

## APPLICATION OF THE "DEFENDER" CONSUMER MODEL

JOHN R. HAUSER AND STEVEN P. GASKIN

*Massachusetts Institute of Technology  
Management Decision Systems Inc.*

This paper examines the feasibility, practicality, and predictive ability of the consumer model which was proposed by Hauser and Shugan (1983). We report results in two product categories, each representing over \$100 million in annual sales. We develop "per dollar" perceptual maps and empirical consumer "taste" distributions. As a first test of the model, we compare the predictive ability of the consumer model in one category to (1) pretest market laboratory measurement models, (2) traditional perceptual mapping procedures, (3) a hybrid model using price as an attribute, and (4) actual market shares in test market cities. In the second product category, we illustrate the application of the quantitative model to augment managerial judgment. Besides developing an empirical version of the 'Defender' consumer model, our analyses raise a number of behavioral hypotheses worth further investigation.  
**(Defensive Strategy; Preference Measurement; Consumer Response)**

### 1. Motivation

In their paper, "Defensive Marketing Strategy", Hauser and Shugan (1983) develop a number of qualitative normative implications on how a firm marketing an established product should defend its profit when facing an attack by a new competitive product. For example, their analyses suggest decreasing budgets for distribution and awareness-advertising while improving the product and repositioning in the direction of the defending product's strength; price should be decreased in unsegmented markets but potentially increased in segmented markets. These implications are derived from a mathematical model of consumer response that assumes heterogeneous consumers maximizing utility in a "per dollar" multi-attributed space. We call this descriptive mathematical model the 'Defender' consumer model.

A key feature of the Defender consumer model is that attributes are measured "per dollar"; for example, laundry detergents might be evaluated with respect to 'efficacy per dollar' and 'mildness per dollar'. However, the "per dollar" assumption is untested and, hence, has become quite controversial in marketing science. See discussions in Rao (1984), Ratchford (1982), Sen (1982), and Gavish, Horsky and Srikanth (1983).

Another key feature is that consumer "tastes", i.e., tradeoffs among the attributes, are assumed heterogeneous. The distribution of these tastes across consumers is estimated by a sum of piecewise uniform distributions. Like "per dollar" maps, this assumption has not been tested empirically prior to this paper.

While Hauser and Shugan's qualitative results require only the existence of "per dollar" perceptual maps and heterogeneous taste distributions, quantitative defensive

strategy results do depend upon our ability to develop adequate empirical representations of the theoretical constructs. Furthermore, an examination of "per dollar" perceptual maps and heterogeneous taste distributions has implications beyond defensive marketing strategy. For example, Lane (1980), Lancaster (1980), and Hauser and Simmie (1981) each assume the existence of "per dollar" perceptual maps in their analyses. Similarly, the debate on the need for heterogeneous preferences is a long standing debate in marketing science. (E.g., see Green and Srinivasan 1978.)

This paper describes in detail an initial application of the Defender consumer model. In particular, we estimate empirically a "per dollar" perceptual map and the corresponding taste distribution. We compare the estimated consumer model to a variety of alternative models and to "actual" consumer behavior. While recognizing that a single empirical test is not sufficient to accept a basic model, we feel that such a test is sufficient to demonstrate that the consumer model is reasonable and worth further exploration.

We also describe a second application in which market research data, managerial judgment, and the Defender consumer model are combined to obtain strategic insight about a new product and a potential defense via the launch of a flanking product.

## 2. Perspective

The primary purpose of this paper is to test the consumer model. If the model is not feasible, if it is onerous to measure and estimate, or if it predicts poorly in our initial test, then we must reexamine its basic assumptions. If the consumer model is feasible, reasonably cost effective, and reasonably accurate in the initial test, then we can proceed to develop the model further and explore its quantitative implications. In either case, we advance our understanding of how to model consumer response.

To apply the consumer model, we follow procedures developed by Hauser and Shugan (1983), making a minor modification to handle single-product evoked sets. (This modification is justified theoretically *prior* to parameter estimation.) We estimate the model based on products which are in the market *prior* to the new entrant. We compare predictions to (1) predictions of two well-documented marketing science models, (2) a hybrid model, and (3) measures of market share taken *after* the new product had entered the test market. The first comparison enables us to understand better alternative marketing science models and assumptions. The second comparison explores the implications of "per dollar" attributes vs. price as an attribute. The third comparison is an initial test of the model's external validity.

For established model comparison, we chose (1) Silk and Urban's (1978) 'Assessor' model and (2) traditional perceptual mapping/preference regression models as described by Urban and Hauser (1980, Chapters 9 and 10).

Assessor has been applied commercially to over 400 products and its predictive accuracy has been examined scientifically by Urban and Katz (1983). Furthermore, it is representative of a class of commercially available 'pretest market models' such as those described in Eskin and Malec (1976), Tauber (1977), Burger, Gundee and Lavidge (1981), Yankelovich, Skelley, and White (1981), and Pringle, Wilson and Brody (1982). Data collection is clearly feasible for models in this class and predictive accuracy is acceptable to many marketing managers. If Defender can predict as well as Assessor, then we have confidence in its predictive ability.

Analysis with traditional perceptual maps is state-of-the-art methodology as recommended by new product development textbooks (Urban and Hauser 1980, Chapters 9, 10; Pressemer 1982, Chapter 5; Wind 1982, Chapter 4; Choffray and Lilien 1980, Chapter 6) as well as basic marketing textbooks (Kotler 1983, Chapter 2). Traditional maps serve as a standard to which "per dollar" perceptual maps can be compared.

Data collection is feasible for traditional maps although their predictive accuracy is not documented as well as for pretest market models.

The hybrid model is included in recognition of the debate in marketing science (e.g., Rao and Gautschi 1982 and Srinivasan 1982) on whether or not price should be treated as an attribute. Since the Defender consumer model differs from most traditional perceptual maps in two ways, "per dollar" scaling and heterogeneous tastes, a comparison of Defender to traditional methods is a simultaneous test of both assumptions. By using heterogeneous tastes with a traditional perceptual map using price as an attribute, we can explore the differential impacts of the two assumptions.

We choose as our measure of predictive accuracy, the share of the new product as measured by SAMI for the test cities. SAMI is based on a reasonably complete universe of warehouse withdrawals and is the market share measure in which the cooperating firm has the most confidence. While no measure is perfect, we feel SAMI is a reasonable benchmark to which to compare Defender's predictions. (For a discussion of other methods to track market share such as Nielson audits and diary panels, see Wind and Lerner 1979.) Remember that this application of Defender is based on data collected prior to test market and, hence, prior to the SAMI measure.

Finally, because this paper is empirical, a number of tradeoffs (sample size, methods of measurement, model estimation, comparison models, measures of external validity, etc.) were made to construct the empirical realization of the theoretical model. We made one set of judgments. Other researchers with different goals and philosophies might have made different empirical judgments. Thus, when we make interpretations of our data we provide as much evidence as confidentiality constraints will allow in the hopes that each reader can make his (her) own interpretations of our analyses.

### 3. Review of the 'Defender' Consumer Model

For ease of exposition, we present the model for two perceptual dimensions and then illustrate its extension to three or more dimensions. (Hauser and Shugan 1983 deal with two dimensions.) This section presents a verbal description of the model. For the interested reader, an appendix contains analytic formulae.

#### *Evoked Set Issues*

As documented in Silk and Urban (1978) consumers vary in the brands that they consider when making a choice. We call a set of brands an evoked set and we observe which consumers use which evoked sets in making a choice. For example, if there are four "Gypsy Moth Tape" products on the market,<sup>1</sup> called Pro-Strip, Cata-Kill, Tree Guard, and Store Brand, then one evoked set is all four products. Another evoked set might be {Pro-Strip, Tree Guard} and another {Tree Guard, Cata-Kill}. With four products there are 15 possible evoked sets.

Defender analyzes each evoked set separately and predicts for each evoked set the share among brands in that evoked set. Brand share in the category is then logically the weighted average across evoked sets where the weights are the fraction of consumers who use that evoked set as their choice set.

Note that because of variation in evoking, a brand may be dominated by another brand yet be predicted by the model to have nonzero category share. Predicted purchases result in the model from purchases by consumers who evoke the dominated brand but not the dominating brand. There are many reasons why they may not evoke the dominating brand; for example, it may be underadvertised and they are not aware

<sup>1</sup>Gypsy Moth Tape is a specialized product used in the forests of New England to protect against insect infestation. Throughout the paper, we use it as a disguise for a real \$100 million category.

of it, or it may not be available where they shop. (Of course, there are other reasons why a brand modeled as dominant might be chosen in the marketplace including omitted variables and model misspecification.)

### Share Within Evoked Sets

Defender assumes (1) that each consumer chooses from his evoked set the product which maximizes utility, (2) that utility is linear in the "per dollar" perceptual dimensions, and (3) that consumers vary in their tastes. Linear utility implies straight-line indifference curves in perceptual space.

For example, suppose that consumer 1 cares only about 'Effective Control'/\$, then his indifference curve will be a vertical line and he will choose Pro-Strip as indicated in Figure 1a. If consumer 2 cares only about 'Ease of Use'/\$, his indifference curve will be a horizontal line and he will choose Tree-Guard as indicated in Figure 1b. Finally, if consumer 3 cares equally about 'Effective Control'/\$ and 'Ease of Use'/\$, his indifference curve will make an angle of  $45^\circ$  with the vertical axis and he will choose Cata-Kill as indicated in Figure 1c.

Clearly, consumers can vary in their tastes, that is, in their willingness to trade off 'Ease of Use' for 'Effective Control'. For two perceptual dimensions, we can represent each consumer by the angle,  $\alpha$ , that his indifference curve makes with the vertical axis.<sup>2</sup> See Figure 1d.

The market share of the brand, say Cata-Kill, will be the percent of consumers whose taste-angles,  $\alpha$ , favor Cata-Kill. Thus, if we know the distribution of  $\alpha$  within the population of consumers who use the evoked set as a choice set, and if we know the perceptual positions of all brands in the evoked set, we can readily compute Cata-Kill's market share. The computation is represented in Figure 2.  $f(\alpha)$  represents the distribution of tastes,  $\alpha$ . All consumers with  $\alpha$ 's between  $\alpha_1$  and  $\alpha_2$  will choose Cata-Kill, hence, the market share will be the shaded area in Figure 2.

