

TESTING COMPETITIVE MARKET STRUCTURES

GLEN L. URBAN, PHILIP L. JOHNSON AND JOHN R. HAUSER

*Massachusetts Institute of Technology
Management Decision Systems, Inc.
Massachusetts Institute of Technology*

An accurate understanding of the structure of competition is important in the formulation of many marketing strategies. For example, in new product launch, product reformulation, or positioning decisions, the strategist wants to know which of his competitors will be most affected and hence most likely to respond. Many marketing science models have been proposed to identify market structure.

In this paper we examine the managerial problem and propose a criterion by which to judge an identified market structure. Basically, our criterion is a quantification of the intuitive managerial criterion that a "submarket" is a useful conceptualization if it identifies which products are most likely to be affected by "our" marketing strategies. We formalize this criterion within the structure of classical hypothesis testing so that a marketing scientist can use statistical statements to evaluate a market structure identified by: (1) behavioral hypotheses, (2) managerial intuition, or (3) market structure identification algorithms.

Mathematically, our criterion is based on probabilities of switching to products in the situation where an individual's most preferred product is not available. 'Submarkets' are said to exist when consumers are statistically more likely to buy again in that 'submarket' than would be predicted based on an aggregate "constant ratio" model. For example, product attributes (e.g., brand, form, size), use situations (e.g., coffee in the morning versus coffee at dinner), and user characteristics (e.g., heavy versus light users) are specified as hypotheses for testing alternate competitive structures.

Measurement and estimation procedures are described and a convergent approach is illustrated. An application of the methodology to the coffee market is presented and managerial implications of six other applications are described briefly.

(Market Structure; Competitive Strategy; Product Line; Entry Opportunities)

Perspective

The modeling of competitive market structure represents an interesting area for research on market behavior and a crucial activity in the formulation of effective marketing strategy. New product development, product policy, and competitive advertising and pricing decisions depend in part upon the identification of which products compete most strongly with one another.

Recently, two lines of research have addressed the structure of competition among products. The first is based on information processing theory and mathematical psychology. Bettman (1971, 1979), Haines (1974), Lussier and Olshavsky (1979), and

Payne (1976) describe the decision processes that *individual* consumers use to attain information, assimilate that information, and utilize it to make product decisions. Such decision processes are often described by sequential processing hierarchies. In mathematical psychology, Tversky (1972) has developed a theory called "elimination-by-aspects" (EBA) in which product characteristics are chosen at random and all products not having those characteristics are eliminated. The process continues until one product remains. While EBA sometimes looks like a decision hierarchy and has often been cited as a basis for aggregate market structure, Tversky and Sattath (1979, p. 540) point out that the process is not entered sequentially, but rather at random. Tversky and Sattath (1979) instead propose a sequential processing rule, called "hierarchical elimination model" (HEM), which applies to very special cases called preference trees. Hauser and Tversky (1983) extend HEM to the general case. However, neither HEM nor EBA is preserved by aggregation. Each needs further assumptions to be applicable to an aggregate market structure.

The second line of research does not model explicitly the individual decision sequence, but rather describes the aggregate nature of competition. Various approaches to defining criteria for competition have been proposed including, among others: (1) switching between brands (Butler 1976; Kalwani and Morrison 1977; Rao and Sabavala 1981; Robinson, Vanhonacker and Bass 1980; Charnes, Cooper, Learner and Phillips 1979; Vanhonacker 1979, 1980; Ehrenberg and Goodhardt 1982); (2) the structure of choice probabilities (McFadden 1980; Batsell 1980); (3) in-use substitution (Stefflre 1972; Day, Shocker and Srivastava 1979; Bourgeois, Haines, and Sommers 1979; Srivastava, Leone, and Shocker 1981; Arabie, Carroll, DeSarbo, and Wind 1981); (4) the segments of consumers who use the product (Frank, Massy, and Wind 1972); and (5) similarity of interpurchase times (Fraser and Bradford 1983). These approaches allow competitive structures to be defined by either product attributes (e.g., form or brand), use similarity, or consumer characteristics (e.g., heavy versus light users).

The information processing/mathematical psychology literature has often been cited as a behavioral motivation for models of aggregate market structure. Although the aggregate descriptions may be consistent with individual behavior, this is rarely true. For example, Tversky and Sattath (1979, p. 552) point out, "It is well known that most probabilistic models (including EBA and the constant ratio rule) are not preserved by aggregation." Even at the individual level, Hauser and Tversky (1983) caution "switching hypotheses are neither implied by nor imply cognitive processing hierarchies." The individual cognitive structure and aggregate market structure methods reflect two different approaches to understanding the bounds of competition.¹

Determining the correct aggregate competitive structure is important to managers. Consider a product line decision. It is often desirable for a firm to have one product in each of the major sectors of the market, and to avoid unnecessary duplication between products. If a firm can identify a sector in which it does not now compete, then it could consider allocating resources to develop a new product for introduction in that sector in order to generate incremental sales and profits. Conversely, the presence of duplication within a sector could lead to the dropping of a product and a consolidation of the product line within that sector. Changes in the product line reflect commitments of millions of dollars in many organizations. Similarly, the success of

¹ Stochastic process models such as Jeuland, Bass, and Wright (1980) are based on assumptions about the aggregate summary probability distribution (e.g. Dirichlet) of individual behavior. Nonetheless, aggregate parameters (e.g., the parameters of the Dirichlet) are estimated and the market is described by those aggregate parameters. We classify this method as an aggregate market structure method.

advertising, pricing or product reformulation strategies in a competitive environment depends in part upon which competitors are most (least) affected by the strategy and hence most (least) likely to consider a competitive response.

Managerial actions will depend on the specific market structure. It is important to have a procedure that can determine if all products in a market compete with each other or if 'submarkets' exist where the level of competition is high within them and low between them.² It is equally important to know how the submarkets are identified if they exist. Are the submarkets characterized by product attributes (decaffeinated versus caffeinated coffee), users characteristics (heavy versus light coffee consumption rates) or uses (morning coffees versus evening coffees)? Finally, it is important that the manager have confidence that the market structure upon which he (she) plans his (her) strategy is a reasonable description of the probable actions by consumers who are affected by his (her) strategy.

This paper pursues the research thrust directed at testing the *aggregate* competitive structure of a market once hypotheses about structure have been generated. It begins with a managerially relevant definition of market structure and formulates the definition as a mathematical statement based on switching when a product is deleted from the market. We develop a statistical test to identify whether a market satisfies that mathematical definition. We use the statistical test in a procedure for testing aggregate competitive structures. We describe alternative measurement procedures to obtain the necessary data for the tests and provide illustrative examples of testing competitive structures based on attribute, use, and user characteristics within the product deletion definitional framework.

We outline managerial use of the statistics and report an empirical application based on laboratory measures for the coffee market. We illustrate empirically how our methodology can test alternative theories of market structure. In this section, we also examine whether our methodology converges with some selected alternative methodologies. Finally, we relate our experience by briefly summarizing the managerial implications of six other applications of our procedures. This paper closes with a brief discussion of research issues.

Definition of Competitive Structure

As an intuitive introduction to the issues of competitive structure consider the following automobile industry situation. In 1981, most of the sporty front wheel drive (FWD) automobiles were sold in the U.S. by foreign manufacturers, e.g., Honda, Saab, Volkswagen. In June 1981, General Motors launched the FWD J-body cars (Chevrolet Cavalier and Pontiac J-2000) which were similar in appearance to the sporty FWD foreign cars. However, the J-body cars were similar in price, interior room, power, and, of course, manufacturer to General Motors' family FWD X-body cars (Chevrolet Citation, and Pontiac Phoenix). If the J-body cars drew customers from sporty FWD imported cars, then General Motors would have been successfully competing against the imported cars. If, instead, the J-body cars drew customers from family FWD domestic cars, then the J-body cars would have cannibalized sales from the X-body cars. Figure 1 is a visual representation of the alternative placements of the

²Testing a 'market' for 'submarkets' presupposes that a 'market' has already been defined. For example, at one level we can test the market, 'coffees,' for submarkets. At another level 'coffees' can be a submarket of 'hot beverages' or even 'beverages.' If the appropriate data are available our tests can be used to address the latter question as well as the former.

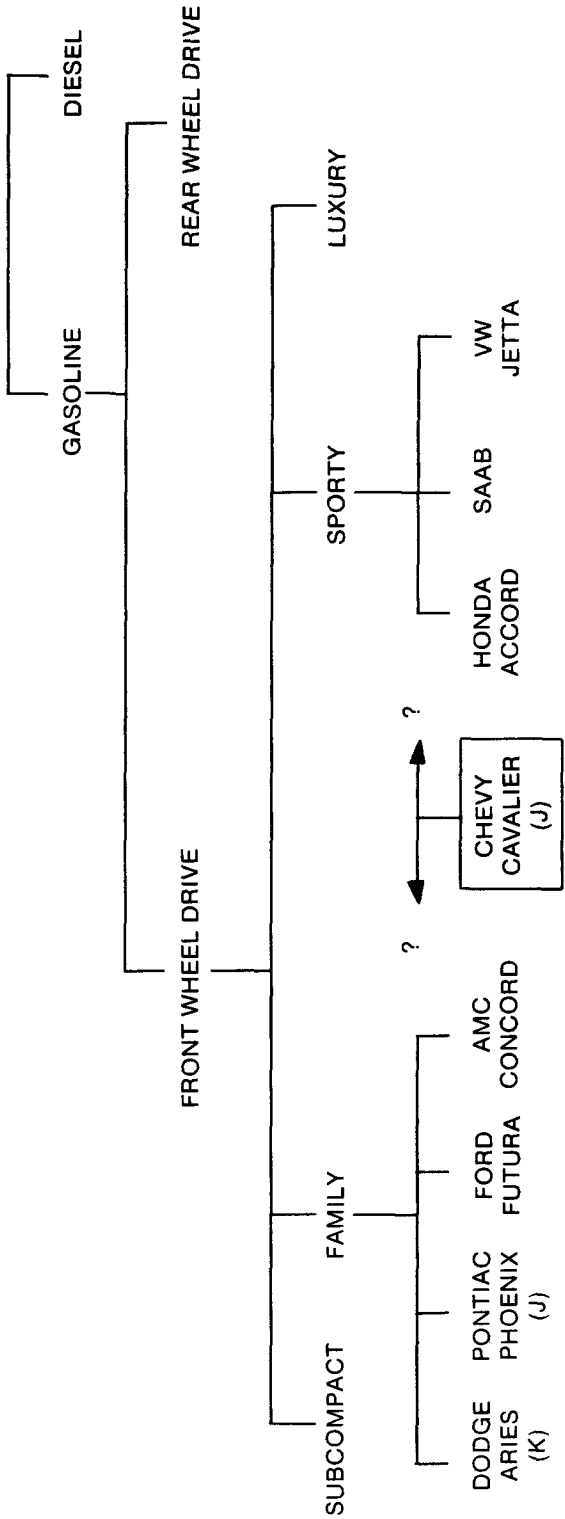


FIGURE 1. Hypothetical Structure of the Automobile Market. (Diesel, Rear Wheel Drive and Subcompact branches may also have sub-branches. Letters in parentheses indicate that other, similar brands are produced by the same manufacturer.)

J-body cars.³ (Figure 1 is simplified for expositional purposes. There are over 160 automobile models on the market.)

