

304

LOW DOSE RISK ESTIMATION USING TWO CONVENIENT FORMS
OF THE GENERALIZED PROBIT MODEL

by

JEFFREY J. RAHENKAMP

B.S., Clarion State College, 1976

A MASTER'S REPORT

submitted in partial fulfillment of the

requirements for the degree

MASTER OF SCIENCE

Department of Statistics

KANSAS STATE UNIVERSITY
Manhattan, Kansas

1980

Approved by

George O. Milliken
Major Professor

SPEC
COLL
LD
2668
.R4
1980
R33
C.2

TABLE OF CONTENTS

	Page
1. INTRODUCTION.	1
1.1 Difficulties encountered in Low Dose Risk Estimations.	1
1.2 Probit type analysis in Bioassay	2
2. SOME COMPETING MODELS IN RISK ESTIMATION.	5
2.1 Mantel-Bryan Procedure	5
2.2 Power Transformation	7
2.3 Generalized Probit Model	7
2.4 Other Techniques	11
3. THE GENERALIZED PROBIT MODEL.	13
3.1 Data Sets.	13
3.2 M.L.E.'s for Generalized Probit Model.	13
3.3 Generalized Probit Procedure	19
3.4 Effective Dose Estimation.	24
4. EMPIRICAL COMPARISONS OF VARIOUS MODELS	31
4.1 V.S.D. for Mantel-Bryan Procedure and Generalized Probit.	31
4.2 Empirical comparisons of Effective Doses	34
4.3 Goodness of fit.	37
5. DISCUSSION	47
REFERENCES	49
ACKNOWLEDGEMENTS	51

**THIS BOOK
CONTAINS
NUMEROUS PAGES
WITH DIAGRAMS
THAT ARE CROOKED
COMPARED TO THE
REST OF THE
INFORMATION ON
THE PAGE.**

**THIS IS AS
RECEIVED FROM
CUSTOMER.**

I. INTRODUCTION

1.1 Difficulties encountered in low dose risk estimation

The inimical problems confronting the statistician involved in low dose risk estimation, especially of a suspected carcinogen, are many. The subject can be a highly emotional one. The cancer mechanism remains poorly understood. This, and a wide variety of types of cancer, leads to a wide variety of "best" ways to mathematically describe a population's tolerance distribution.

A measure of an individual's tolerance (a.k.a. limen, threshold) is that strength of a given stimulus required to produce a response in the experimental unit. The response may be in the form of death or perhaps development of a tumor. The frequency distribution of tolerances is the population's tolerance distribution.

The statistician is also forced into the nebulous area of density functions when faced with the task of determining a "virtually safe" dose (V.S.D.). A V.S.D. is defined to be (Mantel and Bryan [1961]) that dose which gives a negligible increase in the probability of a response above the expected response of the control group. An increase in the neighborhood of 10^{-6} or 10^{-8} is usually used to define virtual safety. An increase of 10^{-8} has been adopted by the Food and Drug Administration (F.D.A.) as a virtually safe level (Federal Register [1977]).

This is extremely soft ground to tread statistically for two reasons. First, serious questions can be raised about any convenient assumptions which appear to hold in the area around the mean or median. Second, even with large sample sizes observing values in the area of interest is impossible. Extrapolation, sometimes over large distances, becomes a

necessity.

Extrapolation of another sort can be a problem. That is, how to equate a V.S.D. for a laboratory mouse to a V.S.D. for a human. Obviously, in any controlled study animals must serve as surrogates. Even here there is disagreement. Guess and Crump [1977], statisticians, state: "There is not much chance of being able to use data from animal tests to set dose levels high enough to be economically useful for food additives, for example, and yet low enough to validly provide assurance that even the risk in animals of the strains tested are as low as 10^{-6} ." Rall [1974], physician, states: "It seems to me that we must in fact use animal tests, today at least, as the basis for prediction of carcinogenic activity." Add to the list of problems those of political and economical nature, and the statistician can be placed in a difficult position.

At the present, faced with much disagreement over the way to describe the cancer mechanism, a probit type analysis seems to be as good as any method for risk estimation. However one that is not limited to normality assumptions may be of special interest. The generalized probit model as suggested by Prentice [1976] appears to be an improvement over the probit and logit models, when presented with apparently nonnormal quantal response data. The more flexible model should provide a better fit and in turn reflect a more accurate V.S.D. estimate.

1.2 Probit type analysis in Bioassay

Bioassay relates to techniques relevant to comparisons between the strengths of alternate but similar biological stimuli. Probit type analysis is appropriate when the response of a subject to a certain strength or level of the stimulus is nominal data: dead or alive, tumor

or no tumor, etc.

An experiment consists of dividing N experimental units into i groups of size n_i , where i is the number of dose levels or number of levels of the stimulus. Let r_i denote the number of experimental units which responds to the stimulus in each group of n_i experimental units. Then the observed probability of a "lifetime" response to dose d_i is $\hat{p}_i = r_i/n_i$, and $P_i = F(d_i)$ denotes the probability of lifetime response due to dose d_i based on the tolerance distribution $F(\cdot)$.

Let z = dose or function of dose and dP = proportion of the population whose tolerances are between z and $z + dz$. Then $dP = f(z)dz$ and $P(z_0) = \int_{-\infty}^{z_0} f(z)dz$ is the proportion of experimental units in the population with tolerances less than z_0 . The probit model assumes that $f(z)$ is a normal distribution; i.e., the tolerances $z \sim N(\mu, \sigma^2)$. Define the normal equivalent deviate as $y = (z - \mu)/\sigma$ where $y \sim N(0, 1)$. Then $P(z_i) = \int_{-\infty}^{y_i} \phi(u)du$, where $\phi(u)$ denotes the standard normal density function. The logit model differs in that the function under the integral, or the tolerance distribution, is assumed to be a logistic density function.

Frequently the transformation $x = \log_{10} z$ or $\log_e z$ is made to help achieve symmetry of the tolerance distribution. This transformation scale is known as the metametric scale and the measure of the dose is called the dose metameter. When the tolerance distribution is normal then plotting y versus x usually results in a straight line. y can then be expressed as $\alpha + \beta x$, where α is the intercept of the line and β is the slope. Given a desired P and estimates for α and β a corresponding log-dose and subsequent dose can be estimated. For example, if $\hat{\alpha} = -1.392$

and $\gamma = 1.88$, then an estimate for the V.S.D. would be the antilog of x where x is the solution to $-1.392 + 1.88 x = -5.612$ where $P(-5.612) = 10^{-8}$. In this case $x = -2.245$ with an antilog of .00569, or V.S.D. = .00569.

A popular "yardstick" used in Bioassay to compare tolerance distributions or effectiveness of the stimulus is the ED50 (a.k.a. effective dose-50, median lethal dose). This is the median of the tolerance distribution or that estimated dose which will elicit a response in 50% of the population.

II. SOME COMPETING MODELS IN RISK ESTIMATION

2.1 Mantel-Bryan Procedure

The (Improved) Mantel-Bryan procedure [1975] is basically a probit approach to determine the V.S.D. estimate. A three parameter model as described by Finney [1947] is used where

a = intercept, b = slope, c = expected response of control
animals

$x = \log_{10}$ (dose) and

$P(a, b, c, x) = c + (1-c)P(a, b, o, x)$ where

$$P(a, b, o, x) = \int_{-\infty}^{a+bx} \phi(u) du.$$

The probit procedure is modified by fixing the slope b at an arbitrarily low value of one. Mantel and Bryan contend that from experience this slope assures a conservative but not unreasonably low V.S.D. estimate. Data at all levels, including any control data, are fitted to get joint maximum likelihood estimates of a and c , with b fixed at one. An upper 99% limit is then found for the intercept a . This limit is the value of a , in conjunction with its corresponding M.L.E. for c , which leads to just a significant reduction in the likelihood function. This criteria relies on the asymptotic distribution of the likelihood ratio function.

Let a be the upper $(1 - \alpha)$ 100% limit on a , then

$$-2 \ln \left[\frac{L(\hat{a}, b=1, \hat{c})}{L(a^*, b=1, \hat{c})} \right] \sim \chi^2_{(1)} \quad . \quad (\text{Theorem 7, pg. 440, Mood, Greybill, Boes [1974]}).$$

or

$$2 [\ln L(\hat{a}) - \ln L(a^*)] \sim \chi^2_{(1)}$$

or

$$[\ln L(\hat{a}) - \ln L(a^*)] \sim \frac{1}{2} \chi^2_{(1)} .$$

Next, the highest dose level is deleted and the procedure is repeated to determine another upper limit on a . The deletion process is repeated, calculating a^* after each deletion, until only the lowest nonzero dose level is used. The lowest of the upper limits is selected; then extrapolating along the line determined by this intercept and $b = 1$, the log (dose) is found (x axis) that corresponds to the normal deviate (y axis) of the desired risk.

The lowest of the upper limits on a results in the highest of the V.S.D. estimates. Mantel and Bryan contend that choosing the lowest upper limit, instead of one of its competitors (one for every nonzero dose level) and consequently the highest V.S.D. is valid because any lower V.S.D. estimates result from the "too conservative" slope of $b = 1$. The higher V.S.D. estimate is due to the "added information" at that particular dose level.

The perfidy of the method with respect to statistics is observed by Crump [Dec. 1977] when he points out three major problems. First, the model typically provides a very poor fit to the data. Second, there has been no cancer mechanism proposed that would lead to the model. Finally, because of the arbitrariness of the model, the claims of conservatism is unjustified.

This procedure is presently used by the F.D.A. to determine virtually safe doses of suspected carcinogens (Federal Registrar [1977]).

2.2 Power Transformation

The probit model assumes the dose-response curve is a normal symmetric sigmoid curve; hence the log transformation is often used in an attempt to achieve symmetry. In an attempt to find a more effective transformation, Egger [1979] has investigated a power transformation, λ ; which forces the response curve into a symmetric sigmoid curve.

Let

$$\begin{aligned} p_i &= \text{probability of a response at} \\ &\quad \text{dose } z_i \\ &= \Phi [\alpha + \beta(z_i^\lambda - 1)/\lambda] \\ &= \Phi (\gamma + \tau z_i^\lambda) \end{aligned}$$

where Φ is the cumulative standard normal density function.

Three techniques for estimating λ are presented by Egger. One is a logistic maximum likelihood estimate. Two are nonparametric techniques, based on the Spearman third moment and on constraints on a set of estimated quantiles. The resulting transformed data can then be analyzed with probit or logit techniques.

2.3 Generalized Probit Model

Another approach to achieving a better fitting model was suggested by Prentice [1976] in the form of a function with two additional shape parameters. In a probit type analysis the function

$$f(w) = \exp(wm_1)[1 + \exp(w)]^{-(m_1+m_2)}/\beta(m_1, m_2)$$

where $\beta(m_1, m_2)$ is the beta function, replaces the standard normal density.

The normal, logistic, extreme maxima, and extreme minima along with other well known distributions are special cases of the above function. The logistic density results when $m_1 = m_2 = 1$; the normal when $m_1, m_2 \rightarrow \infty$; the extreme maxima when $m_1 \rightarrow \infty, m_2 \rightarrow 1$; and the extreme minima when $m_1 = 1, m_2 \rightarrow \infty$. This model should prove more fecund than the more specific models, but finding maximum likelihood estimates for $\underline{\theta}' = [\mu, \sigma, m_1, m_2]$ can be difficult and time consuming.

