

A MEASUREMENT ERROR APPROACH FOR MODELING CONSUMER RISK PREFERENCE*

JEHOSHUA ELIASHBERG AND JOHN R. HAUSER

The Wharton School, University of Pennsylvania, Philadelphia, Pennsylvania 19104
*Sloan School of Management, Massachusetts Institute of Technology, Cambridge,
Massachusetts 02139*

Von Neumann-Morgenstern (vN-M) utility theory is the dominant theoretical model of risk preference. Recently, market researchers have adapted vN-M theory to model consumer risk preference. But, most applications assess utility functions by asking just n questions to specify n parameters. However, any questioning format, especially under market research conditions, introduces measurement error. This paper explores the implications of measurement error on the estimation of the unknown parameters in vN-M utility functions and provides procedures to deal with measurement error.

We assume that the functional form of the utility function, but not its parameters, can be determined a priori through qualitative questioning. We then model measurement error as if question format and other influences cause the consumer to choose the unknown "risk parameter" from a probability distribution and to make his decisions accordingly. We provide procedures to estimate the unknown parameters when the measurement error is either (a) Normal or (b) Exponential.

Uncertainty in risk parameters induces uncertainty in utility and expected utility, and hence uncertainty in choice outcomes. Thus, we derive the induced probability distributions of the consumer's utility and the estimators for the implied probability that an alternative is chosen.

Results are obtained for both the standard decision analysis "preference indifference" question format and for a "revealed preference" format in which the consumer is asked simply to choose between two risky alternatives.

Since uniattribute functions illustrate the essential risk preference properties of vN-M functions, we emphasize uniattribute results. We also provide multiattribute estimation procedures. Numerical examples illustrate the analytical results.

(MARKETING; UTILITY THEORY; RISK MODELING)

1. Perspective

The measurement and modeling of how consumers form preferences among risky alternatives is becoming an important problem in marketing science as researchers begin to focus on purchases of durable goods such as automobiles, home heating systems, home computers, and major appliances. An integral part of such consumer decisions is the choice of a specific product, say a gas furnace, when the attributes of the product, say annual cost and reliability, are not known with certainty.

A number of procedures have been proposed to model consumer risk. For example, Pras and Summers (1978) include the standard deviation of an attribute as a risk measure. Among these procedures is explicit risk assessment with von Neumann-Morgenstern (vN-M) utility functions. vN-M utility functions have the advantages that they are:

(1) theoretically derived from an axiomatic base (von Neumann and Morgenstern 1947, Friedman and Savage 1952, Herstein and Milnor 1953, Jensen 1967, Marschak 1950, and others),

(2) provide a set of practical functional forms derived from testable behavioral assumptions (see review in Keeney and Raiffa 1976), and,

* Accepted by Donald G. Morrison, Special Departmental Editor; received June 9, 1983. This paper has been with the authors 2 months for 2 revisions.

(3) have been applied extensively to model managers' decisions (see extensive reviews in Farquhar 1977 and Keeney and Raiffa 1976).

However, until recently vN-M utility functions have not achieved widespread use in marketing. This reluctance by marketing academics and practitioners stems in part because the question formats can be difficult and because the consumer modeling has not acknowledged measurement error as have more widely accepted techniques such as conjoint analysis (Green and Srinivasan 1978) and logit analysis (McFadden 1980). For example, both Hauser and Urban (1979) and Eliashberg (1980) have successfully modeled consumer preferences and have forecast reasonably well with vN-M theory, but both studies use the decision analysis procedure which requires complex questions to first test behavioral assumptions and then obtain exactly n observations to fit n parameters.

The consumer preference modeling task is different from the decision analysis task. Market research interviews are usually severely limited in time, hence, tradeoffs must be made among interviewee training, assumptions testing, complexity of questions, and the number of questions. Marketing researchers/scientists often prefer to ask more but simpler questions to statistically infer properties and estimate parameters. Such procedures must acknowledge potential measurement error.

More recently, marketing scientists have recognized these issues and have begun to adapt vN-M theory to marketing problems. For example, Ingene (1981) uses a Taylor series expansion to obtain simpler functional forms which are estimable with linear regression. Currim and Sarin (1984) provide two approaches. In the first, they adapt conjoint analysis to vN-M functions and in the second they retain the standard vN-M preference indifference format but use linear programming to minimize the stress of fit. All three approaches have practical merit and indicate the renewed interest in vN-M utility modeling.

In this paper, we take a different approach to the marketing problem. We explicitly acknowledge measurement error, but retain the axiomatic base and powerful, practical functional forms of vN-M theory. In the face of measurement error, we develop procedures to estimate unknown parameters for vN-M utility functions and we examine the implications of such measurement error on the utility functions and the choice outcomes.

2. Conceptualization of Measurement

The primary advantage of vN-M utility theory is its ability to model risk preferences. Basically, products are represented by their attributes and uncertainty (risk) is modeled as a probability distribution over the attributes. The vN-M function assigns a scalar value to every possible outcome of the uncertain attributes such that the consumer will prefer the product which has the maximum expected utility. The axioms imply that such a utility function exists and is unique (subject to a scaling change). The market research task is to obtain an estimate of this function such that expected utility is a reasonable predictor of the consumer's behavior.¹

In general, a vN-M utility function can be an arbitrary function, but research in the last 20 years has identified a set of parametered functions based on reasonable behavioral assumptions. These functions are valuable for market research because they allow us to parameterize, and hence simplify, the estimation problem and because they focus our attention on functional forms that can be justified a priori with a qualitative analysis of the consumer's risk preference.

¹ In marketing research, measurement error exists. Thus, we rarely can predict with certainty and instead forecast choice probabilities. Predictions of choice probabilities require modification of the vN-M axiom system. For one set of revised axioms, see Hauser (1978).

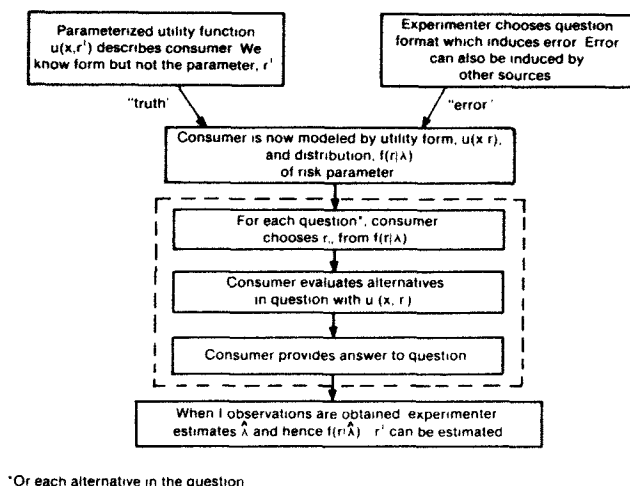


FIGURE 1. Conceptualization of Error Modeling.

We assume that this qualitative analysis has been carried out and that we know the functional form of the consumer's utility function. We do not know and would like to determine the unknown parameter(s) that characterize the degree of the consumer's attitude toward risk. Common functional forms and the interpretations of risk parameters are reviewed in §3. Keeney and Raiffa (1976, pp. 191–193) provide details on the qualitative analysis and Hauser and Urban (1979, Figure C) provide a market research example.

We can conceptualize the market research measurement as shown in Figure 1. We, the experimenter, choose a set of questions. The type of question chosen as well as other factors could well induce biases and errors in the measurement process. For example, Hershey, Kunreuther and Schoemaker (1982) found that the domain of outcomes (e.g., pure loss versus mixed lottery) and the decision context (e.g., abstract versus concrete formulation) may be influential in the observation of the consumer's risk attitude. In our framework, this may influence the parameterized utility function describing the consumer, but for a given utility function, there is some true risk parameter, r^T , and our questioning process induces error when we try to assess r^T . We describe this error by a probability distribution, $f(r|\lambda)$, of the risk parameter, r , where λ is a parameter of the distribution.

We then model the consumer's response as if he chooses a utility function, $u(x, r)$, draws a risk parameter, r_i , from $f(r|\lambda)$ independently for each question,² and provides an answer to the question that is consistent with $u(x, r_i)$. When we obtain I observations, it is our task to estimate $f(r|\lambda)$, or more specifically, λ . If errors are unbiased (zero mean) or if the bias is known, we can then obtain an estimator of r^T .

The assumption of error induced by question format or by other sources such as temporal variation, approximation, etc., and its modeling through random draws of the risk parameter is similar to "random utility" error theories such as Thurstone (1927) or Luce and Suppes (1965), but modified to emphasize the strength of vN-M theory—risk preference.

We note that our model of the consumer's response (dotted box in Figure 1) is a paramorphic model, that is, we assume that the consumer responds as if he follows the postulated procedure. Such details of cognitive response are inherently unobservable (without introducing new observation errors), but serve to provide a modeling frame-

²Or for each product which he evaluates in answering the question.

work with which to represent measurement error. In one interpretation, our assumption acts as a surrogate for explicit modeling of errors due to misspecification of the attributes, exogenous influences, task factors (e.g., problem framing), purchase situation variation, and other unobserved error sources.

To analyze the implications of Figure 1, we investigate a number of issues. (1) We obtain methods to estimate $\hat{\lambda}$, and hence $f(r|\hat{\lambda})$, from data obtained from standard decision analysis indifference questions. (We allow λ to be vector valued.) (2) We obtain methods to estimate $\hat{\lambda}$ from revealed preference questions where the consumer is given two alternatives and asked to choose his most preferred. (3) Since uncertainty in r induces uncertainty in $u(x, r)$, we derive the distribution of utility from the estimated distribution of r . (4) Since uncertainty in utility induces uncertainty in expected utility and hence uncertainty in choice outcomes, we derive expressions for the probability that a given alternative is chosen by the consumer. We investigate these issues for alternatives represented by discrete (Bernoulli) distributions of the attribute, x , and for alternatives represented by continuous distributions (e.g., Normal) of the attribute x . Before we begin the formal development we provide a brief review of vN-M concepts.

3. Review of vN-M Concepts

This section briefly reviews some aspects of vN-M utility theory that are necessary for our analyses. It may be skipped by readers familiar with vN-M theory. For greater detail see Keeney and Raiffa (1976).

Unattributed Functions

Unattributed functions are derived from assumptions about how a consumer's risk preference changes as his "assets" increase. For example, we might expect a consumer to be less concerned about uncertainty of $\pm \$100$ in heating bills if his current base heating bill were \$3,000 than he would be if his base heating bill were \$300. Pratt (1964) proposed a measure, called absolute risk aversion, $R(x)$, of how a consumer's risk attitude varies with his asset level, x . If $u(x)$ is the utility function, $R(x)$ is given by

$$R(x) = - \frac{d^2 u(x)}{dx^2} \bigg/ \frac{du(x)}{dx}. \quad (1)$$

If $R(x)$ is positive, the consumer is risk averse, if $R(x)$ is negative, risk prone, and if $R(x)$ is zero, risk neutral. Larger absolute values of $R(x)$ imply greater risk aversion (proneness).

