

INTERVAL ESTIMATION OF A STIMULUS LEVEL OF
ORDER α IN SENSITIVITY TESTING

Lakshmi Tatikonda

June 1, 1966

INTERVAL ESTIMATION
OF A STIMULUS LEVEL OF
ORDER α in SENSITIVITY
TESTING

by
Lakshmi Tatikonda

This monograph which is essentially the authors doctoral dissertation was prepared with support of Computation and Analysis Division of NASA-MSC (Houston) Contract NAS 9-2619.

TABLE OF CONTENTS

CHAPTER

1.	INTRODUCTION	1 - 9
2.	SEVERAL EXISTING METHODS FOR ANALYZING SENSITIVITY DATA	10-25
3.	A NEW METHOD TO OBTAIN CONFIDENCE LIMITS ON, x_{α} IN UP-AND-DOWN TESTING	26-48
4.	EXPERIMENTAL PROCEDURE TO OBTAIN SAMPLE SIZES IN VARIOUS TECHNIQUES	49-65
5.	COMMENTS ON THE RESULTS OF THE EMPIRICAL STUDY	66-69
	BIBLIOGRAPHY	70-72

Chapter I

INTRODUCTION

Sensitivity testing is a term associated with tests characterized by a sample specimen being subjected to a stimulus of known intensity and the specimen either "responds" to the stimulus or does not "respond" to the stimulus depending on whether some critical physical threshold has or has not been exceeded for that particular sample specimen. That is, for some "stimulus-subject systems" quantitative measurement of the response attributed to the action of the stimulus is impossible or almost impracticable. For example in testing the explosives a common procedure is to drop a weight on specimens of some explosive mixture from various heights and observe whether it explodes or not. There are heights at which some of the specimens will explode and some of them will not. It is assumed that the specimens which will not explode would explode if the weights were dropped from a sufficiently high level. Therefore we suppose that there is a critical height associated with each specimen and there will be "response" or "non response" depending on whether the critical level is or is not exceeded by the intensity imposed by the weight dropped. Thus the population of the specimen is characterized by a continuous variable whose critical height can not be measured [7].

This situation arises in many fields of research. For example insecticides are assayed by assigning batches of insects to standard test preparations and analyzing the relationship between the death rate

and the dose; that is, to observe whether the critical dose for the insect is less than or greater than the selected dose. The same difficulty arises in pharmacuetical research dealing with germicides, anesthetics and similar drugs, in explosives, propellants, detonation devices and armor-piercing projectiles. Perhaps its earliest implementation was in biological studies of dosage mortality and response to drugs [25].

Although the application has been diverse, the sensitivity experiments have many characteristics in common. In true sensitivity experiments it is not possible to take more than one observation on a given specimen. The measurement at any point in the scale destroys the specimen so that a new specimen is required for each measurement. Neither the insect that has died or weakened, nor the explosive having been packed can be used again as a sample. That is, once a test has been made the specimen is altered and so a bonafide result can not be obtained from a second test. A common procedure in this type of experiment is to divide the sample into several groups (usually but not necessarily of the same size) and to test one group at a chosen level, and a second group at a second level and so on.

There are several methods of obtaining and analyzing the above described data. One of the oldest methods is the "probit method." The basis of this method is the linear transformation of the normal curve, with

$$\int_{-\infty}^x \frac{e^{-\frac{1}{2}(t^2)}}{\sqrt{2\pi}} dt$$

being considered as giving percentage response at the level x . The "probit" y is obtained by adding 5 to x in order to avoid negative value in the use of transformation. This transformation makes it possible to represent the relation between the percentage response and the dose as a linear relation, and reduces the problem to one of linear regression [21]. This was developed by Bliss [4] and Fisher [11]. In 1933 Gaddum has showed that with the transformation of percentage effects, what he called "normal equivalent deviations" (n.e.d.), one can plot an "S" shaped curve. It is interesting to note that the history of the "Probit method" goes back to 1860. Fechner, a physiologist used a method which is essentially Gaddum's n.e.d. to express the proportion of trials [21].

In latter years a number of other methods were developed to handle the sensitivity data with or without transformation. Some of the well known methods are the Spearman-Kärber method, the method of extreme effective doses, the Reed-Muench method, the Dragsted-Behrne's method and the Moving average method. A relatively new technique called the "up-and-down method" was developed during World War II and is used in explosives research. A method to obtain and analyze the sensitivity data using up-and-down technique was given by Dixon-and-Mood [7].

But most of these commonly used methods are applicable only in special cases and are based on various assumptions concerning the distribution of the estimators, especially if the confidence

limits are desired. For example in the Spearman-Kärber method one considers the logarithm of the tolerance as being approximately distributed according to a normal density. We define the tolerance distribution as follows.

Definition 1. Let $Y(x)$ be a random variable defined on the closed interval $[a, b]$, then we say that the function $y = M(x)$ is a regression curve of y on x if and only if

$$E[Y(x)] = M(x),$$

where $E[\cdot]$ denotes the expectation of $[\cdot]$.

Definition 2. The regression curve $M(x)$ is said to be a response curve if

$$(c_1) \quad M(x) \text{ is continuous on } [a, b]$$

$$(c_2) \quad 0 \leq M(x) \leq 1$$

$$(c_3) \quad \int_0^1 H[Y(x)] dY = 1 \quad \text{for each } x \text{ on } [a, b]$$

where $H[Y(x)]$ denotes the family of density functions which depends on the parameter x .

A response curve is sometimes called a tolerance distribution by applied statisticians [12].

That is, in Spearman-Kärber method

$$M(z) = \int_{-\infty}^z \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{1}{2}\left(\frac{t-\mu}{\sigma}\right)^2} dt,$$

where $z = \log x$, and x is the stimulus level.

The Reed-Muench, Dragsted-Behrens and the Moving average methods can be applied only when $M(x)$ is symmetrical. Dixon and Mood's method also considers that

$$M(z) = \int_{-\infty}^z \frac{1}{\sqrt{2\pi}\sigma} e^{-\left(\frac{t-\mu}{\sigma}\right)^2} dt$$

where $z = \log x$ or some function of x (where there must be a 1-1 correspondence between z and x), and x , the stimulus level. Most of these methods are relatively effective only when the goal of the tests is to estimate the mean (50 percent-point) or the median. Recently much attention has been given to the problem of estimating the portions of "tails" of the response curves, i.e., to estimate the stimulus levels with either very small or very large response probabilities. We define the response probability as follows:

Definition 3: The statement that y is the probability of response at the stimulus level x means that $Y(x)$ is assigned unity when a response occurs, that is

$$y = \Pr [Y(x) = 1],$$

or equivalently,

$$y = E [Y(x)] = M(x).$$

where $\text{Pr} [\cdot]$ should be read "probability that the event $[\cdot]$ occurs."

In many instances, the experimenter has no clear notion of the structure of the regression curve he wishes to study. In these cases it is not advisable to permit a hypothesis as to the precise shape of the regression curve nor to describe possibly various distribution features. Moreover incorrect assumptions can have serious effects on the design and analysis of the method. Unlike some other types of designs (e.g. factorial) and analysis (e.g., regression), sensitivity procedures depend critically on the accuracy of the distribution assumptions. Thus, for example, if it is assumed that a response function is cumulative normal when it is really, say, uniform, then the design and analysis chosen will be satisfactory only when the median response (that is, the 50 percent stimulus level) is of primary interest. In other regions, corresponding to small or large response probabilities, highly inefficient designs and inaccurate analysis will result. Thus, generally speaking, a parametric approach should be selected only after an extensive amount of previous experience and complete data analysis are available. But often this is not the case, since many sets of sensitivity data are not consistent with the classical response function forms, and it becomes necessary to employ distribution - free (or non-parametric) methods to analyze the sensitivity data.

Fortunately, several distribution-free methods have been developed. Most of these methods in general assume only that the response function is nondecreasing with increasing stimuli. This is the case in general although there are occasional instances in which non-monotonic behavior has been noted. J. B. Gayle [13] has investigated this property by carrying out a larger number of replicate tests to permit a statistical analysis of data. The results indicated that over a considerable range of stimulus levels the frequency of response decreased significantly with increasing stimulus levels.

The most commonly encountered type of sensitivity problem is that of finding at what level of the stimulus variable a given percent response will occur. For example, in biological assay it is often necessary to determine the dose (called LD 50 or ED 50) which is effectively the median of the distribution of responses. Often one is interested in the stress that induces a detonation, say, 95 percent of the time. In each of these situations one is concerned with inverting the relation which gives the probability of response as a function of the stimulus. Thus it is also known as the "inverse response problem".

In general the problem can be stated as follows. Let x be a stimulus level and let the probability of response at x be $M(x)$. That is, let a random variable $y(x)$ take the value unity when a response occurs and zero when a response does not occur.

Hence, for every x , a part $M(x)$ of the population from which the sample specimens are selected will respond when subjected to a stimulus level x , and the remaining part $1 - M(x)$ will not.

Definition 4. The value x_α for each value α such that $0 < \alpha < 1$ is said to be the stimulus level of order α if and only if x_α satisfies the regression equation

$$M(x) = \alpha \quad (1.1)$$

Definition 5. The response curve $M(x)$ is said to be a quantal response curve if and only if the density function $H[y(x)]$ is a Bernoulli probability law [[23] Parzen, p. 218], that is

$$\begin{aligned} H[y(x)] &= [M(x)]^{y(x)} [1 - M(x)]^{1-y(x)}, \quad y(x) = 0, 1 \\ &= 0 \quad \text{elsewhere.} \end{aligned}$$

In 1951 a technique was developed by Robbins and Monro [24], in which one can estimate x_α for a given α without any knowledge of $M(x)$. The method has been called the "Stochastic Approximation Method". The estimation procedure is sequential and distribution free.

In recent years work has been done to construct a confidence interval of preassigned length at a given confidence level. These techniques were developed by R. H. Farrell [10] and P. L. Odell [22]. A generalization of the second technique is developed by the author

and it is described in Chapter III. In Chapter II a brief description of the methods given by Spearman-Kärber, Dixon-Mood, and Farrell is furnished. An empirical study of the sample size, for a fixed confidence level and for a preassigned length of the interval is made. The results of the study and a discussion of the results are given in Chapters IV and V.

Chapter II

SEVERAL EXISTING METHODS FOR ANALYZING SENSITIVITY DATA

In order to gain a measure of efficiency of the method formulated in Chapter III, with respect to those methods already available, three different methods were selected to base the comparative study. These methods are:

- T1. Spearman-Kärber technique [12]
- T2. Dixon-Mood technique [7],
- T3. Farrell's technique [10].

T1 and T2 are well-known methods, while T3 appears not so well-known. A Fourth method formulated by Alexander [25], has appeared recently in an industrial report; however, the theory is not available at this time hence the method could not be included in the comparison. Methods T2 and T3 involve a sequential technique for analyzing the up-and-down data, while method T1 is a fixed design.

Farrell's analysis also includes a way to obtain confidence limits using stochastic approximation [10]; however, we have restricted our study to up-and-down techniques and the Spearman-Kärber technique. A brief description of each of these techniques follow in this chapter.

2.1 Notation and Preliminaries

Though different authors used different notations in their works, a single notation is used throughout this dissertation to facilitate reading and for easier comparison of the results. The

notation is presented as follows:

The symbols $x_1, x_2, \dots, x_n$ denote the stimulus levels at which the specimens are to be tested; y denotes the probability of response at a given stimulus level x . Then the random variable y takes on the value unity or zero. That means,

$$\begin{aligned} Y(x) &= 1 \text{ if the specimen responds and} \\ Y(x) &= 0 \text{ if the specimen does not respond.} \end{aligned} \quad (2.1)$$

Also $E[Y(x)] = M(x)$, where $M(x)$ is defined in Chapter I, and $E[\cdot]$ denotes the expected value of $[\cdot]$.

In all the methods described below, the aim is to construct a confidence interval for various stimulus levels, that is, for a given α to obtain bounds L_1 and L_2 such that

$$P[L_1 \leq x_\alpha \leq L_2] \geq 1 - \beta \quad (2.2)$$

where β is the desired confidence coefficient, and $P[A]$ denotes the probability that the event A occurs. Frequently we will let $P[A|B]$ denote the probability that the event A occurs given that the event B has already occurred.

2.2 Spearman-Kärber Method

Perhaps the oldest technique and for years the most often used for estimating x_{50} is the so-called, Spearman-Kärber method, originally formulated in 1908 by Spearman [26] and later in 1931 by Kärber [18]. The following are the assumptions which restrict the method:

A1. There exist numbers a and b such that $a < b$ and

$$P [Y(x') = 0 \mid x' \leq a] = 1 - \epsilon$$

$$P [Y(x') = 1 \mid x' \geq b] = 1 - \epsilon$$

where $\epsilon > 0$ is arbitrarily small and $x' = \log x$.

A2. The response curve $M(x') = E [Y(x')]$ has the form of the cumulative normal distribution which can be written

$$M(x') = \int_{-\infty}^{x'} \{ 1/(2\pi)^{1/2} \sigma \} \exp \{ (x' - x'_{50})^2 / 2\sigma^2 \} dx'$$

where $\sigma > 0$ and $a < x'_{50} < b$.

