

STATISTICAL INFERENCE FROM COMPLETE
AND INCOMPLETE DATA

A THESIS SUBMITTED TO
THE GRADUATE SCHOOL OF NATURAL AND APPLIED SCIENCES
OF
MIDDLE EAST TECHNICAL UNIVERSITY

BY

OYA CAN MUTAN

IN PARTIAL FULFILLMENT OF THE REQUIREMENTS
FOR
THE DEGREE OF DOCTOR OF PHILOSOPHY
IN
STATISTICS

JANUARY 2010

Approval of the thesis:

**STATISTICAL INFERENCE FROM COMPLETE
AND INCOMPLETE DATA**

submitted by OYA CAN MUTAN in partial fulfillment of the requirements for
the degree of Doctor of Philosophy in Statistics Department, Middle East
Technical University by,

Prof. Dr. Canan Özgen _____
Dean, Graduate School of **Natural and Applied Sciences**

Prof. Dr. H. Öztaş Ayhan _____
Head of Department, **Statistics**

Prof. Dr. Moti Lal Tiku _____
Supervisor, **Statistics Dept., METU**

Examining Committee Members:

Prof. Dr. Ayşen Akkaya _____
Statistics Dept., METU

Prof. Dr. Moti Lal Tiku _____
Statistics Dept., METU

Assoc. Prof. Dr. M. Qamarul Islam _____
Economics Dept., Çankaya University

Assistant Prof. Dr. Barış Sürücü _____
Statistics Dept., METU

Assistant Prof. Dr. Hakan Savaş Sazak _____
Statistics Dept., Ege University

Date: 25.01.2010

I hereby declare that all information in this document has been obtained and presented in accordance with academic rules and ethical conduct. I also declare that, as required by these rules and conduct, I have fully cited and referenced all material and results that are not original to this work.

Oya CAN MUTAN

ABSTRACT

STATISTICAL INFERENCE FROM COMPLETE AND INCOMPLETE DATA

Can Mutan, Oya

Ph.D., Department of Statistics

Supervisor: Prof. Dr. Moti Lal Tiku

January 2010, 145 pages

Let X and Y be two random variables such that Y depends on $X=x$. This is a very common situation in many real life applications. The problem is to estimate the location and scale parameters in the marginal distributions of X and Y and the conditional distribution of Y given $X=x$. We are also interested in estimating the regression coefficient and the correlation coefficient. We have a cost constraint for observing $X=x$, the larger x is the more expensive it becomes. The allowable sample size n is governed by a pre-determined total cost. This can lead to a situation where some of the largest $X=x$ observations cannot be observed (Type II censoring). Two general methods of estimation are available, the method of least squares and the method of maximum likelihood. For most non-normal distributions, however, the latter is analytically and computationally problematic. Instead, we use the method of modified maximum likelihood estimation which is known to be essentially as efficient as the maximum likelihood estimation. The method has a distinct advantage: It yields estimators which are explicit functions of sample observations and are, therefore, analytically and computationally

straightforward. In this thesis specifically, the problem is to evaluate the effect of the largest order statistics $x(i)$ ($i > n-r$) in a random sample of size n (i) on the mean $E(X)$ and variance $V(X)$ of X , (ii) on the cost of observing the x -observations, (iii) on the conditional mean $E(Y|X=x)$ and variance $V(Y|X=x)$ and (iv) on the regression coefficient. It is shown that unduly large x -observations have a detrimental effect on the allowable sample size and the estimators, both least squares and modified maximum likelihood. The advantage of not observing a few largest observations are evaluated. The distributions considered are Weibull, Generalized Logistic and the scaled Student's t .

Key Words: Modified maximum likelihood, Type II censoring, Weibull, Generalized Logistic, scaled Student's t .

ÖZ

TAM VE EKSİK VERİLERDEN İSTATİSTİKSEL SONUÇ ÇIKARSAMA

Can Mutan, Oya

Doktora, İstatistik Bölümü

Tez Yöneticisi: Prof. Dr. Moti Lal Tiku

Ocak 2010, 145 sayfa

X ve Y rassal değişkenlerinden Y 'nin X 'e bağlı olması gerçek hayat uygulamalarında oldukça sık karşılaşılan bir durumdur. Bu tezde böyle durumlar için X ve Y 'nin marjinal dağılımlarındaki ve de $X=x$ iken Y 'nin koşullu dağılımındaki konum ve ölçek parametrelerini, ayrıca regresyon ve korelasyon katsayılarını tahmin etmek hedeflenmiştir. $X=x$ iken bir maliyet kısıtı belirlenmiş, büyük x değerlerinin maliyeti artıracığı varsayılmıştır. Kabul edilebilir örneklem sayısı olan n 'in önceden belirlenmiş bir maliyetle gözlenebileceği belirtilmiş, bu durum beraberinde bazı en büyük x değerlerinin gözlenememesi sonucunu getirmiştir (2. Tip sansürleme). Kullanılan iki genel tahmin yöntemi en küçük kareler ve en çok olabilirlik yöntemleridir. Ancak, ikinci yöntem çoğu normal olmayan dağılım için analitik ve sayısal olarak oldukça problemli sonuçlar vermekte, benzer durumlarda esasen en çok olabilirlik tahmin yöntemi kadar iyi sonuçlar veren uyarlanmış en çok olabilirlik tahmin yönteminin kullanılması uygun olmaktadır. Bu yöntemin belirgin avantajı tahmin edicilerin gözlemlerin açık fonksiyonları cinsinden ifade edilebilmesi ve böylece analitik ve sayısal olarak basitçe hesaplanabilir olmalarıdır. Bu tezde, büyüklüğü n olan rassal bir

örneklem için en büyük sıralı istatistiklerin $x(i)$ ($i > n-r$) (i) X 'in ortalaması $E(X)$ ve varyansı $V(X)$ (ii) x gözlemlerini gözlemenin maliyeti (iii) koşullu ortalama $E(Y|X=x)$ ve varyans $V(Y|X=x)$ ve (iv) regresyon katsayısı üzerindeki etkilerini değerlendirmek amaçlanmıştır. Çok büyük x -gözlemlerinin kabul edilebilir örneklem sayısı ve en küçük kareler ve uyarlanmış en çok olabilirlik tahmin edicileri üzerinde fazlasıyla zararlı etkileri olduğu görülmüş, en büyük birkaç değişkenin kullanılmamasının faydaları değerlendirilmiştir. Dikkate alınan dağılımlar Weibull, Genelleştirilmiş Lojistik ve ölçeklendirilmiş Student-t dağılımıdır.

Anahtar kelimeler: Uyarlanmış en çok olabilirlik, 2. Tip sansürleme, Weibull, Genelleştirilmiş Lojistik, ölçeklendirilmiş Student-t.

To my beloved family ...

ACKNOWLEDGEMENTS

First of all I want to say that I am totally indebted to my supervisor Prof. Dr. Moti Lal Tikü. I feel myself very lucky because I had an opportunity to work with such a great professor like him. As well as his unquestionable outstanding knowledge, his view of life and his human side had influenced me so deeply that I am not the same person I was before I knew him. Now, I am striving to be a better person in every respect. I want to thank him with all my heart for the everything he had done.

I would like to thank Assoc. Prof. Dr. M. Qamarul İslam for following my thesis with deep interest and maintaining a positive attitude throughout this study.

I should also thank Prof. Dr. Ayşen Akkaya for always acting in a supportive manner and being so generous in helping.

I am also grateful to Asst. Prof. Dr. Barış Sürücü for his full support, motivative manner and sincereness.

I want to express my deep appreciation to Asst. Prof. Dr. Hakan Savaş Sazak for his help, encouragement and substantial comments about my thesis.

I also want to thank my dear friend Pelin Özbozkurt. She was with me from the very first moment of this study up to the end. Her existence was so invaluable for me that I will never and ever forget our memories.

I am also thankful to my dear friend Ayça Dönmez. For me, one of the most important gains of this Ph.D. was her, because I can say that we got closer as

a result of this study. She always made me feel that she will be always there to celebrate my happiness and share my sadness.

I would also like to thank the Department of Statistics for providing me a comfortable working environment and necessary equipments and Capital Markets Board of Turkey for encouraging me to complete this thesis.

I want to thank my dear friend Ayhan Topcu for her sincerity and support. She bore witness to all steps of this study and gave me reliance, peace and strength with her true friendship.

I also want to thank my beloved mother Reyhan Can and father Kemal Can for their everlasting support and true love that I feel in each and every step I take throughout my life. The truth is that this work cannot be completed without their support.

I am also thankful to my sister Zeynep Pekcan, my brother-in-law Mert Pekcan and my dear niece Defne Pekcan. They always stand by me with their full support, implicit trust and true love.

I also want to thank my mother-in-law Güler Mutan for her full trust, support, and sincereness.

I owe my loving thanks to my dear husband Serkan Mutan for being my husband, intimate and best friend. I always felt his invisible warm arms around me in the every move I made. He was always gentle, thoughtful and warm-hearted. In my every mood, he was the only one that understood me without any judgements. Moreover, he was the one that makes my each and every single day very special.

Finally, I want to thank my precious daughter Asya... Maybe she is now too young to understand what is going on around her. But I can for sure say that everything is for her... Always to see her happy and to leave a good mark behind...

TABLE OF CONTENTS

ABSTRACT	iv
ÖZ.....	vi
DEDICATION	viii
ACKNOWLEDGEMENTS	ix
TABLE OF CONTENTS	xi
LIST OF TABLES	xiv
LIST OF FIGURES.....	xvii
CHAPTER	
1. INTRODUCTION AND LITERATURE SURVEY	1
1.1. The Bivariate Normal Distribution	3
1.1.1. Estimation of Parameters.....	4
1.1.2. Asymptotic Covariance Matrix	8
1.1.3. Hypothesis Testing	9
1.2. Marginal Distribution is Weibull and Conditional Distribution is Normal	10
1.2.1. Estimation of Parameters.....	11
1.2.2. Fisher Information Matrix	13
1.3. Marginal Distribution is Weibull and Conditional Distribution is Normal with Variance Inversely Proportional to X.....	14
1.3.1. Estimation of Parameters.....	15
1.3.2. Fisher Information Matrix	16
1.4. Marginal Distribution is ‘Censored’ Weibull and Conditional Distribution is Normal	17
1.4.1. Estimation of Parameters (Weibull Marginal).....	18
1.4.2. Sample Information Matrix	20

2. EVALUATING THE EFFECT OF THE LARGEST OBSERVATIONS IN THE MARGINAL DISTRIBUTION	22
2.1. Marginal Distribution is Generalized Logistic and Conditional Distribution is Normal	24
2.1.1. Complete Sample.....	24
2.1.1.1. Estimation of Parameters	24
2.1.1.2. Fisher Information Matrix	29
2.1.2. Censored Sample	30
2.1.2.1. Estimation of Parameters	31
2.1.2.2. Sample Information Matrix.....	35
2.2. Marginal Distribution is Weibull and Conditional Distribution is Generalized Logistic.....	37
2.2.1. Complete Sample.....	37
2.2.1.1. Estimation of Parameters	37
2.2.1.2. Fisher Information Matrix	42
2.2.2. Censored Sample	44
2.2.2.1. Estimation of Parameters	44
2.2.2.2. Sample Information Matrix.....	48
2.3. Marginal Distribution is Weibull and Conditional Distribution is Long – Tailed Symmetric	50
2.3.1. Complete Sample.....	50
2.3.1.1. Estimation of Parameters	50
2.3.1.2. Fisher Information Matrix.....	55
2.3.2. Censored Sample	55
2.3.2.1. Estimation of Parameters	56
2.3.2.2. Sample Information Matrix.....	60
2.4. Marginal and Conditional Distribution is Generalized Logistic.....	61
2.4.1. Complete Sample.....	61
2.4.1.1. Estimation of Parameters	62
2.4.1.2. Fisher Information Matrix.....	68
2.4.1.3. Hypothesis Testing.....	69
2.4.1.4. Bias Corrected Least Square Estimators	70

2.4.2. Censored Sample	72
2.4.2.1. Estimation of Parameters	72
2.4.2.2. Sample Information Matrix	78
3. SIMULATION RESULTS AND ILLUSTRATIVE EXAMPLES	81
3.1. Simulation Results when Marginal Distribution is Generalized Logistic and Conditional Distribution is Normal	82
3.2. Simulation Results when Both Marginal and Conditional Distributions are Generalized Logistic.....	91
3.3. Accuracy of the Fisher Information Matrix	99
3.4. Simulated Powers of the Test Statistic Z	102
3.5. Illustrative Examples	104
3.5.1. Real Life Example 1	104
3.5.2. Real Life Example 2	107
3.5.3. Real Life Example 3	109
4. CONCLUSION	113
REFERENCES	116
APPENDIX	121
PROGRAM WHEN BOTH MARGINAL AND CONDITIONAL DISTRIBUTIONS ARE GENERALIZED LOGISTIC	121
VITA	145

LIST OF TABLES

Table 3.1.1 Simulation results for complete sample when $X \sim GL$, $Y X \sim \text{normal}$ - without outliers.....	83
Table 3.1.2 Simulation results for complete sample when $X \sim GL$, $Y X \sim \text{normal}$ - Dixon's outlier model (2.0 is added to the r of the X -observations)	84
Table 3.1.3 Simulation results for complete sample when $X \sim GL$, $Y X \sim \text{normal}$ - Dixon's outlier model (4.0 is added to the r of the X -observations)	85
Table 3.1.4 Simulation results for complete sample when $X \sim GL$, $Y X \sim \text{normal}$ - Tiku's outlier model (2.0 is added to the largest r of the X -observations)	86
Table 3.1.5 Simulation results for complete sample when $X \sim GL$, $Y X \sim \text{normal}$ - Tiku's outlier model (4.0 is added to the largest r of the X -observations)	87
Table 3.1.6 Simulation results for censored sample when $X \sim GL$, $Y X \sim \text{normal}$ - without outliers.....	88
Table 3.1.7 Simulation results for censored sample when $X \sim GL$, $Y X \sim \text{normal}$ - Dixon's outlier model (2.0 is added to the r of the X -observations)	89
Table 3.1.8 Simulation results for censored sample when $X \sim GL$, $Y X \sim \text{normal}$ - Dixon's outlier model (4.0 is added to the r of the X -observations)	89
Table 3.1.9 Simulation results for censored sample when $X \sim GL$, $Y X \sim \text{normal}$ - Tiku's outlier model (2.0 is added to the largest r of the X -observations).....	90
Table 3.1.10 Simulation results for censored sample when $X \sim GL$, $Y X \sim \text{normal}$ - Tiku's outlier model (4.0 is added to the largest r of the X -observations).....	91
Table 3.2.1 Simulation results for complete sample when $X \sim GL$, $Y X \sim GL$ - without outliers.....	92
Table 3.2.2 Simulation results for complete sample when $X \sim GL$, $Y X \sim GL$ - Dixon's outlier model (2.0 is added to the r of the X -observations)	93

Table 3.2.3 Simulation results for complete sample when $X \sim GL$, $Y X \sim GL$ - Dixon's outlier model (4.0 is added to the r of the X -observations)	94
Table 3.2.4 Simulation results for complete sample when $X \sim GL$, $Y X \sim GL$ - Tiku's outlier model (2.0 is added to the largest r of the X -observations)	95
Table 3.2.5 Simulation results for complete sample when $X \sim GL$, $Y X \sim GL$ - Tiku's outlier model (4.0 is added to the largest r of the X -observations)	96
Table 3.2.6 Simulation results for censored sample when $X \sim GL$, $Y X \sim GL$ - without outliers.....	97
Table 3.2.7 Simulation results for censored sample when $X \sim GL$, $Y X \sim GL$ - Dixon's outlier model (2.0 is added to the r of the X -observations)	97
Table 3.2.8 Simulation results for censored sample when $X \sim GL$, $Y X \sim GL$ - Dixon's outlier model (4.0 is added to the r of the X -observations)	98
Table 3.2.9 Simulation results for censored sample when $X \sim GL$, $Y X \sim GL$ - Tiku's outlier model (2.0 is added to the largest r of the X -observations).....	98
Table 3.2.10 Simulation results for censored sample when $X \sim GL$, $Y X \sim GL$ - Tiku's outlier model (4.0 is added to the largest r of the X -observations).....	99
Table 3.3.1 MVB's for complete sample when $X \sim GL$, $Y X \sim normal$	100
Table 3.3.2 MVB's for censored sample when $X \sim GL$, $Y X \sim normal$ - without outliers	100
Table 3.3.3 MVB's for censored sample when $X \sim GL$, $Y X \sim normal$ - Dixon's outlier model (2.0 is added to the r of the X -observations)	100
Table 3.3.4 MVB's for censored sample when $X \sim GL$, $Y X \sim normal$ - Dixon's outlier model (4.0 is added to the r of the X -observations)	100
Table 3.3.5 MVB's for censored sample when $X \sim GL$, $Y X \sim normal$ - Tiku's outlier model (2.0 is added to the largest r of the X -observations)	101
Table 3.3.6 MVB's for censored sample when $X \sim GL$, $Y X \sim normal$ - Tiku's outlier model (4.0 is added to the largest r of the X -observations)	101
Table 3.3.7 MVB's for complete sample when $X \sim GL$, $Y X \sim GL$	101
Table 3.3.8 MVB's for censored sample when $X \sim GL$, $Y X \sim GL$ - without outliers	101

Table 3.3.9 MVB's for censored sample when $X \sim GL$, $Y X \sim GL$ - Dixon's outlier model (2.0 is added to the r of the X -observations)	102
Table 3.3.10 MVB's for censored sample when $X \sim GL$, $Y X \sim GL$ - Dixon's outlier model (4.0 is added to the r of the X -observations)	102
Table 3.3.11 MVB's for censored sample when $X \sim GL$, $Y X \sim GL$ - Tiku's outlier model (2.0 is added to the largest r of the X -observations).....	102
Table 3.3.12 MVB's for censored sample when $X \sim GL$, $Y X \sim GL$ - Tiku's outlier model (4.0 is added to the largest r of the X -observations).....	102
Table 3.5.1.1. Woolson data.....	104
Table 3.5.1.2. MML and LS estimators for Woolson data.....	106
Table 3.5.1.3. Parametric bootstrap variances of the estimators for Woolson data	106
Table 3.5.2.1. Gross and Clark data	107
Table 3.5.2.2. MML and LS estimators for Gross and Clark data	108
Table 3.5.2.3. Parametric bootstrap variances of the estimators for Gross and Clark data.....	109
Table 3.5.3.1. Firm data	109
Table 3.5.3.2. MML and LS estimators for firm data	111
Table 3.5.3.3. Parametric bootstrap variances of the estimators for the firm data	112

LIST OF FIGURES

Figure 3.4.1. Power curves for complete sample, without outliers and with 10% outliers in X , for $b_1=b_2=0.5$, $n=50$	103
Figure 3.4.2. Power curves for complete samples, without outliers and with 10% outliers in X , for $b_1=b_2=4.0$, $n=50$	103
Figure 3.4.3. Power curves for complete sample, without outliers and with 10 % outliers in X , for $b_1=b_2=8.0$, $n=50$	103
Figure 3.5.1.1. Q-Q plot of X when $b=0.5$	105
Figure 3.5.1.2. Q-Q plot of the standardized residuals for Woolson data.....	105
Figure 3.5.2.1. Q-Q plot of U when $p=0.8$	107
Figure 3.5.2.2. Q-Q plot of the standardized residuals for Gross and Clark data.....	108
Figure 3.5.3.1. Q-Q plot of X when $b=1.0$	110
Figure 3.5.3.2. Q-Q plot of the standardized residuals for the firm data.....	111

CHAPTER 1

INTRODUCTION AND LITERATURE SURVEY

Let X and Y be random variables and Y depends on $X=x$, which is a very common situation in some real life applications. For example,

- (i) X is the white blood count while Y is the survival time of a patient who has leukemia,
- (ii) X is the age while Y is the blood pressure of a patient who has coronary heart disease,
- (iii) X is the areal precipitation while Y is the annual streamflow in a particular locality,
- (iv) X is the storm surge level while Y is the wave height of water in a sea,
- (v) X is the dose of rat-poison given to a rodent while Y is the time taken by the rodent to die,
- (vi) X is the amount of fuel in a nuclear reactor while Y is the energy it produces.

However, although non-normal situations occur frequently in practice, in the literature, most studies are based on normality assumption. When the marginal or conditional distributions are non-normal, application of the very popular maximum likelihood (ML) estimation technique becomes problematic. These

problems are (a) likelihood equations have no explicit solutions and (b) solving them iteratively can lead to unreliable results because of multiple roots, nonconvergence of iterations, or convergence to wrong values; see Barnett (1966), Lee et al. (1980), Puthenpura and Sinha (1986), Vaughan (1992), Qumsiyeh (2007). All these situations bring out the need for statistical methods which give efficient results under non-normal distributions. For these cases, using Tiku's modified maximum likelihood (MML) estimation technique, which linearizes the intractable terms in the likelihood equations, is very useful (Tiku, 1967; Tiku, 1968; Tiku, 1980; Tiku and Suresh, 1992; Vaughan and Tiku, 2000).

Islam et al. (2001), Tiku et al. (2001) and Sazak (2003) perform MML procedure in simple linear regression analysis. In their studies, Islam et al. (2001) and Tiku et al. (2001) consider the case where the explanatory variable is nonstochastic and error terms are coming from nonnormal distributions such as Weibull, Generalized Logistic, Student's t , and short-tailed distributions. Also, Akkaya and Tiku (2008) deal with nonnormal error terms for multiple linear regression with nonstochastic explanatory variables. However, Sazak (2003) states that in real life problems, along with the nonnormal error terms, explanatory variables are usually stochastic and he develops MML methodology for these cases. The distributions that he takes into consideration are Weibull and Generalized Logistic; see Sazak et al. (2006).

Note that the MML estimators are easy to compute and are explicit functions of sample observations and in addition to these, they have the following desired properties;

- (i) Asymptotically, the MML estimators are unbiased and their variances are equal to minimum variance bounds (MVB). In other words, asymptotically, they are fully efficient.
- (ii) MML estimators also perform well for small samples. They have no or negligible bias. Besides, their variances are very

close to MVB and are essentially as efficient as ML estimators.
See Şenoğlu and Tiku (2001, 2002), Vaughan (2002).

The aim of this thesis is to evaluate the effect of the largest order statistics $x_{(i)}$ ($i \geq n-r$) in a random sample of size n

- (i) on the mean $E(X)$ and variance $V(X)$ of X ,
- (ii) on the cost of observing the x -observations, and
- (iii) on the conditional mean $E(Y|X=x)$ and variance $V(Y|X=x)$,
- (iv) on the regression coefficient

for the distributions coming from p -family, Weibull and Generalized Logistic. It is shown that unduly large x -observations have detrimental effects on (i)-(iv). The advantages of not observing a few largest observations on estimation and hypothesis testing are also evaluated.

1.1. The Bivariate Normal Distribution

Suppose that the two random variables (X,Y) are coming from a bivariate normal distribution with parameters $\mu_1, \mu_2, \sigma_1, \sigma_2$ and ρ , where ρ is the correlation coefficient between X and Y . Then, the probability density function is

$$g(x, y) = \frac{1}{2\pi\sigma_1\sigma_2\sqrt{1-\rho^2}} \times \exp\left(-\frac{1}{2(1-\rho^2)}\left(\left(\frac{x-\mu_1}{\sigma_1}\right)^2 + \left(\frac{y-\mu_2}{\sigma_2}\right)^2 - 2\rho\left(\frac{x-\mu_1}{\sigma_1}\right)\left(\frac{y-\mu_2}{\sigma_2}\right)\right)\right) \quad (1.1.1)$$

where $-\infty < x, y < \infty$, $\mu_1, \mu_2 \in \mathbb{R}$, $\sigma_1, \sigma_2 > 0$, $-1 < \rho < 1$.