In an effort to produce a practical and convenient procedure, a closer examination of the Prentice model will be made by concentrating on four models, two being the normal and logistic for comparison purposes. The other two will be where $\underline{\theta}'_1 = [\mu, \sigma, m_1], m_2 = 1$ and $\underline{\theta}'_2 = [\mu, \sigma, m_1], m_1 = 1$; or the M1 model and M2 model respectively. Fixing either m_1 or m_2 model respectively. Fixing either m_1 or m_2 at one greatly simplifies maximum likelihood estimation but still takes advantage of, to some extent, the flexibility of the Prentice model. (See figures 2.31 and 2.32).

The dose metameters will be common logs of the dose. The statistical software package SAS (SAS User's Guide, 1979 Edition) or Statistical Analysis System was utilized to provide easier and faster maximum likelihood estimation for the parameters of interest. SAS also supplied a sum of squares which is asymptotically distributed as a chi-square and provided a measure of goodness of fit.

FIGURE 2.31

ADDED FLEXIBILITY OF M1
 $\mu=2.0$ $\sigma=.3$

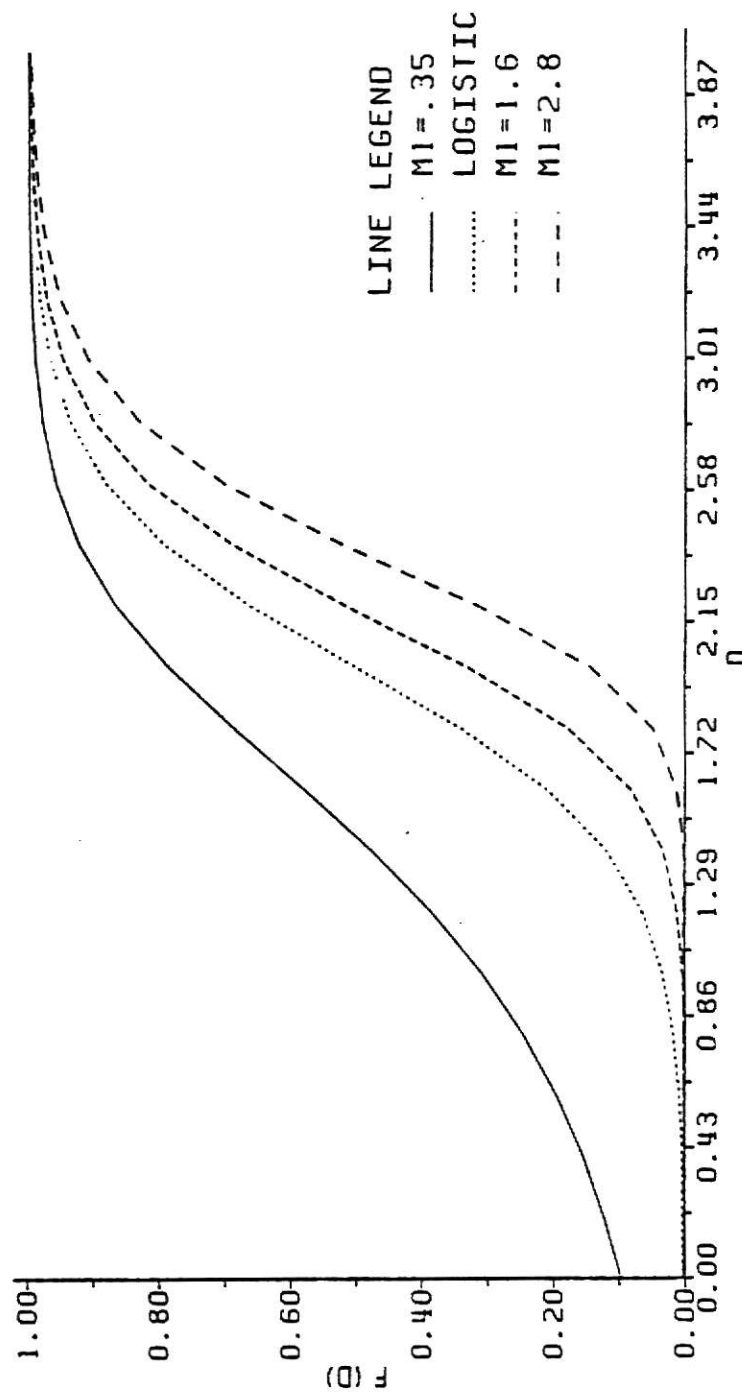
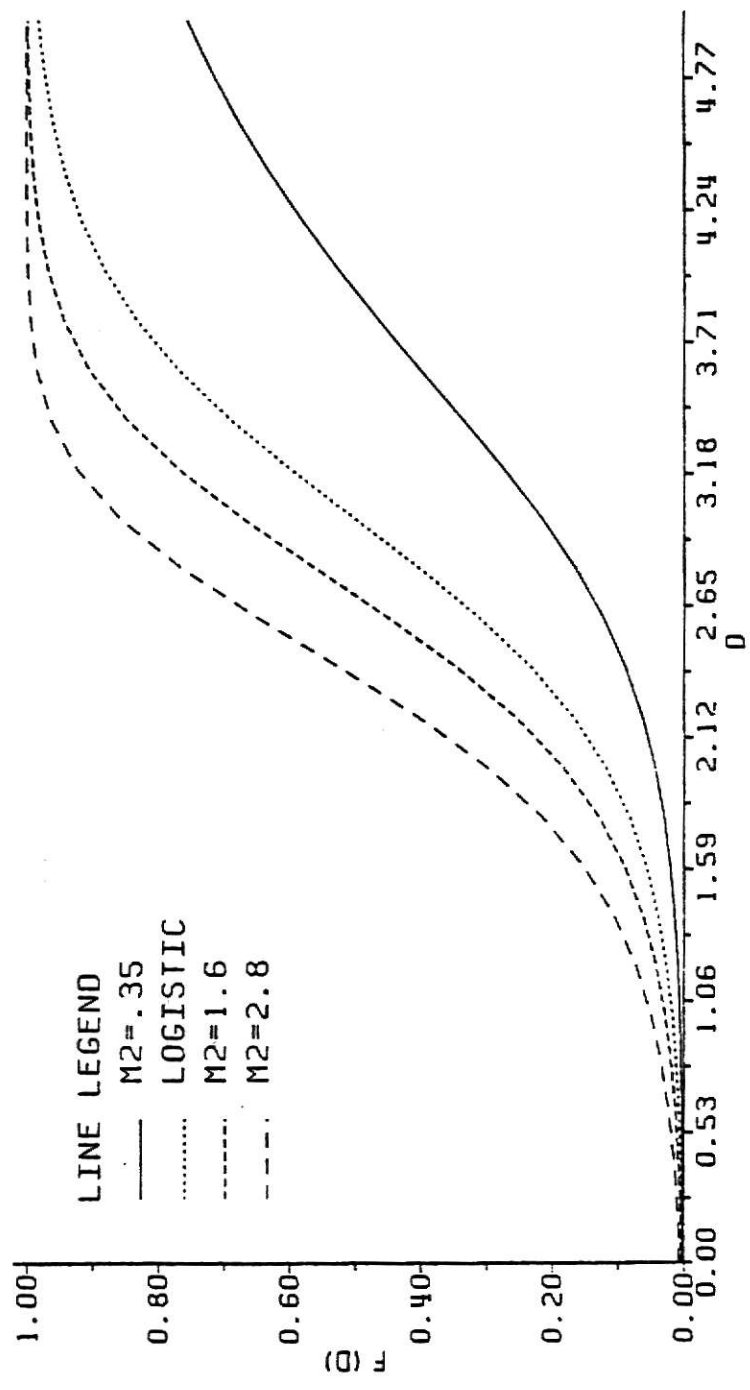


FIGURE 2.32

ADDED FLEXIBILITY OF M2
 $\mu=3.0 \quad \sigma=.5$



2.4 Other Techniques

For those interested in other techniques of risk assessment the following, not considered in this report, will prove helpful: "Can We Use Animal Data to Estimate 'Safe' Doses for Chemical Carcinogens?", (Guess and Crump [1977]); "Statistical Methods for Assessing Risk Following Exposure to Environmental Carcinogens", (Pasternack and Shore [1977]); "Lung Cancer Incidence in Cigarette Smokers: Further Analysis of Doll and Hill's Data for British Physicians", (Whittemore and Altschuler [1976]); "Estimation of 'Safe Doses' in Carcinogenic Experiments", (Hartley and Sielken [1977]); "A Bayesian Approach to Assessing Population Risks from Environmental Carcinogens", (Altschuler [1977]); "Weibull Distributions for Continuous-Carcinogenesis Experiments", (Petro and Lee [1973]), and "Carcinogenic Risk Assessment", (Cornfield [1977]).

TABLE 3.11

Data Set #1 (Copenhaver and Mielke Data Set #1)

Source: Thompson (1947), cited in Copenhaver and Mielke (1977), and Egger (1979)

<u>Dose</u>	<u>n_i</u>	<u>r_i</u>	<u>$\hat{p}_i$</u>
1.0625	50	6	.12
1.1250	50	7	.14
1.2500	50	33	.66
1.5000	50	39	.78
2.0000	50	45	.90
3.0000	50	50	1.00
5.0000	50	50	1.00

Data Set #2 (Copenhaver and Mielke Data Set #10)

Source: Finney (1974), cited in Copenhaver and Mielke (1977), and Egger (1979)

<u>Dose</u>	<u>n_i</u>	<u>r_i</u>	<u>$\hat{p}_i$</u>
2.6358941	49	16	.33
3.0000000	48	18	.38
3.4794154	48	34	.71
3.7894873	49	47	.96
4.0525184	50	47	.94
4.2490096	48	48	1.00

III. THE GENERALIZED PROBIT METHOD

3.1 Data Sets

A technical report by Egger [1979] contains twenty three data sets originally compiled from literature by Copenhaver and Mielke [1977], which lend themselves to probit analysis.

The Probit Procedure (PROC PROBIT) in SAS, using common log of the doses, supplied M.L.E.'s for the probit model which resulted in an excellent fit for many of the data sets (# 6, 9, 11, 19, 21, 22) which the generalized probit in the form of the M1 and M2 models could not improve upon. Some had four or less dose levels (# 2, 5, 12, 13) which made fitting a three parameter model difficult or impossible. Others had small sample sizes (# 3, 8, 15) hampering asymptotic qualities.

Two data sets (# 1, 10) resulted in rejecting the fit of the probit model at the $\alpha = .05$ level as measured by the residual sum of squares which are asymptotically chi-square distributed (see §3.2). These two data sets, shown in Table 3.11, will be used to illustrate the generalized probit method and for empirical comparisons with the probit, logit, Mantel-Bryan, and λ power transformation methods. The doses given are measured on their original metametric scale, not common logs.