Let $x'_1, x'_2, \dots, x'_k$ be k equally spaced levels with $x'_1 < x'_2 < \dots < x'_k$ such that at each level n specimens are to be tested. Let

$$r_i = \sum_{j=1}^n Y_j(x'_i) \quad i = 1, 2, \dots, n \quad (2.3)$$

where $Y_j(x'_i) = 1$ or 0 , according to (2.1). Hence r_i is simply the number of specimens which respond at the i th level, x'_i . Then x'_{50} is estimated by the statistic $\hat{x}'_{50}$ defined as

$$\hat{x}'_{50} = x'_k + d/2 - d \sum_{i=1}^k r_i/n \quad (2.4)$$

where x'_k is the level such that $r_k = n$, and $d = (x'_{i+1} - x'_i) / i = 1, 2, \dots, k-1$, the common difference between the adjoining

stimulus levels. Neither Spearman nor Kärber list any rules for selecting d . If the deviation of x'_{50} from the nearest x'_i is selected at random between $d/2$ and $-d/2$, as is effectively done when the experimenter begins without the knowledge of x'_{50} , then

$$E [\hat{x}'_{50}] = x_{50}$$

and

$$\sigma^2(\hat{x}'_{50}) = d^2 \sum_{i=1}^k M_i (1 - M_i) / n$$

for $n > 4$; $\sigma^2(m)$ denotes the variance of the statistic m and M_i denotes $M(x'_i)$ in A2. Confidence limits for x'_{50} are obtained by assuming

$$A3. \quad (\hat{x}'_{50} - x'_{50}) / \sigma(\hat{x}'_{50}) \sim N(0, 1)$$

where $\sim$ should be read as "distributed according to" and $N(0, 1)$ should be read as "normally distributed with mean zero and variance one". It is clear that the desired confidence limits follow from the statement that

$$P [\hat{x}'_{50} - k_{\beta/2} \sigma(\hat{x}'_{50}) \leq x'_{50} \leq \hat{x}'_{50} + k_{\beta/2} \sigma(\hat{x}'_{50})] \geq 1 - \beta \quad (2.5)$$

where $k_{\beta/2}$ denotes the value of k such that

$$1 - \beta/2 = \int_{-\infty}^k (1/2\pi)^{1/2} \exp \{ -t^2/2 \} dt \quad (2.6)$$

and

$$\sigma^2(\hat{x}'_{50}) = [d^2/n^2 (n-1)] [\sum_{i=1}^k r_i (n - r_i)] \quad (2.7)$$

an unbiased estimate of $\sigma^2(x_{50})$ given by Ervin and Cheesman [17].

An alternative suggested by Gaddum is to use the following formula

$$\sigma^2(\hat{x}'_{50}) = .564 \sigma d/n \quad (2.8)$$

where σ is known and appears in A2. Unfortunately this requires an apriori information about $M(x)$, and is a further restriction. Techniques are available [12] for estimating σ from the data, giving an estimate which can be used for σ in (2.8), however in the simulation (2.7) was used as the true variance of $\hat{x}'_{50}$ and not (2.8).

The confidence interval I'_β for x'_{50} then follows from (2.5) and is

$$I'_\beta = [\hat{x}'_{50} - k_{\beta/2} \hat{\sigma}(\hat{x}'_{50}), \hat{x}'_{50} + k_{\beta/2} \hat{\sigma}(\hat{x}'_{50})] \quad (2.9)$$

Then one can convert these numbers into original stimulus units by taking anti-logarithms giving the desired confidence limits for x_{50} , those being

$$I_\beta = [\text{antilog}(\hat{x}_{50} - k_{\beta/2} \hat{\sigma}(\hat{x}_{50})), \text{antilog}(\hat{x}_{50} + k_{\beta/2} \hat{\sigma}(\hat{x}_{50}))]$$

2.3 Dixon-Mood Method

Chronologically this method is one of the first methods of the so called "up-and-down" tests. In order to perform the up-and-down testing one chooses an initial stimulus level, say x_0 (the best apriori estimate of x_{50}), and a succession of levels $x_1, x_2, \dots$

whose magnitudes exceed x_0 ; and a succession of levels $x_{-1}, x_{-2}, \dots$ whose magnitudes are less than x_0 . The first sample specimen is subjected to a stimulus level, x_0 . If the specimen responds the second specimen is subjected to a stimulus level x_1 . In general if the i th specimen responds (does not respond) to a stimulus level of order j , then the $(i + 1)$ st specimen is subjected to a stimulus level of order $j - 1$ ($j + 1$). It is noted that one would expect the number of responses to be approximately equal to the number of non-responses in such tests. One needs the following restrictions in order to apply this method.

B1. If $x'_0, x'_{+1}, x'_{+2}, \dots$ are the levels such that $x' = g(x)$ where there is a 1 - 1 correspondence between x and x' then

$$M(x') = \int_{-\infty}^{x'} \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{1}{2}\left(\frac{t-x}{\sigma}\right)^2} dt \quad (2.10)$$

B2. One must be able to make a rough estimate of σ prior to the experimentation.

B3. The sample size should be large (about 40 to 50) in order to apply the given analysis.

Dixon and Mood formulated the following statistics to estimate stimulus levels of order α and θ .

$$\hat{x}'_{\alpha} = x' + d\left(\frac{A}{N} \pm \frac{1}{2}\right) + k_{\alpha} s = \hat{x}'_{50} + k_{\alpha} s \quad (2.11)$$

where the $(-\frac{1}{2})$ is used if the number of responses is less than the

number of nonresponses and the $(+\frac{1}{2})$ is used if the number of nonresponses is less than the number of responses.

$$S = (1.62) d \left[\frac{NB - A^2}{N^2} + .029 \right]. \quad (2.12)$$

where the symbols d , A , B , N , k_α and x^* are defined as follows:

$$d = |x'_i - x'_{i+1}|$$

and is assumed to be the same for all $i = 0, \pm 1, \dots$

If $n_0, n_1, n_2, \dots, n_k$ are the observed frequencies at various levels (in increasing order) at which the less frequent event occurs, and x^* is the lowest level at which the less frequent event occurs, then

$$N = \sum_{i=0}^k n_i \quad (2.13)$$

$$A = \sum_{i=0}^k i n_i \quad (2.14)$$

$$B = \sum_{i=0}^k i^2 n_i, \quad (2.15)$$

and k_α is the value of k such that

$$\int_{-\infty}^k \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}t^2} dt = \alpha. \quad (2.16)$$

(2.11) is valid only if

$$\frac{NB - A^2}{N^2} > .3 \quad (2.17)$$

If (2.17) is not true, then one can take the sample size sufficiently large so that (2.17) is true or use the analysis given in [7]. Dixon and Mood have shown that the standard deviation of the estimator

$$\hat{x}'_{50} = x^* + d\left(\frac{A}{N} \pm \frac{1}{2}\right)$$

is approximately

$$G\sigma/\sqrt{N} \quad (2.18)$$

and can be estimated by

$$\hat{\sigma}_{50} = s_m = Gs / \sqrt{N} \quad (2.19)$$

In a similar way the standard deviation of s is estimated by

$$\hat{\sigma}_s = s_s = Hs / \sqrt{N} \quad (2.20)$$

where G and H are empirical values depending on d and σ and are given here in the form of a graph in Fig. 2.1.

$$\text{The variance of } \hat{x}'_{\alpha} \text{ is given by } \text{var}(\hat{x}'_{\alpha}) = \text{var}(\hat{x}'_{50}) + 2k_{\alpha} \text{cov}(\hat{x}'_{50}, s) \quad (2.21)$$

It can be shown that $\text{cov}(\hat{x}'_{50}, s) \approx 0$. Hence we approximate $\text{var}(\hat{x}')$ by

$$\hat{\sigma}_{\alpha}^2 = s_m^2 + k_{\alpha}^2 s_s^2 \quad (2.22)$$

where s_m and s_s are given by (2.19) and (2.20). In order to obtain

the confidence limits a further assumption is made, that being

$$(B4) \quad \frac{\hat{x}'_{\alpha} - x'_{\alpha}}{\hat{\sigma}_{\alpha}} \sim N(0, 1). \quad (2.23)$$

Then the confidence interval follows by the statement that the

$$\Pr [\hat{x}'_{\alpha} - K_{\beta/2} \hat{\sigma}_{\alpha} \leq x'_{\alpha} \leq \hat{x}'_{\alpha} + K_{\beta/2} \hat{\sigma}_{\alpha}] \geq 1 - \beta \quad (2.24)$$

where $K_{\beta/2}$ is defined in (2.6).

Because of the 1 - 1 correspondence between x and x' , the probability statement (2.24) can be written as

$$\Pr [g^{-1} (\hat{x}'_{\alpha} - K_{\beta/2} \hat{\sigma}_{\alpha}) \leq x_{\alpha} \leq g^{-1} (\hat{x}'_{\alpha} + K_{\beta/2} \hat{\sigma}_{\alpha})] \geq 1 - \beta, \quad (2.25)$$

which is the desired confidence interval.

Farrell's Method.

A relatively recent (1962) method which yields non parametric confidence limits for x_{α} is given by Farrell [10]. In this method one does not make any distribution assumptions on $M(x)$, the response function, but needs only the following assumptions.

- C1. $M(x)$ is a monotonic function of the stimulus level x .
- C2. There is an x such that

$$M(x) < \alpha \quad \text{for } x < x_\alpha$$

and

$$M(x) > \alpha \quad \text{for } x > x_\alpha.$$

C3. There exists a family of distribution functions

$\{ G(\cdot, w), w \in \Lambda \}$, where Λ is the finite or infinite open real number interval. Certain knowledge of G is known to the experimenter. That is, G is an independent Bernoulli distribution in the case of up-and-down testing.

C4. For all $x \in (-\infty, \infty)$,

$$M(x) = \int_{-\infty}^{\infty} x \, dG(x, w).$$

The method is quite complicated and will briefly be described as follows: Let $\{ x_n, -\infty < n < \infty \}$ be the sequence of stimulus levels such that $x_{n+1} - x_n = \Delta > 0$ for all n , an integer; where 2Δ is the desired width of the confidence interval, and $\lim_{|n| \rightarrow \infty} |x_n| = \infty$. Again $M(x)$ is the probability of response when the testing is done at the level x .

The initial experiment is made at the level x_0 , and the $(n+1)$ th experiment will be made at the level $x_{N(n+1)}$, where $N(n+1) = N(n) \pm 1$ according as whether there is a non response, or a response at the n th experiment. This can be written in the form

$$N(n+1) = N(n) + g(Z_n, F(x_{N(n)})), \quad (2.26)$$

where

$$g(x,y) = -1 \quad \text{if } x \leq y \quad (2.27)$$

and

$$g(x,y) = +1 \quad \text{if } x > y$$

and Z_n is distributed uniformly on $(0, 1)$.

At this stage, it seems desirable to quote a slight modification of the lemma given in [10]. The modification is made in order to attain compatability with our notation.

Lemma 2.1. Suppose Λ is an open real number interval,

$\{G(\cdot, w), w \in \Lambda\}$ a monotonic family of distributions. If $w \in \Lambda$, $\alpha = \int Y dG(Y, w)$, $\int x^2 dG(x, w) < \infty$, and $\beta > 0$, then there exist real number sequences $\{a_n, n \geq 1\}, \{b_n, n \geq 1\}$ such that $a_n > b_n$ for all $n \geq 1$, and

$$\Pr \left[\text{some } \sum_{i=1}^n Y_i > a_n \right] \leq \beta/4, \quad \text{if } 1 - \alpha < \alpha \quad (2.28)$$

$$\text{and } \Pr \left[\text{some } \sum_{i=1}^n Y_i \leq b_n \right] \leq \beta/4, \quad \text{if } 1 - \alpha \geq \alpha \quad (2.29)$$

$$\text{and } \lim_{n \rightarrow \infty} \frac{a_n}{n} = \lim_{n \rightarrow \infty} \frac{b_n}{n} = \alpha,$$

where $1 - \beta$ is the desired confidence level.

The sequences $\{a_n, n \geq 1\}, \{b_n, n \geq 1\}$ can be constructed as follows

$$\begin{aligned}
 a_n &= n\alpha + an^{\frac{1}{2}} \log(n+1) \\
 b_n &= n\alpha - an^{\frac{1}{2}} \log(n+1),
 \end{aligned}
 \tag{2.30}$$

where a can be found such that

$$\Pr [| S_n - n\alpha | \leq an^{\frac{1}{2}} \log(n+1)] > 1 - \beta/4 \tag{2.31}$$

Next one needs to construct two integer valued random sequences $\{c_n, n \geq 1\}$, $\{d_n, n \geq 1\}$ as follows

$$c(1) = d(1) = 0 \tag{2.32a}$$

of $m \geq 1$, $c(m+1)$ is the least integer n such that

$$N(n) \geq N(c(m)), \tag{2.32b}$$

and $n > c(m)$

and of $m \geq 1$, $d(m+1)$ is the least integer n such that

$$N(n) \leq N(d(m)), \tag{2.32c}$$

and

$$n > d(m)$$

Then there exists two random numbers $\bar{M}$ and $\underline{M}$ where $\bar{M}$ is the least integer n such that

$$\frac{1}{2} \sum_{i=1}^n (1 - g(Z_{N(c(i))}, M(x_{N(c(i))})) \geq a_n \tag{2.33}$$

and $\underline{M}$ is the least integer n such that

$$\frac{1}{2} \sum_{i=1}^n (1 - g(Z_{N(d(i))}, M(x_{Nd(i)}))) \leq b_n \quad (2.34)$$

and

$$\Pr [\underline{M} < \infty] = 1 \quad (2.35)$$

Then

$$\Pr [X_I \geq X_\alpha] \leq \beta/4 \quad (2.36)$$

and

$$\Pr [X_J \leq X_\alpha] \leq \beta/4 \quad (2.37)$$

where $I = N(d(\underline{M}))$ and $J = (c(\underline{M}))$.

In order to show the statements (2.35), (2.36), (2.37) are true, a modified form of a theorem in [10] can be stated as follows.