A related concept is proportional risk aversion, $S(x)$, which measures a consumer's risk preference when consequences are measured in proportion to current assets. For example, if the uncertainty in heating bills were $\pm 10\%$ of the base bill then the proportional risk aversion measure would be appropriate to describe the consumer's risk attitude. If x_0 is the minimum (reference) value of x , then $S(x)$ is given by:

$$S(x) = (x - x_0)R(x). \quad (2)$$

The most common unattributed functional forms are based on constant $R(x)$ or $S(x)$. As Table 1 indicates, constant $R(x)$ implies an exponential function and constant $S(x)$ implies a power function. A third functional form, linear utility, is a special case when $R(x) = S(x) = 0$. This is the risk neutral form which applies when risk does not affect the consumer's decisions.

Other unattributed functional forms are possible, for example, a logarithmic form or a sum of exponential forms, but the three functions in Table 1 are the functional forms that have dominated applications in decision analysis and marketing science. Furthermore, in reviewing 30 applications, Fishburn and Kochenberger (1979) found that the constant $R(x)$ and constant $S(x)$ functional forms fit the data quite well and substantially better than the linear form.

Multiattributed Functions

Multiattributed functions are derived from assumptions about utility and preference independence (or dependence) among attributes. Empirical experience in decision analysis and marketing science has found them to be feasible and useful. We return to the multiattributed issue in §5 where we provide an example based on the commonly used multilinear form.

TABLE I
Common Uniattributed vN-M Utility Functions

Behavioral Assumption	Functional Form	Range of Attribute
1. Constant absolute risk averse ($R(x) = r$)	$1 - e^{-r(x-x_0)}$ $r > 0$	$x_0 \leq x < \infty$
	$\frac{(1 - e^{-r(x-x_0)})}{(1 - e^{-r(x_*-x_0)})}$ $r > 0$	$x_0 \leq x < x_*$
2. Constant proportional risk averse or prone ($S(x) = 1 - r$)	$\frac{(x-x_0)^r}{(x_*-x_0)^r}$ $r \geq 0$	$x_0 \leq x < x_*$
3. Risk neutral (special case of (1) when $r \rightarrow 0$ and (2) when $r \rightarrow 1$.)	$\frac{(x-x_0)}{(x_*-x_0)}$	$x_0 \leq x < x_*$

Note. Functional forms also exist for $r < 0$. For ease of exposition we restrict our analyses to $r > 0$. For constant proportional risk attitude, the utility function is risk averse for $0 < r < 1$ and risk prone for $r > 1$.

Empirical Experience

Neither decision analysts nor marketing scientists have explicitly approached vN-M utility measures as error-laden measures. Meyer and Pratt (1968) provide a procedure for "fairing" deterministically a smooth function through a set of points. Fishburn and Kochenberger (1979) use a minimum mean squared error procedure, and Currim and Sarin (1983) use a conjoint-like procedure and a minimum stress procedure, but none of these authors explicitly models measurement error statistically or examines its implications. We know of no systematic empirical study quantifying measurement error at the individual level.

The only systematic empirical studies of which we are aware relate to variation across individuals. Such studies do not necessarily relate to variation within individuals, but they are appropriate if our theory is to be applied across individuals and they are suggestive of the type of empirical research necessary to examine assumptions within individuals. For example, in one study by Fishburn and Kochenberger (1979), although the intervals are coarse and unequal, 7 of 8 cases frequency distributions "look" either normal or exponential.

4. Single Parameter Uniattributed Utility Functions

Following Fishburn and Kochenberger (1979), we assume that separate parameters are estimated "above target" and "below target," thus we can assume that the utility function is either concave throughout the region, $x_0 \leq x < x_*$, or convex throughout the region. Without loss of generality we assume that the attribute of interest, x , has been scaled such that preference is monotonically increasing in x over the region of estimation. For example, if there were a finite ideal point, say length of an automobile, we either (1) assess separately for the range above and the range below the ideal point or (2) assess with respect to a rescaled attribute such as distance from the ideal point.

Our results are derived at the level of the individual consumer, that is, we assume that any variation in the unknown parameter, r , represents uncertainty in measuring the parameter and/or uncertainty across time and situations. We note, however, that our results can be interpreted for variation across consumers with proper modification in definitions. We begin with an example that illustrates the nature of the problem of interest and its essential characteristics.

An Illustrative Example

Suppose that a consumer is considering replacing his antiquated home heating system with a new oil, gas, electric, or solar system. He is uncertain about unit fuel cost, about heating efficiency, and about weather, thus, the annual savings, x , of the

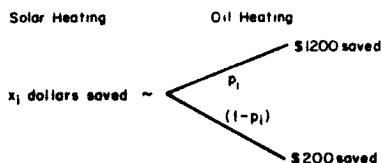


FIGURE 2. Schematic of Uniattribute Lottery Measurement.

new system over the present system is an uncertain outcome. Suppose that he has some prior beliefs about the savings due to each system and that these prior beliefs can be characterized by a probability distribution over the range of $\$200 < x < \1200 . We want to estimate his utility for values of x and to predict his future choices. (Assume for simplicity of exposition that attributes other than x do not affect his choices. §5 relaxes this assumption.)

Using a standard decision analysis lottery questioning format, we ask the lottery question described schematically in Figure 2. The consumer is given a choice between two heating systems. Heating system A , a solar system, has a known savings of x_i dollars. The savings of heating system B , an oil system, are less certain and depend upon the price of oil. If conditions are favorable, the savings are \$1200, and if they are unfavorable, the savings are only \$200. The consumer is asked to specify the likelihood (probability), p_i , of favorable conditions such that he would be indifferent between system A and system B . An example question is given in Appendix 2.

For other market research wordings of this type of question see Hauser and Urban (1977, pp. 593–594) and Eliashberg (1980, pp. 74–75). Alternatively, one might consider using a 0–10 or 0–100 probability scale to elicit p_i . For examples of such scales, see Juster (1966) or Morrison (1979).³

Suppose that from discussions with the consumer we believe a constant proportional risk averse utility function is appropriate. For our problem, the function is:

$$u(x_i, r) = (x_i - 200)^r / (1000)^r. \quad (3)$$

If the vN-M axioms hold, then Figure 2 implies:

$$u(x_i, r) = p_i u(1200, r) + (1 - p_i) u(200, r). \quad (4)$$

Substituting equation (3) in equation (4) yields:

$$(x_i - 200)^r / 1000^r = p_i(1) + (1 - p_i)(0) = p_i. \quad (5)$$

Finally, if there were no errors and we know x_i and p_i , we could obtain r by solving the algebraic relationship in equation (5). That is:

$$r_i = r(x_i, p_i) = \log(p_i) / \log[(x_i - 200) / 1000]. \quad (6)$$

The practice in marketing research is to ask multiple questions as illustrated in Table 2 and to utilize all the information obtained. That is, we could vary x_i and have the consumer specify a p_i for each x_i . We would then use equation (6) to obtain an r_i for each x_i . However, as Table 2 indicates, we are likely to get a different value of r for each question since it is quite unlikely that the consumer will be perfectly consistent in responding to the various questions. The conceptual model in Figure 1 gives us a framework to analyze and interpret the implications of such variation in r_i .

³We note, however, that Morrison (1979) uses an alternative error theory to analyze such probability scales. See also Kalwani and Silk (1982).

TABLE 2

Example of Assessment for the Annual Savings of a Home Heating System

<i>i</i> Measurement	x_i (dollars)	p_i	$r(x_i, p_i)$ (constant proportional risk averse)
1	300	0.29	0.54
2	400	0.46	0.48
3	500	0.53	0.53
4	600	0.64	0.49
5	700	0.72	0.47
6	800	0.75	0.56
7	900	0.83	0.52
8	1000	0.91	0.42
9	1100	0.95	0.49

We begin with maximum likelihood estimators, $\hat{\lambda}$, for λ , when questions are asked in the format of an indifference question. We address revealed preference questions after we derive the necessary analytic tools, i.e., expressions for the distribution of utility and for choice probabilities.

Estimation for Preference Indifference Question Formats

A preference indifference question is a question such as the one described schematically in Figure 2. The experimenter provides x_* , x_0 , and either x_i or p_i ; the consumer answers with a value of p_i (or x_i) such that he is indifferent between the two alternatives.

The experimenter's task is to estimate $\hat{\lambda}$ from I indifference questions. Before we can proceed further, we must make an assumption about the family of distributions, $f(r|\lambda)$. In this paper, we investigate two error distributions: (1) Normal distributions and (2) Exponential distributions.

Normal error theory has the advantage that it is the natural assumption usually made in statistical theory. Its drawback is that r_i can take on any value in the range $(-\infty, \infty)$. However, if the mean is significantly larger than the standard deviation, then negative values of r_i will be extremely rare.

Exponential error is not subject to this problem since we can restrict $r \geq r_0$, i.e.,

$$f(r|\lambda) = (\lambda - r_0)^{-1} \exp[-(r - r_0)/(\lambda - r_0)] \quad \text{for } r \geq r_0.$$

However, exponential error theory does imply an asymmetric distribution with its peak at $r = r_0$ and zero probability for $r < r_0$.

Normal error and Exponential error are clearly quite different theories. Each has its advantages and its disadvantages and, a priori, each reader will have his own favorite theory. We investigate both assumptions in this paper in the belief that these two assumptions are each flexible and together span a broad range of potential shapes for $f(r|\lambda)$.

As it turns out, it is quite simple to obtain the maximum likelihood estimator (MLE) for λ , once an error assumption is made about the shape of $f(r|\lambda)$. (MLE's are important for applied statistics because they are consistent, efficient, and functions of minimal statistics.)

Suppose we ask I questions of the format of Figure 2. That is, for a vector of "certain outcomes," $\mathbf{x} = (x_1, x_2, \dots, x_I)$, we obtain a vector of corresponding "answers," $\mathbf{p} = (p_1, p_2, \dots, p_I)$. Because successive questions are independent, it is easy to show that maximizing the joint probability, $F(\mathbf{p}|\mathbf{x}, \lambda)$, of observing $\mathbf{p}$ given $\mathbf{x}$

and λ is equivalent to maximizing the following log likelihood function:⁴

$$L(\lambda | \mathbf{x}, \mathbf{p}) = \sum_i \log f(r_i | \lambda) \quad (7)$$

where $r_i = r(x_i, p_i)$. The MLE for λ , $\hat{\lambda}$, is the value of λ that maximizes (7). In other words, we simply treat the r_i as data points. This has a number of very practical advantages.

Estimators

If we treat the r_i as data, the MLE's are well known for both Normal error and Exponential error. If μ and σ^2 are the mean and variance of the Normal distribution, then the MLE's⁵ are given by:

$$\hat{\mu} = (1/I) \sum_i r_i, \quad (8)$$

$$\hat{\sigma}^2 = (1/I) \sum_i (r_i - \hat{\mu})^2. \quad (9)$$

For Exponential errors:

$$\hat{\lambda} = (1/I) \sum_i r_i. \quad (10)$$

Furthermore, $\hat{\mu}$ and $\hat{\lambda}$ can be interpreted as the expected ("true") values of r for Normal and Exponential errors, respectively, if we assume induced error is zero-bias.