3.2 Maximum Likelihood Estimation for the Generalized Probit Model

Suppose there are k different dose levels of a stimulus. Dose i , $i = 1, \dots, k$, is administered to n_i subjects with r_i subjects responding. Assuming independence between dose levels, each r_i is independently distributed as a binomial random variable. The probability of a response at dose d_i is

$$P(d_1, \underline{\theta}) = \int_{-\infty}^{y_1} f(w) dw, \quad y_1 = \frac{d_1 - \mu}{\sigma},$$

where

$$f(w) = \exp(wm_1) [1 + \exp(w)]^{-(m_1 + m_2)} / \beta(m_1, m_2)$$

and $\beta(m_1, m_2)$ is the beta function. The parameter of interest is

$$\underline{\theta}' = [\mu, \sigma, m_1, m_2].$$

The likelihood equation for $\underline{\theta}$ is

$$L(\underline{\theta}) = \prod_{i=1}^k \binom{n_i}{r_i} [P(d_i, \underline{\theta})]^{r_i} [1 - P(d_i, \underline{\theta})]^{n_i - r_i}$$

and consequently the log-likelihood equation is

$$\ln L(\underline{\theta}) = c + \sum_{i=1}^k r_i \ln [P(d_i, \underline{\theta})] + \sum_{i=1}^k (n_i - r_i) \ln [1 - P(d_i, \underline{\theta})].$$

Now

$$\frac{\partial \ln L(\underline{\theta})}{\partial \theta_s} = \sum_{i=1}^k r_i \frac{\partial P(d_i, \underline{\theta})}{\partial \theta_s} P(d_i, \underline{\theta})^{-1} + \sum_{i=1}^k (n_i - r_i) \frac{\partial P(d_i, \underline{\theta})}{\partial \theta_s} [1 - P(d_i, \underline{\theta})]^{-1}$$

$$= \sum_{i=1}^k \left[\frac{r_i}{P(d_i, \underline{\theta})} - \frac{n_i - r_i}{1 - P(d_i, \underline{\theta})} \right] \frac{\partial P(d_i, \underline{\theta})}{\partial \theta_s}$$

$$= \sum_{i=1}^k \left[\frac{r_i - n_i P(d_i, \underline{\theta})}{P(d_i, \underline{\theta}) [1 - P(d_i, \underline{\theta})]} \right] \frac{\partial P(d_i, \underline{\theta})}{\partial \theta_s}$$

$$= \sum_{i=1}^k \frac{n_i}{P(d_i, \underline{\theta})[1-P(d_i, \underline{\theta})]} \left[\frac{r_i}{n_i} - P(d_i, \underline{\theta}) \right] \frac{\partial P(d_i, \underline{\theta})}{\partial \theta_s}.$$

Using Leibniz' rule for differentiating an intergal,

$$\frac{\partial P(d_1, \underline{\theta})}{\partial \mu} = \int_{-\infty}^{y_1} \frac{\partial f(w)}{\partial \mu} dw + f(y_1) \frac{\partial y_1}{\partial \mu} - \lim_{a \rightarrow -\infty} f(a) \frac{da}{d\mu}$$

$$= (-1/\sigma)f(y_1),$$

$$\frac{\partial P(d_1, \underline{\theta})}{\partial \sigma} = \int_{-\infty}^{y_1} \frac{\partial f(w)}{\partial \sigma} dw + f(y_1) \frac{\partial y_1}{\partial \sigma} - \lim_{a \rightarrow -\infty} f(a) \frac{da}{d\sigma}$$

$$= (-y_1/\sigma)f(y_1),$$

$$\frac{\partial P(d_1, \underline{\theta})}{\partial m_j} = \int_{-\infty}^{y_1} \frac{\partial f(w)}{\partial m_j} dw + f(y_1) \frac{\partial y_1}{\partial m_j} - \lim_{a \rightarrow -\infty} f(a) \frac{da}{dm_j} \quad j = 1, 2$$

$$= \int_{-\infty}^{y_1} \frac{\partial \ln f(w)}{\partial m_j} f(w) dw \quad j = 1, 2.$$

Fixing either m_1 or $m_2 = 1$ greatly simplifies this last expression.

If $m_2 = 1$, then

$$P(d_1, \underline{\theta}_1) = \left[\frac{\exp(y_1)}{1 + \exp(y_1)} \right]^{m_1}$$

and

$$\frac{\partial P(d_1, \underline{\theta}_1)}{\partial m_1} = P(d_1, \underline{\theta}_1) \ln \left[\frac{\exp(y_1)}{1 + \exp(y_1)} \right].$$

If $m_1 = 1$, then

$$P(d_1, \underline{\theta}_2) = 1 - \left[\frac{\exp(-y_1)}{1 + \exp(-y_1)} \right]^{m_2}$$

and

$$\frac{\partial P(d_1, \underline{\theta}_2)}{\partial m_2} = [1 - P(d_1, \underline{\theta}_2)] \ln \left[\frac{\exp(-y_1)}{1 + \exp(-y_1)} \right].$$

Even with this simplification maximizing the log-likelihood equation to find $\hat{\underline{\theta}}_1' = [\hat{\mu}, \hat{\sigma}, \hat{m}_1]$ for the M1 model or $\hat{\underline{\theta}}_2' = [\hat{\mu}, \hat{\sigma}, \hat{m}_2]$ for the M2 model requires some type of iterative process.

For example, the method of scoring (Rao [1973]) can be used to iteratively find M.L.E.'s. Let $d \ln L / d \underline{\theta}$ be the "efficient score" for $\underline{\theta}$ with dimension s . The M.L.E. for $\underline{\theta}$ is that value for which the efficient score vanishes. Expand the efficient score as a first order Taylor series about some initial estimate $\underline{\theta}_0$ as

$$\frac{d \ln L}{d \underline{\theta}} = \left. \frac{d \ln L}{d \underline{\theta}} \right|_{\underline{\theta} = \underline{\theta}_0} + \left. \frac{d^2 \ln L}{d \underline{\theta}^2} \right|_{\underline{\theta} = \underline{\theta}_0} (\underline{\theta} - \underline{\theta}_0)$$

or

$$\frac{d \ln L}{d \underline{\theta}} = \underline{q} + \underline{A} (\underline{\delta} \underline{\theta})$$

where $[a_{ij}] = \partial \ln L / \partial \theta_i \partial \theta_j$. Setting this expression equal to zero implies $\underline{\delta} \underline{\theta} = -\underline{A}^{-1} \underline{q}$. The next starting value, $\underline{\theta}_1$ is equal to $\underline{\theta}_0 + \underline{\delta} \underline{\theta}$. This process is continued until $\underline{\delta} \underline{\theta}$, the difference between $\underline{\theta}_1$ and the next estimate $\underline{\theta}_{i+1}$ is arbitrarily small. This $\underline{\theta}_{i+1} = \hat{\underline{\theta}}$ is the M.L.E. for $\underline{\theta}$.

If SAS is available, it is quite easy to iteratively maximize the log-likelihood expression for the logit, M1, and M2 models. Let $\hat{p}_i = r_i/n_i$, then $E[\hat{p}_i] = P(d_i, \underline{\theta})$ and $\text{Var}[\hat{p}_i] = P(d_i, \underline{\theta})[1-P(d_i, \underline{\theta})]/n_i$. Hence $\hat{p}_i$ can also be expressed as

$$\hat{p}_i = P(d_i, \underline{\theta}) + \varepsilon_i$$

where

$$E[\varepsilon_i] = 0$$

and

$$\text{Var}[\varepsilon_i] = \frac{P(d_i, \underline{\theta})[1-P(d_i, \underline{\theta})]}{n_i}.$$

The nonlinear procedure in SAS (PROC NLIN), can be used to fit the model

$$\frac{\hat{p}_i}{\sqrt{v_i}} = \frac{P(d_i, \underline{\theta})}{\sqrt{v_i}} + \varepsilon_i^*$$

where $v_i = \text{Var} [\varepsilon_i]$ and $\text{Var} [\varepsilon_i^*] = 1$. PROC NLIN uses an iterative procedure to minimize the residual sum of squares $S(\underline{\theta}) = \sum_1^k [\hat{p}_i - P(d_i, \underline{\theta})]^2 / v_i$. The normal equations for the above model are

$$\frac{\partial S(\underline{\theta})}{\partial \theta_s} = \sum_1^k \frac{n_i}{P(d_i, \underline{\theta})[1-P(d_i, \underline{\theta})]} [\hat{p}_i - P(d_i, \underline{\theta})] \frac{P(d_i, \underline{\theta})}{\partial \theta_s}$$

which are the same as the normal equations for the M.L. technique, thus the solution supplied by SAS is also the solution to the maximum likelihood equations. Included in the SAS output is the model sum of squares, residual sum of squares, asymptotic standard errors of the parameter estimates, and asymptotic correlations.

The sum of squares residual is asymptotically distributed as a chi-square random variable. Considering each dose level separately

$$\left\{ \frac{[r_i - n_i P(d_i, \underline{\theta})]^2}{n_i P(d_i, \underline{\theta})} + \frac{(n_i - r_i - n_i [1-P(d_i, \underline{\theta})])^2}{n_i [1-P(d_i, \underline{\theta})]} \right\} \sim \chi_{(1)}^2 .$$

(Theorem 8 pg 444 Mood, Greybill, Boes[[1974]])

Once again assuming each dose level is independent, thus

$$\begin{aligned}
 & \sum_{i=1}^k \left\{ \frac{[r_i - n_i P(d_i, \underline{\theta})]^2}{n_i P(d_i, \underline{\theta})} + \frac{[-r_i - n_i P(d_i, \underline{\theta})]^2}{n_i [1 - P(d_i, \underline{\theta})]} \right\} \\
 &= \sum_{i=1}^k \left\{ [r_i - n_i P(d_i, \underline{\theta})]^2 \left[\frac{n_i}{n_i P(d_i, \underline{\theta})} + \frac{1}{n_i [1 - P(d_i, \underline{\theta})]} \right] \right\} \\
 &= \sum_{i=1}^k \left\{ \frac{[r_i - n_i P(d_i, \underline{\theta})]^2}{n_i P(d_i, \underline{\theta}) [1 - P(d_i, \underline{\theta})]} \right\} \\
 &= \sum_{i=1}^k \left\{ \frac{n_i [\hat{p}_i - P(d_i, \underline{\theta})]^2}{P(d_i, \underline{\theta}) [1 - P(d_i, \underline{\theta})]} \right\} \sim \chi^2_{(k)} .
 \end{aligned}$$

On substituting the maximum likelihood estimates for the q parameters in $\underline{\theta}$, the asymptotic distribution becomes chi-square with $k-q$ degrees of freedom. (Theorem 9 pg 446 Mood, Graybill, Boes [1974]). This sum of squares is a test statistic for goodness of fit of the proposed model.

3.3 Generalized Probit Procedure

The probit procedure in SAS (PROC PROBIT) is first employed to obtain estimates for μ and σ for the normal tolerance distribution. If there is control data, or if an estimate of natural mortality $\hat{c}$ is desired, it can be incorporated into the model as previously mentioned in §2.2; i.e.,

$$P(d_1, \gamma) = c + (1-c)P(d_1, \underline{\theta})$$

where $\underline{\gamma}' = [\mu, \sigma, m_1, m_2, c]$ and

$$P(d_1, \underline{\theta}) = \int_{-\infty}^{y_1} f(w) dw .$$

Also $\partial P(d_1, \underline{\gamma}) / \partial c = 1 - P(d_1, \underline{\theta})$.