1.1.1.1. Estimation of Parameters

Given a random sample (x_i, y_i) , for $1 \leq i \leq n$, the likelihood function is

$$L \propto \sigma_1^{-n} \sigma_2^{-n} (1 - \rho^2)^{-n/2} \exp \left(-\frac{1}{2(1 - \rho^2)} \left(\sum_{i=1}^n x_{i1}^2 + \sum_{i=1}^n y_{i1}^2 - 2\rho \sum_{i=1}^n x_{i1} y_{i1} \right) \right) \quad (1.1.1.1)$$

where $x_{i1} = \frac{x_i - \mu_1}{\sigma_1}$ and $y_{i1} = \frac{y_i - \mu_2}{\sigma_2}$

Then, the likelihood equations for estimating the parameters are

$$\frac{\partial \ln L}{\partial \mu_1} = \frac{1}{\sigma_1(1 - \rho^2)} \sum_{i=1}^n x_{i1} - \frac{\rho}{\sigma_1(1 - \rho^2)} \sum_{i=1}^n y_{i1} = 0 \quad (1.1.1.2)$$

$$\frac{\partial \ln L}{\partial \mu_2} = \frac{1}{\sigma_2(1 - \rho^2)} \sum_{i=1}^n y_{i1} - \frac{\rho}{\sigma_2(1 - \rho^2)} \sum_{i=1}^n x_{i1} = 0 \quad (1.1.1.3)$$

$$\frac{\partial \ln L}{\partial \sigma_1} = -\frac{n}{\sigma_1} + \frac{1}{\sigma_1(1 - \rho^2)} \sum_{i=1}^n x_{i1}^2 - \frac{\rho}{\sigma_1(1 - \rho^2)} \sum_{i=1}^n x_{i1} y_{i1} = 0 \quad (1.1.1.4)$$

$$\frac{\partial \ln L}{\partial \sigma_2} = -\frac{n}{\sigma_2} + \frac{1}{\sigma_2(1 - \rho^2)} \sum_{i=1}^n y_{i1}^2 - \frac{\rho}{\sigma_2(1 - \rho^2)} \sum_{i=1}^n x_{i1} y_{i1} = 0 \quad (1.1.1.5)$$

and

$$\begin{aligned} \frac{\partial \ln L}{\partial \rho} &= \frac{n\rho}{1 - \rho^2} + \frac{1}{1 - \rho^2} \sum_{i=1}^n x_{i1} y_{i1} - \\ &\quad \frac{\rho}{(1 - \rho^2)^2} \left\{ \sum_{i=1}^n x_{i1}^2 - 2\rho \sum_{i=1}^n x_{i1} y_{i1} + \sum_{i=1}^n y_{i1}^2 \right\} = 0. \end{aligned} \quad (1.1.1.6)$$

The solutions of the equations (1.1.1.2)-(1.1.1.6) are the ML estimators of μ_1 , μ_2 , σ_1 , σ_2 , ρ and they are

$$\begin{aligned}
\hat{\mu}_1 &= \bar{x} = \frac{1}{n} \sum_{i=1}^n x_i, \quad \hat{\mu}_2 = \bar{y} = \frac{1}{n} \sum_{i=1}^n y_i, \\
\hat{\sigma}_1 &= \sqrt{\frac{\sum_{i=1}^n (x_i - \bar{x})^2}{n}}, \quad \hat{\sigma}_2 = \sqrt{\frac{\sum_{i=1}^n (y_i - \bar{y})^2}{n}}, \\
\hat{\rho} &= \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2 \sum_{i=1}^n (y_i - \bar{y})^2}}.
\end{aligned} \tag{1.1.1.7}$$

Note that in the absence of the statistical independence, the joint distribution of two random variables X and Y can be written as the product of a conditional and a marginal distribution. Therefore, $g(x,y)$ is equal to,

$$g(x,y) = f(x).h(y|x). \tag{1.1.1.8}$$

where X and $Y|X$ are statistically independent. In equation (1.1.1.8), if (X,Y) are coming from a bivariate normal distribution, then $f(x)$ and $h(y|x)$ can be written as follows;

$$f(x) = \frac{1}{\sigma_1 \sqrt{2\pi}} \exp\left(-\frac{(x - \mu_1)^2}{2\sigma_1^2}\right) \tag{1.1.1.9}$$

and

$$h(y|x) = \frac{1}{\sigma_{2.1} \sqrt{2\pi}} \exp\left(-\frac{(y - \mu_{y|x})^2}{2\sigma_{2.1}^2}\right) \tag{1.1.1.10}$$

where

$$\begin{aligned}
\sigma_{2.1} &= \sqrt{\sigma_2^2(1-\rho^2)}, \\
\mu_{y|x} &= \mu_{2.1} + \theta x, \\
\mu_{2.1} &= \mu_2 - \theta\mu_1, \\
\theta &= \rho \left(\frac{\sigma_2}{\sigma_1} \right).
\end{aligned} \tag{1.1.1.11}$$

Remark: From (1.1.1.11), it is understood that, the variance of $Y|X$ is less than the variance of Y , if $\rho \neq 0$. Because, $\sigma_{2.1}^2 = \sigma_2^2(1-\rho^2)$ and $(1-\rho^2) \leq 1$. Also, from the formula it is obvious that as ρ increases, the variance of $Y|X$ decreases.

Given a random sample (x_i, y_i) , for $1 \leq i \leq n$, the product of marginal and conditional likelihood functions are

$$L \propto \sigma_1^{-n} \exp \left\{ -\frac{\sum_{i=1}^n (x_i - \mu_1)^2}{2\sigma_1^2} \right\} \sigma_{2.1}^{-n} \exp \left\{ -\frac{\sum_{i=1}^n (y_i - \mu_{2.1} - \theta x_i)^2}{2\sigma_{2.1}^2} \right\}. \tag{1.1.1.12}$$

The log likelihood function is

$$\ln L = -n \ln \sigma_1 - \frac{1}{2} \sum_{i=1}^n z_{i1}^2 - n \ln \sigma_{2.1} - \frac{1}{2} \sum_{i=1}^n z_{i2}^2 \tag{1.1.1.13}$$

where $z_{i1} = \frac{x_i - \mu_1}{\sigma_1}$ and $z_{i2} = \frac{y_i - \mu_{2.1} - \theta x_i}{\sigma_{2.1}}$, $-\infty < z_{i1}, z_{i2} < \infty$.

The likelihood equations for μ_1 , σ_1 , $\mu_{2.1}$, $\sigma_{2.1}$ and θ are

$$\frac{\partial \ln L}{\partial \mu_1} = \frac{1}{\sigma_1} \sum_{i=1}^n z_{i1} = 0 \tag{1.1.1.14}$$

$$\frac{\partial \ln L}{\partial \sigma_1} = -\frac{n}{\sigma_1} + \frac{1}{\sigma_1} \sum_{i=1}^n z_{i1}^2 = 0 \quad (1.1.1.15)$$

$$\frac{\partial \ln L}{\partial \mu_{2,1}} = \frac{1}{\sigma_{2,1}} \sum_{i=1}^n z_{i2} = 0 \quad (1.1.1.16)$$

$$\frac{\partial \ln L}{\partial \sigma_{2,1}} = -\frac{n}{\sigma_{2,1}} + \frac{1}{\sigma_{2,1}} \sum_{i=1}^n z_{i2}^2 = 0 \quad (1.1.1.17)$$

and

$$\frac{\partial \ln L}{\partial \theta} = \frac{1}{\sigma_{2,1}} \sum_{i=1}^n x_i z_{i2} = 0. \quad (1.1.1.18)$$

The estimators, which are the solutions of the likelihood equations are given below:

$$\begin{aligned} \hat{\mu}_1 &= \bar{x}, \quad \hat{\sigma}_1 = \sqrt{\frac{\sum_{i=1}^n (x_i - \bar{x})^2}{n}}, \\ \hat{\mu}_{2,1} &= \bar{y} - \hat{\theta} \bar{x}, \\ \hat{\sigma}_{2,1} &= \sqrt{\frac{\sum_{i=1}^n \{y_i - \bar{y} - \hat{\theta}(x_i - \bar{x})\}^2}{n}} \end{aligned}$$

and

$$\hat{\theta} = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\left(\sum_{i=1}^n (x_i - \bar{x})^2 \right)}. \quad (1.1.1.19)$$

If we compare the results (1.1.1.7) and (1.1.1.19), we see the following:

$$\hat{\mu}_{2.1} = \hat{\mu}_2 - \hat{\theta}\hat{\mu}_1, \quad \hat{\sigma}_{2.1} = \sqrt{\hat{\sigma}_2^2(1 - \hat{\rho}^2)} \quad \text{and} \quad \hat{\theta} = \hat{\rho} \left(\frac{\hat{\sigma}_2}{\hat{\sigma}_1} \right),$$

which are the results of the *invariance property* of ML estimators. Invariance property says that if $\hat{\delta}$ is the ML estimator of some parameter δ and $g(\cdot)$ is a one-to-one function, then $g(\hat{\delta})$ is the ML estimator of $g(\delta)$.

Remark: Note that the invariance property does not necessarily lead to unbiasedness, unless the sample size is large.

1.1.2. Asymptotic Covariance Matrix

The asymptotic covariance matrix is obtained by taking the inverse of Fisher information matrix, which is equal to the negative of the expectation of the second derivative of the log likelihood function with respect to the parameters. Therefore, the Fisher information matrix $I(\mu_1, \sigma_1, \mu_2, \sigma_2, \rho)$ is

$$I = [I_{ij}] = \left[-E \left(\frac{\partial^2 \ln L}{\partial \delta_i \partial \delta_j} \right) \right] \quad i, j = 1, 2, 3, 4, 5, \quad (1.1.2.1)$$

$$\delta_1 = \mu_1, \quad \delta_2 = \sigma_1, \quad \delta_3 = \mu_2, \quad \delta_4 = \sigma_2, \quad \delta_5 = \rho,$$

and the elements of $I^{-1}(\mu_1, \sigma_1, \mu_2, \sigma_2, \rho)$ are given by the following equations,

$$I_{11}^{-1} = \frac{n}{\sigma_1^2(1 - \rho^2)}, \quad I_{12}^{-1} = 0, \quad I_{13}^{-1} = -\frac{n\rho}{\sigma_1\sigma_2(1 - \rho^2)}, \quad I_{14}^{-1} = I_{15}^{-1} = 0,$$

$$I_{22}^{-1} = \frac{n(2 - \rho^2)}{\sigma_1^2(1 - \rho^2)}, \quad I_{23}^{-1} = 0, \quad I_{24}^{-1} = -\frac{n\rho^2}{\sigma_1\sigma_2(1 - \rho^2)}, \quad I_{25}^{-1} = -\frac{n\rho}{\sigma_1},$$

$$I_{33}^{-1} = \frac{n}{\sigma_2^2(1 - \rho^2)}, \quad I_{34}^{-1} = I_{35}^{-1} = 0, \quad I_{44}^{-1} = \frac{n(2 - \rho^2)}{\sigma_2^2(1 - \rho^2)}, \quad I_{45}^{-1} = -\frac{n\rho}{\sigma_2}$$

and

$$I_{55}^{-1} = \frac{n(1 + \rho^2)}{(1 - \rho^2)^2}. \quad (1.1.2.2)$$

If case (1.1.1.9) is taken into consideration, then the Fisher information matrix $I(\mu_1, \sigma_1, \mu_{2.1}, \sigma_{2.1}, \theta)$ is equal to

$$I = [I_{ij}] = \left[-E \left(\frac{\partial^2 \ln L}{\partial \delta_i \partial \delta_j} \right) \right] \quad i, j = 1, 2, 3, 4, 5, \quad (1.1.2.3)$$

$$\delta_1 = \mu_1, \delta_2 = \sigma_1, \delta_3 = \mu_{2.1}, \delta_4 = \sigma_{2.1}, \delta_5 = \theta,$$

and the elements of $I^{-1}(\mu_1, \sigma_1, \mu_{2.1}, \sigma_{2.1}, \theta)$ are

$$I_{11}^{-1} = \frac{n}{\sigma_1^2}, \quad I_{12}^{-1} = I_{13}^{-1} = I_{14}^{-1} = I_{15}^{-1} = 0, \quad I_{22}^{-1} = \frac{2n}{\sigma_1^2}, \quad I_{23}^{-1} = I_{24}^{-1} = I_{25}^{-1} = 0,$$

$$I_{33}^{-1} = \frac{n}{\sigma_{2.1}^2}, \quad I_{34}^{-1} = 0, \quad I_{35}^{-1} = \frac{n\mu_1}{\sigma_{2.1}^2}, \quad I_{44}^{-1} = \frac{2n}{\sigma_{2.1}^2}, \quad I_{45}^{-1} = 0$$

and

$$I_{55}^{-1} = \frac{n(\sigma_1^2 + \mu_1^2)}{\sigma_{2.1}^2}. \quad (1.1.2.4)$$

1.1.3. Hypothesis Testing

Suppose $H_0 : \mu_1 = \mu_2 = 0$ or $H_0 : \mu_1 = \mu_{2.1} = 0$ is to be tested. In order to test these multivariate hypotheses, Hotelling's T^2 statistic, which is a generalization of Student's t statistic is used (Hotelling, 1931):

$$T^2 = n(\bar{X} - \mu_0)S^{-1}(\bar{X} - \mu_0) \quad (1.1.3.1)$$

where

$$\bar{X} = \begin{bmatrix} \bar{X}_1 \\ \bar{X}_2 \\ \cdot \\ \cdot \\ \cdot \\ \bar{X}_p \end{bmatrix}, \mu_0 = \begin{bmatrix} \mu_1^0 \\ \mu_2^0 \\ \cdot \\ \cdot \\ \cdot \\ \mu_p^0 \end{bmatrix}.$$

In equation (1.1.3.1), S and n represent the sample variance-covariance matrix and the sample size, respectively.

Under H_0 and p -variate normality, it is well known that

$$T^2 \sim \frac{p(n-1)}{n-p} F_{(p, n-p)}. \quad (1.1.3.2)$$

Remark: The Hotelling T^2 statistics defined by using $\hat{\mu}_1 - \hat{\mu}_2$, or $\hat{\mu}_1 - \hat{\mu}_{2,1}$ are equal to one another.

1.2. Marginal Distribution is Weibull and Conditional Distribution is Normal

In this section it is assumed that X is coming from Weibull distribution which is widely used in applications (Akkaya and Tiku, 2001; Hand et al., 1994; Cohen and Whitten, 1988; Johnson and Johnson, 1979)

$$f(x) = \frac{p}{\sigma_1^p} x^{p-1} \exp \left\{ - \left(\frac{x}{\sigma_1} \right)^p \right\}, \quad 0 < x < \infty, \sigma_1 > 0, p > 0, \quad (1.2.1)$$

and $Y|X=x$ is assumed to be normal,

$$\frac{(y - \mu_{2.1} - \theta x)}{\sigma_{2.1}} \sim N(0,1).$$

Thus, the density function of $Y|X=x$ is

$$h(y | x) = \frac{1}{\sigma_{2.1} \sqrt{2\pi}} \exp - \left(\frac{1}{2\sigma_{2.1}^2} \{y - \mu_{2.1} - \theta x\}^2 \right), \quad -\infty < y < \infty, \quad (1.2.2)$$

where

$$\begin{aligned} \sigma_{2.1} &= \sqrt{\sigma_2^2 (1 - \rho^2)}, \\ \mu_{2.1} &= \mu_2 - \theta E(X), \end{aligned} \quad (1.2.3)$$

and

$$\theta = \rho \left(\frac{\sigma_2}{\sigma_1} \right).$$

1.2.1. Estimation of Parameters

When X is Weibull and $Y|X=x$ is normal, for a random sample (x_i, y_i) , $1 \leq i \leq n$, the likelihood function is

$$L \propto \frac{p^n}{\sigma_1^{np}} \prod_{i=1}^n x_i^{p-1} \exp \left\{ - \sum_{i=1}^n \left(\frac{x_i}{\sigma_1} \right)^p \right\} \frac{1}{\sigma_{2.1}^n} \exp \left\{ - \frac{1}{2\sigma_{2.1}^2} \sum_{i=1}^n (y_i - \mu_{2.1} - \theta x_i)^2 \right\}. \quad (1.2.1.1)$$

Then,

$$\ln L \propto n \ln p - np \ln \sigma_1 + (p-1) \sum_{i=1}^n \ln x_i - \sum_{i=1}^n z_{i1}^p - n \ln \sigma_{2.1} - \frac{1}{2} \sum_{i=1}^n z_{i2}^2 \quad (1.2.1.2)$$

where

$$z_{i1} = \frac{x_i}{\sigma_1} \text{ and } z_{i2} = \frac{(y_i - \mu_{2.1} - \theta x_i)}{\sigma_{2.1}}.$$

The solutions of the following ML equations (1.2.1.3)-(1.2.1.6) are the ML estimators,

$$\frac{\partial \ln L}{\partial \sigma_1} = -\frac{np}{\sigma_1} + \frac{p}{\sigma_1} \sum_{i=1}^n z_{i1}^p = 0, \quad (1.2.1.3)$$

$$\frac{\partial \ln L}{\partial \mu_{2.1}} = \frac{1}{\sigma_{2.1}} \sum_{i=1}^n z_{i2} = 0, \quad (1.2.1.4)$$

$$\frac{\partial \ln L}{\partial \sigma_{2.1}} = -\frac{n}{\sigma_{2.1}} + \frac{1}{\sigma_{2.1}} \sum_{i=1}^n z_{i2}^2 = 0 \quad (1.2.1.5)$$

and

$$\frac{\partial \ln L}{\partial \theta} = \frac{1}{\sigma_{2.1}} \sum_{i=1}^n x_i z_{i2} = 0. \quad (1.2.1.6)$$

They are

$$\hat{\sigma}_1 = \left(\frac{1}{n} \sum_{i=1}^n x_i^p \right)^{1/p}, \quad \hat{\mu}_{2.1} = \bar{y} - \hat{\theta} \bar{x},$$

$$\hat{\sigma}_{2.1} = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \hat{\mu}_{2.1} - \hat{\theta}x_i)^2} \quad \text{and} \quad \hat{\theta} = \frac{\sum_{i=1}^n (x_i - \bar{x})y_i}{\sum_{i=1}^n (x_i - \bar{x})^2}. \quad (1.2.1.7)$$

Also, from the invariance property of the ML estimators, $\hat{\mu}_2$, $\hat{\sigma}_2$ and $\hat{\rho}$ can be easily obtained and they are given by

$$\hat{\mu}_2 = \hat{\mu}_{2.1} + \hat{\theta} \left(1 + \frac{1}{p}\right) \hat{\sigma}_1, \quad \hat{\sigma}_2 = \sqrt{\hat{\sigma}_{2.1}^2 + \hat{\theta}^2 \hat{\sigma}_1^2} \quad \text{and} \quad \hat{\rho} = \hat{\theta} \frac{\hat{\sigma}_1}{\hat{\sigma}_2}. \quad (1.2.1.8)$$

In Tiku and Akkaya (2004), it is mentioned that in the applications of the Weibull distribution in order to obtain highly efficient estimators, large samples are needed. See also Menon (1963), Harter (1964), Smith (1985).

1.2.2. Fisher Information Matrix

The Fisher information matrix $I(\sigma_1, \mu_{2.1}, \sigma_{2.1}, \theta)$ is

$$I_{(\sigma_1, \mu_{2.1}, \sigma_{2.1}, \theta)} = \begin{bmatrix} np^2/\sigma_1^2 & 0 & 0 & 0 \\ 0 & n/\sigma_{2.1}^2 & 0 & n\sigma_1\Gamma(1+1/p)/\sigma_{2.1}^2 \\ 0 & 0 & 2n/\sigma_{2.1}^2 & 0 \\ 0 & n\sigma_1\Gamma(1+1/p)/\sigma_{2.1}^2 & 0 & n\sigma_1^2\Gamma(1+2/p)/\sigma_{2.1}^2 \end{bmatrix}. \quad (1.2.2.1)$$

Also, the components of $I(\sigma_1, \mu_2, \sigma_2, \rho)$ are given below in equations (1.2.2.2)-(1.2.2.10):

$$I_{11} = \frac{np^2}{\sigma_1^2} + \frac{n\rho^2}{\sigma_1^2(1-\rho^2)} \Gamma\left(1 + \frac{2}{p}\right), \quad (1.2.2.2)$$

$$I_{12} = -\frac{n\rho}{\sigma_1\sigma_2(1-\rho^2)} \Gamma\left(1 + \frac{1}{p}\right), \quad (1.2.2.3)$$

$$I_{13} = \frac{n\rho^2}{\sigma_1\sigma_2(1-\rho^2)}\Gamma^2\left(1+\frac{1}{p}\right) - \frac{n\rho^2}{\sigma_1\sigma_2(1-\rho^2)}\Gamma\left(1+\frac{2}{p}\right), \quad (1.2.2.4)$$

$$I_{14} = \frac{n\rho}{\sigma_1(1-\rho^2)}\Gamma^2\left(1+\frac{1}{p}\right) - \frac{n\rho}{\sigma_1(1-\rho^2)}\Gamma\left(1+\frac{2}{p}\right), \quad (1.2.2.5)$$

$$I_{22} = \frac{n}{\sigma_2^2(1-\rho^2)}, \quad (1.2.2.6)$$

$$I_{23} = I_{24} = 0.0, \quad (1.2.2.7)$$

$$I_{33} = \frac{2n}{\sigma_2^2} - \frac{n\rho^2}{\sigma_2^2(1-\rho^2)}\Gamma^2\left(1+\frac{1}{p}\right) + \frac{n\rho^2}{\sigma_2^2(1-\rho^2)}\Gamma\left(1+\frac{2}{p}\right), \quad (1.2.2.8)$$

$$I_{34} = -\frac{n\rho}{\sigma_2(1-\rho^2)}\Gamma^2\left(1+\frac{1}{p}\right) + \frac{n\rho}{\sigma_2(1-\rho^2)}\Gamma\left(1+\frac{2}{p}\right) - \frac{2n\rho}{\sigma_2(1-\rho^2)} \quad (1.2.2.9)$$

and

$$I_{44} = \frac{2n\rho^2}{(1-\rho^2)^2} - \left(\frac{n}{1-\rho^2}\right)\Gamma^2\left(1+\frac{1}{p}\right) + \left(\frac{n}{1-\rho^2}\right)\Gamma\left(1+\frac{2}{p}\right). \quad (1.2.2.10)$$

1.3. Marginal Distribution is Weibull and Conditional Distribution is Normal with Variance Inversely Proportional to X

The case presented in this section is similar to the situation given in section (1.2). X is coming from Weibull distribution given in equation (1.2.1) and $Y|X=x$ is again assumed to be normal. However, in this case the variance of the conditional distribution is inversely proportional to X , which is a very common situation in some statistical applications. In other words,

$$\frac{\sqrt{x}(y - \mu_{2.1} - \theta x)}{\sigma_{2.1}} \sim N(0,1).$$

Therefore, the density function of $Y|X=x$ is

$$h(y | x) = \frac{\sqrt{x}}{\sigma_{2.1} \sqrt{2\pi}} \exp - \left(\frac{x}{2\sigma_{2.1}^2} \{y - \mu_{2.1} - \theta x\}^2 \right), -\infty < y < \infty. \quad (1.3.1)$$

1.3.1. Estimation of Parameters

Given a random sample (x_i, y_i) $1 \leq i \leq n$, the likelihood function is

$$L \propto \frac{p^n}{\sigma_1^{np}} \prod_{i=1}^n x_i^{p-1} \exp \left\{ - \sum_{i=1}^n \left(\frac{x_i}{\sigma_1} \right)^p \right\} \times$$

$$\frac{1}{\sigma_{2.1}^n} \prod_{i=1}^n \sqrt{x_i} \exp \left\{ - \frac{1}{2\sigma_{2.1}^2} \sum_{i=1}^n x_i (y_i - \mu_{2.1} - \theta x_i)^2 \right\}. \quad (1.3.1.1)$$

Then,

$$\ln L \propto n \ln p - np \ln \sigma_1 + \left(p - 1 + \frac{1}{2} \right) \sum_{i=1}^n \ln x_i - \sum_{i=1}^n z_{i1}^p - n \ln \sigma_{2.1} - \frac{1}{2} \sum_{i=1}^n z_{i2}^2 \quad (1.3.1.2)$$

where

$$z_{i1} = \frac{x_i}{\sigma_1} \text{ and } z_{i2} = \frac{\sqrt{x_i} (y_i - \mu_{2.1} - \theta x_i)}{\sigma_{2.1}}.$$

The ML estimators can be obtained by solving the following equations:

$$\frac{\partial \ln L}{\partial \sigma_1} = -\frac{np}{\sigma_1} + \frac{p}{\sigma_1} \sum_{i=1}^n z_{i1}^p = 0, \quad (1.3.1.3)$$

$$\frac{\partial \ln L}{\partial \mu_{2.1}} = \frac{1}{\sigma_{2.1}} \sum_{i=1}^n \sqrt{x_i} z_{i2} = 0, \quad (1.3.1.4)$$

$$\frac{\partial \ln L}{\partial \sigma_{2,1}} = -\frac{n}{\sigma_{2,1}} + \frac{1}{\sigma_{2,1}} \sum_{i=1}^n z_{i2}^2 = 0 \quad (1.3.1.5)$$

and

$$\frac{\partial \ln L}{\partial \theta} = \frac{1}{\sigma_{2,1}} \sum_{i=1}^n x_i^{3/2} z_{i2} = 0. \quad (1.3.1.6)$$

The ML estimators are

$$\begin{aligned} \hat{\sigma}_1 &= \left(\frac{1}{n} \sum_{i=1}^n x_i^p \right)^{1/p}, \quad \hat{\mu}_{2,1} = \frac{\sum_{i=1}^n x_i y_i - \hat{\theta} \sum_{i=1}^n x_i^2}{\sum_{i=1}^n x_i}, \\ \hat{\sigma}_{2,1} &= \sqrt{\frac{1}{n} \sum_{i=1}^n x_i (y_i - \hat{\mu}_{2,1} - \hat{\theta} x_i)^2} \quad \text{and} \quad \hat{\theta} = \frac{\sum_{i=1}^n x_i \sum_{i=1}^n x_i^2 y_i - \sum_{i=1}^n x_i^2 \sum_{i=1}^n x_i y_i}{\sum_{i=1}^n x_i \sum_{i=1}^n x_i^3 - \left(\sum_{i=1}^n x_i^2 \right)^2}. \end{aligned} \quad (1.3.1.7)$$

1.3.2. Fisher Information Matrix

The components of the $I(\sigma_1, \mu_{2,1}, \sigma_{2,1}, \theta)$ are given by the following equations:

$$I_{11} = \frac{np^2}{\sigma_1^2}, \quad (1.3.2.1)$$

$$I_{12} = I_{13} = I_{14} = 0.0, \quad (1.3.2.2)$$

$$I_{22} = \frac{n\sigma_1}{\sigma_{2,1}^2} \Gamma\left(1 + \frac{1}{p}\right), \quad (1.3.2.3)$$

$$I_{23} = 0.0, \quad (1.3.2.4)$$

$$I_{24} = \frac{n\sigma_1^2}{\sigma_{2.1}^2} \Gamma\left(1 + \frac{2}{p}\right) \quad (1.3.2.5)$$

$$I_{33} = \frac{2n}{\sigma_{2.1}^2}, \quad (1.3.2.6)$$

$$I_{34} = 0.0 \quad (1.3.2.7)$$

and

$$I_{44} = \frac{n\sigma_1^3}{\sigma_{2.1}^2} \Gamma\left(1 + \frac{3}{p}\right). \quad (1.3.2.8)$$

1.4. Marginal Distribution is ‘Censored’ Weibull and Conditional Distribution is Normal

For many real life data it is not always possible to observe the value of each sampling unit and this type of data is called censored data. It is a known fact that working with censored data can be more complicated compared to working with complete data. However, it is necessary for practitioners to learn censoring procedures because this kind of samples can be encountered in many areas (see Tiku, 1981; Lawless, 1982; Sarhan and Greenberg, 1962; Cohen, 1957; Gupta, 1952). For example, during the evaluation of the test scores of the students or while dealing with a life-test experiment. If the observations whose values are smaller or larger than a pre-determined level are censored, this type of censoring is called Type I censoring. On the other hand, if a pre-determined number of smallest or largest observations are censored, or in other words, if instead of censoring unit, censoring limits are random, then such a sample is called Type II censoring (Tiku and Akkaya, 2004; Cohen, 1991; Schneider, 1986). Censoring can be done from left and/or right. Şenoğlu and Tiku (2004) consider parameter estimation and hypothesis testing in experimental design for censored and truncated nonnormal samples by using MML estimation procedure.