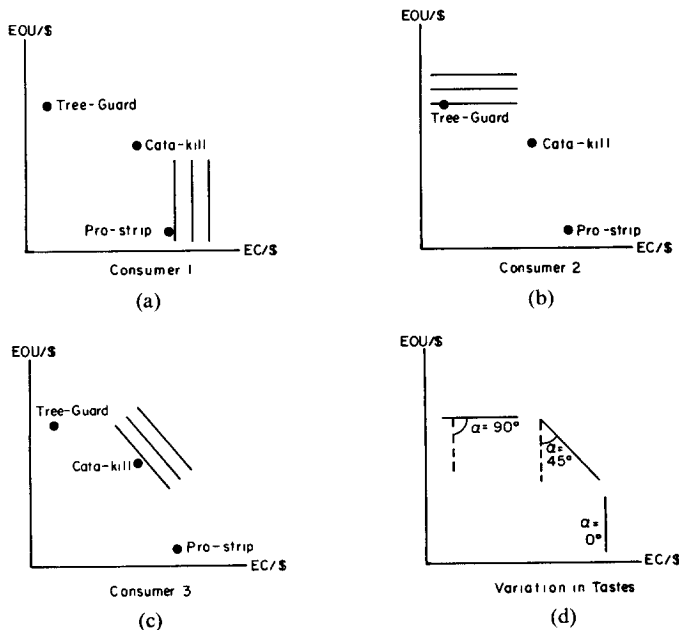


FIGURE 1. Illustration of How Taste Variation Affects Choice.

<sup>2</sup>Technically,  $\alpha$  represents tradeoffs among 'Effective Control' and 'Ease of Use' as well as 'Effective Control'/\$ and 'Ease of Use'/\$ since the 'per dollar' scaling cancels out in computing the tradeoffs.

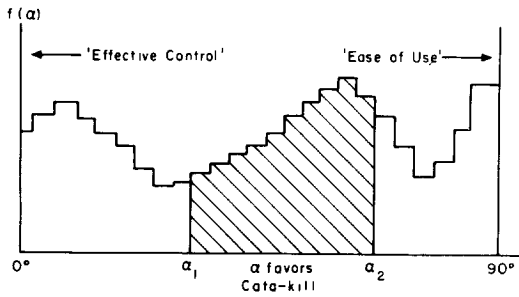


FIGURE 2. Hypothetical Histogram of Consumer Tastes. (Shaded Area Represents Market Share of Cata-Kill.)

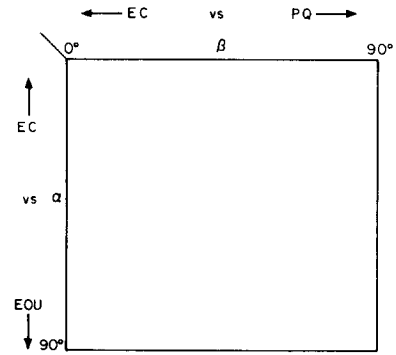


FIGURE 3. Interpretation of  $\alpha$  and  $\beta$  Taste Parameters.

Hauser and Shugan (1983) derive analytic formulae for two dimensions, but for three or more dimensions the formulae become extremely complex. Empirically, we use numerical methods rather than analytic formulae. We divide the range of  $\alpha$  into equal line segments, say  $1^\circ$  each. For each line segment,  $i$ , we use the  $\alpha_i$  for that segment to compute the utility of each product. I.e.,

$$\text{utility of product } j = \left[ \begin{array}{c} \text{'effectiveness'} \\ \text{per dollar} \\ \text{of product } j \end{array} \right] + (\tan \alpha_i) \left[ \begin{array}{c} \text{'ease of use'} \\ \text{per dollar} \\ \text{of product } j \end{array} \right]. \quad (1)$$

The product with the maximum utility, for  $\alpha_i$ , is assigned to the  $i$ th segment and  $f(\alpha_i)$  tells us how many consumers are in that segment. Summing the  $f(\alpha_i)$  across those line segments,  $i$ , where product  $j$  is the maximum utility product, gives us an estimate of the market share of product  $j$ .<sup>3</sup>

This procedure is easily automated and readily generalizes to three or more dimensions. For example, for three dimensions, indifference curves become planes and we require two angles,  $\alpha$  and  $\beta$ , to represent each consumer. The angle  $\alpha$  still represents tradeoffs among 'Ease of Use' and 'Effective Control', while the angle  $\beta$  represents tradeoffs among 'Professional Quality' and 'Effective Control'. See Figure 3. We could also define an angle  $\gamma$  to represent tradeoffs among 'Professional Quality' and 'Ease of Use' but  $\gamma$  is uniquely determined by  $\alpha$  and  $\beta$  and, therefore, redundant. (In particular,  $\tan \gamma = \tan \beta / \tan \alpha$ .)

In three dimensions we divide the  $\alpha - \beta$  feasible region (shown in Figure 3) into equal areas. Empirically, we have found 441 areas work well, that is, a  $21 \times 21$  grid. We extend equation (1) to compute utility:

$$\text{utility of product } j = \left[ \begin{array}{c} \text{'effectiveness'} \\ \text{per dollar of} \\ \text{product } j \end{array} \right] + (\tan \alpha_i) \left[ \begin{array}{c} \text{'Ease of Use'} \\ \text{per dollar} \\ \text{of product } j \end{array} \right] + (\tan \beta_i) \left[ \begin{array}{c} \text{'Professional} \\ \text{Quality'} per} \\ \text{dollar of product } j \end{array} \right]. \quad (2)$$

We assign the maximum utility product to the area corresponding to  $(\alpha_i, \beta_i)$ . The taste distribution,  $f(\alpha_i, \beta_i)$ , tells us how many consumers have tastes represented by  $(\alpha_i, \beta_i)$ . Market share is obtained by summing  $f(\alpha_i, \beta_i)$  across the areas where product  $j$  is the maximum utility product.

<sup>3</sup>As the number of segments goes to infinity, this procedure converges to the Reimann integral of  $f(\alpha)$  and hence to the analytic formula in Hauser and Shugan (1983).

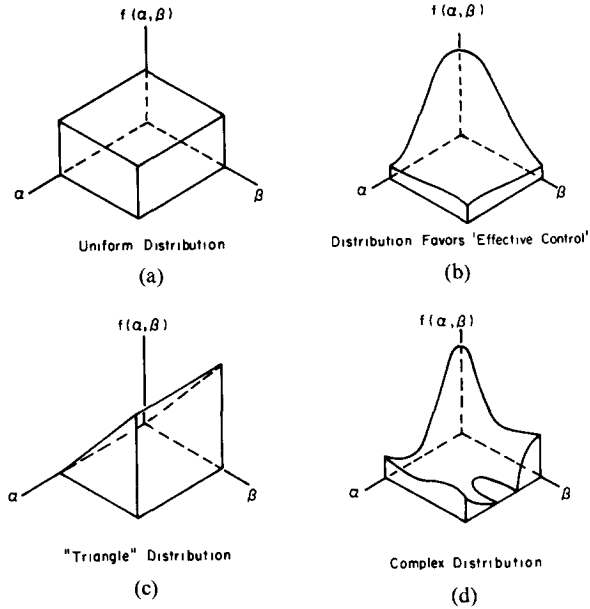


FIGURE 4. Some Alternative Taste Distributions.

Figure 4 provides a visualization of different taste distributions. Figure 4a is a uniform distribution in which all possible taste tradeoffs are equally likely. Figure 4b is a distribution favoring 'Effective Control' over both 'Ease of Use' and 'Professional Quality' while Figure 4c favors 'Professional Quality' over 'Ease of Use' and 'Effective Control' but assumes all possible tradeoffs among 'Effective Control' and 'Ease of Use' are equally likely. Figure 4d is an example of a more complex multimodal distribution.

*Estimation of Taste Distribution*

Equations (1) and (2) are quite simple to use. Market share is readily obtained if we know  $f(\alpha_i)$  or  $f(\alpha, \beta)$  for each area  $i$ . We estimate  $f(\alpha)$  or  $f(\alpha, \beta)$  by adjusting piecewise uniform distributions to fit existing market shares within evoked sets, then summing across evoked sets.

For example, consider a three product, two-dimensional perceptual map such as shown in Figure 1. As drawn, those consumers with angles between  $0^\circ$  and  $30^\circ$  will choose Pro-Strip, consumers with angles between  $30^\circ$  and  $60^\circ$  will choose Cata-Kill, and consumers with angles between  $60^\circ$  and  $90^\circ$  will choose Tree-Guard. If we know that the market shares of Pro-Strip, Cata-Kill, and Tree-Guard are 20%, 50%, and 30%, respectively, then the piecewise uniform approximation of  $f(\alpha)$  shown in Figure 5 will reproduce the existing market shares exactly.

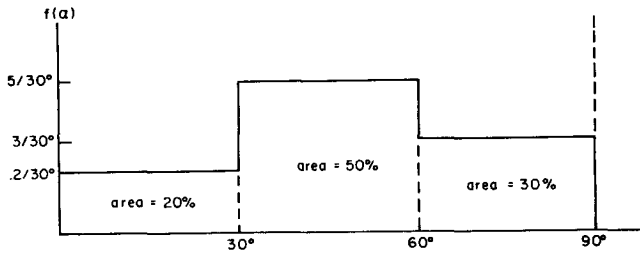


FIGURE 5. Piecewise Uniform Approximation for a Three Product Market.

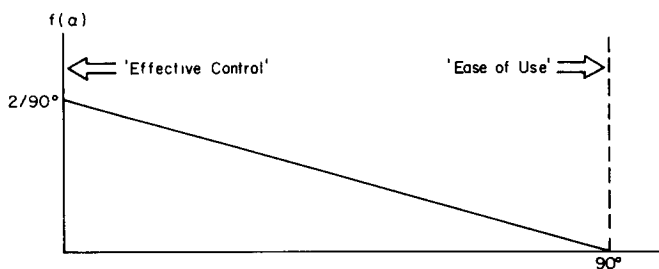


FIGURE 6. Triangle Distribution Favoring ‘Effective Control’ over ‘Ease of Use’.

At first glance, the piecewise uniform approximation appears to be a very coarse approximation, but when we sum these coarse approximations across evoked sets, the taste distribution quickly approaches a more continuous shape as illustrated in Figure 2. In the limit, as the number of products gets large, this procedure clearly converges to the true taste distribution. We do not yet have any analytical expression for finite sample approximation error, but such error does decrease as the number of products increases and we have no reason to believe that the error is biased.

This procedure extends readily to three or more dimensions. For each evoked set we divide the  $\alpha - \beta$  feasible region into mutually exclusive and collectively exhaustive regions according to which product is maximum utility for each  $(\alpha_i, \beta_i)$ . These areas are then adjusted upward or downward to match market shares. Finally, summing across evoked sets gives reasonable approximations to taste distributions such as shown in Figure 4.

