

AD-A108 180

ANG (A H-S) AND ASSOCIATES URBANA IL F/6 12/1
STATISTICAL DECISION METHODS FOR SURVIVABILITY AND VULNERABILITY--ETC(U)
NOV 78 A H ANG, W TANG DNA001-78-C-0277

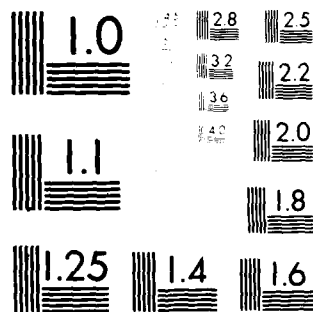
UNCLASSIFIED

DNA-4718F

NL

$$\begin{aligned} & \Delta(\gamma) \leq \\ & \frac{1}{2} \left(\frac{1}{2} \gamma + \frac{1}{2} \gamma \right) \end{aligned}$$

END
DATE
FILMED
01-82
DTIC



MICROCOPY RESOLUTION TEST CHART
NATIONAL BUREAU OF STANDARDS-1963-A

(12)

LEVEL 31

DNA 4718F

AD A108180

STATISTICAL DECISION METHODS FOR SURVIVABILITY AND VULNERABILITY ASSESSMENTS OF STRATEGIC STRUCTURES

A. H-S. Ang
N. H. Tang
A. H-S. Ang and Associates
402 Yankee Ridge Lane
Urbana, Illinois 61801

1 November 1978

Final Report for Period 17 May 1978—31 October 1978

CONTRACT No. DNA 001-78-C-0277

APPROVED FOR PUBLIC RELEASE;
DISTRIBUTION UNLIMITED.

THIS WORK SPONSORED BY THE DEFENSE NUCLEAR AGENCY
UNDER RDT&E RMSS CODE B344078464 V99QAXSC06162 H2590D.

DTIC FILE COPY

Prepared for
Director
DEFENSE NUCLEAR AGENCY
Washington, D. C. 20305

DTIC
ELECTE
DEC 8 1981
S B D

81 12 08 221

Destroy this report when it is no longer
needed. Do not return to sender.

PLEASE NOTIFY THE DEFENSE NUCLEAR AGENCY,
ATTN: STTI, WASHINGTON, D.C. 20305, IF
YOUR ADDRESS IS INCORRECT, IF YOU WISH TO
BE DELETED FROM THE DISTRIBUTION LIST, OR
IF THE ADDRESSEE IS NO LONGER EMPLOYED BY
YOUR ORGANIZATION.



SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

DD FORM 1 JAN 73 1473

EDITION OF 1 NOV 65 IS OBSOLETE

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

UNCLASSIFIED

SECURITY CLASSIFICATION OF THIS PAGE(When Data Entered)

→ function is the central requirement for the implementation of statistical decision concepts to practical strategic decision problems. In this regard, innovative modeling and realistic assumptions of the underlying physical problem for statistical decision analysis are often necessary; for this reason, implementation is seldom straight-forward.

Applications are illustrated for several classes of problems, including sampling and estimation, evaluation of field and laboratory test data, analysis and planning of test programs, and the analysis of a retaliatory system.



UNCLASSIFIED

SECURITY CLASSIFICATION OF THIS PAGE(When Data Entered)

PREFACE

This is the final report of a study conducted for the Strategic Structures Division of the Defense Nuclear Agency under Contract No. 001-78-C-0277 with A. H-S. Ang and Associates. The work was sponsored by the DNA under RDT&E RMSS Code B3440 78464 V99QAXSC06162 H2590D.

This is the second report concerned with the potential applications of probability and statistical concepts in the analysis of survivability/vulnerability of strategic structures. The first report was published in July 1978 as Report NO. DNA 4907F entitled "Approximate Probabilistic Methods for Survivability/Vulnerability Analysis of Strategic Structures."

Capt. M. Moore was the Contract Officer Representative of the project; his assistance throughout the period of the contract is acknowledged.

Accession For	
NTIS GRA&I	☒
DTIC TAB	☐
Unannounced	☐
Justification	
Distribution/	
Availability Codes	
Avail and/or	
Dist	Special
A	

TABLE OF CONTENTS

SECTION	
PREFACE-	1
LIST OF FIGURES-	3
I. INTRODUCTION -	5
II. BASIC STATISTICAL DECISION CONCEPTS-	7
2.1 The Decision Model-	7
2.2 Decision Criteria -	12
2.3 Analysis of a Decision Tree -	13
2.4 Procedure and Requirements for Implementation -	18
III. UTILITY AND UTILITY FUNCTION -	22
3.1 Utility Measure -	22
3.2 Common Types of Utility Functions -	23
3.3 Sensitivity of Expected Utility to Form of Utility Function-	25
3.4 Sensitivity of Decision -	27
IV. DECISION CONCEPTS IN SAMPLING AND ESTIMATION -	31
4.1 Bayesian Sampling -	31
4.2 Bayes' Point Estimator-	34
4.3 Optimal Sample Size -	36
V. EVALUATION OF FIELD AND LABORATORY DATA-	41
5.1 Introductory Remarks-	41
5.2 Field versus Laboratory Tests -	41
5.3 Posterior (updated) Information -	43
5.4 Measure of Information -	44
VI. ANALYSIS AND PLANNING OF TEST PROGRAMS -	48
6.1 Introduction-	48
6.2 Preliminary Considerations-	48
6.3 Measure of Benefit (i.e. Information Gain)-	50
6.4 Optimal Test Strains-	51
6.5 Optimal Number of Test Tunnels-	56
VII. ANALYSIS OF A RETALIATORY SYSTEM -	60
7.1 Premise of Analysis -	60
7.2 The Probability of Attack -	60
7.3 On the Probability of Survival-	66
7.4 A Trade-Off Problem -	67
REFERENCES -	70
APPENDIX A - EVALUATION OF UTILITY IN A DECISION TREE -	71
APPENDIX B - DETERMINATION OF UTILITY FUNCTION -	73

LIST OF FIGURES

<u>Figure</u>		<u>Page</u>
2.1	A decision tree example - - - - -	8
2.2	Decision tree with one decision node- - - - -	14
2.3	Decision tree with subtrees - - - - -	14
2.4	Simplified decision tree- - - - -	16
2.5	Sequential decision analysis- - - - -	16
3.1	Decision tree for example of sensitivity analysis - - - - -	29
3.2	Expected utility as function of probability p - - - - -	30
4.1	Prior, likelihood and posterior functions for parameter μ - -	33
4.2	Selection of point estimator- - - - -	35
4.3	Decision tree for optimal sample size - - - - -	37
6.1	Failure strain of tunnel or rock cavity - - - - -	49
6.2	Failure strain versus tunnel diameter - - - - -	49
6.3	Decision tree for test strain - - - - -	51
6.4	Decision tree for test strains of two tunnels - - - - -	52
6.5	Decision tree for test strains of n tunnels - - - - -	55
6.6	Decision tree for determining number of test tunnels- - - - -	56
7.1	Cluster of four shelters in a "square"- - - - -	62
7.2	Safety factor versus survival probability - - - - -	66
7.3	Probability of failure- - - - -	68

I. INTRODUCTION

Practical decisions are invariably made under conditions of uncertainty; i.e. short of complete information. Indeed, decisions (whether large or small) are seldom made with complete or perfect information. Intelligent decision making obviously requires the availability of good information; the proper analysis of this information, of course, is equally important. In the case of decision situations involving uncertainty, the necessary information may have to be expressed in terms of probability; its analysis would naturally require probabilistic concepts and modeling. Moreover, the criterion or basis for decision, or choice of action, may have to be defined statistically, for example, on the basis of maximum expected utility.

Often, the judgment and preference of a decision maker have an overriding influence on the final choice of action. In some cases, no amount of formal analysis can really influence the intuitive bias of a decision maker. Nevertheless, under conditions of uncertainty, and in complex situations, there are certain analytical tools that may be pertinent and useful for systematically synthesizing available information (even though incomplete) and subjective judgment, taking into consideration any preference or bias of the decision maker. The concepts and tools of statistical decision theory provide the framework for these purposes.

A formal statistical decision analysis will depend on the acceptability of the decision basis, i.e. decision criterion, as well as on the definition and measure of consequence such as a utility function, used in the process. The implementation of the decision concepts summarized herein, therefore, necessarily requires the establishment of a realistic utility function appropriate for a given problem; for this reason, the application of these concepts in real decision problems must be problem-specific. In practice, decisions invariably involve a choice of an action from among a finite number of feasible alternatives. A formal decision analysis, therefore, must include the evaluation of the consequence or payoff associated with each of the feasible alternatives.

In short, decisions can be made with or without any formal analysis; the intuition and judgment of a decision maker often have overriding influence on his choice of alternatives. Under conditions of uncertainty, however, statistical decision theory provides a systematic framework on which the decision making process may be intelligently carried out. This involves the analysis of available information, including judgmental information, and the assessment of the potential consequences of each feasible alternative action. Perfect decisions (i.e. free of the possibility of a wrong action) should not be expected even with an elaborate statistical analysis. A formal decision making process only provides the framework for systematically analysing all available information (including subjective judgments) and for consistently and objectively evaluating the potential consequences, taking into consideration any preferential bias of the

decision maker. Hopefully, the process will minimize the probability of a wrong choice of alternative; at the very least, it may provide additional insight on the problem.

The basic concepts of statistical decision theory are presented; also, the potential roles and applications of these concepts in strategic problems are emphasized. Implementation of statistical decision concepts, however, is seldom straight-forward. Invariably, innovative formulations and modeling of the specific physical problem will be necessary; examples of these are illustrated with special reference to strategic defense problems. As the numerical results will require extensive computer calculations, the illustrations are limited to the discussions of the principles involved and formulations of the necessary models.

II. BASIC STATISTICAL DECISION CONCEPTS

2.1 The Decision Model

Formally, a statistical decision model may be defined in terms of four major components, as follows:

1. The feasible alternatives.
2. The possible outcomes associated with each alternative.
3. The probabilities associated with each possible outcome.
4. The consequence associated with each alternative and a given outcome.

All of these components may be portrayed conveniently in the form of a decision tree. Figure 2.1 illustrates an example of a decision tree applied to a hypothetical targeting decision problem. The decision tree starts out with a decision node (square node in Fig. 2.1), at which point the decision maker identifies the various feasible alternatives. In the case of the example problem of Fig. 2.1, the choices are

- A_1 = deploy a single large weapon
- A_2 = deploy several small weapons
- A_3 = acquire further intelligence data on the enemy's site condition

For each alternative, there is a chance node (circular node in Fig. 2.1) which is followed by several branches denoting a set of possible outcomes associated with this particular action. In this example, the uncertainty lies in the site condition of the enemy's facility, which is assumed (for simplicity) to be either soft soil (θ_1) or hard rock (θ_2). Presumably, a single large weapon would be more effective in destroying a hard rock site, whereas several small weapons will be more cost-effective for a soft soil condition. In alternatives A_1 and A_2 the probabilities of encountering soft soil or hard rock may be assessed using available geological information combined with judgment; these are indicated at the respective branches as 0.3 and 0.7 in Fig. 2.1. In this example, these probabilities are assumed to be invariant with weapon types. Observe that the outcomes at each chance node are mutually exclusive and collectively exhaustive; consequently, the probabilities associated with all the possible outcomes at a chance node should sum up to unity.

A third alternative, A_3 , is included in the decision tree of Fig. 2.1, which involves acquiring further intelligence data on the enemy's site condition. Depending on whether the intelligence data indicates soft soil (z_1) or hard rock (z_2), the decision maker will have to make a decision at node B as to which weapon system to use. Since the intelligence data may be subject to error, the possible outcomes of the subsequent branches will still consist of θ_1 and θ_2 . However, the probabilities of θ_1 and θ_2 will obviously depend on whether the intelligence data favors one or the other outcome. For example, if intelligence data indicates soft soil (i.e. z_1), the probability of actually encountering soft soil at the site will be high, but not necessarily equal to 1.0.

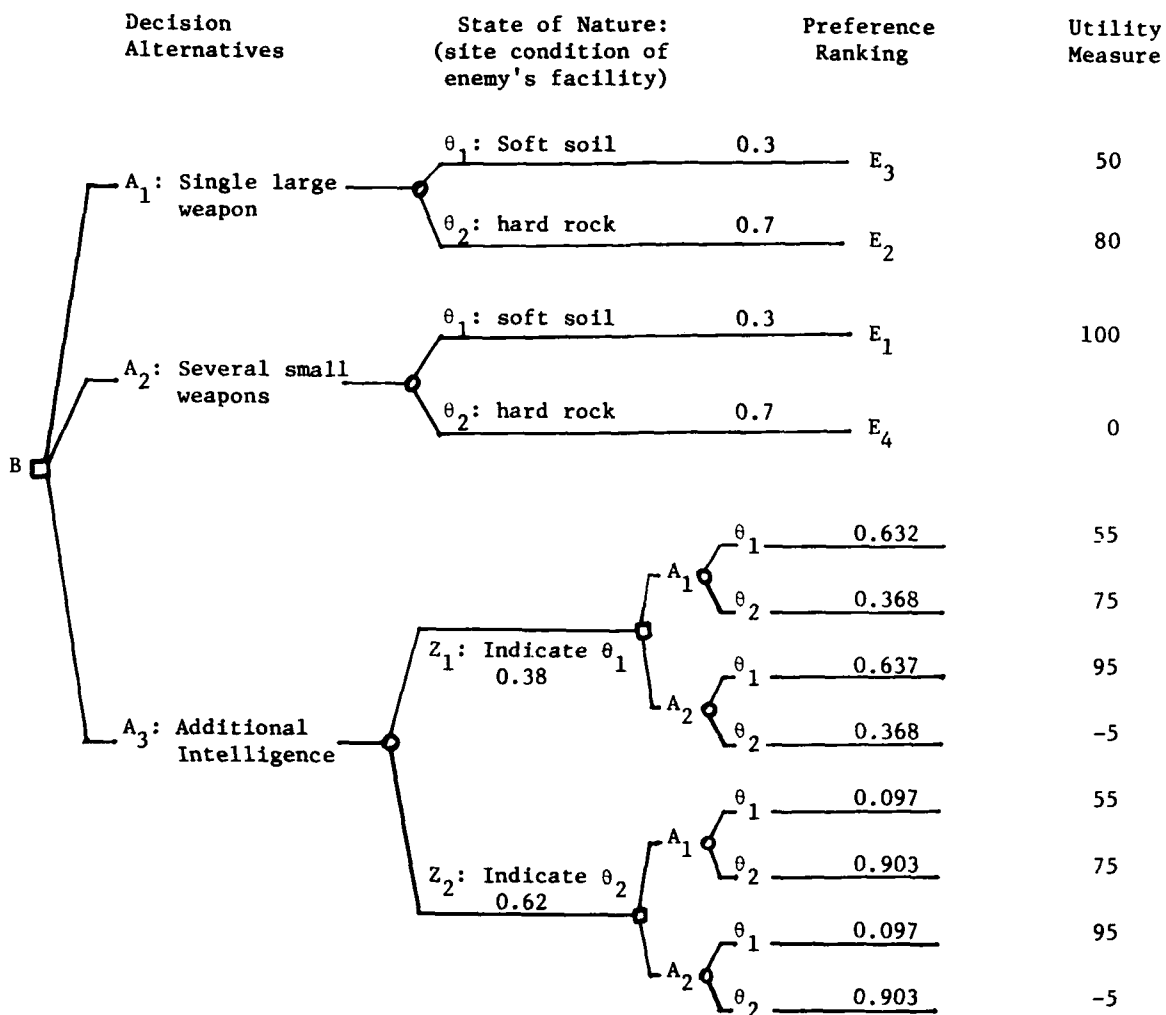


Figure 2.1 A decision tree example

Therefore, conditional probabilities are required here to account for the reliability of the intelligence information.

The consequence associated with each "path" (i.e. a given action combined with a specific outcome) may be measured in terms of a "utility". In general, each path will have its own utility measure as indicated in Fig. 2.1. The utilities for the various paths in Fig. 2.1 were evaluated with the procedure outlined in Appendix A.

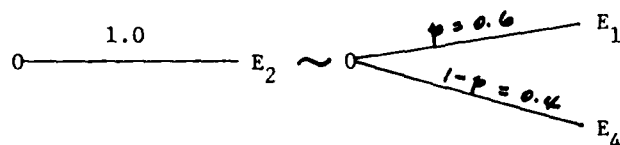
The basic decision in the decision tree of Fig. 2.1 is between alternatives A_1 and A_2 . There are, therefore, essentially four paths that must be ranked in order. The objective of a targeting decision, of course, is to destroy an enemy installation with the optimal weapon. In this regard, assume that a number of small weapons will have maximum destruction if the site of the enemy's installations is of soft soil, whereas using a large weapon on a soft-soil site will result in overkill. With this consideration, the different possibilities are ranked in order of preference as shown in Fig. 2.1, such that $E_1 > E_2 > E_3 > E_4$, in which the symbol $>$ means "is preferred to".

A relative utility value of 100 is then assigned to E_1 and 0 to E_4 ; i.e.,

$$\begin{aligned} u(E_1) &= 100 \\ u(E_4) &= 0 \end{aligned}$$

The relative utility values for E_2 and E_3 are then evaluated by the lottery schemes of Appendix A as follows:

The objective here is to determine the relative utility values for E_2 and E_3 in such a manner that consistency is maintained. For this purpose, consider two lotteries; in one there is certainty of obtaining E_2 , whereas in the other there is a probability p of obtaining E_1 and corresponding probability $(1-p)$ of obtaining E_4 . Clearly, if $p = 1.0$, the second lottery would be preferred; whereas, if $p = 0$ the first lottery would be preferred. Therefore, at some probability p , the two lotteries would be indifferent to the decision maker. Suppose that this is at $p = 0.6$ and $(1-p) = 0.4$, as shown in the following sketch.

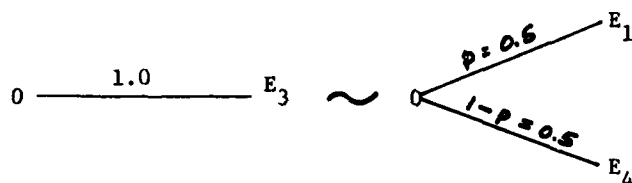


The utility value for E_2 , therefore, is

$$\begin{aligned} u(E_2) &= p u(E_1) + (1-p) u(E_4) \\ &= (0.6 \times 100) + (0.4 \times 0) \\ &= 60 \end{aligned}$$

Similarly, the indifference lotteries for determining the utility for E_3 may be

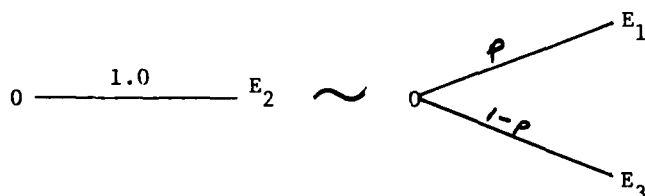
formulated. In this case, suppose that the probability $p = (1-p) = 0.5$. The indifference lotteries for E_3 is, therefore, as shown below.