MLE's are invariant under transformation, that is, if $\hat{\theta}$ is an MLE for θ , then $g(\hat{\theta})$ is the MLE for $g(\theta)$ (Giri 1977, Lemma 5.1.3). Thus, if we interpret $\hat{\mu}$ or $\hat{\lambda}$ as the "true" values of r , then $u(x_i, \hat{\mu})$ or $u(x_i, \hat{\lambda})$ are "true" values of $u(x_i)$ for Normal and Exponential errors, respectively.

Note that equations (8), (9), and (10) apply for both the constant absolute and constant proportional risk averse forms in Table 1. For that matter, they apply for any unattributed utility function for which a unique r_i can be computed for each consumer question. For the generalized version of the binary preference comparison question shown in Figure 2 ($x_0 \leq x_i \leq \infty$ for constant absolute risk aversion and $x_0 \leq x_i \leq x_*$ for constant proportional risk aversion), the inverse functions are given by:

$$\begin{array}{l} \text{Constant} \\ \text{Absolute risk:} \\ \text{Aversion} \end{array} \quad r(x_i, p_i) = - \frac{\log(1 - p_i)}{(x_i - x_0)}, \quad (11)$$

$$\begin{array}{l} \text{Constant} \\ \text{Proportional risk:} \\ \text{Aversion} \end{array} \quad r(x_i, p_i) = - \frac{\log(p_i)}{\log[(x_i - x_0)/(x_* - x_0)]}. \quad (12)$$

When $x_0 \leq x \leq x_*$ for constant absolute risk aversion, the inverse function can be obtained numerically and, for a few special cases, analytically.

Question Format

We derived equation (7) for the case when the x_i 's were specified by the experimenter such that the consumer's answers were the p_i 's. But, by symmetry, it is clear

⁴ $F(\mathbf{p} | \mathbf{x}, \lambda)$ is the product of the $f(r(x_i, p_i) | \lambda)$'s times the Jacobian which is independent of λ .

⁵ Equation (9) is the MLE for $\hat{\sigma}$, but it is biased for finite I . The more commonly used estimator is $(I/(I-1))\hat{\sigma}^2$. Also, if we want to estimate r_0 , its MLE is $\hat{r}_0 = \min_i \{r_i\}$.

that equation (7) applies if the probabilities, p_i 's, are specified and the consumer supplies the certain outcomes, x_i 's. In fact, a modified equation (7) will apply for any question format for which one can obtain an observation of r_i . See Farquhar (1984) for a review of alternative question formats. However, equation (7) does not imply that the experimenter's choice of question format is free from systematic bias. Different formats may induce different biases, e.g., different μ or λ , different magnitudes of error, e.g., different σ^2 for Normal errors, or, for that matter, different types of error, e.g., different $f(r|\lambda)$. But equation (7) does state that once the error assumption is made, equations (8) and (9) or (10) apply independently of the question format.

Statistical Inference

One can test a hypothesis about the "true" value of r . For example, if normal error theory applies and the researcher wishes to test whether the "true" value of r is significantly different from some hypothesized value, r_H , he can use a t -test with $(I - 1)$ degrees of freedom based on the statistic, $(r_H - \hat{\mu})(I - 1)^{1/2}/\hat{\sigma}$. Similar results apply for exponential error theory, except that the sampling distribution for $\hat{\lambda}$ has a gamma density with mean λ_H and variance λ_H^2/n .

An Illustration

Consider the problem in Table 2. Using equations (8), (9), and (10) we estimate $\hat{\mu} = 0.50$ and $\hat{\sigma} = 0.04$ for Normal errors and $\hat{\lambda} = 0.50$ for Exponential error. A chi-square goodness-of-fit test suggests that the data are more likely to be generated from a Normal distribution. A utility function based on $r^T = \hat{\mu} = 0.5$ is shown in Figure 3. For normal error theory, a 95% confidence interval for $\hat{\mu}$ is [0.47, 0.53] and for $\hat{\sigma}$ it is [0.03, 0.08].

Distribution of Utility

If the risk parameter, r , were known with certainty, we could compute $u(x, r)$ for any x and then compute directly the expected utility of a product. However, even with an MLE for the "true" parameter, λ , our knowledge about r is still represented by a random variable with distribution, $f(r|\lambda)$. This uncertainty in r induces uncertainty in $u(x, r)$ for any x . Hence, the expected utility and, ultimately, the choice outcome are random variables. We begin by computing the probability density function of $u(x, r)$. We then examine its implications. For simplicity of analytic exposition we restrict our results to the infinite range ($0 < x < \infty$) constant absolute risk averse utility function and (for exponential errors) to $r > 0$. For constant proportional risk averse utility functions the range can be either finite or infinite and for normal errors r is unrestricted. (The proofs to Propositions 1 and 2 and all subsequent propositions are contained in Appendix 1.)

PROPOSITION 1. *If measurement error is modeled as NORMAL, then the utility functions have lognormal distributions. In particular,*

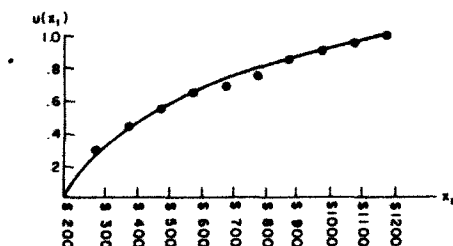


FIGURE 3. Maximum-Likelihood Estimate of Assessed Utility Function.

$u(x, r) \sim \Lambda(-k\hat{\mu}, k\hat{\sigma}^2)$ for constant proportional risk aversion,
 $1 - u(x, r) \sim \Lambda(-x\hat{\mu}, x^2\hat{\sigma}^2)$ for constant absolute risk aversion, where
 $k = \log[(x_* - x_0)/(x - x_0)]$,
 $\Lambda(a, b) \equiv a$ Lognormal distribution with parameters a, b .

PROPOSITION 2. If measurement error is modeled as EXPONENTIAL, then the utility functions have Beta distributions. In particular,

$u(x, r) \sim \text{Beta}(1/\hat{\lambda}k, 1)$ for constant proportional risk aversion,

$u(x, r) \sim \text{Beta}(1, 1/\hat{\lambda}x)$ for constant absolute risk aversion,

where k is defined in Proposition 1, and $\text{Beta}(c, d) \equiv a$ Beta distribution with parameters c, d .

Propositions 1 and 2 are useful for practical applications involving risky and riskless alternatives. For both error theories, the induced distributions on $u(x, r)$ are recognizable distributions with known properties similar to those that arise in quantal choice problems. This will become key as we proceed to forecasts of choice probabilities. Because Lognormal and Beta distributions compound nicely with conjugate distributions (DeGroot 1970) it is possible to obtain analytic results for important distributions of outcomes.

Probability of Choice

If r were known with certainty, the expected utility of each product could be computed and we would simply forecast that the consumer would choose the product with the maximum expected utility. In this case, our forecasting statement would be made categorically, that is, with probability zero or one. Instead, $u(x, r)$ is a random variable with distribution given by Propositions 1 and 2. Hence, the best we can forecast is the probability, P_j ($0 < P_j < 1$) that the consumer will choose product j . That is,

$$P_j = \text{Prob} \left[\int u(x, r) h_j(x) dx \geq \int u(x, r) h_k(x) dx \text{ for } k = 1, 2, \dots, J \right] \quad (13)$$

where $h_j(x)$ is the probability distribution of outcomes for Alternative j .

If we were evaluating riskless alternatives, then (13) becomes a quantal choice problem similar to logit or probit analysis (McFadden 1980) except that we use Lognormal or Beta distributions rather than the double exponential and normal distributions used in logit and probit analyses, respectively.⁶ (See reviews of quantal choice models in Manski and McFadden 1981 and Daganzo 1979.)

Since our focus is on risky alternatives, we examine in detail two important cases of equation (13). We examine binary choices among:

(1) Products whose outcomes are specified by lotteries possessing discrete (Bernoulli) probabilities, and

(2) Products whose outcomes are specified by continuous probability density functions, especially normal distributions.

These cases illustrate the essential ideas behind equation (13). We obtain analytic results for both problems. We leave the problems of other uncertain outcomes and multiple choices for future research, although we point out that, in principle, one could use Propositions 1 and 2 with numerical techniques to compute P_j via equation (13). This would be analogous to the use of numerical techniques in state-of-the-art multiple choice probit analysis (see Daganzo 1979).

⁶Quantal choice problems with lognormal mixtures of double exponential distributions of utility have been studied. See Boyd and Mellman (1980).

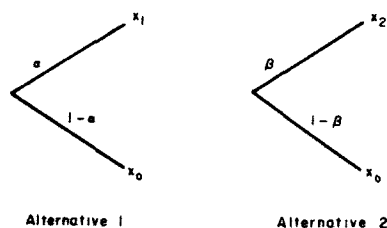


FIGURE 4. Binary Choice Problem.

Binary Choice Between Lotteries

The first consumer choice situation that we consider is characterized as a binary choice problem with dichotomous outcomes illustrated in Figure 4. For example, Alternatives 1 and 2 may be two heating systems heated by two different energy sources and the dichotomous outcomes may represent world events that affect the prices of these sources. Without loss of generality assume $x_1 > x_2$ and $\beta > \alpha$. As before, $x_0 < x_1, x_2 \leq x_*$. (If $x_1 > x_2$ and $\alpha > \beta$, then Alternative 1 would dominate Alternative 2.) This simple choice problem contains the essence of risky choice; the individual must decide between a potentially greater payoff, Alternative 1, and a greater likelihood of the payoff, Alternative 2.

Our objective is to estimate the probability, $\hat{P}_1$, that the consumer will choose Alternative 1, given that we have estimated the parameters, λ , of the probability density function from which the risk parameter, r , is drawn. Before we proceed, we note that, for the binary choice problem presented in Figure 4, measurement errors may be induced once, for the question as a whole, or twice, once for each alternative. This gives rise to two viewpoints (assumptions) regarding the nature of our conceptualizations of how consumers draw r , from $f(r|\lambda)$.

We label these assumptions as single and multiple random draws. Under the single random draw assumption, the consumer draws the corresponding risk parameter only once, and he is consistent in the sense of using the same parameter (and hence, the same utility function) to evaluate all alternatives in his choice set. Under the multiple random draw assumption, the consumer draws the risk parameter every time he evaluates an alternative. The two assumptions imply similar, but slightly different, choice probabilities. We begin with single random draw.