The estimates for μ and σ from PROC PROBIT are now used as starting values, which the user must supply to the nonlinear procedure, to fit the model

$$\frac{\hat{p}_1}{\sqrt{v_1}} = \frac{P(d_1, \underline{\theta})}{\sqrt{v_1}} + \varepsilon_1^*$$

where $P(d_1, \underline{\theta}) = \exp(y_1) / [1 + \exp(y_1)]$ or the logistic cumulative density function. The resulting parameter estimates will be maximum likelihood estimates for the logit model.

These estimates for μ and σ from the logit model in turn serve as starting values or initial estimates to the M1 and M2 models. The M1 model is

$$P(d_1, \underline{\theta}_1) = \left[\frac{\exp(y_1)}{1 + \exp(y_1)} \right]^{m_1}$$

and the M2 model is

$$P(d_1, \theta_2) = 1 - \left| \frac{\exp(-y_1)}{1 + \exp(-y_1)} \right|^{m_2}.$$

A starting value of one for m_1 in the M1 model and m_2 in the M2 model is used to see if either model can improve on the logit fit. See Table 3.31 for an example of the SAS coding needed.

There are four different iterative methods available in the SAS79 version of PROC NLIN. Three require derivatives and one is a non-derivative method. The derivative methods are: a modified Gauss-Newton method, Marquardt method, and gradient or steepest descent method. The nonderivative method, or DUD method is used when no derivatives are supplied by the user. For explanations of the computational procedure of each method see the SAS User's Guide, 1979 Edition, pp 317-318.

In practice, the Gauss-Newton method, which is very similar to the scoring method, was superior when fitting the M1 or M2 model. Using the estimates from the first run with the M1 or M2 model as starting values for another run using the Gauss-Newton or DUD method usually resulted in a slight improvement of fit as measured by a decrease in the residual sum of squares.

For example, using Data Set #1;

Step 1: Probit procedure [using $\log_{10}$ (dose)]

results; $\hat{\mu} = .1164$, $\hat{\sigma} = .10485$

residual SS; 23.0499 (5 df.)

TABLE 3.31

Example of SAS Program for M1 Model

```

DATA CM10;
INPUT D N Z P;
X = LOG10 (D);
C = [threshold probability estimate from control data or PROC PROBIT];
CARDS:
.
.
.
DATA
.
.
.
PROC NLIN METHOD = GAUSS:
PARAMETERS U = [logit estimate] S = [logit estimate] M1 = 1;
BOUNDS S>0 0<M1<57;
Y = (X-U)/S; M2 = 1;
G1 = GAMMA (M1); G2 = GAMMA (M2); G12 = GAMMA (M1 + M2);
BETA = G1 * G2/G12;
PD1 = (EXP(Y)/(1 + EXP(Y))) ** M1;
PD = C + PD1 * (1 - C); W = N/(PD * (1 - PD));
WR = SQRT(W); PK = P * WR; PDW = PD * WR;
MODEL PK = PDW;
F = (((EXP(Y)) ** M1) * (1 + EXP(Y)) ** (-M1 -M2))/BETA
F1 = F * (1 - C);
DER. U = WR * (-F1)/S;
DER. S = WR * (-Y) * F1/S;
DER. M1 = WR * PD1 * LOG (EXP(Y)/(1 + EXP(Y))) * (1 - C);

```

Step 2: Nonlinear Procedure [Logit model, common log trans.]

Method; Gauss

Starting values; $\mu = .1164$, $\sigma = .10485$

results; $\hat{\mu} = .1117$, $\hat{\sigma} = .05781$

residual SS; 21.3715 (5 df.)

Step 3: NLIN [Logit model]

Method; DUD

results; negligible reduction in residual SS

Step 4: NLIN [M1 model]

method; Gauss

starting values; $\mu = .1117$, $\sigma = .05781$, $m_1 = 1$

results; slight reduction in residual SS

Step 5: NLIN [M2 model]

method; Gauss

starting values; $\mu = .1117$, $\sigma = .5781$, $m_2 = 1$

results; $\hat{\mu} = .0267$, $\hat{\sigma} = .01375$, $\hat{m}_2 = .1438$

residual SS; 9.4799 (4 df.)

Step 6: NLIN [M2 model]

method; DUD

starting values; $\mu = .0267$, $\sigma = .01375$, $m_2 = .1438$

results; $\hat{\mu} = .0237$, $\hat{\sigma} = .01349$, $\hat{m}_2 = .1346$

residual SS; 9.3020 (4 df.)

Step 7: NLIN [M2 model]

method; Gauss

results; negligible reduction in residual SS.

Table 3.32 indicates the results, supplied by the SAS output, from fitting Data Set #1 with the probit, logit, and M2 model. Table 3.33 contains the results from fitting Data Set #2 with the provit, logit, and M1 model. In this case, the M1 model outperformed the M2 model. Since there was no control data, c was set equal to zero.

3.4 Effective Dose Estimation

The effective dose- q is that estimated dose which will cause a response in the q^{th} quantile of the population. For example, ED05 is the estimated dose which will produce a response in 5% of the population. Recall in the M2 model, the probability of a response at dose d is given by

$$P(d, \underline{\theta}) = 1 - \left[\frac{\exp(-y)}{1 + \exp(y)} \right]^{m_2}$$

where $y = (x - \mu)/\sigma$ and $x = \log_{10}(\text{dose})$. Fixing $P(d, \underline{\theta})$ at a desired risk level; i.e. $P(d, \underline{\theta}) = q = .10$ and solving for x gives

$$\begin{aligned} 1 - [1 + \exp(y)]^{-m_2} &= q \\ [1 + \exp(y)] &= (1-q)^{-m_2^{-1}} \\ \exp(y) &= (1-q)^{-m_2^{-1}} - 1 \end{aligned}$$

TABLE 3.32

Results for Data Set #1

Probit (PROC PROBIT)

<u>Parameter</u>	<u>Estimate</u>	<u>Asymptotic Covariance</u>		<u>Residual SS (df)</u>
μ	.1164	0.0004219	0.0000272	23.0499 (5)
σ	.10485	0.0000272	0.0006749	

Logit (PROC NLIN)

<u>Parameter</u>	<u>Estimate</u>	<u>Asy. Std. Error</u>	<u>Asymptotic Correlation</u>		<u>Residual SS (df)</u>
μ	.1117	.01896	1.0000	0.1532	21.3715 (5)
σ	.05781	.01536	0.1532	1.0000	

M2 Model (PROC NLIN)

<u>Parameter</u>	<u>Estimate</u>	<u>Asy. Std. Error</u>	<u>Asymptotic Correlation</u>			<u>Residual SS (df)</u>
μ	.0240	0.02416	1.0000	0.8669	0.9013	9.3020 (4)
σ	.01363	0.02114	0.8669	1.0000	0.9932	
m_2	.1364	0.23323	0.9013	0.9932	1.0000	

TABLE 3.33

Results for Data Set #2

Probit (PROC PROBIT)

<u>Parameter</u>	<u>Estimate</u>	<u>Asymptotic Covariance</u>		<u>Residual SS (df)</u>
μ	.4787	-0.0001727	-0.0000785	10.5501 (4)
σ	.07804	-0.0000785	0.0001813	

Logit (PROC NLIN)

<u>Parameter</u>	<u>Estimate</u>	<u>Asy. Std. Error</u>	<u>Asymptotic Correlation</u>		<u>Residual SS (df)</u>
μ	.4779	0.01422	1.0000	-0.3975	11.8448 (4)
σ	.04626	0.00924	-0.3975	1.0000	

M1 Model (PROC NLIN)

<u>Parameter</u>	<u>Estimate</u>	<u>Asy. Std. Error</u>	<u>Asymptotic Correlation</u>		<u>Residual SS (df)</u>
μ	.5780	0.03425	1.0000	-0.8777 -0.9511	5.2952 (3)
σ	.02112	0.01347	-0.8777	1.0000 0.9616	
m_1	.1689	0.15980	-0.9511	0.9616 1.0000	

$$\frac{x-\mu}{\sigma} = \ln [(1-q)^{-m_2} - 1]$$

$$x = \mu + \sigma \ln [(1-q)^{-m_2} - 1] .$$

The antilog of x is ED10. Due to the invariance property of maximum likelihood estimators, substituting $\hat{\mu}$, $\hat{\sigma}$, and $\hat{m}_2$ in the above expression, and taking the antilog of the resulting $\hat{x}$ gives the M.L.E. for the effective dose. Similarly, effective dose estimates for the M1 model are given by the antilog of $\hat{\chi}$ where

$$\hat{\chi} = \hat{\mu} - \hat{\sigma} \ln(q^{-\hat{m}_1} - 1) .$$

The effective dose is a function of $\hat{\theta}_j = [\hat{\mu}, \hat{\sigma}, \hat{m}_j]$, $j = 1, 2$. To find the asymptotic standard error of the effective dose, (Kendall and Stuart Vol. II [1952]) expand $g(\hat{\theta})$ as a first order Taylor series about $\underline{\theta}$ as

$$g(\hat{\theta}) = g(\underline{\theta}) + \frac{\partial f(\hat{\theta})}{\partial \hat{\mu}} \bigg|_{\hat{\theta}=\underline{\theta}} (\hat{\mu}-\mu) + \frac{f(\hat{\theta})}{\partial \hat{\sigma}} \bigg|_{\hat{\theta}=\underline{\theta}} (\hat{\sigma}-\sigma) + \frac{\partial f(\hat{\theta})}{\partial m_j} \bigg|_{\hat{\theta}=\underline{\theta}} (\hat{m}_j - m_j) \quad j=1,2$$

Now

$$\text{Var}[f(\hat{\theta})] = \text{Var}\left[\frac{\partial f(\hat{\theta})}{\partial \hat{\mu}} \hat{\mu} + \frac{\partial f(\hat{\theta})}{\partial \hat{\sigma}} \hat{\sigma} + \frac{\partial f(\hat{\theta})}{\partial \hat{m}_j} \hat{m}_j\right] \quad j=1,2$$

$$= \underline{c}' \underline{\Sigma} \underline{c}$$

where

$$\underline{c}' = \left[\frac{\partial f(\underline{\theta})}{\partial \hat{\mu}}, \frac{\partial f(\underline{\theta})}{\partial \hat{\sigma}}, \frac{\partial f(\underline{\theta})}{\partial \hat{m}_j} \right] \quad j = 1, 2$$

and $\underline{\Sigma}$ is the covariance matrix of $\underline{\theta}$.

For the M2 model

$$\frac{\partial g(\underline{\theta})}{\partial \hat{\mu}} = 1$$

and

$$\frac{\partial g(\underline{\theta})}{\partial \hat{q}} = \ln [(1-q)^{-\hat{m}_2 - 1}]$$

and

$$\frac{\partial g(\underline{\theta})}{\partial \hat{m}_2} = \left[\frac{\partial u}{\partial \hat{m}_2} \right] \left[\frac{\partial f(\underline{\theta})}{\partial u} \right]$$

where $u = [(1-q)^{-\hat{m}_2 - 1} - 1]$. Finally,

$$\frac{\partial g(\underline{\theta})}{\partial \hat{m}_2} = \{ (1-q)^{-\hat{m}_2 - 1} \ln(1-q)^{\hat{m}_2 - 2} \} \left\{ \frac{\sigma}{[(1-q)^{-\hat{m}_2 - 1} - 1]} \right\}$$

$$\hat{\sigma} [1 - (1-q)^{\hat{m}_2 - 1}]^{-1} \ln(1-q)^{\hat{m}_2 - 2} .$$

Similarly, for the M1 model

$$\underline{c}' = [1, -\ln(q^{\hat{m}_1 - 1} - 1), -\hat{\sigma}^2 \hat{m}_1^{-2} [1 - q^{\hat{m}_1 - 1}]^{-1} \ln q] .$$

Although SAS does not supply $\underline{\Sigma}$, it is readily obtainable.