Figure 1 is based on *Consumer Reports* (January/February 1982), but consider for a moment the managerial definition implicit in the above statement. We tend to say that the J-body cars are in the “sporty FWD” ‘submarket’ if they draw share from the Honda Accord, Saab, and Volkswagen Jetta; they are in the “family FWD” ‘submarket’ if they draw share from the Dodge Aries (K), Pontiac Phoenix (X), Ford Futura, and AMC Concord.

However, suppose the J-body cars were popular and capture, say, a market share of 8.8% of the total number of FWD cars sold.⁴ Suppose further that for this example “family FWD” cars account for the majority of the sales of FWD cars, say 96%. Suppose the J-body cars captured all of the consumers who previously would have purchased sporty FWD imported cars (4% of the total market) and 5% of the consumers who previously would have purchased “family FWD” cars (4.8% of the total market— $5\% \times 96\%$). Do we classify the J-car as a “family FWD” car because it draws more consumers from that submarket than the “sporty FWD” submarket? Do we classify the J-car as a “sporty FWD” car because it draws all of the “sporty FWD” sales and only a small fraction of the “family FWD” sales? Do we classify it as neither? Or both? The choice depends on managerial need and how we represent that need through our definition of market structure.

In this paper we choose our definition to reflect what we believe is the key message portrayed by a grouping of products into submarkets. Suppose a macro-market, e.g. “FWD automobiles,” has been defined and suppose we want to ask whether that macro-market contains any submarkets, e.g., “sporty FWD” and “family FWD.”

We must first contemplate an ideal unstructured market. In such an unstructured market, the manager would not expect his product to draw customers equally from all existing products. He would expect popular products to be hurt more than unpopular products, that is, he would expect to draw more customers from high share (of the macro-market) products than from low share products. Thus, we define an ideal unstructured market as a market in which a new product draws its share from existing products in proportion to the market shares of the existing products. In such an unstructured market, the manager would have no need to group products into submarkets since he could predict competitive impact by simply knowing the market shares of all existing products.

On the other hand, if his product were such that it hurt specific identifiable products more so than market shares would predict, he would want to predict this when formulating his market strategy. In our example, the Chevy Cavalier draws more customers from sporty FWD cars than would be predicted by a market share of 8.8%. A market structure grouping should convey this message to the manager. In other words, a market structure (as defined here) tells the manager that he will affect products grouped with his product more so than would be predicted by an ideal unstructured market.

This definition, in our opinion, captures the key managerial requirement of a market partition, that is, that the partitioning of the market into submarkets explains more about consumer behavior than is apparent to the manager from the unpartitioned market. According to this definition, the J-body cars would be classified as “sporty FWD” because they draw all of the previous “sporty FWD” share whereas their market share would predict that they would draw only 8.8% of the “sporty FWD”

³Figure 1 assumes that the structure is stable in the sense that the introduction of the new car does not totally redefine other groupings.

⁴This example is purely illustrative. It does not reflect actual automobile shares.

market. The above definition would reject classification of the J-body cars as "family FWD."

This definition is useful for conceptualization of the managerial issue. But we also want to identify competitive structure before a new product is introduced. Thus we will also operationalize the definition with respect to product deletion. In particular,

A market is defined by a series of submarkets, if, when a product is deleted from a submarket, its former consumers are more likely to buy again in that submarket than would be predicted by market shares.

A moment's reflection reveals that both definitions capture essentially the same phenomenon. We find the second definition easier to operationalize. Both relate directly to the impact of managerial changes in the composition of the product line and to the definition of competition.

Finally, we note that other marketing scientists may define structure in other ways depending upon the managerial problems they face and depending upon what they wish to portray to the manager. We have chosen a specific definition which we feel has intuitive appeal. It is an aggregate definition because the managerial issue it addresses is related to aggregate strategy.

What follows is a deductive mathematical analysis based on this definition.

Statistics

Our approach to statistical testing for the presence of submarkets is based on classical hypothesis testing. We formulate a null hypothesis that reflects the existence of no competitive submarkets. Next we hypothesize a structure. Each hypothesized structure predicts how the market will behave when a product is deleted. To be retained (escape rejection), the hypothesized structure must explain product deletion probabilities better (statistically) than the null hypothesis of no structure. As is the case with all statistical hypothesis testing (e.g. Green and Tull 1978; Morrison 1976), we may:

- (1) retain only one structure as better than the null hypothesis,
- (2) fail to reject the null hypothesis relative to each structure, or
- (3) retain more than one structure.

Managerial action in cases (1) and (2) is clear. We address later what to do in case (3).

We begin with a model to analyze market behavior and formalize what we mean by "no structure."

Aggregate Constant Ratio Model

In order to use the above definition, we must state mathematically what we mean by "would be predicted by market shares." To give this definition rigor we use the aggregate constant ratio model (ACRM). See discussions by Tversky and Sattath (1979, p. 552) and Bell, Keeney, and Little (1975).

Basically ACRM is an aggregate version of what is known as Luce's axiom (1959). However, we caution the reader that ACRM is not an aggregation of individuals who themselves obey an individual level constant ratio (CRM). ACRM is purely a statement about how an aggregation of individuals behaves. It is relevant for our analysis because, like our definition, it applies to the *aggregate* behavior of a group of consumers. For further discussion and examples see Tversky and Sattath (1979, pp. 552-554).

According to ACRM there exist scale values, m_j , for every product, j , in the market. Then, for any submarket (set of products), A , the market share, $P(j|A)$, of product j in

that submarket is given by:⁵

$$P(j|A) = m_j / \sum_{k \in A} m_k \quad (1)$$

where the denominator is simply the sum over all products in the submarket, A. If A is the total market, T, we normalize such that $\sum_{k \in T} m_k = 1$. In this case, m_j also becomes the market share because $P(j|T) = m_j$.

To implement our definitions we introduce some simplified notation. The reader will note that sets and indices of sets are denoted by **boldface type**.

$P_i(j)$ = the overall market share of product j when product i is deleted, $i \neq j$.

s = a set of products, called submarket s .

$P_i(s)$ = the market share of submarket s out of the total market when product i is deleted.

From equation (1) and the above definitions we have:

$$P_i(j) = m_j / \sum_{k \neq i} m_k = m_j / (1 - m_i), \quad (2)$$

$$P_i(s) = \sum_{\substack{j \in s \\ j \neq i}} P_i(j) = \left(\sum_{\substack{j \in s \\ j \neq i}} m_j \right) / (1 - m_i), \quad (3)$$

where the sum in (3) is over all products in submarket s except when i is in submarket s , in which case we delete i from the sum.

For example, if the scale values (m_j) of products 1, 2, 3, and 4 were each 0.25, then $P_1(2) = P_1(3) = P_1(4) = 0.25 / (1 - 0.25) = 0.333$. If the first submarket ($s = 1$) were products 1 and 2, and the second submarket ($s = 2$) were products 3 and 4, then $P_1(s = 1) = P_1(2) = 0.33$ and $P_1(s = 2) = P_1(3) + P_1(4) = 0.67$.

Equations (2) and (3) predict what the new shares will be under the null hypothesis for any specific groupings of products into submarkets. We now compare these aggregate predictions of ACRM to the observed behavior of consumers. Let:

n_i = number of consumers who choose product i when all products are available.

$n_i(j)$ = the number of consumers out of n_i , who formerly chose product i but who now choose product j when product i is deleted from the market.

$n_i(s)$ = the number of consumers out of n_i , who formerly chose product i but who now choose a product from submarket s . (Product i is no longer available.)

If there were no market structure, then ACRM would apply for any grouping of products and we would expect the observed frequencies of purchases to satisfy:⁶

$$n_i(j) / n_i \approx P_i(j) \quad \text{and} \quad (4)$$

$$n_i(s) / n_i \approx P_i(s) \quad (5)$$

where the approximate equality ($\approx$) is due to sampling error for finite n_i .

Now let $\hat{P}_i(s) = n_i(s) / n_i$ be the estimated probability of buying in set s when product i is not available (based on those consumers who previously chose product i). Then,

⁵Equation (1) is similar in structure to that derived by Jeuland, Bass, and Wright (1980, p. 262) under the assumptions that (a) individuals are described by a multinomial process, (b) the parameters of the process are distributed via a Dirichlet distribution across individuals, and (c) brand choice and purchase timing are independent. Thus, our statistical tests can be utilized to examine the market boundaries used with the multinomial-Dirichlet brand switching mode. If 'no structure' is rejected at the market level but applies (is not rejected) at the submarket level, then the multinomial-Dirichlet model may be appropriate for each submarket but not for the aggregate market.

⁶Note that if we consider the entire market, that is, if $s \equiv T$, then equation (5) is an identity since $n_i(T) \equiv n_i$ and $P_i(T) \equiv 1.0$. However, for any $s \neq T$, equation (5) becomes an empirically testable statement.

according to our definition of competitive structure, a structure defined by a specific set of submarkets exists when switching is greater in each submarket than would be predicted by "no structure." Thus, for each submarket, s , in the structure we expect:

$$\hat{P}_i(s) \gtrsim P_i(s) \quad \text{if } i \text{ is in } s, \quad (6)$$

$$\hat{P}_i(s) \lesssim P_i(s) \quad \text{if } i \text{ is not in } s, \quad (7)$$

where approximation is again due to sampling errors. Equations (6) and (7) are now a mathematical interpretation of our definition of competitive structure.

Consider the above example, where $n_1 = 100$; then if $n_1(s = 1) = 34$ ($\hat{P}_1(s = 1) = 0.34$ vs. $P_1(s = 1) = 0.33$) and $n_1(s = 2) = 66$ ($\hat{P}_1(s = 2) = 0.66$ vs. $P_2(s = 2) = 0.67$) we would probably not reject the hypothesis of no structure. On the other hand, if $n_1(s = 1) = 98$ ($\hat{P}_1(s = 1) = 0.98$ vs. $P_1(s = 1) = 0.33$) and $n_1(s = 2) = 2$ ($\hat{P}_1(s = 2) = 0.02$ vs. $P_2(s = 2) = 0.67$) we would probably reject ACRM as a model of the market and, according to equations (6) and (7), say that the hypothesized structure explains observed behavior in a way consistent with our definition.

Note that the criteria we use to test a submarket, that is, the inequalities in equations (6) and (7), depend upon how we define the submarkets. Thus, for "good" groupings into submarkets, equations (6) and (7) will hold, while for other potential groupings they may not hold and, hence, we would reject such groupings as not being superior to 'no groupings.'

If n_i were infinitely large then we could use equations (6) and (7) directly; however, for finite n_i we must recognize sampling errors.

Normal Test

Our null hypothesis of no structure is that the *population* satisfies ACRM for any grouping, thus, if a consumer is drawn randomly from those who previously purchased product i , his probability of choosing from submarket s is $P_i(s)$. If successive draws are independent, then this creates a binomial process for $n_i(s)$ with mean, $n_i P_i(s)$, and variance, $n_i P_i(s)(1 - P_i(s))$. If n_i is sufficiently large, then the Central Limit Theorem applies and the distribution of $n_i(s)$ is given by:⁷

$$n_i(s) \sim N[n_i P_i(s), n_i P_i(s)(1 - P_i(s))] \quad (8)$$

where $\sim N[\mu, \sigma^2]$ means distributed as normal with mean, μ , and variance, σ^2 . We define "success" with respect to the set s rather than with respect to the product j since that is consistent with our qualitative definition. If the m_j are known for all products in the total market, then the $P_i(s)$ are given by equation (3) and the statistical tests are standard one-tail Z -tests based on inequalities (6) and (7).

