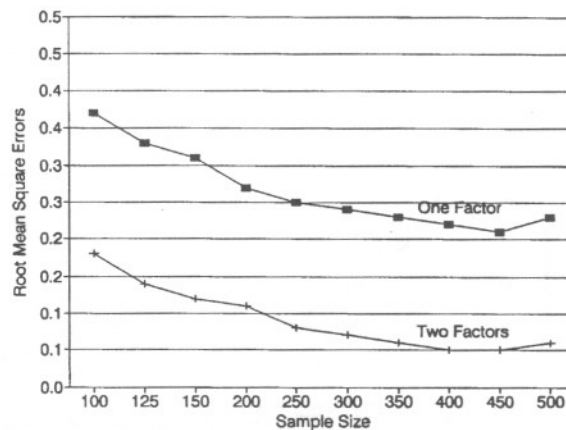


FIGURE 6
ROOT MEAN SQUARE ERRORS FOR
THE TWO-FACTOR MODEL



The figures for the one-factor model show that the model is rejected more times than it ought to be when a 0.05 significance is used. The Muthen Model has been successfully applied elsewhere and has been shown to work in these simulations at the 0.01 significance level. Perhaps the random number generator for the LISCOMP program needs to be tested. The only logical conclusion at this stage is that a good model can sometimes be rejected and it is recommended that the researcher run the model more than once and for different data sets to draw any conclusions.

In summary, the Muthen Model seems to work generally well as the sample size is increased. Note that these samples were run with ten variables only for both one-factor and two-factor models. This was considered appropriate as the ice cream data from IRI was identified with ten flavors and the nature of the data at most would allow two factors to be recovered. "Overfactoring" generally resulted in nonconvergence of the model due to Heywood cases. "Underfactoring" was always possible with the two-factor model as long as the sample size was 100 or more.

Ice Cream Data Work

The ice cream data consisted of 300 households who either chose or rejected the 10 flavors of ice cream on different purchase occasions. Muthen's LISCOMP (1987) program was used to apply the model to the ice cream data. Treating the data as dichotomous and considering 10 variables (flavors), LISCOMP was used to estimate 10 thresholds and conduct an exploratory

factor analysis. The lower and upper limits for the number of factors in the analysis was set at 1 and 5, respectively. GLS estimation was carried out using the full weight matrix approach.

The estimated sample thresholds are as follows:

Flavor	Threshold Value
Vanilla	-.288
Chocolate	.008
Strawberry	.623
Neapolitan	.017
Butter Pecan	6.54
Fruit & Nuts	.761
Heavenly Hash	.253
Danish Toffee	1.527
Mint	1.405
All Others	.994

From the estimations above it can be inferred that vanilla is the most popular flavor as it's estimated threshold value is -.288. This is explained with the following analogy. Just as the normal deviate has a mean of 0, all the ten flavors are also assumed to have means of 0s. However based on the so called response strength of the sample of households, the thresholds (or cutoffs) for each flavor are estimated. If this cutoff is less than 0, the more negative it is the higher is the probability that the households will choose the flavor. When this probability is higher, the flavor is obviously considered more popular. Conversely, the higher the threshold is than 0, the less popular will be the flavor.

Based on this, it can be seen that vanilla, chocolate and neapolitan are the most popular flavors in that order. On the other hand, Danish toffee is the least popular flavor followed by mint (notice that the threshold values are 1.527 and 1.405, respectively). This is very important information for the marketing manager. The thresholds represent how much popular (or desirable) the various flavors are to the households.

Estimated latent roots for the sample correlation matrix are: 3.675, 1.186, 1.066, .967, .784, .600, .598, .469, .383, and .272.

Analysis with one factor produced a chi-square (with 35 df) of 58.819 with a probability of 0.0071. The root

mean square error was .1395. Analysis with two factors produced a chi-square (with 26 df) of 27.043 with a probability of 0.4071. The root mean square error decreased to .0938. The results show a better fit for the two-factor model. The varimax rotated loadings for the two factor solution are as follows:

Flavor	Factor 1	Factor 2
Vanilla	.608	.197
Chocolate	.712	.212
Strawberry	.049	.118
Neapolitan	.430	.438
Butter Pecan	.418	.089
Fruit & Nuts	.437	.342
Heavenly Hash	.618	.155
Danish Toffee	.814	-.026
Mint	.571	.209
All Others	.476	.136

The results show the strawberry and neapolitan flavors to load on Factor 2 while the remainder of the flavors load on Factor 1. Neapolitan ice cream is a combination of vanilla, chocolate and strawberry. As such Factor 2 seems to represent the strawberry flavor. Consequently then Factor 1 should represent all other flavors. This is important information to the marketer in that while all flavors are rather equally desired, the households seemed to prefer anything with strawberry (as neapolitan also has strawberry) differently. Notice from the underlined loadings for the neapolitan flavor that they are almost the same: .430 and .438 on Factors 1 and 2, respectively.