$\underline{\Sigma} = \underline{B}' \underline{A} \underline{B}$ where

$$\underline{B} = \begin{bmatrix} \text{std. error } \hat{\mu} & & \\ & \text{std. error } \hat{\sigma} & \underline{0} \\ \underline{0} & & \text{std. error } \hat{m}_j \end{bmatrix} \quad j = 1, 2$$

and $\underline{A}$ is the asymptotic correlation matrix.

For the first data set, using the logit model

$$\underline{\Sigma}_L = \begin{bmatrix} 0.000360 & 0.000045 \\ 0.000045 & 0.000236 \end{bmatrix} .$$

Using the M2 model

$$\underline{\Sigma}_{M2} = \begin{bmatrix} 0.000584 & 0.000443 & 0.005079 \\ 0.000443 & 0.000447 & 0.004897 \\ 0.005079 & 0.004897 & 0.054396 \end{bmatrix} .$$

For the second data set, using the logit model

$$\underline{\Sigma}_L = \begin{bmatrix} 0.000202 & -0.000052 \\ -0.000052 & 0.000085 \end{bmatrix} .$$

Using the M1 model

$$\underline{\Sigma}_{M1} = \begin{bmatrix} 0.001014 & -0.000353 & -0.004517 \\ -0.000353 & 0.000162 & 0.001825 \\ -0.004517 & 0.001825 & 0.022377 \end{bmatrix} .$$

The BAN (Best Asymptotically Normal) property of M.L.E.'s (Theorem 18 pg. 359; Mood, Greybill, Boes [1974]) implies an estimate of an effective dose, or $\hat{X}$, is asymptotically distributed $N(\mu_X^0, \sigma_X^2)$. Then

$$P[-Z_{\alpha/2} \leq \frac{\hat{X} - \mu_X^0}{\sigma_X} \leq Z_{\alpha/2}] \stackrel{a}{=} 1-\alpha ,$$

where $Z_{\alpha/2}$ is a standard normal deviate corresponding to a probability of $\alpha/2$. This in turn gives

$$P[\hat{X} - Z_{\alpha/2} \sigma_X \leq \mu_X^0 \leq \hat{X} + Z_{\alpha/2} \sigma_X] \stackrel{a}{=} 1-\alpha ,$$

which provides a $(1-\alpha)\%$ asymptotic confidence interval about the effective dose.

IV. EMPIRICAL COMPARISONS OF VARIOUS MODELS

4.1 V.S.D. for Mantel-Bryan Procedure and Generalized Probit

In order to expediate finding a virtually safe dose using the Mantel-Bryan procedure, a different but similar criteria was used to determine an upper 99% limit on the probit intercept a . Instead of finding the value of a that leads to a just significant reduction in the log-likelihood function (see §2.1), the value is found that causes a just significant increase in the sum of squares residual in the weighted β -model.

The nonlinear procedure in SAS is used in the same manner as when fitting the logit or generalized probit model. Here $P(d_i, \theta)$ is the cumulative standard normal density and $y = a + x$ where $x = \log_{10}(\text{dose})$. Recall in this procedure the probit slope b is fixed at one. A parameter c , denoting the probability of a spontaneous response, is included in the model as presented in §2.1 and §3.1. Maximum likelihood estimates are determined for a and c . The sum of squares residual with $k-2$ degrees of freedom (k dose levels, 2 parameters estimated) is noted.

Next, a range of values for a , along with their corresponding $\hat{c}$ values are evaluated in the model until a significant increase in the residual sum of squares is observed. This sum of squares has $k-1$ degrees of freedom, a now being fixed. Recall

$$\text{SS Residual} \approx \chi^2_{(k-q)}$$

where k = dose levels, q = parameters estimated. Also

$$[\text{SS Residual } (a, \hat{c}) - \text{SS Residual } (\hat{a}, \hat{c})] \hat{\sigma}^2 \chi^2_{(1)}.$$

Hence, an increase of $\chi^2_{.99(1)} = 6.635$ reflects a just significant worsening of the fit.

That value for a , say a^* , is the upper 99% confidence limit for a . The highest dose level is now deleted and the above procedure is repeated. The deletion continues until only the lowest nonzero dose level is fit.

For example, fitting the data in Data Set #2 to the Mantel-Bryan model to find $\hat{a}$ and $\hat{c}$ yielded SS Residual = 88.16. A fixed value of $a^* = .064$ along with a corresponding $\hat{c} = .017$ gave SS Residual = 94.8. The highest dose level was deleted and the process repeated.

The V.S.D. is then calculated by using the lowest of the upper limits on a . If a 10^{-8} increase in risk is the virtually safe level desired, then

$$10^{-8} = \int_{-\infty}^{a^* + \log_{10}(\text{dose})} \phi(u) du.$$

This implies $-5.612 = a^* + \log_{10}(\text{dose})$ or the V.S.D. = antilog $(-5.612 - a^*)$. Table 4.11 shows the results for Data Set #2 for finding a^* . In this case $a^* = -0.541$. Using 10^{-8} as the criteria for virtual safety, this gives a V.S.D. of $8.5(10^{-6})$.

A V.S.D. or $ED10^{-8}$ for the logit model, where $\hat{\theta}_L = [.4779, .04626]$ and $q = 10^{-8}$ is given by $\log_{10}(\text{V.S.D.}) = \hat{x}_L = .4779 - .04626 \ln [q^{-1} - 1] = -0.2684$. This implies that the V.S.D. is .5390. See §3.4 for derivations of these expressions. The asymptotic variance for $\hat{x}_L$ is given

TABLE 4.11

Upper limit on the probit intercept for Data Set #2

Dose	a*	$\hat{c}$
1	-0.441	0.035
2	-0.541	0.040
3	-0.346	0.050
4	-0.144	0.045
5	-0.038	0.050

Note: a* calculated using all dose levels less than or equal to level indicated.

by $\underline{c}'\underline{\Sigma}_L\underline{c}$ where $\underline{a}' = 1, 16.1181$ and from §3.4

$$\underline{\Sigma}_L = \begin{bmatrix} 0.000202 & -0.000052 \\ -0.000052 & 0.000086 \end{bmatrix}.$$

A 99% asymptotic confidence interval about X_L would be $[-0.2684 - 2.576(.1551), -0.2684 + 2.576(.1551)]$ or $[-0.6680, 0.1313]$. Since the antilog is a monotone increasing function, this confidence interval in terms of dose rather than $\log_{10}$ (dose) is $[0.2148, 1.3530]$.

Again using the expressions presented in 3.4 the M1 model, where $\underline{\hat{\theta}}'_{M1} = [.5780, .02112, .1689]$, gives $\hat{X}_{M1} = \log_{10} (ED10^{-8}) = -1.4487$. This yields a V.S.D. of .03559. The asymptotic variance of $\hat{X}_{M1}$ is 0.48241. A 99% asymptotic confidence interval about this V.S.D. is $[0.0005782, 0.34048]$.

4.2 Empirical Comparisons of Effective Doses

Figure 4.21 is a plot exhibiting the behavior of the standard errors of the logit and M2 model effective doses-05 through 95, for Data Set #1. In this particular case the M2 standard errors are consistently lower than the logit's, however they gradually overtake the logit's standard errors at both extremes. Figure 4.22 is a similar plot comparing the logit and M1 model for Data Set #2. Here the standard errors for the M1 model are lower for ED45 through ED95, and higher for ED's less than ED45.

Table 4.21 contains the V.S.D. estimate and ED05 - ED95 produced by the logit and M2 models, and their corresponding standard errors for