PROPOSITION 3. *For the binary choice problem with discrete outcomes (Figure 4), under the single random draw assumption, if measurement error is modeled as NORMAL, then:*

$$\hat{P}_1 \approx \Phi[(\hat{\mu} - \kappa \log(\beta/\alpha))/\hat{\sigma}] \quad \text{for constant proportional risk aversion,}$$

$$\hat{P}_1 \approx \begin{cases} \Phi[(r_c - \hat{\mu})/\hat{\sigma}] & \text{if } \alpha x_1 > \beta x_2 \\ 0 & \text{if } \alpha x_1 \leq \beta x_2, \end{cases}$$

where $\kappa^{-1} = \log[(x_1 - x_0)/(x_2 - x_0)],$

and r_c solves the equation $\beta \exp(-r_c x_2) - \alpha \exp(-r_c x_1) = \beta - \alpha$, and $\Phi[\]$ denotes the cumulative distribution function of a normally distributed variate.

PROPOSITION 4. *For the binary choice problem with discrete outcomes (Figure 4), under the single random draw assumption, if measurement error is modeled as EXPO-*

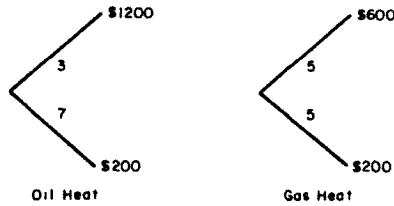


FIGURE 5. Hypothetical Characteristics of the Risk Involved for Two Home Heating Systems.

ENTIAL, then:

$$\hat{P}_1 = [\beta/\alpha]^{-\kappa/\hat{\lambda}} \quad \text{for constant proportional risk aversion,}$$

$$\hat{P}_1 = \begin{cases} 1 - \exp[-r_c/\hat{\lambda}] & \text{if } \alpha x_1 > \beta x_2 \\ & \text{for constant absolute risk aversion (infinite range, } 0 < x < \infty) \\ 0 & \text{if } \alpha x_1 < \beta x_2, \end{cases}$$

where κ and r_c are defined in Proposition 3.

Propositions 3 and 4 are useful results. To illustrate their application, consider the hypothetical alternatives in Figure 5. Alternative 1 is oil heat where the high risk reflects volatile supplies. Alternative 2 is gas heat where the risk reflects only uncertainty in the heating characteristics of the home.

Using the distribution implied by the data in Table 2, $\hat{\mu} = 0.5$, $\hat{\sigma} = 0.04$. From Figure 5, $\alpha = 0.3$, $\beta = 0.5$, $x_0 = 200$, $x_1 = 1200$, and $x_2 = 600$. Assuming a constant proportional risk averse utility function and substituting these values in Proposition 3 yields

$$\hat{P}_1 = \Phi\left[\{0.5 - \log(0.5/0.3)/\log(1000/400)\}/0.04\right] \approx \Phi[-1.44] \approx 0.075$$

for normal error theory. In a marketing forecasting application, we would assign a 0.075 value to the probability that the consumer would choose oil heat.

We now consider multiple random draws. We have been able to obtain analytic results for constant proportional risk averse utility functions.

PROPOSITION 5. *For the binary choice problem with discrete outcomes (Figure 4), under the multiple random draw assumption, for a constant proportional risk averse utility function:*

$$\hat{P}_1 = \Phi\left\{\left[\hat{\mu} - \kappa \log(\beta/\alpha)\right]/\kappa\eta\hat{\sigma}\right\} \quad \text{for NORMAL errors}$$

where $\eta^2 = k_1^2 + k_2^2$, κ is defined in Proposition 3, and

$$k_1 = \log[(x_* - x_0)/(x_1 - x_0)], \quad k_2 = \log[(x_* - x_0)/(x_2 - x_0)],$$

$$\hat{P}_1 = [k_2/(k_1 + k_2)][\beta/\alpha]^{-(1/\hat{\lambda}k_2)} \quad \text{for EXPONENTIAL errors.}$$

It is interesting to compare Proposition 3 (Normal errors, constant proportional risk aversion) to Proposition 5 (Normal errors, constant proportional risk aversion). This comparison illustrates the impact of measurement error on our ability to estimate choice probabilities. Without error, r^T is known and Alternative 1 will be chosen, $P_1 = 1$, whenever $r^T > \kappa \log(\beta/\alpha)$. This corresponds to Propositions 3 and 5 with $\hat{\sigma} \rightarrow 0$. As our uncertainty, $\hat{\sigma}$, about r increases, our ability to predict the consumer's behavior decreases, i.e., $\hat{P}_1$ decreases for $\hat{\mu} > \kappa \log(\beta/\alpha)$. If we compound that error

by allowing the consumer multiple random draws from $f(r|\lambda)$, then our ability to predict is modified still further because we replace $\hat{\sigma}$ by $\kappa\hat{\sigma}$. The differences between Propositions 3 and 5 make clear the implications of our assumption about our knowledge of the consumer's choice process.

One can obtain similar interpretations by comparing Propositions 4 and 5 for exponential errors. The forms are the same, but the constants vary, e.g., κ vs. k_2^{-1} .

Thus, clearly, an "open loop" estimation of probabilities, i.e., use indifference questions to estimate $f(r|\hat{\lambda})$ and Propositions 3 to 5 to estimate $\hat{P}_1$, will depend on the assumptions we make about how uncertainty in $u(x, r)$ affects uncertainty in choice outcomes. On the other hand, a "closed-loop" revealed preference estimation of probabilities will be much less dependent on this assumption. We discuss these issues in detail after we derive the revealed preference estimators. However, we first complete this subsection with estimates for $\hat{P}_1$ for constant absolute risk averse utility functions.

Despite the fact that the Lognormal and Beta distributions are well studied (e.g., Aitchison and Brown 1969, DeGroot 1970, Drake 1967), we have been unable to obtain analytic results for $\hat{P}_1$ with constant absolute risk averse utility. Instead, Proposition 6 relies upon implied integral equations which require numerical techniques.

PROPOSITION 6. *For the binary choice problem with discrete outcomes (Figure 4), under the multiple random draw assumption, for a constant absolute risk averse utility function:*

$$\begin{aligned}\hat{P}_1 &= \text{Prob}[\tilde{\varphi}_2 - \tilde{\varphi}_1 \geq (\beta - \alpha)] \quad \text{for NORMAL errors,} \quad \text{where} \\ \tilde{\varphi} &\sim \Lambda(\log \alpha - x_1 \hat{\mu}, x_1^2 \hat{\sigma}^2), \quad \tilde{\varphi}_2 \sim \Lambda(\log \beta - x_2 \hat{\mu}, x_2^2 \hat{\sigma}^2), \\ \hat{P}_1 &= \text{Prob}[\alpha \tilde{u}_1 - \beta \tilde{u}_2 > 0] \quad \text{for EXPONENTIAL errors,} \quad \text{where} \\ \tilde{u}_l &\sim \text{Beta}(1, 1/\hat{\lambda} x_l) \quad \text{for } l = 1, 2.\end{aligned}$$

Binary Choice Among Alternatives Represented by Continuous Distributions of Outcomes

Although attributes of many consumer products can be represented by lotteries, many other attributes will be represented by continuous distributions such as the Normal distribution. For example, if we buy a new automobile, we might expect that the miles per gallon we actually obtain is best represented by a Normal distribution based on the published EPA estimate.

Proposition 7 derives the probability of choice, $\hat{P}_1$, if the two outcomes are represented by Normal distributions with means and variances, m_1 , v_1^2 and m_2 , v_2^2 , respectively. For simplicity, we state the result only for single random draws with constant absolute risk averse utility functions. Other results are obtainable but some require numerical techniques. Without loss of generality assume $m_1 > m_2$ and $v_1 > v_2$. (If $m_1 > m_2$ and $v_1 < v_2$, then $\hat{P}_1 = 1$.)

PROPOSITION 7. *For binary choice among Normally distributed outcomes, under the single random draw assumption, for a constant absolute risk averse utility function:*

$$\begin{aligned}\hat{P}_1 &= \Phi \left[\frac{2(m_1 - m_2)/(v_1^2 - v_2^2) - \hat{\mu}}{\hat{\sigma}} \right] \quad \text{for NORMAL errors,} \\ \hat{P}_1 &= 1 - \exp \left[-2(m_1 - m_2)/\hat{\lambda}(v_1^2 - v_2^2) \right] \quad \text{for EXPONENTIAL errors.}\end{aligned}$$

We have stated the result explicitly for Normally distributed outcomes, but the key idea of the proof (see Appendix) is to use exponential transforms to simplify the

integral equation for $\hat{P}_1$. We can use the same method to obtain results for any distribution for which the exponential transform is tabled. Tabled functions include the continuous Beta, Cauchy, Chi-square, Erlang, Exponential, Gamma, Laplace, and Uniform distributions as well as some discrete distributions such as the Binomial, Geometric, and Poisson distributions. See tables in Keeney and Raiffa (1976, p. 202) and Drake (1967, pp. 271-276).

For the constant proportional risk averse utility function we can also obtain results by using the Mellin transform, $\int x^s f(x) dx$, for those distributions for which it exists. See tables in Bateman (1954).⁷

Estimation for Revealed Preference Questions

The most commonly used question formats in decision analysis use some form of preference indifference question. However, in marketing such questions have been criticized as too complex. On the other hand, revealed preference questions, where the consumer is asked to choose among (or rank order) alternatives, are very common. For example, conjoint analysis, as reviewed by Green and Srinivasan (1978), uses this form of questioning and is one of the most widely used marketing research procedures. In fact, Currim and Sarin (1984) use a modified conjoint analysis procedure to estimate vN-M like utility functions. Furthermore, revealed preference is one of the most commonly used techniques in transportation demand analysis. See Manski and McFadden (1981).

Since we are addressing a market research issue, we allow the experimenter to choose the question format much as he would choose the fractional factorial design in conjoint analysis. For revealed preference estimation the consumer's task is simple. He is given I pairs of alternatives. Each alternative is described by a probability distribution of outcomes (usually lotteries, but continuous distributions are allowable if they can be described adequately to the consumer). For each pair of alternatives, the consumer is asked to choose the alternative which he prefers. See Appendix 2 for an example question format. Propositions 3 through 7 give us the analytic tools to obtain estimators for $f(r|\lambda)$ from the answers to such questions.

Let $\delta_i = 1$ if the consumer chooses Alternative 1 for the i th pair and let $\delta_i = 0$ if he chooses Alternative 2. Let δ be the vector of δ_i 's. Then the joint probability, $F(\delta|\lambda)$, of observing a particular set of answers, δ , given $f(r|\lambda)$ is given by:

$$F(\delta|\lambda) = \prod_{i=1}^I P_{1i}^{\delta_i} (1 - P_{1i})^{1-\delta_i} \quad (14)$$

where $P_{1i} = P_{1i}(\lambda)$ are determined for each question, i , by the appropriate proposition (i.e., Propositions 3, 4, 5, 6 or 7 or their extensions). For example, for lotteries under the multiple random draw assumption with exponential errors and a constant proportional risk averse utility function, P_{1i} is given by:

$$P_{1i}(\lambda) = [k_{2i}/(k_{1i} + k_{2i})][\beta/\alpha]^{-(1/\lambda_{2i})} \quad (15)$$

In principle, we could form a log-likelihood function based on equations (14) and (15) and then maximize it by numerical techniques to obtain MLE's of $\lambda, \hat{\lambda}$. However, if the experimenter chooses his measurement design carefully, he can obtain practical analytic expressions for $\hat{\lambda}$. In particular, for equation (14), if (1) he chooses α, β, x_1

⁷Exponential transforms are also known as moment generating functions (with the sign of the argument reversed) and Laplace transforms when $x > 0$. The Mellin transform is also known as the factorial moment generating function.