In this thesis, the most common case of censoring, Type II censoring from right, is considered. That is to say, from a random sample of size n , r largest observations ($r > 0$) are censored. That is usually the situation in life-testing experimentation.

In this section, the situation where marginal distribution is ‘censored’ Weibull and conditional distribution is normal is considered. Note that

$$0 \leq x_{(1)} \leq x_{(2)} \leq \dots \leq x_{(n-r)} < \infty,$$

$$\frac{(y - \mu_{2.1} - \theta x)}{\sigma_{2.1}} \sim N(0,1).$$

1.4.1. Estimation of Parameters (Weibull Marginal)

For an ordered sample $(x_{(i)}, y_{[i]})$ $1 \leq i \leq n-r$ (ordered with respect to x_i observations), the likelihood function is given by

$$L \propto \left[\prod_{i=1}^{n-r} f(x_{(i)}) \right] \left[1 - F(x_{(n-r)}) \right]^r \prod_{i=1}^{n-r} h(y_{[i]} | x_{(i)}), \quad (1.4.1.1)$$

$$\left[\prod_{i=1}^{n-r} \frac{p}{\sigma_1^p} x_{(i)}^{p-1} e^{-\left(\frac{x_{(i)}}{\sigma_1}\right)^p} \right] \left[e^{-\left(\frac{x_{(n-r)}}{\sigma_1}\right)^p} \right]^r \times$$

$$\frac{1}{\sigma_{2.1}^{n-r}} e^{-\frac{1}{2\sigma_{2.1}^2} \sum_{i=1}^{n-r} (y_{[i]} - \mu_{2.1} - \theta x_{(i)})^2} \quad (1.4.1.2)$$

where $y_{[i]}$ ’s are the concomitants of $x_{(i)}$ ’s. Then,

$$\ln L = \text{constant} - (n-r)p \ln \sigma_1 + (p-1) \sum_{i=1}^{n-r} \ln x_{(i)} - \sum_{i=1}^{n-r} z_{(i)1}^p - rz_{(n-r)1}^p - (n-r) \ln \sigma_{2.1} - \frac{1}{2} \sum_{i=1}^{n-r} z_{[i]2}^2 \quad (1.4.1.3)$$

where

$$z_{(i)1} = \frac{x_{(i)}}{\sigma_1} \text{ and } z_{[i]2} = \frac{(y[i] - \mu_{2.1} - \theta x_{(i)})}{\sigma_{2.1}}.$$

The solutions of the following maximum likelihood equations are the ML estimators:

$$\frac{\partial \ln L}{\partial \sigma_1} = -\frac{(n-r)p}{\sigma_1} + \frac{p}{\sigma_1} \sum_{i=1}^{n-r} z_{(i)1}^p + \frac{rp}{\sigma_1} z_{(n-r)1}^p = 0, \quad (1.4.1.4)$$

$$\frac{\partial \ln L}{\partial \mu_{2.1}} = \frac{1}{\sigma_{2.1}} \sum_{i=1}^{n-r} z_{[i]2} = 0, \quad (1.4.1.5)$$

$$\frac{\partial \ln L}{\partial \sigma_{2.1}} = -\frac{(n-r)}{\sigma_{2.1}} + \frac{1}{\sigma_{2.1}} \sum_{i=1}^{n-r} z_{[i]2}^2 = 0 \quad (1.4.1.6)$$

and

$$\frac{\partial \ln L}{\partial \theta} = \frac{1}{\sigma_{2.1}} \sum_{i=1}^{n-r} x_{(i)} z_{[i]2} = 0. \quad (1.4.1.7)$$

The ML estimators are

$$\hat{\sigma}_1 = \left[\frac{1}{(n-r)} \left(\sum_{i=1}^{n-r} x_{(i)}^p + rx_{(n-r)}^p \right) \right]^{1/p}, \hat{\mu}_{2.1} = \bar{y}_{[.]} - \hat{\theta} \bar{x}_{(.)},$$

$$\hat{\sigma}_{2,1} = \sqrt{\frac{1}{n-r} \sum_{i=1}^{n-r} (y[i] - \hat{\mu}_{2,1} - \hat{\theta}x_{(i)})^2}, \quad \hat{\theta} = \frac{\sum_{i=1}^{n-r} (x_{(i)} - \bar{x}_{(.)}) (y[i] - \bar{y}[.])}{\sum_{i=1}^{n-r} (x_{(i)} - \bar{x}_{(.)})^2} \quad (1.4.1.8)$$

where

$$\bar{y}[.] = \frac{1}{n-r} \sum_{i=1}^{n-r} y[i] \quad \text{and} \quad \bar{x}_{(.)} = \frac{1}{n-r} \sum_{i=1}^{n-r} x_{(i)}.$$

1.4.2. Sample Information Matrix

Since it is difficult to take the expectation of concomitant terms, instead of obtaining Fisher information matrix $I(\sigma_1, \mu_{2,1}, \sigma_{2,1}, \theta)$, the sample information matrix $\hat{I}$ is calculated, and in Cox (1972, 1975) it is stated that using the sample second derivatives is a good approximation to the Fisher information matrix. $\hat{I}$ is equal to,

$$\hat{I} = [\hat{I}_{ij}] = \left[-\frac{\partial^2 \ln L}{\partial \delta_i \partial \delta_j} \right] \quad i, j = 1, 2, 3, 4, \quad (1.4.2.1)$$

$$\delta_1 = \hat{\sigma}_1, \quad \delta_2 = \hat{\mu}_{2,1}, \quad \delta_3 = \hat{\sigma}_{2,1}, \quad \delta_4 = \hat{\theta}.$$

The components of $\hat{I}$ are given in the following equations:

$$\hat{I}_{11} = -\frac{(n-r)p}{\hat{\sigma}_1^2} + (1+p) \frac{p}{\hat{\sigma}_1^2} \sum_{i=1}^{n-r} z_{(i)1}^p + (1+p) \frac{rp}{\hat{\sigma}_1^2} z_{(n-r)1}^p, \quad (1.4.2.2)$$

$$\hat{I}_{12} = \hat{I}_{13} = \hat{I}_{14} = 0.0, \quad (1.4.2.3)$$

$$\hat{I}_{22} = \frac{(n-r)}{\hat{\sigma}_{2,1}^2}, \quad (1.4.2.4)$$

$$\hat{I}_{23} = \frac{2}{\hat{\sigma}_{2,1}^2} \sum_{i=1}^{n-r} z[i]_2, \quad (1.4.2.5)$$

$$\hat{I}_{24} = \frac{1}{\hat{\sigma}_{2,1}^2} \sum_{i=1}^{n-r} x_{(i)}, \quad (1.4.2.6)$$

$$\hat{I}_{33} = -\frac{(n-r)}{\hat{\sigma}_{2,1}^2} + \frac{3}{\hat{\sigma}_{2,1}^2} \sum_{i=1}^{n-r} z[i]_2^2, \quad (1.4.2.7)$$

$$\hat{I}_{34} = \frac{2}{\hat{\sigma}_{2,1}^2} \sum_{i=1}^{n-r} x_{(i)} z[i]_2, \quad (1.4.2.8)$$

and

$$\hat{I}_{44} = \frac{1}{\hat{\sigma}_{2,1}^2} \sum_{i=1}^{n-r} x_{(i)}^2. \quad (1.4.2.9)$$

In $z(i)_1$ and $z[i]_2$, the unknown parameters are replaced by their estimators.

CHAPTER 2

EVALUATING THE EFFECT OF THE LARGEST OBSERVATIONS IN THE MARGINAL DISTRIBUTION

Suppose (X,Y) have a bivariate distribution and Y depends on X as it is stated in Chapter 1. Let (x_i, y_i) be a random sample. X represents survival and let $x_{(1)} \leq x_{(2)} \leq \dots \leq x_{(n)}$ be the order statistics and

$$t_{(i)} = E\left(\frac{x_{(i)} - \mu_1}{\sigma_1}\right), \quad (2.1)$$

μ_1 and σ_1 being the location and scale parameters in the marginal distribution of X . It is reasonable to assume that the cost of measuring $x_{(i)}$, denoted by c_i , is directly related to its expected value. Specifically,

$$c_i = c + \mu_1 + \sigma_1 t_{(i)} > 0, \quad (2.2)$$

c being a threshold cost. Total cost of measuring n observations is

$$C = nc + n\mu_1 + \sigma_1 \sum_{i=1}^n t_{(i)}. \quad (2.3)$$

Therefore,

$$\hat{n} = \frac{C}{c + \hat{\mu}_1 + \hat{\sigma}_1 \frac{1}{n} \sum_{i=1}^n t(i)}, \quad (2.4)$$

for skew marginal distributions, and

$$\hat{n} = \frac{C}{c + \hat{\mu}_1}, \quad (2.5)$$

for symmetric marginal distributions. $t(i)$'s are approximated by

$$\int_{-\infty}^{t_i} f(z) dz = \frac{i}{n+1} \quad (1 \leq i \leq n)$$

where $f(z)$ is the probability density function of $z = (x - \mu)/\sigma$. The true values of $t(i)$ ($1 \leq i \leq n$) are available for many distributions, however. The equation (2.3) is valid for complete samples. For right censored samples (r largest observations censored),

$$\sum_{i=1}^{n-r} c_i = \left(C - \sum_{i=n-r+1}^n c_i \right) = (n-r)(c + \mu_1) + \sigma_1 \sum_{i=1}^{n-r} t(i),$$

$$C_1 = \left[c + \mu_1 + \frac{\sigma_1}{n-r} \sum_{i=1}^{n-r} t(i) \right] (n-r) \quad (2.6)$$

$$\hat{n} = r + \frac{C_1}{c + \hat{\mu}_1 + \frac{\hat{\sigma}_1}{n-r} \sum_{i=1}^{n-r} t(i)} < r + \frac{C}{c + \hat{\mu}_1 + \frac{\hat{\sigma}_1}{n-r} \sum_{i=1}^{n-r} t(i)}. \quad (2.7)$$

We will give the values of $\hat{n}$ using the permissible total C .

If in a sample, some observations are too influential (e.g., outliers), it will surely have an effect on the estimators $\hat{\mu}_1$, $\hat{\sigma}_1$, $\hat{\mu}_{2.1}$, $\hat{\sigma}_{2.1}$, $\hat{\theta}$, $\hat{\mu}_2$, $\hat{\sigma}_2$ and $\hat{\rho}$. This effect will be evaluated in this Chapter.

2.1. Marginal Distribution is Generalized Logistic and Conditional Distribution is Normal

2.1.1. Complete Sample

In section (2.1.1), X is Generalized Logistic and $Y|X=x$ is normal with density function given in equation (1.2.2).

The probability density function of X is

$$f(x) = \frac{b}{\sigma_1} \frac{\exp\left[-\left(\frac{x - \mu_1}{\sigma_1}\right)\right]}{\left\{1 + \exp\left[-\left(\frac{x - \mu_1}{\sigma_1}\right)\right]\right\}^{b+1}},$$

$$-\infty < x < \infty, -\infty < \mu_1 < \infty, b > 0, \sigma_1 > 0. \quad (2.1.1.1)$$

Note that for $b < 1$, $b = 1$ and $b > 1$, $GL(b, \sigma)$ represents negatively skewed, symmetric and positively skewed distributions, respectively. Therefore, this distribution has a broad application area in real life (Thode, 2002; Tiku and Vaughan, 1997; Agresti, 1996; Hosmer and Lemeshow, 1989; Aitken et al., 1989; Berkson, 1951).

2.1.1.1. Estimation of Parameters

For a random sample (x_i, y_i) $1 \leq i \leq n$, the likelihood function is

$$L \propto \frac{b^n}{\sigma_1^n} \prod_{i=1}^n x_i^{p-1} \frac{e^{-\sum_{i=1}^n \left(\frac{x_i - \mu_1}{\sigma_1} \right)}}{\prod_{i=1}^n \left[1 + e^{-\left(\frac{x_i - \mu_1}{\sigma_1} \right)} \right]^{b+1}} \frac{1}{\sigma_{2.1}^n} \exp \left\{ -\frac{1}{2\sigma_{2.1}^2} \sum_{i=1}^n (y_i - \mu_{2.1} - \theta x_i)^2 \right\} \quad (2.1.1.1.1)$$

Then,

$$\ln L \propto n \ln b - n \ln \sigma_1 + \sum_{i=1}^n z_i - (b+1) \sum_{i=1}^n \ln(1 + e^{-z_i}) - n \ln \sigma_{2.1} - \frac{1}{2\sigma_{2.1}^2} \sum_{i=1}^n e_i^2 \quad (2.1.1.1.2)$$

where

$$z_i = \frac{x_i - \mu_1}{\sigma_1} \text{ and } e_i = y_i - \mu_{2.1} - \theta x_i.$$

The ML estimators are the solutions of the following equations:

$$\frac{\partial \ln L}{\partial \mu_1} = \frac{n}{\sigma_1} - \frac{(b+1)}{\sigma_1} \sum_{i=1}^n \frac{e^{-z_i}}{1 + e^{-z_i}} = 0, \quad (2.1.1.1.3)$$

$$\frac{\partial \ln L}{\partial \sigma_1} = -\frac{n}{\sigma_1} + \frac{1}{\sigma_1} \sum_{i=1}^n z_i - \frac{(b+1)}{\sigma_1} \sum_{i=1}^n z_i \frac{e^{-z_i}}{1 + e^{-z_i}} = 0, \quad (2.1.1.1.4)$$

$$\frac{\partial \ln L}{\partial \mu_{2.1}} = \frac{1}{\sigma_{2.1}^2} \sum_{i=1}^n e_i = 0, \quad (2.1.1.1.5)$$

$$\frac{\partial \ln L}{\partial \sigma_{2.1}} = -\frac{n}{\sigma_{2.1}} + \frac{1}{\sigma_{2.1}^3} \sum_{i=1}^n e_i^2 = 0 \quad (2.1.1.1.6)$$

$$\frac{\partial \ln L}{\partial \theta} = \frac{1}{\sigma_{2.1}^2} \sum_{i=1}^n x_i e_i = 0. \quad (2.1.1.1.7)$$

The ML estimators $\hat{\mu}_{2.1}$, $\hat{\sigma}_{2.1}$ and $\hat{\theta}$ are

$$\hat{\mu}_{2.1} = \frac{1}{n} \sum_{i=1}^n y_i - \frac{\hat{\theta}}{n} \sum_{i=1}^n x_i = \bar{y} - \hat{\theta} \bar{x},$$

$$\hat{\sigma}_{2.1} = \sqrt{\frac{\sum_{i=1}^n (y_i - \hat{\mu}_{2.1} - \hat{\theta} x_i)^2}{n}}$$

and

$$\hat{\theta} = \frac{\sum_{i=1}^n (x_i - \bar{x}) y_i}{\sum_{i=1}^n (x_i - \bar{x})^2}. \quad (2.1.1.1.8)$$

However, because of the function $e^{-z_i} / (1 + e^{-z_i})$, it is very difficult to solve the equations (2.1.1.1.3) and (2.1.1.1.4); ML estimators are, therefore, elusive. Therefore, MML methodology is applied by linearizing this function with the help of the first two terms of Taylor series expansion. For the MML methodology, see Tiku et. al. (1986).

In order to get the MML estimators, first the random sample $x_1, x_2, \dots, x_n$ is arranged in ascending order so that $x_{(1)} \leq x_{(2)} \leq \dots \leq x_{(n)}$. Then, $y[i]$'s, which are the concomitants of $x_{(i)}$'s, are obtained. For $1 \leq i \leq n$, $z(i)$ and $e[i]$ are,

$$z(i) = \frac{x_{(i)} - \mu_1}{\sigma_1} \text{ and } e[i] = y[i] - \mu_{2.1} - \theta x_{(i)}.$$

Note that since complete sums are invariant to ordering, the likelihood equations presented in the equations (2.1.1.1.3) - (2.1.1.1.4) can also be written as

$$\frac{\partial \ln L}{\partial \mu_1} = \frac{n}{\sigma_1} - \frac{(b+1)}{\sigma_1} \sum_{i=1}^n \frac{e^{-z(i)}}{1+e^{-z(i)}} = 0$$

and

$$\frac{\partial \ln L}{\partial \sigma_1} = -\frac{n}{\sigma_1} + \frac{1}{\sigma_1} \sum_{i=1}^n z(i) - \frac{(b+1)}{\sigma_1} \sum_{i=1}^n z(i) \frac{e^{-z(i)}}{1+e^{-z(i)}} = 0. \quad (2.1.1.1.9)$$

If $e^{-z(i)} / (1+e^{-z(i)})$ terms are replaced by $g(z(i))$, the likelihood equations are

$$\frac{\partial \ln L}{\partial \mu_1} = \frac{n}{\sigma_1} - \frac{(b+1)}{\sigma_1} \sum_{i=1}^n g(z(i)) = 0$$

and

$$\frac{\partial \ln L}{\partial \sigma_1} = -\frac{n}{\sigma_1} + \frac{1}{\sigma_1} \sum_{i=1}^n z(i) - \frac{(b+1)}{\sigma_1} \sum_{i=1}^n z(i) g(z(i)) = 0. \quad (2.1.1.1.10)$$

Next step is to linearize the $g(z(i))$ function which is done by using the first two terms of a Taylor series expansion around $E(z(i)) = t(i)$:

$$\begin{aligned} g(z(i)) &\cong g(t(i)) + (z(i) - t(i)) \left. \frac{dg(z(i))}{dz(i)} \right|_{z(i)=t(i)} \\ &= \alpha_i - \beta_i z(i) \quad (1 \leq i \leq n) \end{aligned} \quad (2.1.1.1.11)$$

where

$$\alpha_i = \frac{e^{-t(i)}}{\left(1 + e^{-t(i)}\right)^2} \left(1 + t(i) + e^{-t(i)}\right) \text{ and } \beta_i = \frac{e^{-t(i)}}{\left(1 + e^{-t(i)}\right)^2}; \quad (2.1.1.1.12)$$

$t(i)$ values for $n < 15$ are available in Balakrishnan and Leung (1988). For $n \geq 10$, the solutions of

$$\int_{-\infty}^{t(i)} b \frac{e^{-z}}{\left(1 + e^{-z}\right)^{b+1}} dz = \frac{i}{n+1} \quad (1 \leq i \leq n) \quad (2.1.1.1.13)$$

are the approximate values of $t(i)$ and they are equal to

$$t(i) = -\ln\left(q_i^{-1/b} - 1\right) \text{ and } q_i = i/(n+1). \quad (2.1.1.1.14)$$

We use the values given by (2.1.1.1.14) always. That hardly has any detrimental effect on the efficiencies of the MML estimators.

Then, the modified likelihood equations are

$$\frac{\partial \ln L}{\partial \mu_1} \cong \frac{\partial \ln L^*}{\partial \mu_1} = \frac{n}{\sigma_1} - \frac{(b+1)}{\sigma_1} \sum_{i=1}^n (\alpha_i - \beta_i z(i)) = 0$$

and

$$\frac{\partial \ln L}{\partial \sigma_1} \cong \frac{\partial \ln L^*}{\partial \sigma_1} = -\frac{n}{\sigma_1} + \frac{1}{\sigma_1} \sum_{i=1}^n z(i) - \frac{(b+1)}{\sigma_1} \sum_{i=1}^n z(i) (\alpha_i - \beta_i z(i)) = 0. \quad (2.1.1.1.15)$$

The MML estimators which are the solutions of the equations (2.1.1.1.15) are

$$\hat{\mu}_1 = K + D\hat{\sigma}_1$$

and

$$\hat{\sigma}_1 = \frac{B + \sqrt{B^2 + 4nC}}{2n} \quad (2.1.1.1.16)$$

where

$$\begin{aligned} K &= \frac{1}{m} \sum_{i=1}^n \beta_i x_{(i)}, \quad D = \frac{1}{m} \sum_{i=1}^n \left(\frac{1}{b+1} - \alpha_i \right), \quad m = \sum_{i=1}^n \beta_i, \\ B &= (b+1) \sum_{i=1}^n \left(x_{(i)} - K \right) \left(\frac{1}{b+1} - \alpha_i \right) \\ C &= (b+1) \sum_{i=1}^n \beta_i \left(x_{(i)} - K \right)^2. \end{aligned}$$

Remark: Note that like ML estimators, MML estimators have the invariance property. Therefore, $\hat{\mu}_2$, $\hat{\sigma}_2$ and $\hat{\rho}$ are

$$\hat{\mu}_2 = \hat{\mu}_{2.1} + \hat{\theta} \{ \hat{\mu}_1 + \hat{\sigma}_1 [\psi(b) - \psi(1)] \}, \quad \hat{\sigma}_2 = \sqrt{\hat{\sigma}_{2.1}^2 + \hat{\theta}^2 \hat{\sigma}_1^2} \quad \text{and} \quad \hat{\rho} = \hat{\theta} \frac{\hat{\sigma}_1}{\hat{\sigma}_2}. \quad (2.1.1.1.17)$$

2.1.1.2. Fisher Information Matrix

The components of the information matrix $I(\mu_1, \sigma_1, \mu_{2.1}, \sigma_{2.1}, \theta)$ are given in the equations (2.1.1.2.1) – (2.1.1.2.11):

$$I_{11} = \left(\frac{b}{b+2} \right) \frac{n}{\sigma_1^2}, \quad (2.1.1.2.1)$$

$$I_{12} = \left(\frac{b}{b+2} \right) [\psi(b+1) - \psi(2)] \frac{n}{\sigma_1^2}, \quad (2.1.1.2.2)$$

$$I_{13} = I_{14} = I_{15} = 0, \quad (2.1.1.2.3)$$

$$I_{22} = \frac{n}{\sigma_1^2} + \left(\frac{b}{b+2} \right) [\psi'(b+1) + \psi'(2) + \{\psi(b+1) - \psi(2)\}^2] \frac{n}{\sigma_1^2}, \quad (2.1.1.2.4)$$

$$I_{23} = I_{24} = I_{25} = 0, \quad (2.1.1.2.5)$$

$$I_{33} = \frac{n}{\sigma_{2.1}^2}, \quad (2.1.1.2.6)$$

$$I_{34} = 0, \quad (2.1.1.2.7)$$

$$I_{35} = [\mu_1 + \sigma_1 \{\psi(b) - \psi(1)\}] \frac{n}{\sigma_{2.1}^2}, \quad (2.1.1.2.8)$$

$$I_{44} = \frac{2n}{\sigma_{2.1}^2}, \quad (2.1.1.2.9)$$

$$I_{45} = 0 \quad (2.1.1.2.10)$$

and

$$I_{55} = \left[\mu_1^2 + 2\mu_1\sigma_1 \{\psi(b) - \psi(1)\} + \sigma_1^2 \{ \psi'(b) + \psi'(1) + [\psi(b) - \psi(1)]^2 \} \right] \frac{n}{\sigma_{2.1}^2}. \quad (2.1.1.2.11)$$

2.1.2. Censored Sample

In Section 2.1.2, the sample (x_i, y_i) is ordered with respect to x_i and $(x_{(i)}, y_{[i]})$ $1 \leq i \leq n-r$ is taken into consideration. The marginal distribution considered is ‘censored’ Generalized Logistic while the conditional distribution is normal.