The above describes a reasonable empirical implementation of the Hauser–Shugan analytic formula. However, we have found one modification useful. Suppose that an evoked set consists of but one product, say Pro-Strip, and that that product is very effective but not very easy to use. The Hauser–Shugan procedure would assign a uniform distribution from  $0^\circ$  to  $90^\circ$  to that evoked set. Logically, we would expect consumers who evoke only Pro-Strip to care more about ‘Effectiveness’ than ‘Ease of Use’. Thus, we would expect the  $f(\alpha)$  corresponding to that evoked set to favor ‘Effectiveness’ over ‘Ease of Use’. There are, of course, an infinity of  $f(\alpha)$ ’s to choose from so that ‘Effectiveness’ is favored relative to ‘Ease of Use’. We choose what we believe is the most parsimonious distribution that is consistent with the logical requirement that it favor one dimension over another. In particular, we choose a triangle distribution such as shown in Figure 6.

This procedure is extendable to three or more dimensions. For example, the “triangle” distribution in Figure 4c favors ‘Professional Quality’ and would be used for a product which is very good on ‘Professional Quality’ but not particularly strong on other dimensions.

#### *Forecasting of Evoking*

Hauser and Shugan (1983, p. 349) provide two formulae for forecasting the number of consumers who will be in each evoked set after the new product enters the market. We chose the simpler of the two formulae for our applications. We assume that (1) if a consumer evokes an existing brand before Attack enters the market, he will continue to evoke that brand after Attack is launched, (2) that the probability Attack will be evoked is independent of pre-Attack evoking, and (3) this probability is equal to an advertising index times the distribution index. (One can think of the advertising index as ‘Awareness’ and the distribution index as ‘Availability’, but the model is not limited to these interpretations. What is important is that the product of the advertising and distribution indices gives a reasonable estimate of the probability that the new product is evoked.)

### *Forecasting for a New Product*

To forecast the market share for a new product we first compute "Unadjusted Share", that is, the share that the new product would obtain if everyone evoked it. We do this by first placing the new product on the perceptual map and computing for every  $\alpha_i$  (or  $\alpha_i - \beta_i$  combination) the utility of the new product. For each evoked set we identify that region of the  $\alpha$ -line (or  $\alpha - \beta$  plane) that the new product captures and obtain its adjusted share by adding up  $f(\alpha_i)$  for all  $\alpha_i$  (or  $f(\alpha_i, \beta_i)$  for all  $\alpha_i - \beta_i$ ) that the new product captures. The revised, unadjusted market shares of the existing products are computed analogously by identifying the region of the  $\alpha$ -line (or  $\alpha - \beta$  plane) that they retain within each evoked set.

Note that we use  $f(\alpha)$  (or  $f(\alpha, \beta)$ ) estimated on existing brands only. We assume that the taste distribution is not affected by the introduction of the new product. This is the same assumption implicit in conjoint analysis (Green and Srinivasan 1978), preference regression (Urban and Hauser 1980, Chapter 10), and logit analysis (McFadden 1980).

The actual market share for the new product must recognize that it will not be evoked by everyone, but rather will be evoked in proportion to the advertising and distribution indices. In particular,

$$\text{market share} = (\text{advertising index}) \times (\text{distribution index}) \times (\text{unadjusted share}). \quad (3)$$

For existing products, we recognize that only a fraction evoke the new product and that the remainder of the market is unaffected. In particular, for existing products,

$$\begin{aligned} \text{market share} = & (\text{adv. index}) \times (\text{dist. index}) \times (\text{unadjusted, revised share}) \\ & + [1 - (\text{adv. index})(\text{dist. index})] \times (\text{prior market share}). \end{aligned} \quad (4)$$

### *Forecasting for a Change in Strategy of an Existing Product*

We forecast for existing products in an analogous manner. A change in price or repositioning affects the existing product's position on the map. Based on this new position we recompute utility for every  $\alpha_i$  (or  $\alpha_i$  and  $\beta_i$ ) and use the taste distribution to recompute share within each evoked set. Weighting across evoked sets gives us the new share.

A change in advertising or distribution spending causes more evoking as modeled by the advertising or distribution indices. Based on the new indices we recompute evoking and proceed as above.

### *Controversial Measurement Issues*

Two unresolved measurement issues in the Defender model are the estimation of the consumer taste distribution and the use of "per dollar" perceptual maps.

While Hauser and Shugan (1983) propose the piecewise uniform estimation procedure, it has never before been applied to real data. §5 describes the first application. Since the estimation procedure per se does not depend on "per dollar" perceptual maps we estimate a model using price as a fourth dimension as well as the 'Defender' model which uses "per dollar" perceptual maps.

The more controversial issue is the "per dollar" perceptual map. Factor scores are, at best, interval scaled dimensions. To obtain a "per dollar" perceptual map, we divide the measure of a product's perceptual position by the product's price. However, division assumes that the perceptual dimension is a ratio-scaled measure and that a zero-point, e.g., zero 'Effective Control', can be identified. The existence of a zero-point does not imply that a product will exist with zero 'Effective Control'; after all, even a 1972 Cadillac did not get zero 'miles per gallon' yet 'miles per gallon' is a ratio

scale. A "per dollar" ratio scale requires that positions of real products can be measured relative to some reference and that the consumer's willingness to pay for an improved brand can be measured relative to that reference point.

However, even if a zero-point is identified, there is no assurance that the resulting "ratio-ized" scale will provide a reasonable description of consumer behavior. In fact, we may find that no usable zero-point exists. For the purposes of this paper, we treat this issue as an empirical question and attempt to find a usable zero-point. §6 addresses this issue empirically and, to some extent, theoretically.

#### **4. Data for This Application**

We limit ourselves to the data collected by documented market research procedures. In particular we use data collected for Assessor (Silk and Urban 1978) supplemented with attribute ratings as collected for traditional perceptual maps, e.g., Urban and Hauser (1980, Chapter 9). We observe price directly in the market place. If we can develop a reasonable taste distribution and a feasible "per dollar" perceptual map with these standard data, then a careful, evolutionary, Defender-specific improvement of data collection procedures should be feasible, reasonable, and better.

For our initial applications, we chose two categories in which variety seeking, complicated package size issues, and nonmonotonic attributes do not play a major role. We found many categories satisfying these constraints although we recognize that such issues may need to be faced in other categories. Each category is sold through grocery stores and related retail outlets. Because the firm's defensive strategies derive in part from the Assessor and Defender analyses, we have agreed to disguise the data for publication. The disguising procedure and the measured constructs are described below.

##### *Disguising Procedure*

All comparison statistics such as predictive error are reported without modification. Market share figures are rounded off to the nearest share point. Perceptual dimensions are reported to one significant digit after the decimal point and prices are reported as ratios relative to the lowest price product.

For expositional purposes, we have renamed the first product category, "Gypsy Moth Tape", a product used in the forests of New England to combat insect infestation. We have renamed the attribute dimensions, 'Effective (insect) Control', 'Ease of Use', and 'Professional Quality'. These dimensions make sense for Gypsy Moth Tape and capture the flavor of the disguised category's dimensions. We have renamed the new product, "Attack", and the three dominant defending products, "Pro-Strip", "Cata-Kill", and "Tree-Guard", respectively. "Store Brand" represents private label and generic products. To the best of our knowledge, these names are fictitious, but show some relationship to the perceptual dimensions and products in the disguised category. For the second product category, we simply labeled the products A, B, and C, and perceptual dimensions 1 and 2. (Details on this category are given in §8.)

##### *Data Collected*

The details of Assessor and perceptual mapping data collection are contained in Silk and Urban (1978) and Urban and Hauser (1980, Chapters 9 and 10), respectively. For Defender, we use the following data:

(1) Attribute ratings are obtained on semantic scales for each product in each consumer's evoked set. The evoked set is those products which the consumer has used, has on hand at home, or would seriously consider using. (These scales are not necessarily ratio scales. See §6 for further discussion.)

TABLE 1  
*Disguised Factor Analysis for "Gypsy Moth Tapes"*

|                                  | Factor Loading |
|----------------------------------|----------------|
| 1. <u>Effective Control</u>      |                |
| Protects the trees               | 0.75           |
| Stays on trees                   | 0.71           |
| Strong/Durable                   | 0.54           |
| Good in wet weather              | 0.66           |
| Complete barrier                 | 0.68           |
| Keeps trees healthy              | 0.76           |
| Keeps trees from dying           | 0.77           |
| Clings to trees                  | 0.73           |
| Stays tightly on trees           | 0.66           |
| Stops GM caterpillars            | 0.68           |
| Trees retain leaves              | 0.70           |
| Stays effective all season       | 0.70           |
| (Transparent on trees)           | 0.49           |
| 2. <u>Ease of Use</u>            |                |
| Easy to handle                   | 0.59           |
| Comes in handy dispenser         | 0.52           |
| Comes off roll easily            | 0.72           |
| Does not stick to itself         | 0.59           |
| Easy to find start               | 0.78           |
| Easy to open box                 | 0.66           |
| 3. <u>Professional Quality</u>   |                |
| Difficult to rip                 | 0.52           |
| Made from top quality materials  | 0.59           |
| Weatherproof                     | 0.59           |
| Never fails to stop caterpillars | 0.68           |
| (Can be reused)                  | 0.47           |

(2) Attribute ratings, by consumer, for the new brand are obtained after the consumer has been exposed to the brand.

(3) For each consumer, brand last purchased is recorded. And,

(4) Unit price is observed in the pretest market cities.

The sample size for the "Gypsy Moth Tape" category was 297. Samples were drawn randomly via mall intercept within two pretest market cities as per standard Assessor procedure. See Silk and Urban (1978) and Urban and Katz (1983) for details and discussion of sampling variance.

The attribute ratings are factor analyzed as described in Urban and Hauser (1980, Chapter 9) and Hauser and Koppelman (1979). In the "Gypsy Moth Tape" category, the best solution was three-dimensions explaining 92.6% of the common variance (61% of the total variance). Table 1 is a disguised version of that factor analysis. The factor loadings themselves are not disguised but the names of the dimensions are changed to reflect the fictitious "Gypsy Moth Tape" category.

For Defender, we require only the factor scores for each product as averaged across consumers.<sup>4</sup> For the "Gypsy Moth Tape" category, the standard deviations of the mean scores of major brands varied from 0.04 to 0.09 which is small compared to the range (-0.38 to +0.28) of factor scores used in the perceptual map. Figure 7 is the resulting perceptual map for major "Gypsy Moth Tape" brands.

<sup>4</sup>This version of Defender assumes homogeneous perceptions and heterogeneous tastes. Empirically, this assumption is reasonable, in part, because some of the heterogeneity in perceptions is picked up through heterogeneity in preferences. See discussions in Lancaster (1971), Hauser and Simmie (1981), and §9 of this paper. Research is under way to relax this assumption.