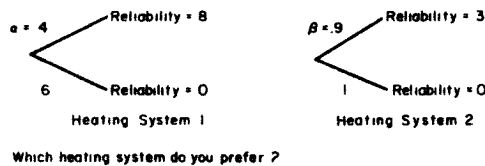


FIGURE 6. Schematic of Revealed Preference Question Corresponding to Proposition 5.

and x_2 for the first question (review Figure 4), if (2) for every subsequent question, i , he chooses α_i and x_{1i} , and if (3) he then chooses β_i and x_{2i} according to the following rule:

$$x_{2i} = x_0 + [(x_2 - x_0)/(x_* - x_0)]^{\gamma_i}(x_* - x_0), \tag{16}$$

$$\beta_i = (\beta/\alpha)^{\gamma_i}\alpha_i, \tag{17}$$

$$\gamma_i = \log[(x_{1i} - x_0)/(x_* - x_0)]/\log[(x_1 - x_0)/(x_* - x_0)], \tag{18}$$

then he can obtain an analytic expression for $\hat{\lambda}$. (Note we have suppressed the subscript 1 on x_{1i} , etc. for the base question, $i = 1$.)

The analytic expression is obtained from the invariance properties of MLE's in the following way. Equations (16) through (18) ensure that for all λ , P_{1i} is constant for all questions. Define I_1 as the number of times the first alternative is chosen, then $F(\delta|\lambda)$ becomes a Binomial distribution for I_1 . The MLE for a Binomial distribution is obtained simply as $\hat{P}_1 = I_1/I$. $\hat{\lambda}$ is obtained by solving the equation $\hat{P}_1(\lambda) = I_1/I$. From equation (15) we obtain

$$\hat{\lambda} = [(1/k_2)\log(\beta/\alpha)]/[\log\{k_2I/(k_1 + k_2)I_1\}] \tag{19}$$

where α , β , k_1 and k_2 are obtained from the reference question.

To illustrate this technique, consider a set of questions in which each alternative is a potential heating system. The attribute of interest is reliability, that is, a 0 to 10 scale indicating how likely it is that the system will not require major repairs during the next five years. (For example, this reliability index might be 10 times the probability that no repair will be required.) One such question is illustrated schematically in Figure 6.

We can then ask 10 questions of this form as indicated in Table 3. (We have rounded β to the nearest 0.05.) We record I_1 , the number of times the consumer prefers Heating System 1.

TABLE 3
Example Experimental Design for
Revealed Preference Questions

Question Number	Heating System 1		Heating System 2	
	x_{1i}	α_i	x_{2i}	β_i
1	9.5	0.5	7.5	0.60
2	9.0	0.5	5.5	0.75
3	8.5	0.5	4.0	0.90
4	8.0	0.4	3.0	0.90
5	7.5	0.3	2.0	0.85
6	7.0	0.2	1.5	0.70
7	6.5	0.2	1.0	0.90
8	9.5	0.6	7.5	0.70
9	9.0	0.6	5.5	0.90
10	9.5	0.7	7.5	0.85

For example, if $I_1 = 2$, then $\hat{\lambda} = 0.44$, and if $I_1 = 3$, then $\hat{\lambda} = 0.61$. We could, of course, obtain better estimates by asking more questions. For example, if $I = 100$ and $I_1 = 21$, then $\hat{\lambda} = 0.45$ and if $I_1 = 22$, then $\hat{\lambda} = 0.47$.

We constructed equations (16) through (19) and Table 3 for Proposition 5, Exponential error. It is also possible to construct experimental designs for Normal errors, for continuous distributions of outcomes, and for constant absolute risk averse utility functions. The experimenter simply chooses the appropriate proposition (or its extension) and derives the condition on α , β , x_1 , and x_2 such that $\hat{P}_1$ is constant for all questions. $\hat{\lambda}$ is the solution to $P_1(\lambda) = I_1/I$. For example, for Proposition 7, we can restrict $(m_1 - m_2)/(v_1^2 - v_2^2)$ to be constant to assure $\hat{P}_1$ is constant. For Normal errors, there are two unknown parameters, hence we must either (1) assume one parameter is known, or (2) ask two sets of clustered questions to obtain two equations in two unknowns. Of course, more parameters will require more questions and the MLE's will depend upon the obtainable sample size.

Revealed Preference vs. Preference Indifference Questions

Each type of question format has its advantages and disadvantages. For example, revealed preference formats are likely to be easier for the respondent to answer, but we require more of them to obtain reasonable estimates of $\hat{\lambda}$ or $\hat{P}_1$.

A more subtle issue is that of "open loop" vs. "closed loop". We saw earlier that the estimate of $\hat{P}_1$, based on the preference indifference format, depends upon our assumption of single or multiple draws. This is because we estimate $f(r|\lambda)$ directly from the preference indifference question and then use $f(r|\lambda)$ to calculate $\hat{P}_1$. We never collect direct information on choice outcomes.

On the other hand, with the revealed preference format, we do collect direct information on choice outcomes. In this case our estimate of $\hat{P}_1$ depends less upon our assumptions. For example, if we use Proposition 3 rather than Proposition 5, our estimate of $\hat{\sigma}$ will be smaller by a factor of $(\kappa\eta)^{-1}$, but $\kappa\eta$ will cancel out when we use the same proposition to forecast $\hat{P}_1$. We can expect similar robustness, but not exact cancellation, with respect to our assumptions on the error distribution.

Such robustness of "closed-loop" revealed preference technique is discussed in the econometric literature. For example, Domencich and McFadden (1975, p. 57) provide a table and discussion illustrating the similarity in probability predictions of the Logit, Probit, and Arctan probability of choice models. The Logit is based on Double-exponential errors, the Probit is based on Normal errors, and the Arctan is based on Cauchy errors.

The preference indifference format has complementary advantages. Because $\hat{\lambda}$ (or $\hat{\mu}$ and $\hat{\sigma}$) are calculated directly from the consumer's answers and not from "solving" $\hat{\lambda} = g(\hat{P}_1)$, estimates of $\hat{\lambda}$ based on the preference indifference format are less likely to be dependent upon our assumptions.

In other words, if our interest is in estimating the purchase probability, $\hat{P}_1$, then the revealed preference format is likely to be better because it uses choice outcome data directly and is robust with respect to our assumptions on the type of error distribution and on single vs. multiple draws. If our interest is in estimating the "true" value of the risk parameter or the distribution of the risk parameter, then the preference indifference format is likely to be better.

Summary of Single Parameter Uniattributed Utility Functions

This completes our analytic discussion of single parameter, uniattributed utility functions. (We discuss empirical issues and assumption testing in §6.)

Occasionally, researchers may wish to use multiple parameter uniattributed utility functions, for example, see Keeney and Raiffa (1976, p. 209). If questions can be

clustered then the multiattributed technique (Procedure 2) discussed in §5 can be applied to multiparameter uniattributed functions. Alternately, one can use a regression approximation⁸ based on a Taylor's series expansion of the risk premium as defined by Pratt (1964) and discussed also by Keeney and Raiffa (1976).

5. Multiattributed Utility Functions

There are many applications in marketing where it is necessary to model decisions involving multiple attributes, each of which is risky. For example, the decision to buy a home heating system might involve reliability as well as annual dollar savings. (Choffray and Lilien 1978 illustrate empirically a multiattributed preference problem for solar air-conditioning.)

If we assume that the qualitative questions to determine the appropriate functional form of the multiattributed utility function have already been asked (e.g., Keeney and Raiffa 1976, pp. 299–301 or Eliashberg 1980, pp. 74–75), then we can estimate multiattributed utility functions in two ways.

Estimation Procedure 1

The first procedure is a two-step procedure which combines the results of §4 with commonly used methods in marketing. In Step 1, we use either preference indifference or revealed preference questions to obtain estimators for uniattributed functions for each attribute. Suppose that the condition of "mutual utility independence" (Keeney 1972) among the attributes has been identified qualitatively, then the multiattributed function, $U(x, w)$ is given by:

$$U(x, w) = \sum_{m=1}^M w_m u(x_m, r_m) + \sum_{m=1}^M \sum_{k=m}^M w_{mk} u(x_m, r_m) u(x_k, r_k) \\ + \text{third order and higher interaction terms.} \quad (20)$$

Note that if the higher order terms are zero (for conditions see Fishburn 1974), then equation (20) reduces to the additive forms commonly assumed in conjoint and logit analyses.

In Step 2, the experimenter then asks either preference indifference (or revealed preference) questions using multiattributed alternatives. Standard conjoint (or logit) analysis procedures can then be used to obtain $\hat{w}$ with $u(x_m, \hat{r}_m)$ rather than x_m as the explanatory variables. ($\hat{r}_m$ is our best estimate of r_m from Step 1.)

Estimation Procedure 1 is an approximation. It is a two-stage procedure with the potential problem of compounding errors from Step 1 to Step 2. However, (1) if such compounding is small relative to other measurement errors, or (2) if the additive form applies and errors are independent and identically distributed (i.i.d.) across attributes, then Estimation Procedure 1 should work well. (If implied distributions of $u(x_m, r_m)$ are i.i.d., then the w_m are equally biased which has no effect since $u(x, w)$ is only unique to a transformation.)

⁸The risk premium, π_i , is the expected value of the lottery minus the certainty equivalent, x_i . Keeney and Raiffa (1976, p. 161) show $\pi_i = (1/2)v_i^2 R(x_i) + \epsilon$ where v_i^2 is the variance of the lottery and ϵ represents higher order terms. Writing $R(x_i, r)$ as a function of r and rearranging terms yields $R_i^0 = R(x_i, r) + \tilde{\epsilon}_i$ where, for Figure 2, $R_i^0 = 2\pi_i/[p_i(x_* - x_i)^2 + (1 - p_i)(x_0 - x_i)^2]$ with $\pi_i = x_i - x_i$ and $x_i = p_i x_* + (1 - p_i)x_0$. Thus, R_i^0 is a function of known data and because the Taylor's series error, ϵ , is included in the measurement error, $\tilde{\epsilon}_i$, we have obtained a regression equation. For example, to combine "absolute" and "proportional" risk aversion, we have $R_i^0 = r_1 + r_2(x_i - x_0)^{-1} + \tilde{\epsilon}_i$ which is linear in the unknown parameters. Note that $u(x, r) = f_1 \exp[-\int R(x, r) dx] + f_2$, where f_1 and f_2 are scaling constants.

Estimation Procedure 2

Estimation Procedure 2 is a one-stage procedure, but is limited by practicality to preference indifference questions. (Revealed preference would require numerical techniques.)