Theorem 2.1 Let $\beta > 0$ be given and $\{a_n, n \geq 1\}$ be a real number sequence defined as in (2.30). Suppose $\{x_n, n \geq 1\}$ be a strictly increasing sequence of real numbers such that $\lim_{n \rightarrow \infty} x_n = \infty$. Let $\{Y_n, n \geq 1\}$ be a sequence of mutually independent and identically distributed Bernoulli random variables such that $\Pr [Y=1] = \alpha$.
Let N be the least integer n such that $\sum_{i=1}^n Y_i \geq a_n$ with N not existing if $\sum_{i=1}^n Y_i < a_n$ for all $n \geq 1$. Then if $\lim_{x \rightarrow \infty} M(x) > \alpha$, then

$$\Pr [X_n \geq X_\alpha] \geq 1 - \beta/4 \quad \text{and} \quad \Pr [N < \infty] = 1.$$

Let P be the total number of observation needed to get $\bar{M}, \underline{M}$, which can be given by the

$$\text{Maximum } [d(\underline{M}), c(\bar{M})] .$$

Define the sequence $P(i)$ as the number of times the sequence $N(0), N(1), N(2), \dots, N(P)$ takes the value i where i is restricted such that $I \leq i \leq J$.

Also define the sequence $M(i,j)$ as the least integer m such that the sequence $N(0), N(1), \dots, N(m)$ takes the value i exactly j times. Note it follows from [15] that the sequence $\{N(n)\}$ takes all the values i infinitely often as n increases.

Then for i in $I \leq i \leq J$, either

$$\frac{1}{2} \sum_{j=1}^{P(i)} (1 - g(Z_{M(i,j)}, F(x_k))) \geq a_{P(i)} \quad (2.38)$$

or

$$\frac{1}{2} \sum_{j=1}^{P(i)} (1 - g(Z_{M(i,j)}, F(x_i))) \leq b_{P(i)} \quad (2.39)$$

except perhaps for one i .

Let I^* be the greatest i in $I \leq i < J$ satisfying (2.39) and J^* be the least i in $I < i \leq J$ satisfying (2.38)

Then

$$\Pr [X_{I^*} \leq X_\alpha \leq X_{J^*}] \geq 1 - \beta$$

and

$$\Pr [|J^* - I^*| \leq 2] = 1 .$$

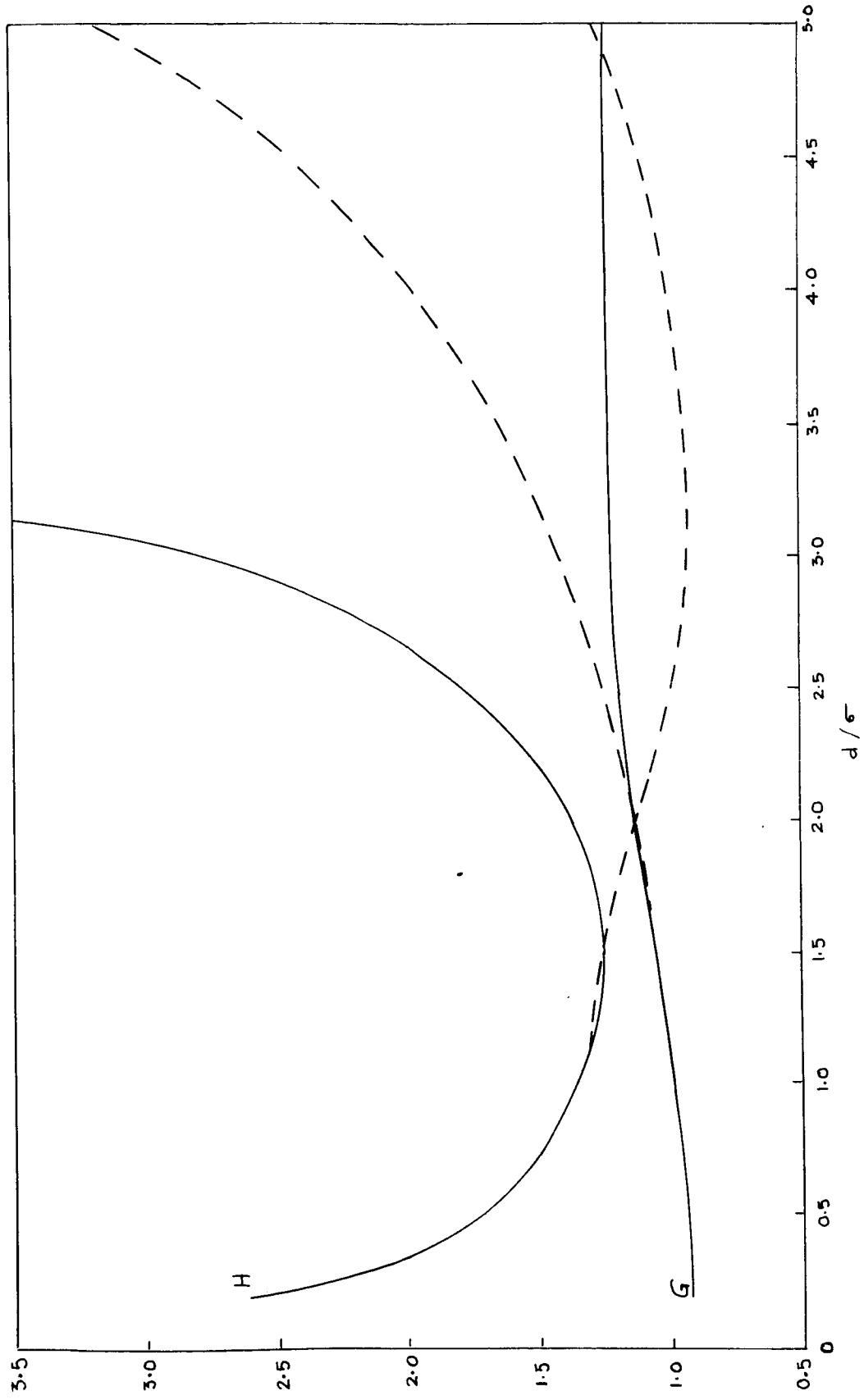


Figure 2-1 G and H Correction Factors for Dixon - Mood Up and Down Test

Chapter III

A NEW METHOD TO OBTAIN CONFIDENCE LIMITS ON

X_α IN UP-AND-DOWN TESTING

The purpose of this chapter is to formulate another method to estimate X_α in a confidence interval estimate such that

- D1. The lower bound, $1 - \beta$, where $0 < 1 - \beta < \beta_{\max} < 1$, of the confidence coefficient is known apriori.
- D2. The upper bound 3δ , for the length of the confidence interval is known apriori.
- D3. A lower bound, C , of $M'(x)$ at $x = x_\alpha$ is known apriori.
- D4. $M(x)$ is an unknown distribution function such that

$$\begin{aligned} \text{a1)} \quad M(x) &= 0, & x &\leq a \\ &= E[Y(x)], & a &\leq x \leq b \\ &= 1, & x &\geq b \end{aligned}$$

where a and b are known apriori and $Y(x)$ defined by (2.1).

$$\text{a2)} \quad M(x_\alpha) = \alpha$$

$$\text{a3)} \quad M'(x) = c' > c > 0$$

for $x_\alpha - \delta/2 < x < x_\alpha + \delta/2$; that is $M(x)$ is linear in a neighborhood of x_α .

$$\begin{aligned} \text{a4)} \quad M'(x) &\text{ exists everywhere except perhaps at} \\ &x = a \text{ or (and) } x = b \end{aligned}$$

Suppose at each level x_i the experimenter takes a sample of size k . Then let

$$\bar{Y}(x_j) = \frac{1}{k} \sum_{i=1}^k Y_i(x_j) / k, \quad (3.1)$$

where Y_i is defined in (2.1).

Corollary (3.1): The x_α which is the solution of $M(x) = \alpha$
is the same as the solution of $E [\bar{Y}(x)] = \alpha$.

Proof: Since the Y_i are from an independent sample

$$\begin{aligned} E [\bar{Y}(x_j)] &= E \left[\frac{1}{k} \sum_{i=1}^k Y_i(x_j) \right] = \frac{1}{k} \left\{ \sum_{i=1}^k E[Y(x_j)] \right\} \\ &= \frac{1}{k} k E[Y(x_j)] = E[Y(x_j)] \end{aligned}$$

and this is true for all x_j .

That is,

$$E [\bar{Y}(x)] = E [Y(x)] = M(x) \quad (3.2)$$

Therefore the value x_α which is the solution of $M(x) = \alpha$ is also the solution of $E [\bar{Y}(x)] = \alpha$.

Thus in order to get a confidence interval for x_α the equation $E [\bar{Y}(x)] = \alpha$ can be statistically solved instead of directly solving the equation $M(x) = \alpha$.

Corollary 3.2: Let $M(x)$ be monotonic non decreasing. Let
 $L_1(x)$ and $L_2(x)$ be the upper and lower tolerance limit functions,
 such that

$$\Pr [L_2(x) < \bar{Y}(x) < L_1(x)] > 1 - 2\beta \quad (3.3)$$

then $L_1(x)$ and $L_2(x)$ are also monotonic increasing.

Proof: We know that

$$\Pr [Y = 1] = M(x) \quad \text{for all } x$$

Therefore

$$\Pr \left[\sum_{i=1}^k Y_i \geq KL_1 \mid M(x_1) \right] = \beta \quad (3.4)$$

If $x_2 > x_1$, then $M(x_2) \geq M(x_1)$ and

therefore

$$\Pr \left[\sum_{i=1}^k Y_i \geq KL_1(x_1) \mid M(x_2) \right] \geq \beta \quad (3.5)$$

Now selecting $KL_1(x_2)$ such that

$$\Pr \left[\sum_{i=1}^k Y_i \geq KL_1(x_2) \mid M(x_2) \right] = \beta$$

(3.5) and (3.6) show that $L_1(x_2) \geq L_1(x)$.

Hence it can be concluded that if $x_2 > x_1$, then $L_1(x_2) \geq L_1(x_1)$ and therefore the tolerance function is monotonic non decreasing.

In a similar way it can be shown that $L_2(x)$ is also monotonic non decreasing. Note that the ranges of $L_1(x)$ and $L_2(x)$ are from zero to unity.

Theorem 3.1. There exists an integer k such that

$$\Pr [L_2(x_\alpha) \leq \bar{Y}(x_\alpha) \leq L_1(x_\alpha)] \geq 1 - 2\beta \quad (3.7)$$

$$\text{where } L_1(x_\alpha) = \alpha + \delta \tan c \quad (3.8)$$

and

$$L_2(x_\alpha) = \alpha - \delta \tan c \quad (3.9)$$

and δ, C are defined in D2 and D3.

Proof: By definition of Y_i , Y_i is a random variable such that $\Pr [Y_i(x) = 1] = M(x)$. That is, $\{Y_i\}$ are independent and identically distributed Bernoulli random variables, and therefore

$k\bar{Y} = \sum_{i=1}^k Y_i$ is distributed as

$$\binom{k}{k\bar{Y}} [M(x)]^{k\bar{Y}} [1 - M(x)]^{k-k\bar{Y}}$$

which is the same as

$$\binom{k}{k\bar{Y}} (\alpha)^{k\bar{Y}} (1 - \alpha)^{k-k\bar{Y}}, \quad (3.10)$$

which is the binomial distribution with parameters k and α , when $x = x_\alpha$. It is important to note that (3.7) can also be written as

$$\Pr [k \cdot L_2(x_\alpha) \leq k\bar{Y}(x_\alpha) \leq kL_1(x_\alpha)] \geq 1 - 2\beta.$$

That is,

$$\Pr [kL_2 \leq \sum_{i=1}^K Y_i \leq kL_1] \geq 1 - 2\beta \quad (3.11)$$

since the distribution of $\sum_{i=1}^K Y_i$ is known. The particular k can be obtained as follows by using binomial probability tables: Select

k such that

$$\frac{[kL_1]}{\sum [kL_2]} \left(\frac{k}{kY} \right) (\alpha)^{kY} (1 - \alpha)^{k-kY} \geq 1 - 2\beta \quad (3.12)$$

Since α and β are known constants in (3.12), k can always be found.

In D4 it is assumed that there exists two known real numbers a and b such that

$$\begin{aligned} Y(x) &= 0 & \text{if } x \leq a \\ &= 1 & \text{if } x \geq b \end{aligned}$$

which implies that

$$\begin{aligned} M(x) &= 0 & \text{if } x \leq a \\ &= 1 & \text{if } x \geq b. \end{aligned}$$

Let the interval (a,b) be partitioned into $(n-1)$ equally spaced abutting intervals where

$$n = \left[\frac{b-a}{\delta} \right] + 1 \quad (3.13)$$

where $[z]$ denotes the largest integer contained in z . Let the n states be defined as follows:

$$S_j = a + (j-1)(b-a)/(n-1), \quad j = 1, 2, \dots, n \quad (3.14)$$

Note that $S_1 = a$ and $S_n = b$.

Theorem 3.2. There exists an i such that the value x_α lies between S_i and S_{i+1} where S_i is defined in (3.14).

Proof: In Fig. (3.1) let P_1 and P_2 be the points where the derivations of $L_1(x)$ and $L_2(x)$ from $M(x)$ are maximum. Let $L'_1(x_\alpha)$ and $L'_2(x_\alpha)$ be the points where the tangent lines at P_1 and P_2 intersect the line $x = x_\alpha$. Note that these tangents will be parallel to $M(x)$ in a neighborhood of x_α , because of the monotonic nature of $M(x)$, $L_1(x)$, and $L_2(x)$ and by the very construction of P_1 and P_2 . Let Q_1 and Q_2 be the points where these tangents intersect the line $y = \alpha$. Let $L_1(x_\alpha)$ and $L_2(x_\alpha)$ be the points where $L_1(x)$ and $L_2(x)$ intersect the $x = x_\alpha$ respectively. Through $L'_1(x_\alpha)$ draw lines making an angle c with the x -axis. Let Q'_1 and Q'_2 be the points where these lines intersect the line $x = x_\alpha$.