Note that our null hypothesis does not assume homogeneous CRM consumers, but is rather an aggregate statement. Heterogeneity is a subtle, complex issue. As illustrated later in this paper (see Tables 1 and 2), it is possible to have every consumer satisfy CRM *yet* have market structure. For example, consider a market where 50% of the consumers satisfy CRM but with probabilities favoring sporty FWD cars and 50% satisfy CRM but with probabilities favoring family FWD cars. Such a market clearly has structure.

However, it is easy to show that equation (8) is an upper bound for heterogeneous

⁷See Drake (1967, pp. 212–221). A good rule of thumb is $n_i > 20$. Note also that the DeMoivre–Laplace Central Limit Theorem requires (suppress s) that $n_i P_i > 3\sigma_i$ and $n_i(1 - P_i) > 3\sigma_i$ where $\sigma_i = [n_i P_i(1 - P_i)]^{1/2}$. If n_i is not sufficiently large to assure that the inequalities are satisfied, we can replace equation (8) with a Poisson approximation, as discussed in Drake (1967, p. 220).

CRM consumers and exact for homogeneous CRM consumers,⁸ thus we feel equation (8) is an appropriate variance formula to represent our *aggregate* definition. That is, it is appropriately conservative and tells us when n_i is large enough to have statistical confidence in the inequality comparisons of equations (6) and (7). For more in-depth discussion of ACRM versus heterogeneous CRM consumers, see Tversky and Sattath (1979).

Summary Statistics

A one-tail Z-test based on equation (8) compares for each product, i , a hypothesized market structure to the null hypothesis of no structure. Because our primary concern is the submarket, s , that contains product i , it is useful to define the following summary statistics:

$$n(s) \equiv \sum_{i \in s} n_i(s), \quad (9)$$

$$n^* \equiv \sum_{s \in T} n(s), \quad (10)$$

where the first sum is over all products in s and the second sum is over all submarkets in the total market, T .

It is easy to write down the distribution of $n(s)$. Remember n_i is a different sample from n_j for $i \neq j$; thus the terms in equations (9) and (10) are independent normal random variables. Therefore, the means, or the variances, are sums of the means, or the variances, in equation (8).⁹

These summary statistics can also be re-expressed as proportions

$$\hat{P}(s) = n(s) / \sum_{i \in s} n_i \quad \text{and} \quad (11)$$

$$\hat{P}^* = n^* / \sum_{i \in T} n_i. \quad (12)$$

The corresponding weighted ACRM proportions are

$$P(s) = \sum_{i \in s} P_i(s) n_i / \sum_{i \in s} n_i \quad \text{and} \quad (13)$$

$$P^* = \sum_{s \in T} \left(P(s) \sum_{i \in s} n_i \right) / \sum_{i \in T} n_i. \quad (14)$$

These aggregate statistics are useful in communicating the results of an evaluation of the overall competitive structure. (Actual statistical tests are made with respect to $n(s)$)

⁸ Suppressing the argument s and the subscript i and indexing consumers by c ,

$$\begin{aligned} \text{var}(n) &= \sum_c P_c(1 - P_c) = n\bar{P} - \sum_c P_c^2 = n\bar{P} - \left[\sum_c (P_c - \bar{P})^2 + n\bar{P}^2 \right] \\ &= n\bar{P} - n\bar{P}^2 - \sum_c (P_c - \bar{P})^2 < n\bar{P} - n\bar{P}^2 = n\bar{P}(1 - \bar{P}). \end{aligned}$$

⁹ In particular,

$$\begin{aligned} n(s) &\sim N \left[\sum_{i \in s} n_i P_i(s), \sum_{i \in s} n_i P_i(s)(1 - P_i(s)) \right] \quad \text{and} \\ n^* &\sim N \left[\sum_{s \in T} \sum_{i \in s} n_i P_i(s), \sum_{s \in T} \sum_{i \in s} n_i P_i(s)(1 - P_i(s)) \right]. \end{aligned}$$

and n^* .) The empirical case represented at the end of this paper demonstrates the use of these summary measures of competition.

Testing Procedure

The previous section provides a means to turn our managerial definition into a statistical statement. Figure 2 summarizes one way in which a marketing scientist can use these statistics to help a manager identify a competitive market structure. We illustrate our procedures later when we apply our methods empirically.

The first step is to generate hypotheses about the nature of submarkets. Exploratory analysis is useful. Applying perceptual mapping or clustering procedures to usage data (e.g. Srivastava, Shocker and Day 1978; Day, Shocker and Srivastava 1979; and Srivastava, Leone and Shocker 1981) could yield hypotheses on submarkets characterized by use occasions. Applying analytical procedures to switching data (e.g. Kalwani and Morrison 1977, or Rubinson, Vanhonacker and Bass 1980) could generate hypotheses characterized by product attributes (e.g. brand or form). Clustering of individual attributes could define user characteristics or hypotheses on competitive sectors (e.g. Frank, Massy and Wind 1972). Managerial judgement may also be utilized to formulate competitive hypotheses.

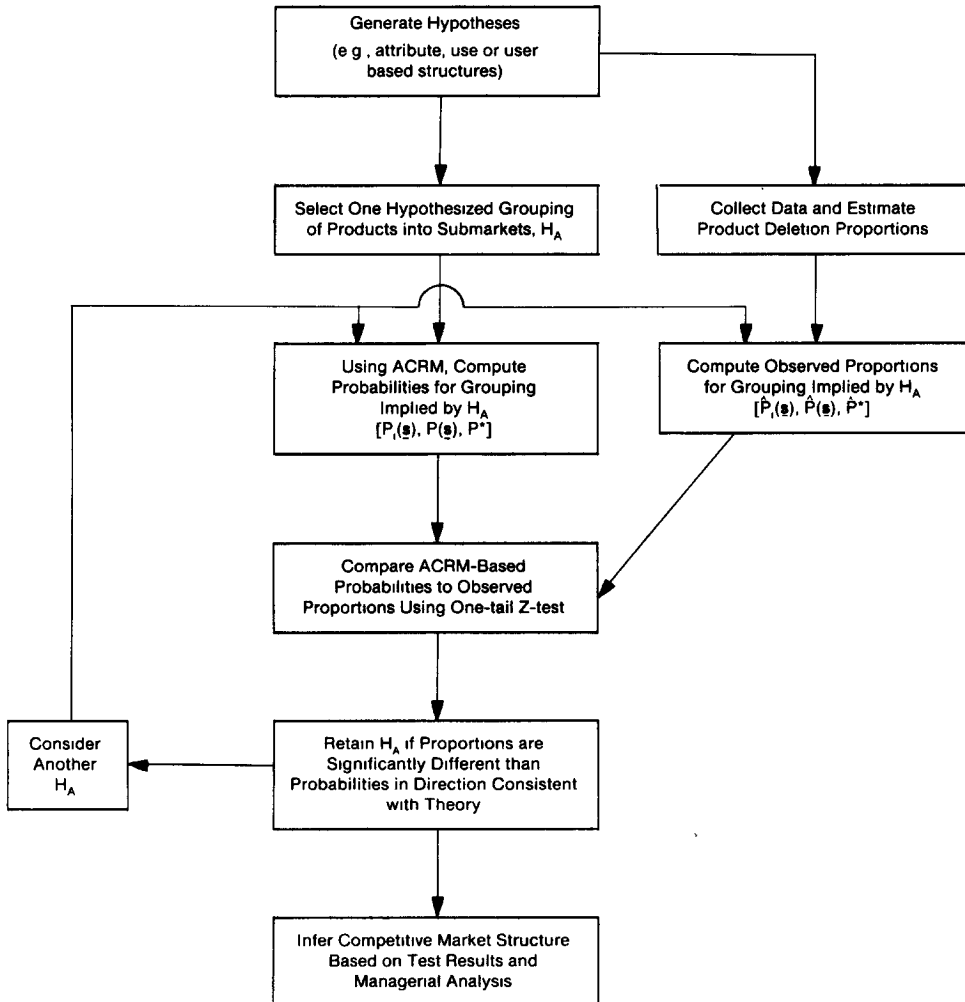


FIGURE 2. Testing Procedure.

Next specific data are obtained to test each hypothesized grouping of products into submarkets. Experimentally controlled forced switching data is one way to estimate directly the forced switching frequencies, $n_i(j)$. We describe others.

Next the hypothesized groupings (H_A) are tested statistically. Since each hypothesized grouping will imply different sets of inequalities as per equations (6) and (7), each hypothesized grouping is tested independently. (The inequalities vary because s varies.) In some cases, one can nest hypotheses, i.e., define ACRM on a submarket and test sub-submarkets. If a hypothesis is not significantly better than no structure at, say, the 10% level, it is eliminated. After the hypotheses are tested, all or some may have been eliminated.

Before selecting the competitive structure for managerial purposes, we may wish to explore for the existence of "compound structures" in the definition of the submarkets.¹⁰ That is, we explore whether consumers can be grouped such that each group is characterized by a different competitive definition (e.g. people over 50 years of age may see the coffee submarkets as decaffeinated and caffeinated, while those under 50 may see the submarkets as ground and instant).

Finally we select the best structure for managerial analysis. If no test rejects the hypothesis of no structure in favor of a specific hypothesized structure, we use the unstructured description. If only one of the hypothesized structures is significant, it is subjected to further evaluation. Because alternative hypotheses have been tested, care should be taken in interpreting the significance level so as not to exploit random error. We recommend re-testing of the chosen hypothesis with saved data or through convergence from separate measures of switching. If the convergent analysis or retest is consistent, the hypothesis may be adopted for managerial analysis. It is clear at this point that the generation of hypotheses must span the relevant set of possibilities if the hypothesis testing procedure is to identify the best competitive definition.

A more difficult case exists if more than one of the attribute, use, or user hypotheses is significant. In this case we rely on managerial judgment. The significance levels of the alternative hypotheses and the results of convergent analysis are useful inputs to the decision, but the final choice will also reflect managerial experience.¹¹ Additional forced switching studies also may be undertaken to collect more evidence for selection of the best competitive description. This would be particularly true if a compound structure were obtained from the exploratory analysis.

Given a competitive structure, managerial analysis could be conducted to assess opportunities for new products, the advisability of dropping products to consolidate product offerings, or to answer questions such as "who is my competition?", or "what universe of products do we use to calculate market share for strategic purposes?"

In the following sections, we describe and illustrate the steps in the testing process and present an empirical application. We begin with data collection and estimation.

¹⁰We have chosen the words "compound structure" to avoid confusion with heterogeneity which is usually taken to mean a continuous probability distribution of varying CRM probabilities. We seek to identify different groups of consumers such that the null hypothesis of no structure is rejected within each consumer group, but for which the hypothesized structure varies across consumer groups.

¹¹We avoid the temptation to choose the highest Z-score since that would be analogous to choosing a regression based on the highest F-score. Such procedures exploit random error. The Z-score provides a guide to the decision but should not be the only criterion. In making a choice it is useful to recognize that the variance of the summary statistic, n^* , is summed over all products and hence depends upon the total sample size. The variances of the more specific submarket statistics, $n(s)$ are summed only over those products in s and hence depend upon $\sum_s n_i$. Since the power of a Z-test increases with sample size, the likelihood of falsely retaining the null test will be similar across alternative hypotheses when using the summary statistics. On the other hand, the likelihood of false retention will be larger for "smaller" submarkets, i.e., smaller $\sum_s n_i$. Thus, care should be exercised when interpreting Z-statistics for small submarkets or when interpreting the Z-statistics testing $n_i(s)$ for small share brands.