On the basis of the latter lottery, the corresponding utility value for E_3 is,

$$\begin{aligned} u(E_3) &= p u(E_1) + (1-p) u(E_4) \\ &= (0.5 \times 100) + (0.5 \times 0) \\ &= 50 \end{aligned}$$

It may be observed that the relative utilities between E_2 and E_3 are consistent in accordance with the preference ranking of E_2 and E_3 . However, consistency must be further verified relative to the utility values of E_1 and E_4 . For this purpose, consider other pairs of lotteries; for example, a lottery with certainty for E_2 and another lottery with probabilities p and $(1-p)$ of obtaining E_1 and E_3 , respectively, as follows:

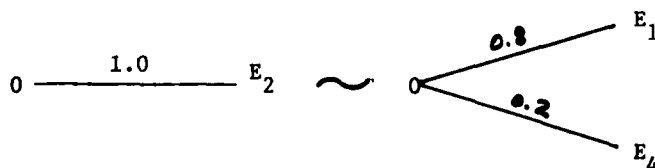


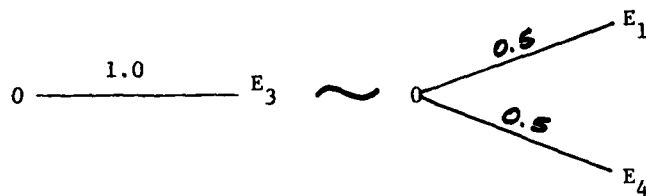
In this case, suppose that the indifference probability is $p = 0.55$ between the two lotteries. Then, the new utility value for E_2 becomes

$$u'(E_2) = (0.55 \times 100) + (0.45 \times 50) = 77.5$$

Since $u(E_2) \neq u'(E_2)$, there is inconsistency in the assignment of relative utilities up to this point.

The entire process may be revised using the following indifference lotteries.



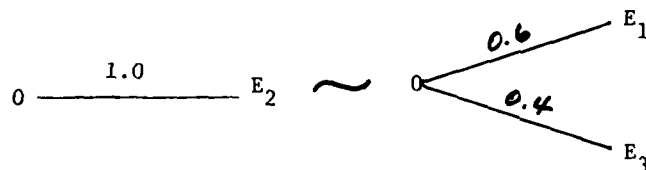


On the basis of the above lotteries,

$$u(E_2) = (0.8 \times 100) + (0.2 \times 0) = 80$$

$$u(E_3) = (0.5 \times 100) + (0.5 \times 0) = 50$$

Again, to verify consistency consider the following indifference lottery:



The new utility value for E_2 then becomes

$$\begin{aligned} u'(E_2) &= p u(E_1) + (1-p) u(E_3) \\ &= (0.6 \times 100) + (0.4 \times 50) \\ &= 80 \end{aligned}$$

Therefore, $u(E_2) = u'(E_2)$. Hence, complete consistency has been achieved.

The utility values for the four paths, therefore, are:

$$\begin{aligned} u(E_1) &\approx 100 \\ u(E_2) &\approx 80 \\ u(E_3) &\approx 50 \\ u(E_4) &\approx 0 \end{aligned}$$

These utility values are also shown in Fig. 1.1 for the first four paths corresponding to alternatives A_1 and A_2 .

With regard to the paths associated with alternative A_3 , the relative utility values should include the cost of acquiring the additional intelligence information. A cost of five utility units is assumed for this purpose. On this basis, the relative utility values for each of the paths associated with alternative A_3 can be obtained by deducting five units from the utility values of the corresponding paths associated with A_1 and A_2 obtained earlier. The results are then those shown in Fig. 2.1.

Generalizations -- The exact configuration of a decision tree will depend on the specific problem. Some trees may have many alternatives whereas others may have many

different possible outcomes. Moreover, instead of discrete branches representing the possible outcomes from a given action, a continuous spectrum of outcomes may sometimes originate from a chance node, such as the possible outcomes defined by the value of a continuous random variable. An alternative could also consist of a sequence of actions and feedbacks, such that the outcome following an earlier action would affect the decision at a subsequent stage. This latter analysis is called a "sequential decision" process and may be represented as a series of decision nodes, branches, chance nodes, etc. in a decision tree.

2.2 Decision Criteria

A formal decision analysis requires a criterion for choice; i.e. a rule for determining what constitutes a "best decision". In this regard, it may be observed that the pay-off values (i.e. utilities) associated with each path depend on the particular outcome which are generally not known deterministically; otherwise, the decision-maker would obviously choose the alternative with the highest utility value. Having constructed and completed the decision tree, there will be multiple utility values, $u_1, u_2, \dots, u_n$, each associated with a particular set of action and outcome, and the corresponding probabilities $p_1, p_2, \dots, p_n$, respectively. This is, of course, a characteristic of decisions under conditions of uncertainty. Depending on the temperament, experience, and degree of risk aversiveness of the decision maker, the "best decision" may mean different things to different decision-makers, and at different times. Nevertheless, certain criteria for decisions may be described as follows:

1. Mini-Max Criterion -- This criterion would suggest that the best decision is the alternative that will minimize the decision-maker's maximum possible loss (or negative utility). In this case, the decision-maker would never venture into anything that may give him substantial positive utility (even though with high probability) as long as there is a finite chance for a substantial loss. This is a pessimistic approach.

2. Maxi-Max Criterion -- With this criterion, the best decision will be to select the alternative that will maximize the decision-maker's maximum possible gain among the alternatives, irrespective of the probability of achieving this maximum gain. Basically, this is an optimistic approach.

The mini-max and maxi-max approaches suggest that decisions are to be made solely on the basis of the utility values; in particular, both approaches ignore the probability values associated with each of the utility values. As a consequence, a considerable amount of information is wasted; moreover, in the long run both the mini-max and maxi-max criteria may accumulate substantial losses.

3. Maximum Expected Monetary Value Criterion -- If the attributes associated with each path in a decision tree can be expressed purely in terms of monetary values, the monetary values of various branches for an alternative could be weighted by their

respective probabilities to obtain the expected monetary value of the alternative; thus

$$E(d_i) = \sum_j p_{ij} d_{ij} \quad (2.1)$$

where d_{ij} denotes the monetary value of the j^{th} branch associated with alternative i , and p_{ij} is the corresponding probability. The best decision then is to choose the alternative with the maximum expected monetary value; i.e.

$$E(d_{\text{opt}}) = \text{Max}_i (\sum_j p_{ij} d_{ij}) \quad (2.2)$$

4. Maximum Expected Utility Criterion -- In the event that monetary value is not the only measure of value (as in the case of a risk averse or risk affirmative decision-maker), or if the attributes describing the value of each path in a decision tree involves non-monetary attributes, a more general approach would be to use utility u_{ij} to denote the consequence associated with a given path. In such a case, the optimal decision would be to choose the alternative with the maximum expected utility; i.e.

$$E(u_{\text{opt}}) = \text{Max}_i (\sum_j p_{ij} u_{ij}) \quad (2.3)$$

where u_{ij} is the utility of the j^{th} branch associated with alternative i .

It may be emphasized that the utility approach described above already accounts for the risk aversiveness of most decision-makers for avoiding, whenever possible, disastrous outcomes; this is accomplished through the assignment of the appropriate utility value or developing the appropriate utility function. Therefore, based on the maximum expected utility criterion, the decision-maker chooses the alternative that would be reasonably acceptable over a wide range of possible outcomes, by properly balancing the true value (gain or loss) of various potential outcomes against their probabilities of occurrence.

2.3 Analysis of a Decision Tree

A basic decision tree consists of one decision node, such as that shown in Fig. 2.2. If the maximum expected utility criterion is used, the expected utility value of each alternative is first computed according to Eq. 2.3 at each chance node, from which the alternative with the maximum expected utility value is the optimal alternative.

Sometimes, a second or subsequent decision is required depending on the outcome of the first decision; in such cases, there will be more than one decision node in a decision tree. An example is shown in Fig. 2.3. In this case, the analysis starts from the right by first identifying the sub-trees, such as B and C as shown in Fig. 2.3. The expected value of each alternative in all the sub-trees are subsequently computed. Assuming that the maximum expected utilities are u_B and u_C , respectively, for sub-trees B and C, the corresponding optimal alternative may then be identified. this point,

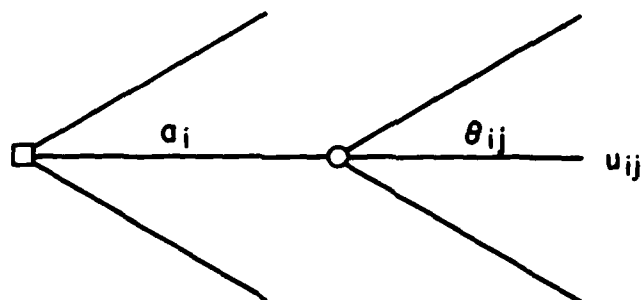


Figure 2.2 Decision tree with one decision node

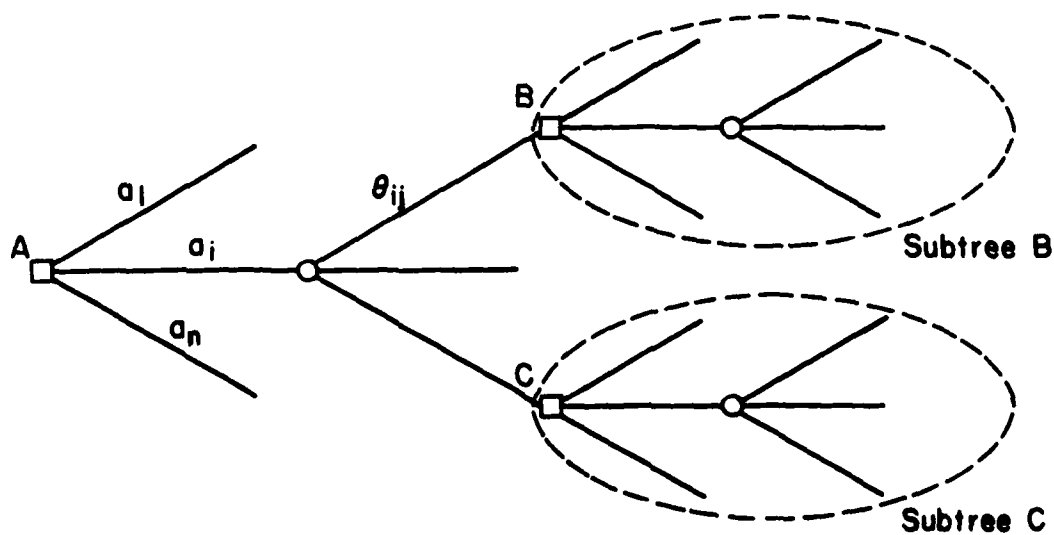


Figure 2.3 Decision tree with subtrees

the decision tree may be simplified as shown in Fig. 2.4, where u_B^* and u_C^* represent the utilities of the entire sub-trees B and C, respectively. The expected utility of each alternative at node A may then be computed and the optimal action selected accordingly.

In a large decision tree, where several decision nodes could appear in a single branch, as in the case of sequential decision analysis, the process of sub-tree analysis will be repeated several times, moving from right to left and advancing from one decision node at a time. Once this backward analysis is completed, the decision process will start from the left. First, the optimal alternative is selected at stage 1; depending on the specific outcome of this first action, the process moves to stage 2. From the results of the backward analysis, the optimal alternative at each decision node, presumably has been identified (as marked in bold lines in Fig. 2.5). In particular, it involves the predetermined optimal alternative at the decision node in this stage (namely, stage 2). Again, depending on the specific outcome at the chance node at this stage, the decision process moves on to stage 3, at which point another predetermined optimal alternative may be selected.

Four types of analysis may be involved in a formal statistical decision analysis; namely:

(1) Prior Analysis -- Literally, a prior analysis involves the calculation of the probabilities and utility values for each path in a decision tree based on existing information and prior assumptions. Any alternative requiring the acquisition of additional information, therefore, will not be included in a prior analysis.

(2) Terminal Analysis -- This involves the analysis and re-evaluation after the necessary experiments or data-gathering schemes have been performed and completed. The probabilities of the possible outcomes are updated accordingly; the remaining analysis are then similar to those of the prior analysis.

(3) Preposterior Analysis -- This primarily addresses the question "should additional information need to be gathered?", and if so which scheme should be used. The decision tree shown in Fig. 2.3 is typical of problems involving preposterior analysis.

(4) Sequential Analysis -- This consists of several decision stages; at each stage, a decision is made on the basis of the observed results of the previous stage. A typical decision tree for this analysis is shown in Fig. 2.5.

The general approach for each of the various analyses described above have been outlined in the previous sections.

Example

The decision tree shown in Fig. 2.1 requires prior as well as preposterior analyses; the latter is associated only with alternative A_3 which requires the acquisition of additional intelligence information before selecting an action, whereas the analysis associated alternatives A_1 and A_2 are strictly prior analysis.

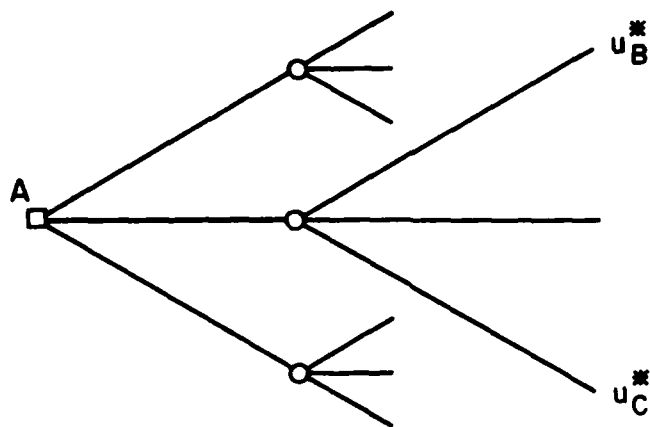


Figure 2.4 Simplified decision tree

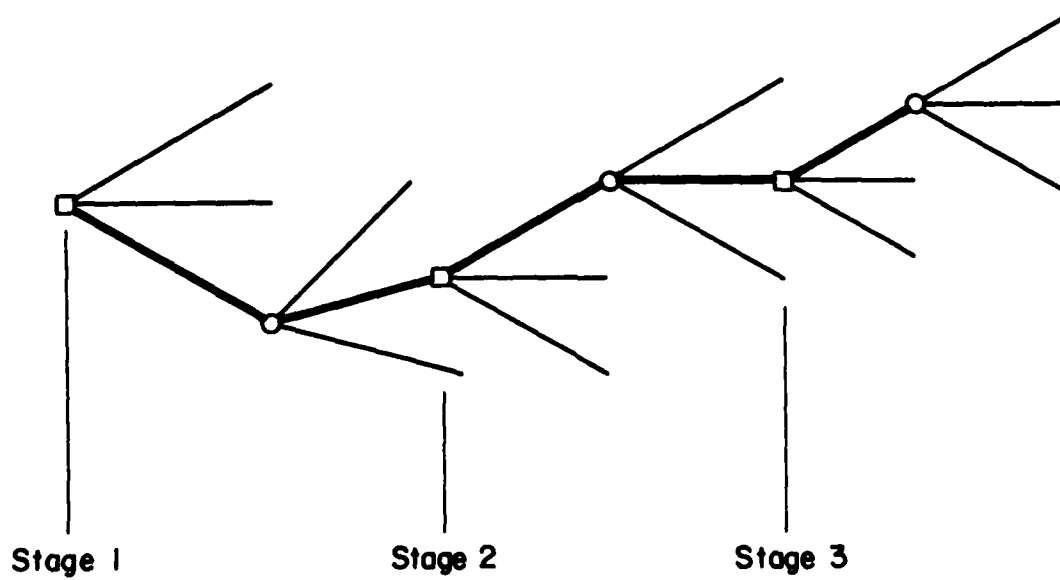


Figure 2.5 Sequential decision analysis

Based on the probability and utility values shown in the decision tree of Fig. 2.1, which are based on existing information and initial assumptions, the expected utilities for alternatives A_1 and A_2 are, respectively, as follows:

$$E(U_{A_1}) = 0.3 \times 50 + 0.7 \times 80 = 71$$

$$E(U_{A_2}) = 0.3 \times 100 + 0.7 \times 0 = 30$$

If no further intelligence information were to be gathered, the maximum expected utility criterion would suggest using alternative A_1 (i.e. deploying a single large weapon in the attack).

However, if the acquisition of additional intelligence information is a viable option available in the decision process, the preposterior analysis for alternative A_3 would proceed from the right of the decision tree as follows:

If the additional intelligence data is z_1 (favoring θ_1), the probabilities are calculated as follows:

$$E(U_{A_1} | z_1) = 0.632 \times 55 + 0.368 \times 75 = 62.36$$

$$E(U_{A_2} | z_1) = 0.632 \times 95 + 0.368 \times (-5) = 58.20$$

whereas if z_2 is the outcome of the additional intelligence gathering, then

$$E(U_{A_1} | z_2) = 0.097 \times 55 + 0.903 \times 75 = 73.06$$

$$E(U_{A_2} | z_2) = 0.097 \times 95 + 0.903 \times (-5) = 4.70$$

implying, therefore, that alternative A_1 (using a single large weapon) would be the preferred action at node B of the decision tree irrespective of the outcome of the additional intelligence gathering. Therefore, the expected utility for action A_3 is (refer to Fig. 2.1)

$$E(U_{A_3}) = 0.38 \times 62.36 + 0.62 \times 73.06 = 68.99$$

Since $E(U_{A_3})$ is slightly less than the expected utility of A_1 , the optimal decision, therefore, is to use a single large weapon in the target plan. In this example, the cost of intelligence gathering was assumed to be -5 utility units; according to the above analysis the additional information is not worth the cost. Of course, if the cost of intelligence is much less than -5 utility units, then it may be worthwhile to acquire additional intelligence information before choosing a weapon system; i.e. alternative A_3 may become optimal.

2.4 Procedure and Requirements in Implementation

As indicated earlier, a statistical decision analysis will consist of the following components: (1) the identification of the feasible alternative actions that require further analysis; (2) the identification of the possible outcomes associated with a given action; (3) the assessment of the probability of occurrence of each of the possible outcomes; (4) the evaluation of the utility associated with each path (i.e. an action and a given outcome).

Identification of Alternatives -- In identifying the feasible alternatives for further decision analysis, one should not simply exhaustively list all the possible actions that may be available. Often, there are options that are clearly not practically feasible or viable; such possible actions are therefore not feasible alternatives. By eliminating the obviously infeasible alternatives, the feasible alternative actions are usually limited in number; these are the ones that require more careful analysis.

Among the feasible alternatives, there may be those that require the acquisition of additional information prior to a final decision. In this latter case, the additional information may involve the use of consultants, or the performance of a field or laboratory test; the latter may also involve decisions on the choice of a test plan.

In the case of a sequential or staged decision process, the alternatives at a subsequent stage may be different from those of the earlier stage. As new information and developments are made available, the available alternatives may change -- the feasible alternatives at an earlier stage may become infeasible at a later stage; conversely, other possible actions which were not feasible at an earlier stage, may become feasible at a later stage.

The Possible Outcomes -- In a well-defined problem, the possible outcomes resulting from a given action should be exhaustively identified. Oftentimes, however, the possible outcomes may not be obvious; for example, they may depend on the opponent's strategy which may not be fully revealed until it is too late. In such cases, of course, the exercise of sound judgment in the identification of the possible outcomes is important.

