



AFRL-RW-EG-TP-2017-001

Analysis of Sensitivity Experiments – An Expanded Primer

Douglas V. Nance

**Air Force Research Laboratory
Munitions Directorate/Ordnance Division
Lethality, Vulnerability and Survivability Branch
(AFRL/RWML)
Eglin AFB, FL 32542-5910**

March 2017

**Distribution A: Approved for public release; distribution unlimited.
Approval Confirmation 96TW-2017-0033 dated 07 February 2017**

**AIR FORCE RESEARCH LABORATORY
MUNITIONS DIRECTORATE**

■ Air Force Materiel Command ■ United States Air Force ■ Eglin Air Force Base, FL 32542

REPORT DOCUMENTATION PAGE				Form Approved OMB No. 0704-0188	
Public reporting burden for this collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing this collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden to Department of Defense, Washington Headquarters Services, Directorate for Information Operations and Reports (0704-0188), 1215 Jefferson Davis Highway, Suite 1204, Arlington, VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to any penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number. PLEASE DO NOT RETURN YOUR FORM TO THE ABOVE ADDRESS.					
1. REPORT DATE (DD-MM-YYYY) 08-03-2017		2. REPORT TYPE Technical Paper		3. DATES COVERED (From - To)	
4. TITLE AND SUBTITLE Analysis of Sensitivity Experiments – An Expanded Primer				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER 62602F	
6. AUTHOR(S) Douglas V. Nance				5d. PROJECT NUMBER 25024504	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER W0TP	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Air Force Research Laboratory, Munitions Directorate Ordnance Division Lethality, Vulnerability, and Survivability Branch (AFRL/RWML) Eglin AFB FL 32542-6810				8. PERFORMING ORGANIZATION REPORT NUMBER AFRL-RW-EG-TP-2017-001	
9. SPONSORING / MONITORING AGENCY NAME(S) AND ADDRESS(ES) Air Force Research Laboratory, Munitions Directorate Ordnance Division Lethality, Vulnerability, and Survivability Branch (AFRL/RWML) Eglin AFB FL 32542-6810 Technical Advisor: Kirk J. Vanden, PhD				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT AFRL-RW-EG-TP-2017-001	
12. DISTRIBUTION / AVAILABILITY STATEMENT Distribution A: Approved for public release; distribution unlimited. Approval Confirmation 96W-2017-0033, dated 07 February 2017					
13. SUPPLEMENTARY NOTES					
14. ABSTRACT This report is expository in nature and is intended to explain the basic concept of a sensitivity experiment as well as the analysis of its data. These experiments are widely applicable to engineering problems that involve binary (pass/fail of go/no-go) outcomes. We begin by introducing the simple idea behind a sensitivity experiment; then we describe the basic "Up and Down" testing method. Particular emphasis is placed upon the parameter this procedure is intended to identify. Next, we describe specific Probit analyses, for analyzing sensitivity test data. A specialized version of this scheme is derived for stable digital computation. Confidence region estimation is discussed along with an analysis of variance. The Logit method is also presented as an alternative to the Probit method. A set of example problems are solved; our results are compared with archival solutions.					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT	18. NUMBER OF PAGES	19a. NAME OF RESPONSIBLE PERSON
a. REPORT	b. ABSTRACT	c. THIS PAGE			Pedro A. Lopez-Fernandez
UNCLASSIFIED	UNCLASSIFIED	UNCLASSIFIED	SAR	96	19b. TELEPHONE NUMBER (include area code) 850-883-2707

This page intentionally left blank

TABLE OF CONTENTS

List of Figures.....	ii
List of Tables.....	iii
Preface.....	iv
Acknowledgements.....	v
Summary.....	v
1 Introduction.....	1
1.1 Sensitivity Tests.....	2
1.2 The Bruceton Testing Method.....	3
1.3 Difficulties with Data Reduction.....	4
2 Technical Approach.....	9
2.1 Maximum Likelihood Equation.....	9
2.2 Dosage Dependent Probabilities.....	10
2.3 Garwood's Method.....	11
2.4 Bernoulli Trial with an Endpoint Probability.....	15
2.5 Newton's Solution Method.....	21
2.6 Minor Variant of the Golub-Grubbs Method.....	23
2.7 Analysis of Confidence in the Estimates.....	28
2.7.1 Success Ratio Confidence Interval.....	30
2.7.2 Success Probability Confidence Interval in the Explanatory Variable...	35
2.7.3 Confidence Analysis for the Mean and Standard Deviation.....	37
2.8 Measures of Variance.....	48
2.9 Logit Fitting Procedure.....	50
3 Test Problems and Results.....	57
3.1 Test Problem 1: Antibiotic Efficacy.....	57
3.2 Test Problem 2: Morbidity of Salmonella Typhimurium.....	62
3.3 Test Problem 3: Arsenic Toxicity.....	66
3.4 Test Problem 4: Dixon/Mood Bruceton First Test Series.....	70
3.5 Test Problem 5: Dixon/Mood Bruceton Second Test Series.....	73
3.6 Test Problem 6: Anti-Pneumococcus Serum.....	76
3.7 Test Problem 7: Aircraft Fastener Failure.....	80
4 Conclusions.....	83
References.....	85

LIST OF FIGURES

<u>Figure</u>	<u>Page</u>
1	Probability Distribution Function for the Standard Normal Distribution.....31
2	Normal Curve Linear Fit Form for Test Problem 1.....58
3	Probability Distributions for Test Problem 1.....59
4	95% Confidence Ellipse for Parameter Estimates α and β in Test Problem 1.....60
5	Normal Curve Linear Fit Form for Test Problem 2.....63
6	Probability Distributions for Test Problem 2.....63
7	95% Confidence Ellipse for Parameter Estimates α and β in Test Problem 2.....64
8	Normal Curve Linear Fit Form for Test Problem 3.....66
9	Probability Distributions for Test Problem 3.....67
10	95% Confidence Ellipse for Parameter Estimates α and β in Test Problem 3.....68
11	Probability Distributions for Test Problem 4.....70
12	95% Confidence Ellipse for Parameter Estimates α and β in Test Problem 4.....72
13	Probability Distributions for Test Problem 5.....73
14	95% Confidence Ellipse for Parameter Estimates α and β in Test Problem 5.....74
15	Linear Logistics Fit Form for Test Problem 6.....76
16	Probability Distributions for Test Problem 6.....77
17	95% Confidence Ellipse for Parameter Estimates α and β in Test Problem 6.....78
18	Linear Logistics Fit Form for Test Problem 7.....80
19	Probability Distributions for Test Problem 7.....81
20	95% Confidence Ellipse for Parameter Estimates α and β in Test Problem 7.....82

LIST OF TABLES

<u>Table</u>	<u>Page</u>
1 Illustration of Responses for a Bruceton Test Series.....	6
2 Summary of the Bruceton Test Results extracted from Table 1.....	6
3 Anti-Pneumococcus Serum Test Data.....	57
4 95% Confidence Interval Endpoints (Model 1) for Test Problem 1.....	60
5 95% Confidence Interval Endpoints (Model 2) for Test Problem 1.....	60
6 Salmonella Typhimurium Morbidity Test Data.....	62
7 95% Confidence Interval Endpoints (Model 1) for Test Problem 2.....	64
8 95% Confidence Interval Endpoints (Model 2) for Test Problem 2.....	64
9 Arsenic Toxicity Test Data.....	66
10 95% Confidence Interval Endpoints (Model 1) for Test Problem 3.....	67
11 95% Confidence Interval Endpoints (Model 2) for Test Problem 3.....	68
12 Bruceton Test Data for Test Problem 4.....	70
13 95% Confidence Interval Endpoints (Model 1) for Test Problem 4.....	71
14 95% Confidence Interval Endpoints (Model 2) for Test Problem 4.....	71
15 Bruceton Test Data for Test Problem 4.....	73
16 95% Confidence Interval Endpoints (Model 1) for Test Problem 5.....	74
17 95% Confidence Interval Endpoints (Model 2) for Test Problem 5.....	74
18 Logistic Test Data for test Problem 6.....	76
19 95% Confidence Interval Endpoints (Model 1) for Test Problem 6.....	77
20 95% Confidence Interval Endpoints (Model 2) for Test Problem 6.....	78
21 Logistic Test Data for Test Problem 7.....	80
22 95% Confidence Interval Endpoints (Model 1) for Test Problem 7.....	81
23 95% Confidence Interval Endpoints (Model 2) for Test Problem 7.....	82

PREFACE

This publication is a revised and expanded edition of a technical report written during the Fall of 2008. That technical report was produced during a flurry of activity hurriedly conducted to evaluate a body of technical work that was relatively new to this organization. Out of practical necessity, the research and code development performed for this project had to be accomplished with haste to meet the constraints of schedule. Years later, it seems that the old report is in need of much improvement; also, a few additional concepts need to be addressed. Specifically, the discussion of confidence regions has been clarified and made more complete. The exposition on confidence ellipses is quite new, and in my view, stands as an important addition to the theory behind sensitivity testing. Also, the Logit method based upon the logistic distribution has been derived in some detail. Additional examples have been added to the exposition, and some effort has been invested in placing a more substantial level of detail within these exercises. Hopefully, this information will be of help to the reader. Naturally, as with any first version manuscript containing lots of equations, the errata list was quite long enumerating typographical mistakes, improperly transcribed equations, etc. A fair amount of work has been performed to correct these mistakes and improve the quality of the derivations. I hope that I have had some success in cleaning up the manuscript, at least in as much as is possible for someone who passed typing class with a score of thirty words per minute with five mistakes.

Statistics is a tough topic, so why expend so much effort on this report? (It is a valid point; this discipline is a tough study, and I have spent a lot of time on it). Well, it is important as is evidenced by its application in the pharmaceutical industry, specifically in qualifying new medications. Of course, there are many other applications for these ideas, particularly for “one-shot” items that absolutely have to work right when they are used. Secondly, this material is very interesting because it integrates different concepts in probability. This integral view of probability forces its practitioner to understand more of the science behind the equations. In many cases, people gravitate toward a deterministic view of the world where the most of one’s day is thought of in terms of absolutes. That is to say, one has for instance a fixed amount of pocket change, 37 minutes left to the end of the workday, or maybe 93 days left until Christmas. When you begin to think in terms of probability, a new paradigm must be adopted. Instead of thinking of these fixed quantities, we must think instead of how these quantities are distributed with respect to one or more variables. Quantities are addressed in terms of average values and probabilities instead of exact values. This idea is important when predicting the behavior of a system with some level of certainty. All measurements are subject to uncertainty, and these uncertainties are propagated through to a system’s output with a modified level of uncertainty. This idea is nicely illustrated by sensitivity analysis where an uncertainty governed by a continuous probability distribution is integrated with the discrete binomial distribution. The statistical properties of the combined distribution are determined by the algorithms described in this paper. The examples are intended to help the reader understand the main concepts.

ACKNOWLEDGEMENTS

The author is grateful to Mr. Craig Schlenter of DetNet Ltd. for posing several insightful questions on the progenitor of this report. These question have served as the motivation for this revision. This effort has been funded by the Munitions Directorate of the Air Force Research laboratory. Without the support of this directorate, this work would not have been possible.

SUMMARY

This report is expository in nature and is intended to explain the basic concept of a sensitivity experiment as well as the analysis of its data. These experiments are widely applicable to engineering problems that involve binary (pass/fail of go/no-go) outcomes. We begin by introducing the simple idea behind a sensitivity experiment; then we describe the basic “Up and Down” testing method. Particular emphasis is placed upon the parameter this procedure is intended to identify. Next, we describe specific Probit analyses, for analyzing sensitivity test data. A specialized version of this scheme is derived for stable digital computation. Confidence region estimation is discussed along with an analysis of variance. The Logit method is also presented as an alternative to the Probit method. A set of example problems are solved; our results are compared with archival solutions.

1 INTRODUCTION

Statistics is a complicated science. Having terrorized dormitories and classrooms across the country, its late night, torturous study sessions have earned it a dubious “GPA-busting” reputation along with the forlorn appellation *Sadistics* among captive students. Although it is a division of mathematics, it stands aside from the mainstream of this discipline. In fact, its more skillful practitioners refer to themselves as *probabilists* instead of mathematicians. Statistics can also incite chagrin among practicing engineers and scientists. In a recent meeting with the author, this assertion was reaffirmed as one professional uttered the phrase, “*I hate statistics!*”

Why does statistics evoke such negative emotion? There is a reasonable answer for this question. Although many professionals are required to use it, they lack the theoretical background and experience required in order to understand how statistics works. It is a branch of mathematics that differs from subjects like calculus or linear algebra. In these subjects, an individual can learn a few basic ideas and apply them in a nearly “cookbook” way to solve problems. Statistics is not so cooperative. Statistical or stochastic theory is deeply buried in advanced mathematical theory involving topics such as *Lesbeque Integration*, *Measure Theory* and *Functional Analysis*. [1] If these topics are well understood, then probabilistic theory is more accessible, and with practice statistics becomes understandable. Unfortunately, few people have the time and patience required to study and master these disciplines. Hence, statistical theory remains arcane, and its mastery eludes all but its most diehard practitioners.

The difficulty associated with mastering statistical inference presents a true dilemma. Statistics is an extremely applied science. It has an extensive number of applications in practically every field of engineering or scientific endeavor. But in order to obtain good computational results, great care is required while converting abstract statistical theories into practice. After having

completed graduate study in stochastic processes, the author now agrees, at least in part, with his venerable instructor and thesis supervisor. Statistics is best taught and understood through the use of clearly explained examples along with carefully planned excursions into probability theory. We have attempted to take this approach in the discussions that follow. He who dares to venture directly into the world of stochastic theory is doomed to become lost, perhaps forever. In other words, when on this safari, you need a guide.

This report is designed to be a guide, of sorts. It focuses on analytical methods used to process binary (Bernoulli) and binomial statistical trials. A single binary trial or experiment has only two possible outcomes, either a *success* or a *failure*. A binomial trial may be thought of as a series of binary trials taken at the same “level”. The result of a binomial trial consisting of n binary experiments is say, p successes and $(n - p)$ failures with a crude probability of success of p/n for a particular binomial trial.[2] A binary trial is a special case; it is a binomial trial with $n = 1$ and p can be either 0 or 1. As it happens, *Sensitivity Tests* are characterized by a mixture or series of binary and binomial trials. For military applications, these tests have a great deal of utility since many items of military hardware may be used only once and may only be evaluated by pass/fail criteria. Many munitions can be thought of in this way. For this reason, we describe the basic features of a sensitivity test in the next section of this report.

1.1 Sensitivity Tests

A general sensitivity test consists of a finite series of either binary or binomial experiments (or a mixture of the two). Each experiment is performed at a *level* defined by distinct value of an explanatory variable.[3] The value of the explanatory variable directly corresponds to the measurable *dosage* for this level. The dosage is related to the *stimulus*, an unmeasurable quantity

that drives the outcome of the experiment (or trial) at the chosen level.[4] The stimulus may be thought of as a mechanistic (physical, biological, etc.) process that results from the dosage. Each binary trial has only two possible outcomes, a success or a failure. If a binomial trial (consisting of n binary trials) is conducted at a given level, then without any loss of generality, we can say that p successes and $(n - p)$ failures result at this level. A basic assumption behind the sensitivity test is that as the dosage (stimulus) increases, then the probability of success also increases. That is to say, at “lower” levels of stimulus, we expect more failed trials. As the stimulus increases at higher levels, we expect more successful trials, up to the point where practically all individual binary trials are successes (or for binomial trials, p equals n).[5] Unfortunately, due to the pass/fail nature of each trial’s outcome, we cannot *precisely* determine the dosage that will result in a success. Instead, all that we can do is select a dosage and then, via test, determine whether the *critical dosage* is higher or lower than the test dosage.[6] That is to say, if the test is a success, the critical dosage is less than or equal to the test dosage. Conversely, if the test is a failure, then the critical dosage is greater than the test dosage. To promote greater understanding, let us consider a basic example of sensitivity testing.

Suppose that a pharmaceutical manufacturer has developed a new antibiotic to fight a particular strain of bacterial pneumonia. As a normal part of the certification procedure, an effective dosage must be estimated for this drug. *The critical dosage* is defined as the mass (say, in milligrams) of the drug that must be administered to eradicate all invading bacteria *with no excess antibiotic remaining in the body*. As you may imagine, this dosage is impossible to measure. Also, the outcome of the test is determined by the stimulus. The stimulus for the problem requires a detailed knowledge at the microscopic scale of how the antibiotic interacts with each individual

bacterium. This information is not known and cannot be determined, so we must obtain an estimate of the effective dosage by sensitivity testing.

To estimate the effective antibiotic dosage, we select a set of dosage levels for testing, for example {0, 10, 20, 30, 40, 50} milligrams. Naturally, this set defines six dosage levels for the sensitivity test. We then select a number of healthy test animals (say, Rhesus monkeys) for use and decide how many are to be tested. In as much as is possible, the test animals (or test articles) should be chosen so that they have the same anatomical and physiological (or physical) characteristics.[7] If so desired, we can formulate the entire sensitivity test series as a sequence of binary trials, one test animal per trial. To conduct a trial in this sequence, we choose a level and infect a test animal with the pneumonia bacteria. Then we wait a pre-specified amount of time and administer the dosage for the chosen level to the test animal. In the period of time after administering the drug, we determine the response; the test animal either lives (a success) or dies (a failure). The response data is collected, and we then move onto the next level and continue testing. At the conclusion of the entire sequence, we have determined a number of successes and failures at each dosage level.

This example brings an important fact to light. As we move between trials, we cannot reuse test articles, (e.g., test animals). Why? Obviously, if a test animal dies at the preceding trial, it cannot be reinfected and tested. If the animal survives the preceding trial, its physiology has been altered by the presence of the antibiotic and by the action of its immune system. The animal's bloodstream, lymphatic system and pleural tissues are inundated with antibodies to the pneumonia. As a result, it cannot be equivalently reinfected, so the results of a subsequent test would be tainted. Once a test article has been exposed to the stimulus, it should not be tested again. The output of a complete sensitivity test series may be represented as a plot of percentage success (p/n) versus the

dosage (or explanatory variable) level. We observe that the change in the success percentage is smaller for low and high dosages while the greatest change occurs near the mean dosage. Hence, we assume that the dosage response curve follows the cumulative normal distribution function.[7] When a sufficient amount of data has been collected, this assumption is testable.[7] With an appropriate number of tests, we can determine the “50% point” for the distribution, the level of dosage that causes a successful response for 50% of all like test articles. This value of dosage is the *mean* for the distribution.[2] We are also interested in how success percentages behave for dosages further away from the mean. The property is termed as *dispersion* and is measured by the *standard deviation*. The primary focus of post-test analysis resides in determining these properties for the distribution.

1.2 The Bruceton Testing Method

The mean μ and standard deviation σ must be estimated accurately if our analysis is to have real meaning. For this reason, the sensitivity test procedure is designed to concentrate measurements around the mean in order to refine the estimate. The Bruceton or “Up and Down” test is commonly used for this purpose.[6] The first step in conducting a Bruceton test is to decide the range of dosages (minimum to maximum) as well as the dosage levels. We must also decide the total number of individual binary trials to be performed during the test series. In practice, we require a minimum of 20 trials, but twice that number is recommended since under theoretical limitations, the effective sample size is usually half of the actual sample size.[6] Table 1 illustrates a notional Bruceton test series. The results of the initial trial are recorded in column T2 while the response for the last trial is recorded in column T21. For the initial trial, we choose a dosage (level 2 in the example) that is believed to be closest to the actual mean. If the outcome of this trial is a

Table 1. Illustration of responses for a Bruceton test series. Series consists of 20 binary trials. The notation "x" denotes a success while "o" denotes a "failure". Six dosage levels are numbered 1 to 6 from low to high.

TRIAL T	T 1	T 2	T 3	T 4	T 5	T 6	T 7	T 8	T 9	T 10	T 11	T 12	T 13	T 14	T 15	T 16	T 17	T 18	T 19	T 20
Lev 1						o														
Lev 2	o				x		o						o							
Lev 3		o		x				o		o		x		o						
Lev 4			x						x		x				o		o			
Lev 5																x		o		o
Lev 6																			x	

Table 2. Summary of the Bruceton test results extracted from Table 1.

Dosage Level	1	2	3	4	5	6
Successes (x)	0	1	2	3	1	1
Failures (o)	1	3	4	2	2	0
Total Number Trials	1	4	6	5	3	1

success, we conduct the next trial at the dosage just below the initial level. On the other hand, if the first trial is a failure, we conduct the second test at the dosage level just above the initial level. The dosage level for the third trial is determined in the same manner but is based upon the outcome of the second trial. Dosage levels for the remaining trials are determined in the same way, but in most cases, we do not use all of the data collected. The reason for excluding some of the data is based upon our desire to resolve the normal mean. As we stated above, Bruceton testing concentrates measurements in the vicinity of the mean dosage. Since the change in dosage level between trials “opposes” the sense of the preceding response, the up and down nature of the test series attempts to center data collection at the location of the mean. Unfortunately, if our initial guess for the mean dosage level is poor, we automatically introduce extraneous data into the test series. Suppose that we have guessed an initial value for the dosage that is too low. Then for two or more tests, we will steadily increase the dosage level and obtain failed responses. For example, see levels 2 and 3 in Table 1. Let us assume that we obtain a success on the third trial (shown in

column 4). This sudden change in the response (indicated by the red block in Table 1) is termed as a *reversal*. It follows that we avoid introducing extraneous “start-up” data by analyzing only information obtained beginning with the second trial. In this way, data collection tends to remain near the mean. The binomial results of this example test series are shown versus dosage level in Table 2.