Data Collection and Estimation Procedures

The construct upon which the Z-statistics are based is $n_i(j)$, the number of consumers who buy product j when product i is deleted from the market. In this section we discuss how one might measure or estimate $n_i(j)$. Greater details on the actual measures are given in the empirical case.

Forced Switching

The most natural method of observing $n_i(j)$ is to first observe the product that an individual most prefers then place him in a choice situation in which his preferred product has been removed from the choice set. We call this experiment forced switching. In our empirical case we use a simulated store to observe each consumer's choice from the appropriately modified choice sets.

Preference Rank

An alternative data collection procedure is to ask consumers to rank order the products they would consider in terms of preference. For a given individual we identify product i as their first ranked product and product j as their second ranked product. Then $n_i(j)$ is estimated as the number of consumers who rank product i first and product j second. To the extent that rank order preference reproduces actual choice, this measurement will provide an estimate of the true $n_i(j)$. Another variant of this approach is to calculate the fraction of people who last purchased their most preferred brand, 2nd most preferred brand, etc. This vector can then be applied to each consumer's rank order brand preferences to provide a probability of purchase for each product in their consideration set. These can be aggregated across consumers to estimate $n_i(j)$.

Logit / Preference Intensity

Suppose that we can estimate for each consumer, c , his probabilities, $P'_c(i)$, of choosing product i . Suppose we can estimate $P'_c(i)$ for all products in consumer c 's consideration set. Suppose further that all individual consumers satisfy CRM.

Then, for consumer, c , the probability, $P_{ci}(j)$, that he chooses product j from the set of products in which product i has been deleted is given by:

$$P_{ci}(j) = P'_c(j) / [1 - P'_c(i)] \quad (15)$$

where $P'_c(j) = 0$ if j is not in c 's consideration set. $n_i(j)$ is then obtained by simple aggregation of conditional probabilities,¹² i.e.:

$$n_i(j) = \sum_c P_{ci}(j) P'_c(i), \quad (16)$$

$$n_i = \sum_c P'_c(i). \quad (17)$$

The first term is the probability that c chooses j second given that he chose i first; the second term is the probability that c chooses i first.

For example, suppose there are four products on the market and two consumers with the CRM probabilities show in Table 1.

For consumer 1, $P_{11}(2) = 0.67$ ($0.67 = 0.4 / (1 - 0.4)$), $P_{11}(3) = 0.17$, $P_{11}(4) = 0.17$ and for consumer 2, $P_{21}(2) = 0.11$, $P_{21}(3) = 0.44$, $P_{21}(4) = 0.44$. The estimate of $n_1(2)$ is

¹² $n_i(j)$ based on equation (16) is an estimate of the $n_i(j)$ that would be observed empirically. Although the sampling variance is different if $P'_c(i)$ varies across c , equation (8) still serves as an upper bound as shown in footnote 8. An alternative hybrid technique uses first preference to compute n_i . Then $n_i(j)$ is computed via equation (15) summed across all consumers who prefer i .

TABLE 1
Example of Individual Probabilities

	Consumer 1	Consumer 2
Product 1	0.4	0.1
Product 2	0.4	0.1
Product 3	0.1	0.4
Product 4	0.1	0.4

then given by:

$$n_1(2) = P_{11}(2)P'_1(1) + P_{21}(2)P'_2(1) = (0.67)(0.4) + (0.11)(0.1) = 0.279,$$

$$n_1 = P'_1(1) + P'_2(1) = (0.4) + (0.1) = 0.5.$$

Thus, $\hat{P}_1(2) = n_1(2)/n_1 = 0.558$. The reader may wish to verify that we obtain an (estimated) force switching matrix in Table 2 which suggests a definite market structure grouping products 1 and 2 as one submarket and products 3 and 4 as the other submarket. This structure occurs because of heterogeneity of preference within the individual CRM models. This example illustrates a case where, although each individual satisfies CRM, the market is structured.

There are many ways to estimate the individual CRM probabilities. The most common method is the multinomial logit model with constant sum paired comparison preference data as the explanatory variables. (Actual choice is the dependent quantal variable in the estimation.) This method assumes that CRM is true over each individual's consideration set. For more details see McFadden (1980), Silk and Urban (1978), and Urban and Hauser (1980, Chapter 11).

Consideration Sets

In some cases the only data available may be an indication of which product is preferred and which products each consumer would consider purchasing. Sometimes we only know the consideration set information. Intuitively we expect that these consideration sets carry much information about what products group together in submarkets. Suppose that consumer c evokes e_c acceptable products where $e_c \geq 2$. Then if we assume consumers are equally likely to buy any product within their consideration sets, we get:

$$P'_c(i) = \begin{cases} 1/e_c & \text{if } i \text{ is considered by consumer } c, \\ 0 & \text{if } i \text{ is not considered by consumer } c. \end{cases} \quad (18)$$

We can use equation (18) with equations (15), (16), and (17).

Alternatively, if we can identify the first preference product, we can use a hybrid approach. That is, we compute $P_{ci}(j)$ with equations (15) and (18), but estimate $n_i(j)$

TABLE 2
Estimated Forced Switching Matrix [$P_i(j)$]

		Product			
Product:	Product 1	Product 1	Product 2	Product 3	Product 4
	Product 1	—	0.56	0.22	0.22
	Product 2	0.56	—	0.22	0.22
	Product 3	0.22	0.22	—	0.56
	Product 4	0.22	0.22	0.56	—

by summing over only those consumers who prefer product i . n_i becomes the number of consumers who prefer product i .

Convergence

The forced choice procedure measures the aggregate proportion of respondents switching to product j when i is not available, while the preference rank, logit, and consideration set procedures calculate these probabilities based on individual choice probabilities. A convergent approach can be used when both experimentally forced choices and individual probabilities are available. The probability of buying again in the submarket ($P_i(s)$) can be estimated by each method and the statistical adequacy assessed. If both methods agree, confidence is increased. If they disagree, an examination of sources of bias may reveal, in a particular application, that one method is preferred. If no bias is identified, the probabilities ($P_i(s)$) can be pooled across methods to obtain a combined assessment of the competitive structure (see application for more details).

Testing Hypotheses

In this section, we formulate and illustrate the statistical testing for hypotheses based on attributes, on use, and on user.

Attribute Hypotheses

We begin with a simplified example of how the normal test would be applied to competitive structures based on attributes of automobiles. Suppose there are only seven specific models (or brands) of automobiles on the market; three diesels (Peugeot 505, Olds Cutlass, and VW Jetta) and four gasoline powered cars (Chevrolet Cavalier, Peugeot 505, Olds Cutlass, and VW Jetta). Table 3 represents a hypothetical data matrix for 100 consumers. The numbers in the first column represent the number of consumers who selected that model of automobile as their first choice, n_i . The numbers in the matrix represent the number of consumers whose first choice is the automobile designated by the row label and whose second choice is the automobile designated by column label, $n_i(j)$. For simplicity in this example, we assume the share of product i is $m_i = n_i / \sum_i n_i$.

Two alternative hypotheses for a competitive structure are "Diesel vs. Gasoline" as shown in Table 4(a) and "model-specific" as shown in Table 4(b). Inspection of the data in Table 3 suggests that "diesel" consumers would stay with diesels and "gasoline" consumers would stay with gasoline powered automobiles if they were

TABLE 3
Hypothetical Forced Switching Matrix [$\hat{P}_i(j)$]

	n_i	Diesel Peugeot	Diesel Cutlass	Diesel Jetta	Gasoline Cavalier	Gasoline Peugeot	Gasoline Cutlass	Gasoline Jetta
Diesel								
Peugeot	10	—*	6	2	1	1	0	0
Cutlass	20	10	—	4	0	2	4	0
Jetta	15	7	2	—	1	1	1	3
Gasoline								
Cavalier	20	0	2	4	—	0	4	10
Peugeot	15	3	1	1	0	—	5	5
Cutlass	10	0	1	0	3	3	—	3
Jetta	10	0	0	2	1	5	2	—

*No repeat purchase is allowed under the forced switching conditions.

TABLE 4
Two Alternative Competitive Structures

(a) "Diesel" vs. "Gasoline"			
<u>Diesel</u>		<u>Gasoline</u>	
Peugeot		Cavalier	
Cutlass		Peugeot	
Jetta		Cutlass	
		Jetta	
(b) "Model-specific"			
<u>Peugeot</u>	<u>Cutlass</u>	<u>Jetta</u>	<u>Cavalier</u>
Gasoline	Gasoline	Gasoline	Gasoline
Diesel	Diesel	Diesel	

forced to switch from their most preferred car. But does this hypothesis hold up statistically?

Table 5 uses the Z-statistics to test the alternative hypotheses. We use the notation $\hat{P}_i(s) = n_i(s)/n_i$ to make the comparison to $P_i(s)$ where s is the submarket that contains product i . Since the cutoff level for a one-tailed significance level of 0.10 is 1.28, Table 5 suggests that the "Diesel vs. Gasoline" market structure is significantly better than no structure at the 0.10 level for all seven products. On the other hand, the Z-tests clearly suggest that the model-specific hierarchy is not significantly better than a hypothesis of no structure. The aggregate tests of $\hat{P}(s)$ versus $P(s)$ also indicate the "Diesel vs. Gasoline" partitioning of the market is significantly better than no structure.

In the above illustration we have used the statistics to test two alternative hypotheses

TABLE 5
Test Statistics for Alternative Structures
(n_i is the number in parentheses)

	"Diesel vs. Gasoline"			"Model-specific"		
	$\hat{P}_i(s)$	$P_i(s)$	Z	$\hat{P}_i(s)$	$P_i(s)$	Z
Peugeot	0.80*	0.39**	2.66***	0.10	0.17	-0.58
Diesel (10)						
Cutlass	0.70	0.31	3.77	0.20	0.13	0.93
Diesel (20)						
Jetta	0.60	0.35	2.02	0.20	0.12	0.95
Diesel (15)						
Cavalier	0.70	0.44	2.34	—	—	—
Gasoline (20)						
Peugeot	0.67	0.47	1.55	0.20	0.12	0.95
Gasoline (15)						
Cutlass	0.90	0.50	2.53	0.10	0.22	-0.92
Gasoline (10)						
Jetta	0.80	0.50	1.89	0.20	0.17	0.25
Gasoline (10)						
	$\hat{P}(s)$	$P(s)$	Z	$\hat{P}(s)$	$P(s)$	Z
Aggregate Test	0.72	0.41	6.10	0.18	0.15	0.84

*Note 0.8 equals fraction of Peugeot choosers who select a diesel Cutlass or Jetta ($0.8 = (2 + 6)/10$ from Table 1).

**by equation 3 ($0.39 = (0.2 + 0.15)/(1 - 0.1)$).

*** $(0.8 - 0.39)/[0.39(1 - 0.39)/10]^{1/2}$.

about market structure. One was retained; the other rejected. However, in some applications more than one hypothesis could be retained. This would be analogous to comparing two regressions based on different variables. It is tempting to accept the structure with the highest Z -statistic (the regression with the highest F), but in doing so one must realize it is not a rigorous use of the statistic. We recommend instead that one exercise the same caution normally exercised when choosing the "best" regression model. See, for example, discussions in Drake (1967), Green and Tull (1978), and Morrison (1976).

One can use the Z -statistic in theory testing mode as illustrated above or in exploratory mode where one uses the statistic to search for the "best" managerial competitive structure. In the latter case we suggest that the chosen hypothesis then be tested with either convergent methods, holdout samples, or both.