Suppose we ask $I \times L$ preference indifference questions where L is the number of parameters, w , to be estimated. Let $x_{pi} = (x_{pi1}, x_{pi2}, \dots, x_{piM})$ be the levels of the M attributes for the certainty equivalent in the p th question in the i th cluster. In assessing $U(x, w)$ we specify either (1) all of x_{pi} or (2) p_{pi} and all but one of the x_{pi} .

If we ask our questions in I clusters of L questions and if a computable vector-valued inverse function, $W(x_i, p_i)$, exists mapping the question set onto the unknown parameters, then the multiattributed problem is simply the multivariate extension of the single parameter unattributed case. (x_i is the matrix with rows x_{pi} and p_i is the vector with elements p_{pi} .) For example, if errors cause $\tilde{w}$ to be distributed with a multivariate normal distribution with mean μ , and covariance matrix, Γ , then the maximum likelihood estimators, $\hat{\mu}$ and $\hat{\Gamma}$, are given by:

$$\hat{\mu} = (1/I) \sum_i w_i, \quad (21)$$

$$\hat{\Gamma} = (1/I) \sum_i [w_i - \hat{\mu}][w_i - \hat{\mu}]' \quad (22)$$

where $w_i = W(x_i, p_i)$ and $(\cdot)'$ denotes transpose.

For a formal proof, see Giri (1977, Chapter 15). As before, we can construct confidence regions with the multivariate extension of a t -test. For example, the appropriate statistic for $\hat{\mu}$ is Hotelling's T^2 statistic (Giri 1977, Chapter 7; and Green 1978, p. 257).

Similar results apply for multivariate exponential error.

Probability of Choice

Choice probabilities can be obtained with equation (13) and numerical techniques. For example, one might use equation (13) by sampling from the multivariate normal distribution, then using the sampled $\tilde{w}$ to compute the expected utility of each option. Choice probabilities are then the percent of times an alternative is chosen in, say, 1000 draws. This computation method is similar to methods used in probit analysis, e.g., Daganzo (1979), and has proven feasible in that context.

Estimation Example

We illustrate Estimation Procedure 2 with a home heating system example. Suppose that besides annual savings, x_1 , the individual is concerned with reliability, x_2 , as measured by 10 times the probability that no repair will be needed each year. We wish to model the consumer's preference by a constant proportional risk averse, multilinear utility function. (This is a two-attribute version of equation (20).)

$$U(x, w) = w_3 u_1(x_1) + w_4 u_2(x_2) + (1 - w_3 - w_4) u_1(x_1) u_2(x_2), \quad (23)$$

$$u_1(x_1) = [(x_1 - 200)/1000]^{w_1}; \quad u_2(x_2) = [x_2/10]^{w_2}.$$

We estimate the distribution of the four unknown parameters, $w = (w_1, w_2, w_3, w_4)$, by asking the lottery question shown in Table 4. In each question, the consumer is asked to give a probability, p_{pi} , such that he is indifferent between a certainty equivalent, (x_{pi1}, x_{pi2}) and a lottery where the system is described by (x''_{pi1}, x''_{pi2}) with probability

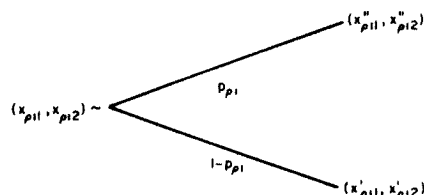


FIGURE 7. Schematic of Multiattribute Lottery.

p_{pi} , and by (x'_{pi1}, x'_{pi2}) with probability, $(1 - p_{pi})$. In other words, the standard lottery shown in Figure 7.

The reader will note that we have designed the questions in Table 4 based on the "corner point" method that allows for easy computation of the inverse function, $W(x_i, p_i)$. That is:

$$w_{1i} = \log(p_{1i})/\log[(x_{1i1} - 200)/1000], \quad w_{2i} = \log(p_{2i})/\log(x_{2i2}/10), \quad (24)$$

$$w_{3i} = p_{3i}/p_{1i}, \quad w_{4i} = p_{4i}/p_{2i}.$$

The corner point method is discussed in Keeney and Raiffa (1976, p. 305) and a market research example with schematics of questions is given in Hauser and Urban (1979, pp. 259-260).

The simplicity of equation (24) is for ease of exposition. Tradeoff questions as well as lotteries can be used and the inverse function can vary with i as long as it is computable for all i . Even with our simplification, the 16 questions provide the experimenter with a variety of questions to be asked. The "answers", p_{pi} , to the lottery

TABLE 4

Example Assessment for the Savings and Reliability of a Home Heating System

Certainty Equivalent				"Win"		"Loss"		Probability	Parameters			
<i>i</i>	ρ	$x_{\rho i1}$	$x_{\rho i2}$	$x''_{\rho i1}$	$x''_{\rho i2}$	$x'_{\rho i1}$	$x'_{\rho i2}$	$p_{\rho i}$	w_{1i}	w_{2i}	w_{3i}	w_{4i}
1	1	400	2	1200	2	200	2	0.45	0.50			
	2	400	2	400	10	400	0	0.60		0.32		
	3	400	0	1200	10	200	0	0.20			0.44	
	4	200	2	1200	10	200	0	0.50				0.83
2	1	600	4	1200	4	200	4	0.65	0.47			
	2	600	4	600	10	600	0	0.75		0.31		
	3	600	0	1200	10	200	0	0.25			0.38	
	4	200	4	1200	10	200	0	0.60				0.80
3	1	800	6	1200	6	200	6	0.75	0.56			
	2	800	6	800	10	800	0	0.80		0.44		
	3	800	0	1200	10	200	0	0.30			0.40	
	4	200	6	1200	10	200	0	0.65				0.81
4	1	1000	8	1200	8	200	8	0.90	0.47			
	2	1000	8	1000	10	100	0	0.95		0.23		
	3	1000	0	1200	10	200	0	0.35			0.39	
	4	200	8	1200	10	200	0	0.75				0.79
$\hat{\mu} =$									0.50	0.33	0.40	0.81

x_1 = savings in dollars

x_2 = reliability index

questions were "constructed" by assuming $w = (0.50, 0.33, 0.40, 0.80)$ and rounding off⁹ to the nearest 0.05.

Examination of Table 4 reveals that the estimated parameters, $\hat{\mu}$, recover the "known" values quite well. The covariance matrix $\hat{\Gamma}$, and the corresponding correlation matrix, $\hat{C}$, can be readily computed with equations (21) and (22).

6. Some Comments on Application and Assumption Testing

In this section we comment briefly on ways to obtain the input data.

Preference Indifference Questions

Preference indifference questions are the most popular means by which to obtain data on consumer risk preference. Appendix 2 provides an example question. For other example questions, see Eliashberg (1980), Hauser and Urban (1977, 1979) and Keeney and Raiffa (1976). See also Hershey, Kunreuther and Schoemaker (1982), Schoemaker and Waid (1982) and Schoemaker (1981) for empirical discussion, interpretation, and suggestions for the type of biases and errors induced when assessing vN-M functions with different question formats.

These are difficult questions to ask. In our own experience we have found it critical to "train" the respondent with warm-up questions and to use multi-colored props and probability wheels to explain the concepts (see Appendix 2). In the initial questions, we iterate in on the indifference point, e.g., indicate a p_i value in which alternative 1 is clearly preferred and a p_i value in which alternative 2 is clearly preferred, then continually narrow the range until the respondent finally indicates it is "too close to call."

Revealed Preference Questions

In a revealed preference question, the respondent is given two lotteries and asked to choose the one he prefers. Because we know of no empirical studies using such questions to assess vN-M functions, we can only speculate on the type of question wording that is appropriate. Appendix 2 provides one such speculation. We expect that the basic task will prove easier for the respondent, but that the concept of a lottery must still be explained carefully. Warm-up questions, props, and attention to potential misunderstandings are likely to be important as we gain experience with this type of question.

7. Future Research

This paper provides the analytic framework to study the effect of measurement error on the modeling of consumer risk preference. But research remains.

Besides the tradeoffs discussed above, we can identify at least four important empirical issues:

(1) We have assumed Normal or Exponential error. If the experimenter wishes to determine the form of $f(r|\lambda)$ empirically, we suggest he use preference indifference questions to obtain r_i for sufficiently large I and plot its histogram. If the histogram is symmetric and unimodal, Normal errors are likely to be the best assumption; if the histogram is unimodal and skewed with $\hat{\sigma} \approx \hat{\mu}$, then Exponential errors are likely to be the best assumption.

(2) We have discussed both single and multiple draw assumptions. If the experimenter wishes to select empirically among the assumptions of single and multiple random draws, we suggest he obtain both preference indifference and revealed preference questions and determine empirically whether Propositions 3 and 5 or Propositions 4 and 5 produce the best match among the alternative question formats.

(3) Empirical evidence to date suggests that preference indifference questions yield reasonable predictive validity. But, theoretically, revealed preference questions should do better for estimates of choice probabilities. An experimenter can test this hypothesis by using preference indifference questions for one sample of consumers and revealed preference questions for another sample of consumers.

⁹Rounding off may not produce independent multivariate Normal error, but it serves to illustrate the technique. The tendency of respondents to round to the nearest 0.05 is an interesting future research question.

(4) Common belief in the decision analysis literature holds that we can identify the appropriate functional form with qualitative assessment. For example, see Farquhar (1984) and Keeney and Raiffa (1976, pp. 188–200) for discussion and examples. An experimenter can partially test this assumption by first using the qualitative techniques and then using our estimators of choice probabilities based on both alternative functional forms. Theory suggests that predictions should be most accurate using the functional form identified qualitatively.

We hope that our analyses encourage researchers to address these and other empirical issues and to identify which assumptions are appropriate under which empirical conditions.

This completes our analysis of the implications of measurement error for modeling consumer risk preference with vN-M utility functions. Our emphasis is on uniattributed single parameter functions since they are most commonly used and illustrate the unique advantage of vN-M utility functions. Our main results are practical and flexible. They enable the experimenter to choose among question formats, error assumptions, and functional forms for the utility function. They provide MLE's for the distributions of risk parameters and for choice probabilities. Furthermore, we have indicated (with references) how one can use numerical techniques for the cases where analytic results are unobtainable.

Besides empirical research, topics such as:

- alternative distributional assumptions for measurement errors,
 - estimators for other functional forms,
 - efficient algorithms for estimators which require numerical techniques (equation (13), Proposition 6, and multiattribute extensions),
 - extensions of Propositions 3 through 7 for multiple choice problems,
 - the analysis of cases when the functional form of the utility function is unknown and must be estimated,
 - the extension of our conceptualization to cases where successive draws from $f(r|\lambda)$ are not independent but rather dependent upon anchoring, heuristics, recency effects, problem framing, etc., and
 - the analysis of the implications of vN-M measurement error for sequential decision making and bargaining solutions,
- are all exciting research questions. We hope our analyses provide a beginning.