From Fig. (3.1) it appears obvious and indeed it is true that

$$L'_2(x_\alpha) \leq L_2(x_\alpha) \quad (3.15)$$

and

$$L'_1(x_\alpha) \geq L_1(x_\alpha) \quad (3.16)$$

Also

$$\frac{M(x_\alpha) - L'_2(x_\alpha)}{\delta_2} = \tan c' \quad (3.17)$$

where c' is the true slope of $M(x)$ which is unknown to the experimenter, and δ_2 is the difference between $M(x_\alpha)$ and Q_2 .

$$\frac{M(x_\alpha) - L_2'(x_\alpha)}{\delta} = \tan c \quad (3.18)$$

From (3.17) and (3.18) one obtains that

$$\frac{M(x_\alpha) - L_2'(x_\alpha)}{\delta_2} \geq \frac{M(x_\alpha) - L_2(x_\alpha)}{\delta}$$

since $\tan c' \geq \tan c$ and since $M(x_\alpha) - L_2'(x_\alpha) \leq M(x) - L_2(x_\alpha)$ it follows that $\delta_2 \leq \delta$. In a similar way it can be shown that $\delta_1 \leq \delta$, where δ_1 is the difference between Q_1 and $M(x_\alpha)$.

Now consider the inequality

$$M(x) - \delta_2 \tan c' \leq L_2(x) \leq M(x) \leq L_1(x) \leq M(x) + \delta_1 \tan c'.$$

By fixing $Y = \alpha$, one can solve the following equations.

$$\begin{aligned} M(x) - \delta_2 \tan c' &= \alpha \\ L_2(x) &= \alpha \\ M(x) &= \alpha \\ L_1(x) &= \alpha \\ M(x) - \delta_1 \tan c' &= \alpha, \end{aligned} \quad (3.19)$$

Because of the monotonicity of the functions in (3.19) their solutions satisfy the following inequality that

$$x_{\alpha} - \delta_1 \leq x(L_1) \leq x_{\alpha} \leq x(L_2) \leq x_{\alpha} + \delta_2$$

since $\delta_2 \leq \delta$ and $\delta_1 \leq \delta$, it follows that

$$| (x_{\alpha} + \delta_2) - (x_{\alpha} - \delta_1) | \leq 2\delta.$$

Therefore the interval $I(x_{\alpha}) = |(x_{\alpha} + \delta_2) - (x_{\alpha} - \delta_1)|$ can cover at the most two states and those two will be S_i and S_{i+1} , which shows that there exists an i such that

$$\Pr [S_i \leq x_{\alpha} \leq S_{i+1}] = 1.$$

The experiment is restricted such that it will be performed at levels $x = S_j$, $j = 1, 2, \dots, n$. We call a trial as an experiment of sample size k at a particular (fixed) level S_j . Before each trial is performed a decision is made to select the experimental level S_j by using the following three rules:

Rule 1. If $\bar{Y}(S_j)$ is greater than α , then the next trial will be performed at the level S_{j-1} .

Rule 2. If $\bar{Y}(S_j)$ is equal to α , then the next trial will be performed at the level S_j .

Rule 3. If $\bar{Y}(S_j)$ is less than α , then the next trial will be performed at the level S_{j+1} . This procedure defines a sequence of experimental levels $\{x_{N(j)}\}$ where $x_{N(1)}$ is the best a priori estimate for x_{α} and $\{N(j)\}$ is defined as

$$\begin{aligned}
N(1) &= 1, \text{ and for } j > 2 \text{ and } j < n \\
N(j+1) &= N(j) + 1, \text{ if } \bar{Y}(x_j) < \alpha \\
N(j+1) &= N(j), \text{ if } \bar{Y}(x_j) = \alpha \\
N(j+1) &= N(j) - 1, \text{ if } \bar{Y}(x_j) > \alpha
\end{aligned} \tag{3.20}$$

Thus, the procedure described by (3.20) defines a random walk with reflecting barriers at the lower and upper limits. On the basis of this fact a theory will be developed to construct the desired confidence interval.

Theorem 3.3. Let P be the transition matrix of the random walk defined by (3.20). Then the elements of P can be defined as follows:

- (1) $P_{1,2} = P_{n,n-1} = 1$
- (2) $P_{i,j} = 0$ for all i and j such that $|i - j| \neq 1$
- (3) $P_{i,i-1} > 0$ and $P_{i,i+1} > 0$
- (4) There exists an $i, i=1, 2, \dots, n$ where $x_\alpha \in (S_i, S_{i+1})$ and real numbers β and $\pi, 0 < \beta < \pi < 1 - \beta < 1$ such that

$$a) \quad 1 = P_{1,2} \geq P_{2,3} \geq \dots \geq P_{i-1,i} \geq 1 - \beta \tag{3.21}$$

$$b) \quad 1 = P_{n,n-1} \geq P_{n-1,n-2} \geq \dots \geq P_{i+2,i+1} \geq 1 - \beta$$

$$c) \quad 0 < P_{2,1} \leq P_{3,2} \leq \dots \leq P_{i-1,i-2} \leq \beta$$

$$d) \quad 0 < P_{n-1,n} \leq P_{n-2,n-1} \leq \dots \leq P_{i+2,i+3} \leq \beta$$

$$e) \quad 1 - \pi \leq P_{i+1,i} \leq 1 - \beta$$

$$f) \quad \beta \leq P_{i+1,i+2} \leq \pi$$

$$g) \quad \pi \leq P_{i,i+1} \leq 1 - \beta$$

$$h) \quad \beta \leq P_{i,i-1} \leq 1 - \pi$$

Proof: We have already mentioned that $L_1(x)$ and $L_2(x)$ are the upper and lower tolerance equations such that

$$\Pr [\bar{Y}(x) \leq L_1(x)] \geq 1 - \beta$$

and

$$\Pr [\bar{Y}(x) \geq L_2(x)] \geq 1 - \beta$$

for all $x \in (a, b)$.

Let $I(x_\alpha)$ be the interval $[x(L_1), x(L_2)]$.

Let $\Pr [\bar{Y} \leq \alpha] = \pi$ and $\Pr [\bar{Y} > \alpha] = 1 - \pi$ given $x = x_\alpha$ (3.22)

where π is a real number in $0 < \pi < 1$. From Fig. (3.1) it is easily seen that

$$\begin{aligned} \Pr [\bar{Y}(x) > \alpha \mid x \in I(x_\alpha) \text{ and } x < x_\alpha] \\ = \Pr [\bar{Y}(x) > L_1(x)] \leq \beta \end{aligned} \quad (3.23)$$

$$\begin{aligned} \Pr [\bar{Y}(x) < \alpha \mid x \in I(x_\alpha) \text{ and } x > x_\alpha] \\ = \Pr [\bar{Y}(x) < L_2(x)] \leq \beta \end{aligned} \quad (3.24)$$

$$\begin{aligned} \Pr [\bar{Y}(x) > \alpha \mid x \in I(x_\alpha) \text{ and } x > x_\alpha] \\ = 1 - \Pr [\bar{Y}(x) > \alpha \mid x \in I(x_\alpha) \text{ and } x < x_\alpha] \geq 1 - \beta \end{aligned} \quad (3.25)$$

$$\begin{aligned} \Pr [\bar{Y}(x) < \alpha \mid x \in I(x_\alpha) \text{ and } x < x_\alpha] \\ = 1 - \Pr [\bar{Y}(x) < \alpha \mid x \in I(x_\alpha) \text{ and } x > x_\alpha] \geq 1 - \beta \end{aligned} \quad (3.26)$$

From Fig. (3.2)

$$\beta \leq \Pr [\bar{Y}(x) > \alpha \mid x \in I(x_\alpha) \text{ and } x < x_\alpha] \leq 1 - \pi \quad (3.27)$$

$$\beta \leq \Pr [\bar{Y}(x) < \alpha \mid x \in I(x_\alpha) \text{ and } x > x_\alpha] \leq \pi \quad (3.28)$$

$$1 - \pi \leq \Pr [\bar{Y}(x) < \alpha \mid x \in I(x_\alpha) \text{ and } x > x_\alpha] \leq 1 - \beta \quad (3.29)$$

$$\pi \leq \Pr [\bar{Y}(x) < \alpha \mid x \in I(x_\alpha) \text{ and } x < x_\alpha] \leq 1 - \beta \quad (3.30)$$

According to the rules 1, 2 and 3, when the experiment is at a state j at the n th trial, then it will be either at $j-1$ th or j th or $j+1$ st state at the $(n+1)$ st trial. If we take the sample size k such that it will satisfy (3.11) and $\bar{Y} \neq \alpha$, then

$$\Pr \left[\begin{array}{c} \text{the process will be at } S_j \\ \text{at the } (n+1)\text{th trial} \end{array} \middle/ \begin{array}{c} \text{it was at } S_j \text{ at} \\ \text{the } n\text{th trial} \end{array} \right] = 0.$$

Let P_{ij} be the probability that the process will be at the state j given it was at the state i , in one trial. That is,

$$\Pr \left[\begin{array}{c} \text{the process will be at } S_j \\ \text{at the } (n+1)\text{th trial} \end{array} \middle/ \begin{array}{c} \text{it was at } S_i \text{ at} \\ \text{the } n\text{th trial} \end{array} \right] = P_{ij}.$$

Since $Y_i(x) = 0$ when $x \leq a$, $\bar{Y}(x) = 0$ when $x \leq a$.

Therefore, $\bar{Y}_i(x \leq a) = 0 < \alpha$. So by rule (3), $P_{12} = 1$.

Similarly since $Y(x) = 1$ for all $x \geq b$, $\bar{Y}(x) = 1$ for all $x \geq b$. Therefore $\Pr [Y(x) > b] = 1 > \alpha$. So by rule (1), $P_{n,n-1} = 1$, which proves (1) in (3.21).

Since the process will be either at S_{j-1} or S_{j+1} after one trial given that it was at S_j before the trial, the

$$P_{i,j} = 0 \quad \text{if} \quad |i-j| \neq 1 \quad \text{which is (2) in (3.21).}$$

By the definition of $P_{i,j}$, $P_{i,i-1} > 0$ and $P_{i,i+1} > 0$.

Rewriting (3.23) . . . (3.26) in terms of P_{ij} one gets

$$0 < P_{j,j-1} \leq \beta \quad \text{where} \quad j = 2, 3, \dots, i-1$$

$$0 < P_{j,j+1} \leq \beta \quad \text{where} \quad j = i+2, i+3, \dots, n-1$$

$$1 - \beta \leq P_{j,j-1} < 1 \quad \text{where} \quad j = i+2, i+3, \dots, n-1$$

$$1 - \beta \leq P_{j,j+1} < 1 \quad \text{where} \quad j = 2, 3, \dots, i-1.$$

Thus the elements of P satisfy (3.21) and the existence of i , $i = 1, 2, \dots, n$ such that $x_\alpha \in (S_i, S_{i+1})$ is shown in theorem (3.2)

In order to get the required number of trials to obtain the desired confidence interval consider the following definitions and theorems [19].

Definition 3.1: A finite Markov chain is a stochastic process which moves through a finite number of states, and for which the probability of entering a certain state after a single step depends only on the last state occupied.

Definition 3.2: An ergodic set of states is a set in which every state can be reached from every other state, and the ergodic set can not be left once it is entered.

Definition 3.3: An ergodic chain is one whose states form a single ergodic set; or --equivalently-- a chain in which it is possible to go from any state to every other state.

Definition 3.4: $S_0, S_1, S_2 \dots$ is a divergent sequence and let

$$t_n = (1/n) \sum_{i=0}^{n-1} S_i$$

if the sequence $t_1, t_2, \dots$ converges to a limit t , then we say that the sequence $\{S_n, n \geq 0\}$ is Cesaro-summable to t .

Definition 3.5: A cyclic chain is an ergodic chain in which each state can be entered at certain periodic intervals.

Definition 3.6: A regular chain is an ergodic chain that is not cyclic.

Theorem 3.4: If P is an ergodic transition matrix,

- a) the sequence $\{P^n, n \geq 0\}$ (Note that n is a positive integer) is Cesaro-summable to A
- b) Each row of A is the same probability vector

$$\gamma = (\gamma_1, \gamma_2 \dots \gamma_n).$$

c) the vector γ is unique fixed probability vector of P .

d) $PA = AP = A$ (the matrix A is called the Cesaro-summable matrix).

Let $v_j^{(n)}$ is the fraction of times in the first n steps that the process will move to the state S_j . The law of large numbers for regular Markov chains can be stated as follows.

Theorem 3.5: Consider a regular Markov chain with limiting vector

$\gamma = (\gamma_1, \gamma_2, \dots, \gamma_n)$. For any initial vector Π ,

$$E\Pi [v_n^{(n)}] \rightarrow \gamma$$

and for every $\epsilon > 0$

$$\Pr [| v_j^{(n)} - \gamma_j | > \epsilon] \rightarrow 0$$

as n tends to infinity.

The transition matrix P defined in (3.21) has the necessary features such that theorems (3.4), (3.5) can be applied. Therefore using theorem (3.4), the following theorem can be stated.