FIGURE 4.21

PLOT OF ASYMPTOTIC STD ERRORS OF ESTIMATES
AT GIVEN PROBABILITIES

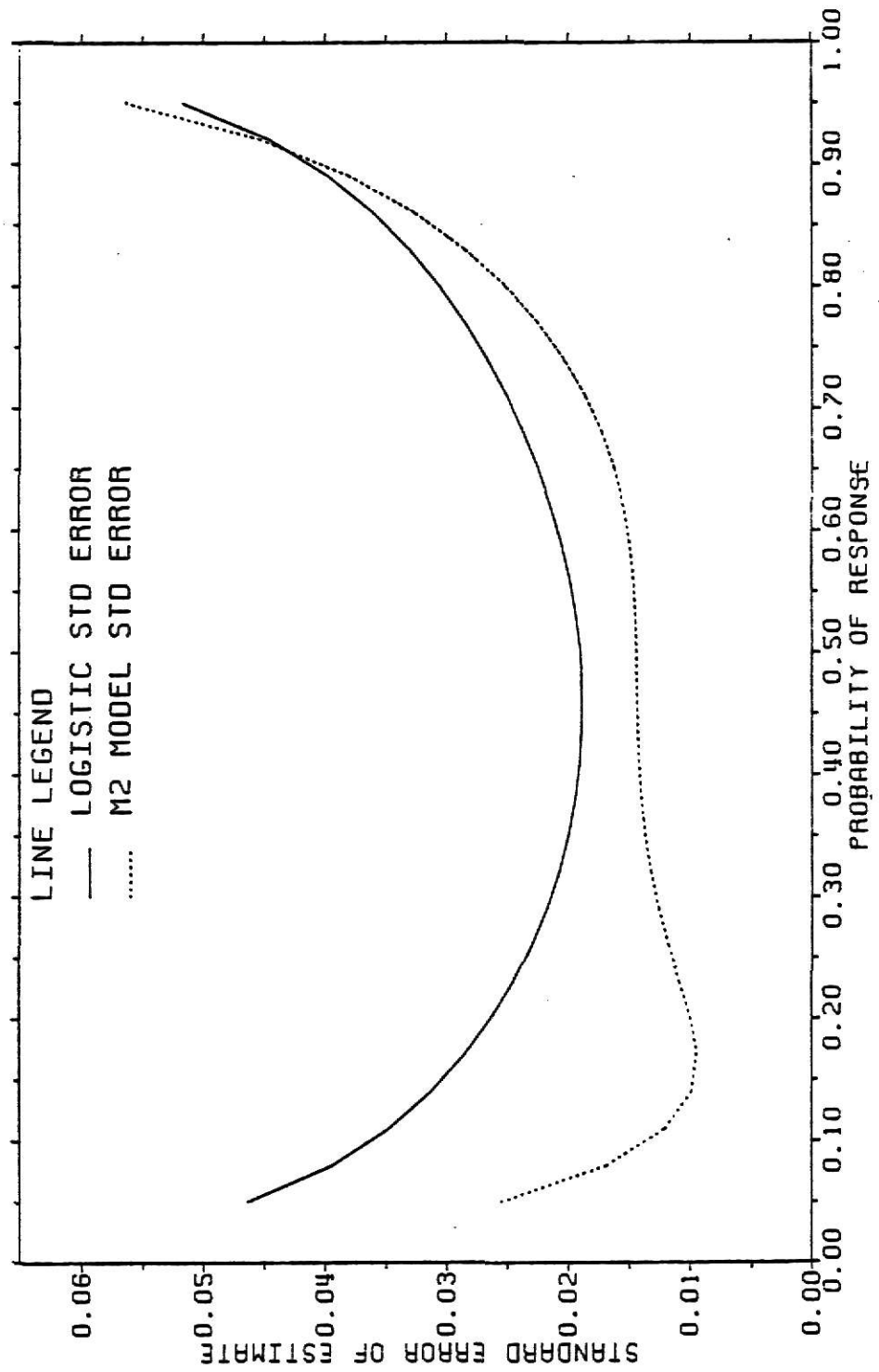
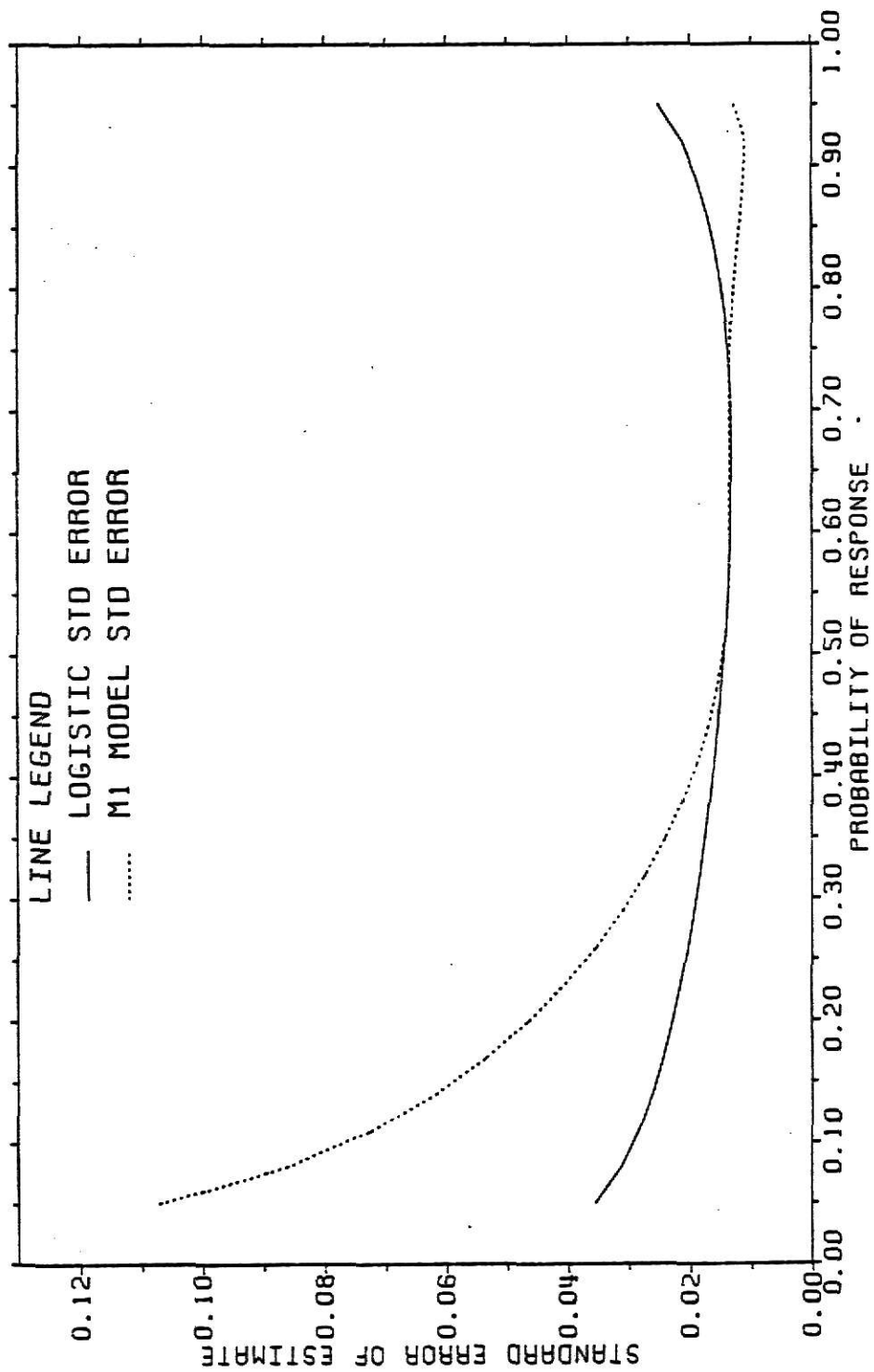


FIGURE 4.22

PLOT OF ASYMPTOTIC STD ERRORS OF ESTIMATES
AT GIVEN PROBABILITIES

Data Set #1. The effective doses are in terms of common logs of the doses. Table 4.22 displays the same information yielded by the logit and M1 models for Data Set #2.

Table 4.23 lists 95% asymptotic confidence intervals for ED05 - ED95 provided by the logit and M2 models for Data Set #1. Here, the ED's have been converted into the original dose metameters, rather than $\log_{10}$ (dose). Table 4.24 contains similar information given by the logit and M1 models for Data Set #2.

Table 4.23 is graphically illustrated in Figures 4.23 and 4.24. Table 4.24 is illustrated in Figures 4.25 and 4.26.

4.3 Goodness of Fit Comparisons

It was shown in §3.2 that the sum of squares residual from the non-linear procedure in SAS, when fitting the weighted $\hat{p}$ model, is an asymptotic chi-square random variable with $k-q$ degrees of freedom; k = number of dose levels, q = number of parameters estimated. This can be then used as a goodness of fit test statistic. Table 4.31 displays the p values, or the probability of obtaining a sum of squares residual of that magnitude given a particular model. The p values for the λ power transformations were taken from Egger [1979].

TABLE 4.21

Estimates of $\log_{10}(\text{dose})$ needed to ellicit a response
in $q \times 100\%$ of the population represented in Data Set #1

q	Logit Estimate	Asymptotic Std. Error	M2 Estimate	Asymptotic Std. Error
10^{-8}	-.8199	.24539	-.1681	.30095
.05	-.0585	.04629	.0133	.02542
.10	-.0153	.03609	.0261	.01326
.15	.0114	.03025	.0353	.00949
.20	.0316	.02626	.0433	.00996
.25	.0482	.02338	.0509	.01152
.30	.0627	.02130	.0586	.01285
.35	.0759	.01988	.0669	.01371
.40	.0883	.01904	.0746	.01416
.45	.1001	.01875	.0834	.01434
.50	.1117	.01897	.0930	.01442
.55	.1233	.01968	.1036	.01462
.60	.1351	.02085	.1153	.01515
.65	.1475	.02248	.1287	.01627
.70	.1607	.02459	.1440	.01816
.75	.1752	.02725	.1622	.02105
.80	.1918	.03061	.1845	.02522
.85	.2120	.03499	.2132	.03118
.90	.2387	.04117	.2536	.04019
.95	.2819	.05165	.3227	.05631

TABLE 4.22

Estimates of $\log_{10}(\text{dose})$ needed to elicit a response
in $q \times 100\%$ of the population represented in Data Set #2

q	Logit Estimate	Asymptotic Std. Error	MI Estimate	Asymptotic Std. Error
10^{-8}	-.2684	.15513	-1.4487	.83340
.05	.3416	.03536	.2013	.10695
.10	.3762	.02905	.2885	.07626
.15	.3976	.02530	.3395	.05852
.20	.4137	.02261	.3756	.04613
.25	.4270	.02049	.4037	.03679
.30	.4387	.01876	.4266	.02948
.35	.4492	.01731	.4460	.02376
.40	.4591	.01608	.4629	.01940
.45	.4686	.01505	.4778	.01632
.50	.4779	.01422	.4912	.01444
.55	.4872	.01359	.5034	.01355
.60	.4967	.01319	.5148	.01333
.65	.5066	.01305	.5255	.01340
.70	.5171	.01323	.5359	.01345
.75	.5288	.01380	.5461	.01325
.80	.5421	.01488	.5565	.01268
.85	.5582	.01667	.5677	.01172
.90	.5796	.01960	.5810	.01082
.95	.6142	.02520	.5999	.01288

TABLE 4.23

95% asymptotic confidence intervals about
ED05 to ED95 for Data Set #1

ED	Logit Estimate	Logit Lower Limit	Logit Upper Limit	M2 Estimate	M2 Lower Limit	M2 Upper Limit
05	0.8740	0.7092	1.0770	1.0312	0.9194	1.1565
10	0.9654	0.8203	1.1362	1.0619	1.0002	1.1274
15	1.0267	0.8957	1.1767	1.0846	1.0391	1.1321
20	1.0754	0.9552	1.2107	1.1048	1.0562	1.1557
25	1.1174	1.0052	1.2417	1.1244	1.0674	1.1844
30	1.1554	1.0495	1.2720	1.1443	1.0798	1.2126
35	1.1910	1.0888	1.3028	1.1651	1.0952	1.2395
40	1.2254	1.1245	1.3353	1.1874	1.1139	1.2658
45	1.2592	1.1571	1.3704	1.2118	1.1359	1.2928
50	1.2933	1.1872	1.4089	1.2389	1.1608	1.3222
55	1.3283	1.2154	1.4517	1.2693	1.1883	1.3559
60	1.3650	1.2424	1.4997	1.3042	1.2180	1.3965
65	1.4044	1.2688	1.5544	1.3448	1.2496	1.4473
70	1.4477	1.2956	1.6176	1.3933	1.2836	1.5123
75	1.4969	1.3237	1.6928	1.4529	1.3212	1.5976
80	1.5554	1.3547	1.7858	1.5292	1.3647	1.7135
85	1.6292	1.3911	1.9079	1.6336	1.4192	1.8805
90	1.7326	1.4588	2.0864	1.7930	1.4956	2.1496
95	1.9138	1.5158	2.4162	2.1023	1.6305	2.7106