Use Hypotheses

Some firms position their products for particular uses and such use categories can become the basis for defining the structure of competition. For example, in the home cleaner market, cleaning the kitchen may be one submarket and cleaning the bathroom another. We would find a competitive structure defined by use if one set of products tends to be used for the kitchen and another set of products for the bathroom.

We can proceed to test use hypotheses in one or two ways. The first way is to collect separate data for each use and follow the procedure in Figure 2 to identify a structure within each use. Since this procedure is a simple extension, we need not illustrate it here.

The second procedure is more complex. Suppose the manager wishes to assign products to use groupings and then test those groupings to determine if consumers stay within groupings when choosing another product for the same use.

The first procedure is a set of multiple structures where, for each use, all products are assigned to submarkets. The second procedure is a single structure where each product is assigned to one group and each group corresponds to a use. We illustrate here the second procedure.

In the second procedure we begin by uniquely assigning products to submarkets. Let n_{iu} = the number of consumers who use product i for use u . One reasonable assignment rule is to assign product i to the use submarket u which maximizes n_{iu} .

We now consider deleting products. If we delete product i (which is assigned to use submarket, u) and ask consumers to consider only use u , then we would hope that consumers are more likely to choose again from products assigned to use submarket u than would be predicted by market share.

Testing this hypothesis mathematically is very similar to testing a market structure defined by brands or by product characteristics. There is one important difference. Suppose there are four cleaning products on the market, Ajax, 409, Top Job, and Mr. Clean. Suppose the above rule assigns Ajax and 409 to bathroom cleaning and assigns Top Job and Mr. Clean to kitchen cleaning. To test a hypothesized market structure we must decide how to deal with consumers whose choice violates the assignment rule. For example consumers who choose Ajax for kitchen cleaning are misclassifications according to our assignments.

For such consumers we choose to remove not Ajax but their most preferred "kitchen" product, where "kitchen" is defined by the procedure described above. In this case, we remove their most preferred of Top Job or Mr. Clean. Note that for all consumers "kitchen" or "bathroom" products are defined by the market structure we are testing. Such "misclassified" consumers will of course still prefer Ajax when either Top Job or Mr. Clean are removed because they preferred Ajax when all products

were available. Thus such consumers will be counted as evidence against the hypothesized structure. This will assure that any statistics we compute will be appropriately conservative. We now turn these verbal statements into mathematical statements. These definitions are similar to those defined earlier except that we now condition on usage, u . Let:

$n_{iu}(j)$ = the number of consumers who would purchase product i from the set of products designated for use u , and who would purchase product j for use u when product i is deleted. Note that $n_{iu}(j)$ is only defined for products, i , contained in usage product set, u .

$n_{iu}(u)$ = the number of consumers who would purchase product i from the product set u for use u , and who would purchase a product, other than i , from usage product set u if i were deleted. Note that $n_{iu}(u) = \sum_{j \in u} n_{iu}(j)$.

Then, if there is no structure and ACRM holds independent of use, we expect

$$\hat{P}_{iu}(u) \approx P_{iu}(u) \quad (19)$$

where $\hat{P}_{iu}(u) = n_{iu}(u)/n_{iu}$ and $P_{iu}(u) = (\sum_{j \in u, j \neq i} m_j)/(1 - m_i)$ where the summation is over all products assigned to use submarket, u , except product i . If submarkets exist we would expect

$$\hat{P}_{iu}(u) > P_{iu}(u) \quad \text{for } i \text{ in } u, \quad \hat{P}_{iu}(u) < P_{iu}(u) \quad \text{for } i \text{ not in } u.$$

The contribution of (19) is the definition of use structure relative to no structure. Statistical testing follows the same procedures as before. We can estimate $n_{iu}(u)$ by any of the four methods discussed above, if we incorporate the following modifications:

(i) data (choice, preference, logit probabilities, consideration) are collected for each use, and

(ii) forced switching is defined with respect to the first choice for each use.

In the above discussion we assigned products to submarkets based on their maximum use. Other assignment rules are possible. The association of products to use submarkets can also be made by factor analysis or cluster analysis of a matrix of the probability of purchase of each product (j) for each use (u). (See Day, Shocker, and Srivastava 1979 and Srivastava, Leone, and Shocker 1981.) In such a case the submarket set u would represent a composite of uses. Managerial assignment could also be used (e.g., products appropriate for breakfast versus those for lunch or for dinner). Note, however, in this composite use procedure we require each product to be assigned only to one submarket. Our first procedure which utilizes submarkets for use allows products to be in multiple submarkets. Consider now an example.

TABLE 6
Illustrative Example of Data to Test Use Competitive Structure

Bathroom Cleaning, $n_{iu}(u)$					
	n_{iu}^*	Ajax	409	Top Job	Mr. Clean
Ajax	250	—	200	25	25
409	150	100	—	25	25
Top Job				—	
Mr. Clean					—
Kitchen Cleaning, $n_{iu}(u)$					
	n_{iu}	Ajax	409	Top Job	Mr. Clean
Ajax		—			
409			—		
Top Job	200	25	25	—	150
Mr. Clean	200	25	25	150	—

* Note. n_{iu} = the number of consumers who would purchase product i for use, u .

TABLE 7
Test Statistics for Usage Submarkets

	$\hat{P}_{iu}(u)$	$P_{iu}(u)$	Z
Ajax (250)	0.80	0.27	17.9
409 (150)	0.67	0.38	7.4
Top Job (200)	0.75	0.33	12.7
Mr. Clean (200)	0.75	0.33	12.7
	$\hat{P}^*$	P^*	Z
Aggregate Test	0.75	0.32	25.8

Table 6 is an illustrative hypothetical example for the cleaning market. There are two uses, **kitchen** and **bathroom**, and four products, Ajax, 409, Top Job, and Mr. Clean. The hypothesis we wish to test is that Ajax and 409 are bathroom products while Top Job and Mr. Clean are kitchen products. Note that **kitchen** forced switching is only defined with respect to kitchen products and that **bathroom** forced switching is only defined with respect to bathroom products.

If the market shares were 0.31 for Ajax, 0.19 for 409, 0.25 for Top Job, and 0.25 for Mr. Clean then we get the Table 7 values for testing the usage submarkets. In this example, a usage based competitive structure explains the data significantly better than a hypothesis of no structure.

User Hypotheses

Many researchers have hypothesized that a market is divided into submarkets by user characteristics. For example, one might hypothesize that heavy beer drinkers (more than 3 beers per day) tend to use Miller Lite, Coors, or Budweiser, while light beer drinkers (2 or fewer beers per day) tend to use Schlitz, Anheuser Busch Natural, or Miller High Life.

To test user based hypotheses about competitive structure we create test statistics analogous to either procedure described for usage based hypotheses. The only difference in the second procedure is that the rows of table analogous to Table 6 are based only on the appropriate user group. For example, if a heavy user prefers Miller Lite, forced switching for him is with respect to Miller Lite. Forced switching with respect to Miller Lite is computed only with respect to heavy users. If a heavy user prefers Schlitz (a light user brand), forced switching is computed with respect to his first choice among heavy user brands. As in the usage statistics, these assignments assure that any statistics are appropriately conservative, e.g., the heavy user who prefers Schlitz will still prefer Schlitz thus providing evidence against a user based competitive structure.

Because the concepts of "use" and "user" are so similar we leave the details on "user" competitive structure to the reader. We note that to collect forced switching data for "use" or "user" competitive structures, it may be necessary to generate hypotheses based on judgment and/or prior research.

Exploratory Analysis of Compound Structure

All tests outlined above are tests of the aggregate competitive structure. Even the "user" competitive structure, which, at first glance, appears to be an assumption of compound structure, is a test of aggregate structure. Products are assigned to submarkets based on user attributes, but the market is described by one set of submarkets.

It is possible that different groupings of consumers could be characterized by different competitive structures. For example, perhaps heavy beer drinkers characterize competition by brand while light beer drinkers characterize competition by use

occasions (e.g., with or without guests). If such a compound hypothesis is formulated then the statistical testing of that hypothesis is straightforward. Simply apply the statistical tests within each grouping of consumers.

Identifying groupings, that is, assigning consumers to groups, is more problematic. Our null hypothesis of no structure is at the aggregate level. So, if we attempt to apply our statistics at the individual level, we may structure the market inappropriately.

We propose two alternative heuristics and caution the reader that these heuristics are developed for exploratory use only and for input to the managerial selection of the best description. Any groupings identified with these statistics must still be subjected to aggregate testing within the grouping. To retain rigor, a different sample should be used for testing from the one used for exploration. The heuristics are given below. Each assumes that a prior set of possible competitive structures has been identified. The raw data are based on the logit/preference intensity or evoked set equations for $P_{ci}(j)$, equations (15) and (18); each heuristic assigns consumers to one of the competitive structures. The rules can be applied for a target product, i , or for a weighted sum across i .

(1) Assign consumers to the competitive structure which maximizes $P_{ci}(s)$ for the submarket which contains i . $P_{ci}(s)$, as well as the definition of (s) , will of course vary across alternative competitive structures.

(2) Compare individual level product deletion probabilities to market level product deletion probabilities computed via ACRM. Assign consumers to the competitive structure with the largest difference, $P_{ci}(s) - P_i(s)$, for the submarket which contains i .

Applying these heuristics allows an exploratory diagnosis of the presence of different structures within subsets of the market.

Managerial Analysis

The final step in our testing procedure is the managerial decision of selecting the best structure to describe the competitive relationship in the market. As we described above, the managerial selection is based on judgment aided by the statistical measures. Managerial analysis is best discussed by example. Thus, we illustrate these issues more fully by describing the application of our testing procedure to the market for coffee. We then discuss the managerial implications of six other applications to date.

Empirical Application

Hypothesis Generation

Past research on market structure indicates several bases for competition in the market for coffee: brand, product attribute, usage, and user characteristic.

The most obvious alternative is a brand structure (see Figure 3a). Discussion with managers and focus groups indicated a candidate product attribute description based on dividing the market into six groups:

1. ground caffeinated
2. ground decaffeinated
3. instant caffeinated regular
4. instant caffeinated freeze dried
5. instant decaffeinated regular
6. instant decaffeinated freeze dried

In formulating a competitive structure based on use occasion, we considered the results of focus group discussions on coffee consumption. From these groups, we identified over 40 possible use scenarios. These were grouped by judgment into nine classes that were felt to span the use environment at a meaningful level of detail. They

were:

1. to start the day/with breakfast
2. between meals/daytime alone
3. between meals with others
4. with lunch
5. with supper
6. dinner with guests
7. in the evening
8. to keep awake in the evening and
9. on weekends

In a sample of 295 users of coffee, each respondent indicated the subset of the nine occasions that applied to him or her and identified the product he/she would evoke for each occasion (see data collection description below for more detail).