Apparently the factors identified have to do with Factor 1 being "Variety of flavors" and Factor 2 being "Strawberry flavor". Further analysis for 3, 4 and 5 factor solutions did not converge. "Overfactoring" was presumably the reason for such results. The results imply that the battery of 10 items has at least two factors.

How can these results be considered for any kind of managerial actions? Can these results be of any value to the marketer? Consider that consumer orientation in marketing requires providing the consumers with products that they prefer. Whether or not the consumer selected a product on a given occasion becomes very important. In this research it was possible to identify two latent variables that seem to influence such purchase behavior.

Based on the overall fit of the two-factor model, the Muthen Model seems to work well with panel data. The results are certainly interpretable. The analysis is very straight forward. The estimation has taken less than 2 minutes on the IBM PC Model 55 SX. The model seems to work well in both simulations and with real data.

CONCLUSION

While the Muthen Model was used with success in psychological research where attitude measurement was of interest, it was not adequately tested in a simulation context. Also the model was never applied to panel data. Muthen's (1978) research was with data from surveys for testing a specific theory. Panel data, as syndicated data, represents purchase histories for groups of individuals and households for a period of time. The Muthen Model here was applied to such data. Under the conditions, the model has yet allowed us to identify the latent variables that seem to produce such purchase behavior.

The Muthen Model has worked well in simulations and with panel data. Further research, that was particularly inadequate, in robustness testing for the model has been conducted. Results show that the model performs well

with sample sizes of at least 100. The model provides a good fit at the 0.01 level of significance. For a two-factor model, 10 variables seem sufficient to produce meaningful results. The model has been shown to be consistent across different sample sizes.

The model can be a useful tool in developing segmentation and positioning strategies. For example, the marketer will not only know how strongly consumers feel about the strawberry flavor in Chicago but, when considering similar data for other regions, can identify segments based on the differences in the strength of the preferences. Promotions can also benefit from the information on the strength of thresholds in the different regions. This research is limited in that only 300 observations were considered for the study. Using larger samples, researchers can hold some data for refitting the model. As more or less variables are considered, the NV/NF ratios are likely to vary. Monte Carlo results suggest that the random number generator of LISCOMP needs testing as the true models were rejected more than the expected number of times. As variables of interest are the flavors of ice cream, the conclusions made in the study seem relevant. The possibilities of this model for marketing users with well designed surveys are numerous.

REFERENCES

- Boomsma, Anne (1985), "Nonconvergence, Improper Solutions, and Starting Values in Lisrel Maximum Likelihood Estimation," *Psychometrika*, 50, 229-42.
- Christofferson, A (1975), "Factor Analysis of Dichotomized Variables," *Psychometrika*, 40, 5-32.
- Coombs, C. H. (1964), *A Theory of Data*, New York: John Wiley & Sons, Inc.
- DeSarbo, Wayne and Donna L. Hoffman (1987), "Constructing MDS Joint Spaces From Binary Choice Data: A Multidimensional Unfolding Threshold Model for Marketing Research," *Journal of Marketing Research*, 24, 40-54.
- Elrod, T. (1988), "Choice Map: Inferring a Product-Market Map From Panel Data," *Marketing Science*, 7, 21-40.
- Green, Paul E. and Wayne S. DeSarbo (1981), "Two Models for Representing Unrestricted Choice Data," *Association for Consumer Research Proceedings*, 309-12.
- Levine, J. H. (1979), "Joint-Space Analysis of 'Pick-Any' Data: Analysis of Choice From an Unconstrained Set of Alternatives," *Psychometrika*, 44, 85-92.
- Massy, William F., D. B. Montgomery and Donald G. Morrison (1970), *Stochastic Models of Buying Behavior*. Cambridge, MA: MIT Press.
- Muthen, Bengt (1978), "Contributions to Factor Analysis of Dichotomous Variables," *Psychometrika*, 43, 551-60.
- (1984), "A General Structural Equation Model With Dichotomous, Ordered Categorical, and Continuous Latent Variable Indicators," *Psychometrika*, 49, 115-132.
- Potthast, Margaret J. (1993), "Confirmatory Factor Analysis of Ordered Categorical Variables With Large Samples," *British Journal of Mathematical and Statistical Psychology*, 46, 273-286.
- Sharma, Shubash, Srinivas Durvasula and William R. Dillon (1989), "Some Results on the Behavior of Alternate Covariance Structure Estimation Procedures in the Presence of Non-Normal Data," *Journal of Marketing Research*, 214-21.

ABOUT THE AUTHORS

S. Prasad Kantamneni (Ph.D., Florida State University) is Assistant Professor of Marketing, Emporia State University.

Charles F. Hofacker (Ph.D., UCLA) is Associate Professor of Marketing at Florida State University.