In the event that the occurrence of a set of possible outcomes is controlled by another process which may be random, an additional chance node could precede the first one, denoting the two levels of outcomes possible in this case. As an example, the parameters of a probability distribution of a random variable may be unknown; thence, the probability of occurrence of a specific value of the random variable will depend on the value of the parameter which is itself a random variable.

Assessment of Probabilities -- Clearly, if more than one outcome is possible subsequent to a given action, the probabilities of occurrence of the various outcomes are important. It may be emphasized that the probabilities are important to denote the relative likelihoods of occurrence of the various possible outcomes, and therefore is the main characteristic of decision making under conditions of uncertainty. The

required probabilities are necessarily estimated or calculated values, which may be evaluated on the following basis:

(1) Based entirely on past observed data-- In this instance, the probabilities of future occurrences of the possible outcomes will be based purely on the statistical observations of the past. The validity of this approach will require a large set of observed data, which is very rarely the case in practice.

(2) Based on subjective judgment of the decision maker -- In this case, the estimated probabilities are purely subjective; consequently, rigorous justification of such probabilities will obviously be difficult. The validity of such probabilities must rely on the quality or credibility of the judgment of the decision maker.

(3) Based on the combination of observed data and subjective judgment -- The Bayes' theorem is the vehicle for combining the two types of information.

Specific methods have been suggested for extracting subjective probabilities; these include the following:

(i) Use of a probability wheel (Spetzler and Stael von Holstein, 1972) and repetitive questioning until the indifference of two lotteries is achieved.

(ii) Establishing the subjective distribution of a random variable by the fractile method (Raiffa, 1968; Schlaefer, 1969), in which the range of a random variable is divided into two equally likely halves, each of which is then subdivided further; the process continues until the required accuracy of the probability calculations in the distribution tail regions is achieved.

(iii) The probability distribution over a given range may be assumed; this could include the uniform and triangular distribution, the latter distribution could be specified with various degrees of skewness to reflect the relative likelihoods over the range.

(iv) Finally, the Delphi method may be used if the subjective judgments and opinions from a group of experts are available. Essentially, this involves the repeated up-dating of individual opinions and estimates in light of the opinions and judgments from the others in the group. In addition to the subjective assessment of probabilities, the Delphi method may be used also for evaluating and assigning utility values.

In order to combine subjective judgment with available data, the probabilities initially obtained on the basis of judgment may be updated in light of the available data; the updating process may be performed on the basis of the Bayes' theorem. Such a situation would occur in the case of alternative A_j in Fig. 2.1; the probabilities available for assessing the reliability of the experiment may be expressed as $P(z_\ell | \theta_j)$ which is the probability of the experiment yielding the outcome z_ℓ given that the actual outcome is θ_j . Since conditional probabilities $P(\theta_j | z_\ell)$ are required, the Bayes' theorem yields

$$P(\theta_j | z_\ell) = \frac{P(z_\ell | \theta_j) P'(\theta_j)}{\sum_j P(z_\ell | \theta_j) P'(\theta_j)} \quad (2.4)$$

where $P'(\theta_j)$ is the prior probability of θ_j which may have to be assessed purely on the basis of judgment, and $P(\theta_j | z_\ell)$ is the updated probability of θ_j following the result of the experiment z_ℓ . Furthermore, the probability that the experiment will yield z_ℓ is given by the theorem of total probability as follows:

$$P(z_\ell) = \sum_j P(z_\ell | \theta_j) P'(\theta_j) \quad (2.5)$$

In the absence of experience or other judgmental information, a diffuse prior probability may be assumed; this implies that $P'(\theta_j)$ is constant for all j . Therefore, no specific bias toward any particular outcome is implicitly assumed with the diffuse prior distribution. A further application of the diffuse prior will be discussed in the section on Bayesian sampling theory.

It might be emphasized that in practice, available observational data is invariably insufficient for decision making purposes, whereas decisions made entirely on the basis of subjective judgment may be difficult to justify and accept. In practice, therefore, some combination of subjective judgment combined with available prior information and observational data is invariably the most viable approach to the calculation of the probabilities for decision making purposes. In short, available information should be used to its fullest extent, supplemented with good judgment and reasonable assumption whenever necessary.

Example

To illustrate the updating process implied in Eq. 2.4, consider the alternative A_3 of the example discussed in Fig. 2.1.

Suppose that the reliability of the intelligence information is 80%; i.e. $P(z_1 | \theta_1) = 0.8$ and $P(z_2 | \theta_2) = 0.8$. Conversely, this also implies that $P(z_1 | \theta_2) = P(z_2 | \theta_1) = 0.2$.

With the prior probability estimates (i.e. may be based purely on judgment) of $P'(\theta_1) = 0.3$ and $P'(\theta_2) = 0.7$, Eq. 2.4 yields

$$\begin{aligned} P(\theta_1 | z_1) &= \frac{P(z_1 | \theta_1) P'(\theta_1)}{P(z_1 | \theta_1) P'(\theta_1) + P(z_1 | \theta_2) P'(\theta_2)} \\ &= \frac{0.8 \times 0.3}{0.8 \times 0.3 + 0.2 \times 0.7} \\ &= 0.632 \end{aligned}$$

and,

$$P(\theta_2|z_1) = 1 - P(\theta_1|z_1) = 0.368$$

Similarly,

$$\begin{aligned} P(\theta_1|z_2) &= \frac{P(z_2|\theta_1) P'(\theta_1)}{P(z_2|\theta_1) P'(\theta_1) + P(z_2|\theta_2) P'(\theta_2)} \\ &= \frac{0.2 \times 0.3}{0.2 \times 0.3 + 0.8 \times 0.7} \\ &= 0.097 \end{aligned}$$

and,

$$P(\theta_2|z_2) = 1 - P(\theta_1|z_2) = 0.903$$

Also, in light of the intelligence data, the probabilities of soft soil or hard rock conditions (i.e. z_1 or z_2) would be calculated according to Eq. 2.5 as follows:

$$\begin{aligned} P(z_1) &= P(z_1|\theta_1) P'(\theta_1) + P(z_1|\theta_2) P'(\theta_2) \\ &= 0.8 \times 0.3 + 0.2 \times 0.7 \\ &= 0.38 \end{aligned}$$

and,

$$\begin{aligned} P(z_2) &= P(z_2|\theta_1) P'(\theta_1) + P(z_2|\theta_2) P'(\theta_2) \\ &= 0.2 \times 0.3 + 0.8 \times 0.7 \\ &= 0.62 \end{aligned}$$

III. UTILITY AND UTILITY FUNCTION

3.1 Utility Measure

Aside from the calculation of the probabilities in a decision tree, the other important requirement is the evaluation of the utility values associated with the various actions and outcomes in a decision tree. For this purpose, a utility function would generally be needed. Unless the consequence of a given action can be expressed purely in monetary terms, a utility function would need to be developed; such a utility function must be suitable for the specific problem under consideration.

The importance of a proper and realistic utility function is central to the application of statistical decision theory. Indeed, the significance of decision theory to practical problems will depend heavily on the successful formulation of realistic and viable utility functions.

A proper utility function must necessarily be problem-specific; it would depend on the gain parameters or loss parameters pertinent to a given problem. For this reason, no single utility function or type of utility function can be applicable to all problems. Nevertheless, there are certain attributes that may be useful to a large class of problems, which may be identified as follows:

1. Utility in Terms of Monetary Value -- For a class of problems, the payoff or consequence from a given action and possible outcome may be expressed in monetary terms. All tangible and intangible attributes are assigned monetary values. In this regard assumptions and value judgment may be necessary for evaluating or assigning monetary costs to all potential consequences, including the cost of a fatality or the cost per unit gain of information.

2. Consequence in Terms of Utility Value -- Monetary value may not be a suitable measure of cost or payoff. For this reason, the concept of utility may be appropriate as a more general measure of payoff, and also for combining several types of attributes. In this regard, two approaches may be considered:

(i) Denote the consequence of each path (i.e. each alternative and a given outcome) by E_i . Rank the results E_i 's in terms of the preference of the decision-maker such that $E_1 > E_2 > \dots > E_n$ where $>$ denotes "is preferred to" and n is the number of paths in the decision tree. Assign the utility values of 1 and 0, respectively, to E_1 and E_n . Conduct an inquiry (see Appendix A) perhaps between the decision analyst and the decision maker, in order to establish a consistent set of utility values for all the other paths E_i , for $i = 2, 3, \dots, n-1$.

(ii) Alternatively, if a scale measure already exists with a specific attribute, such as time or fatalities, or monetary value, a utility function may be established to transform all the different scales into a uniform utility scale. The reason for this transformation is that the original individual scales may not be consistent with the decision-maker's preference. For example, a fatality of 200 may be more than twice

the consequence of a fatality of 100. Considerations such as these may be introduced and should be considered in the formulation of the proper utility function. A procedure for developing such a utility function is summarized in Appendix B.

For a decision problem involving only one attribute, the utility function described above is sufficient for assigning or evaluating the utility value for each path in a decision tree. In a typical case, however, several attributes may be involved. In these cases, a joint (or multi-attribute) utility function is needed for evaluating the expected utility of an alternative action. The determination or formulation of a joint utility function (see Raiffa, 1969, Fishburn, 1970), could be cumbersome especially when the number of attributes becomes large. Assumptions such as the concept of preferential independence and utility independence are often invoked to simplify the assessment and formulation of joint utility functions (Keeney, 1972).

3. Consequence in Terms of Cost-Effectiveness Ratios -- In certain decision problems, the attributes may involve monetary costs and another attribute; the latter may be the number of fatalities saved (in the case of an enemy attack) or the amount of information gained (from performing an experiment), or the damage imposed on an enemy installation (in the case of targeting). In such problems, a suitable measure of utility may be expressed in terms of the unit cost per fatality saved, or unit cost per unit of information gained (e.g. from an experiment).

3.2 Common Types of Utility Functions

Most utility functions are convex; this means that the marginal increase in utility decreases with increasing value of an attribute. The preference behavior of a decision-maker exhibited by a convex utility function is commonly referred to as "risk aversiveness". Most people are risk averse to a certain degree; some may be more risk averse than others. The mathematical forms of utility functions commonly used to model such risk averse behavior would include the following:

1. Exponential Utility Function:

$$U(x) = a + be^{-\gamma x} \quad (3.1)$$

where the parameter γ is the measure of the degree of risk aversion, and a and b are normalization constants. If the utility function is normalized such that $U(0) = 0$ and $U(1) = 1$, then the normalized exponential utility function becomes

$$U(x) = \frac{1}{1-e^{-\gamma}} (1-e^{-\gamma x}) \quad (3.2)$$

2. Logarithmic Type:

$$U(x) = a \ln(x+\beta) + b \quad (3.3)$$

where β is a parameter, generally corresponding to the amount of capital reserve of

the decision-maker; i.e. as β increases, the decision-maker has more utility to spare (such as money), and he becomes less risk averse. A normalized logarithmic utility function may be shown to be:

$$U(x) = \frac{1}{\ln(\frac{1+\beta}{\beta})} [\ln(x+\beta) - \ln \beta] \quad (3.4)$$

3. Quadratic Type

$$U(x) = a(x - 1/2\alpha x^2) + b \quad (3.5)$$

where α is the parameter related to the degree of risk aversion. A normalized quadratic utility function may be shown to be:

$$U(x) = \frac{1}{1-\alpha/2} (x - \frac{1}{2}\alpha x^2) \quad (3.6)$$

The degree of risk aversiveness of a decision-maker is measured by the "risk aversion coefficient,"

$$r(x) = - \frac{u''(x)}{u'(x)} \quad (3.7)$$

where the prime (') denotes derivative with respect to x . This coefficient measures the negative rate of change in curvature of the utility function with respect to a unit change in the slope of the utility function. For the exponential utility function, the risk aversiveness can be shown, using Eq. 3.2 in Eq. 3.7, to be a constant; i.e.

$$r(x) = \gamma. \quad (3.8)$$

Applying Eq. 3.7 with Eqs. 3.4 and 3.6, the coefficients of risk aversion for the normalized logarithmic and quadratic utility functions can be shown to be, respectively,

$$r(x) = \frac{1}{x+\beta} \quad (3.9)$$

and,

$$r(x) = \frac{\alpha}{1-\alpha x} \quad (3.10)$$

which are both functions of x .

Observe that the coefficient of risk aversion does not vary with the attribute in the case of the exponential utility function; whereas in the case of the logarithmic utility function the coefficient of risk aversion decreases with x and in the case of the quadratic utility function the risk aversion increases with x .

The utility function, of course, may also be concave upwards; that is, the marginal increase in utility increases with increasing values of the attribute x . In such cases, the preference behavior of the decision-maker is referred to as risk affinitive. It is believed that this preference behavior is ordinarily not realistic.

3.3 Sensitivity of Expected Utility to Form of Utility Function

Usually, it is difficult to ascertain which type of utility function is most appropriate; e.g. whether it should be of the exponential or quadratic form. The correct choice of the form of the utility function, however, may not be very crucial, especially if the expected utility values are not sensitive to the form of the function.

In order to examine the sensitivity of the expected utility associated with a given action, to the above three forms of utility functions, consider the simple case in which the possible outcomes from an action can be described by the value of a random variable X . In this case, the expected utility of a given action may be expressed as follows.

$$E(U) = \int u(x) f_X(x) dx$$

where $f_X(x)$ is the probability density function of X and $U(x)$ is the utility function. Using the second-order approximation to evaluate the expected utility (see Ang and Tang, 1975), the result is

$$E(U) = u(\bar{x}) + 1/2 \text{Var}(x) + u''(x) \quad (3.11)$$

where $\bar{x}$ and $\text{Var}(x)$ are the mean and variance of the random variable X , and $u''(x)$ is the second derivative of the utility function evaluated at the mean value of X .

Applying Eq. 3.11 to the three types of utility functions described earlier, the second-order approximation of the expected utility becomes, respectively, as follows: For the exponential utility function,

$$E(U) \approx \frac{1}{1-e^{-\gamma}} [1 - e^{-\gamma\bar{x}} - \frac{1}{2} \text{Var}(x) \cdot e^{-\gamma\bar{x}}] \quad (3.12)$$

For the logarithmic utility function,

$$E(U) \approx \frac{1}{\ln(\frac{1+\beta}{\beta})} [\ln(\bar{x}+\beta) - \ln\beta - \frac{1}{2} \text{Var}(x) \cdot \frac{1}{(\bar{x}+\beta)^2}] \quad (3.13)$$

Finally, for the quadratic utility function,

$$E(U) \approx \frac{1}{1-\frac{\alpha}{2}} [\bar{x} - \frac{1}{2} \alpha(\bar{x}^2 + \text{Var } X)] \quad (3.14)$$

In order to compare the expected utility obtained for the three different forms of utility functions, each of the utility functions given in Eqs. 3.2 through 3.6 is calibrated to have the same coefficient of risk aversion at the mean-value of the random variable X . In other words, substituting $\bar{x}$ for X and γ for r in Eqs. 3.9 and 3.10, we obtain

$$\beta = \frac{1}{\gamma} - \bar{x}$$

and,

$$\alpha = \frac{\gamma}{1 + \bar{x} \gamma}$$

A numerical study was performed for the case where $\bar{x} = 0.5$, and $\text{var}(X) = 0.1^2$ and 0.25^2 . The numerical results are summarized in Table 3.1.

From the results shown in Table 3.1, two observations may be deduced, as follows:

1. The expected utility is relatively insensitive to the form of the utility function at a given level of risk aversion; the difference in the expected utility is less than 2% among the three types of utility functions examined herein.

2. The expected utility does not change significantly (at most 20%) over the range of γ (from 0.25 to 1.50) examined herein.

The implication of these observations is that the exact form of the utility function will not be an important factor in the computation of expected utility. Moreover, the risk aversiveness coefficient in the utility function need not be very precise; i.e. any error in the specification of the risk aversiveness coefficient may not cause significant difference in the calculated expected utility. In short, the problem of ascertaining an accurate utility function would not be crucial in the application of statistical decision analysis.

A Related Observation -- As indicated in Eq. 3.11, an approximate expected utility value may be computed on the basis of the mean and variance of the pertinent random variable. This would suggest that the entire probability density function may not be necessary in most decision analysis problems. In practice, the first two statistical moments could be all the information that may be available for a random variable; hence, Eq. 3.11 provides a convenient approximate formula for computing the expected utility of a given action.

Table 3.1 Comparison of expected utility for different forms of utility function

Variance (X)	γ	Expected Utility Value (2nd-order Approximation)		
		Exponential Function	Logarithmic Function	Quadratic Function
0.1^2	0.25	0.530	0.530	0.530
"	0.50	0.563	0.561	0.560
"	1.00	0.618	0.626	0.620
"	1.50	0.672	0.707	0.680
0.25^2	0.25	0.523	0.524	0.523
"	0.50	0.547	0.548	0.547
"	1.00	0.593	0.603	0.594
"	1.50	0.636	0.676	0.641

3.4 Sensitivity of Decision

The question of whether or not the optimal action suggested by a formal decision analysis is sensitive to changes, or possible changes, in the probabilities and/or utilities assigned to the various branches and paths of a decision tree is clearly of interest. In view of the fact that the assignment of probabilities are often based on subjective judgments, the sensitivity of the optimal alternative to variations in the input information clearly deserves some attention.

If the optimal alternative suggested by a formal decision analysis has an expected utility value that is far greater than those associated with the other alternatives, it may be obvious that even significant changes in the input variables would not alter the final results. However, in situations where changes in the input variables could conceivably alter the optimal alternative, the decision analysis could proceed assuming that the probabilities and utilities are variables instead of assigned numbers. An example is presented below demonstrating such a sensitivity analysis for a simple decision tree, in which the pertinent probability factors are treated as variables. It may be observed from this example that the unknown probability (or utility) does not need to be known exactly, so long as it can be estimated to be within a range in which a specific alternative can be shown to be superior.

Example

Consider the hypothetical decision problem of Fig. 2.1. Suppose that the gathering of additional intelligence data is not a viable option. Thus, the only feasible alternatives are A_1 and A_2 . The decision maker believes that there is a high probability that the site will be hard rock; however, the exact value of p is not known. From the decision tree of Fig. 3.1 the expected utility of alternatives A_1 and A_2 are

$$E(U_{A_1}) = 50 (1-p) + 80 p = 50 + 30p$$

$$E(U_{A_2}) = 100 (1-p)$$

Figure 3.2 shows a plot of these two expected utilities as a function of p . It can be observed that for $p \leq 0.385$, $E(U_{A_2}) > E(U_{A_1})$, hence using several small weapons (A_2) is a better alternative; whereas for $p > 0.385$, using a single large weapon (i.e. A_1) is the better alternative. Therefore, even though the exact probability of hard rock at the enemy's site is not known, a knowledge that $p > 0.385$ is sufficient for deciding on a single large weapon.

If several probabilities and utilities have large uncertainties, a conservative approach may be taken to investigate whether the optimal alternative remains valid; that is, the extreme probabilities and utilities may be used in the decision analysis to obtain a conservative choice of action. This may suggest an action that may or may not be the same as the action obtained on the basis of the expected utility criterion.