1.3 Difficulties with Data Reduction

Sensitivity testing does have its caveats. In the first place, the test data must contain a *zone of mixed results*. [8] That is to say, the highest level at which a trial fails must exceed the lowest level observed for a successful trial. If this condition is not satisfied, the equations for estimating σ become inconsistent, and zero is the only value that can be obtained for σ . [9] In the example shown in Table 1, “2” is the lowest level corresponding to a success while “5” is the highest level containing a failure. Hence, for this example, we obtain a zone of mixed results between levels “2” and “5”. The Bruceton test procedure is usually successful in establishing the mixed zone. [8]

Sensitivity testing also possesses certain intrinsic weaknesses. Since it concentrates on resolving the distribution mean, it loses accuracy near the “tails” of the distribution, i.e., the regions near 0% and 100% response. [6] For reliability problems, we are most interested in the region corresponding to response values exceeding 99%. As a result, care must be taken when sensitivity tests are conducted with this purpose in mind. Due diligence must be paid to the structure of the dosage levels and to the number of trials. The chosen data analysis methodology is equally important since the choice of distribution can have a significant effect in the high probability region. In certain cases, the logistic distribution is chosen over the normal distribution due to its more conservative nature in the “tails”, i.e., the tails are longer. [10] This shortcoming emphasizes the importance of correctly calculating confidence intervals for system reliability probabilities in

the tail regions. Secondly, the maximum likelihood estimation (MLE) procedure used to analyze the test data generally requires a large data set. For small data sets, MLE loses “efficiency”; the estimates of μ and σ become biased and acquire excessive variance.[9] Moreover, MLE may not even be able to analyze certain small data sets. A classic example is that of a test series without a mixed zone of results.[9] Note that Dixon and Mood have published the disclaimer: “Measures of reliability may be very misleading if the sample size is less than forty or fifty.”[6] To cope with these potential sources of error, careful test planning must be combined with sound numerical estimation techniques.

2 TECHNICAL APPROACH

For experiments designed with a fixed dosage increment, one may obtain the mean μ and standard deviation σ (or perhaps other linear fitting parameters) by plotting the percentage of successes at each dosage level on probability paper.[7] If this data correlates well with the normal distribution, μ and σ may be extracted graphically from the plot. Unfortunately, for certain types of sensitivity tests, the dosage cannot be precisely controlled, and the data may not exactly conform to the normal distribution. The former difficulty is commonly encountered for system reliability test series. To achieve the best estimates of the mean and standard deviation for these tests, we usually apply generalized MLE procedures.[4]

2.1 Maximum Likelihood Equation

The *Method of Maximum Likelihood* is a powerful estimation procedure that was developed by R.A. Fisher during the first two decades of the Twentieth Century. To adapt this method for analyzing our sensitivity tests, we apply theoretical concepts from probability. Individual binary trials performed in the course of the test series are effectively *independent* from the standpoint of probability. The outcome (or response) of one trial has no effect on the outcome of any other trial. If we envision that the outcome of each trial is represented by its own random variable, then in terms of the dosage (an explanatory variable), these random variables are *identically distributed*. That is to say, they have the same distribution function and parameters (μ and σ). Since the trials are independent, the probability of events occurring in two trials is given by the product of their individual probabilities. Suppose that a total of n_i binary trials are conducted at the i^{th} dosage level. Let the success probability at this level be given by p_i and the failure probability by

$$q_i = 1 - p_i \quad (2.1.1)$$

With these assumptions, the entire test series can be envisioned as a sequence of Bernoulli (pass/fail) trials. At this point in the exposition, we invoke Garwood's mathematical expressions for determining a lethal dosage for a drug.[11] In this lexicon, a *success* is equivalent to a *death* resulting from a dosage while a *failure* is equivalent to *survivor*. Moreover, consider that there are $n_i - s_i$ successes (deaths) at a given dosage level and s_i failures (survivors). As an informational aside, note that the s_i are random variables. The associated likelihood function can be written as

$$\tilde{L} = \prod_{i=1}^N \binom{n_i}{s_i} p^{n_i - s_i} q^{s_i} \quad (2.1.2)$$

where N is the number of dosage levels; the term in the parentheses is a *combination of n_i trials taken s_i at a time*.[11] The combination is used to represent the independence of order for Bernoulli trials defined at the same level. Equation (2.1.2) is the product of binomial probability distributions formulated at the dosage levels.[2] Applying the natural logarithm of the likelihood function lends a great deal of convenience, i.e.,

$$L = \sum_{i=1}^N \left\{ \ln \binom{n_i}{s_i} + (n_i - s_i) \ln(p_i) + s_i \ln(q_i) \right\} \quad (2.1.3)$$

where $L = \ln(\tilde{L})$. By maximizing L , the natural log of the likelihood function, we effectively maximize $\tilde{L}$. Equation (2.1.3) is referred to as the log-likelihood function.

2.2 Dosage Dependent Probabilities

To estimate the “mean dosage”, the dosage defined at the level where the probability equals 0.5, many probabilistic concepts must be brought together. As a result, it is very easy to lose sight of important theoretical details. One may note that we have done nothing to connect the MLE expression to the dosage (or explanatory variable). Equation (2.1.2) is cast in the form of a

binomial probability distribution, but one may note that p_i and q_i are not yet related to the dosage variable. These probabilities are related to the dosage variable through the use of a *link function*. [4] Several link functions are available, but we apply the normal probability distribution function commonly used in the “probit” method. [12] The normal distribution is based upon a continuous probability density function that must be integrated with respect to the explanatory variable (x_i at the i^{th} level) in order to compute p_i where

$$p_i = p(x_i) = \frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^{x_i} \exp\left[-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2\right] dx \quad (2.2.1)$$

Note that by this definition, p_i is the probability that $x \in (-\infty, x_i)$. When (2.2.1) is substituted into (2.1.3), we obtain an expression for the log-likelihood function that explicitly depends on the dosage level. To promote some simplicity, we define the variable

$$t(x) = \frac{x - \mu}{\sigma} \quad (2.2.2)$$

and we can rewrite p_i as

$$p_i(t_i) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{t_i} \exp\left(-\frac{t^2}{2}\right) dt \quad (2.2.3)$$

This equation is now in the form of the *standard normal distribution*. [9] Although (2.1.3), (2.2.2) and (2.2.3) are correct from the standpoint of theory, (2.2.2) must be placed in a different form to support the estimation procedure.

2.3 Garwood's Method

In the course of this research project, the author has encountered a number of different solution procedures for (2.1.3), (2.2.2) and (2.2.3). There is similarity between these methods; here Garwood's method is presented. [11] We begin this discussion with an alternate form for t , i.e.,

$$t = \frac{x - \mu}{\sigma} = \alpha + \beta x \quad (2.3.1)$$

This expression is just a linear polynomial in x , but it offers an advantage from the standpoint of differential calculus. From (2.3.1), it is easy to show that

$$\mu = -\frac{\alpha}{\beta}; \quad \sigma = \frac{1}{\beta} \quad (2.3.2)$$

In the light of (2.1.3) and (2.3.1), critical points, local maxima and minima of L , may be cast in ordered pairs (α, β) ; they may be determined by solving the system of equations:

$$\frac{\partial L}{\partial \alpha} = 0; \quad \frac{\partial L}{\partial \beta} = 0 \quad (2.3.3)$$

By differentiating (2.1.3), we may show that

$$\frac{\partial L}{\partial \theta} = \sum_{i=1}^N \frac{n_i}{p_i q_i} \left(q_i - \frac{s_i}{n_i} \right) \frac{\partial p_i}{\partial \theta} \quad (2.3.4)$$

where $\theta \in (\alpha, \beta)$. In this case, we have noted that

$$\frac{\partial q_i}{\partial \theta} = -\frac{\partial p_i}{\partial \theta} \quad (2.3.5)$$

According to (2.2.3), p_i is a function of t_i , so with the use of the chain rule, we may show that

$$\frac{\partial p_i}{\partial \theta} = \frac{\partial t_i}{\partial \theta} \frac{dp_i}{dt_i} = \frac{\partial t_i}{\partial \theta} p'_i(t_i) \quad (2.3.6)$$

By using the process of differentiation with respect to a parameter to (2.2.3), we have that

$$p'_i(t_i) = \frac{1}{\sqrt{2}} \exp\left(-\frac{t_i^2}{2}\right) = z_i; \quad q'_i(t_i) = -z_i \quad (2.3.7)$$

The derivatives of t_i , in the form of $\partial t_i / \partial \theta$, are calculated as

$$\frac{\partial t_i}{\partial \alpha} = 1; \quad \frac{\partial t_i}{\partial \beta} = x_i \quad (2.3.8)$$

Then (2.3.5), (2.3.6), (2.3.7) and (2.3.8) may be applied to show that

$$\frac{\partial p_i}{\partial \alpha} = z_i; \quad \frac{\partial p_i}{\partial \beta} = x_i z_i; \quad \frac{\partial q_i}{\partial \alpha} = -z_i; \quad \frac{\partial q_i}{\partial \beta} = -x_i z_i \quad (2.3.9)$$

By substituting (2.3.9) into (2.3.4), we obtain

$$\frac{\partial L}{\partial \alpha} = \sum_{i=1}^N \frac{n_i}{p_i q_i} \left(q_i - \frac{s_i}{n_i} \right) z_i \quad (2.3.10)$$

$$\frac{\partial L}{\partial \beta} = \sum_{i=1}^N \frac{n_i}{p_i q_i} \left(q_i - \frac{s_i}{n_i} \right) x_i z_i \quad (2.3.11)$$

Following Garwood [11], define ς_i as

$$\varsigma_i = \varsigma_i(t_i) = \frac{n_i}{p_i q_i} \left(q_i - \frac{s_i}{n_i} \right) z_i \quad (2.3.12)$$

Hence, (2.3.10) and (2.3.11) become

$$\frac{\partial L}{\partial \alpha} = \sum_{i=1}^N \varsigma_i \quad (2.3.13)$$

$$\frac{\partial L}{\partial \beta} = \sum_{i=1}^N \varsigma_i x_i \quad (2.3.14)$$

It is important to derive the second partial derivatives of the log-likelihood function; (2.3.13) and (2.3.14) are helpful in this process. With the use of (2.3.12), we observe that

$$\frac{\partial^2 L}{\partial \alpha^2} = \frac{\partial}{\partial \alpha} \sum_{i=1}^N \varsigma_i = \sum_{i=1}^N \frac{\partial t_i}{\partial \alpha} \frac{d\varsigma_i}{dt_i} = \sum_{i=1}^N \frac{\partial t_i}{\partial \alpha} \varsigma'_i \quad (2.3.15)$$

By using (2.3.8), we obtain

$$\frac{\partial^2 L}{\partial \alpha^2} = \sum_{i=1}^N \varsigma'_i \quad (2.3.16)$$

In a similar manner, the remaining second partial derivatives are derived as follows.

$$\frac{\partial^2 L}{\partial \alpha \partial \beta} = \frac{\partial}{\partial \alpha} \left(\frac{\partial L}{\partial \beta} \right) = \frac{\partial}{\partial \alpha} \left(\sum_{i=1}^N x_i \varsigma_i \right) = \frac{\partial t_i}{\partial \alpha} \sum_{i=1}^N \frac{d}{dt} (x_i \varsigma_i) = \sum_{i=1}^N x_i \varsigma'_i \quad (2.3.17)$$

and

$$\frac{\partial^2 L}{\partial \beta^2} = \frac{\partial}{\partial \beta} \left(\frac{\partial L}{\partial \beta} \right) = \frac{\partial}{\partial \beta} \left(\sum_{i=1}^N x_i \varsigma_i \right) = \sum_{i=1}^N x_i \frac{\partial t_i}{\partial \beta} \frac{d \varsigma_i}{dt} = \sum_{i=1}^N x_i^2 \varsigma'_i \quad (2.3.18)$$

In order to evaluate (2.3.16) through (2.3.18), an expression for ς'_i ; begin with (2.3.12), and apply the quotient rule.

$$\varsigma'_i = \frac{p_i q_i \{n_i z_i (q_i - s_i / n_i)\}' - n_i z_i (q_i - s_i / n_i) (p_i q_i)'}{(p_i q_i)^2} \quad (2.3.19)$$

$$\varsigma'_i = \frac{n_i}{p_i q_i} \left[z_i \left(q_i - \frac{s_i}{n_i} \right)' + z'_i \left(q_i - \frac{s_i}{n_i} \right) \right] - \frac{n_i z_i}{(p_i q_i)^2} \left(q_i - \frac{s_i}{n_i} \right) (p_i q'_i + p'_i q_i) \quad (2.3.20)$$

The application of (2.1.1) and (2.3.7) yields

$$\varsigma'_i = \frac{n_i}{p_i q_i} \left[-z_i^2 + z'_i \left(q_i - \frac{s_i}{n_i} \right) \right] - \frac{n_i z_i}{(p_i q_i)^2} \left(q_i - \frac{s_i}{n_i} \right) (-z_i p_i + z_i q_i) \quad (2.3.21)$$

With careful algebraic simplification, we obtain

$$\varsigma'_i = -\frac{n_i z_i^2}{p_i q_i} + \frac{n_i z'_i}{p_i q_i} \left(q_i - \frac{s_i}{n_i} \right) - \frac{n_i z_i^2}{p_i^2 q_i} \left(q_i - \frac{s_i}{n_i} \right) + \frac{n_i z_i^2}{p_i q_i^2} \left(q_i - \frac{s_i}{n_i} \right) \quad (2.3.22)$$

For the final three terms in (2.3.22), we can extract a common factor, i.e.,

$$\varsigma'_i = -\frac{n_i z_i^2}{p_i q_i} + \frac{n_i z'_i}{p_i q_i} \left(q_i - \frac{s_i}{n_i} \right) \left[\frac{z'_i}{z_i} - \frac{z_i}{p_i} + \frac{z_i}{q_i} \right] \quad (2.3.23)$$

We have values or expressions for each of the terms in (2.3.23) with the exception of z'_i . An equation for this term is easily obtained from (2.3.7). Observe that

$$z'_i = -\frac{t_i}{\sqrt{2\pi}} \exp\left(-\frac{t_i^2}{2}\right) = -t_i z_i \quad (2.3.24)$$

With the substitution of (2.3.23) and (2.3.24) into (2.3.16) through (2.3.18), we arrive at Garwood's formulas for the second partial derivatives of the log-likelihood function.

2.4 Bernoulli Trial with an Endpoint Probability

Garwood's method provides a set of exact formulas for the partial derivatives of the log-likelihood function with respect to parameters α and β . [11] We have successfully employed these formulas to solve test problems, particularly Garwood's examples. Unfortunately, problems arise when this method is applied to sensitivity tests that are comprised of a mixture of binomial and binary trials. For endpoint probabilities, those near zero and unity, terms such as ζ_i and ζ'_i become undefined due the presence of the factor $p_i q_i$ existing in the denominator. For endpoint probabilities, $p_i q_i = 0$; this situation is routinely encountered for levels in the explanatory variable characterized by a single Bernoulli trial (or its binomial equivalent). After observing this difficulty, we elected to revisit the derivation and search for an alternative form of derivatives $\partial L / \partial \alpha$, $\partial L / \partial \beta$, $\partial^2 L / \partial \alpha^2$, $\partial^2 L / \partial \alpha \partial \beta$ and $\partial^2 L / \partial \beta^2$. An alternative formulation is easier to control near the endpoints.

This derivation begins by recalling (2.1.3) for the log-likelihood function.

$$L = \sum_{i=1}^N \left\{ \ln \binom{n_i}{s_i} + (n_i - s_i) \ln(p_i) + s_i \ln(q_i) \right\} \quad (2.4.1)$$

As we differentiate with respect to α , we do so term by term.

$$\frac{\partial L}{\partial \alpha} = \frac{\partial}{\partial \alpha} \sum_{i=1}^N \left[\ln \binom{n_i}{s_i} + (n_i - s_i) \ln p_i + s_i \ln q_i \right] \quad (2.4.2)$$

$$\frac{\partial L}{\partial \alpha} = \sum_{i=1}^N \left[(n_i - s_i) \frac{\partial \ln p_i}{\partial \alpha} + s_i \frac{\partial \ln q_i}{\partial \alpha} \right] \quad (2.4.3)$$

$$\frac{\partial L}{\partial \alpha} = \sum_{i=1}^N \frac{\partial t_i}{\partial \alpha} \left[(n_i - s_i) \frac{1}{p_i} \frac{dp_i}{dt_i} + s_i \frac{1}{q_i} \frac{dq_i}{dt_i} \right] \quad (2.4.4)$$

This expression can be simplified by employing (2.3.7) and (2.3.8) to obtain

$$\frac{\partial L}{\partial \alpha} = \sum_{i=1}^N \left[(n_i - s_i) \frac{z_i}{p_i} - s_i \frac{z_i}{q_i} \right] \quad (2.4.5)$$

By similar means, we can show that

$$\frac{\partial L}{\partial \beta} = \sum_{i=1}^N x_i \left[(n_i - s_i) \frac{z_i}{p_i} - s_i \frac{z_i}{q_i} \right] \quad (2.4.6)$$

The form of equations (2.4.5) and (2.4.6) is quite interesting; the summand contains two terms, and each of these terms contains a ratio of the probability density function to its cumulative distribution function both evaluated at t_i (or x_i under inverse transformation). The behavior of the first (second) term $\sim z_i / p_i$ is questionable as $p_i \rightarrow 0$. Conversely, the behavior of the second term ($\sim z_i / q_i$) requires investigation as $p_i \rightarrow 1$ (or $q_i \rightarrow 0$). If we can assert that z_i / p_i tends to zero as $p_i \rightarrow 0$, we can exclude this term from (2.4.5) in the limit. This issue is addressed by the following proposition.

Proposition: *Let z be the normal probability density function, and let p be its cumulative distribution function both defined in the independent variable t . The ratio z/p is well defined and tends to zero as $p \rightarrow 0$.*

Proof: By examining equations (2.3.7) and (2.2.3), z and p jointly approach zero if and only if $t \rightarrow -\infty$. The ratio z/p may be written as follows.

$$\frac{z}{p}(t) = \frac{\exp(-t^2/2)}{\int_{-\infty}^t \exp(-\tilde{t}^2/2) d\tilde{t}} \quad (\text{P-1})$$

The Taylor series expansion for the exponential function,

$$\exp(x) = 1 + x + \frac{x^2}{2} + \dots = \sum_{i=0}^{\infty} \frac{x^i}{i!} \quad (\text{P-2})$$

is a convergent power series for $x < 0$. Hence, it follows that

$$\exp\left(-\frac{t^2}{2}\right) = 1 - \frac{t^2}{2 \cdot 1!} + \frac{t^4}{4 \cdot 2!} - \frac{t^6}{8 \cdot 3!} + \dots = \sum_{i=0}^n S_i + O(t^{2n}) \quad (\text{P-3})$$

where

$$S_i = \frac{(-1)^i t^{2i}}{2^i i!} \quad (\text{P-4})$$

is also a convergent power series for real valued t . To evaluate p , as occurs in (P-1), we compute the anti-derivative of (P-3) term by term, i.e.,

$$\int_{-\infty}^t \exp\left(-\frac{\tilde{t}^2}{2}\right) d\tilde{t} = t - \frac{t^3}{6 \cdot 1!} + \frac{t^5}{20 \cdot 2!} - \frac{t^7}{56 \cdot 3!} + \dots = t \sum_{i=0}^n \tilde{S}_i + O(t^{2n+1}) \quad (\text{P-5})$$

where

$$\tilde{S}_i = \frac{(-1)^i (t^2)^i}{2^i (2i+1) i!} \quad (\text{P-6})$$

The convergence of this series is confirmed by applying the ratio test. Observe that

$$\left| \frac{\tilde{S}_{i+1}}{\tilde{S}_i} \right| = \left| \frac{(-1)^{i+1} (t^2)^{i+1}}{2^{i+1} [2(i+1)+1] (i+1)!} \right| \cdot \left| \frac{2^i (2i+1) i!}{(-1)^i (t^2)^i} \right| = \frac{(2i+1) t^2}{2i^2 + 5i + 3} \quad (\text{P-7})$$

From this ratio, we conclude that

$$\lim_{i \rightarrow \infty} \left| \frac{\tilde{S}_{i+1}}{\tilde{S}_i} \right| = 0 \quad (\text{P-8})$$

Hence, the power series (P-5) converges for any value of t . Since both of the power series representations for the functions in (P-1) are convergent, both have finite limits, say $f(t)$ for the numerator and $g(t)$ for the denominator, so for any real t

$$\frac{z}{p} = \frac{f(t)}{\tilde{t} g(\tilde{t})^t} = \lim_{\tilde{t} \rightarrow -\infty} \frac{f(t)}{t g(t) - \tilde{t} g(\tilde{t})} = \frac{f(t)}{t g(t)} \quad (\text{P-9})$$

having invoked the properties of the cumulative distribution function as $t \rightarrow -\infty$. Since $f(t)$ and $g(t)$ are the finite limits of series that converge with the same order, we have that

$$\lim_{p \rightarrow 0} \frac{z}{p} = \lim_{p \rightarrow 0} \frac{f(t)}{t g(t)} = 0. \quad \square$$

By using a proof similar to that shown in the preceding proposition, the ratio z/q also approaches zero as $p \rightarrow 1$ ($q \rightarrow 0$). It follows that (2.4.5) and (2.4.6) can be expressed as

$$\frac{\partial L}{\partial \alpha} = \begin{cases} -\sum_{i=1}^N s_i z_i & p_i = 0 \\ \sum_{i=1}^N (n_i - s_i) z_i & q_i = 0 \end{cases} \quad (2.4.7)$$

and

$$\frac{\partial L}{\partial \beta} = \begin{cases} -\sum_{i=1}^N x_i s_i z_i & p_i = 0 \\ \sum_{i=1}^N x_i (n_i - s_i) z_i & q_i = 0 \end{cases} \quad (2.4.8)$$

By using the above equations for endpoint probabilities, $\partial L / \partial \alpha$ and $\partial L / \partial \beta$, derivatives of the log-likelihood function remain well defined for all possible values of p_i and q_i .