Theorem 3.6: Let P be the transition matrix defined in (3.21),

and let the Cesaro-summable matrix be $A = J\gamma$ where $\gamma = (\gamma_1, \gamma_2, \dots, \gamma_n)$ and $J^T = (1, 1, \dots, 1)$ are $1 \times n$ vectors such that

$$\gamma P = \gamma \tag{3.31}$$

$$\gamma J = 1 \tag{3.32}$$

then for all $0 < \beta < \min \left\{ \frac{(16 + \pi^2)^{\frac{1}{2}} - (4 - \pi)}{4\pi}, \pi \right\}$

$$\sum_{j=i-1}^{i+2} \gamma_j \geq 1 - \beta. \quad (3.33)$$

Proof: On solving (3.31) in terms of γ_1 one obtains $\gamma_1 = \gamma_1$ and for $j = 2, 3, \dots, n$

$$\gamma_j = \gamma_1 P_{12} P_{23} \dots P_{j-1,j} / P_{21} P_{32} \dots P_{j,j-1} \quad (3.34)$$

using (3.32) $\sum_{j=1}^n \gamma_j = 1.$

The aim is to find $\sum_{j=i-1}^{i+2} \gamma_j.$

Using (3.34)

$$\begin{aligned} \sum_{j=i-1}^{i+2} \gamma_j &= \frac{\gamma_1 P_{12} \dots P_{i-2,i-1}}{P_{21} \dots P_{i-1,i-2}} + \frac{\gamma_1 P_{12} \dots P_{i-1,i}}{P_{21} \dots P_{i,i-1}} + \frac{\gamma_1 P_{12} \dots P_{i,i+1}}{P_{21} \dots P_{i+1,i}} \\ &+ \frac{\gamma_1 P_{12} \dots P_{i+1,i+2}}{P_{21} \dots P_{i+2,i+1}} \end{aligned}$$

In order to eliminate γ_1 which is an unknown quantity, consider the following ratio

$$r = \frac{\sum_{j=i-1}^{i+2} \gamma_j}{\left(1 - \sum_{j=i-1}^{i+2} \gamma_j\right)}$$

$\sum_{j=i-1}^{i+2} \gamma_j$ can be written as

$$\frac{\gamma_1 P_{12} \cdots P_{i-1,i}}{P_{21} \cdots P_{i-1,i-2}} \left(\frac{1}{P_{i-1,i}} + \frac{1}{P_{i,i-1}} + \frac{P_{i,i+1}}{P_{i,i-1} P_{i+1,i}} + \frac{P_{i,i+1} P_{i+1,i+2}}{P_{i,i-1} P_{i+1,i} P_{i+2,i+1}} \right) \quad (3.35)$$

Using (3.33)

$$1 - \sum_{j=i-1}^{i+2} \gamma_j = \gamma_1 + \gamma_2 + \cdots + \gamma_{i-2} + \gamma_{i-3} + \cdots + \gamma_m$$

which can be put in the form

$$\frac{\gamma_1 P_{12} \cdots P_{i-1,i}}{P_{21} \cdots P_{i-1,i-2}} \left[\left(\frac{1}{P_{i-1,i}} + \frac{P_{i-1,i-2} \cdots P_{21}}{P_{i-2,i-1} \cdots P_{12}} + \cdots + \frac{P_{i-1,i-2}}{P_{i-2,i-1}} \right) + \right. \\ \left. \left(\frac{P_{i,i+1} \cdots P_{i+2,i+3}}{P_{i-1,i} \cdots P_{i+3,i+2}} \right) \left(1 + \frac{P_{i+3,i+4}}{P_{i+4,i+3}} + \cdots + \frac{P_{i+3,i+4} \cdots P_{n-1,n}}{P_{i+4,i+3} \cdots P_{n,n-1}} \right) \right] \quad (3.36)$$

r is the ratio of (3.35) and (3.36) and therefore is free of γ_1 . Hence r can be written as

$$\frac{\left[\frac{1}{P_{i-1,i}} + \frac{1}{P_{i,i-1}} + \frac{P_{i,i+1}}{P_{i-1,i} P_{i+1,i}} + \frac{P_{i,i+1} P_{i+1,i+2}}{P_{i-1,i} P_{i+1,i} P_{i+2,i+1}} \right]}{\frac{1}{P_{i-1,i}} \left[\frac{P_{i-1,i-2} \cdots P_{21}}{P_{i-2,i-1} \cdots P_{12}} + \cdots + \frac{P_{i-1,i-2}}{P_{i-2,i-1}} \right] + \left[\frac{P_{i,i+1} \cdots P_{i+2,i+3}}{P_{i-1,i} \cdots P_{i+3,i+2}} \right]} \\ \left[1 + \frac{P_{i+3,i+4}}{P_{i+4,i+3}} + \cdots + \frac{P_{i+3,i+4} \cdots P_{n-1,n}}{P_{i+4,i+3} \cdots P_{n,n-1}} \right] \quad (3.37)$$

We want to get a lower bound for $\sum_{j=i-1}^{i+2} \gamma_j$. We will get this lower bound as a function of r^* , where r^* is such that $r > r^*$.

$$r = \frac{\sum_{j=i-1}^{i+2} \gamma_j}{1 - \sum_{j=i-1}^{i+2} \gamma_j}$$

which implies that

$$\sum_{j=i-1}^{i+2} \gamma_j = \frac{1}{1 + \frac{1}{r}} \geq \frac{1}{1 + \frac{1}{r^*}}$$

Since we do not know the values of P_{ij} for all $|i-j| = 1$ except $P_{12} = P_{n,n-1} = 1$, it becomes necessary to replace these elements by appropriate bounds. Therefore in order to get a lower bound for r , we will try to minimize the numerator and maximize the denominator of (3.37) using (3.21). Consider first the numerator:

$$\begin{aligned} \text{Numerator} &\geq \frac{1}{1} + \frac{1}{1-\pi} + \frac{\pi}{(1-\pi)(1-\beta)} + \frac{\pi\beta}{(1-\pi)(1-\beta)1} \\ &\geq 1 + \frac{1}{1-\pi} + \frac{\pi}{1-\pi} + \frac{\pi\beta}{1-\pi} \quad \text{since} \quad \frac{1}{1-\beta} > 1 \\ &\geq \frac{1+(1-\pi) + \pi + \pi\beta}{1-\pi} = \frac{2+\pi\beta}{1-\pi} \end{aligned}$$

Similarly, on considering the denominator it follows that

$$\begin{aligned}
 \text{Denominator} &\leq \frac{1}{1-\beta} \left[\left(\frac{\beta}{1-\beta}\right)^{i-2} + \dots + \left(\frac{\beta}{1-\beta}\right) \right] \\
 &+ \frac{(1-\beta) \cdot \pi \cdot \beta}{(1-\beta)(1-\pi)(1-\beta)(1-\beta)} \left[1 + \left(\frac{\beta}{1-\beta}\right) + \dots + \left(\frac{\beta}{1-\beta}\right)^{N-i-2} \right] \\
 &\leq \frac{1}{1-\beta} \cdot \frac{\beta}{1-\beta} \frac{1 - \left(\frac{\beta}{1-\beta}\right)^{i-2}}{1 - \left(\frac{\beta}{1-\beta}\right)} + \frac{\pi\beta}{(1-\pi)(1-\beta)^2} \frac{1 - \left(\frac{\beta}{1-\beta}\right)^{N-i-1}}{1 - \frac{\beta}{1-\beta}} \\
 &\leq \frac{\beta}{(1-\beta)} \frac{1-\beta}{1-2\beta} + \frac{\pi\beta}{(1-\pi)(1-\beta)(1-2\beta)}
 \end{aligned}$$

since $\frac{\beta}{1-\beta} < 1$ and therefore $\left(\frac{\beta}{1-\beta}\right)^x < 1$ where x is a positive integer. That is

$$\begin{aligned}
 \text{Denominator} &\leq \frac{\beta}{(1-\beta)(1-2\beta)} + \frac{\pi\beta}{(1-\beta)(1-2\beta)(1-\pi)} \\
 &\leq \frac{\beta}{(1-\beta)(1-2\beta)(1-\pi)} [1 - \pi + \pi] = \frac{\beta}{(1-\beta)(1-2\beta)(1-\pi)}
 \end{aligned}$$

Therefore

$$\begin{aligned}
 r &> \frac{2 + \pi\beta}{(1-\pi)} \frac{(1-\beta)(1-2\beta)(1-\pi)}{\beta} \\
 &\geq (2 + \pi\beta) (1 - 2\beta) (1 - \beta) / \beta
 \end{aligned}$$

We want to find the range of β such that

$$\sum_{j=i-1}^{i+2} \gamma_j \geq 1 - \beta \quad (3.38)$$

which implies that $r \geq \frac{1 - \beta}{\beta}$. That is,

$$\frac{(2 + \pi\beta)(1 - 2\beta)(1 - \beta)}{\beta} \geq \frac{1 - \beta}{\beta}$$

or

$$(2 + \pi\beta)(1 - 2\beta) - 1 \geq 0$$

or

$$-2\pi\beta^2 - (4 - \pi)\beta + 1 \geq 0.$$

Solving the above quadratic equation for β one obtains that

$$\beta \leq \frac{(4 - \pi) \pm \sqrt{(4 - \pi)^2 + 8\pi}}{-4\pi}$$

But β is positive; therefore

$$\beta \leq \frac{(4 - \pi) - \sqrt{16 + \pi^2}}{-4\pi} = \frac{\sqrt{16 + \pi^2} - (4 - \pi)}{4\pi};$$

and it is obvious that β is real. Since we want $\beta < \pi$, the

required range of β is $0 < \beta < \min \left[\frac{(\sqrt{16 + \pi^2})^{\frac{1}{2}} - (4 - \pi)}{4\pi}, \pi \right]$ (3.39)

The bounds on β will not affect the usual selection of $\beta = 0.05$ or $\beta = 0.10$. Note that the lower bound for $\sum_{j=i-1}^{i+2} \gamma_j$ as formulated by (3.38) is independent of the size of the transition matrix, that is the number of states and the location of the i th state or the location of x_α . Therefore, if the value of β is such that β satisfies (3.39), Theorem 3.5 assures that there exists an M' such that

$$\sum_{j=i-1}^{i+2} \gamma_{kj}^M > 1 - \beta$$

for all k and $M > M'$, where γ_{kj}^M is the k th row, j th column element of P^M . That is, the proportion of times, in the first M trials that the process will move to at least one of the states $S_{i-1}, S_i, S_{i+1}, S_{i+2}$, given $x_\alpha \in (S_i, S_{i+1})$ will be at least $(1 - \beta)$. This can be interpreted as follows: If

$$I = [S_{i-1}, \dots, S_{i+2}]$$

is a random interval and given $S_i \leq x_\alpha \leq S_{i+1}$ then,

$$\Pr [I \text{ covers } x_\alpha \mid x_\alpha \in (S_i, S_{i+1})] \geq 1 - \beta.$$

It is important to note that the interval I can then be constructed in such a way to obtain the desired confidence interval. Since we

choose $|S_{i+1} - S_i| \leq \delta$, $|S_{i+2} - S_{i-1}| \leq 3\delta$ is the desired bound for the length of the confidence interval.

Let M_i be the power of P such that

$$\sum_{j=i-1}^{i+2} \gamma_{mj}^{M_i} > 1 - \beta \quad \text{for all } m = 1, 2, \dots, n.$$

By repeating this until i exhausts its range we get a sequence of numbers $M_1, M_2, \dots, M_{n-1}$.

Let

$$M = \max [M_1, M_2, \dots, M_{n-1}] \quad (3.40)$$

Note that M will give the desired number of trials. Therefore the total number of samples required is $N = KM$. Select an initial estimate for x_α and perform M trials of K observations sequentially by Rules 1, 2, and 3. The experiment ceases at the state S_j (say). Then the desired confidence interval for x_α is given by