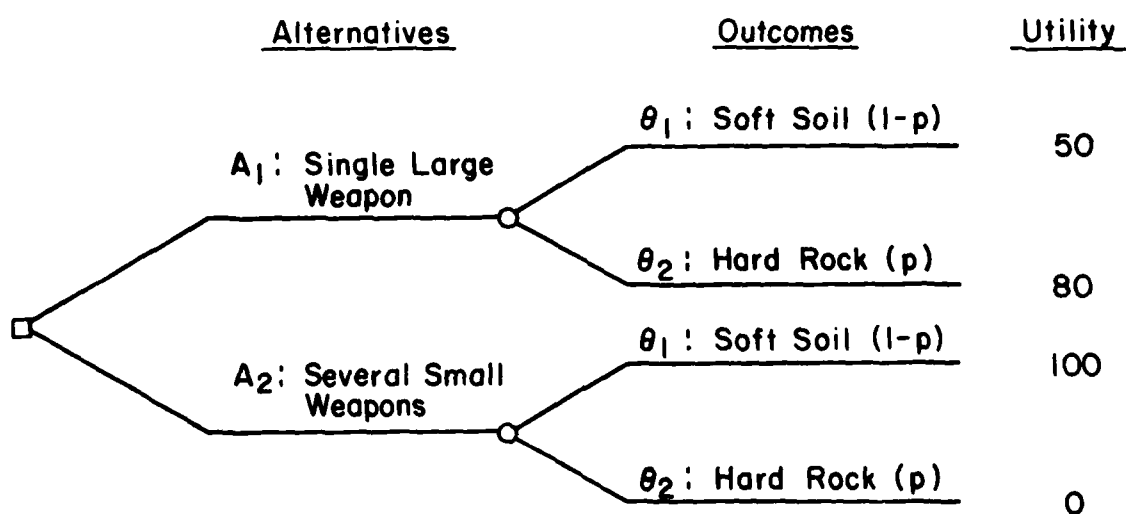


Figure 3.1 Decision tree for example of sensitivity analysis

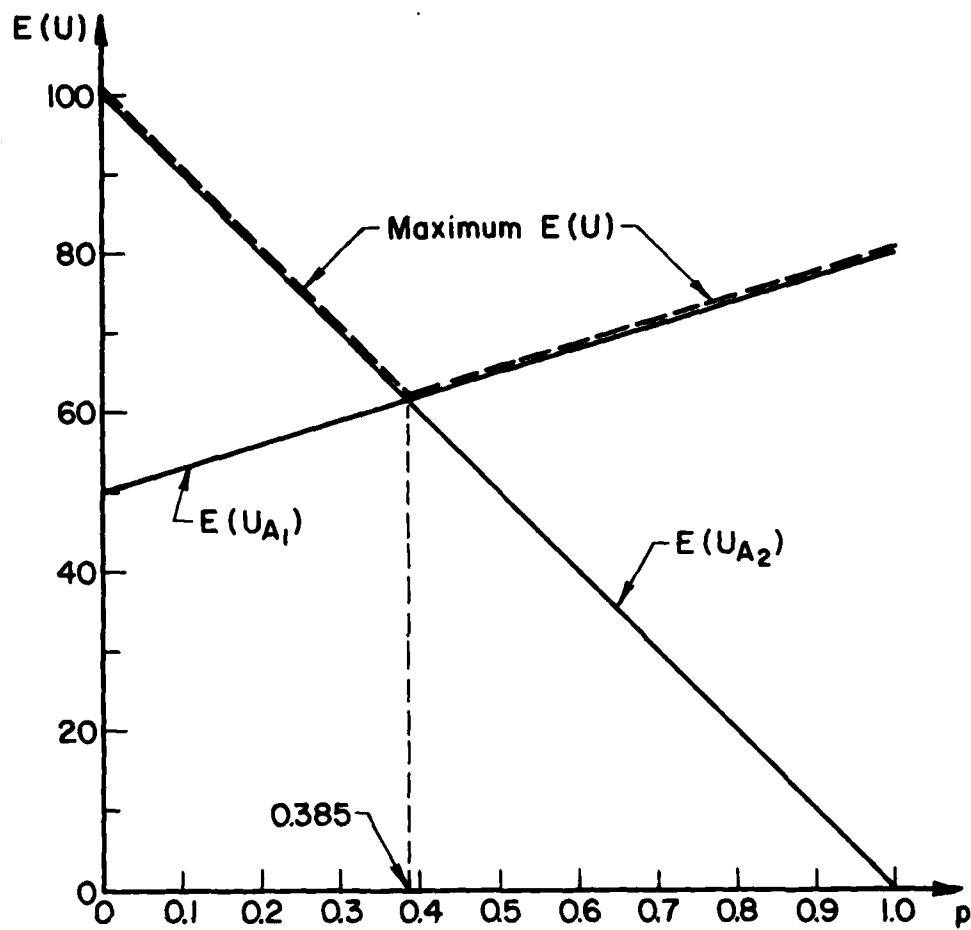


Figure 3.2 Expected utility as function of probability p

IV. DECISION CONCEPTS IN SAMPLING AND ESTIMATION

Statistical sampling may be viewed as a decision problem -- first of all, there is the binary decision of whether or not sampling should be undertaken at all, and if so how extensive should the sampling program be (i.e. what should the sample size be?). In this regard, certain concepts of Bayesian statistics are pertinent.

4.1 Bayesian Sampling

The basic assumption in the Bayesian approach to statistical sampling and estimation is that the parameters of a probability distribution of a random variable are themselves random variables. The uncertainty associated with a given parameter may also be described with a probability density function.

A prior probability density function $f'(\theta)$ for a parameter θ may be prescribed on the basis of judgment, or established on the basis of prior data and information. Observational data may be used to revise or update the prior distribution for the parameter. Based on the Bayes theorem, the updated or (posterior) probability density function for the parameter θ is given by the following (see Ang and Tang, 1975):

$$f''(\theta) = k L(\theta) f'(\theta) \quad (4.1)$$

where:

θ = the parameter under consideration;

$f'(\theta)$ = the prior probability density function of θ , i.e.
prior to the observational data;

$L(\theta)$ = the likelihood function of θ , given as the product of
the density function of the original variable X evaluated
at the observed data points $x_1, x_2, \dots, x_n$; i.e.

$$L(\theta) = \prod_{i=1}^n f_X(x_i | \theta);$$

$f''(\theta)$ = the posterior distribution of θ , i.e. updated in light
of the additional observational data; and

k = a normalization constant to insure that $f''(\theta)$ is a
proper probability density function.

As may be observed in Eq. 4.1, the Bayesian approach allows any prior information on the parameter θ to be incorporated systematically in the final determination of θ .

The prior information could be based on previous sampling results, or indirect measurements, or simply based on subjective judgments. In the event that there is no objective basis for establishing a prior distribution, the diffuse prior (i.e. the uniform distribution between 0 and 1) may be used. In such a case, $f'(\theta)$ is simply a constant, i.e. not a function of θ , and the posterior distribution $f''(\theta)$ becomes

$$f''(\theta) = k L(\theta) = k \prod_{i=1}^n f_X(x_i | \theta) \quad (4.2)$$

which consists of sampling information only, and may be referred to as the "data-based distribution".

Sampling from Gaussian Population -- As an illustration, consider the sampling from a normal population X . Assume that the variance σ^2 is known and the mean-value of the population μ is the only parameter to be estimated from a sample of size n . In this case, the likelihood function can be shown to be (see Ang and Tang, 1975)

$$L(\mu) = N_{\mu}(\bar{x}, \frac{\sigma}{\sqrt{n}})$$

where $N_{\mu}(\bar{x}, \sigma/\sqrt{n})$ denotes the normal probability density function for μ with mean value $\bar{x}$ and standard deviation $\sigma/\sqrt{n}$. $\bar{x}$ is the sample mean of the n data values. In other words,

$$N_{\mu}(\bar{x}, \sigma/\sqrt{n}) = \frac{1}{2\pi\sigma/\sqrt{n}} \exp \left[-\frac{1}{2} \left(\frac{\mu - \bar{x}}{\sigma/\sqrt{n}} \right)^2 \right] \quad (4.3)$$

Hence, the data-based posterior distribution of μ is simply a normal distribution $N_{\mu}(\bar{x}, \sigma/\sqrt{n})$

It can be shown (Ang and Tang, 1975) that the Bayesian distribution of the basic random variable X , incorporating the effect of uncertainty in the parameter μ remains normal with mean $\bar{x}$ and standard deviation $\sqrt{\sigma^2 + \sigma^2/n}$. The sampling uncertainty, denoted by σ^2/n , is added to the basic or inherent variability σ^2 to yield the overall uncertainty.

On the other hand, if a prior distribution of μ is available, e.g. if $f'(\mu)$ is normal $N(\mu', \sigma')$, where μ' and σ' are respectively the prior mean and standard deviation of μ , then the posterior distribution of μ can be shown to be also normal $N_{\mu}(\mu'', \sigma'')$ in which (see Ang and Tang, 1975),

$$\mu'' = \frac{\bar{x}(\sigma')^2 + \mu'(\sigma^2/n)}{(\sigma')^2 + \sigma^2/n} \quad (4.4)$$

and,

$$\sigma'' = \frac{\sigma'(\sigma/\sqrt{n})}{\sqrt{(\sigma')^2 + \sigma^2/n}} \quad (4.5)$$

The results in Eqs. 4.4 and 4.5 can be generalized such that if there are two independent sources of information on the parameter μ , for example as represented by $N_{\mu}(\mu_1, \sigma_1)$ and $N_{\mu}(\mu_2, \sigma_2)$, the two sources can be combined to yield an overall mean-value of the parameter μ as follows:

$$\mu'' = \frac{\mu_1 \sigma_2^2 + \mu_2 \sigma_1^2}{\sigma_1^2 + \sigma_2^2} \quad (4.6)$$

and the corresponding standard deviation of the parameter μ becomes,

$$\sigma'' = \frac{\sigma_1 \sigma_2}{\sqrt{\sigma_1^2 + \sigma_2^2}} \quad (4.7)$$

The Bayesian approach described and illustrated above may be applied also to sampling from populations that are not normal; for such cases, see Ang and Tang (1975).

The relation between the likelihood function and the prior and posterior distributions of the parameter θ is illustrated in Fig. 4.1. Observe that the posterior distribution is sharper than either the prior distribution or the likelihood function. This implies that more information is contained in the posterior distribution than in either the prior or the likelihood function; this is, of course, to be expected as the posterior distribution integrates the information from the prior with that in the likelihood function.

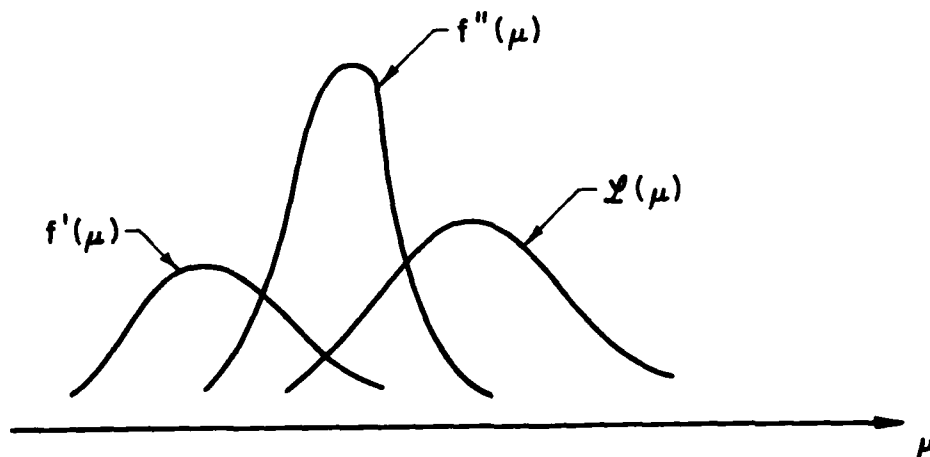


Figure 4.1 Prior, likelihood and posterior functions for parameter μ

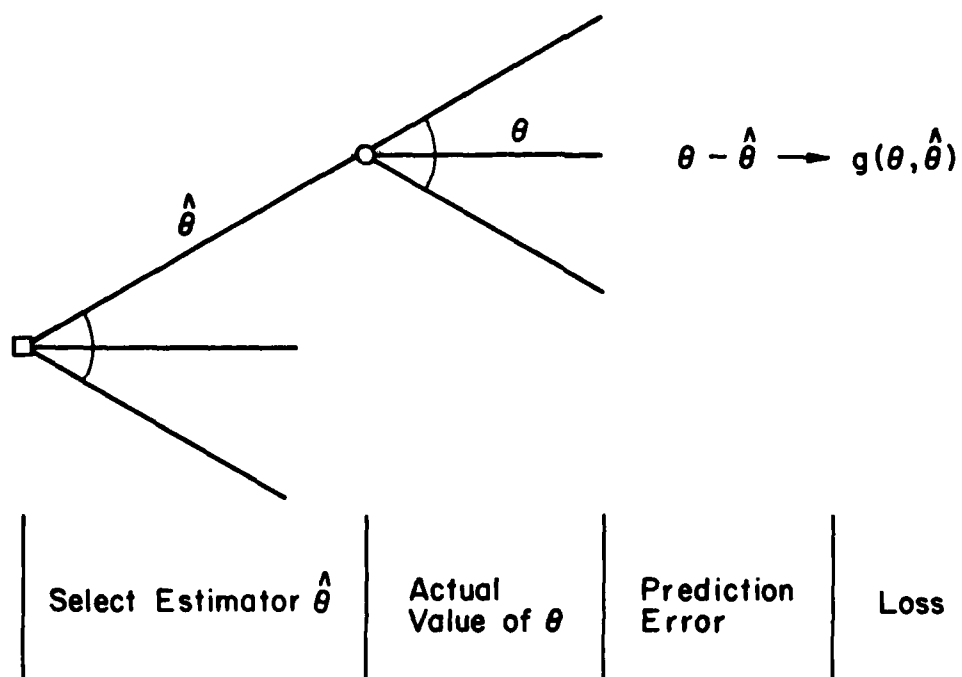


Figure 4.2 Selection of point estimator

4.2 Bayes' Point Estimator

Based on the results of sampling, and any prior information that may be available, the previous section shows that an improved posterior probability distribution for a parameter can be obtained or established. For many practical purposes, however, the point estimator (instead of the distribution) of the parameter is required or is more useful. Such a point estimate may be obtained from a decision analysis.

Consider the decision tree in Fig. 4.2 which depicts a situation in which the point estimator of the parameter θ is to be selected. Because of uncertainty, an estimate of the real value of the parameter θ may contain error; therefore, for a particular choice of the estimator $\hat{\theta}$, a prediction error will result with some associated loss. Modeling this estimation process as a decision process, the Bayes point estimator may be determined such that the expected loss associated with the error in prediction is minimized. Mathematically, if $\hat{\theta}$ is the estimator of a parameter θ whose actual value is described by a distribution $f(\theta)$, the expected loss due to error in prediction is

$$L = \int g(\theta, \hat{\theta}) f(\theta) d\theta \quad (4.8)$$

where $g(\theta, \hat{\theta})$ is the "loss function" for given values of θ and $\hat{\theta}$. The Bayes' estimator is based on the point estimate that minimizes the expected loss; thus, it may be obtained from the following relationship:

$$\frac{\partial L}{\partial \hat{\theta}} = \int \frac{\partial g(\theta, \hat{\theta})}{\partial \hat{\theta}} f(\theta) d\theta = 0 \quad (4.9)$$

Therefore, depending on the form of the loss function, $g(\theta, \hat{\theta})$, the optimal choice for the estimator may be different. For example, if the loss function is quadratic, that is

$$g(\theta, \hat{\theta}) \sim (\theta - \hat{\theta})^2$$

the Bayes estimator can be shown to be the mean-value of θ ; whereas if the loss function is linear, that is

$$g(\theta, \hat{\theta}) \sim (\theta - \hat{\theta})$$

then the Bayes estimator is the median value of θ .

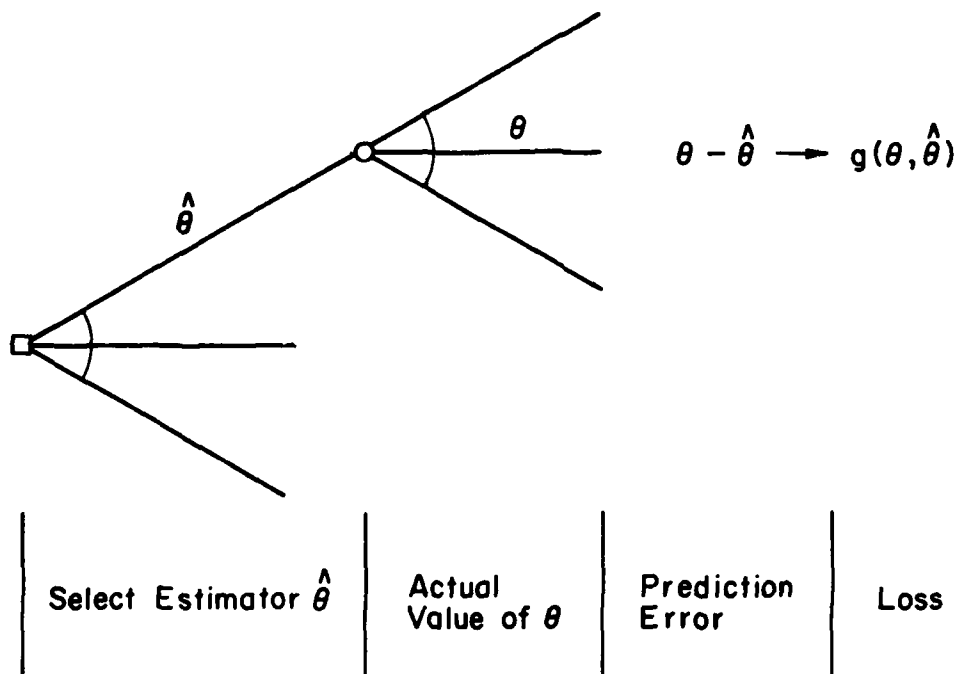


Figure 4.2 Selection of point estimator

4.3 Optimal Sample Size

In addition to the selection of an estimator on the basis of available sample data, other questions that may be addressed pertaining to statistical sampling would include the following.

1. Should a sampling program be implemented?
2. If sampling is necessary, how much sampling should be done; i.e. what should be the sample size?

The problem of determining the optimal sample size may also be modeled as a preposterior analysis. Suppose X is a random variable with parameter θ that needs to be estimated with the sample data. A decision tree that includes the determination of the optimal sample size is shown in Fig. 4.3. Three phases of the decision process may be identified. The first phase of the process pertains to whether or not sampling should be conducted (at node A). If sampling is the preferred action, then a determination of the sample size (at node B) would be needed. After the sample data have been collected, an estimator $\hat{\theta}$ (at node C) is selected. The total loss for each of the paths involving sampling will depend on the sample size n , and the estimated value of the parameter relative to its true value.