We aggregated the occasions into six classes to obtain reasonable sample sizes for each brand and occasion (see Appendix for the data). We factor analyzed the matrix of evoking of brands for specific uses (see Day, Shocker, and Srivastava 1979). Inspection of the matrix of consideration of brands for specific uses shows little variation in the proportions considering a given product across uses (see Appendix). The factor analysis resulted in one dimension of use where most uses loaded on that dimension ($\Lambda_1 = 6.45$ for first factor, $\Lambda_2 = 0.28$ for the second factor). Other analyses were conducted to identify possible use hypotheses. A diary panel record of uses of coffee was collected from the respondents. Each coffee serving was recorded for a one-week period, along with a description of the occasion. 54 percent of the respondents used only one brand over all occasions in the week and 10 percent more used one brand for all uses until it ran out and then switched to a new brand for all subsequent uses. (See Laurent 1978 for a more extensive discussion of this data.) The pantry check indicated 43 percent had only one container of coffee on hand and 60.6 percent had only one container open. There was no strong evidence of usage as the basis of product competition, but an alternative was formulated based on managerial priors and interpretation of the data above. This hypothesis was kept simple and was based on grouping the occasions into two classes by time of day of use (A.M. or P.M.). Brands were assigned to either the A.M. or P.M. group based on whether they were most heavily evoked for A.M. or P.M. use occasions (see Figure 3c).

A competitive structure alternative based on user characteristics was formulated by defining two submarkets based on user purchase rates (heavy—more than one purchase per two weeks, and light—one or fewer purchases per two weeks). Products were assigned to the submarkets where their consideration proportions were highest (see Figure 3d).

The four hypotheses indicated above were subjected to formal statistical testing, based on our product deletion criteria and statistics. We note that other managers and other researchers may have preferred different hypotheses and such hypotheses are viable candidates for future tests.

Data Collection

In accordance with measurement procedures outlined above, 295 users of coffee (those who drink more than one cup of coffee/per day at home) were interviewed in Springfield, Massachusetts, and Indianapolis, Indiana in July and August of 1977. Respondents were interviewed after being recruited in a shopping mall and quotas were set to assure at least 50 respondents used each major type of coffee (ground/instant, caffeinated/decaffeinated, freeze dried). Evoked uses, the products considered for each use, and the last product used were identified for each respondent. Respondents indicated the brands they would consider using for each of the nine usage scenarios that applied to them. Preferences for products for each use were obtained on

Alternative a: Brand

	$\hat{P}^* = 0.07$ $P^* = 0.12$ $Z = -2.7$					
	Maxwell House	Tasters' Choice	Sanka	Brim	Folgers	Nescafe
Brand						
Structure:	$\hat{P}(s) = 0.07$	$\hat{P}(s) = 0.05$	$\hat{P}(s) = 0.08$	$\hat{P}(s) = 0.09$	$\hat{P}(s) = 0.04$	$\hat{P}(s) = 0.06$
No structure:	$P(s) = 0.18$	$P(s) = 0.11$	$P(s) = 0.12$	$P(s) = 0.04$	$P(s) = 0.02$	$P(s) = 0.04$
	$Z = -2.9$	$Z = -1.3$	$Z = -0.8$	$Z = 1.3$	$Z = 0.9$	$Z = -0.5$

Alternative b: Product Attribute

	$\hat{P}^* = 0.32$ $P^* = 0.11$ $Z = 11.3$					
	Ground Caffeinated	Ground Decaffeinated	Instant Caffeinated Regular	Instant Caffeinated Freeze Dried	Instant Decaffeinated Regular	Instant Decaffeinated Freeze Dried
	$\hat{P}(s) = 0.47$	$\hat{P}(s) = 0.19$	$\hat{P}(s) = 0.37$	$\hat{P}(s) = 0.23$	$\hat{P}(s) = 0.28$	$\hat{P}(s) = 0.25$
	$P(s) = 0.11$	$P(s) = 0.03$	$P(s) = 0.13$	$P(s) = 0.09$	$P(s) = 0.09$	$P(s) = 0.13$
	$Z = 8.4$	$Z = 4.1$	$Z = 5.5$	$Z = 3.3$	$Z = 4.1$	$Z = 2.4$

Alternative c: Use

	$\hat{P}^* = 0.40$ $P^* = 0.46$ $Z = -1.5$	
	Morning (A.M.)	Afternoon & Evening (P.M.)
Use Structure:	$\hat{P}(s) = 0.40$	$\hat{P}(s) = 0.41$
No Structure:	$P(s) = 0.44$	$P(s) = 0.48$
	$Z = -1.5$	$Z = -2.8$

Alternative d: Users

	$\hat{P}^* = 0.28$ $P^* = 0.59$ $Z = -8.69$	
	Heavy	Light
User Structure:	$\hat{P}(s) = 0.34$	$\hat{P}(s) = 0.11$
No Structure:	$P(s) = 0.75$	$P(s) = 0.18$
	$Z = -9.3$	$Z = -1.2$

FIGURE 3. Hypothesis Testing for the Coffee Market.

a seven-point scale (extremely well-liked to very much disliked). Brands were rated on 12 product attribute scales. After providing demographic data and answering questions on coffee consumption, respondents were given an opportunity to purchase coffee for their most frequent use with a two dollar coupon they were given as compensation for participating in the interview. When the respondents reached the shelf, they found their first preference product "out of stock." Eighty-five percent of the people made a purchase in the lab and 70 percent noticed their favored brand was missing. At the close of the lab phase, respondents were requested to participate in a usage panel in which they would record in a diary for a week each cup of coffee served at home (when, kind, brand, who present, how many cups, who prepared). Sixty percent

returned a complete diary. Two weeks after the lab, respondents were called back to determine their home inventory of coffee (kinds, brand, open or not, size of package) and recent purchases. The panel and pantry check were conducted in this application to provide additional insight into the effects of the use situation on product choice. In most applications, the evoking and preference by use would be sufficient to determine if use was the best basis for structure in the market.

Hypothesis Testing

We need to calculate the number of respondents who prefer product i and who will purchase in submarket s if brand i is deleted [$n_i(s)$]. We do our initial testing based on the use of $P'_c(i)$ estimated by the preference rank data method outlined above.

The estimated values were good based on the information theoretic test¹³—80 percent of the total uncertainty was explained ($U^2 = 0.795$). The standard deviation between actual and predicted market shares was 0.009. In this application, the McFadden, Train, Tye (1977) residual test indicated that the estimated probabilities [$P'_c(i)$] were not subject to error due to independence of irrelevant alternatives.¹⁴ After comparisons based on preference rank probabilities, the laboratory store shopping data was used for validation.

The first alternative tested is a brand structure (see Figure 3a). The number of respondents, n_i , choosing brand i was based on first preference. The number of respondents buying again in that submarket when their most preferred brand was not available ($n_i(j)$) was then calculated with equation (15) summed across consumers preferring brand i . Summary statistics were calculated by equations (9) through (12). The ACRM probabilities were calculated by equations (2) and (3) where market shares relevant to the sample were used as scale values (m_i) and the aggregate ACRM proportions were calculated by equations (13) and (14). In general respondents were less likely to repeat purchase another variant of brand than the null hypothesis of no structure would predict (see Figure 3a). The brand structure does not adequately describe the market because the probabilities ($\hat{P}, \hat{P}(s)$) are low and the hypothesized brand structure is not significantly better than the null hypothesis of no structure. (Z less than zero for one-tailed test.)

A product attribute description of the market was developed by dividing the products into the six groups described above. This competitive structure was statistically superior to the null hypothesis of no structure in explaining observed market behavior (see Figure 3b). All of the submarket switching probabilities are significantly higher than the probabilities predicted by no structure at the five percent level. The overall test is significant at the one percent level ($Z = 11.3$). Respondents were more likely to buy another brand of the product with a specific attribute than a model of no structure would suggest. For example 0.47 of those who have a first preference for a brand of caffeinated ground coffee would purchase another caffeinated ground coffee if their most preferred brand were not available. The “no structure” probability is 0.11 and this difference is significant at the one percent level ($Z = 8.4$).

The third hypothesis is based on use of coffees (see Figure 3c). The Z -statistics indicate that the use grouping is not significantly better than the null hypothesis of no structure. In fact, since the Z -statistics are negative they suggest qualitatively that the use grouping in Figure 3c actually does worse than the hypothesis of no structure in

¹³ U^2 is based on an information theoretic interpretation of the uncertainty explained by the individual level choice probabilities. The denominator is the uncertainty (entropy) that would be explained by “perfect information,” the numerator is the uncertainty explained by the probabilistic model. Thus $U^2 < 1$. For derivations and examples see Hauser (1978).

¹⁴ An addendum, available from the authors, describes the results of this test and an analysis with hierarchical logit and compound structure based logit.

TABLE 8
Variation in Competitive Structure by Consumer

Dominant Structure	Number in Which One Structure Dominates	$\hat{P}^*$
1. Ground/Instant	84	0.94
2. Caffeinated/Decaffeinated	53	0.94
3. Brand	1	0.85
4. No dominant two-way structure	96	

explaining consumer behavior. Thus this submarket hypothesis does not appear to be a good overall basis for identifying the competitive structure of products in the market for coffee.

The user hypothesis tested here (see Figure 3d), likewise, does not appear to be a good overall basis for defining competitive structure in the market for coffee.

Before we select a structure for managerial use, the data are further analyzed to explore the effects of compound structure, the use of probabilities based on consideration sets, and the application of other market structure methods.

Exploratory Testing of Compound Structure

We explored compound structure by calculating individual probabilities for various submarket structures. The assignment of individuals to the best fitting competitive structure led to the identification of some differences in consumers' views of the competition. Table 8 shows the number of people who (heuristically) fit one structure better than others. The values of $\hat{P}^*$ are high for groups described by ground/instant and caffeinated/decaffeinated structures. Based on our heuristic classification, the compound structure fits better than in the homogeneous case which has a $\hat{P}^* = 0.73$ for an overall ground/instant structure and $\hat{P}^* = 0.64$ for a caffeinated/decaffeinated structure. This is, of course, subject to confirmatory testing since the heuristic procedure is an attempt to maximize $\hat{P}^*$. These two groups may represent significant heterogeneity while the remaining 60 percent of the sample are probably not better described by one substructure than others. (Eighty-four people were equally well described by Alternatives 1 and 2, six people by Alternatives 2 and 3, and six people by Alternatives 1, 2, and 3.) Although some compound structure is evident, most compound structure seems to be simplified versions of Alternative *b*. Thus, we choose the overall attribute structure shown in Figure 3b as an adequate managerial summary description of the market.

Consideration Sets

In our method competition structure can derive significance from consumer preferences and/or the compositions of consideration sets. In order to get an indication of the relative magnitude of these two effects we re-estimated the individual product deletion probabilities assuming equal preference across the products in each consumer consideration set after the most preferred had been removed. See equation (18), hybrid technique. These values are shown in Table 9 for the six submarkets. They are very similar to preference based probabilities and overall significance level drops only slightly ($Z = 11.3$ to $Z = 10.1$). In this case the consideration set contains much of the competitive structure information.