$$\Pr [S_{j-1} \leq x_\alpha \leq S_{j+2}] \geq 1 - \beta$$

and

$$|S_{j+2} - S_{j-1}| \leq 3\delta.$$

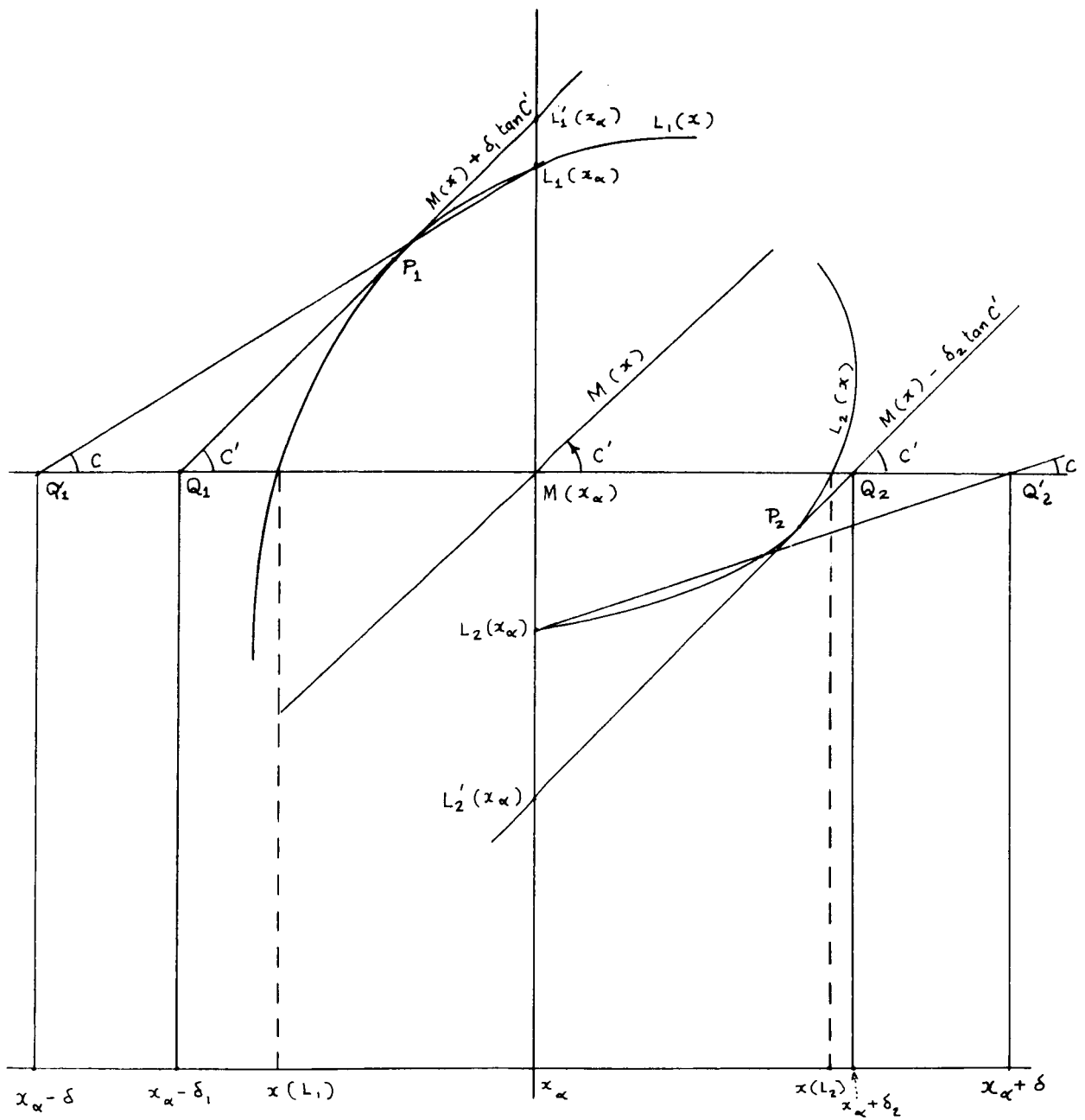


Figure 3.1 The behavior of $M(x)$ in the neighborhood of x_α

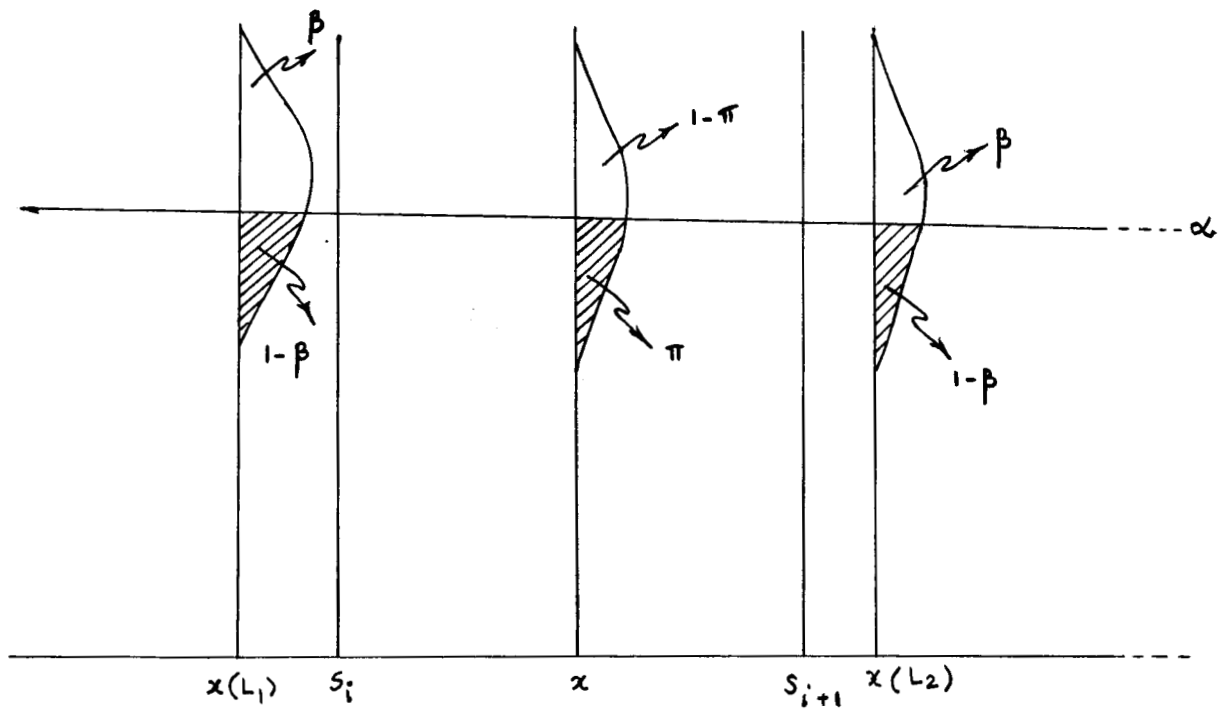


Figure 3.2 The behavior of the distribution of $\bar{Y}(x)$ when $x \in I(x_\alpha)$. Shaded area represents the probability that $\bar{Y}(x)$ is less than α .

Chapter IV
EXPERIMENTAL PROCEDURE TO OBTAIN SAMPLE SIZES IN
VARIOUS TECHNIQUES

Due to the number of uncontrollable parameters involved, perhaps the most practical means available at this time to study the sample size required to get the desired confidence interval, is a "Monte-Carlo Procedure".

"Monte-Carlo Procedures" are often useful in many probabilistic problems. Suppose that we want to study the mortality rate of a given population of insects. One can set up a model for this by comparing random numbers with the response function. Using the random numbers such that they will match with the required statistics, one can generate a random sample from the population which can be analyzed as if it were the data collected in the laboratory. In this way one can obtain the required data without actually performing the experiment, which will be consistent with the real data. There are several ways to generate the desired random numbers [16].

Recall that $M(x)$ is the response function, α a constant such that $0 < \alpha < 1$, β the required confidence coefficient and Δ the desired length of the confidence interval. We considered $\Delta = .2$ and $\beta = .10$.

In practice the form of $M(x)$ is unknown to the experimenter, but it is necessary to define the form of $M(x)$ to perform the sampling scheme.

In this study 6 forms of $M(x)$ were selected arbitrarily, of which two are cumulative normal distributions. These are:

$$\begin{aligned}
 &= 0 && x < 0 \\
 &= 4x^2 && 0 < x < \frac{1}{4} \\
 M_1(x) &= 1 - \frac{4}{3} (1 - x)^2 && \frac{1}{4} < x < 1 \\
 &= 1 && x > 1
 \end{aligned}$$

$$\begin{aligned}
 &= 0 && x < 0 \\
 &= 2x^2 && 0 < x < \frac{1}{2} \\
 M_2(x) &= 1 - 2(1 - x)^2 && \frac{1}{2} < x < 1 \\
 &= 1 && x > 1
 \end{aligned}$$

$$\begin{aligned}
 &= 0 && x < 0 \\
 &= \frac{4}{3} x^2 && 0 < x < \frac{3}{4} \\
 M_3(x) &= 1 - 4(1 - x)^2 && \frac{3}{4} < x < 1 \\
 &= 1 && x > 1
 \end{aligned}$$

$$\begin{aligned}
 &= 0 && x < 0 \\
 &= 3x^2 && 0 < x < \frac{1}{3} \\
 M_4(x) &= 1 - \frac{3}{2}(1 - x)^2 && \frac{1}{3} < x < 1 \\
 &= 1 && x > 1
 \end{aligned}$$

$$M_5(x) = \int_{-\infty}^x \frac{1}{(2\pi)^{\frac{1}{2}}} (\cdot 2) \exp \left\{ -\frac{1}{2} (t/\cdot 2)^2 \right\} dt \quad -\infty < x < \infty$$

$$M_6(x) = \int_{-\infty}^x \frac{1}{(2\pi)^{\frac{1}{2}}} (\cdot 3) \exp \left\{ -\frac{1}{2} (t/\cdot 3)^2 \right\} dt \quad -\infty < x < \infty$$

Sketches of these $M(x)$'s are shown in Fig. 4.1 and Fig. 4.2. The values of x_α for $\alpha = .05, .1, .5, .9, .95$ are given in Table I.

As the procedures to obtain the sample sizes vary for each technique, the procedures are presented separately.

The general procedure is briefly as follows:

- (1) Generate a sequence of random numbers $[Z_i, i = 1, 2, \dots]$ from a uniform distribution on $(0, 1)$.
- (2) Compute $M(x_i)$ given the stimulus level x_i .
- (3) In order to decide whether there is a response or a non-response at the level x_i compare Z_i with $M(x_i)$.
 - (a) If $M(x_i) < Z_i$, consider that there is a nonresponse
 - (b) If $M(x_i) > Z_i$, consider that there is a response.

Spearman-Kärber Method.

In this method at each level the experimenter is required to take n observations. Choose an initial level x_i . Compare $M(x_i)$ with n random numbers $Z_j, j = 1, 2, \dots, n$. Let r_i be the number of responses, and we know that $p_i = r_i/n$. If $p_i > 0$, this procedure is continued at levels below x_i until we get $p_i = 0$, and at levels above x_i until we get $p_i = 1$. We let the level at which $p_i = 0$ be x_1 , the lowest level of the experiment and the level at which $p_i = 1$ as x_k , the highest level of experiment. An example of this is given in Fig. 4.3. The confidence interval for the 50 percent point is constructed for two different values

of d ; (1) $d < 2\sigma$ and (2) $d > 2\sigma$ and for values of n ranging from 3 to 6. Each experiment is repeated ten times. The average of the desired sample sizes and the range of sample values are given in Table II. They are obtained by linearly interpolating between the sample sizes which give the width of the confidence interval slightly larger than 0.2 and slightly smaller than 0.2.

Dixon-Mood Method.

Choose an arbitrary level of experimentation. Let it be x_0 . If $Z_1 > M(x_0)$ take a level above x_0 , that is x_{+1} , and if $Z_1 < M(x_0)$ take a level below x_0 , that is x_{-1} . This procedure can be continued for any number of trials. The confidence intervals are constructed for sample sizes ranging from 20, 30, . . . for two different x_0 (1) near the fifty-percent point and (2) away from the fifty-percent point; and for two different values of d (1) $d < 2\sigma$ and (2) $d > 2\sigma$. Each experiment is repeated 5 times and the average of those values and the range of sample values are given in Table III. These numbers are also obtained by linear interpolation as described in the Spearman-Kärber method.

Farrell's Method.

In this method we choose an arbitrary level x_0 . By definition $N(1) = 0$. If $Z_1 > M(x_{N(1)})$, then $N(2) = N(1) + 1$ and if $Z_1 < M(x_{N(1)})$, then $N(2) = N(1) - 1$. In general the sequence $N(n)$ is constructed as follows. $N(i + 1) = N(i) \pm 1$ according as $Z_i > M(x_{N(i)})$ or $Z_i < M(x_{N(i)})$. The sequences $\{a_n, n > 1\}, \{b_n, n > 1\}, \{c_n, n > 1\}, \{d_n, n > 1\}$ are constructed as described in Chapter II.

The rest of the computation is the same as it is described in Chapter II. This procedure is repeated for various sets of random numbers and the empirical results are given in Table VI. In all these experiments x_0 is taken as the midpoint of the range. That is x_0 is taken as $\frac{1}{2}(b - a)$ where a and b are such that $M(x) = 1$ for $x > b$ and $M(x) = 0$ for $x < a$. In the case of a normal curve a and b are taken as -4σ , 4σ , respectively.

For the particular procedure given by the author, we compute M using (3.39) and k using theorem (3.1). Choosing an arbitrary state S_j , compare Z_j , $j = 1, 2, \dots, k$ with $M(x_i)$.

$$\begin{aligned} Y_j(x_i) &= 0 && \text{if } Z_j > M(x_i) \\ Y_j(x_i) &= 1 && \text{if } Z_j < M(x_i) \end{aligned}$$

Compute

$$\bar{Y}(x_i) = \sum_{j=1}^k Y_j(x_i). \quad \text{If } \bar{Y} < \alpha,$$

then the next trial will be made at the level x_{i+1} . If $\bar{Y} > \alpha$, then the next trial will be made at x_{i-1} . This procedure is repeated M times. If S_j is the level at the end of M trials, then (S_{j-1}, S_{j+2}) is the required confidence interval. These results are given in Table VII.

Parts of the above described calculations were performed with the help of the computer, CDC 1604, The University of Texas Computation Center.