The pertinent decision analysis would start at the last node of the decision tree. The expected loss for a given n , $\{x\}$, and estimate $\hat{\theta}$, is given by weighing the loss over the posterior distribution of θ as follows:

$$E(L|n, \{x\}, \hat{\theta}) = \int L(n, \{x\}, \hat{\theta}, \theta) f''(\theta) d\theta \quad (4.10)$$

At node C, the optimal estimate $\hat{\theta}$ is the value which minimizes the expected loss of Eq. 4.10; that is, from

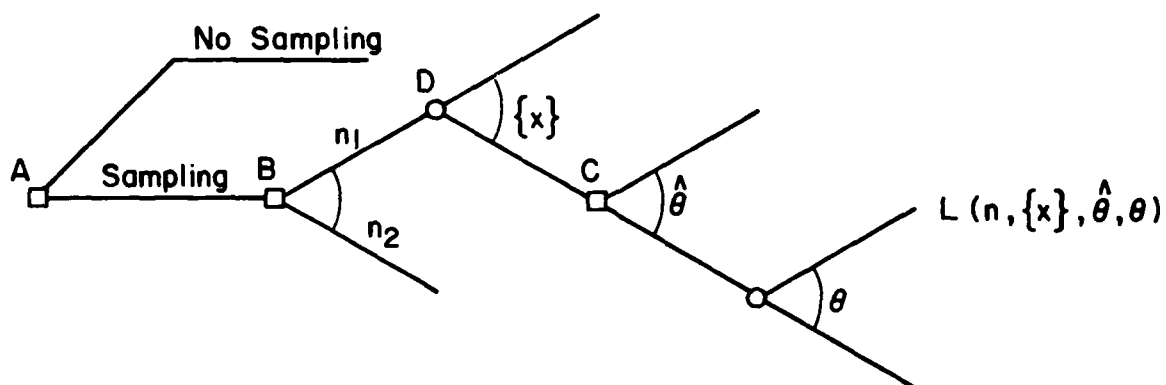
$$\frac{\partial E(L|n, \{x\}, \hat{\theta})}{\partial \hat{\theta}} = 0 \quad (4.11)$$

Eq. 4.11 should yield the optimal estimator, $\hat{\theta}_{opt}$, on the basis of which the expected loss for a given sample size n (at node D) becomes

$$E(L|n) = \int E(L|n, \{x\}, \hat{\theta}_{opt}) f\{x\} d\{x\} \quad (4.12)$$

The required optimal sample size, n_{opt} , may then be evaluated from

$$\frac{dE(L|n)}{dn} = 0 \quad (4.13)$$



Select Sample Size	Observed Sample Data	Select Estimator $\hat{\theta}$	Actual Value of θ	Loss Function
--------------------------	----------------------------	---------------------------------------	--------------------------------	------------------

Figure 4.3 Decision tree for optimal sample size

Example

Suppose that an exploration program is planned for determining the shear velocity of rocks at different sites for the construction of a military system. The mean shear velocity of the rock stratum is of principal interest for engineering purposes; from this information, the elastic modulus and constraint modulus of the rock may be estimated. The shear velocity of the rock may be assumed to follow a normal distribution.

Assume that the loss due to any error in the estimated mean velocity is proportional to the square of the error, and the cost per unit of squared error is related to the cost of performing field measurements. For simplicity, assume also that the c.o.v. of the shear velocity of the rock (representing the inherent variability of the rock stratum plus any error in measurement) is about 15%.

The determination of the optimal number of measurements is of interest. Two cases may be considered; namely,

- (1) no prior information on the mean shear velocity is available;
 - (2) prior information on the mean velocity of similar rocks is available;
- hypothetically suppose that this is $N(4,000 \text{ fps}, 400 \text{ fps})$.

The loss function may be expressed as

$$L(n, \{x\}, \hat{\mu}, \mu) = c(\mu - \hat{\mu})^2 + kn$$

where, μ denotes the mean shear velocity of the rock, which is the parameter to be estimated; $\hat{\mu}$ is the estimator for μ ; c is the cost per unit of squared error in the estimator; and k is the cost of one measurement. This loss function assumes that the loss due to error in $\hat{\mu}$ is quadratic, and the cost of sampling is linear with sample size n . Applying Eq. 4.10,

$$\begin{aligned} E(L|n, \{x\}, \hat{\mu}) &= \int_0^{\infty} [c(\mu - \hat{\mu})^2 + kn] f''(\mu) d\mu \\ &= cE''_{\mu}(\mu - \hat{\mu})^2 + kn \\ &= cE''_{\mu}[(\mu - \mu'') + (\mu'' - \hat{\mu})]^2 + kn \\ &= c\text{Var}''(\mu) + c(\mu'' - \hat{\mu})^2 + kn \end{aligned}$$

where μ'' and $\text{Var}''(\mu)$ are the posterior mean and variance of μ , respectively.

Applying Eq. 4.11,

$$\frac{dE(L|n, \{x\}, \hat{\mu})}{d\hat{\mu}} = 2c(\mu'' - \hat{\mu}) = 0$$

from which the optimal estimator is,

$$(\hat{\mu})_{\text{opt}} = \mu''$$

which means that the optimal estimator is the mean-value of the posterior distribution of μ .

The corresponding expected loss then is

$$\begin{aligned} E[L|n, \{x\}, \hat{\mu}_{\text{opt}}] &= c \text{Var}''(\mu) + (\mu'' - \hat{\mu}_{\text{opt}})^2 + kn \\ &= c \text{Var}''(\mu) + kn \end{aligned}$$

The shear velocity of the rock is Gaussian with known standard deviation σ ; whereas, the posterior variance of μ is given by Eq. 4.5. Hence, the maximum expected loss is,

$$E(L|n, \{x\}, \hat{\mu}_{\text{opt}}) = c \frac{(\sigma')^2 (\sigma^2/n)}{(\sigma')^2 + \sigma^2/n} + kn$$

where μ' and σ' are the prior mean and standard deviation of the mean shear velocity μ .

From Eq. 4.12, the expected loss for a given number of sample measurements n , is thus

$$\begin{aligned} E(L|n) &= \int \left[c \frac{(\sigma')^2 (\sigma^2/n)}{(\sigma')^2 + \sigma^2/n} + kn \right] f(\{x\}) d\{x\} \\ &= c \frac{(\sigma')^2 (\sigma^2/n)}{(\sigma')^2 + (\sigma^2/n)} + kn \end{aligned}$$

For the case in which no prior information on μ is available, the estimate must be based purely on the current field measurements. In such a case, the variance of μ is simply

$$\text{Var}''(\mu) = \sigma^2/n$$

Hence,

$$E(L|n) = c \sigma^2/n + kn$$

The optimal number of measurements is given by the solution to the following equation.

$$\frac{dE(L|n)}{dn} = -c \left(\frac{\sigma^2}{n} \right) + k = 0$$

from which the optimal sample size is,

$$n_{opt} = \sigma \sqrt{c/k}$$

In the present example, the standard deviation of the shear velocity is known to be

$$\sigma = 0.15 \times 4,000 = 600 \text{ fps}$$

Also, assume that the cost ratio $c/k = 3.25 \times 10^{-4}$. Then the optimal sample size is

$$\begin{aligned} n_{opt} &= 600 \sqrt{3.25 \times 10^{-4}} \\ &= 10.82 \text{ or } 11 \end{aligned}$$

Hence, eleven measurements is the optimal sample size for the site exploration program.

In the case where prior information on the mean shear velocity is available, the optimal sample size would be as follows:

$$\begin{aligned} n_{opt} &= \sigma \sqrt{c/k} - \frac{\sigma^2}{(\sigma')^2} \\ &= 10.82 - \frac{600^2}{400^2} \\ &= 8.57 \text{ or } 9 \end{aligned}$$

Hence, in this case the optimal number of measurements is 9, indicating that the prior information is worth approximately two measurements.

V. EVALUATION OF FIELD AND LABORATORY DATA

5.1 Introductory Remarks

The concepts and tools presented in Chapters II through IV may be used in a number of situations of strategic significance. A few of these are illustrated in the following. The potential areas of application, of course, are not limited to those described herein; indeed, the possible areas of applications could be quite broad. In the final analysis, the applications areas are limited only by the imagination and creativity of the user in applying the basic concepts described herein. As alluded to earlier, the application of the general concepts is often not straight-forward, as most significant applications require modeling and conceptual formulations of the underlying physical problem. In this light, the examples of applications described in the sequel are merely to demonstrate the general significance of statistical decision concepts, and to illustrate the applications to specific situations involving strategic planning and decision making.

5.2 Field versus Laboratory Tests

One of the prime objectives of a test program, either field or laboratory, is often to validate or revise/update some prior theoretical or empirical relationship for predicting conditions in the real world, or for extrapolating prior observations beyond the range of available data. In statistical terms, a prediction is usually given in terms of the mean value and associated standard deviation (or coefficient of variation). Invariably, such predictions are imperfect and thus will contain inaccuracy and uncertainties. For example, the predicted or estimated mean-value $\bar{x}$ of a variable X may contain bias, which is a systematic error, such that the correct mean-value would be

$$\mu_X = \bar{v} \bar{x} \quad (5.1)$$

where $\bar{v}$ is the mean bias factor necessary (even if only to be determined subjectively) to obtain the correct mean-value μ_X .

In addition, there may also be statistical error in the estimated mean-value $\bar{x}$. That is, suppose that $\bar{x}$ is the estimated sample mean from a sample of size n . Then, conceptually, if additional samples of the same size were to be obtained, there could conceivably be some scatter in $\bar{x}$; this scatter is given by σ^2/n , which decreases with the sample size n and represents only the random error due to sampling. However, there may be other factors that could contribute additional random errors to $\bar{x}$, which may be represented by the coefficient of variation Δ_v . The c.o.v. Δ_v may be much larger than the sampling error if predictions are based entirely on laboratory and theoretical results; whereas it may be negligible if actual field data were used.

In any case, test results are particularly useful for evaluating or improving the mean bias factor $\bar{v}$ of a prediction, as well as for evaluating the c.o.v. Δ_v . Either field data or laboratory test data, or combinations thereof, are useful for these purposes.

Field Tests -- If field test results are available, they can be used to update any prior information as follows:

Let B = the bias in the prediction of the real world. Then, assuming that the result of a field test is a realization of the real world, the bias is

$$B = X_F / X_P$$

where:

X_F = field observation, which is a realization of the real world, and;

X_P = a prediction of the real world

If the sample size of the field data is n_F , the mean of the individual ratio (x_{Fi}/x_{Pi}) should yield the mean bias

$$\bar{B} = \frac{1}{n_F} \sum_i (x_{Fi}/x_{Pi}) \quad (5.2)$$

Also, from the same set of field data, the error in $\bar{B}$ is.

$$\Omega_{\bar{B}} = \frac{\Omega_B}{\sqrt{n_F}} \quad (5.3)$$

where $\Omega_B = \sigma_B / \bar{B}$, and σ_B is the sample standard deviation of the ratios x_{Fi}/x_{Pi} .

In other words, B is the inaccuracy of the prediction. From n_F sample field data, the mean bias may be estimated as in Eq. 5.2. Also, from the same set of data, the c.o.v. of the mean bias $\bar{B}$ can be estimated as in Eq. 5.3.

Laboratory Tests -- In the case of laboratory test results, the total inaccuracy in the prediction may be divided into two parts. First, the prediction may be bias relative to the laboratory tests; in addition, there may be systematic difference between laboratory and field test results. Therefore, the total inaccuracy of a prediction may be represented as,

$$B = \left(\frac{X_F}{X_L} \right) \left(\frac{X_L}{X_P} \right) = C \times A \quad (5.4)$$

The ratio $C = X_F/X_L$ may be called the "calibration factor" representing the error or bias of laboratory tests, whereas $A = X_L/X_P$ is the ratio of laboratory data to the prediction.

From a set of n_L laboratory test data, the mean bias between a prediction and laboratory tests may be estimated as

$$\bar{A} = \frac{1}{n_L} \sum_i (x_{Li}/x_{Pi}) \quad (5.5)$$

The mean bias of laboratory test results, i.e. $\bar{C}$, may have to be assessed judgmentally, unless there are prior field and laboratory data to permit its estimation. The total mean bias, therefore, is

$$\bar{B} = \bar{C} \times \bar{A} \quad (5.6)$$

From the same set of laboratory data, the c.o.v. of $\bar{A}$ can be estimated as

$$\Omega_{\bar{A}} = \frac{1}{\bar{A}} \sqrt{\sigma_A^2/n_L} \quad (5.7)$$

where σ_A^2 = the sample variance of the ratios (x_{Li}/x_{Pi}) . Thus, the total c.o.v. of $\bar{B}$ becomes

$$\Omega_{\bar{B}} = \sqrt{\Omega_{\bar{C}}^2 + \Omega_{\bar{A}}^2} \quad (5.8)$$

where $\Omega_{\bar{C}}$ is the c.o.v. of $\bar{C}$.

5.3 Posterior (Updated) Information

Assuming that the prior estimates of the prediction error are respectively $\bar{B}'$ and $\Omega_{\bar{B}}'$, the test results (either field or laboratory) may be used to update or improve these prior estimates, obtaining the posterior estimates $\bar{B}''$ and $\Omega_{\bar{B}}''$ as follows:

$$\bar{B}'' = \frac{\bar{B}' (\Omega_{\bar{B}}' \bar{B})^2 + \bar{B} (\Omega_{\bar{B}}' \bar{B}')^2}{(\Omega_{\bar{B}}' \bar{B})^2 + (\Omega_{\bar{B}}' \bar{B}')^2} \quad (5.9)$$

and,

$$\Omega_{\bar{B}}'' = \frac{1}{\bar{B}''} \cdot \frac{(\Omega_{\bar{B}}' \bar{B}) (\Omega_{\bar{B}}' \bar{B}')}{\sqrt{(\Omega_{\bar{B}}' \bar{B})^2 + (\Omega_{\bar{B}}' \bar{B}')^2}} \quad (5.10)$$

where $\bar{B}$ and $\Omega_{\bar{B}}$ are as given by Eqs. 5.2 and 5.3, or Eqs. 5.6 and 5.8.

In the case where the test results consist of both field and laboratory tests, the two types of tests may be combined to give

$$\bar{B} = \frac{\bar{B}_F (\Omega_{\bar{B}_L} \bar{B}_L)^2 + \bar{B}_L (\Omega_{\bar{B}_F} \bar{B}_F)^2}{(\Omega_{\bar{B}_L} \bar{B}_L)^2 + (\Omega_{\bar{B}_F} \bar{B}_F)^2} \quad (5.11)$$

$$\Omega_{\bar{B}} = \frac{1}{\bar{B}} \cdot \frac{(\Omega_{\bar{B}_F} \bar{B}_F) (\Omega_{\bar{B}_L} \bar{B}_L)}{\sqrt{(\Omega_{\bar{B}_F} \bar{B}_F)^2 + (\Omega_{\bar{B}_L} \bar{B}_L)^2}} \quad (5.12)$$

where $\bar{B}_F$ and $\Omega_{\bar{B}_F}$ are given by Eqs. 5.2 and 5.3; whereas $\bar{B}_L$ and $\Omega_{\bar{B}_L}$ are given by Eqs. 5.6 and 5.8. These may then be used to update any prior estimates B' and $\Omega'_{\bar{B}}$; obtaining the corresponding posterior estimates as shown in Eqs. 5.9 and 5.10.

5.4 Measure of Information

In order to define or establish a utility measure for a test plan, some measure of information gained from the experiment is obviously important, in addition to the cost of the experiment. For this purpose, observe the following.

Perhaps of first order importance from a test plan is the evaluation of the mean bias factor of a prediction; in this regard, the utility function may be defined in a quadratic form as follows (or other forms may be prescribed):

$$U|\mu_{\bar{B}}, \bar{B} = -k(\bar{B} - \mu_{\bar{B}})^2 \quad (5.13)$$

where,

$\mu_{\bar{B}}$ = the true mean bias;

$\bar{B}$ = estimated mean bias based on the experimental data.

Assuming that $\mu_{\bar{B}}$ is a random variable, the expected value of Eq. 5.13 with respect to $\mu_{\bar{B}}$ is,

$$E(U|\bar{B}) = -k E_{\mu_{\bar{B}}} (\bar{B} - \mu_{\bar{B}})^2 = k \text{Var}(\bar{B}) = -k(\Omega_{\bar{B}} \bar{B})^2 \quad (5.14)$$

Finally, taking the expected value over all values of $\bar{B}$,

$$E(U) = k E_{\bar{B}} (\Omega_{\bar{B}} \bar{B})^2 = -k \Omega_{\bar{B}}^2 [\text{Var}(\bar{B}) + E^2(\bar{B})] \quad (5.15)$$

Observe that $\text{Var}(\bar{B})$ is a second order term, i.e. $\ll 1.0$; whereas $E^2(\bar{B}) \approx 1$. Hence

$$E(U) = -k \Omega_{\bar{B}}^2. \quad (5.16)$$

Similarly, the utility function may be defined in the form,

$$U|\mu_{\bar{B}}, \bar{B} = -k_1(\bar{B} - \mu_{\bar{B}})^2 + k_2(\mu_{\bar{B}} - 1) \quad (5.17)$$

or

$$U|\mu_{\bar{B}}, \bar{B} = -k_1(\bar{B} - \mu_{\bar{B}})^2 + k_2(\bar{B} - 1) \quad (5.18)$$

In either of these latter cases, it can also be shown that the expected utility is

$$E(U) = -k_1 \Omega_{\bar{B}}^2 \quad (5.19)$$

In short, the above results show that if a quadratic utility function is appropriate, the measure of information content can be defined solely in terms of the c.o.v., $\Omega_{\bar{B}}$

Therefore, the information gained from a test program would be inversely proportional to $\Omega_{\bar{B}}^2$. The cost of a test program obviously is directly proportional to the size of the program, i.e. sample size n . Consequently, the utility of a test program may be defined as a function of the cost per unit of information gained, or $c\Omega_{\bar{B}}^2$. Therefore, the "best" experiment may be the one that minimizes the cost per unit information gain.

Example

For illustration, consider the determination of the bearing capacity of a large site for the design of the foundation of a military facility. For this purpose, assume that the test options include:

- (1) load bearing capacity tests in the field;
- (2) unconfined compression tests of laboratory soil samples obtained from boring, and subsequently calculating the bearing capacity based on the soil parameters obtained from such tests.

Obviously, load-bearing capacity tests in the field are expensive and, therefore, only a limited number of such tests may be performed (if at all). However, more extensive soil samples may be obtained over the entire site.

Suppose that five load tests were performed; the locations of the tests were spaced sufficiently far apart to cover the entire area of the site. At each of the load test locations, assume that one soil sample is also taken which may be considered to be the "control samples". The measured bearing capacity from the load tests, and the corresponding bearing capacity estimated from the control soil samples, may be as follows:

Load Tests	Bearing Capacity x_F	Control Soil Sample	Bearing Capacity x_L	Ratio x_F/x_L
1	3,500 psf	1	4,000 psf	0.875
2	3,000	2	3,000	1.00
3	2,500	3	3,200	0.781
4	4,000	4	3,500	1.143
5	4,300	5	4,500	0.956

For the remainder of the site, only soil samples are taken; say fifteen more were obtained in addition to the control samples, giving results as follows:

Soil Sample	Bearing Capacity, x_L
6 - - - - -	4,300 psf
7 - - - - -	3,600
8 - - - - -	3,800
9 - - - - -	4,700
10 - - - - -	4,200
11 - - - - -	3,900
12 - - - - -	3,300
13 - - - - -	2,900
14 - - - - -	4,800
15 - - - - -	4,000
16 - - - - -	3,700
17 - - - - -	3,600
18 - - - - -	4,100
19 - - - - -	3,900
20 - - - - -	4,500

Assume that the soil type over the entire area of the site is fairly uniform, to permit the use of a uniform bearing capacity for the entire site (if the soil type is not uniform, the area may be divided into subareas). The results of the field and laboratory tests, describe hypothetically above, then can be used as follows.