2.1.2.1. Estimation of Parameters

For a Type II censored sample $(x_{(i)}, y_{[i]})$ ($1 \leq i \leq n-r$), the likelihood function is

$$L \propto \left(\frac{b}{\sigma_1}\right)^{n-r} \frac{e^{-\sum_{i=1}^{n-r} z_{(i)}}}{\prod_{i=1}^{n-r} (1 + e^{-z_{(i)}})^{b+1}} [1 - F(z_{(n-r)})]^r \frac{1}{\sigma_{2.1}^{n-r}} e^{-\frac{1}{2\sigma_{2.1}^2} \sum_{i=1}^{n-r} e_{[i]}^2} \quad (2.1.2.1.1)$$

where

$$z_{(i)} = \frac{x_{(i)} - \mu_1}{\sigma_1}, \quad F(z) = \frac{1}{(1 + e^{-z})^b} \quad \text{and} \quad e_{[i]} = y_{[i]} - \mu_{2.1} - \theta x_{(i)}.$$

Then,

$$\begin{aligned} \ln L \propto & -(n-r) \ln \sigma_1 - \sum_{i=1}^{n-r} z_{(i)} - (b+1) \sum_{i=1}^{n-r} \ln(1 + e^{-z_{(i)}}) + r \ln[1 - F(z_{(n-r)})] \\ & - (n-r) \ln \sigma_{2.1} - \frac{1}{2\sigma_{2.1}^2} \sum_{i=1}^{n-r} e_{[i]}^2. \end{aligned} \quad (2.1.2.1.2)$$

The likelihood equations are given in the equations (2.1.2.1.3) – (2.1.2.1.7):

$$\frac{\partial \ln L}{\partial \mu_1} = \frac{(n-r)}{\sigma_1} - \frac{(b+1)}{\sigma_1} \sum_{i=1}^{n-r} \frac{e^{-z_{(i)}}}{1 + e^{-z_{(i)}}} + \frac{r}{\sigma_1} \frac{f(z_{(n-r)})}{1 - F(z_{(n-r)})} = 0, \quad (2.1.2.1.3)$$

$$\frac{\partial \ln L}{\partial \sigma_1} = -\frac{(n-r)}{\sigma_1} + \frac{1}{\sigma_1} \sum_{i=1}^{n-r} z(i) - \frac{(b+1)}{\sigma_1} \sum_{i=1}^{n-r} z(i) \frac{e^{-z(i)}}{1+e^{-z(i)}} +$$

$$r \frac{z(n-r)}{\sigma_1} \frac{f(z(n-r))}{1-F(z(n-r))} = 0, (2.1.2.1.4)$$

$$\frac{\partial \ln L}{\partial \mu_{2.1}} = \sum_{i=1}^{n-r} \frac{e[i]}{\sigma_{2.1}^2} = 0, (2.1.2.1.5)$$

$$\frac{\partial \ln L}{\partial \sigma_{2.1}} = -\frac{(n-r)}{\sigma_{2.1}} + \sum_{i=1}^{n-r} \frac{e[i]^2}{\sigma_{2.1}^3} = 0 (2.1.2.1.6)$$

and

$$\frac{\partial \ln L}{\partial \theta} = \sum_{i=1}^{n-r} \frac{x(i)e[i]}{\sigma_{2.1}^2} = 0. (2.1.2.1.7)$$

The estimators $\hat{\mu}_{2.1}$, $\hat{\sigma}_{2.1}$ and $\hat{\theta}$ are given below,

$$\hat{\mu}_{2.1} = \frac{\sum_{i=1}^{n-r} y[i]}{n-r} - \hat{\theta} \frac{\sum_{i=1}^{n-r} x(i)}{n-r} = \bar{y}[\cdot] - \hat{\theta} \bar{x}(\cdot),$$

$$\hat{\sigma}_{2.1} = \sqrt{\frac{\sum_{i=1}^{n-r} (y[i] - \hat{\mu}_{2.1} - \hat{\theta} x(i))^2}{n-r}}$$

and

$$\hat{\theta} = \frac{\sum_{i=1}^{n-r} (x(i) - \bar{x}(\cdot))(y[i] - \bar{y}[\cdot])}{\sum_{i=1}^{n-r} (x(i) - \bar{x}(\cdot))^2}. (2.1.2.1.8)$$

However, because of the functions $g_1(z(i)) = \frac{e^{-z(i)}}{1 + e^{-z(i)}}$ and

$$g_2(z_{(n-r)}) = \frac{f(z_{(n-r)})}{1 - F(z_{(n-r)})}, \text{ the remaining likelihood equations given below but}$$

are very difficult to solve:

$$\frac{\partial \ln L}{\partial \mu_1} = \frac{(n-r)}{\sigma_1} - \frac{(b+1)}{\sigma_1} \sum_{i=1}^{n-r} g_1(z(i)) + \frac{r}{\sigma_1} g_2(z_{(n-r)}) = 0$$

and

$$\begin{aligned} \frac{\partial \ln L}{\partial \sigma_1} = & -\frac{(n-r)}{\sigma_1} + \frac{1}{\sigma_1} \sum_{i=1}^{n-r} z(i) - \frac{(b+1)}{\sigma_1} \sum_{i=1}^{n-r} z(i) g_1(z(i)) + \\ & r \frac{z_{(n-r)}}{\sigma_1} g_2(z_{(n-r)}) = 0, \quad (2.1.2.1.9) \end{aligned}$$

Therefore, these functions need to be linearized by using the procedure explained in Section 2.1.1.1. Thus,

$$g_1(z(i)) = \alpha_{1i} - \beta_{1i} z(i) \quad (1 \leq i \leq n-r) \quad (2.1.2.1.10)$$

where α_{1i} 's and β_{1i} 's are given in equations (2.1.1.1.12).

$$g_2(z_{(n-r)}) = \alpha_2 - \beta_2 z_{(n-r)}, \quad (2.1.2.1.11)$$

where

$$\beta_2 = - \left. \frac{dg_2(z_{(n-r)})}{dz_{(n-r)}} \right|_{z_{(n-r)} = t_{(n-r)}},$$

$$= - \frac{f'(z_{(n-r)})[1 - F(z_{(n-r)})] + [f(z_{(n-r)})]^2}{[1 - F(z_{(n-r)})]^2} \Big|_{z_{(n-r)}=t_{(n-r)}}$$

$$\alpha_2 = g_2(z_{(n-r)}) + \beta_2 z_{(n-r)} \text{ at } z_{(n-r)} = t_{(n-r)} \quad (2.1.2.1.12)$$

and

$$t_{(n-r)} = -\ln[q_{(n-r)}^{-1/b} - 1], \quad q_{(n-r)} = 1 - \frac{r}{n+1}.$$

Then, by replacing the equations $g_1(z_{(i)})$ and $g_2(z_{(n-r)})$ by (2.1.2.1.10) and (2.1.2.1.11) respectively, the following modified likelihood equations are obtained,

$$\frac{\partial \ln L}{\partial \mu_1} \cong \frac{\partial \ln L^*}{\partial \mu_1} = \frac{(n-r)}{\sigma_1} - \frac{(b+1)}{\sigma_1} \sum_{i=1}^{n-r} (\alpha_{1i} - \beta_{1i} z_{(i)}) + \frac{r}{\sigma_1} (\alpha_2 - \beta_2 z_{(n-r)}) = 0$$

and

$$\begin{aligned} \frac{\partial \ln L}{\partial \sigma_1} \cong \frac{\partial \ln L^*}{\partial \sigma_1} &= -\frac{(n-r)}{\sigma_1} + \frac{1}{\sigma_1} \sum_{i=1}^{n-r} z_{(i)} - \frac{(b+1)}{\sigma_1} \sum_{i=1}^{n-r} z_{(i)} (\alpha_{1i} - \beta_{1i} z_{(i)}) + \\ &\frac{r}{\sigma_1} z_{(n-r)} (\alpha_2 - \beta_2 z_{(n-r)}) = 0. \end{aligned} \quad (2.1.2.1.13)$$

The MML estimators which are the solutions of the equations (2.1.2.1.13) are given below:

$$\hat{\mu}_1 = K + D\hat{\sigma}_1$$

and

$$\hat{\sigma}_1 = \frac{B + \sqrt{B^2 + 4(n-r)C}}{2(n-r)}. \quad (2.1.2.1.14)$$

Here,

$$K = \frac{(b+1) \sum_{i=1}^{n-r} \beta_{1i} x_{(i)} - r\beta_2 x_{(n-r)}}{(b+1) \sum_{i=1}^{n-r} \beta_{1i} - r\beta_2},$$

$$D = \frac{(n-r) - (b+1) \sum_{i=1}^{n-r} \alpha_{1i} + r\alpha_2}{(b+1) \sum_{i=1}^{n-r} \beta_{1i} - r\beta_2},$$

$$B = (b+1) \sum_{i=1}^{n-r} \left(x_{(i)} - K \left(\frac{1}{b+1} - \alpha_{1i} \right) + r\alpha_2 (x_{(n-r)} - K) \right)$$

and

$$C = (b+1) \sum_{i=1}^{n-r} \beta_{1i} (x_{(i)} - K)^2 - r\beta_2 (x_{(n-r)} - K)^2. \quad (2.1.2.1.15)$$

The MML estimators $\hat{\mu}_2$, $\hat{\sigma}_2$ and $\hat{\rho}$ are the same as in equation (2.1.1.1.17).

2.1.2.2. Sample Information Matrix

Because of the reasons stated in Chapter 1, Fisher information matrix $I(\mu_1, \sigma_1, \mu_{2.1}, \sigma_{2.1}, \theta)$ is intractable. Therefore, instead of I , the sample information matrix $\hat{I}$ is used:

$$\hat{I} = [\hat{I}_{ij}] = \left[-\frac{\partial^2 \ln L}{\partial \delta_i \partial \delta_j} \right] \quad i,j=1,2,3,4,5$$

$$\delta_1 = \hat{\mu}_1, \delta_2 = \hat{\sigma}_1, \delta_3 = \hat{\mu}_{2,1}, \delta_4 = \hat{\sigma}_{2,1}, \delta_5 = \hat{\theta},$$

$$\hat{I}_{11} = \frac{(b+1)}{\hat{\sigma}_1^2} \sum_{i=1}^{n-r} \frac{e^{-z(i)}}{\left\{ 1 + e^{-z(i)} \right\}^2} + \frac{r}{\hat{\sigma}_1^2} \frac{dg_2(z_{(n-r)})}{dz_{(n-r)}}, \quad (2.1.2.2.1)$$

$$\begin{aligned} \hat{I}_{12} = & \frac{(n-r)}{\hat{\sigma}_1^2} - \frac{(b+1)}{\hat{\sigma}_1^2} \sum_{i=1}^{n-r} g_1(z(i)) + \frac{(b+1)}{\hat{\sigma}_1^2} \sum_{i=1}^{n-r} z(i) \frac{e^{-z(i)}}{\left\{ 1 + e^{-z(i)} \right\}^2} + \\ & \frac{r}{\hat{\sigma}_1^2} g_2(z_{(n-r)}) + \frac{r}{\hat{\sigma}_1^2} z_{(n-r)} \frac{dg_2(z_{(n-r)})}{dz_{(n-r)}}, \end{aligned} \quad (2.1.2.2.2)$$

$$\hat{I}_{13} = \hat{I}_{14} = \hat{I}_{15} = 0.0, \quad (2.1.2.2.3)$$

$$\begin{aligned} \hat{I}_{22} = & -\frac{(n-r)}{\hat{\sigma}_1^2} + \frac{2}{\hat{\sigma}_1^2} \sum_{i=1}^{n-r} z(i) - 2 \frac{(b+1)}{\hat{\sigma}_1^2} \sum_{i=1}^{n-r} z(i) g_1(z(i)) + \frac{(b+1)}{\hat{\sigma}_1^2} \sum_{i=1}^{n-r} z(i)^2 \times \\ & \frac{e^{-z(i)}}{\left\{ 1 + e^{-z(i)} \right\}^2} + \frac{2r}{\hat{\sigma}_1^2} z_{(n-r)} g_2(z_{(n-r)}) + \frac{r}{\hat{\sigma}_1^2} z_{(n-r)}^2 \frac{dg_2(z_{(n-r)})}{dz_{(n-r)}}, \end{aligned} \quad (2.1.2.2.4)$$

$$\hat{I}_{23} = \hat{I}_{24} = \hat{I}_{25} = 0.0, \quad (2.1.2.2.5)$$

$$\hat{I}_{33} = \frac{(n-r)}{\hat{\sigma}_{2,1}^2}, \quad (2.1.2.2.6)$$

$$\hat{I}_{34} = \frac{2}{\hat{\sigma}_{2,1}^3} \sum_{i=1}^{n-r} e[i], \quad (2.1.2.2.7)$$

$$\hat{I}_{35} = \frac{1}{\hat{\sigma}_{2,1}^2} \sum_{i=1}^{n-r} x_{(i)}, \quad (2.1.2.2.8)$$

$$\hat{I}_{44} = -\frac{(n-r)}{\hat{\sigma}_{2,1}^2} + \frac{3}{\hat{\sigma}_{2,1}^4} \sum_{i=1}^{n-r} e_{[i]}^2, \quad (2.1.2.2.9)$$

$$\hat{I}_{45} = \frac{2}{\hat{\sigma}_{2,1}^3} \sum_{i=1}^{n-r} x_{(i)} e[i] \quad (2.1.2.2.10)$$

$$\hat{I}_{55} = \frac{1}{\hat{\sigma}_{2.1}^2} \sum_{i=1}^{n-r} x_{(i)}^2. \quad (2.1.2.2.11)$$

In $z(i)$ and $e[i]$, the unknown parameters are replaced by their MML estimators.

2.2. Marginal Distribution is Weibull and Conditional Distribution is Generalized Logistic

2.2.1. Complete Sample

In section (2.2), X is coming from Weibull – for density function of X , see equation (1.2.1) – and the distribution of $Y|X=x$ is assumed to be Generalized Logistic with density function,

$$h(y|x) = \frac{b}{\sigma_{2.1}} \frac{\exp\left[-\left(\frac{y - \mu_{2.1} - \theta x}{\sigma_{2.1}}\right)\right]}{\left\{1 + \exp\left[-\left(\frac{y - \mu_{2.1} - \theta x}{\sigma_{2.1}}\right)\right]\right\}^{b+1}}, \quad -\infty < y < \infty, b > 0, \quad (2.2.1.1)$$

where $\sigma_{2.1}$, $\mu_{2.1}$ and θ are given in equation (1.2.3).

2.2.1.1. Estimation of Parameters

For a random sample (x_i, y_i) $1 \leq i \leq n$, the likelihood function is

$$L \propto \frac{p^n}{\sigma_1^{np}} \prod_{i=1}^n x_i^{p-1} e^{-\sum_{i=1}^n \left(\frac{x_i}{\sigma_1}\right)^p} \frac{b^n}{\sigma_{2.1}^n} \frac{e^{-\frac{1}{\sigma_{2.1}} \sum_{i=1}^n (y_i - \mu_{2.1} - \theta x_i)}}{\prod_{i=1}^n \left[1 + \exp\left(-\frac{y_i - \mu_{2.1} - \theta x_i}{\sigma_{2.1}}\right)\right]^{b+1}} \quad (2.2.1.1.1)$$

$$\ln L \propto -np \ln \sigma_1 + (p-1) \sum_{i=1}^n \ln x_i - \sum_{i=1}^n z_i^p - n \ln \sigma_{2,1} - \frac{1}{\sigma_{2,1}} \sum_{i=1}^n e_i - (b+1) \sum_{i=1}^n \ln \left\{ 1 + e^{-\frac{e_i}{\sigma_{2,1}}} \right\} \quad (2.2.1.1.2)$$

where

$$z_i = \frac{x_i - \mu_1}{\sigma_1} \text{ and } e_i = y_i - \mu_{2,1} - \theta x_i.$$

The likelihood equations are

$$\frac{\partial \ln L}{\partial \sigma_1} = -\frac{np}{\sigma_1} + \frac{p}{\sigma_1} \sum_{i=1}^n z_i^p = 0, \quad (2.2.1.1.3)$$

$$\frac{\partial \ln L}{\partial \mu_{2,1}} = \frac{n}{\sigma_{2,1}} - \frac{(b+1)}{\sigma_{2,1}} \sum_{i=1}^n \left(\frac{e^{-e_i/\sigma_{2,1}}}{1 + e^{-e_i/\sigma_{2,1}}} \right) = 0, \quad (2.2.1.1.4)$$

$$\frac{\partial \ln L}{\partial \sigma_{2,1}} = -\frac{n}{\sigma_{2,1}} + \frac{1}{\sigma_{2,1}^2} \sum_{i=1}^n e_i - \frac{(b+1)}{\sigma_{2,1}^2} \sum_{i=1}^n e_i \left(\frac{e^{-e_i/\sigma_{2,1}}}{1 + e^{-e_i/\sigma_{2,1}}} \right) = 0 \quad (2.2.1.1.5)$$

and

$$\frac{\partial \ln L}{\partial \theta} = \frac{1}{\sigma_{2,1}} \sum_{i=1}^n x_i - \frac{(b+1)}{\sigma_{2,1}} \sum_{i=1}^n x_i \left(\frac{e^{-e_i/\sigma_{2,1}}}{1 + e^{-e_i/\sigma_{2,1}}} \right) = 0. \quad (2.2.1.1.6)$$

The estimator $\hat{\sigma}_1$ is equal to,

$$\hat{\sigma}_1 = \left(\frac{1}{n} \sum_{i=1}^n x_i^p \right)^{1/p}. \quad (2.2.1.1.7)$$

However, because of the function $\frac{e^{-e_i/\sigma_{2.1}}}{(1 + e^{-e_i/\sigma_{2.1}})}$, the likelihood equations

(2.2.1.1.4) - (2.2.1.1.6) are difficult to solve and the corresponding ML estimators cannot be obtained. Therefore, in order to obtain the estimators $\hat{\mu}_{2.1}$, $\hat{\sigma}_{2.1}$ and $\hat{\theta}$, MML methodology is applied.

In order to obtain the corresponding estimators, first w_i 's are defined as $w_i = y_i - \theta x_i$ (for a given θ , $1 \leq i \leq n$). Note that since $\mu_{2.1}$ is a constant $e_i = w_i - \mu_{2.1}$ and w_i have the same order. If a_i is equal to $e_i/\sigma_{2.1}$, then

$$a(i) = \frac{e(i)}{\sigma_{2.1}} = \frac{w(i) - \mu_{2.1}}{\sigma_{2.1}} = \frac{y[i] - \theta x[i] - \mu_{2.1}}{\sigma_{2.1}}.$$

Then, the likelihood equations can be rewritten as

$$\frac{\partial \ln L}{\partial \mu_{2.1}} = \frac{n}{\sigma_{2.1}} - \frac{(b+1)}{\sigma_{2.1}} \sum_{i=1}^n \frac{e^{-a(i)}}{(1 + e^{-a(i)})} = 0,$$

$$\frac{\partial \ln L}{\partial \sigma_{2.1}} = -\frac{n}{\sigma_{2.1}} + \frac{1}{\sigma_{2.1}^2} \sum_{i=1}^n e(i) - \frac{(b+1)}{\sigma_{2.1}^2} \sum_{i=1}^n e(i) \frac{e^{-a(i)}}{(1 + e^{-a(i)})} = 0$$

and

$$\frac{\partial \ln L}{\partial \theta} = \frac{1}{\sigma_{2.1}} \sum_{i=1}^n x[i] - \frac{(b+1)}{\sigma_{2.1}} \sum_{i=1}^n x[i] \frac{e^{-a(i)}}{(1 + e^{-a(i)})} = 0. \quad (2.2.1.1.8)$$

If $\frac{e^{-a(i)}}{(1 + e^{-a(i)})}$ is defined as $g(a(i))$, the likelihood equations become

$$\frac{\partial \ln L}{\partial \mu_{2.1}} = \frac{n}{\sigma_{2.1}} - \frac{(b+1)}{\sigma_{2.1}} \sum_{i=1}^n g(a_{(i)}) = 0,$$

$$\frac{\partial \ln L}{\partial \sigma_{2.1}} = -\frac{n}{\sigma_{2.1}} + \frac{1}{\sigma_{2.1}^2} \sum_{i=1}^n e_{(i)} - \frac{(b+1)}{\sigma_{2.1}^2} \sum_{i=1}^n e_{(i)} g(a_{(i)}) = 0$$

and

$$\frac{\partial \ln L}{\partial \theta} = \frac{1}{\sigma_{2.1}} \sum_{i=1}^n x[i] - \frac{(b+1)}{\sigma_{2.1}} \sum_{i=1}^n x[i] g(a_{(i)}) = 0. \quad (2.2.1.1.9)$$

Now, $g(a_{(i)})$ function is linearized by using the first two terms of Taylor series expansion around $E(a_{(i)}) = t_{(i)}$.

$$\begin{aligned} g(a_{(i)}) &\cong g(t_{(i)}) + (a_{(i)} - t_{(i)}) \left. \frac{dg(a_{(i)})}{da_{(i)}} \right|_{a_{(i)}=t_{(i)}} \\ &= \alpha_i - \beta_i a_{(i)} \quad (1 \leq i \leq n) \end{aligned} \quad (2.2.1.1.10)$$

where α_i , β_i and t_i values are given in equations (2.1.1.1.12) and (2.1.1.1.14).

Then, the modified likelihood equations are

$$\frac{\partial \ln L}{\partial \mu_{2.1}} \cong \frac{\partial \ln L^*}{\partial \mu_{2.1}} = \frac{n}{\sigma_{2.1}} - \frac{(b+1)}{\sigma_{2.1}} \sum_{i=1}^n (\alpha_i - \beta_i a_{(i)}) = 0,$$

$$\frac{\partial \ln L}{\partial \sigma_{2.1}} \cong \frac{\partial \ln L^*}{\partial \sigma_{2.1}} = -\frac{n}{\sigma_{2.1}} + \frac{1}{\sigma_{2.1}^2} \sum_{i=1}^n e_{(i)} - \frac{(b+1)}{\sigma_{2.1}^2} \sum_{i=1}^n e_{(i)} (\alpha_i - \beta_i a_{(i)}) = 0$$

and

$$\frac{\partial \ln L}{\partial \theta} \cong \frac{\partial \ln L^*}{\partial \theta} = \frac{1}{\sigma_{2.1}} \sum_{i=1}^n x[i] - \frac{(b+1)}{\sigma_{2.1}} \sum_{i=1}^n x[i] (\alpha_i - \beta_i a_{(i)}) = 0. \quad (2.2.1.1.11)$$

The solutions of the equations given in (2.2.1.1.11) are the MML estimators and they are,

$$\hat{\mu}_{2.1} = \bar{y}[\cdot] - \hat{\theta} \bar{x}[\cdot] - \frac{\Delta}{m} \hat{\sigma}_{2.1},$$

$$\hat{\theta} = K - D \hat{\sigma}_{2.1}$$

and

$$\hat{\sigma}_{2.1} = \left(-B + \sqrt{B^2 + 4nC} \right) / 2n \quad (2.2.1.1.12)$$

where

$$m = \sum_{i=1}^n \beta_i, \quad \bar{y}[\cdot] = \sum_{i=1}^n \beta_i y[i] / m, \quad \bar{x}[\cdot] = \sum_{i=1}^n \beta_i x[i] / m,$$

$$\Delta_i = \left[\alpha_i - \frac{1}{b+1} \right], \quad \Delta = \sum_{i=1}^n \Delta_i,$$

$$K = \frac{\sum_{i=1}^n \beta_i (x[i] - \bar{x}[\cdot]) y[i]}{\sum_{i=1}^n \beta_i (x[i] - \bar{x}[\cdot])^2}, \quad D = \frac{\sum_{i=1}^n \Delta_i (x[i] - \bar{x}[\cdot])}{\sum_{i=1}^n \beta_i (x[i] - \bar{x}[\cdot])^2},$$

$$B = (b+1) \sum_{i=1}^n \Delta_i \{ (y[i] - \bar{y}[\cdot]) - K(x[i] - \bar{x}[\cdot]) \}$$

and

$$C = (b+1) \left\{ \sum_{i=1}^n \beta_i (y[i] - \bar{y}[\cdot])^2 - K \sum_{i=1}^n \beta_i (x[i] - \bar{x}[\cdot]) y[i] \right\}. \quad (2.2.1.1.13)$$

Also, $\hat{\mu}_2$, $\hat{\sigma}_2$ and $\hat{\rho}$ are given by

$$\hat{\mu}_2 = \hat{\mu}_{2.1} + \hat{\theta}\hat{\sigma}_1\Gamma\left(1+\frac{1}{p}\right), \quad \hat{\sigma}_2 = \sqrt{\hat{\sigma}_{2.1}^2 + \hat{\theta}^2\hat{\sigma}_1^2} \quad \text{and} \quad \hat{\rho} = \hat{\theta}\frac{\hat{\sigma}_1}{\hat{\sigma}_2}. \quad (2.2.1.1.14)$$

Remark: It may be noted that the MML estimator $\hat{\rho}$ is quite different from the Pearson correlation coefficient, given in equation (1.1.1.7). The latter can have substantial bias even for large n (Sazak et al, 2006, p. 82). Therefore, it should be used with caution.