Appendix 1. Proofs to Propositions

PROOF 1. By definition, if z is a normal variable with mean, μ , and variance, σ^2 , and if $z = \log y$, then y is a lognormal random variable with parameters μ and σ^2 , designated by $y \sim \Lambda(\mu, \sigma^2)$. See Aitchison and Brown (1969). For constant proportional risk aversion, $r(x, u) = -k^{-1} \log u$ or $\log u = -kr(x, u)$. If $r(x, u) \sim N(\hat{\mu}, \hat{\sigma}^2)$, then $-kr(x, u) \sim N(-k\hat{\mu}, k^2\hat{\sigma}^2)$ which is our first result. For constant absolute risk aversion $\log(1 - u) = -xr(x, u)$ yielding the second result.

PROOF 2. Restriction to $r > 0$ implies $r_0 = 0$. $f(r|\lambda)$ induces a distribution on u , $g(u|\lambda)$, according to the following transformation formula: $g(u|\lambda) = f(r(x, u)|\lambda) |\partial r(x, u)/\partial u|$. See Mood and Graybill (1963, p. 224). For exponential error, $f(r|\lambda) = \lambda^{-1} \exp(-r/\lambda)$, and for constant proportional risk aversion $r(x, u) = -k^{-1} \log u$ and $|\partial r/\partial u| = 1/ku$. Substituting in the transformation formula yields $g(u|\lambda) = (\lambda k)^{-1} u^{(1/\lambda k)-1}$ which we recognize as a Beta distribution with parameters $(1/\lambda k)$ and 1. A Beta distribution with parameters c, d is proportional to $u^{c-1}(1-u)^{d-1}$. For constant absolute risk aversion $r(x, u) = -(1/x) \log(1 - u)$ and $|\partial r/\partial u| = 1/[x(1 - u)]$. Substituting in the transformation formula yields $g(u|\lambda) = (\lambda x)^{-1} (1 - u)^{(1/\lambda x)-1}$ which we recognize again as a Beta distribution with parameters 1 and $(1/\lambda x)$.

PROOF 3. We scale $u(x, r)$ such that $u(x_0, r) = 0$ and $u(x_*, r) = 1$. Then, for the binary choice problem $\hat{P}_1 = \text{Prob}[au(x_1, \tilde{r}) > bu(x_2, \tilde{r})]$ where $\tilde{r}$ indicates random variable. Substituting for the constant proportional risk aversion utility function, $u(x, r) = (x - x_0)^r / (x_* - x_0)^r$. This yields that

$$\hat{P}_1 = \text{Prob}[\tilde{r} > \log(\beta/a)] / \log[(x_1 - x_0)/(x_2 - x_0)] = \text{Prob}[\tilde{r} > \kappa \log(\beta/a)].$$

Recognizing that $r \sim N(\hat{\mu}, \hat{\sigma}^2)$ and $\Phi[(\hat{\mu} - z)/\hat{\sigma}] = \text{Prob}[\tilde{r} > z]$ yields the result. The result is only approxi-

mate since we ignore $r < 0$ which occurs with low probability when μ is significantly greater than σ . Now substituting the infinite range constant absolute risk aversion utility function, $u(x, r) = 1 - e^{-rx}$, into $au(x_1, r) = bu(x_2, r)$ yields the equation for r_c . Note that as $r \rightarrow \infty$, $u(x, r) \rightarrow 1$, and Alternative 2 will be preferred since $\beta > \alpha$. As $r \rightarrow 0$, Alternative 1 will be preferred if $\alpha x_1 > \beta x_2$ since $u(x, r)$ approaches linearity. Thus, if there is only one solution to the equation for r_c , $\hat{P}_1 = \text{Prob}[0 < \tilde{r} < r_c]$. For $\alpha x_1 > \beta x_2$, there is only one solution to the equation for $r_c > 0$. We provide a proof of this fact in Lemma 1, below. If $\alpha x_1 < \beta x_2$, then $au(x_1, r) < bu(x_2, r)$ for $r > 0$, hence $\hat{P}_1 = 0$.

LEMMA 1. Assume $\beta > \alpha$ and $x_1 > x_2$, then the equation $\beta \exp(-r_c x_2) - \alpha \exp(-r_c x_1) = \beta - \alpha$ has at most one solution for $r_c > 0$.

PROOF. First, rewrite the equation in functional form:

$$E(r) = \alpha[1 - \exp(-rx_1)] - \beta[1 - \exp(-rx_2)] \quad (\text{A1})$$

recognizing $x_1 > x_2$ and $\beta > \alpha$. Alternative 1 will be chosen whenever $E(r) > 0$.

By a Taylor expansion $E(r) \approx r(\alpha x_1 - \beta x_2)$ as $r \rightarrow 0$. Let $E(0^+) = \lim_{r \rightarrow 0^+} E(r)$ and let $E(\infty) = \lim_{r \rightarrow \infty} E(r)$. Then $E(0^+) > 0$ if $\alpha x_1 > \beta x_2$ and $E(0^+) < 0$ if $\alpha x_1 < \beta x_2$. By direct substitution $E(\infty) = \alpha - \beta < 0$ since $\alpha < \beta$.

Now differentiate $E(r)$ yielding $E'(r) = dE(r)/dr = \alpha x_1 \exp(-rx_1) - \beta x_2 \exp(-rx_2)$. Setting the derivative equal to zero yields $r^* = (\log \alpha x_1 - \log \beta x_2)/(x_1 - x_2)$. Since $x_1 > x_2$, $r^* > 0$ iff $\alpha x_1 > \beta x_2$. Furthermore, $E'(0^+) = \alpha x_1 - \beta x_2$ thus $E'(0^+) > 0$ iff $\alpha x_1 > \beta x_2$.

Assume $\alpha x_1 < \beta x_2$, then $E(0) < 0$, $E(\infty) < 0$, and $E(r)$ is monotonically decreasing in the range $(0, \infty)$. Thus, there is no solution to (A1) for $r_c > 0$. If $\alpha x_1 = \beta x_2$, the only solution for $r > 0$ is $r_c = 0$.

Assume $\alpha x_1 > \beta x_2$, then $E(0^+) > 0$, $E'(0^+) > 0$, and $r^* > 0$. Thus, $E(r) > 0$ for $r < r^*$. Now $E(r^*) > 0$, $E(\infty) < 0$, and $E(r)$ is monotonically decreasing in the range (r^*, ∞) . Thus, there is exactly one solution to (A1) for $r_c > 0$ and it occurs in the range (r^*, ∞) . Note that we have also proven that $E(r) > 0$ for $r \in (0, r_c)$ and $E(r) < 0$ for $r \in (r_c, \infty)$, thus Alternative 1 can only be chosen when r is in the range $(0, r_c)$.

PROOF 4. The results follow the same arguments used in Proposition 3 except $\text{Prob}[\tilde{r} > z] = \exp(-z/\hat{\lambda})$. The result is exact since $\text{Prob}[r < 0] = 0$.

PROOF 5. Alternative 1 will be chosen if $au(x_1, \tilde{r}_1) > bu(x_2, \tilde{r}_2)$. Consider first normal error theory. Rearranging terms, this condition becomes $\log u(x_1, \tilde{r}_1) - \log u(x_2, \tilde{r}_2) > \log(\beta/\alpha)$. Using Proposition 1, the left-hand side of the inequality is distributed as $N[\hat{\mu}(k_2 - k_1), \hat{\sigma}^2(k_1^2 + k_2^2)]$ and the result follows from the recognition that $k_2 - k_1 = \log[(x_1 - x_0)/(x_2 - x_0)] = \kappa^{-1}$.

Now consider exponential error theory. Again rearranging terms indicates that Alternative 1 will be chosen if $u(x_1, \tilde{r}_1)/u(x_2, \tilde{r}_2) > \beta/\alpha$. Let $\tilde{u}_i = u(x_i, \tilde{r}_i)$ then by Proposition 2 and the assumption of independent densities the joint p.d.f., $g(u_1, u_2)$, is given by:

$$g(u_1, u_2) = (\hat{\lambda}^2 k_1 k_2)^{-1} (u_1)^{(1/\hat{\lambda} k_1) - 1} (u_2)^{(1/\hat{\lambda} k_2) - 1}.$$

Define $\tilde{z} = \tilde{u}_1/\tilde{u}_2$ and $\tilde{t} = \tilde{u}_2$ then the joint p.d.f. of $\tilde{z}$ and $\tilde{t}$ is obtained using a Jacobian transformation: $f_{\tilde{z}, \tilde{t}}(z, t) = q_1 q_2 t^{q_1 - 1} (t)^{q_1 + q_2 - 1}$ where $q_i = (1/\hat{\lambda} k_i)$. Integrating out t yields the marginal distribution for z :

$$f_z(z) = \begin{cases} \frac{q_1 q_2}{q_1 + q_2} z^{q_1 - 1}, & 0 < z < 1, \\ \frac{q_1 q_2}{q_1 + q_2} z^{-q_2 - 1}, & z > 1. \end{cases}$$

Since $(\beta/\alpha) > 1$,

$$\hat{P}_1 = \text{Prob}[z > \beta/\alpha] = \int_{\beta/\alpha}^{\infty} f_z(z) dz = [q_1/(q_1 + q_2)](\beta/\alpha)^{-q_2}.$$

Finally, substituting $q_i = (1/\hat{\lambda} k_i)$ into the above expression yields the result.

PROOF 6. The proof is similar to that of Proposition 5. For NORMAL errors we use the limited reproductive properties of the lognormal distribution (Aitchison and Brown 1969) and some algebra. See Barouch and Kaufman (1976) for issues involving sums of lognormal random variables.

PROOF 7. First we recognize that the expected utility for a constant absolute risk aversion utility function with outcomes described by $f(x)$ is given by

$$E(u) = \int u(x, r) f(x) dx = 1 - \int e^{-rx} f(x) dx = 1 - M(r)$$

where $M(r)$ is the exponential transform of $f(x)$. See Keeney and Raiffa (1976, p. 201), Drake (1967, Chapter 3). For the Normal distribution, $M(r) = \exp(-rm + r^2 v^2/2)$. Thus,

$$\hat{P}_1 = \text{Prob}[1 - \exp(-rm_1 + r^2 v_1^2/2) > 1 - \exp(-rm_2 + r^2 v_2^2/2)].$$

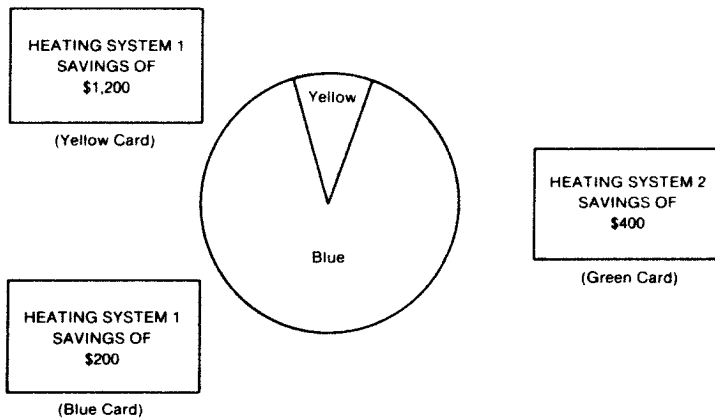


FIGURE A1. Layout of Props for Warm-up Questions.