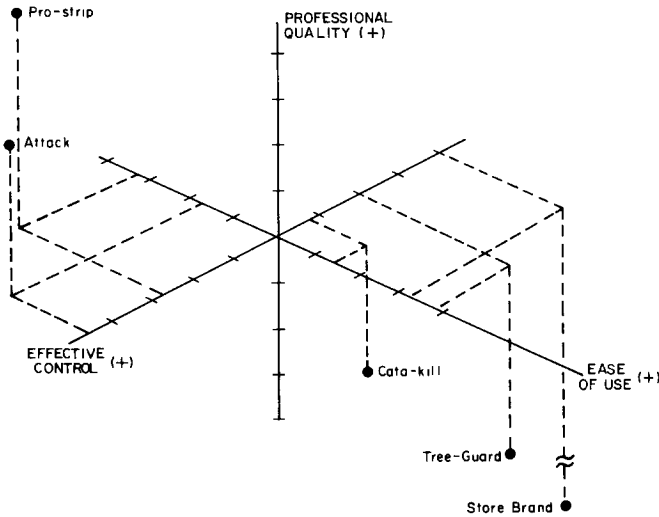


FIGURE 7. Perceptual Map for Category (Disguised).

For comparison to traditional perceptual maps, we used preference regression with constant sum paired comparison preference measures for all pairs of brands in each consumer's evoked set as the dependent measures. Following standard procedure, the explanatory variables were the factor scores representing each consumer's perceptions. See Urban and Hauser (1980, Chapter 10). The importance weights were 0.48, 0.38, and 0.14 for 'Effective Control', 'Ease of Use', and 'Professional Quality', respectively. Forecasting procedures are explained in the next section.

For comparison to Assessor, we recorded the market shares and awareness and availability forecasts as contained in the final Assessor reports provided to the firm. Forecasting procedures are detailed in Silk and Urban (1978).

### 5. Estimation of Taste Distribution

Assume for a moment that a zero-point for the perceptual map has been identified at 'Effective Control' =  $-0.3$ , 'Ease of Use' =  $-0.2$ , and 'Professional Quality' =  $-0.4$ . (Details are given in the next section.) This zero-point assures, as least, that all brands have positive perceptual scores. (Positive scores are necessary to assure that market share is declining in price for taste angles,  $\alpha$  and  $\beta$ , between  $0^\circ$  and  $90^\circ$ .) We also rescale all perceptual dimensions such that "more is better", for example 'Difficulty of Use' could be rescaled to 'Ease of Use'.<sup>5</sup>

Setting the price of store brand equal to 1.0, the relative prices of "Pro-Strip", "Cata-Kill", and "Tree-Guard" are approximately 2.9, 1.3, and 1.2, respectively. The new product "Attack" came in as a premium priced product with relative price approximately 2.9. A "per dollar" perceptual map based on this zero-point and these prices is shown in Figure 8.

Based on Figure 8, we compute utility for each  $\alpha - \beta$  combination and we identify for each evoked set those areas of the  $\alpha - \beta$  region where each product has the highest utility. Adjusting piecewise uniform distributions as per §3 and summing across

<sup>5</sup> Ideal points would require more complex scaling. For example, 'Sweetness' could become the 'Right Amount of Sweetness'. In our applications, we carefully word the raw semantic scales such that either "more is better" or "less is better".

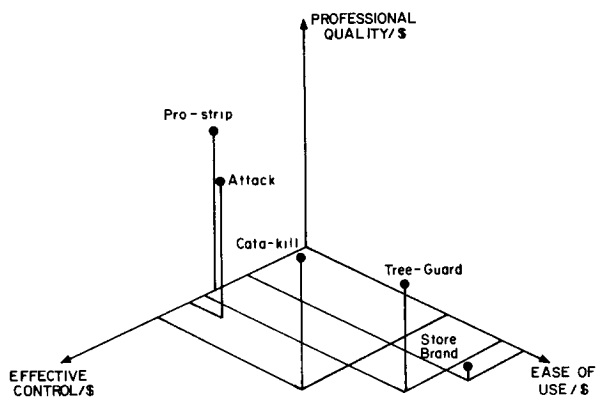


FIGURE 8. "Per Dollar" Perceptual Map (Disguised).

evoked sets we obtain the taste distribution,  $f(\alpha, \beta)$ . Since four products imply 15 evoked sets, the resulting taste distribution is reasonably smooth as shown in Figure 9. Figure 9 also shows that portion of the taste distribution captured (before attack) by "Pro-Strip", "Cata-Kill", and "Tree-Guard", respectively.

For convenience of exposition, we have not shown the portion of the taste distribution captured by "Store Brand", but it can be obtained readily from Figure 9. Since market shares must sum to 1.0, the portions captured by each brand must sum to  $f(\alpha, \beta)$ . Empirically, we have found that physical models of Figure 9, which fit together as a puzzle to form  $f(\alpha, \beta)$ , are valuable guides to help brand managers visualize their markets.

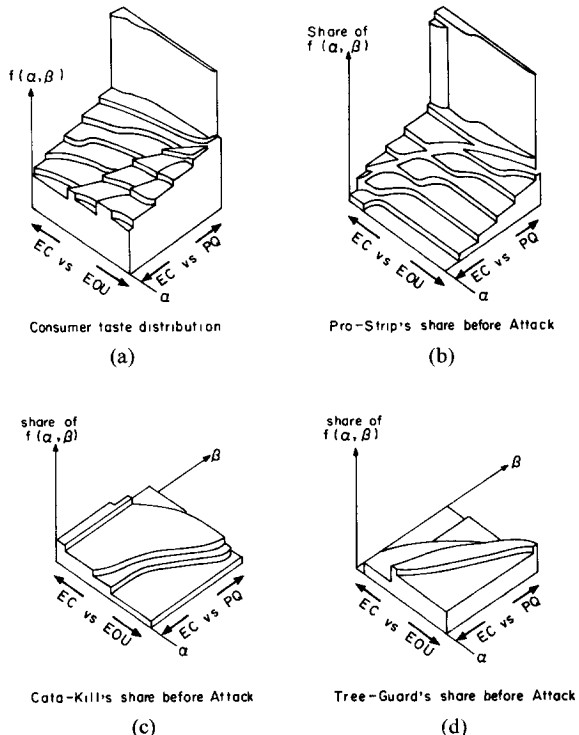


FIGURE 9. Consumer Taste Distribution,  $f(\alpha, \beta)$ , Representing Tradeoffs among Attributes; Shares of Pro-Strip, Cata-Kill and Tree-Guard Before Attack.

### Interpretation

Examine Figure 9a. Tradeoffs among ‘Effective Control’ and ‘Ease of Use’ are approximately uniformly distributed, but tradeoffs among ‘Effective Control’ and ‘Professional Quality’ slope toward ‘Professional Quality’ with a “tower” at extreme ‘Professional Quality’. We call this tower a ‘Professional Quality’ segment because it represents those consumers who put a very high weight on ‘Professional Quality’.

Now examine Figure 9b which illustrates that portion of the taste distribution representing consumers who choose “Pro-Strip”. Because “Pro-Strip” is clearly the best brand (before “Attack”) on ‘Professional Quality’ (review Figure 8), it captures almost the entire ‘Professional Quality’ segment. It also captures other portions of the taste distribution representing consumers that evoke only “Pro-Strip” or only “Pro-Strip” and one or two other brands. Similarly, Figures 9c and 9d illustrate that “Cata-Kill” and “Tree-Guard” capture more central portions of the taste distribution. Finally, “Store-Brand” (not shown) is left with that portion of the taste distribution favoring ‘Ease of Use’.

When “Attack” enters, the market changes. Based on the perceptual map in Figure 8, we expect that while “Attack” does well on ‘Professional Quality’, it does not do as well as Pro-Strip. We compute Attack’s unadjusted share by recomputing utility for each  $\alpha_i - \beta_i$  pair and summing across  $i$  as indicated in §3. Indeed, as Figure 10 illustrates, at 100% evoking, Attack captures a central portion of the taste distribution leaving the ‘Professional Quality’ segment to Pro-Strip. However, Attack (at 100% evoking) does hurt Pro-Strip by stripping away the more moderate consumers that Pro-Strip used to capture. (Compare Figure 9b to Figure 10b.) Of course, at less than 100% evoking, Attack’s share will be scaled down in proportion to its evoking and Pro-Strip’s share will be a weighted combination of its before and after shares (Figures 9b and 10b). See equations (3) and (4).

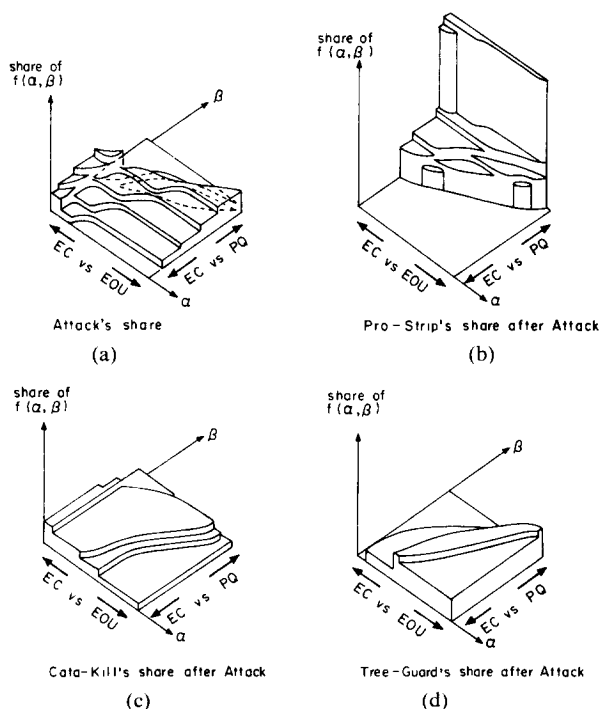


FIGURE 10. The Market after Attack Enters. (a) Attack's Share, (b) Pro-Strip's Share after Attack, (c) Cata-Kill's Share after Attack, and (d) Tree-Guard's Share after Attack. All Shares Are Unadjusted for Awareness and Availability.

According to Figures 9 and 10, Pro-Strip has retained its 'Professional Quality' franchise. If it is to regain share it must regain it from consumers with more moderate taste tradeoffs. However, strategically, it must maintain its 'Professional Quality' image to preempt a repositioning challenge by Attack. Figure 10 also indicates which consumers Cata-Kill and Tree-Guard will lose. Although they lose fewer consumers, they cannot ignore Attack's entry.