From the field test results, the mean bearing capacity is

$$\bar{x}_F = 3460 \text{ psf};$$

and the corresponding c.o.v. is

$$\Omega_{X_F} = 0.21$$

Also, the uncertainty (c.o.v.) in the estimated mean bearing capacity is

$$\Omega_{\bar{X}_F} = \frac{0.21}{\sqrt{5}} = 0.09$$

On the basis of the laboratory soil samples, the mean bearing capacity is calculated to be

$$\bar{x}_L = 3875 \text{ psf};$$

and the corresponding c.o.v. is,

$$\Omega_{X_L} = 0.14$$

from which the c.o.v. of $\bar{x}_L$ is,

$$\Omega_{\bar{X}_L} = \frac{0.14}{\sqrt{20}} = 0.03.$$

Comparing the bearing capacity estimated using the control soil samples with the corresponding field test results, it is obvious that x_L tends to overestimate the actual value of the bearing capacity at the site. In other words, there is a systematic bias in the laboratory-based bearing capacity estimates; i.e. the correct bearing capacity may be expressed as

$$X = C X_L$$

where C is the bias in the laboratory-based bearing capacity estimate for the site.

From the ratios of X_F/X_L given in the above table, the mean value of C is

$$\bar{C} = 0.951$$

and the associated c.o.v. is

$$s_C = \frac{0.136}{0.951} = 0.143$$

whereas the c.o.v. of $\bar{C}$ is,

$$\Delta_C = \frac{0.143}{\sqrt{5}} = 0.064.$$

Thus, the total uncertainty underlying the laboratory-based mean bearing capacity is

$$\Omega_{\bar{X}} = \sqrt{\Delta_C^2 + \Omega_{\bar{X}_L}^2} = \sqrt{0.064^2 + 0.03^2} = \underline{0.071}$$

The load test data as well as the unconfined compression laboratory test data described above are useful for estimating the bearing capacity for the site in question. That is, both sets of data can be combined in accordance with Eqs. 5.11 and 5.12, to give a composite estimate of the bearing capacity for the site as follows:

$$\begin{aligned} \bar{X} &= \frac{3460(0.071 \times 0.951 \times 3875)^2 + 0.951 \times 3875(0.09 \times 3460)^2}{(0.071 \times 0.951 \times 3875)^2 + (0.09 \times 3460)^2} \\ &= 3592 \text{ psf.} \end{aligned}$$

$$\begin{aligned} \Omega_{\bar{X}} &= \frac{1}{3592} \cdot \frac{(0.09 \times 3460)(0.071 \times 0.951 \times 3875)}{\sqrt{(0.09 \times 3460)^2 + (0.071 \times 0.951 \times 3875)^2}} \\ &= 0.06. \end{aligned}$$

VI. ANALYSIS AND PLANNING OF TEST PROGRAMS

6.1 Introduction

In test programs that are extremely costly, such as the Mighty Epic and the Diablo Hawk test programs, the need for a systematic framework for planning the experiments is clear. The purpose and objective of any test program, of course, is to obtain information to improve the existing state of information and knowledge.

A test plan, therefore, should be to obtain the maximum benefit, which may be defined in terms of information gained or improvements accruing from the tests.

The basic concepts of statistical decision may be used to advantage in the planning of test programs; however, the application of these concepts require modeling of the specific problem underlying a particular test program. In view of this, the models developed herein will refer to test programs similar to the Mighty Epic or the Diablo Hawk programs.

In particular, referring to test programs such as the Diablo Hawk, the programs involve the testing of specified diameter tunnels subjected to nuclear weapons effects. Decisions, therefore, are required relative to the choice of diameters of the tunnels to be tested, the location of the tunnels relative to the blast point which may be given in terms of the maximum strain at which the tunnels may be subjected, as well as the number of tunnels (or sections of tunnels) of a given diameter.

An approach to this problem is developed and formulated below. Before describing the approach, consider the following.

6.2 Preliminary Consideration

Presumably, the failure strain ϵ_f of a tunnel with given diameter is random as shown in Fig. 6.1. Moreover, it may be assumed that the failure strain depends on the diameter of the tunnel; that is, there is a size effect which may be a regression relation (could be nonlinear between the mean failure strain and the diameter of the tunnel), as illustrated hypothetically in Fig. 6.2.

Suppose that there is some prior information on the failure strain of tunnels with given diameter d , to establish its probability distribution $f_{\epsilon_f}(\epsilon)$, with mean μ_{ϵ_f} and σ_{ϵ_f} . Furthermore, assume that σ_{ϵ_f} is known, whereas μ_{ϵ_f} is a random variable with prior distribution $f'_{\mu_{\epsilon_f}}(x)$, whose mean value is μ' and standard deviation σ' .

Then, if one tunnel of diameter d is tested at a given strain ϵ , the result can be used to revise and update the prior probability distribution $f'_{\mu_{\epsilon_f}}(x)$ as follows: If the tunnel survives the test at strain ϵ , i.e. $\epsilon_f \geq \epsilon$, the posterior distribution for μ_{ϵ_f} is,

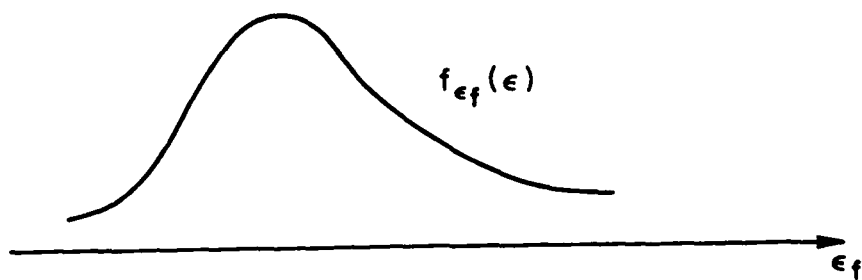


Figure 6.1 Failure strain of tunnel or rock cavity

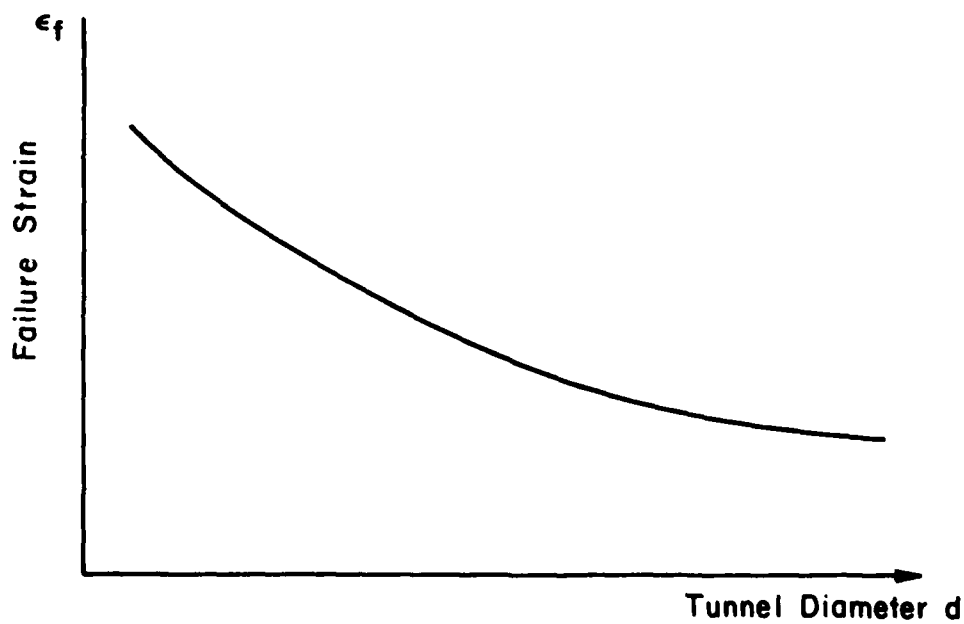


Figure 6.2 Failure strain versus tunnel diameter

$$\begin{aligned}
f''_{\mu_\epsilon}(x|\epsilon_f \geq \epsilon) &= P(\mu_{\epsilon_f} = x | \epsilon_f \geq \epsilon) \\
&= \frac{P(\mu_{\epsilon_f} = x \cap \epsilon_f \geq \epsilon)}{P(\epsilon_f \geq \epsilon)} \\
&= \frac{P(\epsilon_f \geq \epsilon | \mu_{\epsilon_f} = x) P(\mu_{\epsilon_f} = x)}{P(\epsilon_f \geq \epsilon)} \\
&= \frac{[1 - F_{\epsilon_f}(\epsilon | \mu_{\epsilon_f} = x)]}{[1 - F_{\epsilon_f}(\epsilon)]} f'_{\mu_\epsilon}(x) \quad (6.1)
\end{aligned}$$

whereas, if the tunnel fails at the test strain ϵ , i.e. $\epsilon_f < \epsilon$, then the corresponding posterior distribution for μ_{ϵ_f} would be

$$\begin{aligned}
f''_{\mu_\epsilon}(x|\epsilon_f < \epsilon) &= P(\mu_{\epsilon_f} = x | \epsilon_f < \epsilon) \\
&= \frac{F_{\epsilon_f}(\epsilon | \mu_{\epsilon_f} = x)}{F_{\epsilon_f}(\epsilon)} f'_{\mu_\epsilon}(x) \quad (6.2)
\end{aligned}$$

The updated mean and standard deviation of the mean failure strain then become

$$\mu'' = \int_0^\infty x \cdot f''_{\mu_\epsilon}(x) dx \quad (6.3)$$

$$\sigma'' = \left[\int_0^\infty (x - \mu'')^2 f''_{\mu_\epsilon}(x) dx \right]^{1/2} \quad (6.4)$$

6.3 Measure of Benefit (i.e. Information Gain)

Ordinarily, there may be prior information on the mean failure strain, say μ' , and on its uncertainty σ' (even if only based on subjective judgments).

The results of a test or tests can then be used to update the prior information to obtain the posterior mean μ'' and associated uncertainty σ'' . The benefit of an experiment, therefore, may be a function of the difference between μ'' and σ'' relative to the prior values μ' and σ' . Consequently, in defining a measure to represent the benefit of a test, consider the following:

If no tests were conducted, the prior statistics μ' and σ' will be used in the planning and design of a system. Therefore, if there are large errors in μ' or in σ' , the loss will correspondingly be large, whereas if these priors are accurate (i.e.

small error), the resulting loss in the use of this prior information directly in design will be correspondingly small. On this premise, the benefit that may be gained from a test or tests can be defined as follows.

$$U(\mu'', \sigma'') = a \left(\frac{\mu'' - \mu'}{\mu'} \right)^2 + \left(\frac{\sigma'^2 - \sigma''^2}{\sigma'^2} \right) \quad (6.5)$$

where μ'' and σ'' are the updated mean and standard deviation after the test data have been obtained.

The second term in Eq. 6.5, $\left(\frac{\sigma'^2 - \sigma''^2}{\sigma'^2} \right)$, is the reduction in uncertainty accruing from the test results; whereas the first term, $\left(\frac{\mu'' - \mu'}{\mu'} \right)^2$ represents the reduction in the loss that would otherwise incur had there been no tests conducted.

An optimal test program may then be developed to obtain the maximum potential benefit, meaning maximum U .

6.4 Optimal Test Strains

Suppose that only a single tunnel is to be tested at a strain level ϵ . The question then is "at what maximum strain ϵ should the tunnel be tested in order to derive the maximum benefit from the test?"

The decision variable in this case obviously is ϵ . If a tunnel is subjected to a test strain ϵ , two possibilities may occur -- the tunnel may survive the strain level ϵ (event S) or it may fail (event F). The simple decision tree in this case may be represented as shown in Fig. 6.3.

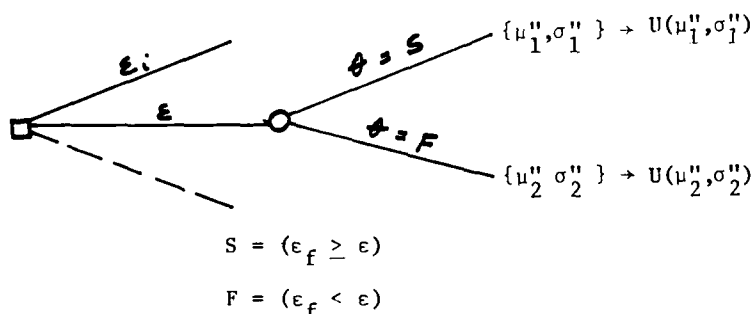


Figure 6.3 Decision tree for test strain

For one tunnel, the results of the test may either be that the tunnel survives (S) or fails (F) when subjected to the load that induces the strain ϵ . Depending on whether the tunnel survives or fails at the strain level ϵ , the corresponding posterior mean and standard deviation may be denoted as (μ_1'', σ_1'') and (μ_2'', σ_2'') . The corresponding utility would be $U(\mu_1'', \sigma_1'')$ and $U(\mu_2'', \sigma_2'')$ as given by Eq. 6.5.

Hence, the expected utility of a test at strain level ϵ is

$$E(U|\epsilon) = U(\mu_1'', \sigma_1'') \cdot P(\epsilon_f \geq \epsilon) + U(\mu_2'', \sigma_2'') P(\epsilon_f < \epsilon) \quad (6.6)$$

in which the posterior quantities μ'' and σ'' may be obtained through Eqs. 6.3 and 6.4.

The optimal test strain, ϵ_{opt} , may then be obtained on the basis of the maximum expected utility criterion, as follows:

$$\frac{dE(U|\epsilon)}{d\epsilon} = 0 \quad (6.7)$$

This means that if a single tunnel of a given diameter were to be tested, it should be placed at such a distance from the blast point at which the induced maximum strain is ϵ_{opt} .

Generalizations -- The basic procedure described above for determining the optimal test strain for one tunnel can be extended to any number of tunnels of the same diameter. Consider first the case of two tunnels.

If two tunnels are to be tested, the strains ϵ_1 and ϵ_2 that the tunnels should be subjected to can be determined also on the basis of maximum expected utility, in the sense of Eq. 6.7.

The decision tree in this case, is as shown in Fig. 6.4.

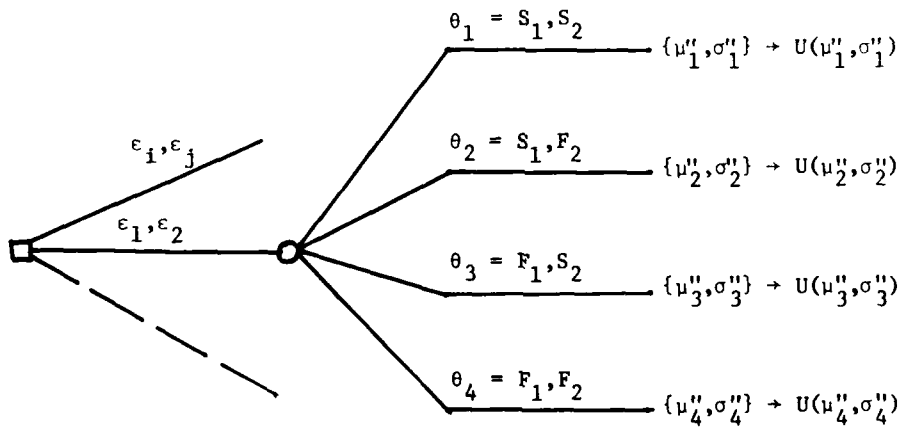


Figure 6.4 Decision tree for test strains of two tunnels.

As shown in Fig. 6.4, the test results for the two tunnels could conceivably be that both tunnels survived or one of them survives whereas the other fails, or both tunnels fail, at strains ϵ_1 and ϵ_2 , respectively; i.e. there are four possible outcomes, θ_1 , θ_2 , θ_3 , θ_4 . Corresponding to each of these four possible outcomes of the set of two tunnels, the resulting posterior means and standard deviations are μ_i'' and σ_i'' , $i = 1, 2, 3, 4$.

The expected utility at the test strains ϵ_1 and ϵ_2 is, therefore,

$$\begin{aligned} E(U|\epsilon_1, \epsilon_2) = & U(\mu_1'', \sigma_1'') P(S_1, S_2) + U(\mu_2'', \sigma_2'') P(S_1, F_2) \\ & + U(\mu_3'', \sigma_3'') P(F_1, S_2) + U(\mu_4'', \sigma_4'') P(F_1, F_2) \end{aligned} \quad (6.8)$$

where, assuming that the outcomes between the two tunnels are statistically independent,

$$P(S_1, S_2) = P(\epsilon_f \geq \epsilon_1) P(\epsilon_f \geq \epsilon_2)$$

$$P(S_1, F_2) = P(\epsilon_f \geq \epsilon_1) P(\epsilon_f < \epsilon_2)$$

$$P(F_1, S_2) = P(\epsilon_f < \epsilon_1) P(\epsilon_f \geq \epsilon_2)$$

$$P(F_1, F_2) = P(\epsilon_f < \epsilon_1) P(\epsilon_f < \epsilon_2)$$

in which;

$$P(\epsilon_f < \epsilon_1) = \int_0^{\epsilon_1} f_{\epsilon_f}(x) dx$$

where $f_{\epsilon_f}(x)$ is the probability density function of the failure strain as depicted in Fig. 6.1.

And,

$$P(\epsilon_f \geq \epsilon_1) = 1 - P(\epsilon_f < \epsilon_1).$$

Also,

$$\begin{aligned} \mu_j'' &= \int_0^{\infty} x f_j''(x) dx \\ \sigma_j''^2 &= \int_0^{\infty} (x - \mu_j'')^2 f_j''(x) dx; \quad j = 1, 2, 3, 4 \end{aligned}$$

where:

$$f_1''(x) = P(\mu_{\epsilon_f} = x | \epsilon_f > \epsilon_1, \epsilon_f > \epsilon_2)$$

$$= \frac{P(\mu_{\epsilon_f} = x \cap \epsilon_f > \epsilon_1, \epsilon_f > \epsilon_2)}{P(\epsilon_f > \epsilon_1, \epsilon_f > \epsilon_2)}$$

Assuming that the outcomes of the two test tunnels are statistically independent,

$$f_1''(x) = \frac{P(\epsilon_f \geq \epsilon_1 | \mu_{\epsilon_f} = x) P(\epsilon_f \geq \epsilon_2 | \mu_{\epsilon_f} = x)}{P(\epsilon_f > \epsilon_1) P(\epsilon_f > \epsilon_2)} f'(x)$$

where $f'(x)$ is the prior distribution of the mean failure strain μ_{ϵ_f} .

Similarly,

$$f_2''(x) = \frac{P(\epsilon_f \geq \epsilon_1 | \mu_{\epsilon_f} = x) P(\epsilon_f < \epsilon_2 | \mu_{\epsilon_f} = x)}{P(\epsilon_f \geq \epsilon_1) P(\epsilon_f < \epsilon_2)} f'(x)$$

$$f_3''(x) = \frac{P(\epsilon_f < \epsilon_1 | \mu_{\epsilon_f} = x) P(\epsilon_f \geq \epsilon_2 | \mu_{\epsilon_f} = x)}{P(\epsilon_f < \epsilon_1) P(\epsilon_f \geq \epsilon_2)} f'(x)$$

$$f_4''(x) = \frac{P(\epsilon_f < \epsilon_1 | \mu_{\epsilon_f} = x) P(\epsilon_f < \epsilon_2 | \mu_{\epsilon_f} = x)}{P(\epsilon_f < \epsilon_1) P(\epsilon_f < \epsilon_2)} f'(x)$$

The optimal test strains, $(\epsilon_1, \epsilon_2)_{opt}$, may then be obtained on the basis of the maximum expected utility as follows:

$$\frac{\partial E(U | \epsilon_1, \epsilon_2)}{\partial \epsilon_1} = 0$$

$$\frac{\partial E(U | \epsilon_1, \epsilon_2)}{\partial \epsilon_2} = 0 \quad (6.9)$$

The generalization of the above formulations to n tunnels of the same diameter, tested at strains $\epsilon_1, \epsilon_2, \dots, \epsilon_n$, may be portrayed with the decision tree shown in Fig. 6.5 below.