In addition to the Jacobian matrix, we must also use the matrix of second partial derivatives or Hessian matrix. The reason for its necessity will be made clear in the next section. We would

like to construct the second partial derivatives in terms of the ratios z_i / p_i and z_i / q_i . To derive

$\partial^2 L / \partial \alpha^2$, we differentiate (2.4.5), i.e.,

$$\frac{\partial^2 L}{\partial \alpha^2} = \sum_{i=1}^N \left[(n_i - s_i) \frac{\partial}{\partial \alpha} \left(\frac{z_i}{p_i} \right) - s_i \frac{\partial}{\partial \alpha} \left(\frac{z_i}{q_i} \right) \right] \quad (2.4.9)$$

With some careful mathematics and the use of (2.3.7) and (2.3.24), we can show that

$$\frac{\partial}{\partial \alpha} \left(\frac{z_i}{p_i} \right) = \frac{\partial t_i}{\partial \alpha} \frac{d}{dt_i} \left(\frac{z_i}{p_i} \right) = -t_i \left(\frac{z_i}{p_i} \right) - \left(\frac{z_i}{p_i} \right)^2 \quad (2.4.10)$$

$$\frac{\partial}{\partial \alpha} \left(\frac{z_i}{q_i} \right) = \frac{\partial t_i}{\partial \alpha} \frac{d}{dt_i} \left(\frac{z_i}{q_i} \right) = -t_i \left(\frac{z_i}{q_i} \right) + \left(\frac{z_i}{q_i} \right)^2 \quad (2.4.11)$$

After substituting (2.4.10) and (2.4.11) into (2.4.9), we obtain

$$\frac{\partial^2 L}{\partial \alpha^2} = \sum_{i=1}^N \left[(n_i - s_i) \left(\frac{z_i}{p_i} \right) \left(-t_i - \frac{z_i}{p_i} \right) - s_i \left(\frac{z_i}{q_i} \right) \left(-t_i + \frac{z_i}{q_i} \right) \right] \quad (2.4.12)$$

Note that (2.4.12) is characterized by the presence of the ratios z_i / p_i and z_i / q_i but no other terms containing z_i and p_i . The mixed partial derivative is developed as follows.

$$\frac{\partial^2 L}{\partial \alpha \partial \beta} = \sum_{i=1}^N \left[(n_i - s_i) \frac{\partial}{\partial \beta} \left(\frac{z_i}{p_i} \right) - s_i \frac{\partial}{\partial \beta} \left(\frac{z_i}{q_i} \right) \right] \quad (2.4.13)$$

By carefully differentiating and using (2.3.7) and (2.3.8), we can show that

$$\frac{\partial}{\partial \beta} \left(\frac{z_i}{p_i} \right) = \frac{\partial t_i}{\partial \beta} \frac{d}{dt_i} \left(\frac{z_i}{p_i} \right) = -x_i \left(t_i \left(\frac{z_i}{p_i} \right) + \left(\frac{z_i}{p_i} \right)^2 \right) \quad (2.4.14)$$

$$\frac{\partial}{\partial \beta} \left(\frac{z_i}{q_i} \right) = \frac{\partial t_i}{\partial \beta} \frac{d}{dt_i} \left(\frac{z_i}{q_i} \right) = -x_i \left(t_i \left(\frac{z_i}{q_i} \right) - \left(\frac{z_i}{q_i} \right)^2 \right) \quad (2.4.15)$$

By substituting (2.4.14) and (2.4.15) into (2.4.13), the mixed partial derivative is revealed as

$$\frac{\partial^2 L}{\partial \alpha \partial \beta} = - \sum_{i=1}^N x_i \left[(n_i - s_i) \left(\frac{z_i}{p_i} \right) \left(t_i + \frac{z_i}{p_i} \right) - s_i \left(\frac{z_i}{q_i} \right) \left(t_i - \frac{z_i}{q_i} \right) \right] \quad (2.4.16)$$

Note that the ratios found in (2.4.12) recur. If (2.4.6) is differentiated, the remaining second order partial derivative may be calculated.

$$\frac{\partial^2 L}{\partial \beta^2} = \frac{\partial}{\partial \beta} \left[\sum_{i=1}^N x_i \left\{ (n_i - s_i) \frac{z_i}{p_i} - s_i \frac{z_i}{q_i} \right\} \right] \quad (2.4.17)$$

By simplifying, we have that

$$\frac{\partial^2 L}{\partial \beta^2} = \sum_{i=1}^N x_i \left[(n_i - s_i) \frac{\partial}{\partial \beta} \left(\frac{z_i}{p_i} \right) - s_i \frac{\partial}{\partial \beta} \left(\frac{z_i}{q_i} \right) \right] \quad (2.4.18)$$

Substitution of (2.4.14) and (2.4.15), produces

$$\frac{\partial^2 L}{\partial \beta^2} = \sum_{i=1}^N x_i \left[(n_i - s_i) \left\{ -x_i \left(t_i \left(\frac{z_i}{p_i} \right) + \left(\frac{z_i}{p_i} \right)^2 \right) \right\} - s_i \left\{ -x_i \left(t_i \left(\frac{z_i}{q_i} \right) - \left(\frac{z_i}{q_i} \right)^2 \right) \right\} \right] \quad (2.4.19)$$

$$\frac{\partial^2 L}{\partial \beta^2} = - \sum_{i=1}^N x_i^2 \left[(n_i - s_i) \left(\frac{z_i}{p_i} \right) \left(t_i + \frac{z_i}{p_i} \right) - s_i \left(\frac{z_i}{q_i} \right) \left(t_i - \frac{z_i}{q_i} \right) \right] \quad (2.4.20)$$

To recap, this section presents an alternative formulation of Garwood's method that is useful in evaluating endpoint probabilities. At these endpoints, the second order partial derivatives may be computed as follows.

$$\frac{\partial^2 L}{\partial \alpha^2} = \begin{cases} \sum_{i=1}^N s_i z_i (t_i - z_i) & p_i = 0 \\ - \sum_{i=1}^N (n_i - s_i) z_i (t_i + z_i) & q_i = 0 \end{cases} \quad (2.4.21)$$

$$\frac{\partial^2 L}{\partial \alpha \partial \beta} = \begin{cases} \sum_{i=1}^N x_i s_i z_i (t_i - z_i) & p_i = 0 \\ - \sum_{i=1}^N x_i (n_i - s_i) z_i (t_i + z_i) & q_i = 0 \end{cases} \quad (2.4.22)$$

$$\frac{\partial^2 L}{\partial \beta^2} = \begin{cases} \sum_{i=1}^N x_i^2 s_i z_i (t_i - z_i) & p_i = 0 \\ -\sum_{i=1}^N x_i^2 (n_i - s_i) z_i (t_i + z_i) & q_i = 0 \end{cases} \quad (2.4.23)$$

For the endpoint probabilities ($p = 1, q = 0$) and ($p = 0, q = 1$), equations (2.4.7), (2.4.8), (2.4.21), (2.4.22) and (2.4.23) are instrumental in supporting the procedure for estimating μ and σ .

2.5 Newton's Solution Method

In Sections 2.3 and 2.4, it was demonstrated that we may fit sensitivity test data to the normal distribution through Fisher's Method of Maximum Likelihood.[13] Moreover, the procedure requires that we estimate two parameters, the distribution mean and standard deviation, μ and σ , respectively. To do so, we must determine α and β satisfying (2.3.3). Estimates of μ and σ are obtained by solving equations (2.3.3) cast either in Garwood's form (Section 2.3) or an alternative form (Section 2.4). The fitting procedure takes the form of an iterative scheme; rewrite α and β as follows.

$$\alpha = \alpha_0 + \Delta\alpha ; \quad \beta = \beta_0 + \Delta\beta \quad (2.5.1)$$

Substitute (2.5.1) into (2.3.3); when a Taylor series is constructed in variables α and β , observe that

$$\frac{\partial L}{\partial \alpha}(\alpha, \beta) = \frac{\partial L}{\partial \alpha}(\alpha_0, \beta_0) + \Delta\alpha \frac{\partial^2 L}{\partial \alpha^2}(\alpha_0, \beta_0) + \Delta\beta \frac{\partial^2 L}{\partial \alpha \partial \beta}(\alpha_0, \beta_0) + O(\Delta\alpha^2, \Delta\alpha \Delta\beta) \quad (2.5.2)$$

$$\frac{\partial L}{\partial \beta}(\alpha, \beta) = \frac{\partial L}{\partial \beta}(\alpha_0, \beta_0) + \Delta\alpha \frac{\partial^2 L}{\partial \alpha \partial \beta}(\alpha_0, \beta_0) + \Delta\beta \frac{\partial^2 L}{\partial \beta^2}(\alpha_0, \beta_0) + O(\Delta\alpha \Delta\beta, \Delta\beta^2) \quad (2.5.3)$$

An iterative estimation scheme may be retaining only the first order terms in (2.5.2) and (2.5.3).

In matrix form, we have that

$$\begin{bmatrix} \frac{\partial^2 L}{\partial \alpha^2}(\alpha_0, \beta_0) & \frac{\partial^2 L}{\partial \alpha \partial \beta}(\alpha_0, \beta_0) \\ \frac{\partial^2 L}{\partial \alpha \partial \beta}(\alpha_0, \beta_0) & \frac{\partial^2 L}{\partial \beta^2}(\alpha_0, \beta_0) \end{bmatrix} \begin{bmatrix} \Delta \alpha \\ \Delta \beta \end{bmatrix} = - \begin{bmatrix} \frac{\partial L}{\partial \alpha}(\alpha_0, \beta_0) \\ \frac{\partial L}{\partial \beta}(\alpha_0, \beta_0) \end{bmatrix} \quad (2.5.4)$$

The 2 x 2 matrix on the left side of (2.5.4) is $\mathbf{H}$, the Hessian matrix, and elements of the Jacobian matrix (a vector $\vec{J}$ in this case) resides on the right side. The pair (α_0, β_0) is designated as a starting guess; then the pair $(\Delta \alpha, \Delta \beta)$ is an increment that when added to (α_0, β_0) ostensibly forms an improved estimate, a new iterative value for (α, β) , and by transformation (2.3.2), a new estimate for (μ, σ) . The system (2.5.4) may be solved for the increments by inverting $\mathbf{H}$ and multiplying from the left. The inverse 2 x 2 matrix $\mathbf{H}^{-1}$ is easily obtained as long as

$$|\mathbf{H}(\alpha_0, \beta_0)| \neq 0 \quad (2.5.5)$$

The updated increment $(\Delta \alpha, \Delta \beta)$ is calculated from the equation

$$[\Delta \alpha, \Delta \beta]^T = \mathbf{H}^{-1}(\alpha_0, \beta_0) \bullet \vec{J}(\alpha_0, \beta_0) \quad (2.5.6)$$

It follows that

$$[\alpha, \beta]^T = [\alpha_0, \beta_0]^T + \mathbf{H}^{-1}(\alpha_0, \beta_0) \bullet \vec{J}(\alpha_0, \beta_0) \quad (2.5.7)$$

where $[\alpha, \beta]^T$ is a vector that contains the refined parameter estimates. As with other Newton approximation methods, (2.5.7) is easily structured as an iterative scheme. After a number of iterations, when this scheme has converged to the desired level of accuracy, the distribution mean and standard deviation can be computed from (2.3.2).

2.6 Minor Variant of the Golub-Grubbs Method

Garwood's method for probit analysis is highlighted in the preceding sections. This method is based upon a linear fit (2.3.1) of the probit. Recall that the probit is the argument to the normal probability integral [8]

$$p\left(\frac{x-\mu}{\sigma}\right) = \int_{-\infty}^{\frac{x-\mu}{\sigma}} \frac{1}{\sqrt{2\pi}} \exp\left(-\frac{t^2}{2}\right) dt \quad (2.6.1)$$

the probit being defined as

$$t = \frac{x-\mu}{\sigma} \quad (2.6.2)$$

The maximum likelihood method can be applied with the form (2.6.2). For this variant of the method, let p_i be the success probability while s_i is the number of successes out of n_i trials. In this case, the log-likelihood function is written as

$$L = \sum_{i=1}^N \left\{ \ln \binom{n_i}{s_i} + s_i \ln(p_i) + (n_i - s_i) \ln(q_i) \right\} \quad (2.6.3)$$

To determine the value of the explanatory variable x for which the log-likelihood function is maximized, it is necessary to differentiate (2.6.3) with respect to μ and σ , i.e.,

$$\frac{\partial L}{\partial \mu} = \sum_{i=1}^N s_i \frac{\partial}{\partial \mu} (\ln p_i) + (n_i - s_i) \frac{\partial}{\partial \mu} (\ln q_i) \quad (2.6.4)$$

$$\frac{\partial L}{\partial \sigma} = \sum_{i=1}^N s_i \frac{\partial}{\partial \sigma} (\ln p_i) + (n_i - s_i) \frac{\partial}{\partial \sigma} (\ln q_i) \quad (2.6.5)$$

Applying the chain rule along with the derivative of the logarithm,

$$\frac{\partial L}{\partial \mu} = \sum_{i=1}^N \frac{s_i}{p_i} \frac{\partial p_i}{\partial \mu} + \frac{(n_i - s_i)}{q_i} \frac{\partial p_i}{\partial \sigma} \quad (2.6.6)$$

$$\frac{\partial L}{\partial \sigma} = \sum_{i=1}^N \frac{s_i}{p_i} \frac{\partial p_i}{\partial \sigma} + \frac{(n_i - s_i)}{q_i} \frac{\partial p_i}{\partial \sigma} \quad (2.6.7)$$

To evaluate (2.6.4) and (2.6.5), the following additional derivatives are required. From (2.6.2), we have that

$$\frac{\partial t_i}{\partial \mu} = \frac{\partial}{\partial \mu} \left(\frac{x_i - \mu}{\sigma} \right) = -\frac{1}{\sigma} \quad (2.6.8)$$

$$\frac{\partial t_i}{\partial \sigma} = \frac{\partial}{\partial \sigma} \left(\frac{x_i - \mu}{\sigma} \right) = (x_i - \mu) \left(-\frac{1}{\sigma^2} \right) = -\frac{t_i}{\sigma} \quad (2.6.9)$$

Additional derivatives are shown in (2.3.7) and (2.3.24); therefore,

$$\frac{\partial p_i}{\partial \mu} = \frac{dp_i}{dt_i} \frac{\partial t_i}{\partial \mu} = -\frac{z_i}{\sigma} ; \quad \frac{\partial q_i}{\partial \mu} = \frac{z_i}{\sigma} \quad (2.6.10)$$

$$\frac{\partial p_i}{\partial \sigma} = \frac{dp_i}{dt_i} \frac{\partial t_i}{\partial \sigma} = -\frac{z_i t_i}{\sigma} ; \quad \frac{\partial q_i}{\partial \sigma} = \frac{z_i t_i}{\sigma} \quad (2.6.11)$$

Since $q_i = 1 - p_i$,

$$\frac{\partial q_i}{\partial \mu} = \frac{z_i}{\sigma} ; \quad \frac{\partial q_i}{\partial \sigma} = \frac{z_i t_i}{\sigma} \quad (2.6.12)$$

By substituting these derivatives into (2.6.6) and (2.6.7),

$$\frac{\partial L}{\partial \mu} = \frac{1}{\sigma} \sum_{i=1}^N \left(-s_i \frac{z_i}{p_i} + (n_i - s_i) \frac{z_i}{q_i} \right) \quad (2.6.13)$$

$$\frac{\partial L}{\partial \sigma} = \frac{1}{\sigma} \sum_{i=1}^N \left(-s_i t_i \frac{z_i}{p_i} + (n_i - s_i) t_i \frac{z_i}{q_i} \right) \quad (2.6.14)$$

A maximum occurs in the log-likelihood function for the pair (μ, σ) when

$$\frac{\partial L}{\partial \mu} = \frac{\partial L}{\partial \sigma} = 0 \quad (2.6.15)$$

It is also necessary to determine the second partial derivatives $\frac{\partial^2 L}{\partial \mu^2}$, $\frac{\partial^2 L}{\partial \mu \partial \sigma}$ and $\frac{\partial^2 L}{\partial \sigma^2}$. Based upon the ratios in (2.6.13) and (2.6.14), the following ordinary derivatives are needed.

$$\frac{d}{dt_i} \left(\frac{z_i}{p_i} \right) = \left(\frac{z_i}{p_i} \right)' = - \left(t_i \left(\frac{z_i}{p_i} \right) + \left(\frac{z_i}{p_i} \right)^2 \right) \quad (2.6.16)$$

$$\frac{d}{dt_i} \left(\frac{z_i}{q_i} \right) = \left(\frac{z_i}{q_i} \right)' = - \left(t_i \left(\frac{z_i}{q_i} \right) - \left(\frac{z_i}{q_i} \right)^2 \right) \quad (2.6.17)$$

$$\frac{d}{dt_i} \left(\frac{t_i z_i}{p_i} \right) = \left(\frac{t_i z_i}{p_i} \right)' = \frac{z_i}{p_i} - t_i^2 \left(\frac{z_i}{p_i} \right) - t_i \left(\frac{z_i}{p_i} \right)^2 \quad (2.6.18)$$

$$\frac{d}{dt_i} \left(\frac{t_i z_i}{q_i} \right) = \left(\frac{t_i z_i}{q_i} \right)' = \frac{z_i}{q_i} - t_i^2 \left(\frac{z_i}{q_i} \right) + t_i \left(\frac{z_i}{q_i} \right)^2 \quad (2.6.19)$$

Notice that the ratios z_i/p_i and z_i/q_i repeatedly occur in these derivatives. With all of these derivatives in mind, the second order partial derivatives may be derived as follows.

$$\frac{\partial^2 L}{\partial \mu^2} = \frac{1}{\sigma} \sum_{i=1}^N \left[-s_i \frac{\partial}{\partial \mu} \left(\frac{z_i}{p_i} \right) + (n_i - s_i) \frac{\partial}{\partial \mu} \left(\frac{z_i}{q_i} \right) \right] \quad (2.6.20)$$

$$\frac{\partial^2 L}{\partial \mu^2} = \frac{1}{\sigma} \sum_{i=1}^N \left[s_i \left(\frac{z_i}{p_i} \right)' - (n_i - s_i) \left(\frac{z_i}{q_i} \right)' \right] \frac{\partial t_i}{\partial \mu} \quad (2.6.21)$$

By substituting (2.6.8), (2.6.16) and (2.6.17),

$$\frac{\partial^2 L}{\partial \mu^2} = \frac{1}{\sigma^2} \sum_{i=1}^N \left[-s_i \left(t_i \left(\frac{z_i}{p_i} \right) + \left(\frac{z_i}{p_i} \right)^2 \right) + (n_i - s_i) \left(t_i \left(\frac{z_i}{q_i} \right) - \left(\frac{z_i}{q_i} \right)^2 \right) \right] \quad (2.6.22)$$

Similarly, the mixed partial derivative can be developed as follows. Differentiate (2.6.14) with respect to μ .

$$\frac{\partial^2 L}{\partial \mu \partial \sigma} = \frac{1}{\sigma} \sum_{i=1}^N \left[-s_i \frac{\partial}{\partial \mu} \left(\frac{t_i z_i}{p_i} \right) + (n_i - s_i) \frac{\partial}{\partial \mu} \left(\frac{t_i z_i}{q_i} \right) \right] \quad (2.6.23)$$

Hence,

$$\frac{\partial^2 L}{\partial \mu \partial \sigma} = \frac{1}{\sigma} \sum_{i=1}^N \left[-s_i \left(\frac{t_i z_i}{p_i} \right)' + (n_i - s_i) \left(\frac{t_i z_i}{q_i} \right)' \right] \frac{\partial t_i}{\partial \mu} \quad (2.6.24)$$

and by applying (2.6.8),

$$\frac{\partial^2 L}{\partial \mu \partial \sigma} = \frac{1}{\sigma^2} \sum_{i=1}^N \left[s_i \left(\frac{t_i z_i}{p_i} \right)' - (n_i - s_i) \left(\frac{t_i z_i}{q_i} \right)' \right] \quad (2.6.25)$$

Substituting (2.6.18) and (2.6.19), we obtain

$$\frac{\partial^2 L}{\partial \mu \partial \sigma} = \frac{1}{\sigma^2} \sum_{i=1}^N \left[s_i \left(\left(\frac{z_i}{p_i} \right) - t_i^2 \left(\frac{z_i}{p_i} \right) - t_i \left(\frac{z_i}{p_i} \right)^2 \right) - (n_i - s_i) \left(\left(\frac{z_i}{q_i} \right) - t_i^2 \left(\frac{z_i}{q_i} \right) + t_i \left(\frac{z_i}{q_i} \right)^2 \right) \right] \quad (2.6.26)$$

The final partial derivative requires a more tedious mathematical development. By differentiating (2.2.7), we obtain

$$\frac{\partial^2 L}{\partial \sigma^2} = \frac{\partial}{\partial \sigma} \left[\frac{1}{\sigma} \sum_{i=1}^N \left(-s_i t_i \left(\frac{z_i}{p_i} \right) + (n_i - s_i) t_i \left(\frac{z_i}{q_i} \right) \right) \right] \quad (2.6.27)$$

Application of the product rule produces

$$\begin{aligned} \frac{\partial^2 L}{\partial \sigma^2} = \frac{d}{d\sigma} \left(\frac{1}{\sigma} \right) \sum_{i=1}^N \left[-s_i t_i \left(\frac{z_i}{p_i} \right) + (n_i - s_i) t_i \left(\frac{z_i}{q_i} \right) \right] \\ + \frac{1}{\sigma} \sum_{i=1}^N \left[\left(-s_i \left(\frac{t_i z_i}{p_i} \right)' + (n_i - s_i) \left(\frac{t_i z_i}{q_i} \right)' \right) \right] \frac{\partial t_i}{\partial \sigma} \end{aligned} \quad (2.6.28)$$

Simplifying with the use of (2.6.9), we obtain

$$\begin{aligned} \frac{\partial^2 L}{\partial \sigma^2} = & \frac{1}{\sigma^2} \sum_{i=1}^N \left[s_i t_i \left(\frac{z_i}{p_i} \right) - (n_i - s_i) t_i \left(\frac{z_i}{q_i} \right) \right] \\ & + \frac{1}{\sigma^2} \sum_{i=1}^N t_i \left[s_i \left(\frac{t_i z_i}{p_i} \right)' - (n_i - s_i) \left(\frac{t_i z_i}{q_i} \right)' \right] \end{aligned} \quad (2.6.29)$$

By grouping similar terms,

$$\frac{\partial^2 L}{\partial \sigma^2} = \frac{1}{\sigma^2} \sum_{i=1}^N \left[s_i t_i \left(\frac{z_i}{p_i} + \left(\frac{t_i z_i}{p_i} \right)' \right) - (n_i - s_i) t_i \left(\frac{z_i}{q_i} + \left(\frac{t_i z_i}{q_i} \right)' \right) \right] \quad (2.6.30)$$

An iterative scheme may be constructed for this method in the same manner as was applied for Garwood's procedure (2.5.2) through (2.5.4). Equations (2.6.13), (2.6.14), (2.6.22), (2.6.26), and (2.6.30) can be used to populate the iterative linear system

$$\begin{bmatrix} \frac{\partial^2 L}{\partial \mu^2}(\mu_0, \sigma_0) & \frac{\partial^2 L}{\partial \mu \partial \sigma}(\mu_0, \sigma_0) \\ \frac{\partial^2 L}{\partial \mu \partial \sigma}(\mu_0, \sigma_0) & \frac{\partial^2 L}{\partial \sigma^2}(\mu_0, \sigma_0) \end{bmatrix} \begin{bmatrix} \Delta \mu \\ \Delta \sigma \end{bmatrix} = - \begin{bmatrix} \frac{\partial L}{\partial \mu}(\mu_0, \sigma_0) \\ \frac{\partial L}{\partial \sigma}(\mu_0, \sigma_0) \end{bmatrix} \quad (2.6.31)$$

where

$$\begin{aligned} \Delta \mu &= \mu - \mu_0 \\ \Delta \sigma &= \sigma - \sigma_0 \end{aligned} \quad (2.6.32)$$

The components of this system a common factor of $1/\sigma$. Since the absolute value of σ can be quite small, it is numerically advantageous to multiply the matrix and vector in (2.6.31) by σ^2 . Doing so has the effect of increasing the magnitude of the matrix components, particularly the diagonal entries. Accomplishing this change enhances the convergence of the iterative scheme.