#### *Comparison of Forecasts*

Because  $f(\alpha, \beta)$  is a probability distribution, Attack's unadjusted share is the volume (Figure 10a) of  $f(\alpha, \beta)$  that it captures. Analytically, this is a 17% share. This is within one standard deviation (based on Urban and Katz 1983) of the Assessor unadjusted prediction of 19%. For comparison, we forecast based on preference regression and the traditional perceptual map in Figure 7. Based on standard procedures<sup>6</sup> we forecast an unadjusted share of 43% for Attack, much greater than the share predicted by either Assessor or Defender. Such a large share is not surprising if we place an ideal vector on Figure 7 as implied by the preference regression weights. Such an ideal vector would be shaded away from 'Ease of Use' toward 'Effective Control' and low on 'Professional Quality'. Attack does well relative to that ideal vector and, hence, traditional analysis predicts a high share for Attack. The share is lower based on Defender because the taste distributions in Figures 9 and 10 suggest a large 'Professional Quality' segment that Attack does not capture.<sup>7</sup>

Defender differs from traditional analysis because of the taste distribution and the "per dollar" map. We performed two additional analyses in an attempt to identify their differential effects. First, we modified preference regression to include price as a fourth attribute. A priori we expect this to lower Attack's share since Attack is a high priced product. However, although price has a significant coefficient, the coefficient is positive, probably because the higher priced brands, which are also better in the perceptual dimensions, get higher market share.<sup>8</sup> Thus, preference regression with price does not do as well as traditional analysis (using Assessor as a standard) and does much worse than Defender. See Table 2.

TABLE 2  
*Predicted Unadjusted Market Shares (Disguised Product Names)*

|                                      | Attack | Pro-Strip | Cata-Kill | Tree-Guard | Store Brand |
|--------------------------------------|--------|-----------|-----------|------------|-------------|
| Pre-Attack                           | —      | 0.41      | 0.24      | 0.24       | 0.12        |
| Assessor                             | 0.19   | 0.37      | 0.18      | 0.18       | 0.08        |
| Defender* (Per dollar maps)          | 0.17   | 0.30      | 0.23      | 0.21       | 0.09        |
| Traditional Perceptual Maps          | 0.43   | 0.26      | 0.13      | 0.14       | 0.05        |
| (With price)                         | 0.55   | 0.29      | 0.05      | 0.09       | 0.02        |
| Defender-like (traditional maps)     | 0.60   | 0.19      | 0.03      | 0.15       | 0.04        |
| Defender-like (price as a dimension) | 0.28   | 0.23      | 0.17      | 0.21       | 0.11        |

\*Zero-point =  $(-0.3, -0.2, -0.4)$ .

<sup>6</sup>Following established procedures (Urban and Hauser 1980, Chapters 10 and 11), we use the importance weights, 0.48, 0.38, and 0.14, and the perceptions for *each consumer* to compute the utility of each brand,  $j$ , for *each consumer*,  $c$ . The market share of a brand,  $j$ , is the percent of consumers for whom  $j$  is the maximum utility brand. See appendix for equations. Note that Figure 7 is just average perceptions, the forecasting procedure uses each consumer's perceptions.

<sup>7</sup>One can interpret preference regression as a normal distribution for  $\tan \alpha$  and  $\tan \beta$  centered at the angles corresponding to the importance weights. I.e.,  $\tan \alpha = (0.38)/(0.48)$  and  $\tan \beta = (0.14)/(0.48)$  or  $\alpha = 38^\circ$  and  $\beta = 16^\circ$ . The taste distribution in  $\alpha - \beta$  space would be transformed, but would not have a 'Professional Quality' segment as identified by Defender.

<sup>8</sup>See also discussion of price in preference regression in Rao and Gautschi (1982) and Srinivasan (1982).

We also undertook Defender-like analyses based on (1) a perceptual map without price and (2) a perceptual map with price as a fourth dimension. That is, we estimated a taste distribution,  $f(\alpha, \beta)$  or  $f(\alpha, \beta, \gamma)$ , based on the perceptual maps and forecast using the procedures of §3. Since we do not divide by price,<sup>9</sup> the ranking of utility via equation (2) does not depend on the zero-point. The forecasts are shown in Table 2.

Defender-like analysis without price as either a dimension or a scale factor does poorly because high priced brands tend to dominate low priced brands.

Defender-like analysis with price as a fourth dimension is much more interesting. Although the forecast (unadjusted) share of Attack is nine percentage points larger than the Assessor forecast the forecast is more moderate than traditional analysis. Interestingly, the marginal distribution of  $f(\alpha, \beta)$ , integrating  $\gamma$  out of  $f(\alpha, \beta, \gamma)$ , is similar to the taste distribution in Figure 9a. That is, tradeoffs among 'Effectiveness' and 'Ease of Use' are approximately uniformly distributed, and tradeoffs among 'Effectiveness' and 'Professional Quality' slant upwards toward 'Professional Quality' to achieve their maximum at extreme  $\beta$ . However, the 'Professional Quality' segment is not quite as pronounced as in Figure 9a. (One reason the forecast is nine points larger is because  $f(\alpha, \beta, \gamma)$  does not identify a pronounced 'Professional Quality' segment.) Tradeoffs among price, 'Effectiveness', and 'Ease of Use' are approximately uniformly distributed and tradeoffs among price and 'Professional Quality' favor 'Professional Quality'.

It appears, from this initial application, that Defender predictions differ from traditional predictions because of both the method of estimating taste variation and the use of "per dollar" perceptual maps.

For this initial application, Defender appears to reproduce Assessor's forecast better than any of the other four models tested. External validity against SAMI share is examined in §7. We caution the reader that this is but an initial test and it would be unfair to reject any of the models without extensive experience across product categories. Furthermore, models not included in our test such as conjoint analysis could also provide reasonable forecasts. However, we can infer that the Defender predictions are reasonable and that the Defender consumer model is worth further study.

Table 2 also compares the post-Attack predictions for all brands in the category. Comparing Defender and Assessor predictions, we see that Defender predicts a greater draw from Pro-Strip than does Assessor. As illustrated in Figures 9 and 10, this makes intuitive sense since Attack is positioned to draw moderate consumers who will accept a less 'Professional Quality' product if it is slightly easier to use and more effective than Pro-Strip. Assessor does not model this phenomenon explicitly. In test market, Attack did indeed draw more heavily from Pro-Strip than from other brands.

#### *A Comment*

At this point the reader may wonder why Defender should be developed if pretest market models (e.g., Assessor) already fulfill the predictive function. First, the Defender prediction requires only that we measure the new product's perceptual position and observe its price. Defender does not require the extensive laboratory measures that are required by pretest market models. (However, subject to further validation we recommend the use of convergent methods of prediction.) Second, and more importantly, the goal of our predictive test is not to establish a better forecasting model, but to investigate the reasonableness of the Defender consumer model. If the consumer model can predict well in at least one category, then we have more confidence in the

<sup>9</sup>Of course, when price is a fourth dimension, we enter it as a negative number so that "more is better". Also, the dimensions in the 4-dimensional extension of equation (2) are no longer "per dollar" when price is a fourth attribute. An alternative model might use  $(1/\text{price})$  as a dimension.

strategy implications that are based on Figures 9 and 10. Finally, the issues of "per dollar" perceptual maps and heterogeneous taste distributions are interesting scientifically independent of normative managerial considerations.

## 6. Ratio Scaling of "Per Dollar" Perceptual Maps

The estimates of the preference distribution,  $f(\alpha, \beta)$ , and the predictive accuracy of Defender will vary depending upon the reference zero-point chosen. This section examines the sensitivity of these estimates and predictions as the zero-point is varied. We also examine the sensitivity of the predictions to the choice of the preference distribution.

### *Feasible Region*

In order to ensure that all products in the "per dollar" perceptual map have positive scores on each dimension, and thus the market share is a decreasing function of price, we must choose a zero-point which is below the minimum value among brands along each dimension of the traditional perceptual map. For "Gypsy Moth Tape", these minimum values are  $(-0.21, -0.17, -0.38)$ , respectively, for 'Effective Control', 'Ease of Use', and 'Professional Quality'. The zero-point selected for the analyses in §5 was chosen to be within the feasible region, but not right on the border of the feasible region. These decisions were made *prior* to the predictive test.

### *Sensitivity*

We were somewhat surprised that an arbitrarily chosen zero-point did as well as it did. After all, it is not guaranteed that a zero-point will exist for which  $f(\alpha, \beta)$  can be chosen to fit market shares of existing brands within evoked sets. Predictive ability is certainly not guaranteed.

We systematically varied the zero-point, re-estimated  $f(\alpha, \beta)$  for each zero-point, and re-predicted "Attack's" market share. The results are summarized in Table 3 for the feasible region. Table 3 indicates that the predictions vary, but not dramatically, as we vary the zero-point within the feasible region. (Predictions vary from 19% to 13% for the feasible zero-point in Table 3.) Predictions do vary dramatically outside the feasible region, 0.13 to 0.35, as might be expected. Interestingly, had we chosen the point  $(-0.21, -0.17, -0.38)$ , which is the maximum allowable point in the feasible region, we would have predicted 19.5% which is even closer to the Assessor prediction than our *a priori* conservative selection.

TABLE 3  
Sensitivity to Zero-Point  
(Predicted Share of 'Attack' as a Function of the Zero-Point)

| $Z_0 = -0.4$ |              | $X_0 = -0.4$ | $X_0 = -0.3$ | $X_0 = -0.2$ |
|--------------|--------------|--------------|--------------|--------------|
|              | $Y_0 = -0.2$ | 0.15         | 0.17         | 0.19         |
|              | $Y_0 = -0.3$ | 0.15         | 0.17         | 0.19         |
|              | $Y_0 = -0.4$ | 0.14         | 0.16         | 0.18         |
| $Z_0 = -0.5$ |              | $X_0 = -0.4$ | $X_0 = -0.3$ | $X_0 = -0.2$ |
|              | $Y_0 = -0.2$ | 0.13         | 0.15         | 0.17         |
|              | $Y_0 = -0.3$ | 0.13         | 0.14         | 0.16         |
|              | $Y_0 = -0.4$ | 0.13         | 0.14         | 0.16         |

$X_0$  = zero-point for 'effective control'

$Y_0$  = zero-point for 'ease of use'

$Z_0$  = zero-point for 'professional quality'

Thus it appears that, at least for this one product category, choosing a zero-point reasonably close to the boundary of the feasible region allows us to (1) fit the market response to existing brands, and (2) predict the market share of the new brand as well as Assessor. Furthermore, predictions are reasonably insensitive to the choice of the zero-point as long as it is close to the boundary of the feasible region (upper right of Table 3). We turn now to the brief discussion of possible theoretical explanations for this empirical phenomenon.