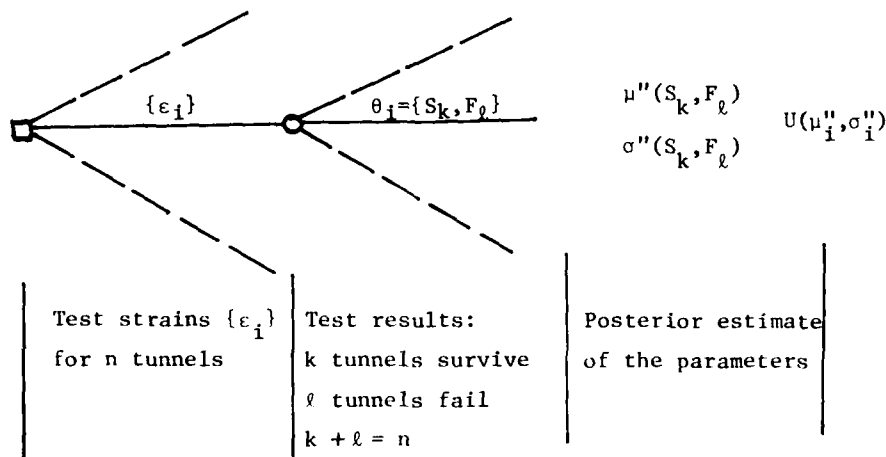


Figure 6.5 Decision tree for test strains of n tunnels

The expected utility then may be represented as,

$$E(U|\{\epsilon_i\}) = \sum_{\text{all } i} P(\theta_i) \cdot U(\mu''_i, \sigma''_i) \quad (6.10)$$

in which

$$P(\theta_i) = \prod_{i \in k} (P(\epsilon_f \geq \epsilon_i)) \cdot \prod_{i \in l} P(\epsilon_f < \epsilon_i)$$

and,

$$\mu''_i = \int_0^{\infty} x f''_{\theta_i}(x) dx$$

$$\sigma''_i = \int_0^{\infty} (x - \mu'')^2 f''_{\theta_i}(x) dx$$

Thus, the utility associated with outcome θ_i is,

$$U(\mu''_i, \sigma''_i) = a \left(\frac{\mu''_i - \mu'_i}{\mu'_i} \right)^2 + \left(\frac{\sigma'^2_1 - \sigma''^2_1}{\sigma'^2_1} \right)$$

from which the expected utility of Eq. 6.10, and the optimal set of test strains $\{\epsilon_i\}_{\text{opt}}$ may be obtained from

$$\frac{\partial E(U|\{\epsilon_i\})}{\partial \epsilon_i} = 0 \quad (6.11)$$

6.5 Optimal Number of Test Tunnels

In order to determine the optimal number of tunnels (of a given diameter) in a test program, the consideration of the cost per tunnel must be included. Assuming that for a given number of tunnels with the same diameter, the optimal test strains have been determined as described above, and introducing the costs of each tunnel, c , the expected utility per unit cost for n tunnels can be expressed as,

$$E(U^*|\{\epsilon_i\}_{opt}, n) = \frac{1}{nc} E(U|\{\epsilon_i\}_{opt}) \quad (6.12)$$

that is, for n tunnels of a given diameter, the optimal test plan is the one with $\{\epsilon_i\}_{opt}$ as determined from Eq. 11. Eq. 6.12 then gives the expected utility per unit cost for this set of tunnels.

Again, on the basis of the maximum expected utility criterion, the optimal number of tunnels of a given diameter, n_{opt} , may then be obtained from,

$$\frac{\partial E(U^*|\{\epsilon_i\}_{opt}, n)}{\partial n} = 0 \quad (6.13)$$

The decision tree would appear as shown in Figure 6.6 below.

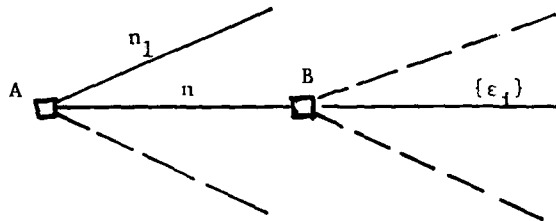


Figure 6.6 Decision tree for determining number of test tunnels.

Referring to Fig. 6.6 the optimal set of test strains is obtained at node B; whereas, at node A, the optimal number of tunnels (of the same diameter) are determined.

The above analyses, therefore, would yield the optimal number of test tunnels of a given diameter, and a corresponding set of optimal strain levels at which the various tunnels should be tested. The same analysis may be performed for tunnels of other diameters. The results may then be combined with those of the other test tunnels (i.e., of other diameters) to obtain the expected utility (or information gain) per unit cost. Repeated evaluations of several test plans, each consisting of different combinations of tunnel diameters, should provide a basis for identifying the optimal

test plan on the basis of maximum expected utility per unit cost.

Example

The concepts formulated above are difficult to illustrate numerically; much of the numerical calculations will involve extensive numerical integrations making computer calculations necessary. Nevertheless, the following formulation and discussion may serve to clarify the concepts developed herein.

Suppose the diameter of full-size prototype tunnels in the field will be twenty feet. Because of expense, the maximum diameter of test tunnels may have to be limited to ten feet; moreover, the results obtained for 10-foot diameter tunnels may be considered to approach those of the full-size 20-foot tunnels. Even for 10-foot tunnels, the expense of testing tunnels of this size may be so high that the number of such tunnels must also be limited; the use of smaller diameter tunnels in a test plan, therefore, may be more cost-effective.

Consider the case in which two sizes of tunnels are to be used in a test program; namely, tunnels with two-foot and ten-foot diameters. Furthermore, assume that only two 10-foot diameter tunnels will be tested, whereas a larger number of 2-foot tunnels will be economically possible. A decision on the strain levels at which the various tunnels should be tested is required in this case.

Consider first the 10-foot tunnels, at test strain levels ϵ_1 and ϵ_2 for the two tunnels, the possible results or outcome from the tests of these two tunnels will be as follows:

$$(S_1, S_2), (S_1, F_2), (F_1, S_2), (F_1, F_2)$$

i.e. one or both of the tunnels may survive at the specified test strains ϵ_1 and ϵ_2 , or none of them may survive.

Suppose also that from prior observational data, supplemented with judgment, the mean failure strain of ten-foot diameter tunnels is estimated to be μ' , and standard deviation σ' . However, the mean failure strain contains uncertainty; hence, it may be assumed to be a random variable with normal distribution; i.e.

$$f'_\mu(x) = N_\mu(\mu', \sigma')$$

After the test results, the posterior statistics of the failure strain may be updated or revised depending on which outcome is realized. For example, if tunnel No. 1 survives the test strain ϵ_1 whereas tunnel No. 2 fails at ϵ_2 , i.e. the outcome is (S_1, F_2) , then the posterior statistics for the mean failure strain of 10-foot tunnels are:

$$\mu'' = \int x f''_\mu(x) dx$$

$$\sigma''^2 = \int (x - \mu'')^2 f''_\mu(x) dx$$

where,

$$f''_{\mu}(x) = P(\mu = x | \epsilon_f \geq \epsilon_1, \epsilon_f < \epsilon_2)$$

Assuming that the outcomes between the two tunnels are statistically independent,

$$f''_{\mu}(x) = \frac{P(\epsilon_f \geq \epsilon_1 | \mu = x) P(\epsilon_f < \epsilon_2 | \mu = x)}{P(S_1) P(F_2)} \cdot f'_{\mu}(x)$$

$$= \frac{[1 - \Phi(\frac{\epsilon_1 - \mu}{\sigma_{\epsilon_f}})] \Phi(\frac{\epsilon_2 - \mu}{\sigma_{\epsilon_f}})}{P(S_1) \cdot P(F_2)} f'_{\mu}(x)$$

in which,

$$P(S_1) = P(\epsilon_f \geq \epsilon_1) = 1 - \Phi\left(\frac{\epsilon_1 - \mu}{\sqrt{\sigma_{\epsilon_f}^2 + \sigma'^2}}\right)$$

$$P(F_2) = P(\epsilon_f < \epsilon_2) = \Phi\left(\frac{\epsilon_2 - \mu}{\sqrt{\sigma_{\epsilon_f}^2 + \sigma'^2}}\right)$$

and $\Phi(-)$ is the cumulative probability of the standard normal distribution. The posterior distribution f'' for the other outcomes can be similarly determined.

The above results, of course, pertain to tunnels with 10-foot diameter. For purposes of discussion, suppose that the above analysis yield posterior estimates for the mean failure strain of 10-foot tunnels to be

$$\mu'' = 2 \times 10^{-3}$$

and

$$\sigma'' = 0.4 \times 10^{-3}$$

On the basis of the results for 10-foot tunnels, the mean failure strain for the full-size 20-foot diameter tunnels may be determined as follows:

$$\epsilon_{f20} = A_{10} \epsilon_{f10}$$

where A is the bias factor, principally to take account of any size effect (the statistics of A may be determined from prior experience). The mean failure strain for 20-foot tunnels then may be expressed as

$$\mu''_{20} = \bar{A} \mu''_{10}$$

Assuming $\bar{A} = 0.95$ and $\delta_{\bar{A}} = 0.4$,

$$\mu''_{20} \approx 0.95 (2 \times 10^{-3}) = 1.90 \times 10^{-3}$$

and the associated c.o.v. for the mean failure strain of 20-foot tunnels is

$$\begin{aligned} \delta_{\mu_{20}} &= \sqrt{\delta_{\bar{A}}^2 + (\sigma''/\mu'')^2} \\ &= \sqrt{0.4^2 + 0.2^2} \\ &= 0.45 \end{aligned}$$

From which,

$$\sigma''_{\mu_{20}} = 0.45 (1.90 \times 10^{-3}) = 0.86 \times 10^{-3}$$

At this point, if the prior information is μ'_{20} and σ'_{20} , which may be results from earlier tests, then the gain in information (or utility) accruable from the test of the two 10-foot diameter tunnels becomes

$$U = a \left(\frac{\mu''_{20} - \mu'_{20}}{\mu'_{20}} \right)^2 + \left(\frac{\sigma'^2_{20} - \sigma''^2_{20}}{\sigma'^2_{20}} \right)$$

Observe that, implicitly, U is a function of ϵ_1 and ϵ_2 . Hence, the optimal test strains may be obtained from maximizing U , i.e.

$$\frac{\partial U}{\partial \epsilon_1} = 0$$

and,

$$\frac{\partial U}{\partial \epsilon_2} = 0$$

The results would be the optimal strains that should be applied to the two 10-foot tunnels. Computer calculations will be necessary to obtain these results.

VII. ANALYSIS OF A RETALIATORY SYSTEM

7.1 Premise of Analysis

Another potential application of statistical decision concepts is in the evaluation of a new retaliatory system, such as the MX system which is currently under consideration and development. (For this discussion, consider a vertical shelter configuration.)

As a retaliatory system, the principal objective is to insure a minimum level of retaliation after the first strike by a potential adversary. In the case of the MX system, this may be the number of surviving missiles following an all out enemy attack.

The number of missiles that may be estimated to survive an enemy attack will depend on the following:

- (1) The number of shelters in the total system.
- (2) The probability that a "loaded" shelter will be targeted or under attack by an enemy warhead.
- (3) The probability that a missile will survive if it is under attack.

Mathematically, the number of surviving missiles may be expressed as

$$NSM = P(S|A) P(A) \times N_m$$

where:

NSM = number of surviving missiles after an enemy attack;

N_m = total number of missiles in the MX system

$P(A)$ = probability that a "loaded" shelter will be under attack;

$P(S|A)$ = probability of a missile surviving the attack A.

The probability of attack, $P(A)$, will depend on the attack strategy deployed by the enemy; i.e. it is a conditional probability $P(A|D_i)$, where D_i is the strategy used by the enemy, such as the "one-on-one" or the "walk" scheme. In the "one-on-one" scheme, the enemy attacks with one warhead for each shelter; whereas, in the "walk" scheme, one warhead is aimed at a cluster of shelters.

On the other hand, the probability of survival, $P(S|A)$, is a function of the hardness and the weapon-placement accuracy of the enemy's warhead. Spacing of the shelters may also be relevant if the enemy deploys the "walk" scheme; however, spacing may not be pertinent for the "one-on-one" scheme.

7.2 The Probability of Attack

Depending on the enemy's arsenal, there may be more than one option for his attack strategy. In other words, there may be several potential threats to a retaliatory system. Accordingly, the attack probability, $P(A|D_i)$, will further depend on the attack strategy deployed by the enemy; moreover, this will obviously also depend on the quality of the enemy's intelligence.

The One-on-One Scheme -- One of the threats to the MX system is if the enemy chooses to deploy one warhead per given shelter. In this case, the probability that a "loaded" shelter will be under attack is a function of the enemy's intelligence and the number of successful RV's (re-entry vehicles) in the attack.

For the one-on-one threat, the probability that a given missile may be under attack can be developed as follows:

Let N_e = number of successful enemy RV's in the attack;
 N_m = total number of missiles in the MX system; i.e. the number of shelters that are "loaded".
 N_s = total number of vertical shelters in the system.
 β = a measure of enemy intelligence, expressed in terms of a fraction of loaded shelters that are known to or can be pin-pointed by the enemy.

Observe that if the enemy has 0 intelligence, and chooses to deploy the one-and-one attack strategy, then this is equivalent to targeting the missiles at random among the N_s vertical shelters. This means

$$\beta = \frac{N_m}{N_s}.$$

The probability of attack, if the enemy uses the one-on-one scheme, therefore, is given by

$$P(A|1/1) = P(N_e/N_m)$$

For example, if the total number of shelters is 5,000, and among these, 250 are loaded, then assuming that half the "loaded" shelters can be identified by the enemy then

$$\beta = \frac{250}{2,500} = 1/10$$

In such a case, the probability of attack becomes

$$P(A|1/1) = \frac{1}{10}(N_e/N_m).$$

Therefore, if the enemy can successfully deliver 250 RV's, the probability of attack of a missile is,

$$P(A|1/1) = 0.1 \times 250/250 = 0.10$$

The "Walk" Scheme -- In the case of the "Walk" scheme, the enemy would presumably aim his RV's at a cluster of shelters; that is, several (e.g. 3 or 4) shelters may be subject to the overpressure from a single weapon burst.

Again, suppose that there is a total of N_s shelters, among which N_m are "loaded". For purposes of this discussion, assume that the shelters are arranged in clusters of 4 in a square as shown in Fig. 7.1.

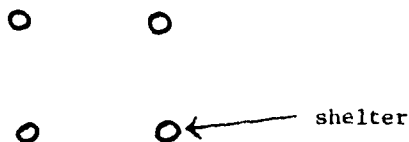


Figure 7.1 Cluster of four shelters in a "square"

Then,

$$\frac{N_m}{\beta N_s} = \text{proportion of the total number of shelters that the enemy estimates are loaded.}$$

For a given target point or "square", the number of corners in the square that the enemy would estimate are loaded can be described as follows:

$$P(Y = y) = \binom{4}{y} \left(\frac{N_m}{\beta N_s}\right)^y \left(1 - \frac{N_m}{\beta N_s}\right)^{4-y}$$

in which, Y = the number of corners that the enemy will think are loaded. The expected value of Y is

$$E(Y) = 4 \left(\frac{N_m}{\beta N_s}\right)$$

Let X = number of corners that are actually loaded. The expected value of X is,

$$E(X) = \beta E(Y) = 4 \left(\frac{N_m}{N_s}\right)$$

Obviously, the enemy can be expected to target only those squares that have at least one loaded shelter; i.e. $Y \geq 1$. On this assumption, the number of missiles in a square, however, may be estimated only in terms of its probability as follows:

$$P(Y = y | Y \geq 1) = \frac{1}{P(Y \geq 1)} \left[\binom{4}{y} \left(\frac{N_m}{\beta N_s}\right)^y \left(1 - \frac{N_m}{\beta N_s}\right)^{4-y} \right]; \quad \text{for } y \geq 1.$$

whereas, the corresponding expected value of Y is

$$E(Y | Y \geq 1) = \sum_{y=1}^4 y P(Y = y | Y \geq 1)$$

Alternatively, of course, the enemy may target only those squares with $Y \geq 2$, or $Y \geq 3$, or only those squares with $Y = 4$. In these latter cases, the corresponding

conditional probabilities and expected values of Y are respectively as follows:

For squares with $Y \geq 2$:

$$P(Y = y | Y \geq 2) = \frac{1}{P(Y \geq 2)} \binom{4}{y} \left(\frac{N_m}{\beta N_s}\right)^y \left(1 - \frac{N_m}{\beta N_s}\right)^{4-y}; \quad y \geq 2$$

$$E(Y | Y \geq 2) = \sum_{y=2}^4 y P(Y = y | Y \geq 2)$$

For squares with $Y \geq 3$:

$$P(Y = y | Y \geq 3) = \frac{1}{P(Y \geq 3)} \binom{4}{y} \left(\frac{N_m}{\beta N_s}\right)^y \left(1 - \frac{N_m}{\beta N_s}\right)^{4-y}; \quad y \geq 3$$

$$E(Y | Y \geq 3) = \sum_{y=3}^4 y P(Y = y | Y \geq 3)$$

Finally for squares with $Y = 4$:

$$P(Y = 4 | Y = 4) = 1.0$$

$$E(Y | Y = 4) = 4.0$$

The conditional expected value of X, therefore, is

$$E(X | Y \geq y) = \beta E(Y | Y \geq y); \quad y = 1, 2, 3, 4$$

Again, assuming that each cluster consists of four shelters in a square, the total number of squares in the entire MX system, therefore, is

$$\frac{N_s}{4}$$

then, the number of squares with $(Y \geq 1)$ is,

$$\text{Number of squares with } Y \geq 1 = [P(Y = 1) + P(Y = 2) + P(Y = 3) + P(Y = 4)] \times \frac{N_s}{4}$$

$$\text{Number of squares with } (Y \geq 2) = [P(Y=2) + P(Y=3) + P(Y=4)] \times \frac{N_s}{4}$$

$$\text{Number of squares with } (Y \geq 3) = [P(Y=3) + P(Y=4)] \times \frac{N_s}{4}$$

and,

$$\text{Number of squares with } (Y = 4) = P(Y=4) \times \frac{N_s}{4}$$

Therefore, if the enemy targets only those squares with $(Y \geq 2)$, the number of successful RV's needed by the enemy will be

$$[P(Y=2) + P(Y=3) + P(Y=4)] \frac{N_s}{4}.$$

In which case, the number of loaded shelters that will be under attack will be

$$E(X|Y \geq 2) \times (\text{No. of squares with } Y \geq 2)$$

Hence, the probability that a given missile will be under attack becomes

$$P(A|\text{Walk}) = \frac{E(X|Y \geq 2) \times (\text{No. of squares with } Y \geq 2)}{N_m}$$

Numerical Example

Suppose that the MX system consists of 4600 shelters, among which 200 of them are loaded with intercontinental missiles. It is probably reasonable to assume that this information is known by the enemy. However, at any given time, the enemy would not know exactly where the 200 missiles may be placed among the 4600 shelters. To illustrate the above model, assume that the intelligence measure of the enemy is $\beta = 1/10$. In this case, $\beta' = \frac{200}{4600} = \frac{1}{23}$ would be equivalent to the absence of intelligence information.