2.7 Analysis of Confidence in the Estimates

One tool for analyzing the level of believability associated with the estimate of an unknown parameter is the confidence region.[14] The confidence region represents one of the more arcane concepts in statistics, so it is best illustrated for a single explanatory variable say, the “ideal” dosage for an antibiotic. When the ideal dosage is administered, it should manifest a concentration in living tissue that is sufficient for the eradication of the pathogenic bacteria. Given the level of uncertainty in dosage measurement, administration and effectiveness, we would like to estimate an effective range for dosage quantity and then to attach to it a probabilistic level of assurance. The dosage is a quantity represented in one dimension, so in terms of the success probability, its range is designated as a *confidence interval*, and the mathematical process used to determine the interval is known as an *interval estimate*. The end points of the interval act as *random variables* while the statistical parameter theoretically contained within the interval is regarded as *fixed*, but *unknown*. [13] The concept of the confidence interval can be confusing, so it is instructive to explore why this concept can present difficulties.

Consider a physical or biological process that exhibits some randomness or “noise”. Suppose that the behavior of this process can be characterized, in part, by an unknown, yet fixed *non-random* parameter θ . A good example of this parameter is the mean of a statistical distribution. The parameter exists, but its exact value cannot be known. Instead, θ must be estimated via experimentation, perhaps by using a series of independent experimental trials. To differentiate it from the unknown exact value, the estimated value of θ is denoted as $\hat{\theta}$. Based also on information provided by the experiments, an interval estimate $[\theta_L, \theta_H]$ associated with $\hat{\theta}$ may be calculated. For this interval, the investigator would compute a probability p for the interval and then make the following assertion.

CLAIM: With 100p % certainty, the exact value of the parameter θ is contained within $[\theta_L, \theta_H]$.

It is in consideration of this assertion that the difficulty arises. Although it sounds reasonable, this claim is false. Why? Procedurally, the interval endpoints (values of random variables) are calculated based upon the value of $\hat{\theta}$. Thus, at the conclusion of the experiments (a series of Bernoulli or binomial trials), $\hat{\theta}$ is computed; then θ_L and θ_H are calculated directly. As they pertain to a single set of experiments $\hat{\theta}$, θ_L and θ_H are fixed. Since they have been calculated, these numbers have no random behavior. Therefore, we cannot assign a probability to any of these values, so the claim made above cannot be true. For that reason, it is important to understand the true probabilistic interpretation for this situation.

Recall that in the mathematics of probability, the interval endpoints are represented by random variables. Denote these random variables as Θ_L and Θ_R . An initial test provides a first estimate for these values (and for $\hat{\theta}$). If we repeat the same experiment, that is, the series of trials, a number of times, we obtain a sequence of values for $\hat{\theta}$, Θ_L and for Θ_R . For any one of these tests, there is nothing random left to constitute a random interval endpoint. The probabilistic aspects of the interval emerge only when we view the series of repeated tests in a larger sense. That is, each test becomes a separate realization. In this sense, the interval endpoints fluctuate randomly. If 100 tests are performed, then we obtain 100 random intervals $[\Theta_L, \Theta_R]$. Therefore, randomness is effectively restored to the notion of the interval. With this interpretation, one may wonder as to how the faulty claim made above might be replaced by a statistically sound assertion. Suppose that we are interested in the parameter θ , and consider a 95% confidence interval associated with θ . The statistically sound assertion is made as follows.

PROPER CLAIM: Suppose that an unknown statistical parameter θ is estimated in 100 separate, repeatable experiments. If 95% confidence interval endpoints are calculated for each experiment, then the true value of the parameter θ will be contained in 95 of the 100 intervals. The probability that the true value of θ lies within the interval computed for an individual test is either one or zero.

In the paragraphs that follow, the mathematical development of a confidence interval is presented. This exposition is not rigorous, but the practical mathematics behind interval construction is emphasized. Elementary estimation theory is drawn upon as a core resource.

2.7.1 Success Ratio Confidence Interval

For the success ratio (or success probability) confidence interval, the success probability $\hat{p}$ is estimated. The associated interval $[p_L, p_R]$ is determined with a confidence level of C where

$$C = 100(1-\alpha)\%, \quad 0 < \alpha < 1$$

If 100 separate, but identically constructed tests are performed, then 100 estimates of $\hat{p}$ are computed along with 100 intervals $[p_L, p_R]$. Hence, by invoking the proper claim made in the previous section, $100(1-\alpha)$ of these intervals will contain the true success probability. The interval endpoints may be derived by using theory associated with the binomial distribution.[15] The starting point for this derivation is the DeMoivre-Laplace Theorem where it is assumed that n , the number of binomial trials conducted for the experiment, is large. Suppose that Y is the number of successful trials and that p is the unknown success probability; then the DeMoivre-Laplace theorem states that the ratio z defined by

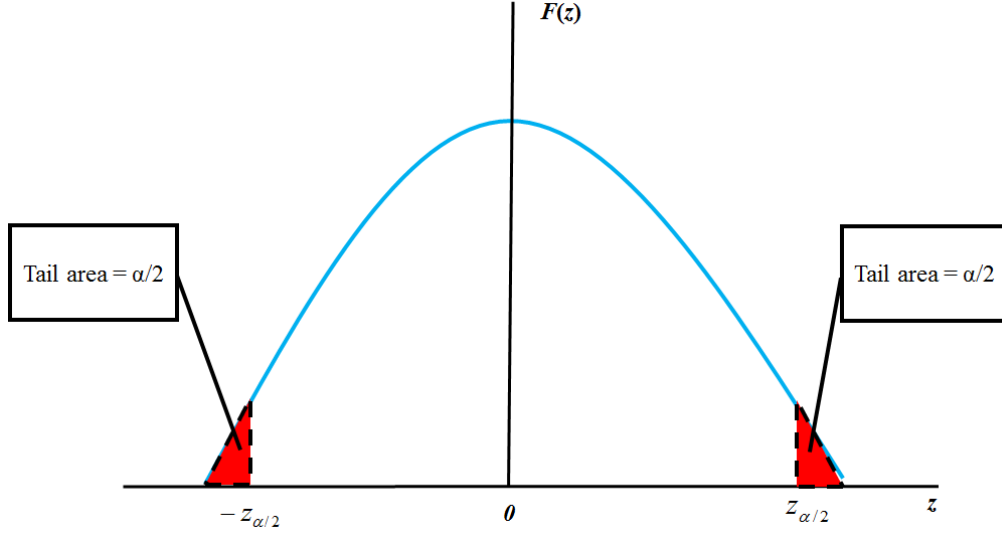


Figure 1 Probability Distribution Function for the Standard Normal Distribution

$$z = \frac{\frac{Y}{n} - p}{\sqrt{\frac{p(1-p)}{n}}} \quad (2.7.1.1)$$

tends toward the standard normal distribution $N(0,1)$, i.e.,

$$f(z) = \frac{1}{\sqrt{2\pi}} \exp\left(-\frac{z^2}{2}\right) \quad (2.7.1.2)$$

In terms of probability, this theorem implies that

$$P\left[-z_{\alpha/2} < \frac{\frac{Y}{n} - p}{\sqrt{\frac{p(1-p)}{n}}} < z_{\alpha/2}\right] \doteq 1 - \alpha \quad (2.7.1.3)$$

The similarity of (2.7.1.1) to (2.7.1.2) can be illustrated by realizing that for the binomial distribution, $\mu = np$ and $\sigma^2 = np(1-p)$. [15] With these substitutions, (2.7.1.1) can be rewritten as

$$z = \frac{Y - \mu}{\sigma} \quad (2.7.1.4)$$

This expression is identical to (2.2.2), the argument of the standard normal distribution. Denote $\hat{p} = Y / n$ as the estimate of the success probability; then (2.7.1.3) becomes

$$P \left[-z_{\alpha/2} < \frac{\hat{p} - p}{\sqrt{\frac{p(1-p)}{n}}} < z_{\alpha/2} \right] \doteq 1 - \alpha \quad (2.7.1.5)$$

$N(0,1)$ is a symmetric function about its zero mean, so a typical confidence interval $(-z_{\alpha/2}, z_{\alpha/2})$ is also symmetric about the mean. The probability associated with this interval is obtained by integrating the density function (2.6.1.2) over $(-z_{\alpha/2}, z_{\alpha/2})$. As is illustrated in Figure 1, the regions $(-\infty, -z_{\alpha/2})$ and $(z_{\alpha/2}, \infty)$ form the “tails” of the distribution. The tails are important when determining the probability or confidence coefficient associated with the interval. As is shown in Figure 1, the probability (area under the curve) for the tail regions is α . Then the confidence coefficient for the interval is defined as $100(1-\alpha)\%$. For example, consider a confidence coefficient of 95%; the attendant value of α is 0.05. It follows that the probability in each tail region is $\alpha/2 = 0.025$. The left interval endpoint on the explanatory z axis for $N(0,1)$ can now be determined by solving for $-z_{\alpha/2}$ in the cumulative distribution function equation

$$\frac{1}{\sqrt{2\pi}} \int_{-\infty}^{-z_{\alpha/2}} \exp\left(-\frac{z^2}{2}\right) dz = \frac{\alpha}{2} \quad (2.7.1.6)$$

Taking (2.6.1.5) into careful consideration, p is identified as the unknown probability while $\hat{p}$ is its estimate. By algebraically solving for p in (2.7.1.3), an interval estimate for p can be obtained consistent with the interval probability $1 - \alpha$. Collett [4] derives two expressions for the interval; the first is derived as follows. Begin by rewriting (2.6.1.5).

$$\frac{|\hat{p} - p|}{\sqrt{\frac{p(1-p)}{n}}} \leq z_{\alpha/2} \quad (2.7.1.7)$$

The bounds of the interval are identified by selecting the equality in (2.7.1.7); by rearranging terms, we obtain

$$(\hat{p} - p)^2 - z_{\alpha/2}^2 \left(\frac{p(1-p)}{n} \right) = 0 \quad (2.7.1.8)$$

Additional algebraic manipulation produces the quadratic equation

$$p^2 \left(1 + \frac{z_{\alpha/2}^2}{n} \right) - p \left(2\hat{p} + \frac{z_{\alpha/2}^2}{n} \right) + \hat{p}^2 = 0 \quad (2.7.1.9)$$

From an analysis of its form, (2.7.1.9) defines a parabola.[15] The roots of this equation delimit the endpoints of the confidence interval. Formulas for the endpoints are calculated by using the quadratic formula. After some simplification, we obtain

$$p = \frac{\hat{p} + \frac{z_{\alpha/2}^2}{2n} \pm z_{\alpha/2} \sqrt{\frac{\hat{p}(1-\hat{p})}{n} + \frac{z_{\alpha/2}^2}{4n^2}}}{1 + \frac{z_{\alpha/2}^2}{n}} \quad (2.7.1.10)$$

This formula contains the so-called standard error, $se(\hat{p})$ for the binomial distribution [4], i.e.,

$$se(\hat{p}) = \frac{\hat{p}(1-\hat{p})}{n} \quad (2.7.1.11)$$

With the use of (2.7.1.10) and (2.7.1.11), the interval endpoints p_{Low} and p_{High} are written as

$$p_{Low} = \frac{\hat{p} + \frac{z_{\alpha/2}^2}{2n} - z_{\alpha/2} \sqrt{se(\hat{p}) + \frac{z_{\alpha/2}^2}{4n}}}{1 + \frac{z_{\alpha/2}^2}{n}} \quad (2.7.1.12)$$

$$p_{High} = \frac{\hat{p} + \frac{z_{\alpha/2}^2}{2n} + z_{\alpha/2} \sqrt{\text{se}(\hat{p}) + \frac{z_{\alpha/2}^2}{4n}}}{1 + \frac{z_{\alpha/2}^2}{n}} \quad (2.7.1.13)$$

These formulas deserve added commentary. Recall that the confidence coefficient is based upon $N(0,1)$ and the symmetric interval $(-z_{\alpha/2}, z_{\alpha/2})$. The interval delineated by (2.7.1.12) and (2.7.1.13) is symmetric about the quantity

$$\tilde{p} = \frac{\hat{p} + \frac{z_{\alpha/2}^2}{2n}}{1 + \frac{z_{\alpha/2}^2}{n}} \quad (2.7.1.14)$$

where $\tilde{p}$ can be thought of as an observed and rather crude success ratio probability. It is interesting to note that $\tilde{p}$ does not lie at the center of the interval. Not only has this probability been translated off of the estimate, it has also been scaled. Recall that the DeMoivre-Laplace theorem states that the binomial distribution tends toward the normal distribution when n becomes large. Equation (2.7.1.14) is consistent with the conditions of this theorem. Note that $\tilde{p}$ approaches $\hat{p}$ as n becomes very large. Equations (2.7.1.12) through (2.7.1.14) delineate a first model (denoted as Model 1) for the endpoints of the confidence interval for the success probability.

The second model for the success probability is derived by recalling (2.7.1.7). In this case, the right hand side of the equation is evaluated by substituting $\hat{p}$, the estimate for the success probability. By rewriting this expression as an equality at the endpoint, we obtain

$$p - \hat{p} = \pm z_{\alpha/2} \sqrt{\frac{\hat{p}(1 - \hat{p})}{n}} \quad (2.7.1.15)$$

Hence, with the use of (2.7.1.11), we obtain the second model (denoted as Model 2), i.e.,

$$p_{Low} = \hat{p} - z_{\alpha/2} \text{se}(\hat{p}) \quad (2.7.1.16)$$

$$p_{High} = \hat{p} - z_{\alpha/2} \text{se}(\hat{p}) \quad (2.7.1.17)$$

Model 2 is more commonly used for delineating the endpoints of this confidence interval.

2.7.2 Success Probability Confidence Interval in the Explanatory Variable

For sensitivity testing, the binomial distribution remains a critical mathematical construct for determining the statistical distribution of the success probability for a random system whether it be physical or biological. However, this distribution alone lacks the ability to connect the system's response to an explanatory variable. This deficiency is remedied by replacing $\hat{p}$ with a probability calculated from the cumulative distribution function developed from the fitting procedure described in Section 2.3. This function depends upon the distribution produced by fitting to the test data. Therefore, it isomorphically maps to the explanatory variable through the distribution function. This result allows confidence intervals to be defined on the explanatory variable axis. By doing so, success probabilities can be calculated in terms of the explanatory variable for the standard normal distribution via (2.2.2) and (2.2.3). To illustrate this concept, let us revisit the example scenario described in Sections 2.7 and 2.7.1.

To recap, the issue under consideration is determining confidence intervals for drug efficacy. It is important to determine the ideal dosage required for an antibiotic to achieve a curative blood concentration with a chosen level of reliability. The dosage (perhaps expressed in milligrams) serves as the explanatory variable. In the determination of dosage-mortality curves, the natural logarithm of the dosage is sometimes used. Sensitivity tests are conducted at a set of discrete dosage levels in order to calibrate a normal distribution function for the drug's success. A fairly common practice is to estimate the drug dosage that will cure 99% of the test subjects.

Accordingly, dosage confidence intervals may be computed for a confidence coefficient of 95%. Starting with the theory presented in the preceding section, we can construct this interval.

Upon revisiting equations (2.7.1.12) and (2.7.1.13) for Model 1 or (2.7.1.16) and (2.7.1.17) for Model 2, note the presence of the term $\hat{p}$, an estimate of the exact success probability. Based upon our fit to the normal distribution, endpoint probabilities p_{Low} and p_{High} can be mapped onto corresponding arguments t_{Low} and t_{High} for the standard normal distribution $N(0,1)$, i.e.,

$$p_{Low} = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{t_{Low}} \exp\left(-\frac{t^2}{2}\right) dt \quad (2.7.2.1)$$

$$p_{High} = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{t_{High}} \exp\left(-\frac{t^2}{2}\right) dt \quad (2.7.2.2)$$

The arguments t_{Low} and t_{High} are functions of the endpoint probabilities and can be extracted from (2.7.2.1) and (2.7.2.2) via numerical solution procedures. The associated probits are related to the dosage interval endpoints x_{Low} and x_{High} through the equations

$$t_{Low} = \frac{x_{Low} - \mu}{\sigma}; \quad t_{High} = \frac{x_{High} - \mu}{\sigma} \quad (2.7.2.3)$$

where μ and σ are the maximum likelihood estimates for the mean and standard deviation obtained by the techniques described in Section 2.4. On the other hand, if a linear fitting form is used as in Section 2.3, the following forms are applied in lieu of (2.7.2.3).

$$t_{Low} = a + b x_{Low}; \quad t_{High} = a + b x_{High} \quad (2.7.2.4)$$

The values t_{Low} and t_{High} are implicit functions of their corresponding probabilities. In fact, within the context of archived literature, this function is denoted as the “Probit”. Therefore, (2.7.2.1) can

be expressed as

$$t_{Low} = \text{Probit} (p_{Low}) \quad (2.7.2.5)$$

A similar formula may easily be written for t_{High} .

2.7.3 Confidence Analysis for the Mean and Standard Deviation

The content of this section constitutes an upgrade to the corresponding sections in the predecessor of this report.[16] The analysis shown here is a modified version of the derivations presented in the DiDonato's excellent report.[17] As is true for the maximum likelihood analyses described above, the crux of the probit fit resides in the interplay between the binomial distribution (a discrete probability model) and the normal distribution (a continuous probability model). For this reason, when expected values are computed for the log-likelihood function, expectation operators must be applied for both the discrete and continuous distributions. To assess the level of confidence in the mean and standard deviation, begin with the binomial distribution, i.e.,

$$f_i(s; \mu, \sigma) = \binom{n_i}{s} p^{(n_i-s)} q^s \quad (2.7.3.1)$$

where $n = 0, 1, \dots, n$. Recall that probability $p(x; \mu, \sigma)$ is given by the normal distribution with $q = 1 - p$; note that x is the explanatory variable for the normal probabilities. The explanatory variable is not an active participant in this analysis. Define the functions G and H as follows.

$$G_i = \frac{\partial}{\partial \mu} [\ln f_i(s; \mu, \sigma)] \quad (2.7.3.2)$$

$$H_i = \frac{\partial}{\partial \sigma} [\ln f_i(s; \mu, \sigma)] \quad (2.7.3.3)$$

At this point, it is desirable to take the expected values of G and H . Begin by treating f as a continuous function of s ; although s takes on discrete values, we can regard s in the sense of a distribution. Hence,

$$E(G_i(s; \mu, \sigma)) = E_d \left(\int_s G_i(s; \mu, \sigma) f_i(s; \mu, \sigma) ds \right) \quad (2.7.3.4)$$

where E_d represents the discrete expectation. Substituting for G ,

$$E(G_i(s; \mu, \sigma)) = E_d \left(\int_s \frac{\partial}{\partial \mu} (\ln f_i(s; \mu, \sigma)) f_i(s; \mu, \sigma) ds \right) \quad (2.7.3.5)$$

The natural logarithm has special properties under differentiation, i.e.,

$$E(G_i(s; \mu, \sigma)) = E_d \left(\int_s \frac{1}{f_i(s; \mu, \sigma)} \frac{\partial f_i(s; \mu, \sigma)}{\partial \mu} f_i(s; \mu, \sigma) ds \right) \quad (2.7.3.6)$$

Clearly, we have that

$$E(G_i(s; \mu, \sigma)) = E_d \left(\int_s \frac{\partial f_i(s; \mu, \sigma)}{\partial \mu} ds \right) \quad (2.7.3.7)$$

Consider the integrand and perform the differentiation.

$$\frac{\partial f_i}{\partial \mu} = \frac{\partial}{\partial \mu} \left[\binom{n_i}{s} p^{n_i-s} q^s \right] = \binom{n_i}{s} \frac{\partial}{\partial \mu} [p^{n_i-s} q^s] \quad (2.7.3.8)$$

$$\frac{\partial f_i}{\partial \mu} = \binom{n_i}{s} \left[\frac{\partial}{\partial \mu} (p^{n_i-s}) q^s + p^{n_i-s} \frac{\partial}{\partial \mu} (q^s) \right] \quad (2.7.3.9)$$

By differentiating the probabilities, noting $q = 1 - p$ and collecting terms, we obtain

$$\frac{\partial f_i}{\partial \mu} = \binom{n_i}{s} p^{n_i-s-1} q^{s-1} [(n_i - s)q - sp] \frac{\partial p}{\partial \mu} \quad (2.7.3.10)$$

$$E_d \left(\frac{\partial f}{\partial \mu} \right) = \binom{n}{s} p^{E_d[n-s]-1} q^{E_d[s]-1} [E_d(n-s)q - E_d(s)p] \frac{\partial p}{\partial \mu} \quad (2.7.3.11)$$

By nature of the binomial distribution, both $(n - s)$ and S are discrete random variables. Moreover,

$$E_d(n_i - s) = np ; \quad E_d(s) = n_i q \quad (2.7.3.12)$$

Substitution in (2.7.3.10) yields

$$E_d\left(\frac{\partial f_i}{\partial \mu}\right) = \binom{n_i}{s} p^{n_i p - 1} q^{n_i q - 1} [n_i p q - n_i q p] \frac{\partial p}{\partial \mu} = 0 \quad (2.7.3.13)$$

By backing these results into (2.7.3.7), we can conclude that

$$E(G_i) = 0 \quad (2.7.3.14)$$

It is also necessary to determine the expected value of H_i , this calculation is easily accomplished because the only difference is that the derivative of f_i is taken with respect to σ instead of μ .