### *Anchoring Effect*

For "Gypsy Moth Tape" Defender predicts best if we choose the zero-point to be near the maximum allowable point in the feasible region. At this point, we do not know whether this phenomenon is specific to the product category or whether it is a generalizable behavioral phenomenon.

If it is a generalizable phenomenon, then it raises an interesting set of strategies in which a firm can launch a "decoy" brand to shift the zero-point and perhaps increase the share of another of the firm's products.<sup>10</sup> Such decoying phenomena have been established experimentally in marketing science. For example, see Huber, Payne, and Puto (1982) and Huber and Puto (1983). In fact, Huber, Payne, and Puto suggest that the (dominated) decoy brand anchors perceptual dimensions and that other brands are then measured relative to the decoy.

The Huber-Payne-Puto anchoring effect explains the predictive ability of the maximum feasible zero-point by suggesting that consumers evaluate products relative to the worst product along each dimension. Such an anchoring effect is also consistent with the framing theories of Tversky and Kahneman (1981). See also discussions of price referents in Rao and Gautschi (1982) and Rao and Weiss (1982).

### *Uniform Distribution*

The Defender model requires us to estimate a taste distribution,  $f(\alpha, \beta)$ . We wondered how sensitive predictive results were to variations in this taste distribution. For example, how badly would predictions deteriorate if we used a uniform distribution (as in Figure 4a) rather than the appropriate  $f(\alpha, \beta)$ ?

First, we tried a uniform distribution without any model adjustments and found that we could not even fit existing shares. Reviewing the literature, we recognized that traditional models which assume a priori taste distributions all require "brand specific constants", that is, constants that are added to the utility of each existing product. For example, both logit analysis, which assumes double exponential taste distributions, and preference regression, which assumes Normal taste distributions, require brand specific constants for consistent estimators. See discussion in Coslett (1982). For a related viewpoint, see Srinivasan (1982). Thus, for a uniform taste distribution, we felt it was appropriate to include brand specific constants in the Defender model.

Analogous to logit and preference regression procedures, we selected brand specific constants for each of the existing brands by fitting the Defender model to existing markets shares. It was feasible to find brand specific constants. However, to predict, we need to forecast the brand specific constant for Attack.<sup>11</sup> Since this is equivalent to forecasting market share, we conclude that, at least for "Gypsy Moth Tape", the taste

<sup>10</sup>A decoy brand is a brand that is worse than all existing brands on one or more perceptual dimensions. In a true decoy strategy, the firm would expect that few if any consumers would purchase the decoy brand, but that the decoy brand would make the firm's existing brands look better and, hence, increase their share.

<sup>11</sup>The problem of estimating brand specific constants for new brands is a recurring problem in logit analysis. The issue of brand specific constants discussed above is not different qualitatively from that faced in logit analysis.

distribution contains important information about the product category and is, therefore, necessary to the model. A uniform distribution is not sufficient.

The empirical importance of the taste distribution  $f(\alpha, \beta)$  is satisfying since Hauser and Shugan (1983) allocate considerable theoretical effort to investigating the impact of the taste distribution. For example, a uniform distribution implies that the optimal defensive price response is to decrease price, but a multi-modal taste distribution may imply a price increase.

### *Summary*

Based on the above sensitivity analyses, for the product category under test, we posit that:

- (1) With the appropriate taste distribution, Defender predicts well for zero-points close to the maximum feasible values.
- (2) Predictions are reasonably insensitive to the choice of a zero-point if it is close to the maximum feasible value.
- (3) Predictions are sensitive to the choice of the taste distribution.

We expect some of these propositions to generalize while others will need to be modified as we gain further experience in a variety of product categories. At present, they are empirical propositions based on one empirical test.

## **7. Comparison to SAMI Data**

We are encouraged by the ability of Defender to match the predictive ability of Assessor for unadjusted share. However, actual share is based on adjustments due to awareness and availability (actually, advertising and distribution indices). For our purposes, we use the awareness and availability forecasts contained in the Assessor report as appropriate indices for the Defender consumer model. This is an empirical assumption subject to future research. As described in detail in Silk and Urban (1978), these forecasts are judgmental based on the firm's and the analyst's experience in the product category and based on expected levels of advertising and distribution budgets for the new product. For a discussion related to the effect of error in these forecasts see Urban and Katz (1983, p. 224). For Attack, these estimates were 0.7 and 0.6, respectively, yielding an adjusted share forecast of 7.1 percent ( $0.071 = 0.7 \times 0.6 \times 0.17$ ).

Defender and Assessor use "last brand purchased" as the raw data to estimate market shares within evoked sets. For Assessor, at least, models based on this raw data seem to forecast reasonably well the test market shares as reported by brand managers.

In our case, the managers of Pro-Strip felt that the SAMI shares were the best measures of test market shares. SAMI shares are the share figures they use in their own strategic planning. However, like any empirical measure, SAMI measures are not perfect. SAMI measures a virtual universe of warehouse withdrawals for large grocery stores, but may under represent drug stores and mass merchandisers. For "Gypsy Moth Tape", the latter is a small fraction of sales and the Pro-Strip managers felt that brand shares through the minor channels would be similar to those for the large grocery stores.

The "actual" shares available to us from test market are SAMI measures of volume share and of dollar share. At first, one might expect the appropriate measure is volume share, but in this category Attack and Pro-Strip tend to be reused occasionally whereas Store Brand requires slightly more product per use. Thus, volume share will be less than last brand purchased for Attack and Pro-Strip and more for Store Brand. Dollar share corrects for some of this measurement bias since the reusable brands cost more and Store Brand costs less. Together, volume share and dollar share bracket "last

TABLE 4  
*Comparison of Forecast Market Shares to  
 Test Market Results*

|                                       | Share |
|---------------------------------------|-------|
| <u>Predictions :</u>                  |       |
| Assessor                              | 8%    |
| Defender                              | 7%    |
| Traditional Perceptual Maps           | 18%   |
| Defender-like with Price as Dimension | 12%   |
| <u>Test Market Results*</u>           |       |
| SAMI dollar share                     | 8%    |
| SAMI volume share                     | 7%    |

\* Test market shares are four-week SAMI shares measured one year after data collection. All shares are rounded to the nearest percentage point for confidentiality. Awareness and availability are assumed to be 0.7 and 0.6, respectively, as per Assessor report.

brand purchased" share. Subject to these considerations we feel the reader will find comparisons to SAMI shares interesting.

Table 4 reports both the volume and dollar SAMI shares for the two test market cities one year after the initial data collection. For ease of comparison and confidentiality, we have averaged across the two cities.

Examining Table 4, we see that the SAMI shares are 7% for volume share and 8% for dollar share. Thus, the corresponding "last brand purchased" share would be in the range of 7% to 8%. Both Defender and Assessor forecast shares in this range whereas traditional perceptual maps are off by over a factor of 2. Defender-like analysis with price as a dimension does better with a forecast of 12%. We note that in other product categories predictions may not be as close as the predictions in Table 4. Only comparisons across a large number of categories can truly assess external validity.

Urban and Katz (1982) report a standard deviation in predictive errors of Assessor of about 2.0 percentage points. We expect Defender to be in that range. Based on this standard deviation, there is a 40% chance that predictions are within one percentage point of "actual". If favorable roundoff truncations are considered, there is a 40%–68% chance that predictions are within one percentage point of "actual". Thus, it appears reasonable that the accuracy reported in Table 4 is compatible with published error bounds on Assessor.

In summary, Defender appears to make predictions within the range of "actual" share. Thus, subject to future validation tests in other product categories, we feel the Defender consumer model is a reasonable marketing science model.

## 8. Application to a Second Product Category

After applying Defender analysis to the "Gypsy Moth Tape" category, the manufacturers of Pro-Strip asked us if we could apply Defender to another category in which they were under attack. This application illustrates the use of Defender analysis to guide and expand managerial judgment.

### *Managerial Problem*

In spring of 1981, a major \$100 million product category, which we call *A-B-C*, was about to be attacked by a new product, which we call *N*. Product *N* had been in test market for nine months and its share had been oscillating between 10 and 14 percent. In that test market the share of brand *C*, the market leader, was down between 2 and 6

percent.<sup>12</sup> In response, the managers of brand *C* commissioned a number of market research studies including Assessor and perceptual mapping. The Assessor analysis suggested that brand *N* could maintain a 12 percent adjusted share in a national launch. The perceptual map suggested that brand *N* was positioned between brands *B* and *C* that, qualitatively, a 2–6 percent draw from *C* could be maintained in a national launch.<sup>13</sup>

In response, the manufacturers of brand *C* began work on a defensive flanker, which we will call brand *D*. By fall 1983, brand *N* had not yet been launched in the national market, but the manufacturers of *C* judged a national launch to be imminent. Brand *D* was not yet ready for a full Assessor study, but the brand *C* managers wanted an early reading on the potential impact of brand *D*.

### Defender Map

Despite the lag of two and one-half years from the original data collection, a time lag in which the market may have changed, we were asked to retrofit the 1981 perceptual map. Furthermore, in 1983 there was some concern by the brand *C* managers with the apparent position of brand *B* in the 1981 map, which they felt was overstated relative to the current market position of brand *B*. Thus, based on the

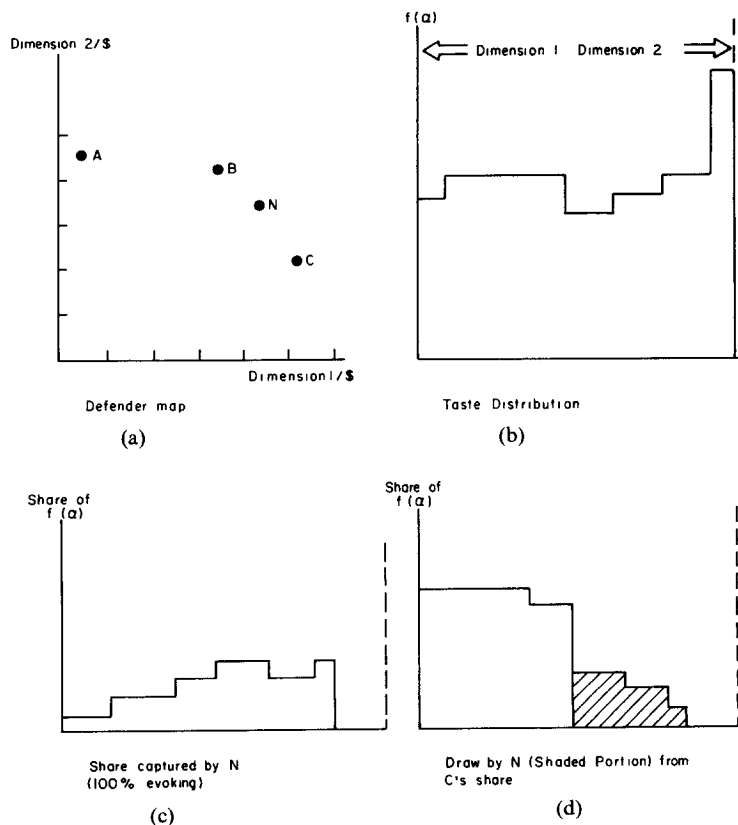


FIGURE 11. Application in A-B-C Category. (a) Defender Map, (b) Total Distribution, (c) Share Captured by *N*, and (d) Draw by *N* from *C*'s Share.