2.2.1.2. Fisher Information Matrix

Note that if

$$h(z) = b \frac{e^{-z}}{(1+e^{-z})^{b+1}}, \quad b > 0, \quad -\infty < z < \infty, \quad (2.2.1.2.1)$$

then

$$\begin{aligned} E(Z) &= \psi(b) - \psi(1), \\ E(Z^2) &= \psi'(b) + \psi'(1) + \{\psi(b) - \psi(1)\}^2, \\ E\left[\frac{e^{-Z}}{1+e^{-Z}}\right] &= \frac{1}{(b+1)}, \\ E\left[Z \frac{e^{-Z}}{1+e^{-Z}}\right] &= \frac{1}{(b+1)} \{\psi(b) - \psi(2)\}, \\ E\left[\frac{e^{-Z}}{(1+e^{-Z})^2}\right] &= \frac{b}{(b+1)(b+2)}, \\ E\left[Z \frac{e^{-Z}}{(1+e^{-Z})^2}\right] &= \frac{b}{(b+1)(b+2)} \{\psi(b+1) - \psi(2)\} \end{aligned}$$

and

$$E \left[Z^2 \frac{e^{-Z}}{(1+e^{-Z})^2} \right] = \frac{b}{(b+1)(b+2)} \left[\psi'(b+1) + \psi'(2) + \{\psi(b+1) - \psi(2)\}^2 \right]. \quad (2.2.1.2.2)$$

For some particular b 's, the values of $\psi(b)$ and $\psi'(b)$ are given in Tiku and Akkaya (2004). For detailed information about psi-function, see Abramowitz and Stegun (1985).

According to the information given in equation (2.2.1.2.2), the components of the $I(\sigma_1, \mu_{2.1}, \sigma_{2.1}, \theta)$ are given below:

$$I_{11} = \frac{np^2}{\sigma_1^2}, \quad (2.2.1.2.3)$$

$$I_{12} = I_{13} = I_{14} = 0, \quad (2.2.1.2.4)$$

$$I_{22} = \frac{nb}{(b+2)\sigma_{2.1}^2}, \quad (2.2.1.2.5)$$

$$I_{23} = \frac{nb}{(b+2)\sigma_{2.1}^2} \{\psi(b+1) - \psi(2)\}, \quad (2.2.1.2.6)$$

$$I_{24} = \frac{nb}{(b+2)\sigma_{2.1}^2} \Gamma\left(1 + \frac{1}{p}\right) \sigma_1, \quad (2.2.1.2.7)$$

$$I_{33} = \frac{n}{\sigma_{2.1}^2} + \frac{nb}{(b+2)\sigma_{2.1}^2} \left[\psi'(b+1) + \psi'(2) + \{\psi(b+1) - \psi(2)\}^2 \right], \quad (2.2.1.2.8)$$

$$I_{34} = \frac{nb}{(b+2)\sigma_{2.1}^2} \{\psi(b+1) - \psi(2)\} \Gamma\left(1 + \frac{1}{p}\right) \sigma_1 \quad (2.2.1.2.9)$$

$$I_{44} = \frac{nb}{(b+2)\sigma_{2.1}^2} \Gamma\left(1 + \frac{2}{p}\right) \sigma_1^2. \quad (2.2.1.2.10)$$

2.2.2. Censored Sample

In Section 2.2.2, the estimation procedure of the parameters and the sample information matrix are presented when the marginal distribution is ‘censored’ Weibull and the conditonal distribution is Generalized Logistic.

2.2.2.1. Estimation of Parameters

The likelihood function for a a Type II censored sample $(x_{(i)}, y_{[i]})$ $(1 \leq i \leq n-r)$ is

$$L \propto \left[\prod_{i=1}^{n-r} \frac{p}{\sigma_1^p} x_{(i)}^{p-1} e^{-\left(\frac{x_{(i)}}{\sigma_1}\right)^p} \right] \left[e^{-\left(\frac{x_{(n-r)}}{\sigma_1}\right)^p} \right]^r \times$$

$$\left(\frac{b}{\sigma_{2.1}} \right)^{n-r} \frac{e^{-\sum_{i=1}^{n-r} \left(\frac{y_{[i]} - \mu_{2.1} - \theta x_{(i)}}{\sigma_{2.1}} \right)}}{\prod_{i=1}^{n-r} \left[1 + e^{-\frac{y_{[i]} - \mu_{2.1} - \theta x_{(i)}}{\sigma_{2.1}}} \right]^{b+1}}, \quad (2.2.2.1.1)$$

where $y_{[i]}$ ’s are the concomitants of $x_{(i)}$ ’s. Then,

$$\ln L \propto -(n-r)p \ln \sigma_1 + (p-1) \sum_{i=1}^{n-r} \ln x_{(i)} - \sum_{i=1}^{n-r} z_{(i)}^p - r(z_{(n-r)})^p$$

$$-(n-r) \ln \sigma_{2.1} - \frac{1}{2} \sum_{i=1}^{n-r} e_{[i]} - (b+1) \sum_{i=1}^{n-r} \ln \left\{ 1 + e^{-e_{[i]}/\sigma_{2.1}} \right\} \quad (2.2.2.1.2)$$

where

$$z(i) = \frac{x(i)}{\sigma_1}, \text{ and } e[i] = y[i] - \mu_{2.1} - \theta x(i).$$

The ML estimators are the solutions of the following likelihood equations;

$$\frac{\partial \ln L}{\partial \sigma_1} = -\frac{(n-r)p}{\sigma_1} + \frac{p}{\sigma_1} \sum_{i=1}^{n-r} z(i)^p + \frac{rp}{\sigma_1} z_{(n-r)}^p = 0, \quad (2.2.2.1.3)$$

$$\frac{\partial \ln L}{\partial \mu_{2.1}} = \frac{(n-r)}{\sigma_{2.1}} - \frac{(b+1)}{\sigma_{2.1}} \sum_{i=1}^{n-r} \frac{e^{-e[i]/\sigma_{2.1}}}{\left(1 + e^{-e[i]/\sigma_{2.1}}\right)} = 0, \quad (2.2.2.1.4)$$

$$\frac{\partial \ln L}{\partial \sigma_{2.1}} = -\frac{(n-r)}{\sigma_{2.1}} + \frac{1}{\sigma_{2.1}^2} \sum_{i=1}^{n-r} e[i] - \frac{(b+1)}{\sigma_{2.1}^2} \sum_{i=1}^{n-r} e[i] \frac{e^{-e[i]/\sigma_{2.1}}}{\left(1 + e^{-e[i]/\sigma_{2.1}}\right)} = 0 \quad (2.2.2.1.5)$$

and

$$\frac{\partial \ln L}{\partial \theta} = \frac{1}{\sigma_{2.1}} \sum_{i=1}^{n-r} x(i) - \frac{(b+1)}{\sigma_{2.1}} \sum_{i=1}^{n-r} x(i) \frac{e^{-e[i]/\sigma_{2.1}}}{\left(1 + e^{-e[i]/\sigma_{2.1}}\right)} = 0, \quad (2.2.2.1.6)$$

The ML estimator of σ_1 is

$$\hat{\sigma}_1 = \left[\frac{1}{(n-r)} \left(\sum_{i=1}^{n-r} x(i)^p + rx_{(n-r)}^p \right) \right]^{1/p}. \quad (2.2.2.1.7)$$

However, because of the function $\frac{e^{-e[i]/\sigma_{2.1}}}{\left(1 + e^{-e[i]/\sigma_{2.1}}\right)}$, the equations

(2.2.2.1.4) – (2.2.2.1.6) are difficult to solve. Therefore, to obtain the estimators, MML methodology is applied. In order to get the MML estimators, first $a[i]$'s are defined as $e[i]/\sigma_{2.1}$, then they are rearranged in ascending order and the corresponding concomitant terms, $(y[i], x[i])$'s, are obtained.

After applying the procedure explained above, the likelihood equations can be rewritten as

$$\frac{\partial \ln L}{\partial \mu_{2,1}} = \frac{(n-r)}{\sigma_{2,1}} - \frac{(b+1)}{\sigma_{2,1}} \sum_{i=1}^{n-r} \frac{e^{-a(i)}}{\left(1 + e^{-a(i)}\right)} = 0,$$

$$\frac{\partial \ln L}{\partial \sigma_{2,1}} = -\frac{(n-r)}{\sigma_{2,1}} + \frac{1}{\sigma_{2,1}} \sum_{i=1}^{n-r} a(i) - \frac{(b+1)}{\sigma_{2,1}} \sum_{i=1}^{n-r} a(i) \frac{e^{-a(i)}}{\left(1 + e^{-a(i)}\right)} = 0$$

and

$$\frac{\partial \ln L}{\partial \theta} = \frac{1}{\sigma_{2,1}} \sum_{i=1}^{n-r} x[i] - \frac{(b+1)}{\sigma_{2,1}} \sum_{i=1}^{n-r} x[i] \frac{e^{-a(i)}}{\left(1 + e^{-a(i)}\right)} = 0. \quad (2.2.2.1.8)$$

If $g(a(i))$ function is defined as $\frac{e^{-a(i)}}{\left(1 + e^{-a(i)}\right)}$, then the equations given in

(2.2.2.1.8) become

$$\frac{\partial \ln L}{\partial \mu_{2,1}} = \frac{(n-r)}{\sigma_{2,1}} - \frac{(b+1)}{\sigma_{2,1}} \sum_{i=1}^{n-r} g(a(i)) = 0,$$

$$\frac{\partial \ln L}{\partial \sigma_{2,1}} = -\frac{(n-r)}{\sigma_{2,1}} + \frac{1}{\sigma_{2,1}} \sum_{i=1}^{n-r} a(i) - \frac{(b+1)}{\sigma_{2,1}} \sum_{i=1}^{n-r} a(i) g(a(i)) = 0$$

and

$$\frac{\partial \ln L}{\partial \theta} = \frac{1}{\sigma_{2,1}} \sum_{i=1}^{n-r} x[i] - \frac{(b+1)}{\sigma_{2,1}} \sum_{i=1}^{n-r} x[i] g(a(i)) = 0. \quad (2.2.2.1.9)$$

Then, $g(a_{(i)})$ function is linearized by using the the procedure explained in Section 2.1.1.1 and

$$g(a_{(i)}) \cong \alpha_i - \beta_i a_{(i)}, \quad (1 \leq i \leq n-r) \quad (2.2.2.1.10)$$

where α_i 's and β_i 's are given in the equations (2.1.1.1.11).

After the linearization procedure the modified likelihood equations are obtained and they are

$$\begin{aligned} \frac{\partial \ln L}{\partial \mu_{2.1}} &\cong \frac{\partial \ln L^*}{\partial \mu_{2.1}} = \frac{(n-r)}{\sigma_{2.1}} - \frac{(b+1)}{\sigma_{2.1}} \sum_{i=1}^{n-r} (\alpha_i - \beta_i a_{(i)}) = 0, \\ \frac{\partial \ln L}{\partial \sigma_{2.1}} &\cong \frac{\partial \ln L^*}{\partial \sigma_{2.1}} = -\frac{(n-r)}{\sigma_{2.1}} + \frac{1}{\sigma_{2.1}} \sum_{i=1}^n a_{(i)} - \frac{(b+1)}{\sigma_{2.1}} \sum_{i=1}^{n-r} a_{(i)} (\alpha_i - \beta_i a_{(i)}) = 0 \end{aligned}$$

and

$$\frac{\partial \ln L}{\partial \theta} \cong \frac{\partial \ln L^*}{\partial \theta} = \frac{1}{\sigma_{2.1}} \sum_{i=1}^{n-r} x_{[i]} - \frac{(b+1)}{\sigma_{2.1}} \sum_{i=1}^{n-r} x_{[i]} (\alpha_i - \beta_i a_{(i)}) = 0. \quad (2.2.2.1.11)$$

Solving the equations given in (2.2.2.1.11) the following MML estimators are obtained:

$$\begin{aligned} \hat{\mu}_{2.1} &= \bar{y}[\cdot] - \hat{\theta}[\cdot] - \frac{\Delta}{m} \hat{\sigma}_{2.1}, \\ \hat{\theta} &= K - D \hat{\sigma}_{2.1} \end{aligned}$$

and

$$\hat{\sigma}_{2.1} = \left(-B + \sqrt{B^2 + 4(n-r)C} \right) / 2(n-r) \quad (2.2.2.1.12)$$

where

$$\begin{aligned}
m &= \sum_{i=1}^{n-r} \beta_i, \quad \bar{y}[\cdot] = \sum_{i=1}^{n-r} \beta_i y[i] / m, \quad \bar{x}[\cdot] = \sum_{i=1}^{n-r} \beta_i x[i] / m, \\
\Delta_i &= \left[\alpha_i - \frac{1}{b+1} \right], \quad \Delta = \sum_{i=1}^{n-r} \Delta_i, \\
K &= \frac{\sum_{i=1}^{n-r} \beta_i (x[i] - \bar{x}[\cdot]) y[i]}{\sum_{i=1}^{n-r} \beta_i (x[i] - \bar{x}[\cdot])^2}, \quad D = \frac{\sum_{i=1}^{n-r} \Delta_i (x[i] - \bar{x}[\cdot])}{\sum_{i=1}^{n-r} \beta_i (x[i] - \bar{x}[\cdot])^2}, \\
B &= (b+1) \sum_{i=1}^{n-r} \Delta_i \{ (y[i] - \bar{y}[\cdot]) - K(x[i] - \bar{x}[\cdot]) \} \\
C &= (b+1) \left\{ \sum_{i=1}^{n-r} \beta_i (y[i] - \bar{y}[\cdot])^2 - K \sum_{i=1}^{n-r} \beta_i (x[i] - \bar{x}[\cdot]) y[i] \right\}. \quad (2.2.2.1.13)
\end{aligned}$$

Also, for the MML estimators $\hat{\mu}_2$, $\hat{\sigma}_2$ and $\hat{\rho}$, see the equation (2.2.1.1.14).

2.2.2.2. Sample Information Matrix

Because of the reasons stated in Chapter 1, Fisher information matrix $I(\sigma_1, \mu_{2.1}, \sigma_{2.1}, \theta)$ cannot be obtained and instead of I , the sample information matrix $\hat{I}$ is calculated, and the components of $\hat{I}$ are given in the following equations:

$$\begin{aligned}
\hat{I} = [\hat{I}_{ij}] &= \left[-\frac{\partial^2 \ln L}{\partial \delta_i \partial \delta_j} \right] \quad i, j=1, 2, 3, 4 \\
\delta_1 &= \hat{\sigma}_1, \quad \delta_2 = \hat{\mu}_{2.1}, \quad \delta_3 = \hat{\sigma}_{2.1}, \quad \delta_4 = \hat{\theta},
\end{aligned}$$

$$\hat{I}_{11} = -\frac{(n-r)p}{\hat{\sigma}_1^2} + \frac{p}{\hat{\sigma}_1^2} \sum_{i=1}^{n-r} z_{(i)}^p + \frac{p^2}{\hat{\sigma}_1^2} \sum_{i=1}^{n-r} z_{(i)}^p + \frac{rp}{\hat{\sigma}_1^2} z_{(n-r)}^p + \frac{rp^2}{\hat{\sigma}_1^2} z_{(n-r)}^p, \quad (2.2.2.2.1)$$

$$\hat{I}_{12} = \hat{I}_{13} = \hat{I}_{14} = 0, \quad (2.2.2.2.2)$$

$$\hat{I}_{22} = \frac{(b+1)}{\hat{\sigma}_{2.1}^2} \sum_{i=1}^{n-r} \frac{e^{-a[i]}}{\left(1 + e^{-a[i]}\right)^2}, \quad (2.2.2.2.3)$$

$$\hat{I}_{23} = \frac{(n-r)}{\hat{\sigma}_{2.1}^2} - \frac{(b+1)}{\hat{\sigma}_{2.1}^2} \sum_{i=1}^{n-r} \frac{e^{-a[i]}}{\left(1 + e^{-a[i]}\right)} + \frac{(b+1)}{\hat{\sigma}_{2.1}^2} \sum_{i=1}^{n-r} a[i] \frac{e^{-a[i]}}{\left(1 + e^{-a[i]}\right)^2}, \quad (2.2.2.2.4)$$

$$\hat{I}_{24} = \frac{(b+1)}{\hat{\sigma}_{2.1}^2} \sum_{i=1}^{n-r} x_{(i)} \frac{e^{-a[i]}}{\left(1 + e^{-a[i]}\right)^2}, \quad (2.2.2.2.5)$$

$$\begin{aligned} \hat{I}_{33} = & -\frac{(n-r)}{\hat{\sigma}_{2.1}^2} + \frac{2}{\hat{\sigma}_{2.1}^2} \sum_{i=1}^{n-r} a[i] - \frac{2(b+1)}{\hat{\sigma}_{2.1}^2} \sum_{i=1}^{n-r} a[i] \frac{e^{-a[i]}}{\left(1 + e^{-a[i]}\right)} \\ & + \frac{(b+1)}{\hat{\sigma}_{2.1}^2} \sum_{i=1}^{n-r} a_{[i]}^2 \frac{e^{-a[i]}}{\left(1 + e^{-a[i]}\right)^2}, \end{aligned} \quad (2.2.2.2.6)$$

$$\begin{aligned} \hat{I}_{34} = & \frac{1}{\hat{\sigma}_{2.1}^2} \sum_{i=1}^{n-r} x_{(i)} - \frac{(b+1)}{\hat{\sigma}_{2.1}^2} \sum_{i=1}^{n-r} x_{(i)} \frac{e^{-a[i]}}{\left(1 + e^{-a[i]}\right)} + \frac{(b+1)}{\hat{\sigma}_{2.1}^2} \times \\ & \sum_{i=1}^{n-r} x_{(i)} a[i] \frac{e^{-a[i]}}{\left(1 + e^{-a[i]}\right)^2} \end{aligned} \quad (2.2.2.2.7)$$

$$\hat{I}_{44} = \frac{(b+1)}{\hat{\sigma}_{2.1}^2} \sum_{i=1}^{n-r} x_{(i)}^2 \frac{e^{-a[i]}}{\left(1 + e^{-a[i]}\right)^2}. \quad (2.2.2.2.8)$$

In $z_{(i)}$ and $a[i]$, the unknown parameters are replaced by their MML estimators.

2.3. Marginal Distribution is Weibull and Conditional Distribution is Long – Tailed Symmetric

2.3.1. Complete Sample

In section (2.3) X is coming from Weibull distribution with probability density function,

$$f(x) = \frac{p_0}{\sigma_1^{p_0}} x^{p_0-1} \exp \left\{ - \left(\frac{x}{\sigma_1} \right)^{p_0} \right\}, \quad 0 < x < \infty, \quad \sigma_1 > 0, \quad p_0 > 0, \quad (2.3.1.1)$$

while the distribution of $Y|X=x$ long-tailed symmetric.

$$h(y|x) = \frac{1}{\sigma_{2.1} \sqrt{k} \beta \left(\frac{1}{2}, p - \frac{1}{2} \right)} \left\{ 1 + \frac{(y - \mu_{2.1} - \theta x)^2}{k \sigma_{2.1}^2} \right\}^{-p}, \quad -\infty < y < \infty. \quad (2.3.1.2)$$

In equation (2.3.1.2), $k=2p-3$, $p \geq 2$. $\sigma_{2.1}$, $\mu_{2.1}$ and θ are given in equation (1.2.3). And note that the distribution of $t = \sqrt{(v/k)}(y - \mu_{2.1} - \theta x) / \sigma_{2.1}$ is Student t with $v=2p-1$ degrees of freedom.

2.3.1.1. Estimation of Parameters

For a random sample (x_i, y_i) $1 \leq i \leq n$, the likelihood and the log-likelihood functions are

$$L \propto \frac{p_0^n}{\sigma_1^{np_0}} \prod_{i=1}^n x_i^{p_0-1} e^{-\sum_{i=1}^n \left(\frac{x_i}{\sigma_1} \right)^{p_0}} \frac{1}{\sigma_{2.1}^n} \prod_{i=1}^n \left\{ 1 + \frac{(y_i - \mu_{2.1} - \theta x_i)^2}{k \sigma_{2.1}^2} \right\}^{-p} \quad (2.3.1.1.1)$$

$$\ln L \propto -np_0 \ln \sigma_1 + (p_0 - 1) \sum_{i=1}^n \ln x_i - \sum_{i=1}^n z_i^{p_0} - n \ln \sigma_{2.1} - p \sum_{i=1}^n \ln \left\{ 1 + \frac{e_i^2}{k\sigma_{2.1}^2} \right\}, \quad (2.3.1.1.2)$$

where

$$z_i = \frac{x_i}{\sigma_1} \text{ and } e_i = y_i - \mu_{2.1} - \theta x_i.$$

The likelihood equations are

$$\frac{\partial \ln L}{\partial \sigma_1} = -\frac{np_0}{\sigma_1} + \frac{p_0}{\sigma_1} \sum_{i=1}^n z_i^{p_0} = 0, \quad (2.3.1.1.3)$$

$$\frac{\partial \ln L}{\partial \mu_{2.1}} = \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^n \left(\frac{e_i / \sigma_{2.1}}{1 + e_i^2 / k\sigma_{2.1}^2} \right) = 0, \quad (2.3.1.1.4)$$

$$\frac{\partial \ln L}{\partial \sigma_{2.1}} = -\frac{n}{\sigma_{2.1}} + \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^n \left(\frac{e_i^2 / \sigma_{2.1}^2}{1 + e_i^2 / k\sigma_{2.1}^2} \right) = 0 \quad (2.3.1.1.5)$$

and

$$\frac{\partial \ln L}{\partial \theta} = \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^n \left(\frac{x_i e_i / \sigma_{2.1}}{1 + e_i^2 / k\sigma_{2.1}^2} \right) = 0. \quad (2.3.1.1.6)$$

$\hat{\sigma}_1$ is obtained from equation (2.3.1.1.3) and is equal to

$$\hat{\sigma}_1 = \left\{ \frac{1}{n} \sum_{i=1}^n x_i^{p_0} \right\}^{1/p_0}. \quad (2.3.1.1.7)$$

However, because of the function $\left(\frac{e_i/\sigma_{2.1}}{1+e_i^2/k\sigma_{2.1}^2}\right)$, the equations (2.3.1.1.4)

– (2.3.1.1.6) are difficult to solve. In order to obtain the corresponding MML estimators, first define w_i and a_i , as in section (2.2.1.1). Then, rewrite the likelihood equations as

$$\begin{aligned}\frac{\partial \ln L}{\partial \mu_{2.1}} &= \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^n \frac{a(i)}{1 + \frac{a(i)^2}{k}} = 0, \\ \frac{\partial \ln L}{\partial \sigma_{2.1}} &= -\frac{n}{\sigma_{2.1}} + \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^n \frac{a(i)^2}{1 + \frac{a(i)^2}{k}} = 0 \\ \frac{\partial \ln L}{\partial \theta} &= \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^n x[i] \frac{a(i)}{1 + \frac{a(i)^2}{k}} = 0.\end{aligned}\tag{2.3.1.1.8}$$

If $a/\left(1 + \frac{a^2}{k}\right)$ is defined as the function $g(a)$, then the likelihood

equations become

$$\begin{aligned}\frac{\partial \ln L}{\partial \mu_{2.1}} &= \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^n g(a(i)) = 0, \\ \frac{\partial \ln L}{\partial \sigma_{2.1}} &= -\frac{n}{\sigma_{2.1}} + \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^n a(i)g(a(i)) = 0 \\ \frac{\partial \ln L}{\partial \theta} &= \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^n x[i]g(a(i)) = 0.\end{aligned}\tag{2.3.1.1.9}$$

Linearized $g(a(i))$ function is given in the following equation where

$$t(i) = E(a(i)):$$

$$\begin{aligned}
g(a_{(i)}) &\cong g(t_{(i)}) + (a_{(i)} - t_{(i)}) \left. \frac{dg(a_{(i)})}{da_{(i)}} \right|_{a_{(i)}=t_{(i)}} \\
&= \alpha_i + \beta_i a_{(i)} \quad (1 \leq i \leq n)
\end{aligned} \tag{2.3.1.1.10}$$

where

$$\alpha_i = \frac{2 \frac{t_{(i)}^3}{k}}{\left\{ 1 + \frac{t_{(i)}^2}{k} \right\}^2} \quad \text{and} \quad \beta_i = \frac{1 - \frac{t_{(i)}^2}{k}}{\left\{ 1 + \frac{t_{(i)}^2}{k} \right\}^2}.$$

Note that $t_{(i)}$ values can be conveniently obtained by using an IMSL subroutine. However, note that when $\beta_1 < 0$, which can be observed for small p and large n values, instead of α_i and β_i , α_i^* and β_i^* are used. Because β_i 's should be greater than 0 for $\hat{\sigma}$ to be positive and real. For α_i^* and β_i^* , see Islam and Tiku (2004).