TABLE I

x_{α} FOR VARIOUS α AND FOR DIFFERENT $M(x)$

$M \quad \alpha$	.05	.10	.50	.90	.95
M1	.1118	.15811	.38713	.72614	.80635
M2	.15811	.22361	.5000	.7764	.84189
M3	.19365	.27386	.61237	.84189	.8882
M4	.12884	.18248	.4227	.7419	.8175
M5	-.3290	-.25632	0	.25632	.3290
M6	-.4935	-.38448	0	.38448	.4935

TABLE II

Averages of 10 Different Trials for Spearman-Kärber's Method, Their Minimum and Maximum

$\alpha = .50$

Confidence level = .10

Width of the interval = .20

	M_1			M_2			M_3			M_4			M_5			M_6		
	min.	aver.	max.	min.	aver.	max.	min.	aver.	max.	min.	aver.	max.	min.	aver.	max.	min.	aver.	max.
$d < 25$	9	18	23	16	21	43	11	24	42	21	30	41	13	18	31	11	28	57
$d \geq 25$	27	30	43	51	50	54	42	46	47	43	51	53	26	33	54	111	123	151

TABLE III

Averages of 5 Different Trials for Dixon-Mood's Method and Their Maximum and Minimum

Confidence level = .10

Width of the interval = .20

[1 = $d < 25$, x_0 near 50% point, 2 = $d < 25$, x_0 far from the 50% point, 3 = $d > 25$]

$\frac{x}{n}$	$\alpha = .05$			$\alpha = .10$			$\alpha = .50$			$\alpha = .90$			$\alpha = .95$			
	min.	aver.	max.	min.	aver.	max.	min.	aver.	max.	min.	aver.	max.	min.	aver.	max.	
M_1	1	42	54	95	31	37	62	7	11	18	31	37	42	42	34	95
	2	38	61	128	26	52	90	7	14	21	26	52	90	38	61	128
	3	69	80	105	47	62	82	7	17	19	47	62	82	69	80	105
M_2	1	44	60	102	27	51	96	7	16	18	27	51	96	44	60	102
	2	36	72	93	24	53	62	7	21	33	24	53	62	36	72	93
	3	78	96	107	38	53	68	8	18	22	38	53	68	78	96	107
M_3	1	43	36	76	32	37	65	5	12	16	32	37	65	43	56	76
	2	43	61	87	29	46	50	8	15	26	29	46	50	43	61	87
	3	55	75	85	37	54	58	12	18	20	37	56	58	55	75	85
M_4	1	42	67	86	33	41	61	9	15	17	33	41	61	42	67	86
	2	44	68	98	21	55	69	16	21	28	21	55	69	44	68	98
	3	56	77	83	37	51	59	11	15	17	37	51	59	56	77	83
M_5	1	19	39	53	16	35	42	7	15	21	16	35	42	19	39	53
	2	24	48	62	20	37	51	7	17	22	20	37	51	24	48	62
	3	49	81	91	23	58	69	9	15	30	23	58	69	49	81	91
M_6	1	61	81	107	41	70	92	13	20	31	41	70	92	61	81	107
	2	78	99	134	52	71	102	15	28	42	52	71	102	78	99	134
	3	96	146	192	65	94	117	18	37	46	65	94	117	96	146	192

TABLE IV

NUMBER OF TRIALS WHICH DOES NOT COVER
THE REQUIRED VALUE IN THE CASE OF
NON NORMAL DISTRIBUTIONS

		$\alpha=.05$	$\alpha=.10$	$\alpha=.50$	$\alpha=.90$	$\alpha=.95$
M1	$d < 2\sigma$ x_0 near the 50% pt.	2	0	0	3	5
	$d < 2\sigma$ x_0 far from the 50% pt.	2	1	0	1	2
	$d \geq 2\sigma$	0	3	0	0	0
M2	$d < 2\sigma$; x_0 near the 50% pt.	0	0	0	1	0
	$d < 2\sigma$; x_0 far from the 50% pt.	1	4	0	0	1
	$d \geq 2\sigma$	0	0	0	2	1
M3	$d < 2\sigma$ x_0 near the 50% pt.	3	4	0	1	1
	$d < 2\sigma$ x_0 far from the 50% pt.	1	1	0	1	1
	$d \geq 2\sigma$	2	2	0	1	2
M4	$d < 2\sigma$; x_0 near the 50% pt.	0	2	0	0	3
	$d < 2\sigma$; x_0 far from the 50% pt.	2	4	0	3	3
	$d \geq 2\sigma$	0	0	0	0	0

TABLE V

AVERAGE NUMBER OF TRIALS WHICH DOES NOT COVER
THE REQUIRED VALUE IN THE CASE OF
NON NORMAL DISTRIBUTIONS

M \ α	.05	.10	.50	.90	.95
M1	.267	.267	0	.267	.467
M2	.067	.267	0	.200	.133
M3	.534	.467	0	.200	.267
M4	.133	.267	0	.133	.267

TABLE VI

Farrell's Method - Averages of 5 Different Trials and Their Maximum and Minimum

	$\alpha = .05$			$\alpha = .10$			$\alpha = .50$			$\alpha = .90$			$\alpha = .95$		
	min.	aver.	max	min	aver.	max.	min.	aver	max	min	aver	max	min	aver.	max
M_1	21462	23771	28014	3250	5055	6328	37	60	132	1562	2812	6371	12150	14782	19363
M_2	13561	17791	21759	596	2456	6687	28	25	142	5775	7574	12617	13357	20712	24785
M_3	28866	34614	42512	5124	16282	27156	23	42	63	1037	4380	8788	8187	15659	20703
M_4	10634	16882	23677	2623	3308	5200	37	33	48	2646	6772	14329	19077	24416	29620
M_5	22308	31998	36730	276	2770	5467	16	26	41	561	2059	6612	30334	34551	42487
M_6	27934	39524	49579	3034	4767	9202	32	37	43	2003	4697	8304	29357	35835	46867

TABLE VII

FIXED SAMPLE SIZES FOR THE TATIKONDA METHOD

Confidence Level = .10 Width of the Interval = .20

M \ α	.05	.10	.50	.90	.95
M1	473	371	473	1100	1100
M2	840	720	330	720	840
M3	840	720	600	360	330
M4	600	480	363	990	990
M5	1298	1102	926	1102	1298
M6	2600	2496	1664	2496	2600

TABLE VIII

A TABLE TO COMPARE MAXIMUM SAMPLE SIZES OF VARIOUS METHODS

M(x)	Method	$\alpha=.05$	$\alpha=.10$	$\alpha=.50$	$\alpha=.90$	$\alpha=.95$	No.
M1	Farrell	23771	5055	60	2812	14782	10
	Tatikonda	473	371	437	1100	1100	Fixed
	Dixon-Mood	80	62	17	62	80	5
	Spearman-Kärber -		--	30	--	--	10
M2	Farrell	17791	2456	25	7574	20712	
	Tatikonda	840	720	330	720	840	
	Dixon-Mood	96	53	21	53	96	
	Spearman-Kär.	--	--	50	--	--	
M3	Farrell	34614	16282	42	4380	15659	
	Tatikonda	840	720	600	360	330	
	Dixon-Mood	75	54	18	54	75	
	Spearman-Kär.	--	--	46	--	--	
M4	Farrell	16882	3308	33	6772	24416	
	Tatikonda	600	480	363	990	990	
	Dixon-Mood	77	51	21	51	77	
	Spearman-Kär.	--	--	51	--	--	
M5	Farrell	31998	2770	26	2059	34551	
	Tatikonda	1298	1102	926	1102	1298	
	Dixon-Mood	81	51	17	51	81	
	Spearman-Kär.	--	--	33	--	--	
M6	Farrell	39524	4767	37	4607	35835	
	Tatikonda	2600	2296	1664	2696	2600	
	Dixon-Mood	146	94	37	94	146	
	Spearman-Kär.	--	--	123	--	--	

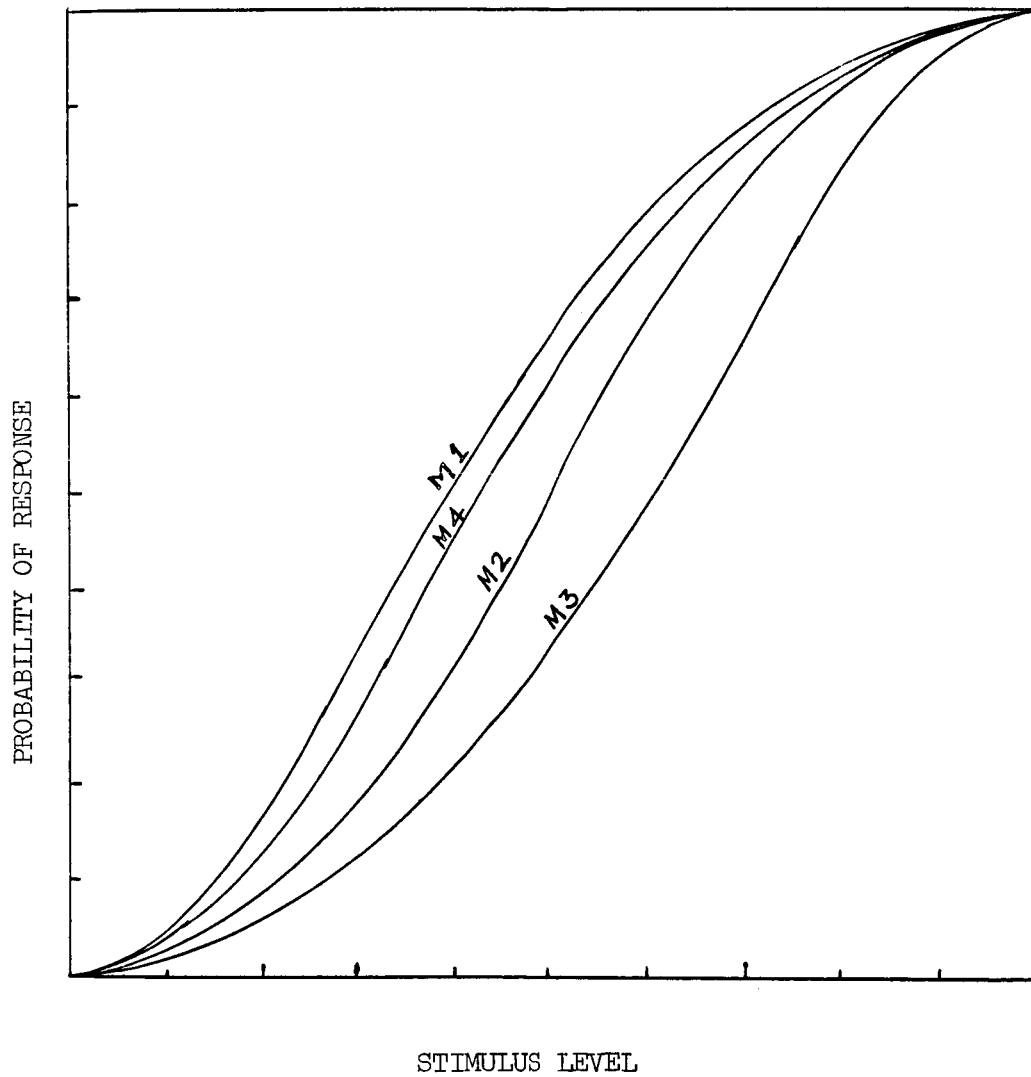


Figure 4.1 Non-normal Response Curves

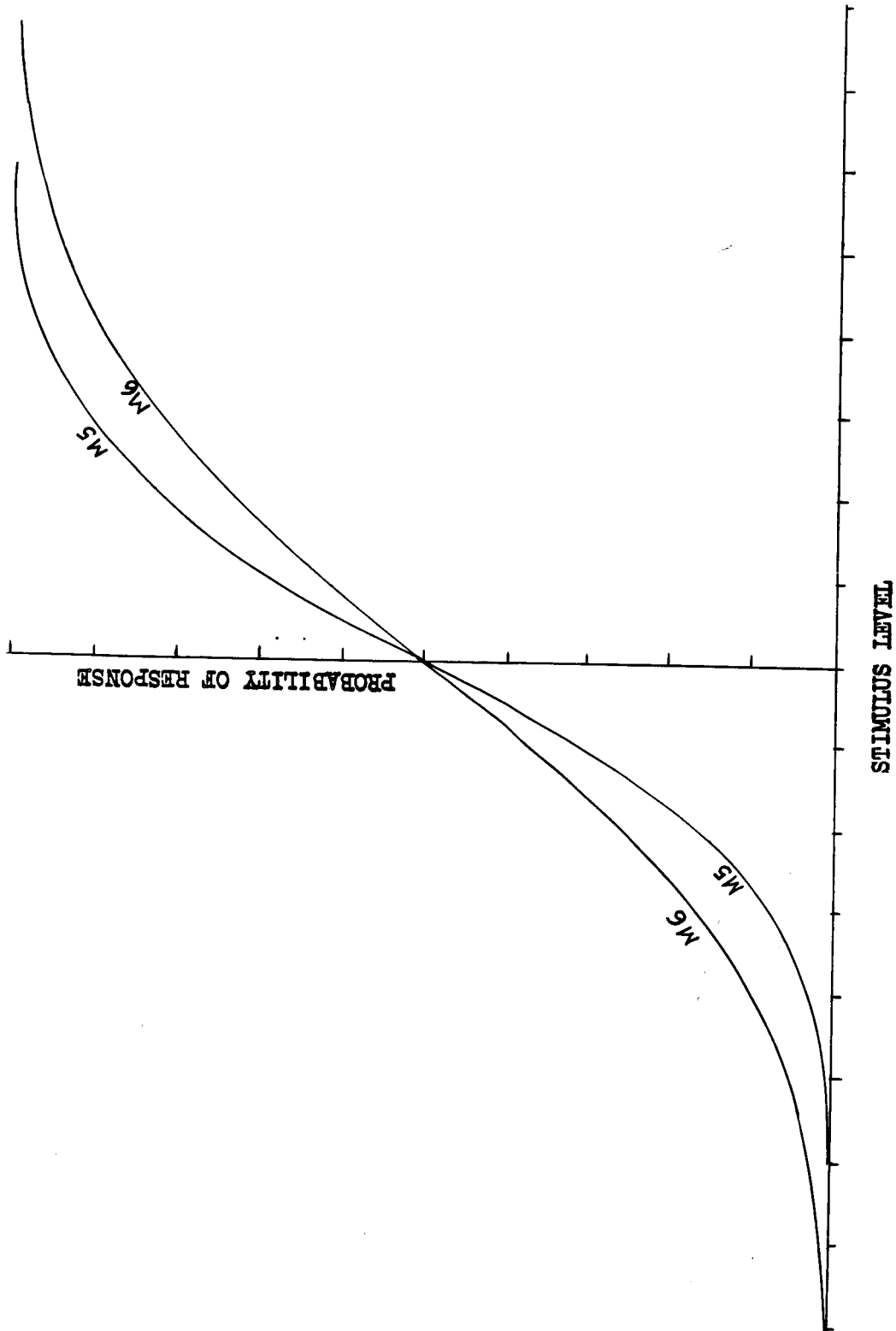


Figure 4.2 Cumulative Normal Response Curves

Stimulus level x_i	No. of specimens Tested (n)	No. of Responses (r)	Prob. of Response p_i
1.0	5	5	1.0
.8	5	3	0.6
.6	5	2	0.4
.4	5	1	0.2
.2	5	1	0.2
0	5	0	0.1