<sup>12</sup>Shares were oscillating due to promotion and trial/repeat phenomena. Such oscillation is normal for a test market.

<sup>13</sup>The perceptual map was based on factor analysis of 17 semantic scales. The factor solution yielded two dimensions explaining 95.7 of the common variance (45.6% of the total variance). We label the dimensions 1 and 2 for confidentiality.

original map, secondary market research data, and the combined judgment of the management and analysis teams we produced the map in Figure 11a as representative of the 1983 (test) market. The relative prices are 1.0, 1.02, 2.5, and 2.4 for brands *A*, *B*, *C* and *N*, respectively.

Because the map was modified by the judgment of managers who were aware of both the Assessor and the test market results, analyses based on this map cannot be considered tests of external validity. We present these analyses to illustrate the merging of quantitative Defender analyses with qualitative managerial judgment.

### *Analyses*

Following the procedures outlined in §3 we use the map in Figure 11a and estimates of the market shares of *A*, *B*, and *C* (as obtained by "last brand purchased" in Assessor) for every evoked set to estimate the taste distribution. The empirical taste distribution,  $f(\alpha)$ , is shown in Figure 11b. It is close to a uniform distribution but with a minor 'dimension 2' segment.

Based on the position of brand *N* and the taste distribution, we calculate that the portion of the taste distribution captured by brand *N* is 12 percent and that the draw from brand *C* is 4 percent. Since the map was drawn by managers and analysts implicitly aware of these results, these predictions indicate simply that we have been successful in developing an analytic representation of the test market results.

Fifteen alternative scenarios were simulated in which share, draw, cannibalization, and profit were calculated with the Defender model. Some scenarios involved launching brand *D*, some involved repositioning brand *C*, and some involved a combination of these strategies. Countermoves by brand *N* were also simulated. The analyses suggested that brand *D* should not be launched since most of its sales would come from brand *C*. However, the manufacturers of brand *C* could defend their market successfully by matching any price cuts by brand *N* and by repositioning brand *C* to maintain dominance on dimension 1.

The management team was satisfied that the Defender map captured their intuition in a usable way. They are now defending their market based on these recommendations.

## **9. Conclusions, Hypotheses, and Future Directions**

This completes our initial test of the Defender consumer model. Based on our results, we feel that the consumer model has the potential to predict as well as existing state-of-the-art models and to have good external validity. Because the consumer model is based on empirically observed marketing phenomena, is derived from axiomatic economic theory, predicts as well as highly refined pretest market models in at least one category, and shows reasonable external validity, we posit that the Defender model is an adequate representation of aggregate consumer response.

Furthermore, because we limited ourselves to standard data collection procedures, our results suggest that the Defender consumer model is feasible and within existing market research data collection budgets.

However, the defensive strategy research stream is far from complete. Our analyses raise as many questions as they answer. The remainder of this section highlights the issues that we feel are most important.

### *Behavioral Hypothesis: Anchoring*

Tests in the "Gypsy Moth Tape" product category suggest that the best zero-points are near the position of the worst product along each dimension. This hypothesis is intuitively appealing and is consistent with experiments and theories in consumer

behavior. Perhaps it can begin to explain why we are able empirically to “ratio-ize” what theoretically should be an interval scale.

### *Multiple Dimensions*

Both Attack and product *N* entered the market with relative strengths on more than one dimension. Since Attack was the fifth major product in its market and product *N* was the fourth major product in its category, these multiple-dimensional attacks are consistent with the economic theories of Lane (1980) which suggest that such late entrants in a category use such a “central” attack. Such multiple dimensional attacks are the result of earlier products choosing positions to capture the largest possible portion of the taste distribution. As a result, if new technologies are not discovered, later products will find it more difficult to obtain the same market share as early entrants. This phenomenon provides an alternative explanation to the “rewards to first entrant” theories in marketing (Urban et al. 1984) and in economics (Schmalensee 1982).

### *Heterogeneity*

Traditional perceptual maps, which use preference regression or logit analysis, model perceptions as heterogeneous. Assessor and stochastic preference theory (Bass 1974) model preferences as heterogeneous. Conjoint analysis models the consumer taste distribution as heterogeneous. All predict well under the right circumstances. In reality, we know that perceptions, tastes, and even choice rules are heterogeneous. In each of these models either perceptions or preferences is modeled as heterogeneous while the other is viewed as homogeneous. Such a modeling assumption is necessary because a fully heterogeneous model, that is, a model having unique tastes *and* perceptions for every consumer, likely would be overspecified in the sense of having more parameters than data.

Defender (in this application) assumes perceptions are homogeneous despite the fact that we know they are distributed with standard errors equal to approximately 10% of the range. We feel that the model predicts market behavior as well as it does because the taste distribution captures some of the heterogeneity of perceptions as well as tastes. See similar speculation in Lancaster (1971). Similarly, models such as preference regression and logit may do well despite assuming homogeneous tastes because perceptual heterogeneity captures some of the heterogeneity in tastes. This phenomenon is related to issues of aggregation in econometrics. For a general discussion, see Stoker (1982). Because heterogeneity is debated frequently in marketing science, we believe this issue is worth further analytical and empirical investigation. For example, we are now extending the Defender consumer model to include normally distributed heterogeneous tastes as suggested by Hauser and Simmie (1981, pp. 42–44).

### *Future Directions*

The most controversial assumption in the Defender consumer model is the “per dollar” perceptual map. This paper has begun to address that issue as well as the procedure to estimate taste distributions. This is a beginning. The next step is to develop a full normative application including response functions to predict awareness, availability, and perceptual position as a function of dollar spending by the defending firm. In theory, response functions are feasible using a variety of techniques suggested by Little (1975) and others. A related question is the empirical viability of the Assessor awareness and availability indices as estimates of evoking. We are now expanding the consumer model to include empirical response functions, consumer response to promotion, and the dynamic effects of short-term and long-term response.

Besides the two applications reported in this paper, the Defender consumer model

has been applied (in the U.S.) to the OTC analgesic category (Halloran and Silver 1983), to decision support software (Elkin and Borschberg 1983), and to another OTC health care product. It has been applied in Japan to a food product and a health and beauty aid. However, validations and Assessor comparisons are not yet available in these categories.

Future research includes investigation of the behavioral hypotheses, validation in more product categories, further validation of Defender's forecasts of draw from existing brands, and improvements in data collection.<sup>14</sup>

**Acknowledgement.** We wish to thank Glen Urban, Al Silk, John Little and Leigh McAlister of MIT, Robert Klein, Sue Bass and Peter Guadagni of Management Decision Systems (MDS), and Karl Irons of the Test Marketing Group, Inc. for their advice, insights, and support throughout this project. The data were supplied by an unnamed sponsor and by Management Decision Systems. The necessary computer support was provided by Management Decision Systems. This paper has benefitted from seminars at MIT, MDS, University of Connecticut, University of Osaka (Japan), ORSA/TIMS Orlando, Marketing Science Institute Special Mini-Conference, Kao Soap Ltd. (Japan), Johnson and Johnson, Ltd. (Japan), and numerous unnamed consumer product companies in the United States.

#### Appendix. Equations of Defender Model

We state the equations for three perceptual dimensions. Generalization is as described in the text. See Hauser and Shugan (1983) for analytic formulae of two perceptual dimensions.

##### Notation

Let  $j$  index the products,  $n$  indicate the new product,  $l$  index the evoked sets, and  $i$  index the  $N$  equal areas of the  $\alpha - \beta$  feasible region. Let  $x_j, y_j, z_j$  indicate the "per dollar" positions of product  $j$  (or the new product) on each of the three perceptual dimensions, let  $\alpha_i, \beta_i$  be representative values of  $\alpha$  and  $\beta$  for the  $i$ th region, and let  $u_j(i)$  be the utility of brand  $j$  for  $i$ th area of the  $\alpha - \beta$  region. Let  $f_i(\alpha_i, \beta_i)$ , or more simply  $f_i(i)$ , be the fraction of consumers using evoked set  $i$  who are represented by taste angles  $\alpha_i$  and  $\beta_i$ . Let  $A$  be the index set of all products, let  $A_l$  be the index set of all products belonging to evoked set  $l$ , and  $S_l$  be the probability that a randomly chosen consumer will choose from evoked set  $l$ . Let  $m_{j|l}$  and  $M_{j|l}$  be the predicted and actual market shares of product  $j$  in evoked set  $l$  and let  $m_j$  and  $M_j$  be the overall market shares of product  $j$ .

##### Shared within Evoked Sets

Utility is defined by equation (2) in the text. In our notation this becomes:

$$u_j(i) = x_j + (\tan \alpha_i)y_j + (\tan \beta_i)z_j. \quad (A1)$$

Let  $\delta_j(i, l)$  be an indicator variable such that

$$\delta_j(i, l) = \begin{cases} 1 & \text{if } u_j(i) > u_k(i) \quad \text{for all } k \in A_l, \\ 0 & \text{otherwise.} \end{cases} \quad (A2)$$

For the new product,  $n$ , define  $A_l^* = A_l \cup \{n\}$  and define  $\delta_j^*(i, l)$  as above replacing  $A_l$  by  $A_l^*$ . Further define

$$\delta_n(i, l) = \begin{cases} 1 & \text{if } u_n(i) > u_j(i) \quad \text{for all } j \in A_l, \\ 0 & \text{otherwise.} \end{cases} \quad (A3)$$

Then, the market share of an existing product is given by:

$$m_j = \sum_l S_l \sum_i f_i(i) \delta_j(i, l). \quad (A4)$$

##### Estimation of Taste Distribution

Consider first evoked sets with two or more products. In these evoked sets,

$$f_i(\alpha_i, \beta_i) = \sum_j \left[ M_{j|l} \delta_j(i, l) / \sum_k \delta_j(k, l) \right]. \quad (A5)$$

Now consider evoked sets which contain only one product. Let  $g_i(\alpha_i)$  and  $h_i(\beta_i)$  be the marginals such that

<sup>14</sup>This paper was received May 1983 and was with the authors for 2 revisions.