The process is the same, and the result corresponding to (2.7.3.14) is

$$E_d\left(\frac{\partial f_i}{\partial \sigma}\right) = \binom{n_i}{s} p^{n_i p - 1} q^{n_i q - 1} [n_i p q - n_i q p] \frac{\partial p}{\partial \sigma} = 0 \quad (2.7.3.15)$$

Hence,

$$E(H_i) = 0 \quad (2.7.3.16)$$

Mirroring the maximum likelihood analysis, we the variances for G_i and H_i , then the covariance of G_i and H_i . Begin with the variance of G_i , i.e.,

$$\text{Var}(G_i) = \sigma_{G_i}^2 \quad (2.7.3.17)$$

By definition with the use of (2.7.3.14),

$$\sigma_{G_i}^2 = E(G_i^2) - (E(G_i))^2 = E(G_i^2) \quad (2.7.3.18)$$

$$E(G_i^2) = \int_s G_i^2(s; \mu, \sigma) f_i(s; \mu, \sigma) ds \quad (2.7.3.19)$$

Substituting for G and then evaluating the derivative of the logarithm,

$$E(G_i^2) = \int_s \left(\frac{\partial}{\partial \mu} \ln f_i(s; \mu, \sigma) \right)^2 f_i(s; \mu, \sigma) ds \quad (2.7.3.20)$$

$$E(G_i^2) = \int_s \frac{1}{f_i^2(s; \mu, \sigma)} \left(\frac{\partial f_i(s; \mu, \sigma)}{\partial \mu} \right)^2 f_i(s; \mu, \sigma) ds \quad (2.7.3.21)$$

Thus, the initial result is

$$E(G_i^2) = \int_s \frac{1}{f_i} \left(\frac{\partial f_i}{\partial \mu} \right)^2 ds \quad (2.7.3.22)$$

Now consider the partial derivative with respect to μ of $E(G_i)$, i.e.,

$$\frac{\partial E(G_i)}{\partial \mu} = \frac{\partial}{\partial \mu} \int_s G_i(s; \mu, \sigma) f_i(s; \mu, \sigma) ds \quad (2.7.3.23)$$

Recall (2.7.3.14) and interchange the order of differentiation and integration; note that the domain of integration is finite, and the differentiation is with respect to a parameter not the variable of integration.

$$0 = \int_s \frac{\partial}{\partial \mu} [G_i(s; \mu, \sigma) f_i(s; \mu, \sigma)] ds \quad (2.7.3.24)$$

By applying the product rule,

$$0 = \int_s \frac{\partial G_i(s; \mu, \sigma)}{\partial \mu} f_i(s; \mu, \sigma) + G_i(s; \mu, \sigma) \frac{\partial f_i(s; \mu, \sigma)}{\partial \mu} ds \quad (2.7.3.25)$$

Using (2.7.3.2) to substitute for G_i ,

$$0 = \int_s \frac{\partial^2 \ln f_i}{\partial \mu^2} f_i(s; \mu, \sigma) ds + \int_s \frac{\partial \ln f_i}{\partial \mu} \frac{\partial f_i(s; \mu, \sigma)}{\partial \mu} ds \quad (2.7.3.26)$$

Rearrangement along with the properties of the logarithm yields

$$\int_s \frac{1}{f_i} \left(\frac{\partial f_i(s; \mu, \sigma)}{\partial \mu} \right)^2 ds = - \int_s \frac{\partial^2 \ln f_i}{\partial \mu^2} f_i(s; \mu, \sigma) ds \quad (2.7.3.27)$$

Upon substitution into (2.7.3.22) we learn that

$$E(G_i^2) = - \int_s \frac{\partial^2 \ln f_i}{\partial \mu^2} f_i(s; \mu, \sigma) ds \quad (2.7.3.28)$$

By definition of the expectation,

$$E(G_i^2) = -E\left(\frac{\partial^2 \ln f_i}{\partial \mu^2}\right) \quad (2.7.3.29)$$

The variance of H_i is computed via a similar analysis; the only real change is that we differentiate with respect to σ instead of μ . Therefore,

$$\sigma_{H_i}^2 = E(H_i^2) = -E\left(\frac{\partial^2 \ln f_i}{\partial \sigma^2}\right) \quad (2.7.3.30)$$

It is also necessary to determine a form for the covariance of G_i and H_i . This expression is derived as follows.

$$\text{Cov}(G_i, H_i) = \sigma_{G_i H_i} = \int_s \frac{\partial \ln f_i}{\partial \mu} \frac{\partial \ln f_i}{\partial \sigma} f_i(s; \mu, \sigma) ds \quad (2.7.3.31)$$

As before, the properties of the natural logarithm are applied.

$$\sigma_{G_i H_i} = \int_s \frac{1}{f_i} \frac{\partial f_i}{\partial \mu} \frac{1}{f_i} \frac{\partial f_i}{\partial \sigma} f_i(s; \mu, \sigma) ds = \int_s \frac{1}{f_i(s; \mu, \sigma)} \frac{\partial f_i}{\partial \mu} \frac{\partial f_i}{\partial \sigma} ds \quad (2.7.3.32)$$

Recall (2.7.3.2) and differentiate it with respect to σ , i.e.,

$$E(G_i) = \int_s G_i f_i(s; \mu, \sigma) ds \quad (2.7.3.33)$$

Using (2.7.3.14) and interchanging the order of integration and differentiation,

$$0 = \frac{\partial E(G_i)}{\partial \sigma} = \frac{\partial}{\partial \sigma} \int_s G_i f_i(s; \mu, \sigma) ds = \int_s \frac{\partial (G_i f_i(s; \mu, \sigma))}{\partial \sigma} ds \quad (2.7.3.34)$$

By applying the product rule,

$$0 = \int_s \frac{\partial G_i}{\partial \sigma} f_i(s; \mu, \sigma) + G_i(s; \mu, \sigma) \frac{\partial f_i}{\partial \sigma} ds \quad (2.7.3.35)$$

Substituting for G_i , we obtain

$$0 = \int_s \frac{\partial^2 \ln f_i}{\partial \mu \partial \sigma} f_i(s; \mu, \sigma) ds + \int_s \frac{\partial \ln f_i}{\partial \mu} \frac{\partial f_i}{\partial \sigma} ds \quad (2.7.3.36)$$

Again, the logarithm's properties are applied to obtain

$$0 = \int_s \frac{\partial^2 \ln f_i}{\partial \mu \partial \sigma} f_i(s; \mu, \sigma) ds + \int_s \frac{1}{f_i} \frac{\partial f_i}{\partial \mu} \frac{\partial f_i}{\partial \sigma} ds \quad (2.7.3.37)$$

Hence,

$$\int_s \frac{1}{f_i} \frac{\partial f_i}{\partial \mu} \frac{\partial f_i}{\partial \sigma} ds = - \int_s \frac{\partial^2 \ln f_i}{\partial \mu \partial \sigma} f_i(s; \mu, \sigma) ds \quad (2.7.3.38)$$

Substituting this result into (2.7.3.32), we conclude that

$$\sigma_{G_i H_i} = - \int_s \frac{\partial^2 \ln f_i}{\partial \mu \partial \sigma} f_i(s; \mu, \sigma) ds = -E \left(\frac{\partial^2 \ln f_i}{\partial \mu \partial \sigma} \right) \quad (2.7.3.39)$$

having used the definition of the expectation. Equation (2.7.3.32) demonstrates that this covariance is completely symmetric, i.e.,

$$\text{Cov}(G_i, H_i) = \sigma_{G_i H_i} = \sigma_{H_i G_i} = \text{Cov}(H_i, G_i) \quad (2.7.3.40)$$

This development is quite lengthy, but at this point, we can begin to roll up the results shown above into a form pointed toward a confidence assessment. Recall (2.7.3.1), the i^{th} factor of the likelihood function; it follows that the entire likelihood function is written as

$$F = \prod_{i=1}^N f_i(s; \mu, \sigma) \quad (2.7.3.41)$$

Recall that the index i ranges over the N levels taken on by the explanatory variable. Accordingly, the log-likelihood function is expressed as

$$L = \ln F = \sum_{i=1}^N \ln f_i(s; \mu, \sigma) \quad (2.7.3.42)$$

Equation (2.7.3.42) is an important result because, via the properties of the logarithm, the product of probabilities in the likelihood function has been transformed into a sum of N log-likelihood factors. Using (2.7.3.2) and (2.7.3.3), partial derivatives (taken with respect to the fitting parameters) of the log-likelihood function can be written as

$$\frac{\partial L}{\partial \mu} = \sum_{i=1}^N \frac{\partial \ln f_i(s; \mu, \sigma)}{\partial \mu} = \sum_{i=1}^N G_i(s; \mu, \sigma) \quad (2.7.3.43)$$

$$\frac{\partial L}{\partial \sigma} = \sum_{i=1}^N \frac{\partial \ln f_i(s; \mu, \sigma)}{\partial \sigma} = \sum_{i=1}^N H_i(s; \mu, \sigma) \quad (2.7.3.44)$$

These derivatives can be organized in the form of a vector

$$\left(\frac{\partial L}{\partial \mu}, \frac{\partial L}{\partial \sigma} \right) = \left(\sum_{i=1}^N G_i, \sum_{i=1}^N H_i \right) \quad (2.7.3.45)$$

For large sample sizes, that being defined by the n_i used at each level of the explanatory variable,

the vector (2.7.3.45) is distributed in multinormal distribution $N(\vec{0}, A^{-1})$ where

$$A^{-1} = \begin{bmatrix} E\left(-\frac{\partial^2 L}{\partial \mu^2}\right) & E\left(-\frac{\partial^2 L}{\partial \mu \partial \sigma}\right) \\ E\left(-\frac{\partial^2 L}{\partial \mu \partial \sigma}\right) & E\left(-\frac{\partial^2 L}{\partial \sigma^2}\right) \end{bmatrix} = \begin{bmatrix} \sum_{i=1}^N \sigma_{G_i}^2 & \sum_{i=1}^N \sigma_{G_i H_i} \\ \sum_{i=1}^N \sigma_{G_i H_i} & \sum_{i=1}^N \sigma_{H_i}^2 \end{bmatrix} = \begin{bmatrix} A_{\mu\mu} & A_{\mu\sigma} \\ A_{\mu\sigma} & A_{\sigma\sigma} \end{bmatrix} \quad (2.7.3.46)$$

A^{-1} is denoted as the covariance matrix.[17] Using this information, the departure of the parameter estimates from their true values can be expressed in the form of a Taylor cast in two variables, i.e.,

$$\frac{\partial L}{\partial \mu}(\bar{\mu}, \bar{\sigma}) = \frac{\partial L}{\partial \mu}(\mu_0, \sigma_0) + \frac{\partial^2 L}{\partial \mu^2}(\mu_1, \sigma_1)(\bar{\mu} - \mu_0) + \frac{\partial^2 L}{\partial \mu \partial \sigma}(\mu_1, \sigma_1)(\bar{\sigma} - \sigma_0) \quad (2.7.3.47)$$

$$\frac{\partial L}{\partial \sigma}(\bar{\mu}, \bar{\sigma}) = \frac{\partial L}{\partial \sigma}(\mu_0, \sigma_0) + \frac{\partial^2 L}{\partial \mu \partial \sigma}(\mu_2, \sigma_2)(\bar{\mu} - \mu_0) + \frac{\partial^2 L}{\partial \sigma^2}(\mu_2, \sigma_2)(\bar{\sigma} - \sigma_0) \quad (2.7.3.48)$$

These expressions are presented as equalities in deference to the standard form for the Taylor series. Instead of expanding the second order partial derivatives around the point (μ_0, σ_0) , they are centered at two points (μ_1, σ_1) and (μ_2, σ_2) that lie along the line segment between $(\bar{\mu}, \bar{\sigma})$ and (μ_0, σ_0) . Although these points are not known *a priori*, they both converge to $(\bar{\mu}, \bar{\sigma})$ with probability one. Moreover, the elements of (2.7.3.46) converge to the expected values with probability one.[17] At $(\bar{\mu}, \bar{\sigma})$, the first order partial derivatives vanish, so (2.7.3.47) and (2.7.3.48) are rewritten as

$$\frac{\partial L}{\partial \mu}(\mu_0, \sigma_0) = -\frac{\partial^2 L}{\partial \mu^2}(\mu_1, \sigma_1)(\bar{\mu} - \mu_0) - \frac{\partial^2 L}{\partial \mu \partial \sigma}(\mu_1, \sigma_1)(\bar{\sigma} - \sigma_0) \quad (2.7.3.49)$$

$$\frac{\partial L}{\partial \sigma}(\mu_0, \sigma_0) = -\frac{\partial^2 L}{\partial \mu \partial \sigma}(\mu_2, \sigma_2)(\bar{\mu} - \mu_0) - \frac{\partial^2 L}{\partial \sigma^2}(\mu_2, \sigma_2)(\bar{\sigma} - \sigma_0) \quad (2.7.3.50)$$

Finally, at this point of the development, we can address the level of confidence in the estimates $\bar{\mu}$ and $\bar{\sigma}$. To do so, it is instructive to write the linear system (2.7.3.49) and (2.7.3.50) in matrix-vector form. Define vectors as shown below:

$$\bar{\theta} = \begin{bmatrix} \bar{\mu} - \mu_0 \\ \bar{\sigma} - \sigma_0 \end{bmatrix}; \quad \bar{\lambda} = \begin{bmatrix} \frac{\partial L}{\partial \mu}(\mu_0, \sigma_0) \\ \frac{\partial L}{\partial \sigma}(\mu_0, \sigma_0) \end{bmatrix} \quad (2.7.3.51)$$

The (2.7.3.49) and (2.7.3.50) can be written as

$$\bar{\lambda} = A^{-1} \bullet \bar{\theta} \quad (2.7.3.52)$$

or

$$\bar{\theta} = A \bullet \bar{\lambda} \quad (2.7.3.53)$$

The vector $\vec{\theta}$ has an approximate bivariate multinormal distribution, that is $N(0, A)$. The matrix A is denoted as the asymptotic variance/covariance matrix. To fix a level of confidence to the estimate $(\bar{\mu}, \bar{\sigma})$, consider the quadratic form

$$Q = \vec{\theta}^T \bullet A^{-1} \bullet \vec{\theta} \quad (2.7.3.54)$$

For a vector $\vec{X}$ with distribution $N(v, V)$; the associated joint density function $f(X)$ has the general form

$$f(\vec{X}) = \frac{1}{(2\pi)^{M/2} |V|} \exp \left[-\frac{1}{2} (X - v) \bullet V^{-1} \bullet (X - v)^T \right] \quad (2.7.3.55)$$

This density function has a χ^2 (Chi-squared) distribution with m degrees of freedom where m is the number of components in the random vector X . [14] For our specific concern (2.7.3.54), there are two degrees of freedom: (μ, σ) for the Golub-Grubbs model and (α, β) for the Garwood model. For the former model with Q provided by the χ^2 distribution, (2.7.3.54) can be expanded to obtain

$$A_{\mu\mu}(\mu - \bar{\mu})^2 + 2A_{\mu\sigma}(\mu - \bar{\mu})(\sigma - \bar{\sigma}) + A_{\sigma\sigma}(\sigma - \bar{\sigma})^2 = \chi^2_{1-\gamma} \quad (2.7.3.56)$$

In this equation, the calculation is centered at the parameter pair estimate $(\bar{\mu}, \bar{\sigma})$, and so are the matrix elements $A_{\mu\mu}$, $A_{\mu\sigma}$ and $A_{\sigma\sigma}$. The locus defining the boundary for the confidence assessment is given by the ordered pairs (μ, σ) , a continuum. The position of this locus is controlled by the term on the right hand side, the tail of the χ^2 distribution for confidence level γ . The χ^2 term is given by

$$P(\chi^2 \leq 1 - \gamma) = \int_0^{\chi^2} \frac{1}{2} \exp\left(-\frac{w}{2}\right) dw \quad (2.7.3.57)$$

Hence,

$$\chi_{1-\gamma}^2 = -2\ln(1 - \Pr(\chi^2 \leq 1 - \gamma)) \quad (2.7.3.58)$$

As an illustration, for a confidence level of $\gamma = 0.95$, $1 - \gamma = 0.05$, the value of $\chi_{1-\gamma}^2$ is ~ 0.1025 . If one cares to apply some analytic geometry, (2.7.3.56) can be shown to be an ellipse in the plane. The ellipse can be rotated into a standard position, and the lengths of its semi-major and semi-minor axes can be calculated. Doing so is all well and good, and it is even instructive, but it adds complication through the rotation. Instead, we can simply solve the quadratic equation for one parameter of the other. If we solve for σ , we obtain

$$\sigma = \bar{\sigma} - \frac{A_{\mu\mu}}{A_{\sigma\sigma}}(\mu - \bar{\mu}) \pm \sqrt{\frac{\chi_{1-\gamma}^2}{A_{\sigma\sigma}} + \left(\frac{A_{\mu\sigma}^2}{A_{\sigma\sigma}^2} - \frac{A_{\mu\mu}}{A_{\sigma\sigma}}\right)(\mu - \bar{\mu})^2} \quad (2.7.3.59)$$

To plot this ellipse, we must determine the domain of the independent variable μ . To do so, the argument of the radical must be non-negative, i.e.,

$$\frac{\chi_{1-\gamma}^2}{A_{\sigma\sigma}} + \left(\frac{A_{\mu\sigma}^2}{A_{\sigma\sigma}^2} - \frac{A_{\mu\mu}}{A_{\sigma\sigma}}\right)(\mu - \bar{\mu})^2 \geq 0 \quad (2.7.3.60)$$

The endpoints of the interval for μ are determined by solving (2.7.3.60) for the equality. These endpoints are:

$$\mu_{\min} = \bar{\mu} - \sqrt{\frac{\chi_{1-\gamma}^2 A_{\sigma\sigma}}{A_{\mu\sigma}^2 - A_{\mu\mu} A_{\sigma\sigma}}} \quad (2.7.3.61)$$

$$\mu_{\max} = \bar{\mu} + \sqrt{\frac{\chi_{1-\gamma}^2 A_{\sigma\sigma}}{A_{\mu\sigma}^2 - A_{\mu\mu} A_{\sigma\sigma}}} \quad (2.7.3.62)$$

Therefore, $\mu \in [\mu_{\min}, \mu_{\max}]$. Within this domain, (2.7.3.59) defines two branches of the ellipse in σ that are easily plotted so long as the radical is made to vanish at the endpoints. Note that floating point arithmetic tends to create extremely tiny negative numbers under the radical. The argument must be set to zero in these cases.

After the pages and pages of mathematics required to reach this point in our development, it is of great importance to interpret the meaning of this confidence assessment. The confidence ellipse exists in the parameters plane whereas the confidence interval consists of only 2 discrete points. DiDonato [17] claims that the true parameter pair lies within the ellipse with a probability of γ . This claim lies in disagreement with the frequentist interpretation of probability. At the conclusion of a single experiment consisting of n levels in an explanatory variable, a single parameter pair is calculated; then the elliptical locus is calculated. These calculations are entirely deterministic; there is nothing random about them. The only way that randomness is restored is through repeated experiments. That is to say, the entire experiment must be repeated say, M times. Each repeat of the experiment, becomes a separate realization of the experiment. Moreover, each repeat of the experiment may use different levels and number of levels; it may also use different numbers of trials at each level. In this sense, randomness is restored to the ensemble of experiments and to the locus of points constituting the confidence ellipse. With this interpretation, if 10 experiments are performed with the subsequent calculation of 90% confidence ellipses, then 9 of the 10 ellipses will contain the true parameter pair. For any of the individual experiments, then the probability that the experiment's confidence ellipse contains the true parameter pair is either one or zero. Taken by itself, the single experiment's ellipse only provides the analyst with a sense of the extent of scatter in the parameter pair. Scheffe [17] echoes this interpretation while also citing

the former claim, but that latter interpretation is rigorous from frequentist theory and is adhered to in this report.

2.8 Measures of Variance

In this section, we address measures of variance associated with our fitting procedure. As it happens, estimates of the mean μ and standard deviation σ for the probability distribution, have their own variances $\text{Var}(\mu)$ and $\text{Var}(\sigma)$. Also, we can imagine that μ and σ vary in a correlated manner in the probability distribution. This eventuality implies that the covariance of μ and σ may be non-zero. As it happens, these estimated variance values are given as entries in A^{var} , the asymptotic variance-covariance matrix, the inverse of the information matrix. The information matrix A_{info} is given by

$$A_{\text{info}} = -E \begin{bmatrix} \frac{\partial^2 L}{\partial \mu^2} & \frac{\partial^2 L}{\partial \mu \partial \sigma} \\ \frac{\partial^2 L}{\partial \mu \partial \sigma} & \frac{\partial^2 L}{\partial \sigma^2} \end{bmatrix} \quad (2.8.1)$$

The asymptotic variance-covariance matrix is the inverse of the information matrix, i.e.,

$$A^{\text{var}} = \begin{bmatrix} -E \left(\frac{\partial^2 L}{\partial \mu^2} \right) & -E \left(\frac{\partial^2 L}{\partial \mu \partial \sigma} \right) \\ -E \left(\frac{\partial^2 L}{\partial \mu \partial \sigma} \right) & -E \left(\frac{\partial^2 L}{\partial \sigma^2} \right) \end{bmatrix}^{-1} \quad (2.8.2)$$

The second partial derivatives are calculated during the estimation procedure, as in Section 2.6. It is easy to compute their expected values by using properties of the binomial probability distribution. The elements of the asymptotic variance-covariance matrix are interpreted as follows.

$$A^{\text{var}} = \begin{bmatrix} \text{Var}(\mu) & \text{Cov}(\mu, \sigma) \\ \text{Cov}(\mu, \sigma) & \text{Var}(\sigma) \end{bmatrix} \quad (2.8.3)$$

This matrix contains the desired variance estimates for the parameters produced the fitting procedure.