Fig. (4.3) EXAMPLE OF A DATA FOR SPEARMAN-KÄRBER METHOD

		No. of 1's	No. of 0's
1.0	1	1	0
.8	0 1 1 1 1	4	1
.6	1 1 1 1 0 1 1 0 0 1 1 1 0 1 1	11	4
.4	1 0 0 1 1 0 1 0 0 1 1 1 1 0 1 0 0	12	11
.2	0 0 0 0 1 0 0 0 1 0 0 0 0 0	2	12
0	0 0	0	2
TOTAL		30	30

Fig. (4.4) EXAMPLE OF A DATA FOR DIXON-MOOD METHOD

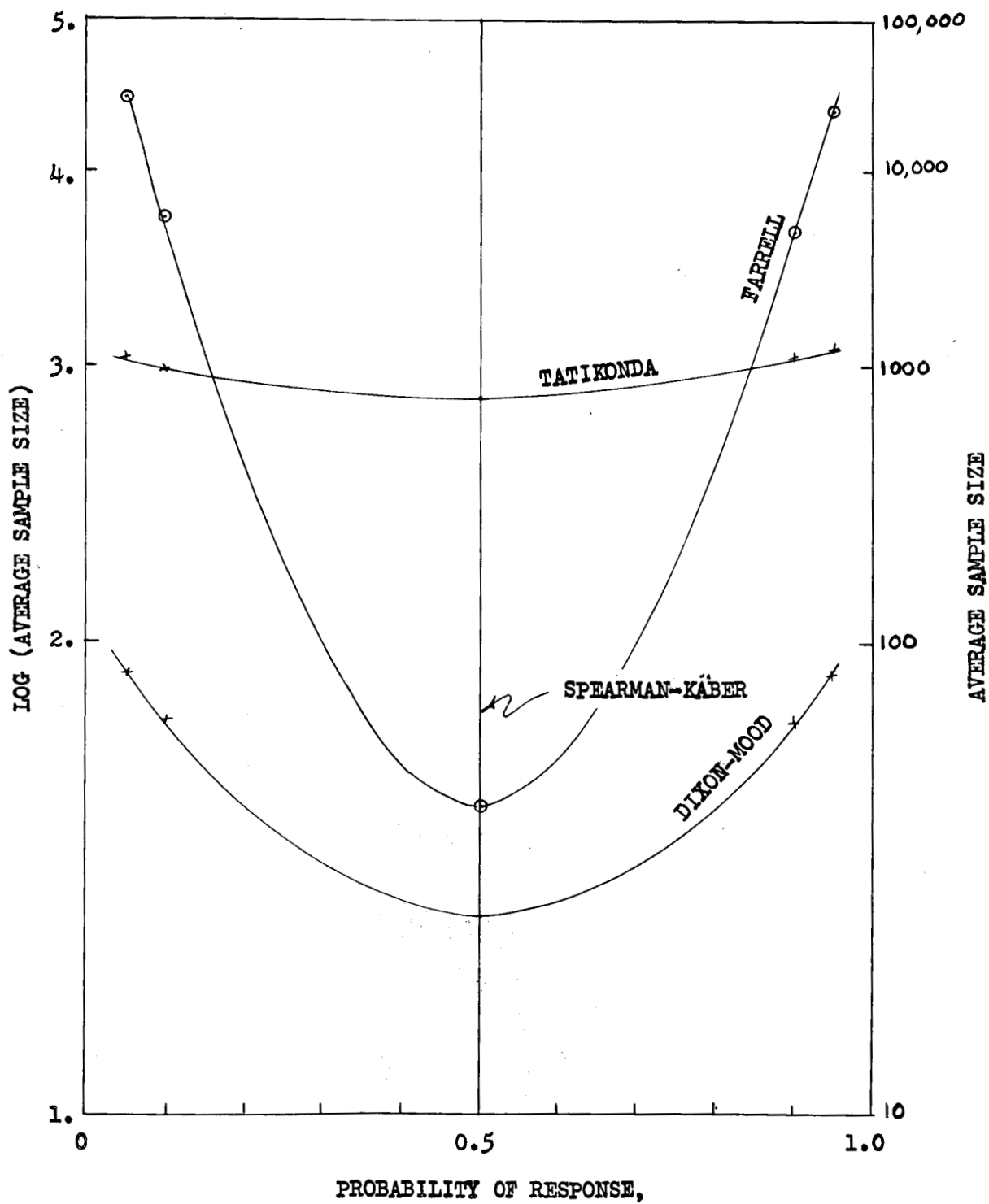


Figure 4.5 Sample size versus probability of response

Chapter V

COMMENTS ON THE RESULTS OF THE EMPIRICAL STUDY

A significant result of the empirical study seems to be that the two nonparametric methods given by Farrell and the author involve large sample sizes especially when the tail points are of interest. In the methods given by Spearman-Kärber and Dixon-Mood, the sample sizes are not large but it is assumed that some function of the stimulus has cumulative normal distribution. In these two methods the sample size increases with d , the difference between the normalized levels. In the Spearman-Kärber method, the sample size when $d > 2\sigma$ increased two to four times the sample size when $d < 2\sigma$. Similar behavior is noted in Dixon-Mood's method and even though the sample size required to obtain the desired confidence interval for the 50 percent point varies little, there is significant increase in the sample size required at the tail points. The sample size does not vary significantly when the starting levels are chosen near the 50 percent level or away from the 50 percent level. Probably the reason for this is as follows: When we choose a level far from the 50 percent level, the experimentation will be made at steadily decreasing levels or increasing levels depending on whether the first experiment is made at a level larger than the 50 percent level or smaller than the 50 percent point. From then onwards the process will continue as if the initial experiment is made near the 50 percent point. In estimating closely the 50 percent point, one needs some additional experiments which are finite in number and depend on the

distance that the initial level is from the true 50 percent point. In the results given in Table II, the initial level is chosen at the 97 percent level (that is, about $\mu + 2\sigma$ in the case of a normal curve). The restriction B3 in Chapter II, which says that the sample should be about 40 to 50 in order to use the method seems to be unnecessary, as the estimations obtained when the sample sizes are 20, 30 appeared reliable. This is also consistent with the results shown by Brownlee, Hodges, and Rosenblatt [5].

On studying the Tatikonda techniques it appears that an advantage is gained due to the distribution-free assumption, concerning the response curve $M(x)$, which may be important in many experiments. Unfortunately, the sample sizes are very large compared to two of the methods. Among the two nonparametric methods, the one proposed by the author seems to be better with respect to sample size than Farrell's method. The analysis and computation are much simpler in the method proposed by the author when compared to that given by Farrell. Farrell's analysis is not as easy to understand as the other methods and the computation process is a bit tedious. This method depends mainly on the fact that the sequence $N(0), N(1), \dots, N(n)$ takes the value i infinitely many times as n infinitely large, which is a result proved by Harris [15]. But as the process proceeds the difference between the members in the sequences $\{c(n), n > 1\}, \{d(n), n > 1\}$ becomes very large, and thus makes the total sample size large. It is sometimes noted that the difference between two consecutive members of these sequences is as large as 1000.

In general the up-and-down method does not seem to be efficient for the end points. That is, it is not an efficient method for estimating the small or large percentage points unless the normality assumptions are made. Another obvious disadvantage of this general method is that each specimen s must be tested separately. This very well could be the reason why the method is not used in many quantal response experiments [5]. The total time required for computation is likely to be prohibitive unless the response to the stimulus is immediate, as for example, in sensitivity testing of explosives. But in tests of insecticides, for example, a large group of insects can sometimes be treated as a single one. In large experiments of this kind any advantage of the up-and-down method may be outweighed by the requirements of single test. It is not necessary that the total time required to run an up-and-down series is n times the time required to conduct some other non-sequential experiment with n trials. When the up-and-down is made sequentially for a sufficient number of times, it is possible to make an estimate for the required stimulus level with the guaranteed accuracy, regardless of the initial guess at that point (but in the case of normality it depends on the guess for σ). The estimates in the up-and-down always exist, while in some other methods, this need not be true for small size.

Considering the application of Spearman-Kärber method and Dixon-Mood method to response functions which are non-normal, it seems that the analysis given by Spearman-Kärber still holds for non-normal response functions, with the assumption A3. While considering the Dixon-Mood

method the assumption B4 is basic. It is noted that in this particular study sometimes the computed intervals did not contain the true value. These values are given in Table IV. These numbers indicate the number of times the interval failed to include the true value in the total of 5 trials. The average number of these failures are given in Table V. The interval always includes the 50 percent point, whether the response curve is cumulative normal or not. From the results shown in Table IV, it seems to be that the number of times the interval includes the desired point increases when $d \geq 2\sigma$. The reason for this may be the increase in sample size when $d \geq 2\sigma$. The average of the sample sizes required by different techniques is given in the form of a graph. The graph shows that the sample size required by Dixon-Mood method and Spearman-Kärber method is almost negligible compared to the sample size required by Farrell's method for the end points. The sample size required by the method given by the author is nearly a constant. That is, it does not vary much from the 50 percent point to the end points. It depends strongly on the slope of the curve at the stimulus level being estimated, where as, in the other methods the sample size depends on the percentage point being estimated. Another difference between the method given by the author and the competing methods studies here, is that in all competing methods, the sample sizes are random numbers, while sample size associated with the author's method are fixed.

BIBLIOGRAPHY

- (1) Bartlett, M. S., "A Modified Probit Technique for Small Probabilities", Supp. J. R. Stat. Soc. 8 (1946), pp. 173-177.
- (2) Berkson, J., "Review of Probit Analysis by D. J. Finney", American Statistical Association Journal, 43 (1948), pp. 148-153.
- (3) Berkson, J., "A Statistically Precise and Relatively Simple Method of Estimating the Bioassay with Quantal Response", American Statistical Association Journal, 48 (1953), pp. 565-599.
- (4) Bliss, C. I., "The Calculation of the Dosage Mortality Curve", Annals of Applied Biology, 22 (1935), pp. 134-166.
- (5) Brownlee, K. A., Hodges, J. L. and Rosenblatt, M. "The Up-and-Down Method with Small Samples", American Stat. Assoc. Journal 48, (1953), pp. 262-277.
- (6) Derman, C., "Non Parametric Up-and-Down Experimentation", Annals of Math. Stat., 28 (1957), pp. 795-798.
- (7) Dixon, W. J. and Mood, A. M., "A Method for Obtaining and Analyzing Sensitivity Data", American Stat. Assoc. Journal, 43 (1948), pp. 109-126.
- (8) Dvoretzky, "On Stochastic Approximation", Proc. 3rd Berkeley Sym. on Math. Stat. and Prob., J. Neyman (ed.), Univ. of Calif. Press, (1956), pp. 39-55.

- (9) Epstein, B. and Churchman, C. W., "On the Statistics of Sensitivity Data", Annals of Math.Stat., 15 (1944), pp. 90-96.
- (10) Farrell, R. H., "Bounded Length Confidence Intervals for the Zero of a Regression Function", Annals of Math. Stat., 33 (1962), pp. 237-247.
- (11) Fisher, R. A., *ibid.*, p. 149.
- (12) Finny, D. J., Statistical Methods in Biological Assay, Hofner (1964), pp. 507-530.
- (13) Gayle, J. B., "Nonmonotonicity on Sensitivity Test Data", NASA TMX-53194.
- (14) Guttman, L. and Guttman, R., "An Illustration of the Use of Stochastic Approximation", Biometrics, 15 (1959), pp. 551-559.
- (15) Harris, T. E., "First Passage and Recurrence Distribution", Trans. Amer. Math. Soc., 73 (1952), pp. 471-488.
- (16) Hull, T. E., and Dobell, A. R., "Random Number Generators", SIAM Review 4, 3 (1962), pp. 230-254.
- (17) Irwin, J. O. and Cheesman, E. A., "An Approximate Method of Determining the Median Effective Dose and its Error in the Case of Quantal Response", Journal of Hygiene, 39 (1939), pp. 574-580.
- (18) Kärber, G., Beitrage Zur Kollektiven Behandlung Pharmakologischer Reihenversuche. Archiv für Experimentelle Pathologie und Pharmakologie, (1931), 162, pp. 480-487.

- (19) Kiefer, J. and Wolfowitz, J., "Stochastic Estimation of the Max of a Regression Function", Annals of Math. Stat., 23 (1952), pp. 737-744.
- (20) Kemeny, J. G., and Snell, J. L., Finite Markov Chains, New York, D. Van Nostrand Company, Inc. (1960).
- (21) Miller, Lloyd C. "Biological Assays Involving Quantal Responses", Annals of The New York Academy of Sciences, 52. Art. 6, pp. 69-
- (22) Odell, P. L., Stochastic Approximation and Non-Parametric Interval Estimation in Sensitivity Testing which Involves Quantal Response Data, (1962) Unpublished Doctoral Thesis, Oklahoma State University.
- (23) Parzen, E., Modern Probability Theory and its Applications , John Wiley & Sons, Inc. (1962).
- (24) Robbins, H. and Monro, S., "A Stochastic Approximation Method", Annals of Math. Stat., 22, (1951), pp. 400-407.
- (25) Rothman, D., Alexander, M. J., and Zimmerman, J. M., "The Design and Analysis of Sensitivity Experiments", Final Report for NASA Contract, MAS8-11061, DCN1-3, 84-80257-01 (1F) CPB02-1200-63, May (1965).
- (26) Spearman, C., " The Method of Right and Wrong Cases", British Journal of Psychology, 2 (1911-1912), pp. 227-242.