The modified likelihood equations which are obtained after the linearization procedure are:

$$\begin{aligned}
\frac{\partial \ln L}{\partial \mu_{2.1}} &\cong \frac{\partial \ln L^*}{\partial \mu_{2.1}} = \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^n \{\alpha_i + \beta_i a_{(i)}\} = 0, \\
\frac{\partial \ln L}{\partial \sigma_{2.1}} &\cong \frac{\partial \ln L^*}{\partial \sigma_{2.1}} = -\frac{n}{\sigma_{2.1}} + \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^n a_{(i)} \{\alpha_i + \beta_i a_{(i)}\} = 0
\end{aligned}$$

and

$$\frac{\partial \ln L}{\partial \theta} \cong \frac{\partial \ln L^*}{\partial \theta} = \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^n x[i] \{ \alpha_i + \beta_i a_{(i)} \} = 0. \quad (2.3.1.1.11)$$

The MML estimators which are the solutions of the equations given in (2.3.1.1.11) are,

$$\begin{aligned} \hat{\mu}_{2.1} &= \bar{y}[\cdot] - \hat{\theta} \bar{x}[\cdot], \\ \hat{\theta} &= K + D \hat{\sigma}_{2.1} \end{aligned}$$

and

$$\hat{\sigma}_{2.1} = \left\{ B + \sqrt{B^2 + 4nC} \right\} / 2n,$$

where

$$\begin{aligned} \bar{y}[\cdot] &= \frac{\sum_{i=1}^n \beta_i y[i]}{m}, \quad \bar{x}[\cdot] = \frac{\sum_{i=1}^n \beta_i x[i]}{m}, \quad m = \sum_{i=1}^n \beta_i, \\ K &= \frac{\sum_{i=1}^n \beta_i (x[i] - \bar{x}[\cdot]) y[i]}{\sum_{i=1}^n \beta_i (x[i] - \bar{x}[\cdot])^2}, \quad D = \frac{\sum_{i=1}^n \alpha_i x[i]}{\sum_{i=1}^n \beta_i (x[i] - \bar{x}[\cdot])^2}, \\ B &= \frac{2p}{k} \sum_{i=1}^n \alpha_i \{ (y[i] - \bar{y}[\cdot]) - K(x[i] - \bar{x}[\cdot]) \} \end{aligned}$$

and

$$C = \frac{2p}{k} \left\{ \sum_{i=1}^n \beta_i (y[i] - \bar{y}[\cdot])^2 - K \sum_{i=1}^n \beta_i (x[i] - \bar{x}[\cdot]) y[i] \right\}. \quad (2.3.1.1.12)$$

Also, $\hat{\mu}_2$, $\hat{\sigma}_2$ and $\hat{\rho}$ are obtained as

$$\hat{\mu}_2 = \hat{\mu}_{2.1} + \hat{\theta}\hat{\sigma}_1\Gamma\left(1 + \frac{1}{p_0}\right), \quad \hat{\sigma}_2 = \sqrt{\hat{\sigma}_{2.1}^2 + \hat{\theta}^2\hat{\sigma}_1^2}, \quad \hat{\rho} = \hat{\theta}\frac{\hat{\sigma}_1}{\hat{\sigma}_2} \quad (2.3.1.1.13)$$

2.3.1.2. Fisher Information Matrix

The Fisher information matrix $I(\sigma_1, \mu_{2.1}, \sigma_{2.1}, \theta)$ has components which are given in (2.3.1.2.1) – (2.3.1.2.8),

$$I_{11} = \frac{np_0^2}{\sigma_1^2}, \quad (2.3.1.2.1)$$

$$I_{12} = I_{13} = I_{14} = 0, \quad (2.3.1.2.2)$$

$$I_{22} = \frac{p(p-1/2)}{(p+1)(p-3/2)} \frac{n}{\sigma_{2.1}^2}, \quad (2.3.1.2.3)$$

$$I_{23} = 0, \quad (2.3.1.2.4)$$

$$I_{24} = \frac{p(p-1/2)}{(p+1)(p-3/2)} \Gamma\left(1 + \frac{1}{p_0}\right) \frac{n}{\sigma_{2.1}^2} \sigma_1, \quad (2.3.1.2.5)$$

$$I_{33} = \frac{2(p-1/2)}{(p+1)} \frac{n}{\sigma_{2.1}^2}, \quad (2.3.1.2.6)$$

$$I_{34} = 0 \quad (2.3.1.2.7)$$

$$I_{44} = \frac{p(p-1/2)}{(p+1)(p-3/2)} \Gamma\left(1 + \frac{2}{p_0}\right) \frac{n}{\sigma_{2.1}^2} \sigma_1^2. \quad (2.3.1.2.8)$$

2.3.2. Censored Sample

In this section, the situation where the marginal distribution is ‘censored’ Weibull and the conditional distribution is long-tailed symmetric is taken into consideration.

2.3.2.1. Estimation of Parameters

The likelihood function for a Type II censored sample $(x_{(i)}, y_{[i]})$ ($1 \leq i \leq n-r$) is

$$L \propto \left[\prod_{i=1}^{n-r} \frac{p_0}{\sigma_1^{p_0}} x_{(i)}^{p_0-1} e^{-\left(\frac{x_{(i)}}{\sigma_1}\right)^{p_0}} \right] \left[e^{-\left(\frac{x_{(n-r)}}{\sigma_1}\right)^{p_0}} \right]^r \times \frac{1}{\sigma_{2.1}^{n-r}} \prod_{i=1}^{n-r} \left[1 + \frac{(y_{[i]} - \mu_{2.1} - \theta x_{(i)})^2}{k \sigma_{2.1}^2} \right]^{-p} \quad (2.3.2.1.1)$$

where $y_{[i]}$'s are the concomitants of $x_{(i)}$'s. The log likelihood function is

$$\ln L \propto -(n-r)p_0 \ln \sigma_1 + (p_0-1) \sum_{i=1}^{n-r} \ln x_{(i)} - \sum_{i=1}^{n-r} z_{(i)}^{p_0} - r z_{(n-r)}^{p_0} - (n-r) \ln \sigma_{2.1} - p \sum_{i=1}^{n-r} \ln \left(1 + \frac{e_{[i]}^2}{k \sigma_{2.1}^2} \right) \quad (2.3.2.1.2)$$

where

$$z_{(i)} = \frac{x_{(i)}}{\sigma_1} \text{ and } e_{[i]} = y_{[i]} - \mu_{2.1} - \theta x_{(i)}.$$

Here, the likelihood equations are

$$\frac{\partial \ln L}{\partial \sigma_1} = -\frac{(n-r)p_0}{\sigma_1} + \frac{p_0}{\sigma_1} \sum_{i=1}^{n-r} z_{(i)}^{p_0} + \frac{r p_0}{\sigma_1} z_{(n-r)}^{p_0} = 0, \quad (2.3.2.1.3)$$

$$\frac{\partial \ln L}{\partial \mu_{2.1}} = \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^{n-r} \left(\frac{e_{[i]}/\sigma_{2.1}}{1 + e_{[i]}^2/k\sigma_{2.1}^2} \right) = 0, \quad (2.3.2.1.4)$$

$$\frac{\partial \ln L}{\partial \sigma_{2.1}} = -\frac{(n-r)}{\sigma_{2.1}} + \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^{n-r} \left(\frac{e_{[i]}^2/\sigma_{2.1}^2}{1 + e_{[i]}^2/k\sigma_{2.1}^2} \right) = 0 \quad (2.3.2.1.5)$$

and

$$\frac{\partial \ln L}{\partial \theta} = \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^{n-r} \frac{x_{(i)}e_{[i]}/\sigma_{2.1}}{1 + e_{[i]}^2/k\sigma_{2.1}^2} = 0. \quad (2.3.2.1.6)$$

The solution of the equation (2.3.2.1.3) is $\hat{\sigma}_1$ and is equal to

$$\hat{\sigma}_1 = \left\{ \frac{\sum_{i=1}^{n-r} x_{(i)}^{p_0} + rx_{(n-r)}^{p_0}}{n-r} \right\}^{1/p_0}. \quad (2.3.2.1.7)$$

However, because of the function $\frac{e_{[i]}/\sigma_{2.1}}{1 + e_{[i]}^2/k\sigma_{2.1}^2}$, the equations

(2.3.2.1.4) – (2.3.2.1.6) are difficult to solve. Therefore, the estimators are obtained by using the MML methodology; a_i terms and the corresponding $(y_{[i]}, x_{[i]})$'s are obtained by the procedure explained in section (2.2.2.1). Then, the likelihood equations are written as

$$\frac{\partial \ln L}{\partial \mu_{2.1}} = \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^{n-r} \frac{a(i)}{1 + \frac{a(i)^2}{k}} = 0,$$

$$\begin{aligned}\frac{\partial \ln L}{\partial \sigma_{2.1}} &= -\frac{(n-r)}{\sigma_{2.1}} + \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^{n-r} \frac{a_{(i)}^2}{1 + \frac{a_{(i)}^2}{k}} = 0 \\ \frac{\partial \ln L}{\partial \theta} &= \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^{n-r} x_{[i]} \frac{a_{(i)}}{1 + \frac{a_{(i)}^2}{k}} = 0.\end{aligned}\tag{2.3.2.1.8}$$

If in (2.3.2.1.8), $a_{(i)} / \left(1 + \frac{a_{(i)}^2}{k}\right)$ is denoted by $g(a_{(i)})$, the likelihood

equations become:

$$\begin{aligned}\frac{\partial \ln L}{\partial \mu_{2.1}} &= \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^{n-r} g(a_{(i)}) = 0, \\ \frac{\partial \ln L}{\partial \sigma_{2.1}} &= -\frac{(n-r)}{\sigma_{2.1}} + \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^{n-r} a_{(i)} g(a_{(i)}) = 0\end{aligned}$$

and

$$\frac{\partial \ln L}{\partial \theta} = \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^{n-r} x_{[i]} g(a_{(i)}) = 0.\tag{2.3.2.1.9}$$

After linearizing the $g(a_{(i)})$ functions, the modified likelihood equations are written as in equations (2.3.2.1.10). In order to see the linearized $g(a_{(i)})$ function and the corresponding (α_i, β_i) terms, see equation (2.3.1.1.10).

$$\begin{aligned}\frac{\partial \ln L}{\partial \mu_{2.1}} &\cong \frac{\partial \ln L^*}{\partial \mu_{2.1}} = \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^{n-r} \{\alpha_i + \beta_i a_{(i)}\} = 0, \\ \frac{\partial \ln L}{\partial \sigma_{2.1}} &\cong \frac{\partial \ln L^*}{\partial \sigma_{2.1}} = -\frac{(n-r)}{\sigma_{2.1}} + \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^{n-r} a_{(i)} \{\alpha_i + \beta_i a_{(i)}\} = 0\end{aligned}$$

and

$$\frac{\partial \ln L}{\partial \theta} \cong \frac{\partial \ln L^*}{\partial \theta} = \frac{2p}{k\sigma_{2.1}} \sum_{i=1}^{n-r} x[i] \{\alpha_i + \beta_i a(i)\} = 0. \quad (2.3.2.1.10)$$

The solutions of the equations displayed in (2.3.2.1.10) are the following MML estimators:

$$\hat{\mu}_{2.1} = \bar{y}[\cdot] - \hat{\theta} \bar{x}[\cdot] - \hat{\sigma}_{2.1} \sum_{i=1}^{n-r} \frac{\alpha_i}{m},$$

$$\hat{\theta} = K + D \hat{\sigma}_{2.1}$$

and

$$\hat{\sigma}_{2.1} = \left\{ B + \sqrt{B^2 + 4(n-r)C} \right\} / 2(n-r)$$

where

$$\bar{y}[\cdot] = \frac{\sum_{i=1}^{n-r} \beta_i y[i]}{m}, \quad \bar{x}[\cdot] = \frac{\sum_{i=1}^{n-r} \beta_i x[i]}{m}, \quad m = \sum_{i=1}^{n-r} \beta_i,$$

$$K = \frac{\sum_{i=1}^{n-r} \beta_i (x[i] - \bar{x}[\cdot]) y[i]}{\sum_{i=1}^{n-r} \beta_i (x[i] - \bar{x}[\cdot])^2}, \quad D = \frac{\sum_{i=1}^{n-r} \alpha_i (x[i] - \bar{x}[\cdot])}{\sum_{i=1}^{n-r} \beta_i (x[i] - \bar{x}[\cdot])^2},$$

$$B = \frac{2p}{k} \sum_{i=1}^{n-r} \alpha_i \{ (y[i] - \bar{y}[\cdot]) - K(x[i] - \bar{x}[\cdot]) \}$$

and

$$C = \frac{2p}{k} \left\{ \sum_{i=1}^{n-r} \beta_i (y[i] - \bar{y}[\cdot])^2 - K \sum_{i=1}^{n-r} \beta_i (x[i] - \bar{x}[\cdot]) y[i] \right\}. \quad (2.3.2.1.11)$$

Also, for the MML estimators of μ_2 , σ_2 and ρ , see equation (2.3.1.1.13).

2.3.2.2. Sample Information Matrix

The components of the sample information matrix $\hat{I}$ are given in the equations (2.3.2.2.1) – (2.3.2.2.8).

$$\hat{I} = [\hat{I}_{ij}] = \left[-\frac{\partial^2 \ln L}{\partial \delta_i \partial \delta_j} \right] \quad i, j = 1, 2, 3, 4$$

$$\delta_1 = \hat{\sigma}_1, \quad \delta_2 = \hat{\mu}_{2.1}, \quad \delta_3 = \hat{\sigma}_{2.1}, \quad \delta_4 = \hat{\theta},$$

$$\hat{I}_{11} = -\frac{(n-r)p_0}{\hat{\sigma}_1^2} + \frac{p_0}{\hat{\sigma}_1^2} \sum_{i=1}^{n-r} z_{(i)}^{p_0} + \frac{p_0^2}{\hat{\sigma}_1^2} \sum_{i=1}^{n-r} z_{(i)}^{p_0} + \frac{rp_0}{\hat{\sigma}_1^2} z_{(n-r)}^{p_0} + \frac{rp_0^2}{\hat{\sigma}_1^2} z_{(n-r)}^{p_0} \quad (2.3.2.2.1)$$

$$\hat{I}_{12} = \hat{I}_{13} = \hat{I}_{14} = 0, \quad (2.3.2.2.2)$$

$$\hat{I}_{22} = \frac{2p}{k\hat{\sigma}_{2.1}^2} \sum_{i=1}^{n-r} \frac{\left(1 - a_{[i]}^2/k\right)}{\left(1 + a_{[i]}^2/k\right)^2}, \quad (2.3.2.2.3)$$

$$\hat{I}_{23} = \frac{2p}{k\hat{\sigma}_{2.1}^2} \left\{ \sum_{i=1}^{n-r} \frac{a[i]}{\left(1 + a_{[i]}^2/k\right)} + \sum_{i=1}^{n-r} a[i] \frac{\left(1 - a_{[i]}^2/k\right)}{\left(1 + a_{[i]}^2/k\right)^2} \right\}, \quad (2.3.2.2.4)$$

$$\hat{I}_{24} = \frac{2p}{k\hat{\sigma}_{2.1}^2} \sum_{i=1}^{n-r} x_{(i)} \frac{\left(1 - a_{[i]}^2/k\right)}{\left(1 + a_{[i]}^2/k\right)^2}, \quad (2.3.2.2.5)$$

$$\hat{I}_{33} = -\frac{n-r}{\hat{\sigma}_{2.1}^2} + \frac{2p}{k\hat{\sigma}_{2.1}^2} \left\{ \sum_{i=1}^{n-r} \frac{a_{[i]}^2}{\left(1 + a_{[i]}^2/k\right)} + 2 \sum_{i=1}^{n-r} \frac{a_{[i]}^2}{\left(1 + a_{[i]}^2/k\right)^2} \right\}, \quad (2.3.2.2.6)$$

$$\hat{I}_{34} = \frac{4p}{k\hat{\sigma}_{2.1}^2} \sum_{i=1}^{n-r} x_{(i)} \frac{a_{[i]}}{\left(1 + a_{[i]}^2/k\right)^2} \quad (2.3.2.2.7)$$

$$\hat{I}_{44} = \frac{2p}{k\hat{\sigma}_{2.1}^2} \sum_{i=1}^{n-r} x_{(i)}^2 \frac{\left(1 - a_{[i]}^2/k\right)}{\left(1 + a_{[i]}^2/k\right)^2}. \quad (2.3.2.2.8)$$

In $z_{(i)}$ and $a_{[i]}$, the unknown parameters are replaced by their MML estimators.

2.4. Marginal and Conditional Distribution is Generalized Logistic

2.4.1. Complete Sample

In section (2.4) both X and $Y|X=x$ are coming from Generalized Logistic distribution. The probability density function of X is

$$f(x) = \frac{b_1}{\sigma_1} \frac{\exp\left[-\left(\frac{x - \mu_1}{\sigma_1}\right)\right]}{\left\{1 + \exp\left[-\left(\frac{x - \mu_1}{\sigma_1}\right)\right]\right\}^{b_1+1}}, \quad -\infty < x < \infty, \quad b_1 > 0, \quad \sigma_1 > 0 \quad (2.4.1.1)$$

and the probability density function of $Y|X=x$ is

$$h(y | x) = \frac{b_2}{\sigma_{2.1}} \frac{\exp \left[- \left(\frac{y - \mu_{2.1} - \theta x}{\sigma_{2.1}} \right) \right]}{\left\{ 1 + \exp \left[- \left(\frac{y - \mu_{2.1} - \theta x}{\sigma_{2.1}} \right) \right] \right\}^{b_2 + 1}}, \quad -\infty < y < \infty, \quad b_2 > 0$$

(2.4.1.2)

where $\sigma_{2.1}$, $\mu_{2.1}$ and θ are given in equation (1.2.3).

2.4.1.1. Estimation of Parameters

For a random sample (x_i, y_i) ($1 \leq i \leq n$), the likelihood function is

$$L = \left(\frac{b_1}{\sigma_1} \right)^n \frac{e^{-\sum_{i=1}^n \left(\frac{x_i - \mu_1}{\sigma_1} \right)}}{\prod_{i=1}^n \left[1 + e^{-\left(\frac{x_i - \mu_1}{\sigma_1} \right)} \right]^{b_1 + 1}} \times \left(\frac{b_2}{\sigma_{2.1}} \right)^n \frac{e^{-\frac{1}{\sigma_{2.1}} \sum_{i=1}^n (y_i - \mu_{2.1} - \theta x_i)}}{\prod_{i=1}^n \left[1 + e^{-\left(\frac{y_i - \mu_{2.1} - \theta x_i}{\sigma_{2.1}} \right)} \right]^{b_2 + 1}}$$

(2.4.1.1.1)

and the loglikelihood function is,

$$\begin{aligned} \ln L = & n \ln b_1 - n \ln \sigma_1 - \sum_{i=1}^n z_i - (b_1 + 1) \sum_{i=1}^n \ln \left(1 + e^{-z_i} \right) + n \ln b_2 - n \ln \sigma_{2.1} \\ & - \frac{1}{\sigma_{2.1}} \sum_{i=1}^n e_i - (b_2 + 1) \sum_{i=1}^n \ln \left(1 + e^{-e_i / \sigma_{2.1}} \right) \end{aligned} \quad (2.4.1.1.2)$$

where

$$z_i = \frac{x_i - \mu_1}{\sigma_1} \quad \text{and} \quad e_i = y_i - \mu_{2.1} - \theta x_i.$$

The likelihood equations are

$$\frac{\partial \ln L}{\partial \mu_1} = \frac{n}{\sigma_1} - \frac{(b_1 + 1)}{\sigma_1} \sum_{i=1}^n \left(\frac{e^{-z_i}}{1 + e^{-z_i}} \right) = 0, \quad (2.4.1.1.3)$$

$$\frac{\partial \ln L}{\partial \sigma_1} = -\frac{n}{\sigma_1} + \frac{1}{\sigma_1} \sum_{i=1}^n z_i - \frac{(b_1 + 1)}{\sigma_1} \sum_{i=1}^n z_i \left(\frac{e^{-z_i}}{1 + e^{-z_i}} \right) = 0, \quad (2.4.1.1.4)$$

$$\frac{\partial \ln L}{\partial \mu_{2.1}} = \frac{n}{\sigma_{2.1}} - \frac{(b_2 + 1)}{\sigma_{2.1}} \sum_{i=1}^n \left[\frac{e^{-e_i/\sigma_{2.1}}}{1 + e^{-e_i/\sigma_{2.1}}} \right] = 0, \quad (2.4.1.1.5)$$

$$\frac{\partial \ln L}{\partial \sigma_{2.1}} = -\frac{n}{\sigma_{2.1}} + \frac{1}{\sigma_{2.1}^2} \sum_{i=1}^n e_i - \frac{(b_2 + 1)}{\sigma_{2.1}^2} \sum_{i=1}^n e_i \left(\frac{e^{-e_i/\sigma_{2.1}}}{1 + e^{-e_i/\sigma_{2.1}}} \right) = 0 \quad (2.4.1.1.6)$$

and

$$\frac{\partial \ln L}{\partial \theta} = \frac{1}{\sigma_{2.1}} \sum_{i=1}^n x_i - \frac{(b_2 + 1)}{\sigma_{2.1}} \sum_{i=1}^n x_i \left(\frac{e^{-e_i/\sigma_{2.1}}}{1 + e^{-e_i/\sigma_{2.1}}} \right) = 0. \quad (2.4.1.1.7)$$

However, the equations (2.4.1.1.3) - (2.4.1.1.4) and (2.4.1.1.5) - (2.4.1.1.7)

are difficult to solve because of the functions $\frac{e^{-z_i}}{1 + e^{-z_i}}$ and $\frac{e^{-\frac{e_i}{\sigma_{2.1}}}}{1 + e^{-\frac{e_i}{\sigma_{2.1}}}}$,

respectively. Therefore, MML methodology is used.