Comparison to Other Methods

The final phase of exploratory analysis of the data was the examination of the results of applying two alternative methods. Kalwani (1979) has applied hierarchical

TABLE 9
Probabilities Based on Consideration Sets— $\hat{P}(s)$

	With Preference	Consideration Set Only
Ground/Caffeinated	0.47	0.44
Ground/Decaffeinated	0.19	0.18
Instant/Caffeinated/Regular	0.37	0.38
Instant/Caffeinated/Freeze Dried	0.23	0.16
Instant/Decaffeinated/Regular	0.28	0.27
Instant/Decaffeinated/Freeze Dried	0.25	0.22
Overall	0.32	0.30

clustering procedures to our forced switching data. He found the best definition to be ground versus instant coffee with a second level division into caffeinated and decaffeinated variants. He tested several measures of proximity based on switching: (1) "last purchase" from "purchase previous to last purchase," (2) "last purchase" to "next purchase" (obtained in call-back), and (3) forced switching in laboratory. In all clustering, the cophenetic correlations were high (0.814 to 0.887). No statistical basis for determining the best partitioning was available, but Kalwani felt that the results based on forced switching were superior since they tended to cluster products together that shared a common attribute and agreed with managerial priors. His results are consistent with the findings reported in our hypothesis testing. Kalwani also attempted to apply the Hendry model to the data. First he found the switching matrix did not meet the required equilibrium assumption. The observed switching was much below the theoretical switching (overall theoretical $k_w = 0.53$, empirical $k_w = 0.28$). Next he calculated values of the switching constant k_w for possible partitions. The reported values of the theoretical k_w 's were very similar for all alternatives (0.48 to 0.53) and uniformly above the observed values (0.12 to 0.28). He could not reach any conclusion on the hierarchical form based on the Hendry methodology applied to this data. Although the Hendry method was not informative here, it does not reject an attribute-based partitioning and Kalwani's clustering analysis is consistent with the product based on submarkets shown in Figure 3b.

Managerial Analysis

"Best" Structure. The first task is to define the best structure for competition. The formal hypothesis testing resulted in only one hypothesis that could explain the observed switching significantly better than the null hypothesis of no competitive structure. The exploratory analysis indicates compound structure is not likely to be a serious managerial concern and that the analysis of the data by a clustering methodology is consistent with the product attribute analysis. An analysis of the forced choice shopping data collected in the laboratory store allows a test of convergent validation of the measures used to test the product attribute hypothesis (see Figure 3b). In Table 10 the probabilities from the preference analysis, shopping measures, the pooled probabilities, and the "no structure" probabilities are shown. The shopping probabilities are substantially greater than the "no structure" values and the product attribute structure is significantly better than the null model based on both the shopping ($Z = 9.4$) and pooled probabilities ($Z = 14.8$). Based on the pooled data for the attribute submarkets, we examined the significance of an attribute structure for each brand where at least 10 consumers preferred most that brand (see Table 11). In nine of the ten cases, the tests indicate significance at least at the 10 percent level. Thus, from the perspective of the significant brands, we select the best description of the competitive structure of the

TABLE 10
Pooling of Laboratory Shopping Forced Switching and Individual Logit-Based Probabilities

Submarket	$\hat{P}(s)$			No Structure ($P(s)$)	Z-Statistic Pooled Data Versus No Structure
	Individual	Shopping	Pooled		
Ground/Caffeinated	0.47	0.47	0.47	0.11	11.4
Ground/Decaffeinated	0.19	0.20	0.20	0.03	5.1
Instant/Caffeinated/Regular	0.37	0.28	0.33	0.13	6.2
Instant/Caffeinated/Freeze Dried	0.23	0.25	0.24	0.09	4.9
Instant/Decaffeinated/Regular	0.28	0.18	0.24	0.09	4.2
Instant/Decaffeinated/Freeze Dried	0.25	0.32	0.28	0.13	4.1
Aggregate Z Statistic	11.3	9.4	14.8		

coffee market based on product attributes for managerial use in determining marketing strategy.

Managerial Implications. The appropriate product strategy for a company depends, in part, upon which products the firm now offers. Table 12 shows the existing brands (1977) of three major manufacturers. If Nestlé is considered, there is a gap in the product line coverage since they now offer no ground coffees. Our analysis suggests ground coffee is a separate market from instant. Substitution between a Nestlé ground coffee and Nestlé brands of instant coffee would be low based on this definition of competition. Our tests suggest that Nestlé should reconsider its coverage of submarkets. In 1977 Nestlé had a high share in the instant submarkets, but a zero market share in the ground coffee submarket. For strategic analysis, Nestlé might wish to calculate share with respect to the instant submarket rather than with respect to the entire coffee market.

The development of a new ground coffee may be a major opportunity for Nestlé. Its proven ability to market instant coffees suggests it has the capability to advertise, promote, and distribute a ground coffee. This opportunity is suggested by our analysis and should be subjected to further analysis. Perceptual maps of the ground,

TABLE 11
*Brand Level Significance Tests for Attribute Hierarchy**

	$\hat{P}_i(s)$	$P_i(s)$	Z
Ground/Caffeinated			
Maxwell House	0.49	0.10	11.4
Chock Full O'Nuts	0.60	0.14	5.3
Instant/Caffeinated/Regular			
Instant Maxwell House	0.33	0.10	6.5
Nescafe	0.37	0.19	2.1
Instant/Caffeinated/Freeze Dried			
Maxim	0.30	0.14	2.5
Tasters' Choice	0.21	0.06	4.5
Instant/Decaffeinated/Regular			
Nescafe	0.18	0.16	0.2
Sanka	0.25	0.08	4.9
Instant/Decaffeinated/Freeze Dried			
Sanka	0.31	0.14	2.3
Brim	0.35	0.14	2.8

*Table includes only brands where more than 10 consumers preferred most that brand.

TABLE 12
Brand Offerings by Selected Firms (1977)

	Ground		Instant			
	Caffeinated	Decaffeinated	Caffeinated		Decaffeinated	
			Regular	Freeze-Dried	Regular	Freeze-Dried
Nestle	No Brand	No Brand	Nescafe	Tasters' Choice	Nescafe	Tasters' Choice
General Foods	Maxwell House Yuban	Sanka Brim	Maxwell House Yuban	Maxim	Sanka	Sanka Brim
Proctor & Gamble	Folger's	No Brand	Folger's	No Brand	(High Point)	No Brand

() = in test market at time of study

caffeinated and decaffeinated submarkets could be drawn to more specifically define positioning opportunities and the potential of a new entry can be estimated. See examples in Urban, Johnson and Brudnick (1981) and Urban, Carter and Mucha (1983). These opportunity identification activities are useful first steps in the development of a new product offering. For more details on how to balance these considerations with sales potential, penetration, scale, input, reward, risk, and match to the organization's capabilities, see Urban and Hauser (1980, Chapter 5).

General Foods has a different strategic problem. Each submarket is covered—General Foods already has many coffee offerings. If Maxwell House and Yuban (Sanka and Brim) are not well differentiated within their submarkets, then General Foods may have too many offerings and should consider consolidating its brands. It could also direct its product development efforts at creating a new submarket. The methodology we propose only describes the current market situation. A new product can fit into a submarket or, possibly, create a new submarket. For example, perhaps General Foods could develop fully brewed coffee that is frozen in plastic cups and can be heated in a microwave oven. This may add a new submarket to the structure. This is an opportunity that could be pursued by product development and testing. As the idea is converted into a concept, product, and complete design, data would be collected to confirm the strategic opportunity and balance it against other considerations.

Proctor and Gamble has recently entered the coffee market with the Folgers brand. Their strategy appears to be to build their business by sequentially entering each submarket in order of the sales potential of the submarkets. High Point is now in the national market. A next entry for them to evaluate could be a freeze dried instant. The opportunity of a ground decaffeinated coffee could also be evaluated.

Other Applications

The coffee example shows how the testing procedure can be used to define the basis for market share, identify opportunities for new products, and provide a structure to evaluate product strategy. We therefore call the model, measurement, and testing system PRODEGY (Product Strategy). It has been applied seven times in the past and other applications are now in process.

In one case on home cleaning products alternative competitive market hypotheses included forms of the products (e.g., spray, liquid, powder, foam, aerosol), user locations (e.g., kitchen, bathroom), user tasks (e.g., glass, counters, floor, metal

chrome), user intensity (e.g., light cleaning, heavy cleaning), and individual attributes (e.g., “sophisticated” cleaners who use many special purpose products versus “basic” cleaners who use a few generic products). The selected submarket structure was a “use” partitioning. It reflected a new perspective on the market and indicated a major gap in the firm’s product line. Previously, four attempts had been made to develop a new brand in this product class, but all failed based on pre-test market analysis. The new opportunity identified by PRODEGY was pursued and it succeeded in pre-test market analysis. It is now being successfully test marketed and the brand manager credits the competitive market structure as a major contributor to the product’s success.

In application to the beer market, brand, product form, and usage level structures were evaluated. The best partitioning was found to be consistent with the firm’s previous beliefs and reflected an empirical confirmation of its strategic assumption on competitive boundaries. The empirical confirmation of the firm’s prior views increased confidence in their belief that they were adequately covering the existing market for beer. They turned their development attention to creating a new partition in the market by searching for a new form of the product (e.g., super light beer).

In a study of detergents, a new competitive brand recently introduced in the market was placed in the structure. PRODEGY indicated to the firm that a pourable powder detergent was not creating a new subsection in the competitive structure. The new product was perceived as competing in an existing product form partition. The firm responded by repositioning their existing brand rather than spending large sums of money to introduce a major brand against the new entrant.

In the final consumer goods case, food products were studied. It was found that the “pre-packaged” and “deli” brands reflected separate submarkets in the category. The firm previously had refrained from selling a “deli” product because they felt it would compete with their line of “prepackaged” goods. After the study indicated this was not true, they began development of an entrant into the deli market with the prospect of major sales increases. Perceptual maps indicated the firm had the product strength to compete in the deli segment and concept testing was begun to test the transferability of the firm’s brand name from the pre-packaged to deli submarkets.

Two applications have been made to industrial products. In a study of financial decision support system software for a firm selling a sophisticated software product, the “no structure” hypothesis could not be rejected and the firm’s strategic space was described by a single perceptual map (Burrow and Burns 1982 and Borschberg and Elkins 1983). Financial planners viewed the major dimensions as “power” (math capability, large data bases, financial functions and consolidation capabilities), “ease of use” (manipulate data and develop models easily, understandable documentation, and easy to learn), and “vendor quality” (reputation and support). Planners considered both simple programs such as Visicalc or Supercalc and complex programs such as Express or XSIM for each specific use (e.g., long-range planning, budgeting, and financial reporting). A map was drawn by graphing the attribute levels per dollar. This suggested that complex products would not be successful unless they had both good power and ease of use per dollar. The firm commenced a program to improve its ease of use per dollar by developing menu driven products which used subsets of the full software systems capabilities and reduced its prices. It also committed to monitoring the market to see if partitions would develop as users became more knowledgeable and their problem solving needs increased.

Another industrial application was to heart pacemakers (Kasinkas 1982). In this application the customer was the doctor. For each of three patient symptomologies doctors indicated which brands and types of pacemakers they would consider (a list of 40 products was supplied) and their probability of implanting that device. Direct

estimates of probabilities of purchase (Juster 1966 and Morrison 1979) and the assumption of equal probabilities over the consideration set were used to estimate the product deletion probabilities. The use of laboratory forced switching was not feasible in this study. The pacemaker study found submarkets defined by product attributes (e.g. programmable versus nonprogrammable units) as a significantly better description of the market than "no structure." The implication for the manufacturers was not to drop its older simple nonprogrammable units because a submarket persists for it and to continue to develop more advanced units to capture the sophisticated submarket.

Summary and Research Issues

This paper has presented a model and measurement methodology to estimate the structure of competition. In seven applications it produced encouraging statistical significance and managerial insight; however, several issues deserve further research.

One definition has been proposed in this paper for the description of competitive structure of a market. A statistical test was derived based on this definition and several procedures were suggested to collect data for the tests. We provide empirical applications based on our methods, but it would be useful to conduct comprehensive comparative empirical studies across additional methods (e.g., overlapping clustering) and alternative data sources (e.g., panel and UPC data). This is a subject for future research.