FIGURE 4.23
95% LOGISTIC C.I. ABOUT
ED05 TO ED95

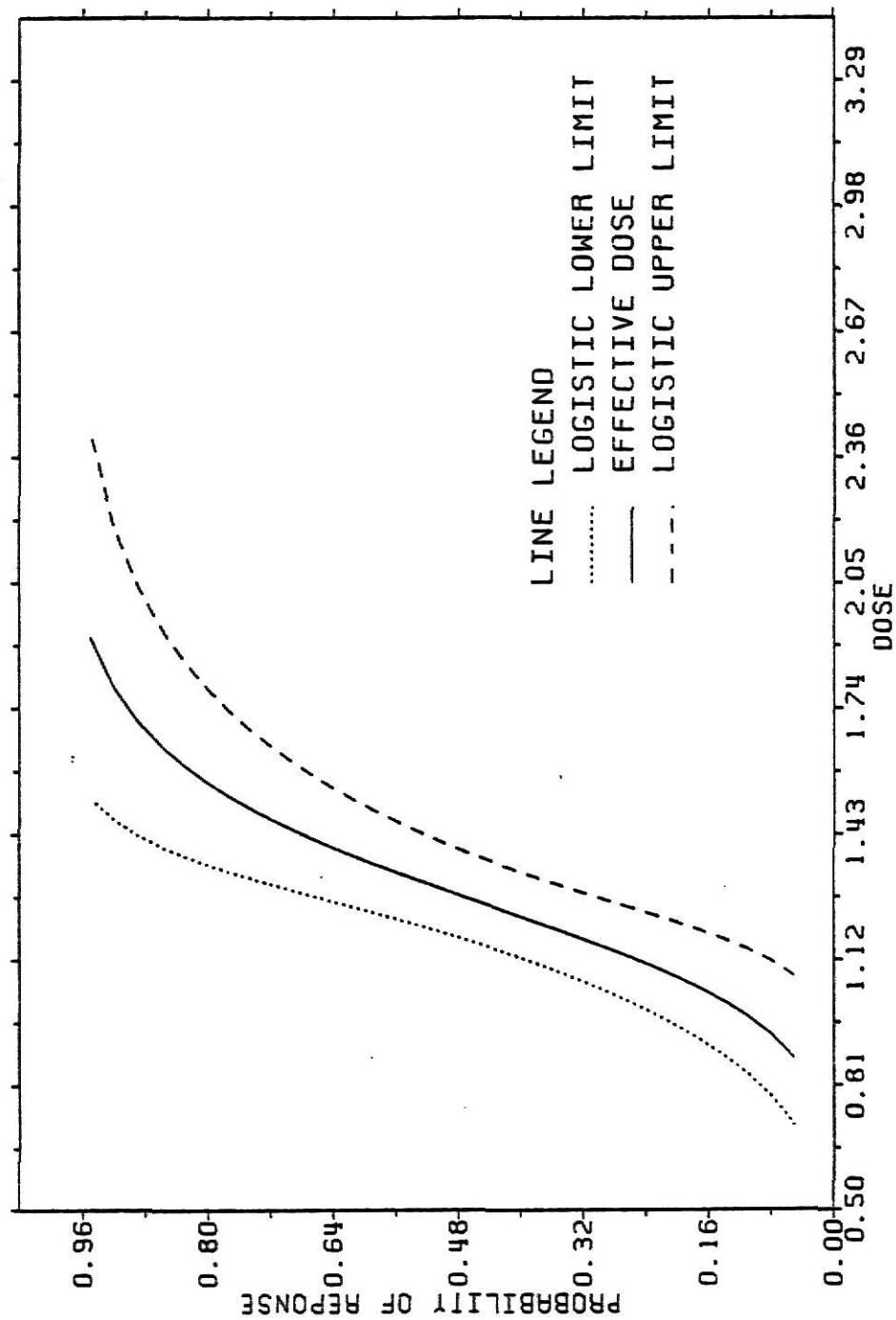


FIGURE 4.24

95% M2 MODEL C.I. ABOUT
ED05 TO ED95

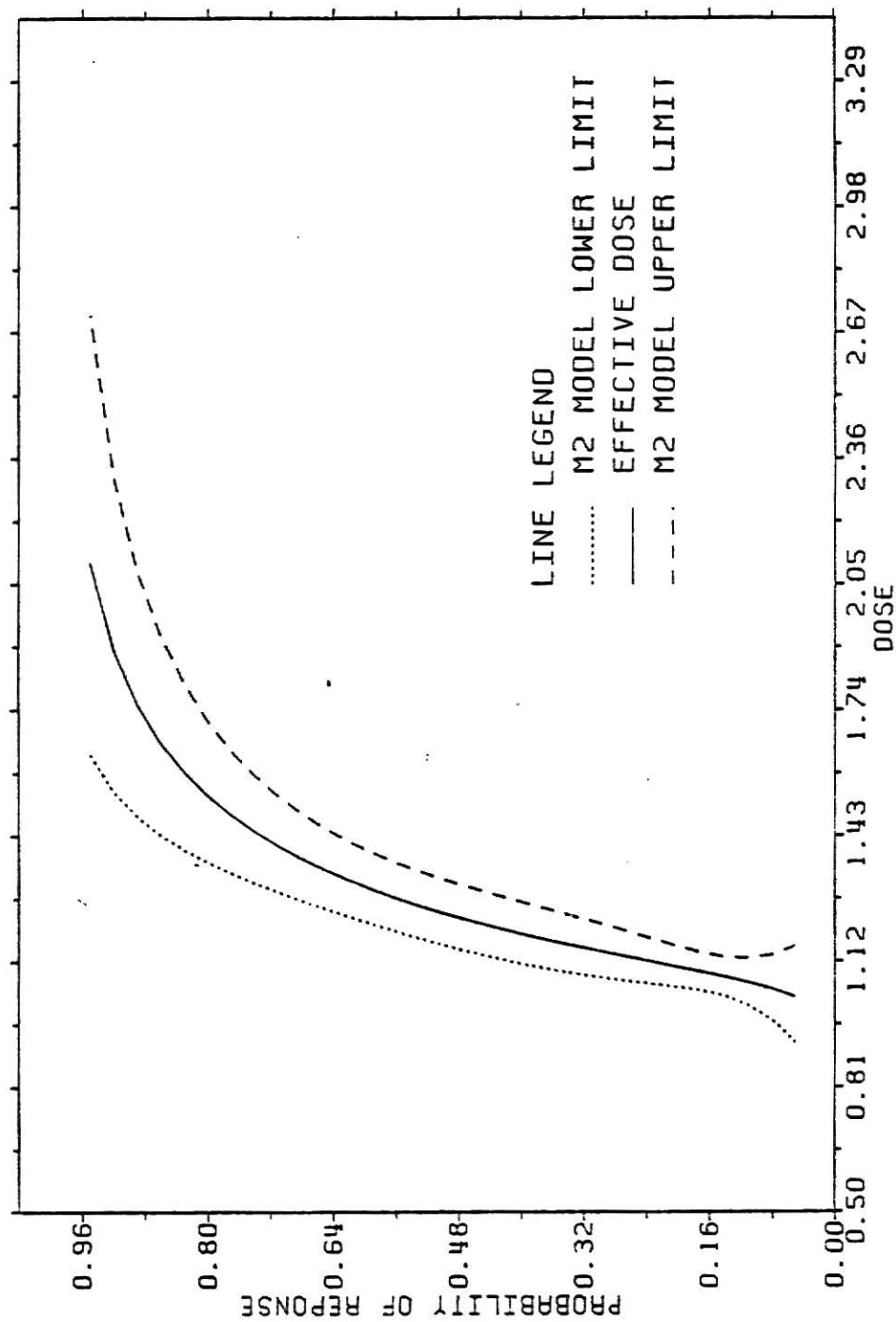


TABLE 4.24

95% asymptotic confidence intervals about
ED05 to ED95 for Data Set #2

ED	Logit Estimate	Logit Lower Limit	Logit Upper Limit	M1 Estimate	M1 Lower Limit	M1 Upper Limit
05	2.1957	1.8719	2.5755	1.5897	0.9811	2.5759
10	2.3778	2.0856	2.7109	1.9430	1.3772	2.7412
15	2.4980	2.2284	2.8002	2.1850	1.6779	2.8454
20	2.5925	2.3410	2.8710	2.3748	1.9284	2.9245
25	2.6732	2.4371	2.9323	2.5333	2.1458	2.9908
30	2.7458	2.5229	2.9884	2.6707	2.3380	3.0508
35	2.8134	2.6020	3.0420	2.7928	2.5082	3.1089
40	2.8782	2.6768	3.0949	2.9032	2.6598	3.1689
45	2.9418	2.7486	3.1486	3.0046	2.7912	3.2343
50	3.0054	2.8186	3.2046	3.0988	2.9033	3.3075
55	3.0704	2.8877	3.2646	3.1875	2.9984	3.3885
60	3.1382	2.9568	3.3306	3.2720	3.0810	3.4750
65	3.2104	3.0268	3.4052	3.3539	3.1570	3.5630
70	3.2895	3.0989	3.4919	3.4346	3.2323	3.6495
75	3.3788	3.1748	3.5960	3.5162	3.3120	3.7330
80	3.4841	3.2577	3.7261	3.6017	3.4013	3.8138
85	3.6159	3.3538	3.8984	3.6961	3.5057	3.8969
90	3.7987	3.4768	4.1503	3.8106	3.6290	4.0014
95	4.1136	3.6715	4.6090	3.9804	3.7556	4.2185