On the other hand, suppose that the linear fitting form is used as in Section 2.3, i.e.,

$$t = \alpha + \beta x \quad (2.8.4)$$

The information matrix elements cast in terms of α and β can, by transformation, be expressed in terms of μ and σ . The information matrix expressed in terms of α and β is written as

$$A_{\text{info}} = -E \begin{bmatrix} \frac{\partial^2 L}{\partial \alpha^2} & \frac{\partial^2 L}{\partial \alpha \partial \beta} \\ \frac{\partial^2 L}{\partial \alpha \partial \beta} & \frac{\partial^2 L}{\partial \beta^2} \end{bmatrix} \quad (2.8.5)$$

Recall that the transformation (2.3.2) has the general form

$$\alpha = \alpha(\mu, \sigma); \quad \beta = \beta(\mu, \sigma) \quad (2.8.6)$$

By using the chain rule for partial differentiation, we observe that

$$\frac{\partial}{\partial \mu} = \frac{\partial \alpha}{\partial \mu} \frac{\partial}{\partial \alpha} + \frac{\partial \beta}{\partial \mu} \frac{\partial}{\partial \beta} \quad (2.8.7)$$

$$\frac{\partial}{\partial \sigma} = \frac{\partial \alpha}{\partial \sigma} \frac{\partial}{\partial \alpha} + \frac{\partial \beta}{\partial \sigma} \frac{\partial}{\partial \beta} \quad (2.8.8)$$

It is a tedious mathematical exercise, but the repeated application of (2.8.7) and (2.8.8) leads to

$$\begin{aligned} \frac{\partial^2 L}{\partial \mu^2} = & \left[\frac{\partial \alpha}{\partial \mu} \frac{\partial^2 L}{\partial \alpha^2} + \frac{\partial \beta}{\partial \mu} \frac{\partial^2 L}{\partial \alpha \partial \beta} \right] \frac{\partial \alpha}{\partial \mu} + \frac{\partial L}{\partial \alpha} \frac{\partial^2 \alpha}{\partial \mu^2} \\ & + \left[\frac{\partial \alpha}{\partial \mu} \frac{\partial^2 L}{\partial \alpha \partial \beta} + \frac{\partial \beta}{\partial \mu} \frac{\partial^2 L}{\partial \beta^2} \right] \frac{\partial \beta}{\partial \mu} + \frac{\partial L}{\partial \beta} \frac{\partial^2 \beta}{\partial \mu^2} \end{aligned} \quad (2.8.9)$$

$$\begin{aligned} \frac{\partial^2 L}{\partial \mu \partial \sigma} = & \left[\frac{\partial \alpha}{\partial \sigma} \frac{\partial^2 L}{\partial \alpha^2} + \frac{\partial \beta}{\partial \sigma} \frac{\partial^2 L}{\partial \alpha \partial \beta} \right] \frac{\partial \alpha}{\partial \mu} + \frac{\partial L}{\partial \alpha} \frac{\partial^2 \alpha}{\partial \mu \partial \sigma} \\ & + \left[\frac{\partial \alpha}{\partial \sigma} \frac{\partial^2 L}{\partial \alpha \partial \beta} + \frac{\partial \beta}{\partial \sigma} \frac{\partial^2 L}{\partial \beta^2} \right] \frac{\partial \beta}{\partial \mu} + \frac{\partial L}{\partial \beta} \frac{\partial^2 \beta}{\partial \mu \partial \sigma} \end{aligned} \quad (2.8.10)$$

$$\begin{aligned} \frac{\partial^2 L}{\partial \sigma^2} = & \left[\frac{\partial \alpha}{\partial \sigma} \frac{\partial^2 L}{\partial \alpha^2} + \frac{\partial \beta}{\partial \sigma} \frac{\partial^2 L}{\partial \alpha \partial \beta} \right] \frac{\partial \alpha}{\partial \sigma} + \frac{\partial L}{\partial \alpha} \frac{\partial^2 \alpha}{\partial \sigma^2} \\ & + \left[\frac{\partial \alpha}{\partial \sigma} \frac{\partial^2 L}{\partial \alpha \partial \beta} + \frac{\partial \beta}{\partial \sigma} \frac{\partial^2 L}{\partial \beta^2} \right] \frac{\partial \beta}{\partial \sigma} + \frac{\partial L}{\partial \beta} \frac{\partial^2 \beta}{\partial \sigma^2} \end{aligned} \quad (2.8.11)$$

Equations (2.8.9) through (2.8.11) are the elements of A_{info} , and they are easily computed with the use of the following coordinate transformation derivatives.

$$\frac{\partial \alpha}{\partial \mu} = -\frac{1}{\sigma}; \quad \frac{\partial \alpha}{\partial \sigma} = \frac{\mu}{\sigma^2}; \quad \frac{\partial \beta}{\partial \mu} = 0; \quad \frac{\partial \beta}{\partial \sigma} = -\frac{1}{\sigma^2} \quad (2.8.12)$$

$$\frac{\partial^2 \alpha}{\partial \mu^2} = 0; \quad \frac{\partial^2 \alpha}{\partial \mu \partial \sigma} = \frac{1}{\sigma^2}; \quad \frac{\partial^2 \alpha}{\partial \sigma^2} = -\frac{2\mu}{\sigma^3} \quad (2.8.13)$$

$$\frac{\partial^2 \beta}{\partial \mu^2} = 0; \quad \frac{\partial^2 \beta}{\partial \mu \partial \sigma} = 0; \quad \frac{\partial^2 \beta}{\partial \sigma^2} = \frac{2}{\sigma^3} \quad (2.8.14)$$

Equations (2.8.9) through (2.8.14) provide the required information matrix entries needed in order to compute the asymptotic variance-covariance matrix. For purposes of calculation, (2.8.9) through (2.8.11) can be simplified by noting that $\partial L / \partial \alpha$ and $\partial L / \partial \beta$ are zero at the point of maximum likelihood.

2.9 Logit Fitting Procedure

The probit analysis discussed earlier in this report is based upon the normal distribution. In some cases, it is desirable to use a distribution that performs a bit differently in the tail region. The

logistic distribution is often applied for linear fitting. This distribution is computed as the natural logarithm of the odds ratio, i.e.,

$$\text{logit}(p) = \ln\left(\frac{p}{1-p}\right) = \alpha + \beta x \quad (2.9.1)$$

In this formula, p is the success probability, so the odds ratio is clearly present. The explanatory variable is x while α and β are the unknown parameters. Rearrangement of (2.9.1) allows calculation of the probability; specifically,

$$P(x; \alpha, \beta) = \frac{\exp(\alpha + \beta x)}{1 + \exp(\alpha + \beta x)} \quad (2.9.2)$$

Note, unlike the probit method, the computation of this probability does not require the use of numerical quadrature. The probability density function is calculated by differentiating (2.9.2) with respect to x . Hence,

$$f(x) = \frac{d p(x; \alpha, \beta)}{dx} = \frac{\beta \exp(\alpha + \beta x)}{[1 + \exp(\alpha + \beta x)]^2} \quad (2.9.3)$$

To elucidate some of the properties of this distribution, we can substitute $t = -\alpha - \beta x$ then rewrite (2.9.2) as

$$\ln\left(\frac{p}{1-p}\right) = -t \quad (2.9.4)$$

$$p(t) = \frac{\exp(-t)}{1 + \exp(-t)} \quad (2.9.5)$$

The complementary probability is

$$q(t) = \frac{1}{1 + \exp(-t)} \quad (2.9.6)$$

For application of this distribution, t can be substituted for $-t$. However, for the proof below, (2.9.5) must be left unaltered. The expected value for this distribution is zero, i.e., $E(t) = 0$. This property is shown by following [18]. By definition,

$$E(t) = \int_{-\infty}^{\infty} t f(t) dt = \int_{-\infty}^{\infty} \frac{t \exp(-t)}{[1 + \exp(-t)]^2} dt \quad (2.9.7)$$

This integral can be rewritten as follows.

$$E(t) = \int_{-\infty}^0 \frac{t \exp(-t)}{[1 + \exp(-t)]^2} dt + \int_0^{\infty} \frac{t \exp(-t)}{[1 + \exp(-t)]^2} dt \quad (2.9.8)$$

In the leftmost integral, make the substitution $t = -y$.

$$E(t) = \int_0^{\infty} \frac{-y \exp(y)}{[1 + \exp(y)]^2} dy + \int_0^{\infty} \frac{t \exp(-t)}{[1 + \exp(-t)]^2} dt \quad (2.9.9)$$

By expanding the quadratic term, we obtain

$$E(t) = \int_0^{\infty} \frac{-y \exp(y)}{1 + 2 \exp(y) + \exp(2y)} dy + \int_0^{\infty} \frac{t \exp(-t)}{1 + 2 \exp(-t) + \exp(-2t)} dt \quad (2.9.10)$$

If we multiply the integrand of the left integral by $\exp(-y)/\exp(-y)$ and the integrand of the right integral by $\exp(t)/\exp(t)$, we have that

$$E(t) = \int_0^{\infty} \frac{-z}{\exp(-z) + 2 + \exp(z)} dz + \int_0^{\infty} \frac{z}{\exp(z) + 2 + \exp(-z)} dz \quad (2.9.11)$$

Since z is a dummy variable of integration, these two integrals are of the same magnitude but opposite sign; thus,

$$E(t) = 0 \quad (2.9.12)$$

For the linear form shown in (2.9.1), (2.9.11) implies that

$$E(\alpha + \beta x) = -\frac{\alpha}{\beta} \quad (2.9.13)$$

The variance for this distribution can also be calculated. By definition, and with the use of (2.9.12),

$$\text{Var}(t) = E(t^2) - (E(t))^2 = E(t^2) \quad (2.9.14)$$

This expectation may be evaluated in the same manner. Observe that

$$E(t^2) = \int_{-\infty}^{\infty} \frac{t^2 \exp(-t)}{[1 + \exp(-t)]^2} dt = \int_{-\infty}^0 \frac{t^2 \exp(-t)}{[1 + \exp(-t)]^2} dt + \int_0^{\infty} \frac{t^2 \exp(-t)}{[1 + \exp(-t)]^2} dt \quad (2.9.15)$$

Again, the substitution $t = -y$ is employed in the integral over negative t , i.e.,

$$\int_{-\infty}^0 \frac{t^2 \exp(-t)}{[1 + \exp(-t)]^2} dt = \int_0^{\infty} \frac{(-y)^2 \exp(y)}{[1 + \exp(y)]^2} dy = \int_0^{\infty} \frac{y^2 \exp(y)}{[\exp(y)(\exp(-y) + 1)]^2} dy \quad (2.9.16)$$

Simplification shows that

$$\int_{-\infty}^0 \frac{t^2 \exp(-t)}{[1 + \exp(-t)]^2} dt = \int_0^{\infty} \frac{y^2 \exp(-y)}{[(\exp(-y) + 1)]^2} dy \quad (2.9.17)$$

This integral is the same as the remaining integral over positive t in (2.9.15), so

$$E(t^2) = 2 \int_0^{\infty} \frac{t^2 \exp(-t)}{[1 + \exp(-t)]^2} dt \quad (2.9.18)$$

To evaluate this integral, let $y = \exp(-t)$, with this substitution the integrand in (2.9.18) becomes

$$\frac{y}{(1+y)^2} = \sum_{n=1}^{\infty} (-1)^{n-1} n y^n \quad (2.9.19)$$

The series expansion appearing in (2.9.18) is derived by applying polynomial division to

$1/(1+y)^2$ and then multiplying by the z remaining in the numerator. Therefore,

$$\frac{\exp(-t)}{[1 + \exp(-t)]^2} = \sum_{n=1}^{\infty} (-1)^{n-1} n (\exp(-t))^n = \sum_{n=1}^{\infty} (-1)^{n-1} n \exp(-nt) \quad (2.9.20)$$

Substituting this result into (2.9.17),

$$E(t^2) = 2 \sum_{n=1}^{\infty} (-1)^{n-1} n \int_0^{\infty} t^2 \exp(-nt) dt \quad (2.9.21)$$

Integration by parts can be used to evaluate the integral in (2.9.20), i.e.,

$$\int_0^{\infty} t^2 \exp(-nt) dt = \frac{2}{n^3} \quad (2.9.22)$$

With this result, we learn that

$$E(t^2) = 2 \sum_{n=1}^{\infty} (-1)^{n-1} n \left(\frac{2}{n^3} \right) = 4 \sum_{n=1}^{\infty} \frac{(-1)^{n-1}}{n^2} \quad (2.9.23)$$

The convergent series in (2.9.22) evaluates to a value of $\pi^2 / 12$, so

$$E(t^2) = \frac{\pi^2}{3} \quad (2.9.24)$$

If we return this result to the original logistic form (2.9.3), we see that

$$\text{Var}(\alpha + \beta x) = E[(\alpha + \beta x)^2] = \frac{\beta^2 \pi}{3} \quad (2.9.25)$$

Equations (2.9.12) and (2.9.24) convey basic properties of the logistic distribution. It can be employed for analyzing sensitivity test data. The likelihood equation, applied again, is

$$l = \prod_{i=1}^N \binom{n_i}{s_i} p_i^{s_i} q_i^{n_i - s_i} \quad (2.9.26)$$

so the log-likelihood function is

$$L = \sum_{i=1}^N \binom{n_i}{s_i} + s_i \ln p_i + (n_i - s_i) \ln q_i \quad (2.9.27)$$

The log-likelihood function is maximized by computing the derivatives

$$\frac{\partial L}{\partial \alpha} = \sum_{i=1}^N \frac{s_i}{p_i} \frac{\partial p_i}{\partial \alpha} + \frac{(n_i - s_i)}{q_i} \frac{\partial q_i}{\partial \alpha} \quad (2.9.28)$$

$$\frac{\partial L}{\partial \beta} = \sum_{i=1}^N \frac{s_i}{p_i} \frac{\partial p_i}{\partial \beta} + \frac{(n_i - s_i)}{q_i} \frac{\partial q_i}{\partial \beta} \quad (2.9.29)$$

For this application, the success probability has the form

$$p_i = \frac{\exp(t_i)}{1 + \exp(t_i)} \quad (2.9.30)$$

For readers who have read through previous discussions in this report, the terms in the two equations above will appear very familiar, and indeed, they are. For $t_i = \alpha + \beta x_i$, we can use the derivatives provided in Section 3. It is particularly important to note that $dp_i / dt_i = z_i$, i.e.,

$$z_i = \frac{\exp(t_i)}{(1 + \exp(t_i))^2} \quad (2.9.31)$$

With these derivatives in mind, (2.9.27) and (2.9.28) are rewritten as

$$\frac{\partial L}{\partial \alpha} = \sum_{i=1}^N s_i \frac{z_i}{p_i} - (n_i - s_i) \frac{z_i}{q_i} \quad (2.9.32)$$

$$\frac{\partial L}{\partial \beta} = \sum_{i=1}^N s_i x_i \frac{z_i}{p_i} - (n_i - s_i) x_i \frac{z_i}{q_i} \quad (2.9.33)$$

The attendant second partial derivatives are expressed as

$$\frac{\partial^2 L}{\partial \alpha^2} = \sum_{i=1}^N s_i \left(\frac{z_i}{p_i} \right)' - (n_i - s_i) \left(\frac{z_i}{q_i} \right)' \quad (2.9.34)$$

$$\frac{\partial^2 L}{\partial \alpha \partial \beta} = \sum_{i=1}^N s_i x_i \left(\frac{z_i}{p_i} \right)' - (n_i - s_i) x_i \left(\frac{z_i}{q_i} \right)' \quad (2.9.35)$$

$$\frac{\partial^2 L}{\partial \beta^2} = \sum_{i=1}^N s_i x_i^2 \left(\frac{z_i}{p_i} \right)' - (n_i - s_i) x_i^2 \left(\frac{z_i}{q_i} \right)' \quad (2.9.36)$$

The prime ' notation represents the ordinary derivative d / dt_i . The derivatives of ratios occurring above are

$$\left(\frac{z_i}{p_i} \right)' = \frac{d}{dt_i} \left(\frac{z_i}{p_i} \right) = \frac{p_i z_i' - p_i' z_i}{p_i^2} = \frac{z_i'}{p_i} - \frac{z_i^2}{p_i^2} \quad (2.9.37)$$

$$\left(\frac{z_i}{q_i}\right)' = \frac{d}{dt_i} \left(\frac{z_i}{q_i}\right) = \frac{q_i z_i' - q_i' z_i}{q_i^2} = \frac{z_i'}{q_i} + \frac{z_i^2}{q_i^2} \quad (2.9.38)$$

Equations (2.9.32) through (2.9.38) may be used in a Newton iterative scheme as described in Section 2.5.

The logit procedure (as described above and upon convergence) renders an estimate of the parameters α and β . Naturally, questions arise as to the certainty of these estimates. Given that the framework for the success probability still resides in the binomial distribution, the concept of the confidence interval for the success probability is retained here altered only by the application of the logistic probability formulas provided above.[4] As in the case of the probit model (based on the normal distribution), it is desirable to obtain a bivariate confidence relationship between α and β . Given that the logit method utilizes the same mathematical underpinnings, i.e., the information and asymptotic variance/covariance matrices, one might suspect that the same methodology leading to the computation of the confidence ellipse can be applied. It is a promising idea indeed, but a troublesome question arises. The theoretical development described in Section 2.7.3 is based not only on the normal distribution but, more importantly, on its connection to the chi-squared distribution that governs the uncertainty in parameters μ and σ via a special quadratic form.[14] This quadratic form involves both the information matrix and the parameter increments. The logistic distribution differs from the normal distribution, so the connection to the chi-squared distribution is not immediately evident. Fortunately, research has shown that the logistic distribution correlates well enough with the normal distribution to assert that the uncertainty for the bivariate parameters is distributed chi-squared.[19,20]. For this reason, confidence ellipses can be calculated for logit parameter estimates in the same way as for probit estimates.

3 TEST PROBLEMS AND RESULTS

In the preceding sections, some of the probability theory supporting the analysis of sensitivity test data has been described. The mathematics employed is quite complicated and much time is required in order to gain a working level of knowledge in this discipline. Mired in theory, it is easy to lose touch with the practical aspects of the calculations. For this reason, this chapter of the report presents the setup and documents the results associated with a series of basic test problems. To enhance understanding and permit brevity, we have selected problems from Garwood [11], Dixon/Mood [6] and Collett [4]. These problems are useful in testing the algorithms discussed above in Chapter 2. In and of themselves, these problem sets constitute a reasonable practicum for didactic purposes, but they do not address all of the potential scenarios for sensitivity testing. That noble goal remains on the horizon, but perhaps it can now be glimpsed in clearer focus.

3.1 Test Problem 1: Antibiotic Efficacy

This example addresses a series of serum testing trials associated with an anti-pneumonia drug. Selected doses of the drug are administered to groups of mice. The data is presented as a series of five binomial trials, but the data can also be represented in terms of a sequence of binary trials (each treated mouse constitutes a binary trial at a particular dosage level). The data is provided in Table 1, and the drug dosage is measured in cubic centimeters.[11] Each individual dosage X_D is administered to 40 mice, and after being infected with pneumonia, the number of

Table 3 Anti-Pneumonococcus Serum Test Data

Dosage (X_D , cc)	$\tilde{X}_D$	Successes out of 40
0.000625	-2	33
0.00125	-1	22
0.0025	0	8
0.005	1	5
0.01	2	2

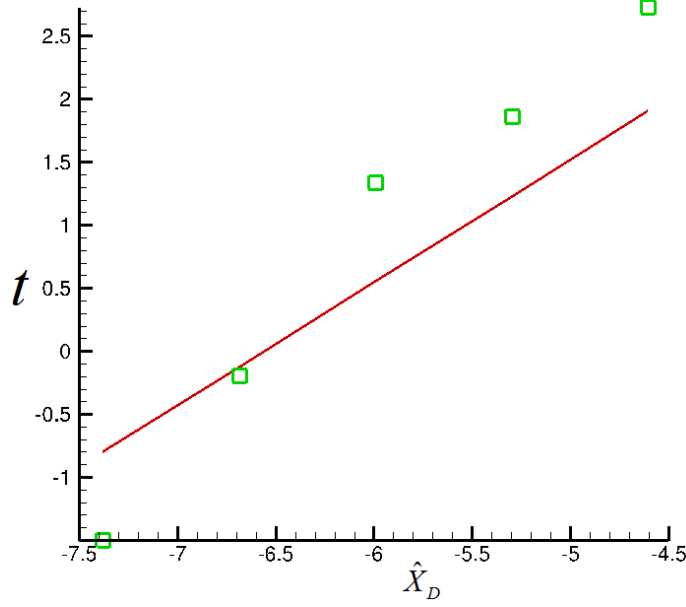


Figure 2 Normal Curve Linear Fit Form for Test Problem 1

expired mice is counted and recorded (an expired mouse is denoted as a success in this context).