Therefore,

$$\frac{N_m}{\beta N_s} = \frac{200}{0.1 \times 4600} = 0.435$$

from which,

$$P(Y = 0) = (0.435)^0 (1 - 0.435)^4 = 0.102$$

$$P(Y = 1) = 4(0.435) (1 - 0.435)^3 = 0.314$$

$$P(Y = 2) = \frac{4!}{2!2!} (0.435)^2 (1 - 0.435)^2 = 0.362$$

$$P(Y = 3) = \frac{4!}{3!} (0.435)^3 (1 - 0.435) = 0.186$$

$$P(Y = 4) = (0.435)^4 = 0.036$$

Then, the number of squares with $Y \geq 1 = (1 - 0.102) \frac{4600}{4} = 1033$.

Similarly;

$$\text{number of squares with } (Y \geq 2) = (1 - 0.102 - 0.314) 1150 = 672$$

$$\text{number of squares with } (Y \geq 3) = (0.186 + 0.036) 1150 = 255$$

$$\text{number of squares with } (Y = 4) = 0.036 \times 1150 = 41.$$

Also,

$$\begin{aligned} E(Y|Y \geq 1) &= \frac{1}{0.898} (1 \times 0.314 + 2 \times 0.362 + 3 \times 0.186 + 4 \times 0.036) \\ &= 1.94 \end{aligned}$$

Similarly,

$$E(Y|Y \geq 2) = 2.44$$

$$E(Y|Y \geq 3) = 3.16$$

$$E(Y|Y = 4) = 4.00$$

and,

$$E(X|Y \geq 1) = 1/10 \times 1.94 = 0.194$$

$$E(X|Y \geq 2) = 1/10 \times 2.44 = 0.244$$

$$E(X|Y \geq 3) = 1/10 \times 3.16 = 0.316$$

$$E(X|Y = 4) = 1/10 \times 4.00 = 0.400$$

The last four figures given above are simply the expected number of corners with $Y \geq 1$, $Y \geq 2$, $Y \geq 3$, and $Y = 4$; these will, of course, not be realized in practice.

If the enemy chooses to attack only those squares containing at least two loaded shelters, i.e. $Y \geq 2$, then the number of RV's that he must successfully deploy is 672. If the enemy has this capability, the expected number of missiles in the MX system that will be under attack is $672 \times 0.244 = 164$. Therefore, the proportion among the 200 missiles that will be under attack is

$$P(A|\text{walk}) = \frac{164}{200} = 0.82$$

On the other hand, if the enemy has limited capability, say less than 300 RV's, he may have to target only those squares in the MX system with $Y \geq 3$; according to the above calculations, the number of successful RV's required in this case will be 255. Accordingly, the expected number of missiles in the MX system that will be under attack is $255 \times 0.316 = 81$. Then, the proportion of the active missiles that will be under attack is

$$P(A|\text{walk}) = \frac{81}{200} = 0.41$$

7.3 On the Probability of Survival

To achieve or insure a given survival probability, the necessary hardness (e.g. expressed in terms of overpressure capacity) must be designed into the individual shelter, as well as of its components such as the shock isolation system and the closure. The required hardness may be determined by applying the appropriate "factor of safety" which is simply a factor that may be used to amplify the weapon-induced overpressure in order to determine the required over capacity of the shelter for insuring the desired survival probability. For this purpose, the appropriate relationship between survival probability and the median safety factor ought to be developed, which may be in the form shown schematically in Fig. 7.2.

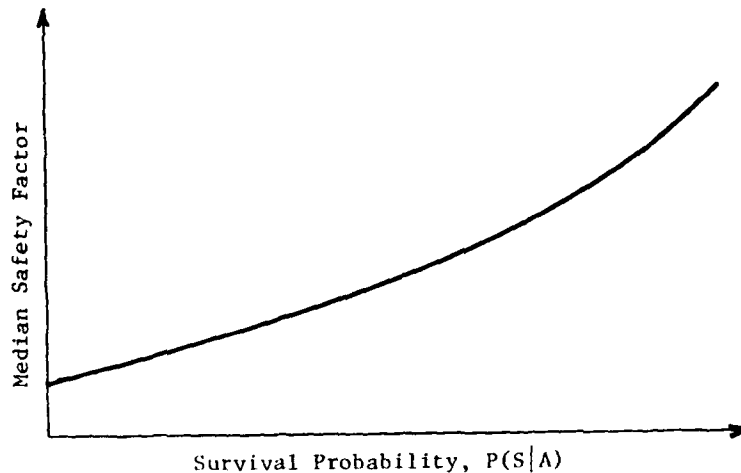


Figure 7.2 Safety factor versus survival probability

That is, hardness may be defined as the median overpressure capacity of a shelter or of its associated components. In this context, it is a function of the desired probability of survival against a given overpressure environment, as indicated in Fig. 7.2. Conversely, the probability of survival is a function of the relative positions between the PDF (probability density function) of the applied overpressure and the PDF of the overpressure capacity of a shelter as shown in Fig. 7.3. The relative positions may be given in terms of the ratio of the median overpressure capacity to the median of the applied overpressure, $\bar{R}/\bar{S}$, which may be called the median safety factor. Additionally, the survival probability is also a function of the uncertainty in the applied environment as well as in the overpressure capacity of the shelter. The uncertainties in the applied overpressure will include those associated with the following factors:

- inaccuracy in the overpressure-distance relationship;
- error in the estimation of the coupling effect;
- the determination of the effect of the height-of-burst;
- the analysis of the effect of soil-structure interaction, among others.

The uncertainties in the overpressure capacity of a shelter facility will include those arising from the following:

- inaccuracy in the determination of the strength of vertical concrete/steel cylinder;
- the determination of the strength of closure;
- the determination of the capacity of a shock isolation system;
- error in the calculation of the response, etc.

7.4 A Trade-Off Problem

If a missile is under attack, its probability of survival under a "one on one" threat would be lower than that under a "walk" scheme. On the other hand, for the same number of RV's and the same enemy intelligence, the probability that a missile will be under attack is correspondingly lower with the "one on one" scheme compared with the "walk" scheme, especially if the shelters are spaced sufficiently far apart.

In short, the number of surviving missiles is a function of the probability of attack as well as the conditional probability of surviving an attack. If there is a relative likelihood that the enemy may employ one or the other schemes, say with probabilities p_o and p_w in the case of the "one-on-one" and "walk" schemes such that p_o and $p_w = 1.0$, then the expected NSM is

$$NSM = \{ [1 - P(F|A_o) P(A_o)] p_o + [1 - P(F|A_w) P(A_w)] p_w \} N_m$$

where $P(F|A) =$ conditional probability of failure given an attack.

In the case of the "one-on-one" scheme the probability of attack $P(A_o)$ is simply a function of N_e , N_m and β . Implicitly, this probability is also a function of the total number of shelters N_s , since β is implicitly a function of N_s .

The probability of survival of a given missile, once it is under attack, will be fairly low in the case of the "one-on-one" attack scheme. Its chance of surviving an enemy RV will be largely due to the weapon inaccuracy, i.e. the CEP. Furthermore, the probability of survival, $P(S|A_o)$, will not be affected by the spacings between the shelters if the enemy deploys the "one-on-one" scheme.

However, in the case of the "walk" scheme, the probability of attack $P(A_w)$, is a function of N_e , N_m , N_s , and β . That is,

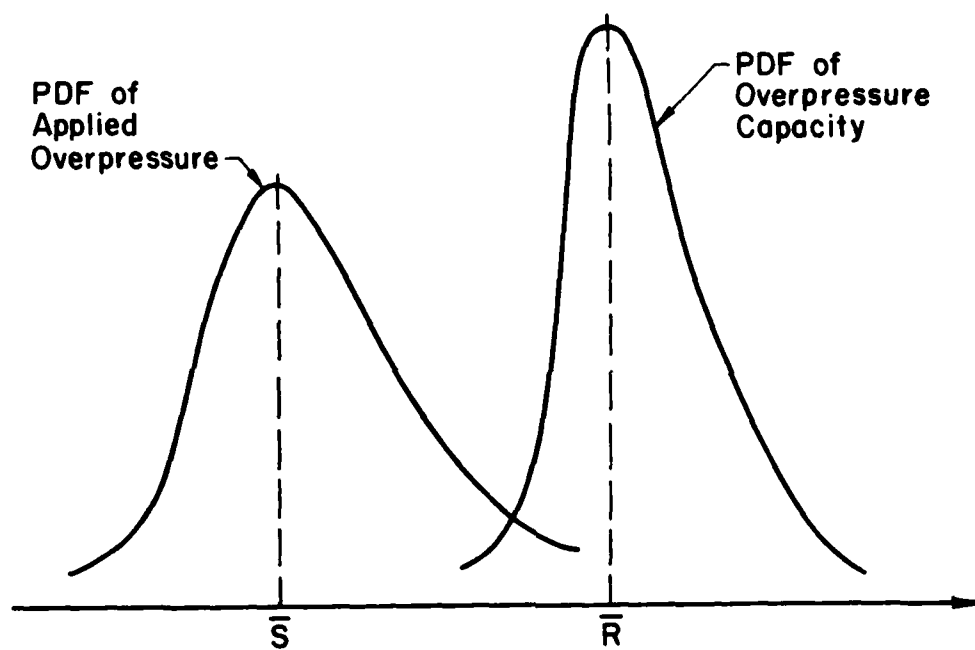


Figure 7.3 Probability of failure

$$P(A_w) = g(N_e, N_m, N_s, \beta)$$

Increasing the number of shelters, N_s , will decrease $P(A_w)$, and thus increases the NSM; this would mean using closer spacings among the shelters for the same land area or the same total cost. However, for a given shelter hardness, the probability of survival, of course, will decrease with closer spacings. The spacing may be increased by using fewer shelters; this, will have the effect of increasing $P(A_w)$ but will also increase $P(S|A_w)$.

Increasing the hardness of the shelter (i.e. increasing its overpressure capacity) against a given weapon yield will also increase the probability of survival $P(S|A_w)$.

In otherwords, in the case of the "walk" scheme, there is an optimal expected NSM between using large spacing (say 9,000 ft) with fewer shelters vs. using closer spacing (say 4,000 ft) with correspondingly larger number of shelters and higher level of shelter hardness. That is, for the "walk" scheme, increasing the hardness as well as increasing the spacings between shelters will increase the survival probability; however, within a finite land area or limited budget, larger spacings between the shelters will mean smaller number of shelters and thus increasing the probability of attack $P(A_w)$ of the missiles.

In this light, there is therefore a trade-off between spacing and hardness in order to achieve the maximum possible NSM (the number of surviving missiles), within a given total budget for the MX system. In order to perform such a trade-off study between these major factors, the necessary information must be developed. Such information would include the following.

(1) The relationship between the probability of attack $P(A_w)$ as a function of the variables N_e , N_m , N_s , and β . These relationships may then be used to develop the relationships of $P(A_w)$ as a function of the shelter spacing.

(2) In the case of the probability of survival, $P(S|A_w)$, increasing the spacing between shelters will correspondingly decrease the overpressure on a given shelter for a given weapon yield, and thus increases the probability of survival. Also, increasing the hardness of a shelter will increase the survival probability since the capacity of the shelter to withstand the applied overpressure will correspondingly increase. These relationships should also be developed.

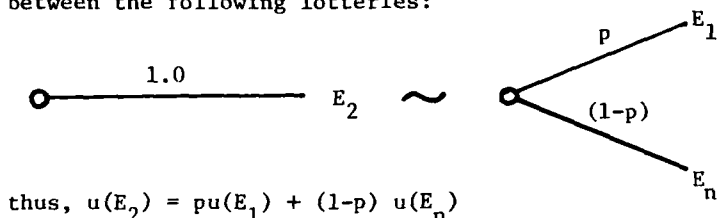
REFERENCES

- Ang, A. H-S. and Tang, W. H., Probability Concepts in Engineering Planning and Design, Vol. I - Basic Principles, John Wiley & Sons, 1975.
- Fishborn, P. C., Utility Theory for Decision Making, Vol. 18, Publication in Operation Research Series, edited by David B. Hertz, John Wiley & Sons, Inc., New York, 1970.
- Keeney, R. C., "Utility Functions for Multi-Attributed Consequences" Management Science, Vol. 18, pp. 276-287, 1972.
- Raiffa, H., Decision Analysis, Addison-Wesley, 1968.
- Raiffa, H., "Preferences for Multi-Attributed Alternatives," RM-5868 - DOT/RC, Rand Corp., Santa Monica, Calif., April 1969.
- Schlaifer, R., Analysis of Decisions Under Uncertainty, McGraw-Hill, 1969.
- Spetzler, C. S. and Stahl von Holstein, C. S., "Probability Encoding in Decision Analysis," Management Science, 22, 1975, pp. 340-358.

APPENDIX A: EVALUATION OF UTILITY IN A DECISION TREE

The procedure is as follows:

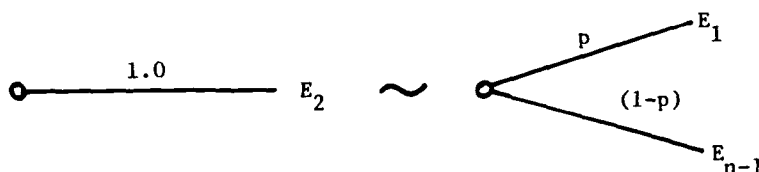
- (i) Rank the events involved in the decision tree in order of preference such that $E_1 > E_2 > \dots > E_n$; where the symbol $>$ means "is preferred to."
- (ii) Assign $u(E_1) = 1.0$ and $u(E_n) = 0$.
- (iii) Establish the value of p such that the decision maker is indifferent between the following lotteries:



$$\text{thus, } u(E_2) = pu(E_1) + (1-p)u(E_n)$$

$$= p$$

- (iv) Repeat step (iii) $n-3$ times with E_2 replaced each time by $E_3, \dots, E_{n-1}$ respectively; the value of p obtained for different E_1 will generally be different.
- (v) At this step, a set of utilities for $E_1, E_2, \dots, E_n$ has been determined. However, to provide cross checking on these values, one may start another series of inquiry using $u(E_1)$ and $u(E_{n-1})$ as the set of new reference points and determine a new utility value for E_2 , namely $u'(E_2)$, such that indifference is achieved as follows:



$$\text{where } u'(E_2) = pu(E_1) + (1-p)u(E_{n-1})$$

$$= p + (1-p)u(E_{n-1})$$

If the utility values are consistent, $u'(E_2)$ should be equal to $u(E_2)$ as determined earlier in step (iii).

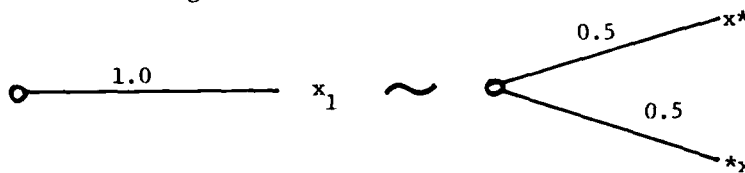
- (vi) Repeat step (v) $n-4$ times with E_2 replaced each time by $E_3, \dots, E_{n-2}$ respectively.

If any inconsistencies are found in the above procedure, repeat the questioning process until all utility values agree satisfactorily.

APPENDIX B: DETERMINATION OF UTILITY FUNCTION

A procedure for determining a utility function is as follows:

- (i) Identify the range of values of the specific attribute covered by the decision analysis.
- (ii) Assign utility values of 1.0 and 0, respectively, to the utilities of the two extremes.
- (iii) The utility of an intermediate value of the attribute may be determined from the value of X such that indifference is achieved by the decision maker on the following lotteries:

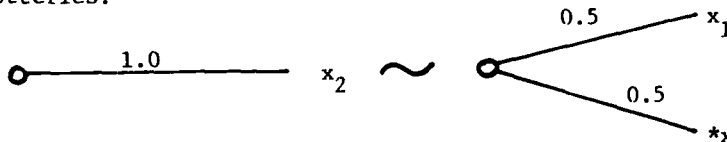


$$\text{Hence, } u(X_1) = 0.5u(X^*) + 0.5 u(*X)$$

$$= 0.5 \times 1 + 0.5 \times 0 = 0.5$$

where X^* and $*X$ are, respectively, the extreme values of the attributes with utilities 1.0 and 0, respectively.

- (iv) Repeat step (iii) by replacing X^* by X_1 ; the utility of another value of the attribute could be obtained from the following pair of indifference lotteries:



$$\text{Hence, } u(X_2) = 0.5 u(X_1) + 0.5 u(*X)$$

$$= 0.5 \times 0.5 + 0.5 \times 0 = 0.25$$

- (v) Repeat the above procedure by varying the values of the attribute on the 50-50 lottery, to obtain the utilities for other attribute values.
- (vi) Plot the utility values as a function of the attribute value X , and fit a curve through the set of points. The resultant curve is the utility function.

PRECEDING PAGE BLANK-NOT FILLED

DISTRIBUTION LIST

DEPARTMENT OF DEFENSE

Defense Communications Agcy
2 cy ATTN: Code C-661, T. Trenkle

Defense Nuclear Agcy
2 cy ATTN: SPSS, B. Smith
4 cy ATTN: TITL

Defense Technical Info Ctr
12 cy ATTN: DD

DEPARTMENT OF THE ARMY

US Army Engr Waterways Exper Sta
2 cy ATTN: P. Mlakar

DEPARTMENT OF THE NAVY

Naval Surface Weapons Ctr
2 cy ATTN: W. McDonald
2 cy ATTN: R. Barash

DEPARTMENT OF ENERGY CONTRACTORS

Lawrence Livermore National Lab
2 cy ATTN: L-531, D. Morgan

DEPARTMENT OF DEFENSE CONTRACTORS

Agbabian Associates
2 cy ATTN: C. Bagge

Boeing Co
2 cy ATTN: S. Irack

Boeing Computer Services, Inc
2 cy ATTN: R. Erickson

DEPARTMENT OF DEFENSE CONTRACTORS (Continued)

IIT Research Inst
2 cy ATTN: A. Longinow

University of Illinois
2 cy ATTN: Y. Lin

University of Illinois
4 cy ATTN: A. H-S Ang
4 cy ATTN: W. Tang

J. H. Wiggins Co, Inc
2 cy ATTN: J. Collins

Kaman Tempo
ATTN: DASIAC

Pacific-Sierra Research Corp
ATTN: H. Brode

R & D Associates
ATTN: P. Haas

R & D Associates
2 cy ATTN: H. Cooper

Rand Corp
2 cy ATTN: A. Laupa

Science Applications, Inc
2 cy ATTN: G. Binninger

TRW Defense & Space Sys Grp
2 cy ATTN: P. Dai

PRECEDING PAGE BLANK-NOT FILMED