FIGURE 4.25
95% LOGISTIC C.I. ABOUT
ED05 TO ED95

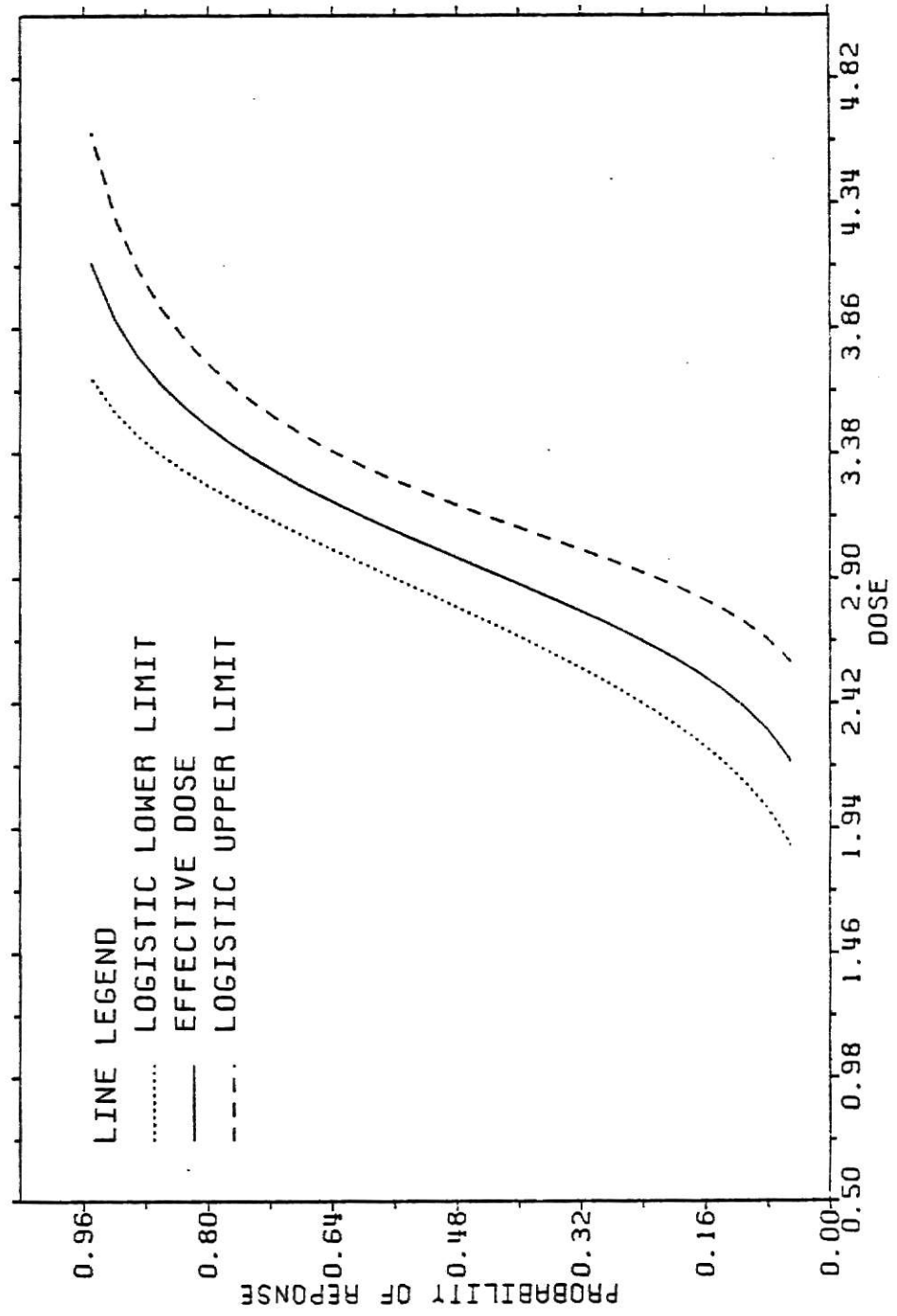


FIGURE 4.26

95% M1 MODEL C.I. ABOUT
ED05 TO ED95

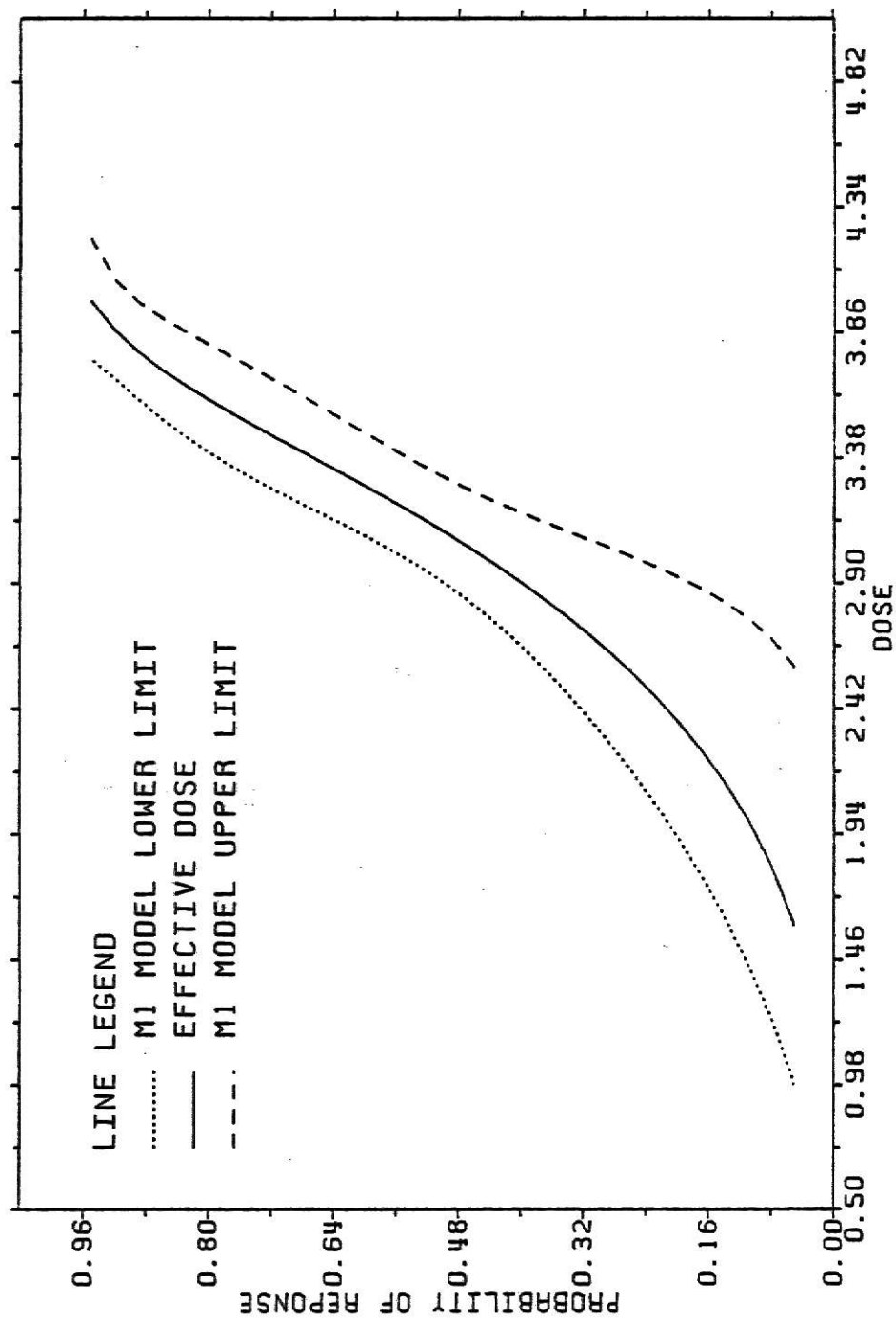


TABLE 4.31

Goodness of Fit p values

	Data Set #1		Data Set #2	
Model	p value	df.	p value	df.
Probit	.00033	5	.0321	4
Logit	.00069	5	.0185	4
χ^2 -logit	.0056	3	.1405	2
χ^2_{H} -logit	.0112	4	.1373	2
Generalized Probit	.0540	4	.1543	3

p = prob chi-square > SSRes.

V. DISCUSSION

Low dose risk assessment, especially involving a potential carcinogen, can possess a certain appeal along with its many problems. The main problems, the lack of knowledge concerning the cancer mechanism and the virtually impossible task of obtaining observed values in the area of interest, make it difficult to evaluate the "correctness" of a low dose risk assessment.

Armed with this lack of knowledge, quite a variety of procedures have emerged in the literature to attack the question, "What is a safe dose of this potential carcinogenic agent?". The F.D.A., the regulatory arm of the government concerning such agents that we are exposed to involuntarily, has adopted the Mantel-Bryan procedure. The typically poor fit to the dose-response curve, as illustrated by a sum of squares residual of 88.16 compared to 5.295 given by the M1 model for Data Set #2, intuitively raises questions of the validity of V.S.D. estimates from a statistical standpoint.

Schneiderman [1974] calls attention to the fact that the classical models used in Bioassay work, the probit and logit, are usually indistinguishable over the range of observed values; however, when extrapolating to a virtually safe level differences often do appear with the probit model generally giving the highest V.S.D. estimate. This incongruity leaves both results suspect. The M1 and M2 models did differ from both the logit and probit models in the range of observed values. The sum of squares residual derived for the first data set were 23.0499 and 21.3715 for the probit and logit models respectively, and 9.302 for

the M2 model. Similarly, sum of squares residuals of 10.5501, 11.8448, and 5.2952 for the probit, logit, and M1 models respectively resulted for the second data set.

The more flexible generalized probit model, in the form of the M1 or M2 model, when presented with an apparently nonnormal or nonsymmetric sigmoid response curve, gave a much better fit. This improvement came at a price. The asymptotic standard errors of effective dose estimates at the lower end of the response curve are larger for both the M1 and M2 models than for the logit model. Perhaps this increase is due to the necessity of an additional parameter estimate, the relatively high asymptotic correlations, and the flexibility of the model. Finding a better approximation for these standard errors may be a worthwhile area for future investigation.

The generalized probit model is not being presented here as a substitute for either the probit or logit model, but rather an extension, for they are simply special cases of the generalized probit model. SAS makes the M1 and M2 versions practical to use. A lower 99% limit on an effective dose at a desired "virtually safe" level has a basis in traditional statistical concepts and techniques, making claims of conservatism and precision justifiable.

REFERENCES

1. Altschuler, B. (1977). A Bayesian Approach to Assessing Population Risks from Environmental Carcinogens, Environmental Health, SIAM, 31-44.
2. Copenhaver, T. W. and P. W. Mielke (1977). Quantal Analysis: A Generalized Tolerance Distribution for Quantal Response Assay Analysis, presented at Joint Statistical Meetings, Chicago, Illinois, August 1977.
3. Cornfield, J. (1977). Carcinogenic Risk Assessment, Science, 198:693-699.
4. Crump, K. S. (Dec. 1977). Response to Salsburg's Open Query, Biometrics, 33:752-755.
5. Egger, M. J. (1979). Power Transformations to Achieve Symmetry in Quantal Bioassay, Technical Report No. 47, Division of Biostatistics, Stanford Univ.
6. Federal Register (22 Feb. 1977). 42 (No. 35): 10412-10422.
7. Finney, D. J. (1947). Probit Analysis, London: Cambridge Univ. Press.
8. Guess, H. A. and K. S. Crump (1977). Can We Use Animal Data to Estimate Safe Doses for Chemical Carcinogens?, Environmental Health, SIAM, 13-30.
9. Hartley, H. O. and R. L. Sielken (1977). Estimation of Safe Doses in Carcinogenic Experiments, Biometrics, 33:1-30.
10. Kendall, M. G. and A. Stuart (1952). The Advanced Theory of Statistics, Vol. II, New York: Hafner.
11. Mantel, N. and W. R. Bryan (1961). Safety Testing of Carcinogenic Agents, Journal of the National Cancer Institute, 27(2): 455-470.
12. Mantel, N., et al. (1975). Improved Mantel-Bryan Procedure for Safety Testing of Carcinogens, Cancer Research, 35:865-872.
13. Mood, A. M., F. A. Greybill and D. C. Boes (1974). Introduction to the Theory of Statistics, New York: McGraw-Hill.
14. Pasternack, B. S. and R. E. Shove (1977). Statistical Methods for Assessing Risk Following Exposure to Environmental Carcinogens, Environmental Health, SIAM, 49-71.

15. Petro, R. and P. Lee (1973). Weibull Distributions for Continuous-Carcinogenesis Experiments, Biometrics, 29:457-470.
16. Prentice, R. L. (1975). Discrimination Among Some Parametric Models, Biometrika, 62:607-614.
17. Prentice, R. L. (1976). A Generalization of the Probit and Logit Methods for Dose Response Curves, Biometrics, 32:761-768.
18. Rall, D. P. (1974). Problems of Low Doses of Carcinogens, Journal Washington Academy of Sciences, 64(2):63-68.
19. Rao, C. R. (1973). Linear Statistical Inference and Its Applications, New York: Wiley.
20. Schneiderman, M. A. (1974). Safe Dose? Problem of the Statistician in the World of Trans-Science, Journal Washington Academy of Sciences, 64(2):68-78.

ACKNOWLEDGEMENTS

A brief note of thanks goes to Dr. George Milliken for the sharing of his time and expertise. It also extends to Drs. Hess, Fryer, and Kemp.

The completion of my studies is also due to the financial support of the Department of Statistics made possible by Dr. Arthur Dayton and the unending moral support given by parents, family, and friends.

LOW DOSE RISK ESTIMATION USING TWO CONVENIENT FORMS
OF THE GENERALIZED PROBIT MODEL

by

JEFFREY J. RAHENKAMP

B.S., Clarion State College, 1976

AN ABSTRACT OF A MASTER'S REPORT

submitted in partial fulfillment of the

requirements for the degree

MASTER OF SCIENCE

Department of Statistics

KANSAS STATE UNIVERSITY

1980

Various approaches to the problem of low dose risk assessment involving potential carcinogenic agents are presented and discussed. Two convenient versions of the generalized probit model are more thoroughly investigated. The computer software package SAS is employed to facilitate maximum likelihood estimation of the parameters in the generalized probit model. Empirical comparisons are made between competing models using two trial data sets from the literature. These comparisons are in the areas of goodness of fit, "virtually safe" dose estimates, effective dose estimates, and asymptotic confidence intervals about the effective doses.