In this case, it is advantageous to transform the dosage levels via the natural logarithm. The dosage

levels shown in Table 3 are cleverly selected to promote logarithmic mapping. Let $\hat{X}_D = \ln(X_D)$;

then the transformed dosages may be written as

$$\hat{X}_D^n = (n-1)\ln(2) + \ln(X_D^1), \quad n = 2, \dots, 5 \quad (3.1.1)$$

Garwood [11] applies a final transformation to the data in order to center the distribution at

$\tilde{X}_D = 0$. If there are an odd number n of data (as in this case), then let m be the natural number

with $n = 2m + 1$. The distribution is centered on the discrete interval $[-m, m]$ with the

transformation

$$\tilde{X}_D^m = \left[2 \left(\frac{\hat{X}_D^m - \hat{X}_D^1}{\hat{X}_D^5 - \hat{X}_D^1} \right) - 1 \right] m, \quad m = 1, \dots, 5 \quad (3.1.2)$$

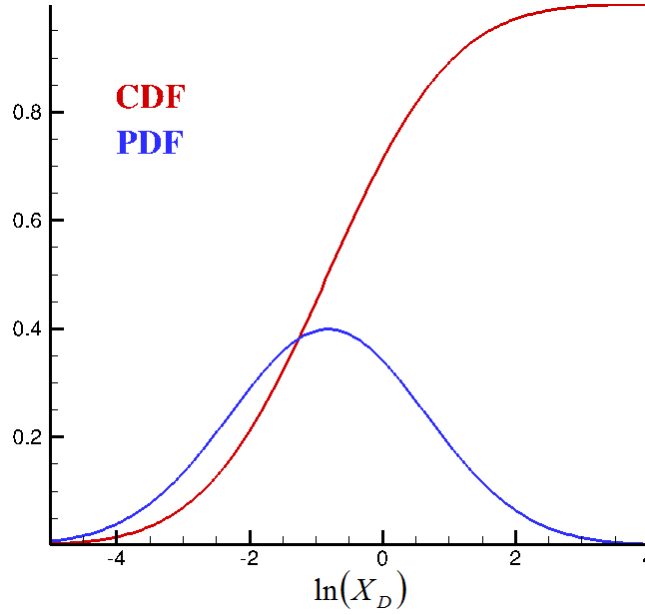


Figure 3 Probability Distributions for Test Problem 1

To perform the fitting procedure, Garwood's implementation of the method of maximum likelihood (Section 2.3) is employed with the linear form $\alpha + \beta x$ and the normal distribution. To start the iterative procedure, the initial estimates for α and β are

$$\alpha = 2; \quad \beta = 1 \quad (3.1.3)$$

The iterative scheme converges quickly to the final values of

$$\alpha = 0.554428; \quad \beta = 0.676081 \quad (3.1.4)$$

These values agree extremely well with the archived estimates in Garwood. The linear fit for this test problem is shown in Figure 2, and the probability density and cumulative distribution functions are shown in Figure 3. In accordance with the binomial model (Model 1), the 95% confidence interval endpoints in terms of the logarithmic dosage for probability estimates are provided in Table 4. Confidence interval (95%) endpoints computed using Model 2, the standard error method, are provided in Table 5. Also, the 95% confidence ellipse for the parameter estimates for α and

Table 4 95% Confidence Interval Endpoints (Model 1) for Test Problem 1

Left Endpoint $\ln(X_D)$	Probability Estimate	Right Endpoint $\ln(X_D)$
-8.520631	0.100000	-7.612530
-7.607901	0.200000	-7.201041
-7.266414	0.300000	-6.889720
-6.979877	0.400000	-6.616004
-6.715589	0.500000	-6.354047
-6.454251	0.600000	-6.085607
-6.177621	0.700000	-5.789374
-5.857550	0.800000	-5.424880
-5.419810	0.900000	-4.850462

Table 5 95% Confidence Interval Endpoints (Model 2) for Test Problem 1

Left Endpoint $\ln(X_D)$	Probability Estimate	Right Endpoint $\ln(X_D)$
-8.167631	0.100000	-7.660243
-7.646688	0.200000	-7.233966
-7.295645	0.300000	-6.917792
-7.004956	0.400000	-6.642483
-6.738876	0.500000	-6.380899
-6.477293	0.600000	-6.114819
-6.201983	0.700000	-5.824130
-5.885810	0.800000	-5.473088
-5.120757	0.900000	-4.449925

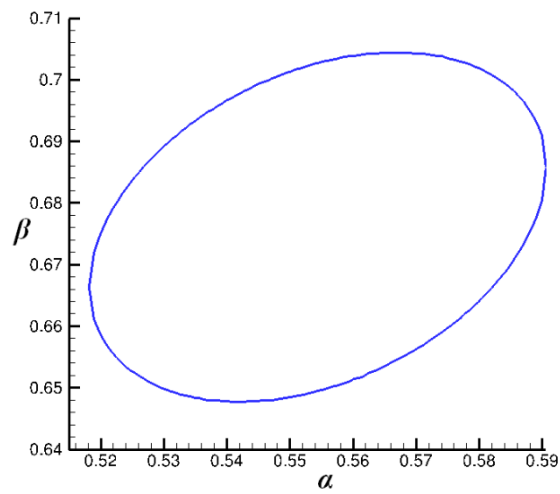


Figure 4 95% Confidence Ellipse for Parameter Estimates α and β in Test Problem 1

β is included as Figure 4. Variances have also been computed for the estimated mean and standard deviation. It is to realize that these parameters are calculated for the transformed explanatory variable, $\tilde{X}_D \in [-2, 2]$. The estimated mean and standard deviation are computed as -0.820 and 1.479, respectively. On the same axis, their standard deviations are respectively estimated to be 6.367 and 5.416.

3.2 Test Problem 2: Morbidity of *Salmonella Typhimurium*

This test problem examines the morbidity of the bacterium *Salmonella Typhimurium* when administered to laboratory mice. A reagent containing a fraction of the toxic bacteria was prepared and administered to seven groups of mice.[21] Each group consisted of five mice. In this case, the number of survivors is reported for each group in Table 6.

Table 6 *Salmonella Typhimurium* Morbidity Test Data

Dosage (mg)	$\tilde{X}_D$	Successes out of 5
0.0625	-3	4
0.125	-2	3
0.25	-1	2
0.5	0	0
1.0	1	0
2.0	2	0
4.0	3	0

The calculations entail performing a fit to the normal curve via a linear argument $(\alpha + \beta x)$. Again, Garwood's method is employed as in Section 2.3. From a casual inspection of Table 6, it is evident that the dosage values, after natural logarithmic transformation, conform to equation (3.1.1). The transformed dosage locus $\tilde{X}_D$ is then obtained through the use of (3.1.2) with m set to three. As for the previous test problem, the starting values for α and β are kept the same as in equation (3.1.3). Garwood's algorithm converges to the final values below.

$$\alpha = 1.51676; \quad \beta = 0.86344 \quad (3.2.1)$$

As in the previous case, the value of α has not been adjusted by the addition of five probits. The addition used for classical analyses, does not alter the results. The results shown here agree with the archived estimates with excellent accuracy. The linear fit for this test problem is presented in Figure 5 and is compared against the empirical data. The attendant probability density (PDF) and

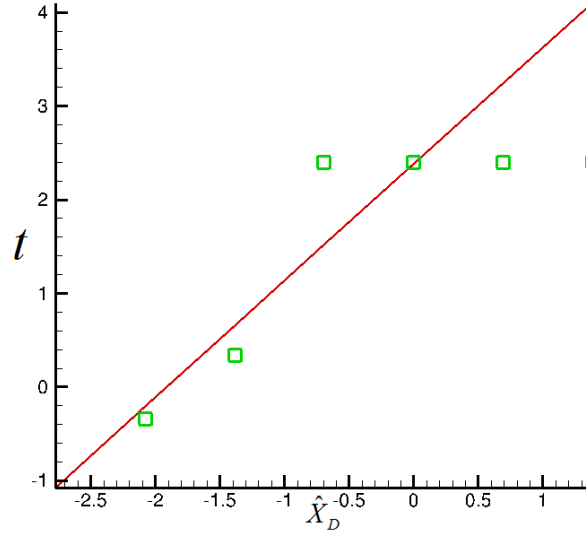


Figure 5 Normal Curve Linear Fit Form for Test Problem 2

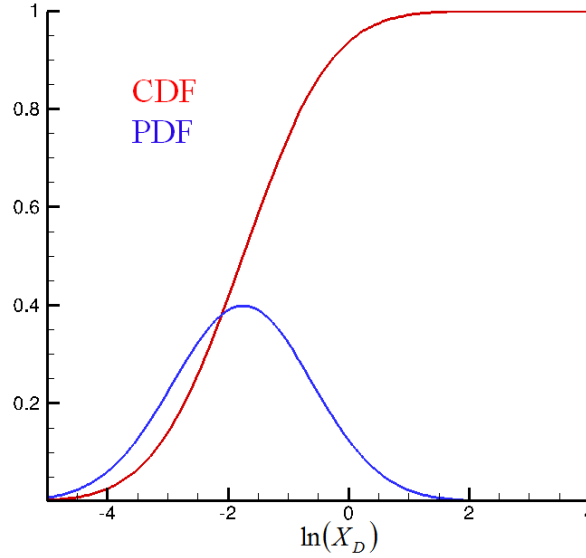


Figure 6 Probability Distributions for Test Problem 2

cumulative distribution (CDF) functions are presented in Figure 6. The endpoints of the probability 95% confidence interval for this test problem are provided in Table 7 for Confidence Model 1. Similar information is provided in Table 8 for Model 2, the standard error model. Again both of these intervals are calculated for the 95% confidence level. The notation --- indicates that the predicted endpoint falls out of the domain. Finally, the 95% confidence ellipse computed for

Table 7 95% Confidence Interval Endpoints (Model 1) for Test Problem 2

Left Endpoint $\ln(X_D)$	Probability Estimate	Right Endpoint $\ln(X_D)$
-3.303797	0.100000	-2.406570
-2.889236	0.200000	-2.117651
-2.605373	0.300000	-1.877926
-2.370206	0.400000	-1.652407
-2.155195	0.500000	-1.420311
-1.943868	0.600000	-1.156728
-1.720968	0.700000	-0.802384
-1.463119	0.800000	0.442143
-1.107982	0.900000	---

Table 8 95% Confidence Interval Endpoints (Model 2) for Test Problem 2

Left Endpoint $\ln(X_D)$	Probability Estimate	Right Endpoint $\ln(X_D)$
-4.506577	0.100000	-2.588143
-3.110754	0.200000	-2.258334
-2.749067	0.300000	-2.007947
-2.483716	0.400000	-1.784878
-2.254286	0.500000	-1.567229
-2.036637	0.600000	-1.337799
-1.813568	0.700000	-1.072448
-1.563181	0.800000	-0.710761
-1.233372	0.900000	0.685063

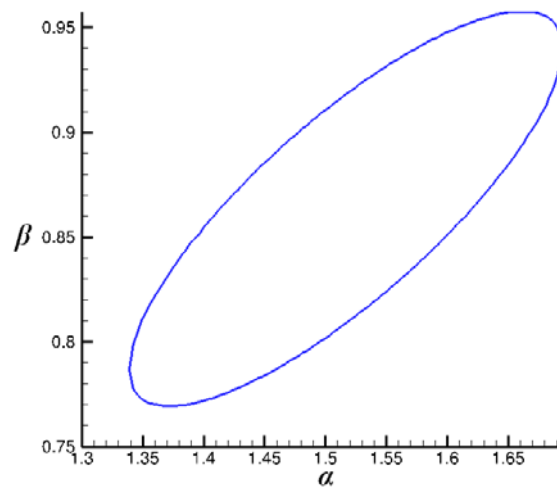


Figure 7 95% Confidence Ellipse for Parameter Estimates α and β in Test Problem 2

estimates α and β is shown in Figure 7. The mean and standard deviation for the fitted distribution are calculated as -1.756 and 1.158, respectively, while the corresponding standard deviations for these estimates are 2.675 and 2.583. Again, these values are computed for the fully transformed explanatory variable $\tilde{X}_D \in [-3, 3]$.

3.3 Test Problem 3: Arsenic Toxicity

The second test problem examines the susceptibility of brine shrimp to arsenic in liquid solution.[22] As the concentration of arsenic is increased through a geometric progression, we expect fewer shrimp to survive. The data for this test series is provided in Table 9. The true concentrations for arsenic are not provided, so the uniform values $\tilde{X}_D$ are used as listed.

Table 9 Arsenic Toxicity Test Data

Solution	$\tilde{X}_D$	Successes
C	-3	8
D	-2	8
E	-1	6
F	0	5
G	1	5
H	2	1
I	3	0

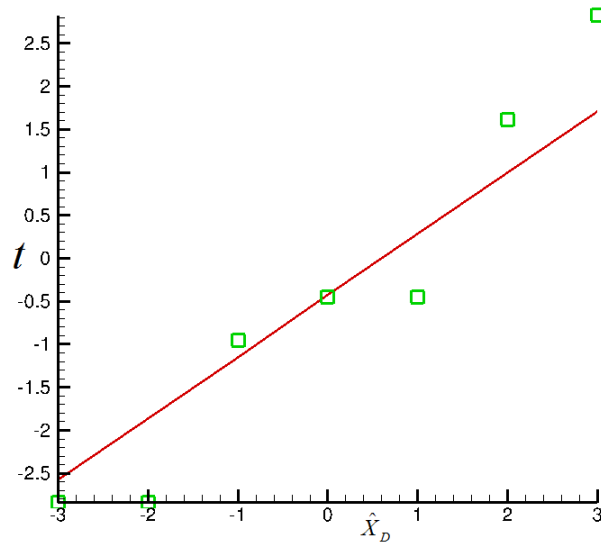


Figure 8 Normal Curve Linear Fit Form for Test Problem 3

In this case, a response is defined as a brine shrimp surviving immersion in the solution. Our starting guess for the estimation routine is given by (3.2.1), the same values as used in the preceding test problem. The converged parameter estimates are

$$\alpha = -0.434165; \quad \beta = 0.712834 \quad (3.3.1)$$

These values constitute an excellent match for the archived estimates, and the final linear fit is shown in Figure 8. Since no actual values are specified for the solution concentrations, the same dosage levels occurring in the preceding test problem are used here. Probability density and

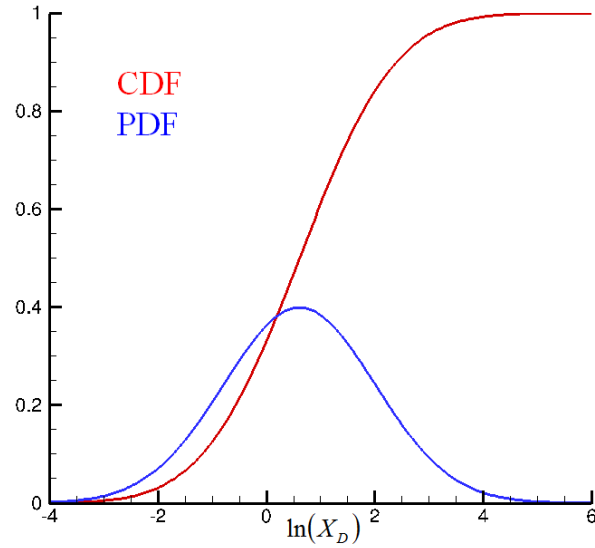


Figure 9 Probability Distributions for Test Problem 3

Table 10 95% Confidence Interval Endpoints (Model 1) for Test Problem 3

Left Endpoint $\tilde{X}_D$	Probability Estimate	Right Endpoint $\tilde{X}_D$
-1.884088	0.100000	-1.021735
-1.393996	0.200000	-0.660039
-1.056911	0.300000	-0.370771
-0.776747	0.400000	-0.106607
-0.520028	0.500000	0.155914
-0.267356	0.600000	0.438264
-0.000706	0.700000	0.775710
0.307552	0.800000	1.279043
0.730878	0.900000	---

cumulative distribution functions are easily calculated for this problem; these plots are shown in Figure 9. On comparing the parameter estimates to the corresponding values in Garwood, one may

Table 11 95% Confidence Interval Endpoints (Model 2) for Test Problem 3

Left Endpoint $\tilde{X}_D$	Probability Estimate	Right Endpoint $\tilde{X}_D$
-2.240194	0.100000	-1.166359
-1.544010	0.200000	-0.767613
-1.161147	0.300000	-0.467232
-0.861901	0.400000	-0.201910
-0.596125	0.500000	0.054181
-0.340034	0.600000	0.319958
-0.074711	0.700000	0.619203
0.225669	0.800000	1.002066
0.624415	0.900000	1.698250

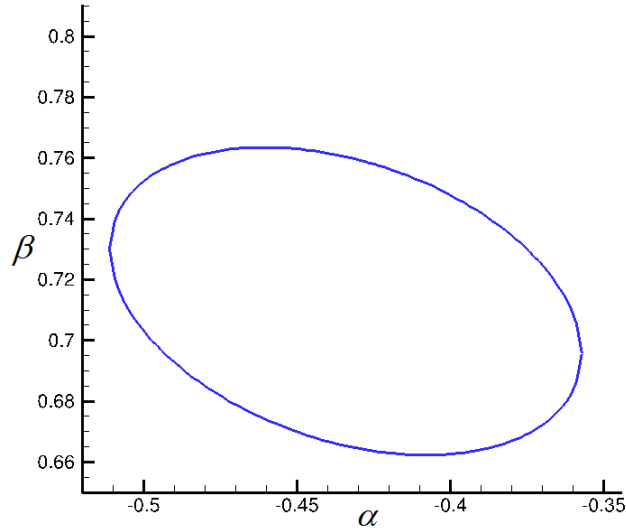


Figure 10 95% Confidence Ellipse for Parameter Estimates α and β in Test Problem 3

note that the probit quantity of 5 has not been added to the estimate of α shown in (3.3.1).[11] With this adjustment to α , the agreement is excellent. In accordance with Confidence Model 1, success probability confidence intervals are provided in Table 10. Corresponding confidence intervals computed by using Model 2 are shown in Table 11. The 95% confidence ellipse for the estimates of α and β is presented in Figure 10. The mean and standard deviation for this distribution are 0.609 and 1.402, respectively. These parameters are computed on $\tilde{X}_D$ the transformed

explanatory variable axis along with their respective standard deviation estimates of 3.159 and 3.221.

3.4 Test Problem 4: Dixon/Mood Bruceton First Test Series

The focus of this test problem is a Bruceton or Up and Down test. As described earlier, the test consists of a set of Bernoulli trials; each trial has either a pass or fail outcome. This problem is taken from Dixon and Mood and is comprised of 60 trials conducted over five (dosage) levels in the explanatory variable. The test data can be formulated in terms of five Binomial trials. The data in this form is provided in Table 12.

Table 12 Bruceton Test Data for Test Problem 4

Dosage Level	Number of Trials	Successes
2.0	1	1
1.7	10	10
1.4	27	18
1.1	20	2
0.8	2	0

For this test problem, the algorithm described in Section 2.6 is employed to estimate the mean and standard deviation for the normal probability distribution. Initial values for the mean μ and standard deviation σ are set as follows:

$$\mu_0 = 1.4 \quad \sigma_0 = 0.18 \quad (3.4.1)$$

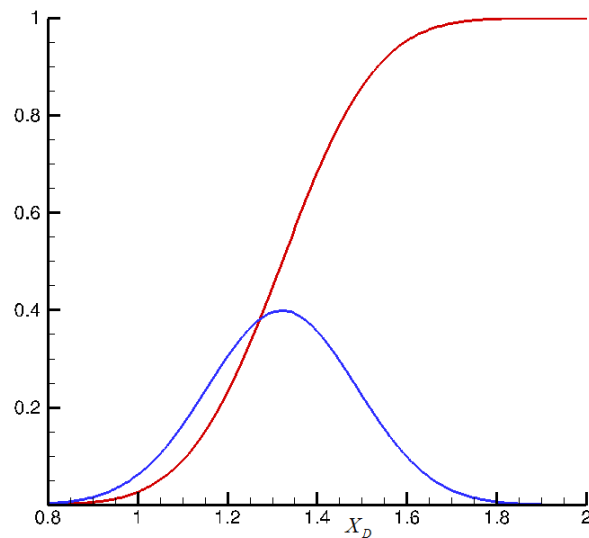


Figure 11 Probability Distributions for Test Problem 4

The iterative scheme is convergent rendering the final estimates

$$\mu = 1.3216; \quad \sigma = 0.1667 \quad (3.4.2)$$

With these values in mind, the probability distributions for this test problem are shown in Figure 11. 95% confidence intervals for the success probability according to Confidence Model 1 are documented in Table 13. Similar information for Confidence Model 2 is contained in Table 14. The confidence ellipse for parameters μ and σ is shown in Figure 12. The standard deviations

Table 13 95% Confidence Interval Endpoints (Model 1) for Test Problem 4

Left Endpoint X_D	Probability Estimate	Right Endpoint X_D
1.0467469	0.100000	1.1897100
1.1305151	0.200000	1.2520503
1.1881752	0.300000	1.3016885
1.2361231	0.400000	1.3468707
1.2800741	0.500000	1.3916125
1.3233422	0.600000	1.4394931
1.3690090	0.700000	1.4961876
1.4217988	0.800000	1.5784431
1.4942660	0.900000	---

Table 14 95% Confidence Interval Endpoints (Model 2) for Test Problem 4

Left Endpoint X_D	Probability Estimate	Right Endpoint X_D
0.9921593	0.100000	1.1663933
1.1067824	0.200000	1.2347979
1.1715661	0.300000	1.2862759
1.2225015	0.400000	1.3316982
1.2678664	0.500000	1.3754852
1.3116534	0.600000	1.4208501
1.3570757	0.700000	1.4717855
1.4085537	0.800000	1.5365692
1.4769583	0.900000	1.6511923

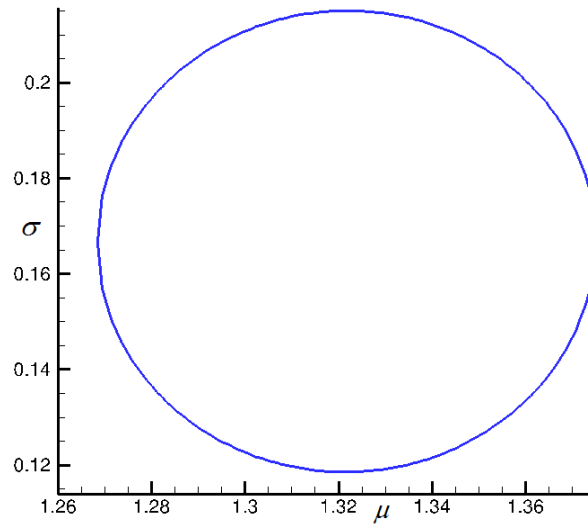


Figure 12 95% Confidence Ellipse for Parameter Estimates μ and σ in Test Problem 4

for μ and σ are computed from the asymptotic variance-covariance as 0.1664 and 0.1505, respectively.

3.5 Test Problem 5: Dixon/Mood Bruceton Second Test Series

This section of the report discusses the second Mood/Dixon Bruceton test problem. As for the preceding case, this test entails 60 trials taken over four levels of the explanatory variable. The test data, compiled as a series of four binomial trials, is provided in Table 15.