In order to obtain the corresponding estimators, first w_i 's are defined as $w_i = y_i - \theta x_i$ (for a given θ , $1 \leq i \leq n$). Then, both $x_{(i)}$ and $w_{(i)}$ values are ordered in ascending order. Note that since $\mu_{2.1}$ is a constant, $e_{(i)} = w_{(i)} - \mu_{2.1}$ and $w_{(i)}$ have the same order. Similarly, since μ_1 is a constant and σ_1 is positive, $x_{(i)}$ and $z_{(i)}$ have the same order. Also, if a_i is defined as $e_i/\sigma_{2.1}$, then

$$a(i) = \frac{e(i)}{\sigma_{2.1}} = \frac{w(i) - \mu_{2.1}}{\sigma_{2.1}} = \frac{y[i] - \theta x[i] - \mu_{2.1}}{\sigma_{2.1}}.$$

Now, the likelihood equations can be rewritten as

$$\frac{\partial \ln L}{\partial \mu_1} = \frac{n}{\sigma_1} - \frac{(b_1 + 1)}{\sigma_1} \sum_{i=1}^n \frac{e^{-z(i)}}{\left(1 + e^{-z(i)}\right)} = 0,$$

$$\frac{\partial \ln L}{\partial \sigma_1} = -\frac{n}{\sigma_1} + \frac{1}{\sigma_1} \sum_{i=1}^n z(i) - \frac{(b_1 + 1)}{\sigma_1} \sum_{i=1}^n z(i) \frac{e^{-z(i)}}{\left(1 + e^{-z(i)}\right)} = 0,$$

$$\frac{\partial \ln L}{\partial \mu_{2.1}} = \frac{n}{\sigma_{2.1}} - \frac{(b_2 + 1)}{\sigma_{2.1}} \sum_{i=1}^n \frac{e^{-a(i)}}{\left(1 + e^{-a(i)}\right)} = 0,$$

$$\frac{\partial \ln L}{\partial \sigma_{2.1}} = -\frac{n}{\sigma_{2.1}} + \frac{1}{\sigma_{2.1}^2} \sum_{i=1}^n e(i) - \frac{(b_2 + 1)}{\sigma_{2.1}^2} \sum_{i=1}^n e(i) \frac{e^{-a(i)}}{\left(1 + e^{-a(i)}\right)} = 0$$

and

$$\frac{\partial \ln L}{\partial \theta} = \frac{1}{\sigma_{2.1}} \sum_{i=1}^n x[i] - \frac{(b_2 + 1)}{\sigma_{2.1}} \sum_{i=1}^n x[i] \frac{e^{-a(i)}}{\left(1 + e^{-a(i)}\right)} = 0. \quad (2.4.1.1.8)$$

If $\frac{e^{-z(i)}}{1 + e^{-z(i)}}$ and $\frac{e^{-a(i)}}{\left(1 + e^{-a(i)}\right)}$ are replaced by $g_1(z(i))$ and $g_2(a(i))$, the

likelihood functions become

$$\frac{\partial \ln L}{\partial \mu_1} = \frac{n}{\sigma_1} - \frac{(b_1 + 1)}{\sigma_1} \sum_{i=1}^n g_1(z(i)) = 0,$$

$$\begin{aligned}\frac{\partial \ln L}{\partial \sigma_1} &= -\frac{n}{\sigma_1} + \frac{1}{\sigma_1} \sum_{i=1}^n z(i) - \frac{(b_1+1)}{\sigma_1} \sum_{i=1}^n z(i) g_1(z(i)) = 0, \\ \frac{\partial \ln L}{\partial \mu_{2,1}} &= \frac{n}{\sigma_{2,1}} - \frac{(b_2+1)}{\sigma_{2,1}} \sum_{i=1}^n g_2(a(i)) = 0, \\ \frac{\partial \ln L}{\partial \sigma_{2,1}} &= -\frac{n}{\sigma_{2,1}} + \frac{1}{\sigma_{2,1}^2} \sum_{i=1}^n e(i) - \frac{(b_2+1)}{\sigma_{2,1}^2} \sum_{i=1}^n e(i) g_2(a(i)) = 0\end{aligned}$$

and

$$\frac{\partial \ln L}{\partial \theta} = \frac{1}{\sigma_{2,1}} \sum_{i=1}^n x[i] - \frac{(b_2+1)}{\sigma_{2,1}} \sum_{i=1}^n x[i] g_2(a(i)) = 0. \quad (2.4.1.1.9)$$

The next step is linearizing the functions $g_1(z(i))$ around $E(z(i)) = t_{1(i)}$ and $g_2(a(i))$ around $E(a(i)) = t_{2(i)}$:

$$\begin{aligned}g_1(z(i)) &\cong g_1(t_{1(i)}) + (z(i) - t_{1(i)}) \left. \frac{dg_1(z(i))}{dz(i)} \right|_{z(i)=t_{1(i)}} \\ &= \alpha_{1i} - \beta_{1i} z(i) \quad (1 \leq i \leq n)\end{aligned} \quad (2.4.1.1.10)$$

where

$$\alpha_{1i} = \frac{e^{-t_{1(i)}}}{\left(1 + e^{-t_{1(i)}}\right)^2} \left(1 + t_{1(i)} + e^{-t_{1(i)}}\right), \quad \beta_{1i} = \frac{e^{-t_{1(i)}}}{\left(1 + e^{-t_{1(i)}}\right)^2}$$

and

$$t_{1(i)} = -\ln \left[q_i^{-1/b_1} - 1 \right], \quad q_i = \frac{i}{n+1}. \quad (2.4.1.1.11)$$

Similarly,

$$\begin{aligned}
g_2(a_{(i)}) &\cong g_2(t_{2(i)}) + (a_{(i)} - t_{2(i)}) \frac{dg_2(a_{(i)})}{da_{(i)}} \Big|_{a_{(i)}=t_{2(i)}} \\
&= \alpha_{2i} - \beta_{2i}a_{(i)} \quad (1 \leq i \leq n)
\end{aligned} \tag{2.4.1.1.12}$$

where

$$\alpha_{2i} = \frac{e^{-t_{2(i)}}}{\left(1 + e^{-t_{2(i)}}\right)^2} \left(1 + t_{2(i)} + e^{-t_{2(i)}}\right), \quad \beta_{2i} = \frac{e^{-t_{2(i)}}}{\left(1 + e^{-t_{2(i)}}\right)^2},$$

and

$$t_{2(i)} = -\ln \left[q_i^{-1/b_2} - 1 \right], \quad q_i = \frac{i}{n+1}. \tag{2.4.1.1.13}$$

Thus, the modified maximum likelihood equations are obtained as

$$\begin{aligned}
\frac{\partial \ln L}{\partial \mu_1} &\cong \frac{\partial \ln L^*}{\partial \mu_1} = \frac{n}{\sigma_1} - \frac{(b_1+1)}{\sigma_1} \sum_{i=1}^n (\alpha_{1i} - \beta_{1i}z_{(i)}) = 0, \\
\frac{\partial \ln L}{\partial \sigma_1} &\cong \frac{\partial \ln L^*}{\partial \sigma_1} = -\frac{n}{\sigma_1} + \frac{1}{\sigma_1} \sum_{i=1}^n z_{(i)} - \frac{(b_1+1)}{\sigma_1} \sum_{i=1}^n z_{(i)} (\alpha_{1i} - \beta_{1i}z_{(i)}) = 0, \\
\frac{\partial \ln L}{\partial \mu_{2.1}} &\cong \frac{\partial \ln L^*}{\partial \mu_{2.1}} = \frac{n}{\sigma_{2.1}} - \frac{(b_2+1)}{\sigma_{2.1}} \sum_{i=1}^n (\alpha_{2i} - \beta_{2i}a_{(i)}) = 0, \\
\frac{\partial \ln L}{\partial \sigma_{2.1}} &\cong \frac{\partial \ln L^*}{\partial \sigma_{2.1}} = -\frac{n}{\sigma_{2.1}} + \frac{1}{\sigma_{2.1}^2} \sum_{i=1}^n e_{(i)} - \frac{(b_2+1)}{\sigma_{2.1}^2} \sum_{i=1}^n e_{(i)} (\alpha_{2i} - \beta_{2i}a_{(i)}) = 0 \\
\frac{\partial \ln L}{\partial \theta} &\cong \frac{\partial \ln L^*}{\partial \theta} = \frac{1}{\sigma_{2.1}} \sum_{i=1}^n x_{[i]} - \frac{(b_2+1)}{\sigma_{2.1}} \sum_{i=1}^n x_{[i]} (\alpha_{2i} - \beta_{2i}a_{(i)}) = 0.
\end{aligned} \tag{2.4.1.1.14}$$

The MML estimators which are the solutions of the equations (2.4.1.1.14) are given below,

$$\begin{aligned}\hat{\mu}_1 &= K_1 + D_1 \hat{\sigma}_1, \\ \hat{\sigma}_1 &= \frac{B_1 + \sqrt{B_1^2 + 4nC_1}}{2n}, \\ \hat{\mu}_{2.1} &= \bar{y}[\cdot] - \hat{\theta}[\cdot] - \frac{\Delta}{m_2} \hat{\sigma}_{2.1}, \\ \hat{\sigma}_{2.1} &= \frac{-B_2 + \sqrt{B_2^2 + 4nC_2}}{2n}\end{aligned}$$

and

$$\hat{\theta} = K_2 - D_2 \hat{\sigma}_{2.1}, \quad (2.4.1.1.15)$$

where

$$\begin{aligned}K_1 &= \frac{1}{m_1} \sum_{i=1}^n \beta_{1i} x_{(i)}, \quad D_1 = \frac{1}{m_1} \sum_{i=1}^n \left[\frac{1}{(b_1 + 1)} - \alpha_{1i} \right], \quad m_1 = \sum_{i=1}^n \beta_{1i}, \\ B_1 &= (b_1 + 1) \sum_{i=1}^n (x_{(i)} - K_1) \left[\frac{1}{(b_1 + 1)} - \alpha_{1i} \right], \quad C_1 = (b_1 + 1) \sum_{i=1}^n \beta_{1i} (x_{(i)} - K_1)^2, \\ B_2 &= (b_2 + 1) \sum_{i=1}^n \Delta_i \{y[i] - \bar{y}[\cdot] - K_2(x[i] - \bar{x}[\cdot])\}, \\ C_2 &= (b_2 + 1) \left\{ \sum_{i=1}^n \beta_{2i} (y[i] - \bar{y}[\cdot])^2 - K_2 \sum_{i=1}^n \beta_{2i} (x[i] - \bar{x}[\cdot]) y[i] \right\}, \\ K_2 &= \frac{\sum_{i=1}^n \beta_{2i} (x[i] - \bar{x}[\cdot]) y[i]}{\sum_{i=1}^n \beta_{2i} (x[i] - \bar{x}[\cdot])^2} \quad \text{and} \quad D_2 = \frac{\sum_{i=1}^n \Delta_i (x[i] - \bar{x}[\cdot])}{\sum_{i=1}^n \beta_{2i} (x[i] - \bar{x}[\cdot])^2},\end{aligned}$$

$$\bar{x}[\cdot] = \frac{1}{m_2} \sum_{i=1}^n \beta_{2i} x[i], \quad \bar{y}[\cdot] = \frac{1}{m_2} \sum_{i=1}^n \beta_{2i} y[i], \quad \Delta_i = \alpha_{2i} - \frac{1}{(b_2 + 1)},$$

$$\Delta = \sum_{i=1}^n \Delta_i \quad \text{and} \quad m_2 = \sum_{i=1}^n \beta_{2i}. \quad (2.4.1.1.16)$$

Also, $\hat{\mu}_2$, $\hat{\sigma}_2$ and $\hat{\rho}$ are obtained as

$$\hat{\mu}_2 = \hat{\mu}_{2.1} + \hat{\theta} \{ \hat{\mu}_1 + \hat{\sigma}_1 [\psi(b_1) - \psi(1)] \}, \quad \hat{\sigma}_2 = \sqrt{\hat{\sigma}_{2.1}^2 + \hat{\theta}^2 \hat{\sigma}_1^2} \quad \text{and} \quad \hat{\rho} = \hat{\theta} \frac{\hat{\sigma}_1}{\hat{\sigma}_2}. \quad (2.4.1.1.17)$$

The iteration procedure in order to compute the MML estimators are explained in Sazak (2003).

2.4.1.2. Fisher Information Matrix

The components of the $I(\mu_1, \sigma_1, \mu_{2.1}, \sigma_{2.1}, \theta)$ are given below:

$$I_{11} = \frac{n}{\sigma_1^2} \left(\frac{b_1}{b_1 + 2} \right), \quad (2.4.1.2.1)$$

$$I_{12} = \frac{n}{\sigma_1^2} \left(\frac{b_1}{b_1 + 2} \right) [\psi(b_1 + 1) - \psi(2)], \quad (2.4.1.2.2)$$

$$I_{13} = I_{14} = I_{15} = 0, \quad (2.4.1.2.3)$$

$$I_{22} = \frac{n}{\sigma_1^2} + \frac{nb_1}{(b_1 + 2)\sigma_1^2} \left[\psi'(b_1 + 1) + \psi'(2) + \{ \psi(b_1 + 1) - \psi(2) \}^2 \right], \quad (2.4.1.2.4)$$

$$I_{23} = I_{24} = I_{25} = 0, \quad (2.4.1.2.5)$$

$$I_{33} = \frac{nb_2}{(b_2 + 2)\sigma_{2.1}^2}, \quad (2.4.1.2.6)$$

$$I_{34} = \frac{nb_2}{(b_2 + 2)\sigma_{2.1}^2} \{ \psi(b_2 + 1) - \psi(2) \}, \quad (2.4.1.2.7)$$

$$I_{35} = \frac{nb_2}{(b_2 + 2)\sigma_{2.1}^2} [\mu_1 + \sigma_1 \{\psi(b_1) - \psi(1)\}], \quad (2.4.1.2.8)$$

$$I_{44} = \frac{n}{\sigma_{2.1}^2} + \frac{nb_2}{(b_2 + 2)\sigma_{2.1}^2} \left[\psi'(b_2 + 1) + \psi'(2) + \{\psi(b_2 + 1) - \psi(2)\}^2 \right], \quad (2.4.1.2.9)$$

$$I_{45} = \frac{nb_2}{(b_2 + 2)\sigma_{2.1}^2} \{\psi(b_2 + 1) - \psi(2)\} [\mu_1 + \sigma_1 \{\psi(b_1) - \psi(1)\}] \quad (2.4.1.2.10)$$

and

$$I_{55} = \frac{nb_2}{(b_2 + 2)\sigma_{2.1}^2} \times \left[\mu_1^2 + 2\mu_1\sigma_1 \{\psi(b_1) - \psi(1)\} + \{\psi'(b_1) + \psi'(1) + [\psi(b_1) - \psi(1)]^2\} \sigma_1^2 \right]. \quad (2.4.1.2.11)$$

2.4.1.3. Hypothesis Testing

The parameter θ is an indicator of the linear relationship between X and Y . Therefore, testing $H_0 : \theta = 0$ against $H_1 : \theta > 0$ is of great importance. The corresponding test statistic is

$$Z = \frac{\hat{\theta}}{\sqrt{V(\hat{\theta})}}, \quad (2.4.1.3.1)$$

where $V(\hat{\theta})$ is the variance of $\hat{\theta}$ under H_0 and its asymptotic value is obtained by taking the inverse of the Fisher information matrix $I(\mu_1, \sigma_1, \mu_{2.1}, \sigma_{2.1}, \theta)$.

Note that the null distribution of the test statistic Z is standard normal and large Z values mean the rejection of H_0 . The power properties of the Z -test will be discussed later.

2.4.1.4. Bias Corrected Least Square Estimators

Location and scale least square (LS) estimators should be corrected for bias, because instead of estimating the location and scale parameters, they estimate the mean and the variance of the distribution. Therefore, in this section, the bias correction procedure will be presented, when both marginal and conditional distribution are Generalized Logistic. Similar adjustments need to be done for other distributions.

Now, $E(X)$ and $V(X)$ in the Generalized Logistic distribution are

$$E(X) = \mu_1 + \sigma_1 \{\psi(b_1) - \psi(1)\}$$

and

$$V(X) = \sigma_1^2 \{\psi'(b_1) - \psi'(1)\}. \quad (2.4.1.4.1)$$

Therefore, the bias corrected LS estimators, $\hat{\mu}_{1,LS}$ and $\hat{\sigma}_{1,LS}$ are

$$\hat{\mu}_{1,LS} = \bar{x} - \hat{\sigma}_{1,LS} \{\psi(b_1) - \psi(1)\}$$

and

$$\hat{\sigma}_{1,LS} = \frac{s_x}{\sqrt{\psi'(b_1) + \psi'(1)}} \quad (2.4.1.4.2)$$

where

$$s_x = \sqrt{\sum_{i=1}^n (x_i - \bar{x})^2 / (n-1)}.$$

Similarly, $\bar{w}$ and s_w^2 are the LS estimators of $E(w)$ and $V(w)$ (remember that $w_i = y_i - \theta x_i$), and

$$E(w) = \mu_{2.1} + \sigma_{2.1} \{\psi(b_2) - \psi(1)\},$$

and

$$V(w) = \sigma_{2.1}^2 \{\psi'(b_2) - \psi'(1)\}. \quad (2.4.1.4.3)$$

Therefore, the bias corrected LS estimators, $\hat{\mu}_{2.1,LS}$ and $\hat{\sigma}_{2.1,LS}$, are

$$\begin{aligned} \hat{\mu}_{2.1,LS} &= \bar{w} - \hat{\sigma}_{2.1,LS} \{\psi(b_2) - \psi(1)\}, \\ \hat{\sigma}_{2.1,LS} &= \frac{s_w}{\sqrt{\psi'(b_2) + \psi'(1)}}, \end{aligned} \quad (2.4.1.4.4)$$

where

$$\bar{w} = \bar{y} - \left(\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y}) / s_x^2 \right) \bar{x} \text{ and } s_w = \sqrt{\sum_{i=1}^n (w_i - \bar{w})^2 / (n-1)}.$$

Also, for Y ($Y = \mu_{2.1} + \theta x + e$, $e = w - \mu_{2.1}$), $\bar{y}$ and s_y^2 are the LS estimators of $E(Y)$ and $V(Y)$, which are given in (2.4.1.4.5):

$$\begin{aligned} E(Y) &= \mu_2 + \sigma_2 \sqrt{1 - \rho^2} \{\psi(b_2) - \psi(1)\} \\ V(Y) &= \sigma_2^2 \left[\rho^2 \{\psi'(b_1) + \psi'(1)\} + (1 - \rho^2) \{\psi'(b_2) + \psi'(1)\} \right]. \end{aligned} \quad (2.4.1.4.5)$$

The bias corrected LS estimators, $\hat{\mu}_{2,LS}$ and $\hat{\sigma}_{2,LS}$, are

$$\begin{aligned}\hat{\mu}_{2,LS} &= \bar{y} - \hat{\sigma}_{2,LS} \sqrt{1 - \hat{\rho}_{LS}^2} \{\psi(b_2) - \psi(1)\} \\ \hat{\sigma}_{2,LS} &= \frac{s_y}{\sqrt{\hat{\rho}_{LS}^2 \{\psi'(b_1) + \psi'(1)\} + (1 - \hat{\rho}_{LS}^2) \{\psi'(b_2) + \psi'(1)\}}},\end{aligned}\quad (2.4.1.4.6)$$

where,

$$s_y = \sqrt{\sum_{i=1}^n (y_i - \bar{y})^2 / (n-1)}.$$

2.4.2. Censored Sample

In this section, the marginal distribution is coming from ‘censored’ Generalized Logistic while the conditional distribution is coming from Generalized Logistic.

2.4.2.1. Estimation of Parameters

The likelihood function for a Type II censored sample $(x_{(i)}, y_{[i]})$ ($1 \leq i \leq n-r$) is

$$\begin{aligned}L \propto \left(\frac{b_1}{\sigma_1}\right)^{n-r} \prod_{i=1}^{n-r} \left[\frac{e^{-z(i)}}{\left\{1 + e^{-z(i)}\right\}^{b_1+1}} \right] \left[1 - F(z_{(n-r)})\right]^r \times \\ \left(\frac{b_2}{\sigma_{2.1}}\right)^{n-r} \frac{e^{-\frac{1}{\sigma_{2.1}} \sum_{i=1}^{n-r} e_{[i]}}}{\prod_{i=1}^{n-r} \left[1 + e^{-\frac{e_{[i]}}{\sigma_{2.1}}}\right]^{b_2+1}}\end{aligned}\quad (2.4.2.1.1)$$

where

$$z(i) = \frac{x(i) - \mu_1}{\sigma_1}, \quad F(z) = \frac{1}{(1 + e^{-z})^b} \quad \text{and} \quad e[i] = y[i] - \mu_{2.1} - \theta x(i).$$

The log likelihood function is

$$\begin{aligned} \ln L \propto & -(n-r) \ln \sigma_1 + \sum_{i=1}^{n-r} z(i) - (b_1 + 1) \sum_{i=1}^{n-r} \ln(1 + e^{-z(i)}) + r \ln[1 - F(z_{(n-r)})] \\ & - (n-r) \ln \sigma_{2.1} - \frac{1}{\sigma_{2.1}} \sum_{i=1}^{n-r} e[i] - (b_2 + 1) \sum_{i=1}^{n-r} \ln \left(1 + e^{-\frac{e[i]}{\sigma_{2.1}}} \right). \end{aligned} \quad (2.4.2.1.2)$$

The likelihood equations are

$$\frac{\partial \ln L}{\partial \mu_1} = \frac{(n-r)}{\sigma_1} - \frac{(b_1 + 1)}{\sigma_1} \sum_{i=1}^{n-r} \left(\frac{e^{-z(i)}}{1 + e^{-z(i)}} \right) + \frac{r}{\sigma_1} \frac{f(z_{(n-r)})}{[1 - F(z_{(n-r)})]} = 0, \quad (2.4.2.1.3)$$

$$\begin{aligned} \frac{\partial \ln L}{\partial \sigma_1} = & -\frac{(n-r)}{\sigma_1} + \frac{1}{\sigma_1} \sum_{i=1}^{n-r} z(i) - \frac{(b_1 + 1)}{\sigma_1} \sum_{i=1}^{n-r} z(i) \frac{e^{-z(i)}}{1 + e^{-z(i)}} + \\ & \frac{r}{\sigma_1} z_{(n-r)} \frac{f(z_{(n-r)})}{1 - F(z_{(n-r)})} = 0, \end{aligned} \quad (2.4.2.1.4)$$

$$\frac{\partial \ln L}{\partial \mu_{2.1}} = \frac{(n-r)}{\sigma_{2.1}} - \frac{(b_2 + 1)}{\sigma_{2.1}} \sum_{i=1}^{n-r} \frac{e^{-e[i]/\sigma_{2.1}}}{[1 + e^{-e[i]/\sigma_{2.1}}]} = 0, \quad (2.4.2.1.5)$$

$$\begin{aligned} \frac{\partial \ln L}{\partial \sigma_{2.1}} = & -\frac{(n-r)}{\sigma_{2.1}} + \frac{1}{\sigma_{2.1}^2} \sum_{i=1}^{n-r} e[i] - \frac{(b_2 + 1)}{\sigma_{2.1}^2} \sum_{i=1}^{n-r} e[i] \frac{e^{-e[i]/\sigma_{2.1}}}{[1 + e^{-e[i]/\sigma_{2.1}}]} = 0 \\ & (2.4.2.1.6) \end{aligned}$$

and

$$\frac{\partial \ln L}{\partial \theta} = \frac{1}{\sigma_{2.1}} \sum_{i=1}^n x(i) - \frac{(b_2 + 1)}{\sigma_{2.1}} \sum_{i=1}^{n-r} x(i) \frac{e^{-e[i]/\sigma_{2.1}}}{\left[1 + e^{-e[i]/\sigma_{2.1}}\right]} = 0. \quad (2.4.2.1.7)$$

However, because of the functions $\frac{e^{-z(i)}}{1 + e^{-z(i)}}$, $\frac{f(z_{(n-r)})}{1 - F(z_{(n-r)})}$ and

$\frac{e^{-e[i]/\sigma_{2.1}}}{\left(1 + e^{-e[i]/\sigma_{2.1}}\right)}$, the equations (2.4.2.1.3) – (2.4.2.1.7) are difficult to solve.

Therefore, the estimators are obtained by using the MML methodology. First, a_i terms and the corresponding concomitants $(y[i], x[i])$'s are determined by the

procedure explained in section (2.2.2.1). Then, the functions $\frac{e^{-z(i)}}{1 + e^{-z(i)}}$,

$\frac{f(z_{(n-r)})}{1 - F(z_{(n-r)})}$ and $\frac{e^{-a(i)}}{\left(1 + e^{-a(i)}\right)}$ are replaced by $g_1(z(i))$, $g_2(z_{(n-r)})$ and

$g_3(a(i))$ respectively. After that, the likelihood equations become,

$$\frac{\partial \ln L}{\partial \mu_1} = \frac{(n-r)}{\sigma_1} - \frac{(b_1 + 1)}{\sigma_1} \sum_{i=1}^{n-r} g_1(z(i)) + \frac{r}{\sigma_1} g_2(z_{(n-r)}) = 0,$$

$$\frac{\partial \ln L}{\partial \sigma_1} = -\frac{(n-r)}{\sigma_1} + \frac{1}{\sigma_1} \sum_{i=1}^{n-r} z(i) - \frac{(b_1 + 1)}{\sigma_1} \sum_{i=1}^{n-r} z(i) g_1(z(i)) +$$

$$\frac{r}{\sigma_1} z_{(n-r)} g_2(z_{(n-r)}) = 0,$$

$$\frac{\partial \ln L}{\partial \mu_{2.1}} = \frac{(n-r)}{\sigma_{2.1}} - \frac{(b_2 + 1)}{\sigma_{2.1}} \sum_{i=1}^{n-r} g_3(a(i)) = 0,$$

$$\frac{\partial \ln L}{\partial \sigma_{2.1}} = -\frac{(n-r)}{\sigma_{2.1}} + \frac{1}{\sigma_{2.1}^2} \sum_{i=1}^{n-r} e(i) - \frac{(b_2 + 1)}{\sigma_{2.1}^2} \sum_{i=1}^{n-r} e(i) g_3(a(i)) = 0$$

$$\frac{\partial \ln L}{\partial \theta} = \frac{1}{\sigma_{2.1}} \sum_{i=1}^n x[i] - \frac{(b_2 + 1)}{\sigma_{2.1}} \sum_{i=1}^{n-r} x[i] g_3(a_{(i)}) = 0. \quad (2.4.2.1.8)$$

The linearized forms of $g_1(z_{(i)})$, $g_2(z_{(n-r)})$ and $g_3(a_{(i)})$ functions and the corresponding (α, β) values are given in the following equations:

$$\begin{aligned} g_1(z_{(i)}) &\cong g_1(t_{1(i)}) + (z_{(i)} - t_{1(i)}) \left. \frac{dg_1(z_{(i)})}{dz_{(i)}} \right|_{z_{(i)}=t_{1(i)}} \\ &= \alpha_{1i} - \beta_{1i} z_{(i)} \quad (1 \leq i \leq n-r) \end{aligned} \quad (2.4.2.1.9)$$

where

α_{1i} , β_{1i} and $t_{1(i)}$ are given in (2.4.1.1.11).