Simplifying yields $\hat{P}_1 = \text{Prob}[\tilde{r} < 2(m_1 - m_2)/(v_1^2 - v_2^2)]$. Finally, substituting the appropriate $f(r|\lambda)$ yields the result.

Appendix 2. Example Questions

Warm-Up

This set of questions allows you to express how important you feel it is that you are certain about the savings you achieve with your new heating system. Most people find these questions difficult but interesting to answer, but feel it is important to express their feelings on this aspect of their preference.

To better understand this question, imagine that you are not sure about the savings you will achieve with your new heating system. In fact, these savings can be between \$200 and \$1,200. The exact savings you achieve will be determined by a game of chance.

Imagine that someone is going to spin this wheel. [Interviewer lays out the wheel and cards, as shown in Figure A1.] If it comes up yellow (chances are 1 in 10) your savings will be \$1,200. If it comes up blue, your savings will be only \$200. We will call this "game of chance" 'Heating System 1'.

Imagine that you are given a choice between 'Heating System 1' and a guaranteed savings of \$400. We call the guaranteed savings 'Heating System 2'. Which system would you prefer? [Most people would prefer the guarantee of \$400.]¹⁰

Let us now assume we improve the odds on savings \$1,200 to 9 in 10 for 'Heating System 1'. [Interviewer changes the size of the yellow area to represent odds of 9 in 10.] Now if the wheel comes up yellow you save \$1,200, if it comes up blue you save only \$200. Now which heating system do you prefer? [In this case most people prefer to take the chance with 'Heating System 1'.]

Preference Indifference Format

Now imagine you are allowed to set the odds for 'Heating System 1'. If the yellow area is very small you prefer the guarantee of \$400; if it is very large you would be willing to take a chance on the \$1,200 savings. Try to find some "indifference" setting in which you can not make up your mind between 'Heating System 1' and Heating System 2'.

That is, set the yellow area of the wheel for 'Heating System 1' such that if it were larger you would take a chance on 'Heating System 1' and if it were smaller you would prefer the guarantee of \$400 given by 'Heating System 2'.

[After the respondent sets the wheel, the interviewer checks the setting by challenging the respondent to make a choice. If the respondent can make a choice, the interviewer encourages him to modify his probability. The interviewer keeps iterating the question until the respondent is truly indifferent. At this point, the interviewer records the answer and proceeds to the next preference indifference question.]

Revealed Preference Format

Now imagine that you are given the choice between two games of chance representing heating systems A and B, respectively.

¹⁰In a personal interview, we also allow for the rare individual who is strongly risk seeking and prefers the lottery.

In 'Heating System A', the size of the yellow area is 30% of the total. If the yellow area comes up, your savings are \$1,200. If the blue area comes up, you save only \$200.

In 'Heating System B', the size of the yellow area is 50% of the total. If the yellow area comes up, your savings are \$600. If the blue area comes up, you save only \$200.

[The respondent is shown schematics such as Figure 5 or a modified Figure A1 to represent Heating Systems A and B.]

Which system do you prefer, 'Heating System A' or 'Heating System B'?

[Note. The warm-up section can be modified to introduce the respondent to choices between the lotteries. We suspect that after the first or second revealed preference question, subsequent questions can be streamlined.

These example questions are for a personal interview. A mail questionnaire will require modification in format. Pretesting will improve the questions for each application.]

References

- AITCHISON, J. AND J. C. BROWN, *The Lognormal Distribution with Special Reference to its Uses in Economics*, Cambridge University Press, England, 1969.
- BAROUCH, E. AND G. M. KAUFMAN, "On Sums of Lognormal Random Variables," Working Paper #831-76, Sloan School of Management, M.I.T., Cambridge, Mass., February 1976.
- BATEMEN, H. (Bateman Manuscript Project), *Tables of Integral Transforms*, McGraw Hill Book Co., Inc., New York, 1954.
- BOYD, J. H. AND R. E. MELLMAN, "The Effect of Fuel Economy Standards on the U.S. Automobile Market: A Hedonic Demand Analysis," *Transportation Res.*, 14A, 5-6 (October-December 1980), 367-378.
- CHOFFRAY, J. J. AND G. L. LILJEN, "The Market for Solar Cooling: Perceptions, Response, and Strategy Implications," *Studies in Management Sci.*, 10 (1978), 209-226.
- CURRIM, I. AND R. SARIN, "A Comparative Evaluation of Multiattribute Consumer Preference Models," *Management Sci.*, 30 (May 1984), 543-561.
- DAGANZO, C., *Multinomial Probit: The Theory and its Applications to Demand Forecasting*, Academic Press, New York, 1979.
- DEGROOT, M. H., *Optimal Statistical Decisions*, McGraw Hill, New York, 1970.
- DOMENCICH, T. A. AND D. MCFADDEN, *Urban Travel Demand: A Behavioral Analysis*, American Elsevier Publishing Co., New York, 1975.
- DRAKE, A. W., *Fundamentals of Applied Probability Theory*, McGraw Hill Book Co., New York, 1967.
- ELIASHBERG, J., "Consumer Preference Judgments: An Exposition with Empirical Applications," *Management Sci.*, 26, 1 (January 1980), 60-77.
- FARQUHAR, P. H., "A Survey of Multiattribute Utility Theory and Applications," *Studies in Management Sci.*, 6 (1977), 59-89.
- , "Utility Assessment Methods," Working Paper, University of California, Davis, November 1982, *Management Sci.*, 30 (November 1984).
- FISHBURN, P. C., "von Neumann-Morgenstern Utility Functions on Two Attributes," *Oper. Res.*, 22, 1 (January-February 1974), 35-45.
- AND G. A. KOCHENBERGER, "Two-Piece von Neumann-Morgenstern Utility Functions," *Decision Sci.*, 10 (1979), 503-518.
- FRIEDMAN, M. AND L. J. SAVAGE, "The Expected-Utility Hypothesis and the Measurability of Utility," *J. Political Econom.*, 60 (1952), 463-474.
- GIRI, N. C., *Multivariate Statistical Inference*, Academic Press, New York, 1977.
- GREEN, P. E., *Analyzing Multivariate Data*, The Dryden Press, Hinsdale, Ill., 1978.
- AND V. SRINIVASAN, "Conjoint Analysis in Consumer Research: Issues and Outlook," *J. Consumer Res.*, 5, 2 (September 1978), 103-123.
- HAUSER, J. R., "Consumer Preference Axioms: Behavioral Postulates for Describing and Predicting Stochastic Choice," *Management Sci.*, 24 (September 1978), 1131-1341.
- AND G. L. URBAN, "A Normative Methodology for Modeling Consumer Response to Innovation," *Oper. Res.*, 25, 4 (July-August 1977), 579-629.
- AND ———, "Direct Assessment of Consumer Utility Functions: von Neumann-Morgenstern Theory Applied to Marketing," *J. Consumer Res.*, 5 (March 1979), 251-262.
- HERSHEY, J. C., H. C. KUNREUTHER AND P. J. H. SCHOEMAKER, "Sources of Bias in Assessment Procedures for Utility Functions," *Management Sci.*, 28 (August 1982), 936-954.
- HERSTEIN, I. N. AND J. MILNOR, "An Axiomatic Approach to Measurable Utility," *Econometrica*, 21 (1953), 291-297.
- INGENE, C. A., "Risk Management by Consumers," Working Paper, University of Texas at Dallas, Richardson, October 1981.
- JENSEN, N. E., "An Introduction to Bernoullian Utility Theory. I. Utility Functions," *Swedish J. Econom.*, 69 (1967), 163-183.

- JUSTER, F. T., "Consumer Buying Intentions and Purchase Probability: An Experiment in Survey Design," *J. Amer. Statist. Assoc.*, 61 (September 1966), 658-696.
- KALWANI, M. U. AND A. J. SILK, "On the Reliability and Predictive Validity of Purchase and Intention Measures," *Marketing Sci.*, 1, 3 (Summer 1982), 243-285.
- KEENEY, R. L., "Utility Functions for Multiattributed Consequences," *Management Sci.*, 18 (1972), 276-287.
- AND H. RAIFFA, *Decisions with Multiple Objectives: Preferences and Value Tradeoffs*, John Wiley & Sons, New York, 1976.
- LUCE, R. D. AND P. SUPPES, "Preference, Utility, and Subjective Probability," in R. D. Luce, R. R. Bush, and E. Galanter (Eds.), *Handbook of Mathematical Psychology*, John Wiley & Sons, New York, 1965, 249-410.
- MANSKI, C. F. AND D. MCFADDEN, *Structural Analysis of Discrete Data with Econometric Applications*, MIT Press, Cambridge, Mass., 1981.
- MARSCHAK, J., "Rational Behavior, Uncertain Prospects, and Measurable Utility," *Econometrica*, 18 (1950), 111-141.
- MCFADDEN, D., "Quantal Choice Analysis: A Survey," *Ann. Econom. and Social Measurement*, (1976), 363-390.
- , "Econometric Models for Probabilistic Choice Among Products," *J. Business*, 53, 3, 2 (July 1980), 513-530.
- MEYER, R. F. AND J. W. PRATT, "The Consistent Assessment and Fairing of Preference Functions," *IEEE Trans. Systems Sci. Cybernetic.*, SSC-4, 3 (September 1968), 270-278.
- MOOD, A. M. AND F. A. GRAYBILL, *Introduction to the Theory of Statistics*, McGraw-Hill Book Co., New York, 1963.
- MORRISON, D. G., "Purchase Intentions and Purchase Behavior," *J. Marketing*, 43 (Spring 1979), 65-74.
- PRAS, B. AND J. O. SUMMERS, "Perceived Risk and Composition Models for Multiattribute Decisions," *J. Marketing Res.*, 15 (August 1978), 429-437.
- PRATT, J. W., "Risk Aversion in the Small and in the Large," *Econometrica*, 32 (1964), 122-136.
- SCHOEMAKER, P. J. H., "Behavioral Issues in Multiattribute Utility Modeling and Decision Analysis," *Organizations: Multiple Agents with Multiple Criteria*, J. Morse (Ed.), Springer-Verlag, New York, 1981, 447-464.
- AND C. C. WAID, "An Experimental Comparison of Different Approaches to Determining Weights in Additive Utility Models," *Management Sci.*, 28 (February 1982), 182-196.
- THURSTONE, L., "A Law of Comparative Judgment," *Psychological Rev.*, 34 (1927), 273-286.
- VON NEUMANN, J. AND O. MORGENTHAU, *Theory of Games and Economic Behavior*, Princeton University Press, Princeton, N.J., 1947.

Copyright 1985, by INFORMS, all rights reserved. Copyright of Management Science is the property of INFORMS: Institute for Operations Research and its content may not be copied or emailed to multiple sites or posted to a listserv without the copyright holder's express written permission. However, users may print, download, or email articles for individual use.