Table 15 Bruceton Test Data for Test Problem 5

Dosage Level	Number of Trials	Successes
1.9	3	3
1.5	29	27
1.9	27	1
0.1	1	0

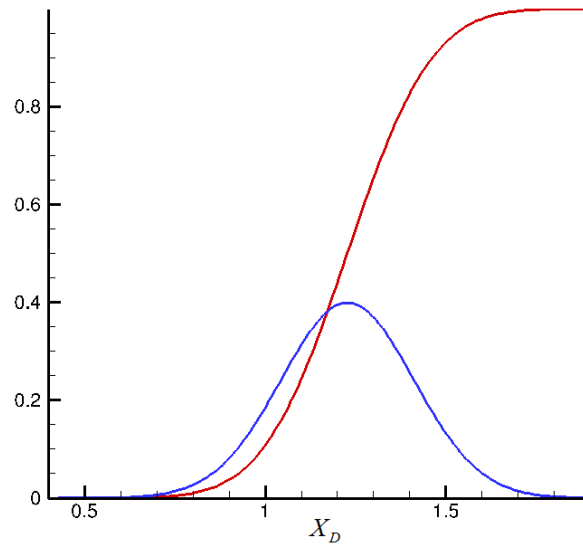


Figure 13 Probability Distributions for Test Problem 5

The initial estimates for μ and σ are the same as in (3.4.1). The iterative scheme converges to the final values of

$$\mu = 1.22775; \quad \sigma = 0.18347 \quad (3.5.1)$$

Accordingly the probability density and cumulative distribution functions for the normal distribution fit to the data is shown in Figure 13. Confidence intervals for the success probability are listed in Table 16 for Confidence Model 1 and in Table 17 for Confidence Model 2. The

Table 16 95% Confidence Interval Endpoints (Model 1) for Test Problem 5

Left Endpoint X_D	Probability Estimate	Right Endpoint X_D
0.9252993	0.100000	1.0825744
1.0174535	0.200000	1.1511556
1.0808859	0.300000	1.2057630
1.1336339	0.400000	1.2554684
1.1819848	0.500000	1.3046893
1.2295845	0.600000	1.3573632
1.2798229	0.700000	1.4197333
1.3378975	0.800000	1.5102234
1.4176194	0.900000	---

Table 17 95% Confidence Interval Endpoints (Model 2) for Test Problem 5

Left Endpoint X_D	Probability Estimate	Right Endpoint X_D
0.8652470	0.100000	1.0569234
0.9913450	0.200000	1.1321760
1.0626141	0.300000	1.1888074
1.1186486	0.400000	1.2387769
1.1685550	0.500000	1.2869475
1.2167255	0.600000	1.3368538
1.2666950	0.700000	1.3928884
1.3233265	0.800000	1.4641575
1.3985790	0.900000	1.5902555

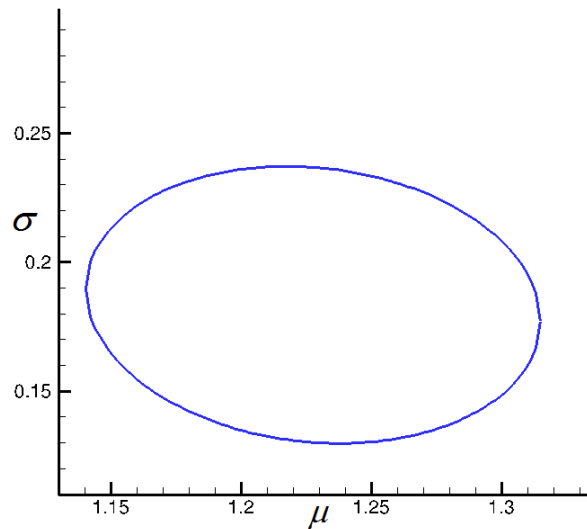


Figure 14 95% Confidence Ellipse Computed for Parameter Estimates μ and σ in Test Problem 5

confidence ellipse for the estimated values of μ and σ is shown in Figure 14. Standard deviation approximations have been calculated for the estimated μ and σ . Respectively, these values are 0.27238 and 0.16754.

3.6 Test Problem 6: Anti-Pneumococcus Serum

In the structure of this report, this test problem is new. This report's predecessor [16] did not discuss the logistic probability distribution. Yet, this statistical distribution is widely applied, so it has been included in this revised report. This test problem is taken from Collett [4] and represents a means of testing antibiotic dosage efficacy. In this clinical test series, five antibiotic dosage levels are selected. For each level, 40 mice are infected with pneumococcal bacterial and then injected with the dosage for the level. Then the number of mice that do not survive after a seven days waiting period at that level are recorded. The binomial data for this test series is presented in Table 18. A success in this case is interpreted as a non-survivor. The logistic distribution is used to fit

Table 18 Logistic Test Data for Test Problem 6

Dosage Level	Number of Trials	Successes
0.0028	40	35
0.0056	40	21
0.0112	40	9
0.0225	40	6
0.0450	40	1

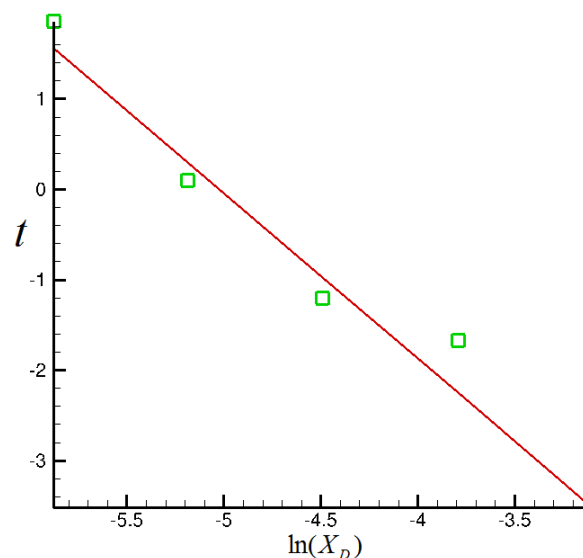


Figure 15 Linear Logistics Fit Form for Test Problem 6

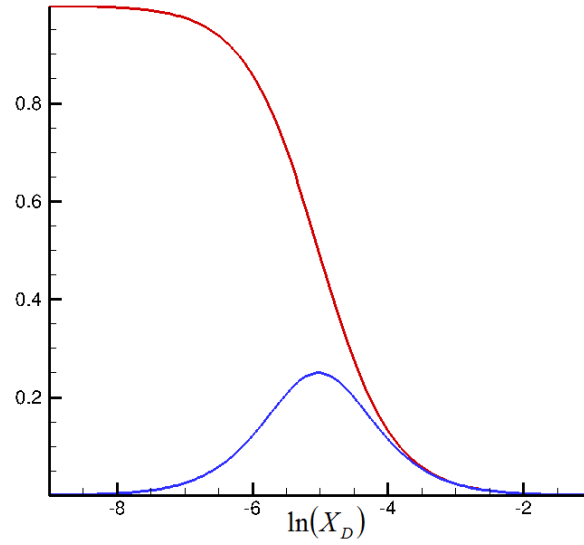


Figure 16 Probability Distributions for Test Problem 6

Table 19 95% Confidence Interval Endpoints (Model 1) for Test Problem 6

Left Endpoint $\ln(X_D)$	Probability Estimate	Right Endpoint $\ln(X_D)$
-3.4005	0.100000	-4.3251
-3.9572	0.200000	-4.6660
-4.2990	0.300000	-4.9281
-4.5654	0.400000	-5.1651
-4.8004	0.500000	-5.4031
-5.0276	0.600000	-5.6671
-5.2675	0.700000	-6.0009
-5.5514	0.800000	-6.5509
-5.9659	0.900000	---

this data with the natural logarithm of the dosage as an explanatory variable. Moreover, the linear fitting form (2.9.1) is employed. At the start of the iterative scheme, the following values are assumed for α and β .

$$\alpha_0 = -9.19; \quad \beta_0 = -1.83 \quad (3.6.1)$$

Table 20 95% Confidence Interval Endpoints (Model 2) for Test Problem 6

Left Endpoint $\ln(X_D)$	Probability Estimate	Right Endpoint $\ln(X_D)$
-2.9558	0.100000	-4.1853
-3.7993	0.200000	-4.5689
-4.1985	0.300000	-4.8433
-4.4875	0.400000	-5.0815
-4.7328	0.500000	-5.3123
-4.9637	0.600000	-5.5576
-5.2019	0.700000	-5.8466
-5.4763	0.800000	-6.2458
-5.8598	0.900000	-7.0893

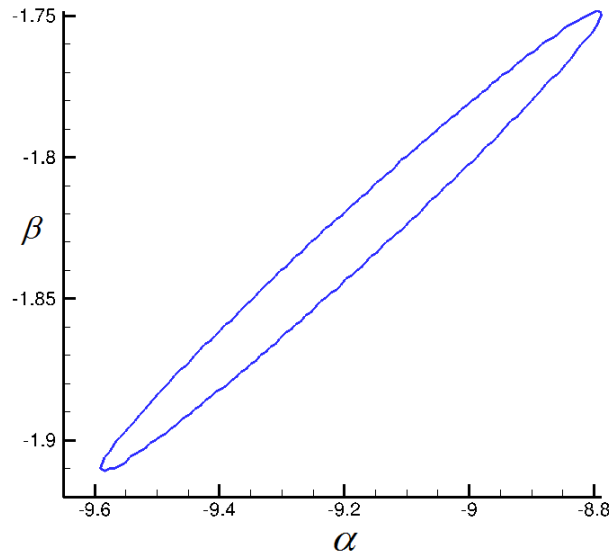


Figure 17 95% Confidence Ellipse for Parameter Estimates α and β in Test Problem 6

The iterative scheme converges to the final parameter estimates of

$$\alpha = -9.1893; \quad \beta = -1.8296 \quad (3.6.2)$$

The attendant curve fit is shown in Figure 15 while the probability density and cumulative distribution functions are shown in Figure 16. Based upon Confidence Model 1, the confidence interval endpoints for success probability are listed in Table 19. Similar information for

Confidence Model 2 is provided in Table 20. The confidence ellipse for this test problem is presented in Figure 17. The mean and standard deviation for this distribution are easily computed from the linear form (2.3.2) as

$$\mu = -5.02256; \quad \sigma = -0.54656 \quad (3.6.3)$$

based upon the logarithm of the explanatory variable. From the asymptotic variance-covariance matrix, the standard deviations for the parameter estimates (3.6.3) are computed as 9.69016 and 13.40990, respectively.

3.7 Test Problem 7: Aircraft Fastener Failure

This study involves the failure of aircraft fasteners under compressive loading. Series of tests were performed on collections of fasteners for ten different compressive loads measured in psi. For each fastener collection (or lot), a certain number of fasteners fail resulting in a binomial series with load as the explanatory variable. Table 21 contains the data for the series of fastener tests. For this test case, the load is directly used as the explanatory variable without transformation. As in the previous test problem, the logistic distribution is applied. The resulting linear fit is

Table 21 Logistic Test Data for Test Problem 7

Load (PSI)	Lot Size	Number Failing
2500	50	10
2700	70	17
2900	100	30
3100	60	21
3300	40	18
3500	85	43
3700	90	54
3900	50	33
4100	80	60
4300	65	51

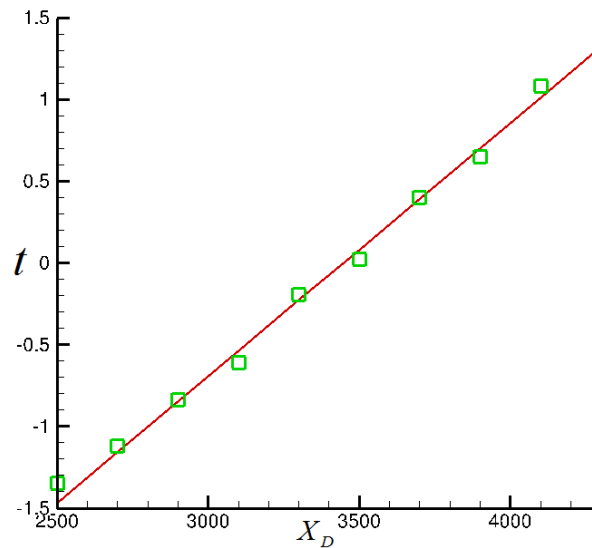


Figure 18 Linear Logistics Fit Form for Test Problem 7

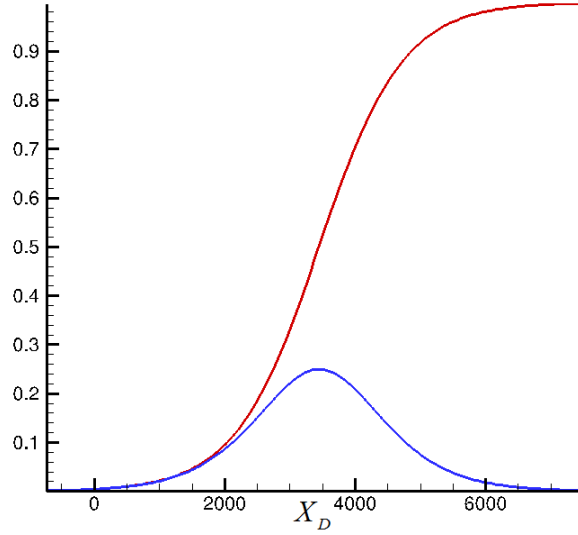


Figure 19 Probability Distributions for Test Problem 7

shown in Figure 18. The probability density and cumulative distribution functions for the logistic distribution are presented in Figure 19. In accordance with Confidence Model 1, the success ratio confidence interval endpoints are provided in Table 22. Corresponding endpoints for Model 2 are provided in Table 23. The confidence ellipse for distribution parameters α and β is shown in

Table 22 95% Confidence Interval Endpoints (Model 1) for Test Problem 7

Left Endpoint X_D	Probability Estimate	Right Endpoint X_D
1745.1869	0.100000	2341.7176
2343.7843	0.200000	2794.1246
2721.8920	0.300000	3116.8377
3022.1416	0.400000	3393.5079
3290.9217	0.500000	3657.2709
3553.8612	0.600000	3931.5218
3834.3142	0.700000	4244.8644
4168.6696	0.800000	4655.5562
4657.8710	0.900000	5387.4483

Figure 20. As is evident, the fitting procedure functions quite well rendering the parameter estimates of

$$\alpha = -5.33971 \text{ and } \beta = 1.54843 \times 10^{-3} \quad (3.7.1)$$

Table 23 95% Confidence Interval Endpoints (Model 2) for Test Problem 7

Left Endpoint X_D	Probability Estimate	Right Endpoint X_D
1650.0926	0.100000	2285.9406
2298.4357	0.200000	2758.7448
2690.1111	0.300000	3087.8241
2995.9883	0.400000	3366.5277
3267.1356	0.500000	3629.7845
3530.3924	0.600000	3900.9318
3809.0960	0.700000	4206.8090
4138.1752	0.800000	4598.4844
4610.9794	0.900000	5246.8275

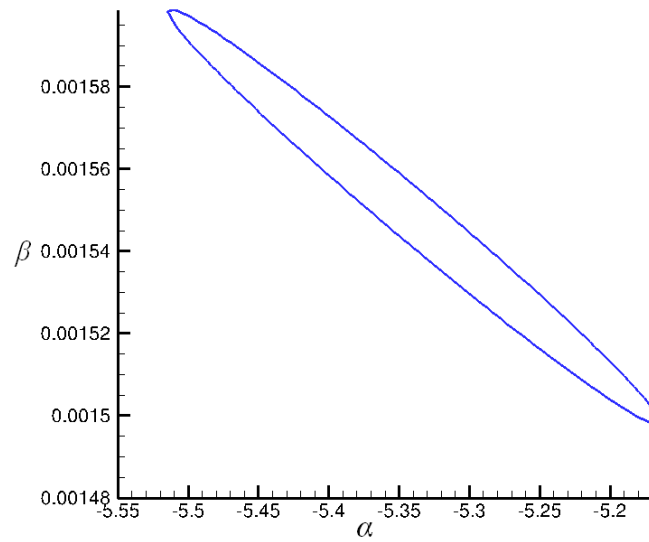


Figure 2 95% Confidence Ellipse for Parameter Estimates α and β in Test Problem 7

The corresponding mean and standard deviation for this distribution are computed as

$$\mu = 3448.460 \text{ and } \sigma = 645.813 \quad (3.7.2)$$

A variance can be calculated for these two parameters, i.e., 1.86844×10^{-2} and 1.52363×10^{-2} , respectively.

4 CONCLUSIONS

In this report, we have presented a discussion of sensitivity (or Go/No-Go) testing from basic principles. The prototypical set of sensitivity experiments is based upon the Bruceton test procedure, or the “Up and Down Method”, used to estimate the statistical parameters associated with a series of binary trials. The outcome of each binary trial is a random variable, and these random variables are independent and identically distributed. For the “Probit” method, the mean and standard deviation are estimated by fitting data to a normal distribution. For pharmaceutical applications, the so-called “dosage-mortality curve” is produced in the form of a linear fit. The behavior of the normal curve is then used to determine the probability of success associated with an explanatory variable of choice. For the Logit method, the logistic distribution is used since it is a bit more conservative near the tails, yet it looks a lot like the normal distribution. The set-up and execution of the Bruceton test procedure have been discussed for an example scenario. The importance of locating the mean is stressed as well as ensuring that extraneous data far removed from the mean is excluded from the analysis. Moreover, issues surrounding the efficiency of the data have been discussed. That is to say, we must obtain a sufficiently large sample in order for our analysis techniques and assumptions to apply. Algorithms for data analysis have been discussed in detail, particularly on the topics of Probit analysis, Logit analysis and confidence estimation. We have shown the set of equations required to fit go-no go data to the normal distribution and to the logistic distribution. These methods have proven capable and they readily address binomial trials on nonuniformly separated levels. The fundamental solution scheme discussed here is based upon Garwood’s method, yet we have presented an alternative to this method that can easily incorporate both binary and binomial trials. The algorithm structure works for both the Probit and Logit distributions. The alternate algorithm is stable and remains well-

defined in every case. From the standpoint of sampling, we have discussed interval estimation for sensitivity tests. Algorithms have been suggested for the qualified estimation of probability confidence intervals as well as those for the mean and variance. We have also derived equations that allow the variance and covariance to be estimated for the normal mean and variance produced by our fitting procedure. Seven classic test problems have been chosen from engineering and biological science to serve as practical examples. These problems have been solved by directly coding the algorithms presented in this report. In each case, our results have achieved excellent agreement with archival solutions. Confidence regions have been computed for these cases. We have also discussed shortfalls that exist within our process for estimating intervals and for estimating elements in the asymptotic variance-covariance matrix.

Sensitivity testing is very important, not only to the government sector, but also to the commercial sector, namely the pharmaceutical industry. The efficacies and effective dosages for medications are largely determined through this type of testing and analysis. In most cases, the statistical distributions produced by the analysis are used to make important management decisions. The investment in the development of a single medication can cost hundreds of millions of dollars. The cost of qualifying explosive systems is also high. For these reasons, it is important that we attain a good understanding of the probability and statistics. It is very easy to become mired in either of these fields. These divisions of mathematics undergird sensitivity testing and analysis. System reliability is of paramount importance for protecting both the investment of funding and human life. Failing to accurately estimate the reliability of a modern drug or engineering system can have grave consequences.

REFERENCES

1. Billingsley, P., *Probability and Measure*, 2nd Ed., Wiley, New York, New York, 1986.
2. Ross, S.M., *Introduction to Probability Models*, 4th Ed., Academic Press, San Diego, CA, 1989.
3. Langlie, H.J., A Reliability Test Method for “One-Shot” Items, *Proceedings of the Eighth Conference on the Design of Experiments in Army Research Development and Testing*, Aeronutronic Division, Ford Motor Company, ADP014612 in ADA419759, 1963.
4. Collett, D., *Modelling Binary Data*, 2nd Ed., *Texts in Statistical Science*, Chapman & Hall/CRC, 2003.
5. Crow, E.L., Davis, F.A. and Maxfield, M.W., “Statistics Manual with Examples Taken from Ordnance Development”, NAVORD Report 3369, Research Department, U.S. Naval Ordnance Test Station, September 1955.
6. Dixon, W.J. and Mood, A.M., “A method for obtaining and analyzing sensitivity data”, *Journal of the American Statistical Association*, Vol. 43, 1948, pp. 109-126.
7. Bliss, C.I., “The calculation of the dosage-mortality curve”, *Annals of Applied Biology*, Vol. 22, 1935, pp. 134-167.
8. Golub, A. and Grubbs, F.E., “Analysis of sensitivity experiments when the levels of stimulus cannot be controlled”, *American Statistical Association Journal*, Vol. 57, 1956, pp. 257-265.
9. Banerjee, K.S., “On the Efficiency of Sensitivity Experiments Analyzed by the Maximum Likelihood Estimation Procedure Under the Cumulative Normal Response”, Technical Report ARBRL-TR-02269, ADA092349, US Army Armament Research and Development Command, Ballistic Research Laboratory, Aberdeen Proving Ground, MD, September 1980.
10. Bowling, S.R., Khasawneh, M.T., Kaewkuekool, S. and Cho, B.R., “A logistic approximation to the cumulative normal distribution”, *Journal of Industrial Engineering and Management*, Vol. 2, No. 1, 2009, pp. 114-127.
11. Garwood, F., “The application of maximum likelihood to dosage - mortality curves”, *Biometrika*, Vol. 23, 1941, pp. 46-58.
12. Cornfield, J. and Mantel, N., “Some new aspects of the application of maximum likelihood to the calculation of the dosage response curve”, *American Statistical Association Journal*, Vol. 45, 1950, pp. 181-210.
13. Hogg, R.V. and Craig, A.T., *Introduction to Mathematical Statistics*, 3rd Ed., Macmillan Publishing, New York, 1970.
14. Scheffé, H. *The Analysis of Variance*, Wiley-Interscience, New York, New York, 1959.

15. Larsen, R.J. and Marx, M.L., *Introduction to Mathematical Statistics and its Applications*. Prentice-Hall, New Jersey, 1981.
16. Nance, D.V., Analysis of Sensitivity Experiments – A Primer, AFRL-RW-EG-TR-2009-7005, Technical Report, Air Force Research Laboratory, 2008.
17. DiDonato, A.R. and Jarnagin, M.P., Jr., Use of the Maximum Likelihood Method Under Quantal Responses for Estimating the Parameters of a Normal Distribution and its Application to an Armor Penetration Problem, TR-2846, Technical Report, Naval Weapons Laboratory, AD-762399, National Technical Information Service, 1972.
18. “Compute Variance of Logistic Distribution”, Web Article Addressing Question 1267635, <http://math.stackexchange.com>.
19. Mazumdar, S., “Monte Carlo Methods for Confidence Bands in Nonlinear Regression”, Master’s Thesis, University of North Florida, 1995.
20. Khorasani, F. and Milliken, G.A., “Simultaneously confidence bands for nonlinear regression models”, *Communications in Statistics: Theory and Methods*, Vol. 11, No. 11, 1982, pp. 1241-1253.
21. MacKenzie, G.M., Pike, R.M. and Swinney, R.E., “Virulence of Salmonella Typhimurium – II. Studies of the polysaccharide antigens of virulent and avirulent strains”, *Journal of Bacteriology*, Vol. 40, No. 2, 1940, pp. 197-214.
22. Fisher, R.A. and Yates, F., *Statistical Tables for Biological, Agricultural and Medical Research*, Oliver & Boyd, Edinburgh & London, 1938.

DISTRIBUTION LIST

AFRL-RW-EG-TP-2017-001

*Defense Technical Info Center
8725 John J. Kingman Rd Ste 0944
Fort Belvoir VA 22060-6218

AFRL/RWML (1)
AFRL/RWORR (STINFO Office) (1)