$f_i(\alpha_i, \beta_i) = g_i(\alpha_i)h_i(\beta_i)$ . Notice the implicit parsimonious independence assumption for singleton evoked sets. Then:

$$g_i(\alpha_i) = \begin{cases} 2(1 - \alpha_i/90^\circ)/90^\circ & \text{if for } j \in A_i, x_j \geq x_k \text{ for all } k \in A \text{ and } y_j < y_k \text{ for some } k \in A, \\ 2\alpha_i/(90^\circ)^2 & \text{if for } j \in A_i, y_j \geq y_k \text{ for all } k \in A \text{ and } x_j < x_k \text{ for some } k \in A, \\ (N)^{-1/2} & \text{otherwise.} \end{cases}$$

$h_i(\beta_i)$  is defined analogously to  $g_i(\alpha_i)$  replacing  $y_j$  by  $z_j$ . Finally,

$$f(\alpha_i, \beta_i) = \sum_l f_l(\alpha_i, \beta_i)S_l. \quad (\text{A6})$$

#### Forecasting for the New Product

Let  $a_n$  and  $d_n$  be the advertising and distribution indices for the new product. Then the market share of the new product is given by:

$$m_n = a_n d_n \sum_l S_l \sum_i f_l(i) \delta_n(i, l). \quad (\text{A7})$$

For existing products, the market share after attack is given by:

$$m_j^* = a_n d_n \sum_l S_l \sum_i f_l(i) \delta_j^*(i, l) + (1 - a_n d_n) \sum_l S_l \sum_i f_l(i) \delta_j(i, l). \quad (\text{A8})$$

For changes in existing products' strategies we define  $a_j^0$ ,  $d_j^0$ ,  $a_j^*$ , and  $d_j^*$  for before and after advertising and distribution indices and repeat equations (A1)–(A4), (A7), and (A8) treating the modified product,  $j^*$ , as we treated the new product,  $n$ . Evoking probabilities  $S_j^*$  for  $A_j^*$  such that  $j^* \in A_j^*$  are modified upward (or downward) by  $a_j^* d_j^* / a_j^0 d_j^0$ . Evoking for  $A_j$  such that  $j^* \notin A_j$  is modified downward (or upward) such that  $\sum_l S_l^* = 1$ . The notation is cumbersome but this procedure is easy to implement numerically.

#### Preference Regression

Let  $c$  index consumers and  $C$  be the number of consumers. Then consumer  $c$ 's perceptions of the  $j$ th product are  $x_{cj}$ ,  $y_{cj}$ , and  $z_{cj}$ , respectively. Let  $p_{cj}$  be consumer  $c$ 's preference for the  $j$ th product. Then the preference regression equation is:

$$p_{cj} = w_x x_{cj} + w_y y_{cj} + w_z z_{cj} + \text{error} \quad (\text{A9})$$

where  $w_x$ ,  $w_y$ , and  $w_z$  are obtained via regression across  $c$  and  $j$ . For prediction, we define estimated utilities,  $\hat{u}_{cj}$  and  $\hat{u}_{cn}$  as derived from equation (A9) and define a new indicator variable

$$\Delta_{cj} = \begin{cases} 1 & \text{if } \hat{u}_{cj} \geq \hat{u}_{cn} \text{ for } j, n \text{ evoked by } c, \\ 0 & \text{otherwise.} \end{cases}$$

Estimated market share is given by

$$m_j = (1/C) \sum_c \Delta_{cj}. \quad (\text{A10})$$

#### References

- Bass, F. M. (1974), "The Theory of Stochastic Preference and Brand Switching," *Journal of Marketing Research*, 2 (February), 1–20.
- Burger, P., C. H. Gundee and R. Lavidge (1981), "COMP: A Comprehensive System for the Evaluation of New Products," in Y. Wind, V. Mahajan, and R. N. Cardozo, Eds., *New Product Forecasting*, Lexington, Mass.: Lexington Books, 269–284.
- Choffray, J. M. and G. L. Lilien (1980), *Market Planning for New Industrial Products*, New York: John Wiley & Sons.
- Coslett, S. R. (1982), "Efficient Estimation of Discrete-Choice Models," in C. Manski and D. McFadden, Eds., *Structural Analysis of Discrete Data With Econometric Applications*, Cambridge, Mass: M.I.T. Press, 51–110.
- Elkin, W. B. and A. Borschberg (1983), "Defensive Marketing Strategy: Its Application to a Financial Decision Support System," S. M. Thesis, Sloan School of Management, M.I.T., Cambridge, Mass., June.
- Eskin, G. J. and J. Malec (1976), "A Model for Estimating Sales Potential Prior to Test Market," *Proceedings of American Marketing Association Educators' Conference*, Chicago: American Marketing Association, 230–233.
- Gaskin, S. P. (1983), "Defender: The Test and Application of a Defensive Marketing Model," S. M. Thesis, Sloan School of Management, M.I.T., Cambridge, Mass. February.

- Gavish, B., D. Horsky and K. Srikanth (1983), "An Approach to Optimal Positioning of a New Product," *Management Science*, 29, 11 (November), 1277-1297.
- Green, P. E. and V. Srinivasan (1978), "Conjoint Analysis in Consumer Research: Issues and Outlook," *Journal of Consumer Research*, 5, 2 (September), 103-123.
- Halloran, M. J. and M. A. Silver (1983), "Defensive Marketing Strategy: Empirical Applications," S. M. Thesis, Sloan School of Management, M.I.T., Cambridge, Mass., June.
- Hauser, J. R. and F. S. Koppelman (1979), "Alternative Perceptual Mapping Techniques: Relative Accuracy and Usefulness," *Journal of Marketing*, 16 (November), 495-506.
- and S. M. Shugan (1983), "Defensive Marketing Strategies," *Marketing Science*, 2, 4 (Fall), 319-360.
- and P. Simmie (1981), "Profit Maximizing Perceptual Positions: An Integrated Theory for the Selection of Product Features and Price," *Management Science*, 27, 1 (January), 33-56.
- Huber, J., J. W. Payne and C. Puto (1982), "Adding Asymmetrically Dominated Alternatives: Violations of Regularity and the Similarity Hypothesis," *Journal of Consumer Research*, 9, 1 (June), 90-98.
- and C. Puto (1983), "Market Boundaries and Product Choice: Illustrating Attraction and Substitution Effects," *Journal of Consumer Research*, 10, 1 (June), 31-44.
- Kotler, P. (1983), *Principles of Marketing*, 2nd Ed., Englewood Cliffs, N.J.: Prentice-Hall.
- Lancaster, K. J. (1971), *Consumer Demand: A New Approach*, New York: Columbia University Press.
- (1980), "Competition and Product Variety," *The Journal of Business*, 53, 3, 2 (July), S79-S104.
- Lane, W. J. (1980), "Product Differentiation in a Market with Endogenous Sequential Entry," *Bell Journal of Economics*, II, 1 (Spring), 237-260.
- Little, J. D. C. (1975), "BRANDAID: A Marketing Mix Model, Structure, Implementation, Calibration, and Case Study," *Operations Research*, 23, 4 (July-August), 628-673.
- McFadden, D. (1980), "Econometric Models for Probabilistic Choice Among Products," *Journal of Business*, 53, 3, II (July), S13-S29.
- Pessemier, E. A. (1982), *Product Management: Strategy and Organization*, 2nd Ed., New York: John Wiley & Sons.
- Pringle, L. G., R. D. Wilson and E. I. Brody (1982), "NEWS: A Decision Oriented Model for New Product Analysis and Forecasting," *Marketing Science*, 1, 1 (Winter), 1-30.
- Rao, V. R. (1984), "Pricing Research in Marketing: The State-of-the-Art," *Journal of Business*, 57, 1 (January), S39-S60.
- and D. A. Gauscht (1982), "The Role of Price in Individual Utility Judgments: Development and Empirical Validation of Alternative Models," in L. McAlister, Ed., *Choice Models for Buyer Behavior*, Greenwich, Conn.: JAI Press, 57-80.
- and E. N. Weiss (1982), "Resource Allocation Problems with Interval Scale Coefficients," Working Paper, Cornell University, Ithaca, N.Y., December.
- Ratchford, B. T. (1982), "Economic Approaches to the Study of Market Structure and their Implications for Marketing Analysis," in R. K. Srivastava and A. D. Shocker, Eds., *Analytic Approaches to Product and Marketing Planning. The Second Conference*, Cambridge, Mass.: Marketing Science Institute, 60-78.
- Schmalensee, R. (1982), "Product Differentiation Advantages of Pioneering Brands," *American Economic Review*, 72, 349-365.
- Sen, S. (1982), "Issues in Optimal Product Design," in R. K. Srivastava and A. D. Shocker, Eds., *Analytic Approaches to Product and Marketing Planning: The Second Conference*, Cambridge, Mass.: Marketing Science Institute, 265-274.
- Srinivasan, V. (1902), "Comments on the Role of Price in Individual Utility Judgments," in L. McAlister, Ed., *Choice Models for Buyer Behavior*, Greenwich, Conn.: JAI Press, Inc., 81-90.
- Stoker, T. M. (1982), "The Use of Cross-Sectional Data to Characterize Macro Functions," *Journal of the American Statistical Association*, 77, 378 (June), 369-380.
- Tauber, E. M. (1977), "Forecasting Sales Prior to Test Market," *Journal of Marketing*, 41, 1 (January), 80-84.
- Tversky, A. and D. Kahneman (1981), "The Framing of Decisions and the Psychology of Choice," *Science*, 211, 4481 (January 30), 453-458.
- Urban, G. L., T. Carter, S. Gaskin and S. Mucha (1984), "Market Share Rewards to Pioneering Brands: An Empirical Analysis and Strategic Implications," Working Paper, M.I.T., Cambridge, Mass., January.
- Urban, G. L. and J. R. Hauser (1980), *Design and Marketing of New Products*, Englewood Cliffs, N.J.: Prentice-Hall.
- and G. M. Katz (1983), "Pretest-Market Models Validation and Managerial Implications," *Journal of Marketing Research*, 20, 3 (August), 221-234.
- Wind, Y. L. (1982), *Product Policy: Concepts, Methods, and Strategy*, Reading, Mass.: Addison-Wesley.
- and D. Lerner (1979), "On the Measurement of Purchase Data: Surveys vs. Purchase Diaries," *Journal of Marketing Research*, 16, 1 (February), 39-47.
- Yankelovich, Skelly and White, Inc. (1981), "LTM Estimating Procedures," in Y. Wind, B. Mahajan, and R. N. Cardozo, Eds., *New Product Forecasting*, Lexington, Mass.: Lexington Books, 249-268.