A technical issue to be resolved in the model is what to do if only one product defines a competitive submarket. In this case, the model criterion based on forced switching is ambiguous. Although a single product submarket cannot be tested formally, intuitively one would be indicated if consumers would consider that product as the only acceptable alternative for a specific use. In terms of our measures, a consideration set of one brand and a large proportion of respondents refusing to buy in the laboratory store would indicate this condition. Research may improve procedures for identifying and testing single product branches, but, in practice, we rarely observe such submarkets because competition usually develops quickly if the initial product is successful.

The empirical success of using consideration sets to estimate individual probabilities opens the possibility of monitoring the competitive structure over time through low cost telephone surveys or UPC panels. The ability to represent the dynamics of the competitive relationships would be very useful in strategy formulation. It could allow low cost measurement of the emergence of new submarkets.

New applications are underway in several markets to further assess the empirical adequacy of the model and its managerial relevance. Our early applications suggest that the proposed model could be a useful tool to aid in market strategy formulation.¹⁵

Acknowledgement. We would like to acknowledge the very valuable comments on our work received from Al Silk, Len Lodish, and Manu Kalwani. Richard Brudnick did the empirical analysis of the compound structure. We wish to thank the editors and reviewers of *Marketing Science*. In particular, the proof in Footnote 8 is due to a reviewer.

Appendix. Use and Brand Consideration by Occasion (Coffees)

In interviews with 295 coffee drinkers (greater than one cup per day), 808 uses were evoked across six major use classes (average is 2.7 uses per person). The table shows the proportion of the respondents who evoked each use and the proportion who would consider a given brand for those who evoked a given use. For example, 8.5 percent of the respondents who evoked breakfast as a use would consider Brim (Instant) as a product for this use.

¹⁵This paper was received July 1981 and has been with the authors for 3 revisions.

	Breakfast	Day Alone	Day Others Present	Lunch	Supper	Evening
Percent Evoking Use	96.3	41.7	23.4	30.2	43.7	43.1
Percent Who Would Consider Brand:						
Brim (Instant)	8.5	6.5	12.1	7.9	8.4	7.1
Folgers (Instant)	6.3	6.5	3.0	6.7	6.7	6.3
Folgers (Ground)	7.7	9.8	9.1	3.4	5.9	7.9
Maxwell House (Instant)	32.0	26.0	40.9	36.0	33.6	27.6
Nescafe	10.6	11.4	18.2	9.0	8.4	15.0
Nescafe (Decaf)	10.6	8.9	6.1	12.4	10.1	7.1
Maxim	12.0	15.4	7.6	13.5	13.4	11.8
Tasters' Choice	20.1	18.7	24.2	20.2	18.5	22.0
Tasters' Choice (Decaf)	13.0	17.1	16.7	20.2	17.6	17.3
Sanka (Instant)	22.9	22.0	18.2	21.3	20.2	24.4
Sanka (Freeze Dried)	7.7	8.1	9.1	12.4	11.8	3.9
Sanka (Ground)	6.3	7.3	12.1	5.6	9.2	4.7
Chock Full O'Nuts	10.6	8.9	10.6	6.7	8.4	9.4
Hills Brothers	7.0	4.1	1.5	3.4	6.7	1.6
Maxwell House (Ground)	28.9	32.5	24.2	29.2	32.8	26.0

References

- Arabie, P., J. D. Carroll, W. De Sarbo and J. Wind (1981), "Overlapping Clusters: A New Method for Product Positioning," *Journal of Marketing Research*, 18, 3 (August), 310-317.
- Batsell, R. R. (1980), "A Market Share Model Which Simultaneously Captures the Effects of Utility and Substitutability," Working Paper 80-007, Wharton School, University of Pennsylvania, Philadelphia, Pa., March.
- Bell, D. E., R. C. Keeney, and J. D. C. Little (1975), "A Market Share Theorem," *Journal of Marketing Research*, 12, 2 (May), 136-41.
- Bettman, J. R. (1979), *An Information Processing Theory of Consumer Choice*, Reading, Mass: Addison-Wesley.
- (1971), "Toward a Statistics for Consumer Decision Nets," *Journal of Consumer Research*, 1, 1 (June), 51-62.
- Borschberg, A. and W. Elkins (1983), "Defensive Marketing Strategy: Its Application to a Financial Decision Support System," unpublished Master's Thesis, Alfred P. Sloan School of Management, MIT.
- Bourgeois, J. C., G. H. Haines and M. S. Sommers (1979), "Defining Brand Space," Working Paper, University of Toronto, Toronto, Canada.
- Burrow, B. and I. Burns (1982), "A Strategic Analysis of the Market for Financial Decision Support Software," unpublished Master's Thesis, Alfred P. Sloan School of Management, MIT.
- Butler, D. H. (1976), "Development of Statistical Marketing Models," in *Speaking of Hendry*, Croton-On-Hudson: Hendry Corporation.
- Charnes, A., W. W. Cooper, D. B. Learner and F. Y. Phillips (1979), "A Theorem and Approach to Market Segmentation," Working Paper, University of Texas, Austin, November.
- Day, G. S., A. D. Shocker and K. Srivastava (1979), "Consumer Oriented Approaches to Identifying Product Markets," *Journal of Marketing*, 43, 4 (Fall), 8-20.
- Drake, A. W. (1967), *Fundamentals of Applied Probability Theory*, New York: McGraw-Hill Book Co.
- Ehrenberg, A. S. C. and G. J. Goodhardt (1982), "The Dirichlet: A Comprehensive Model of Buying Behavior," Working Paper, London Business School and City University, London, May.
- Frank, R. E., W. F. Massy and Y. Wind (1972), *Market Segmentation*, Englewood Cliffs, N.J.: Prentice-Hall.
- Fraser, C. and J. W. Bradford (1983), "Competitive Market Structure: Principal Partitioning of Revealed Substitutabilities," *Journal of Consumer Research*, 10, 1 (June), 15-30.
- Green, P. E. and D. S. Tull (1978), *Research for Marketing Decisions*, 4th Ed., Englewood Cliffs, N.J.: Prentice-Hall.
- Haines, G. H. (1974), "Process Models of Consumer Decision Making," in G. D. Hughes and M. L. Ray

- (Eds.), *Buyer/Consumer Information Processing*, Chapel Hill, N.C.: University of North Carolina Press.
- Hauser, J. R. (1978), "Testing the Accuracy, Usefulness, and Significance of Probabilistic Choice Models: An Information Theoretic Approach," *Operations Research*, 26, 3 (May-June), 406-421.
- and A. Tversky (1983), "Agendas and Choice Probabilities," Presentation at Conference on Choice Modeling (Duke University, June 1981) and Working Paper, MIT, revised May.
- Jeuland, A. P., F. M. Bass and C. P. Wright (1980), "A Multibrand Stochastic Model Compounding Heterogeneous Erlang Timing and Multinomial Choice Processes," *Operations Research*, 28, 2 (March-April), 255-277.
- Juster, F. T. (1966), "Consumer Buying Intentions and Purchase Probability: An Experiment in Survey Design," *Journal of American Statistical Association*, 61, 658-696.
- Kalwani, M. U. (1979), "Analysis of Competition in Consumer Markets: A Hierarchical Approach," Working Paper, Alfred P. Sloan School of Management, MIT, Cambridge, Mass.
- and D. Morrison (1977), "A Parsimonious Description of the Hendry System," *Management Science*, 23, 5 (January), 467-477.
- Kasinkas, M. (1982), "Competitive Structures in Industrial Markets: A Hierarchical Modeling Approach to Segmentation in the U.S. Pacemaker Market," Master's Thesis, Alfred P. Sloan School of Management, MIT.
- Laurent, G. (1978), "A Study of Multiple Variant Consumption for Frequently Purchased Consumer Products," unpublished Ph.D. Thesis, Alfred P. Sloan School of Management, MIT, Cambridge, Mass.
- Luce, R. D. (1959), *Individual Choice Behavior*, New York: John Wiley & Sons, Inc.
- Lussier, D. A. and R. W. Olshavsky (1979), "Task Complexity and Contingent Processing in Brand Choice," *Journal of Consumer Research*, 6 (September), 154-165.
- McFadden, D. (1980), "Econometric Models of Probabilistic Choice," in *Structural Analysis of Discrete Data with Econometric Applications*, F. Manski and D. McFadden (eds.), Cambridge, Mass.: MIT Press.
- , K. Train and W. Tye (1977), "An Application of Diagnostic Tests for the Independence of Irrelevant Alternatives Property of the Multinomial Logit Model," *Transportation Research Record*, 637, 39-45.
- Morrison, D. F. (1976), *Multivariate Statistic Methods*, 2nd Ed., McGraw-Hill Book Co.: New York.
- Morrison, D. G. (1979), "Purchase Intentions and Purchase Behavior," *Journal of Marketing*, 43, 2 (Spring), 65-74.
- Payne, J. W. (1976), "Task Complexity and Contingent Processing in Decision Making: An Information Search and Protocol Analysis," *Organization Behavior and Human Performance*, Vol. 16 (June), 355-387.
- Rao, V. R. and D. J. Sabavala (1981), "Inferences of Hierarchical Choice Processes from Panel Data," *Journal of Consumer Research*, 8, 1 (June), 85-96.
- Rubinson, J. R., W. F. Vanhonacker and F. M. Bass (1980), "On 'A Parsimonious Description of the Hendry System'," *Management Science*, 26, 2 (February), 215-226.
- Silk, A. J. and G. L. Urban (1978), "Pre-Test Market Evaluation of New Packaged Goods: A Model and Measurement Methodology," *Journal of Marketing Research*, 15, 2 (May), 171-191.
- Srivastava, R. K., R. P. Leone and A. D. Shocker (1981), "Market Structure Analysis: Hierarchical Clustering of Products Based on Substitution in Use," *Journal of Marketing*, 45, 3 (September), 38-48.
- , A. D. Shocker and G. S. Day (1978), "An Exploratory Study of the Influence of Usage Situation on Perceptions of Product Markets," in H. K. Hunt, (Ed.), *Advances in Consumer Research*, Vol. 5, Ann Arbor: Association for Consumer Research, 37-8.
- Stefflre, V. (1972), "Some Applications of Multidimensional Scaling to Social Science Problems," in *Multidimensional Scaling: Theory and Applications in the Behavioral Sciences*, Vol. 2, A. K. Romney, R. N. Shepard and S. B. Nerlove (Eds.), New York: Seminar Press.
- Tversky, A. (1972), "Elimination by Aspects: A Theory of Choice," *Psychological Review*, 79, 4, 281-299.
- and S. Sattath (1979), "Preference Trees," *Psychological Review*, 86, 6, 542-573.
- Urban, G. L. and J. R. Hauser (1980), *Design and Marketing of New Products*, Englewood Cliffs, N.J.: Prentice-Hall.
- , P. L. Johnson, and R. H. Brudnick (1981), "Market Entry Strategy Formulation," Working Paper 1103-80, Alfred P. Sloan School of Management, MIT, January.
- , T. Carter and Z. Mucha (1983), "Market Share Rewards to Pioneering Brands: An Exploratory Empirical Analysis," in H. Thomas and R. Gardner (Eds.), *Marketing Strategy*, John Wiley: New York.
- Vanhonacker, W. R. (1979), "Brand Switching and Market Partitioning: An Empirical Study of Brand Competition," unpublished doctoral dissertation, Purdue University, West Lafayette, Ind., December.
- (1980), "The Hendry Partitioning Methodology: Theory and Applications," Working Paper 288A, Columbia University, New York, January.