Similarly,

$$\begin{aligned} g_2(z_{(n-r)}) &\cong g_2(t_{2(n-r)}) + (z_{(n-r)} - t_{2(n-r)}) \left. \frac{dg_2(z_{(n-r)})}{dz_{(n-r)}} \right|_{z_{(n-r)}=t_{2(n-r)}} \\ &= \alpha_2 - \beta_2 z_{(n-r)}, \end{aligned} \quad (2.4.2.1.10)$$

where

$$\beta_2 = - \left. \frac{f'(z_{(n-r)}) [1 - F(z_{(n-r)})] + [f(z_{(n-r)})]^2}{[1 - F(z_{(n-r)})]^2} \right|_{z_{(n-r)}=t_{2(n-r)}},$$

$$\alpha_2 = g_2(z_{(n-r)}) + \beta_2 z_{(n-r)} \text{ at } z_{(n-r)} = t_{2(n-r)},$$

$$t_{2(n-r)} = -\ln \left[q_{2(n-r)}^{-1/b_1} - 1 \right], \quad q_{2(n-r)} = 1 - \frac{r}{n+1}. \quad (2.4.2.1.11)$$

Also,

$$\begin{aligned}
g_3(a_{(i)}) &\cong g_3(t_{3(i)}) + (a_{(i)} - t_{3(i)}) \frac{dg_3(a_{(i)})}{da_{(i)}} \Big|_{a_{(i)}=t_{3(i)}} \\
&= \alpha_{3i} - \beta_{3i} a_{(i)} \quad (1 \leq i \leq n-r)
\end{aligned} \tag{2.4.2.1.12}$$

where

$$\alpha_{3i} = \frac{e^{-t_{3(i)}}}{\left(1 + e^{-t_{3(i)}}\right)^2} \left(1 + t_{3(i)} + e^{-t_{3(i)}}\right), \quad \beta_{3i} = \frac{e^{-t_{3(i)}}}{\left(1 + e^{-t_{3(i)}}\right)^2},$$

and

$$t_{3(i)} = -\ln \left[q_{3i}^{-1/b_2} - 1 \right], \quad q_{3i} = \frac{i}{n+1}. \tag{2.4.2.1.13}$$

The modified likelihood equations are,

$$\begin{aligned}
\frac{\partial \ln L}{\partial \mu_1} &\cong \frac{\partial \ln L^*}{\partial \mu_1} = \frac{(n-r)}{\sigma_1} - \frac{(b_1+1)}{\sigma_1} \sum_{i=1}^{n-r} (\alpha_{1i} - \beta_{1i} z_{(i)}) \\
&\quad + \frac{r}{\sigma_1} (\alpha_2 - \beta_2 z_{(n-r)}) = 0,
\end{aligned}$$

$$\begin{aligned}
\frac{\partial \ln L}{\partial \sigma_1} &\cong \frac{\partial \ln L^*}{\partial \sigma_1} = -\frac{(n-r)}{\sigma_1} + \frac{1}{\sigma_1} \sum_{i=1}^{n-r} z_{(i)} - \frac{(b_1+1)}{\sigma_1} \sum_{i=1}^{n-r} z_{(i)} (\alpha_{1i} - \beta_{1i} z_{(i)}) \\
&\quad + \frac{r}{\sigma_1} z_{(n-r)} (\alpha_2 - \beta_2 z_{(n-r)}) = 0,
\end{aligned}$$

$$\frac{\partial \ln L}{\partial \mu_{2.1}} \cong \frac{\partial \ln L^*}{\partial \mu_{2.1}} = \frac{(n-r)}{\sigma_{2.1}} - \frac{(b_2+1)}{\sigma_{2.1}} \sum_{i=1}^{n-r} (\alpha_{3i} - \beta_{3i} a_{(i)}) = 0,$$

$$\frac{\partial \ln L}{\partial \sigma_{2.1}} \cong \frac{\partial \ln L^*}{\partial \sigma_{2.1}} = -\frac{(n-r)}{\sigma_{2.1}} + \frac{1}{\sigma_{2.1}^2} \sum_{i=1}^{n-r} e_{(i)} - \frac{(b_2+1)}{\sigma_{2.1}^2} \sum_{i=1}^{n-r} e_{(i)} (\alpha_{3i} - \beta_{3i} a_{(i)}) = 0$$

and

$$\frac{\partial \ln L}{\partial \theta} \cong \frac{\partial \ln L^*}{\partial \theta} = \frac{1}{\sigma_{2.1}} \sum_{i=1}^n x_{[i]} - \frac{(b_2+1)}{\sigma_{2.1}} \sum_{i=1}^{n-r} x_{[i]} (\alpha_{3i} - \beta_{3i} a_{(i)}) = 0. \quad (2.4.2.1.14)$$

The MML estimators which are the solutions of the equations given in (2.4.2.1.14) are

$$\begin{aligned} \hat{\mu}_1 &= K_1 + D_1 \hat{\sigma}_1, \\ \hat{\sigma}_1 &= \left\{ B_1 + \sqrt{B_1^2 + 4(n-r)C_1} \right\} / 2(n-r), \\ \hat{\mu}_{2.1} &= \bar{y}[\cdot] - \hat{\theta}[\cdot] - \frac{\Delta}{m_2} \hat{\sigma}_{2.1}, \\ \hat{\theta} &= K_2 - D_2 \hat{\sigma}_{2.1} \end{aligned}$$

and

$$\hat{\sigma}_{2.1} = \left\{ -B_2 + \sqrt{B_2^2 + 4(n-r)C_2} \right\} / 2(n-r)$$

where

$$\begin{aligned}
K_1 &= \left\{ (b_1 + 1) \sum_{i=1}^{n-r} \beta_{1i} x_{(i)} - r \beta_2 x_{(n-r)} \right\} / \left\{ (b_1 + 1) \sum_{i=1}^{n-r} \beta_{1i} - r \beta_2 \right\}, \\
D_1 &= \left\{ (n-r) - (b_1 + 1) \sum_{i=1}^{n-r} \alpha_{1i} + r \alpha_2 \right\} / \left\{ (b_1 + 1) \sum_{i=1}^{n-r} \beta_{1i} - r \beta_2 \right\}, \\
B_1 &= (b_1 + 1) \sum_{i=1}^{n-r} \left(x_{(i)} - K_1 \right) \left(\frac{1}{b_1 + 1} - \alpha_{1i} \right) + r \alpha_2 \left(x_{(n-r)} - K_1 \right), \\
C_1 &= (b_1 + 1) \sum_{i=1}^{n-r} \beta_{1i} \left(x_{(i)} - K_1 \right)^2 - r \beta_2 \left(x_{(n-r)} - K_1 \right)^2, \\
K_2 &= \sum_{i=1}^{n-r} \beta_{3i} \left(x_{[i]} - \bar{x}_{[.]} \right) y_{[i]} / \sum_{i=1}^{n-r} \beta_{3i} \left(x_{[i]} - \bar{x}_{[.]} \right)^2, \\
D_2 &= \sum_{i=1}^{n-r} \Delta_i \left(x_{[i]} - \bar{x}_{[.]} \right) / \sum_{i=1}^{n-r} \beta_{3i} \left(x_{[i]} - \bar{x}_{[.]} \right)^2, \\
B_2 &= (b_2 + 1) \sum_{i=1}^{n-r} \Delta_i \left\{ y_{[i]} - \bar{y}_{[.]} - K_2 \left(x_{[i]} - \bar{x}_{[.]} \right) \right\} \\
C_2 &= (b_2 + 1) \left\{ \sum_{i=1}^{n-r} \beta_{3i} \left(y_{[i]} - \bar{y}_{[.]} \right)^2 - K_2 \sum_{i=1}^{n-r} \beta_{3i} \left(x_{[i]} - \bar{x}_{[.]} \right) y_{[i]} \right\}; \\
\bar{y}_{[.]} &= \frac{1}{m_3} \sum_{i=1}^{n-r} \beta_{3i} y_{[i]}, \quad \bar{x}_{[.]} = \frac{1}{m_3} \sum_{i=1}^{n-r} \beta_{3i} x_{[i]}, \quad m_3 = \sum_{i=1}^{n-r} \beta_{3i},
\end{aligned}$$

and

$$\Delta_i = \alpha_{3i} - \frac{1}{b_2 + 1} \quad \text{and} \quad \Delta = \sum_{i=1}^{n-r} \Delta_i. \quad (2.4.2.1.15)$$

Also, for the MML estimators of μ_2 , σ_2 and ρ , see the equation (2.4.1.1.17).

2.4.2.2. Sample Information Matrix

Since Fisher information matrix $I(\mu_1, \sigma_1, \mu_{2,1}, \sigma_{2,1}, \theta)$ is intractable, instead of I , the sample information matrix $\hat{I}$ is used. The components of $\hat{I}$ is given in the following equations:

$$\hat{I} = [\hat{I}_{ij}] = \left[-\frac{\partial^2 \ln L}{\partial \delta_i \partial \delta_j} \right] \quad i, j = 1, 2, 3, 4, 5$$

$$\delta_1 = \hat{\mu}_1, \delta_2 = \hat{\sigma}_1, \delta_3 = \hat{\mu}_{2,1}, \delta_4 = \hat{\sigma}_{2,1}, \delta_5 = \hat{\theta},$$

$$\hat{I}_{11} = \frac{(b_1 + 1)}{\hat{\sigma}_1^2} \sum_{i=1}^{n-r} \frac{e^{-z(i)}}{\left\{ 1 + e^{-z(i)} \right\}^2} + \frac{r}{\hat{\sigma}_1^2} \frac{dg_2(z_{(n-r)})}{dz_{(n-r)}}, \quad (2.4.2.2.1)$$

$$\begin{aligned} \hat{I}_{12} = & \frac{(n-r)}{\hat{\sigma}_1^2} - \frac{(b_1 + 1)}{\hat{\sigma}_1^2} \sum_{i=1}^{n-r} g_1(z(i)) + \frac{(b_1 + 1)}{\hat{\sigma}_1^2} \sum_{i=1}^{n-r} z(i) \frac{e^{-z(i)}}{\left\{ 1 + e^{-z(i)} \right\}^2} + \\ & \frac{r}{\hat{\sigma}_1^2} g_2(z_{(n-r)}) + \frac{r}{\hat{\sigma}_1^2} z_{(n-r)} \frac{dg_2(z_{(n-r)})}{dz_{(n-r)}}, \end{aligned} \quad (2.4.2.2.2)$$

$$\hat{I}_{13} = \hat{I}_{14} = \hat{I}_{15} = 0.0, \quad (2.4.2.2.3)$$

$$\begin{aligned} \hat{I}_{22} = & -\frac{(n-r)}{\hat{\sigma}_1^2} + \frac{2}{\hat{\sigma}_1^2} \sum_{i=1}^{n-r} z(i) - 2 \frac{(b_1 + 1)}{\hat{\sigma}_1^2} \sum_{i=1}^{n-r} z(i) g_1(z(i)) + \frac{(b_1 + 1)}{\hat{\sigma}_1^2} \sum_{i=1}^{n-r} z(i)^2 \times \\ & \frac{e^{-z(i)}}{\left\{ 1 + e^{-z(i)} \right\}^2} + \frac{2r}{\hat{\sigma}_1^2} z_{(n-r)} g_2(z_{(n-r)}) + \frac{r}{\hat{\sigma}_1^2} z_{(n-r)}^2 \frac{dg_2(z_{(n-r)})}{dz_{(n-r)}}, \end{aligned} \quad (2.4.2.2.4)$$

$$\hat{I}_{23} = \hat{I}_{24} = \hat{I}_{25} = 0.0, \quad (2.4.2.2.5)$$

$$\hat{I}_{33} = \frac{(b_2 + 1)}{\hat{\sigma}_{2,1}^2} \sum_{i=1}^{n-r} \frac{e^{-a[i]}}{\left(1 + e^{-a[i]} \right)^2}, \quad (2.4.2.2.3)$$

$$\begin{aligned} \hat{I}_{34} = & \frac{(n-r)}{\hat{\sigma}_{2,1}^2} - \frac{(b_2 + 1)}{\hat{\sigma}_{2,1}^2} \sum_{i=1}^{n-r} \frac{e^{-a[i]}}{\left(1 + e^{-a[i]} \right)} + \frac{(b_2 + 1)}{\hat{\sigma}_{2,1}^2} \sum_{i=1}^{n-r} a[i] \frac{e^{-a[i]}}{\left(1 + e^{-a[i]} \right)^2}, \\ & (2.4.2.2.4) \end{aligned}$$

$$\hat{I}_{35} = \frac{(b_2 + 1)}{\hat{\sigma}_{2,1}^2} \sum_{i=1}^{n-r} x(i) \frac{e^{-a[i]}}{\left(1 + e^{-a[i]} \right)^2}, \quad (2.4.2.2.5)$$

$$\hat{I}_{44} = -\frac{(n-r)}{\hat{\sigma}_{2.1}^2} + \frac{2}{\hat{\sigma}_{2.1}^2} \sum_{i=1}^{n-r} a[i] - \frac{2(b_2+1)}{\hat{\sigma}_{2.1}^2} \sum_{i=1}^{n-r} a[i] \frac{e^{-a[i]}}{\left(1+e^{-a[i]}\right)} + \frac{(b_2+1)}{\hat{\sigma}_{2.1}^2} \sum_{i=1}^{n-r} a[i]^2 \frac{e^{-a[i]}}{\left(1+e^{-a[i]}\right)^2}, \quad (2.4.2.2.6)$$

$$\hat{I}_{45} = \frac{1}{\hat{\sigma}_{2.1}^2} \sum_{i=1}^{n-r} x(i) - \frac{(b_2+1)}{\hat{\sigma}_{2.1}^2} \sum_{i=1}^{n-r} x(i) \frac{e^{-a[i]}}{\left(1+e^{-a[i]}\right)} + \frac{(b_2+1)}{\hat{\sigma}_{2.1}^2} \times \sum_{i=1}^{n-r} x(i) a[i] \frac{e^{-a[i]}}{\left(1+e^{-a[i]}\right)^2} \quad (2.4.2.2.7)$$

and

$$\hat{I}_{55} = \frac{(b_2+1)}{\hat{\sigma}_{2.1}^2} \sum_{i=1}^{n-r} x(i)^2 \frac{e^{-a[i]}}{\left(1+e^{-a[i]}\right)^2}. \quad (2.4.2.2.8)$$

In $z(i)$ and $a[i]$, the unknown parameters are replaced by their MML estimators.

Comment: What is very interesting indeed is that the MML estimators have the same forms irrespective of the underlying marginal and conditional distributions. It is also clear why these problems remained unsolved because earlier authors tried maximum likelihood estimation which is enormously problematic in the situations we have considered in this thesis. Modified maximum likelihood estimation has made it possible to solve these difficult problems.

Remark: The Fisher information matrices or the sample information matrices can be used to determine the variances (and covariances) of the MMLE's. They give accurate values even for sample sizes as small as $n=20$. This is illustrated in the next chapter, e.g., compare the values given in Tables 3.1.1 and 3.3.1.

CHAPTER 3

SIMULATION RESULTS AND ILLUSTRATIVE EXAMPLES

In this chapter, the simulation results for both complete and censored samples are given. The results are performed for 100,000/ n Monte Carlo runs and n is equal to 20, 50 and 100. The case when marginal distribution is Generalized Logistic and conditional distribution is normal, also the case when both marginal and conditional distributions are Generalized Logistic are taken into consideration. However, the simulations can also be done for other location – scale distributions.

In the censoring situation, samples from the assumed marginal distribution are censored and r (censoring number) is to be $\text{int}[0.10*n]$. In the simulation study, the effects of the largest observations (in the marginal distribution) on the sample size and on the estimators are evaluated. Outliers are created according to Tiku's and Dixon's outlier model.

Note that, in the simulations C_0 , which is the total permissible cost, is taken as $nc + n\mu_1 + \sigma_1 \sum_{i=1}^n t_{(i)}$. It is assumed that $c + \mu_1 + \sigma_1 t_{(1)} > 0$, as expected.

This is essentially saying that the cost of observing the i^{th} ordered observation is directly related to its expected value.

In Dixon's outlier model, in order to create outliers, a constant is added to the r of the X -observations, which are randomly selected. However, in Tiku's

outlier model, observations are ordered and a constant is added to the largest r of the observations. In other words, by using Tiku's outlier model, creating outliers are guaranteed (See Dixon, 1950; 1953; Hawkins, 1977; Tiku, 1975; 1977).

It should also be mentioned that in the complete sample simulations, the performance of the MML estimators are compared with the bias corrected LS estimators by looking at their relative efficiencies. Note that no bias correction is applied to MML estimators since the sample sizes are large enough and MML estimators are self bias correcting. Relative efficiencies of the LS estimators depending on the variance (var) and mean square error ($mse=var+bias^2$) are calculated. The formula for the corresponding relative efficiencies are given below,

$$RE(LS) \text{ var} = [\text{var}(MMLE)/\text{var}(LSE)]*100$$

$$RE(LS) \text{ mse} = [mse(MMLE)/mse(LSE)]*100$$

3.1. Simulation Results when Marginal Distribution is Generalized Logistic and Conditional Distribution is Normal

In this section, the simulation results for both complete and censored samples are given, when the marginal distribution is Generalized Logistic and the conditional distribution is normal (Section 2.1). Also, via the simulations, the effects of the outliers in X , which are created according to Dixon's and Tiku's outlier model are discussed. During this evaluation, the estimated sample sizes and the performance of the estimators, $\hat{\mu}_1$, $\hat{\sigma}_1$, $\hat{\mu}_{2,1}$, $\hat{\sigma}_{2,1}$, $\hat{\theta}$, $\hat{\mu}_2$, $\hat{\sigma}_2$ and $\hat{\rho}$ are evaluated. In the simulations $b, \mu_1, \sigma_1, \mu_{2,1}, \sigma_{2,1}, \theta, \mu_2, \sigma_2$, and ρ are taken to be 4.0, 0.0, 1.0, 0.0, 0.8660, 0.5, 0.9167, 1.0, and 0.5, respectively. The estimated sample sizes for complete and censored samples are determined subject to the permissible cost and the formulas for them are,

$$\hat{n}_{comp} = \text{int} \left[\frac{C_0}{c + \hat{\mu}_1 + \hat{\sigma}_1 \{\psi(b) - \psi(1)\}} \right], \quad \hat{n}_{cens} = r + \frac{C_0}{c + \hat{\mu}_1 + \frac{\hat{\sigma}_1}{n-r} \sum_{i=1}^{n-r} t(i)}. \quad (3.1.1)$$

Here, c is a threshold cost so that the denominator of the equation given in (3.1.1) is always greater than zero.

Table 3.1.1 Simulation results for complete sample when $X \sim GL$, $Y|X \sim \text{normal}$ - without outliers

n=20									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.064	0.972	-0.001	0.811	0.498	0.920	0.959	0.502	20
mean									
LSE	0.048	0.972	-0.001	0.855	0.498	0.912	0.935	0.605	20
variance									
MMLE	0.115	0.033	0.114	0.019	0.024	0.061	0.022	0.022	13.309
variance									
LSE	0.137	0.044	0.114	0.021	0.024	0.061	0.021	0.025	14.204
RE (LS)									
var	83.8	75.3	100.0	89.8	100.0	99.8	104.3	87.8	93.7
RE (LS)									
mse	85.5	75.7	100.0	103.9	100.0	99.8	93.4	60.6	
n=50									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.039	0.985	0.009	0.844	0.498	0.927	0.982	0.498	50
mean									
LSE	0.032	0.985	0.009	0.862	0.498	0.924	0.943	0.615	50
variance									
MMLE	0.045	0.013	0.047	0.007	0.009	0.026	0.009	0.008	29.510
variance									
LSE	0.057	0.018	0.047	0.008	0.009	0.026	0.008	0.010	30.606
RE (LS)									
var	78.9	71.2	100.0	96.1	100.0	99.6	106.3	86.3	96.4
RE (LS)									
mse	80.2	71.6	100.0	102.1	100.0	99.8	78.8	36.3	
n=100									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.007	1.000	0.001	0.855	0.502	0.925	0.995	0.504	100
mean									
LSE	0.006	0.999	0.001	0.864	0.502	0.923	0.950	0.624	100
variance									
MMLE	0.022	0.007	0.023	0.004	0.004	0.012	0.004	0.004	53.727
variance									
LSE	0.027	0.009	0.023	0.004	0.004	0.012	0.004	0.005	56.444
RE (LS)									
var	79.6	69.9	100.0	97.4	100.0	97.5	107.9	87.5	95.2
RE (LS)									
mse	79.6	69.9	100.0	100.3	100.0	97.7	65.7	20.8	

Table 3.1.2 Simulation results for complete sample when $X \sim GL$, $Y|X \sim \text{normal}$ - Dixon's outlier model (2.0 is added to the r of the X -observations)

n=20									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.082	1.066	-0.002	0.807	0.500	1.016	0.980	0.540	18
mean									
LSE	0.052	1.076	-0.002	0.850	0.500	1.010	0.956	0.647	18
variance									
MMLE	0.122	0.034	0.114	0.019	0.020	0.062	0.022	0.020	8.532
variance									
LSE	0.149	0.045	0.114	0.021	0.020	0.062	0.021	0.021	8.825
RE (LS)									
var	81.6	74.0	100.0	90.1	100.0	100.3	103.3	92.9	96.7
RE (LS)									
mse	84.5	74.2	100.0	105.5	100.0	101.9	96.3	49.8	
n=50									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.059	1.077	0.009	0.844	0.498	1.022	1.005	0.532	45
mean									
LSE	0.054	1.079	0.009	0.861	0.498	1.021	0.964	0.650	45
variance									
MMLE	0.050	0.013	0.043	0.007	0.007	0.024	0.009	0.007	19.871
variance									
LSE	0.061	0.018	0.043	0.007	0.007	0.024	0.008	0.008	20.089
RE (LS)									
var	81.4	71.3	100.0	95.9	100.0	99.6	106.3	91.1	98.9
RE (LS)									
mse	83.2	77.4	100.0	102.4	100.0	100.1	91.6	27.2	
n=100									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.048	1.081	-0.002	0.854	0.500	1.013	1.013	0.533	90
mean									
LSE	0.048	1.083	-0.002	0.863	0.500	1.015	0.967	0.654	90
variance									
MMLE	0.024	0.007	0.024	0.004	0.004	0.012	0.004	0.004	37.678
variance									
LSE	0.029	0.009	0.024	0.004	0.004	0.012	0.004	0.004	38.301
RE (LS)									
var	82.6	75.3	100.0	97.4	100.0	99.2	107.3	94.9	98.4
RE (LS)									
mse	83.8	84.4	100.0	100.8	100.0	98.3	88.4	17.2	

In Tables 3.1.1 – 3.1.5, complete sample results are given while in Tables 3.1.6 - 3.1.10, censored sample results are displayed. Since the conditional distribution is normal, it is seen that the MML and LS estimators, $\hat{\mu}_{2.1}$, $\hat{\sigma}_{2.1}$ and $\hat{\theta}$, which are coming from the conditional distribution perform similarly.

In Table 3.1.1, the performance of the estimators are displayed for complete sample without outliers. It is seen that MML estimators have negligible

bias. The MML estimators $\hat{\mu}_1$, $\hat{\sigma}_1$, $\hat{\mu}_2$, $\hat{\sigma}_2$ and $\hat{\rho}$ are more efficient than LS estimators and their efficiencies increase as sample size increases. Also, considering the estimated sample sizes, it is observed that the estimated sample sizes are equal to the actual values.

Table 3.1.3 Simulation results for complete sample when $X \sim GL$, $Y|X \sim \text{normal}$ - Dixon's outlier model (4.0 is added to the r of the X -observations)

n=20									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	-0.034	1.224	-0.002	0.811	0.502	1.107	1.028	0.595	16
mean									
LSE	-0.187	1.320	-0.002	0.854	0.502	1.118	1.025	0.724	16
variance									
MMLE	0.128	0.033	0.100	0.019	0.012	0.061	0.021	0.014	6.140
variance									
LSE	0.173	0.046	0.100	0.021	0.012	0.061	0.022	0.012	5.748
RE (LS)									
var	73.6	72.1	100.0	89.8	100.0	100.8	95.9	114.6	106.8
RE (LS)									
mse	61.9	56.1	100.0	104.1	100.0	96.0	96.7	36.9	
n=50									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	-0.069	1.232	0.005	0.843	0.499	1.098	1.048	0.586	42
mean									
LSE	-0.190	1.320	0.005	0.861	0.499	1.118	1.029	0.725	41
variance									
MMLE	0.047	0.014	0.038	0.008	0.005	0.024	0.009	0.006	14.561
variance									
LSE	0.068	0.020	0.038	0.008	0.005	0.024	0.009	0.005	13.373
RE (LS)									
var	69.0	68.5	100.0	96.2	100.0	100.4	95.6	114.6	108.9
RE (LS)									
mse	49.6	55.1	100.0	102.3	100.0	88.2	111.1	23.2	
n=100									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	-0.082	1.232	0.001	0.857	0.500	1.089	1.057	0.582	84
mean									
LSE	-0.205	1.326	0.001	0.865	0.500	1.113	1.034	0.727	82
variance									
MMLE	0.024	0.006	0.018	0.004	0.002	0.013	0.004	0.003	30.426
variance									
LSE	0.034	0.009	0.018	0.004	0.002	0.013	0.004	0.002	27.325
RE (LS)									
var	69.7	68.8	100.0	97.4	100.0	101.6	100.0	122.7	111.3
RE (LS)									
mse	40.0	52.1	100.0	99.7	100.0	82.9	139.7	17.5	

Tables 3.1.2 and 3.1.3 show that outliers in X , created with respect to the Dixon's outlier model, affect the estimators significantly. The performance of the

estimators, especially the LS estimators, deteriorate which means they have larger bias and variance. Also, considering the efficiencies, it is seen that the MML estimators $\hat{\mu}_1$, $\hat{\sigma}_1$, $\hat{\mu}_2$, $\hat{\sigma}_2$ and $\hat{\rho}$ are more advantageous than LS estimators.

Also, the estimated sample sizes are smaller than the actual ones. Therefore, it is understood that the outliers in X make it impossible to observe the required sample sizes subject to fixed cost.

Table 3.1.4 Simulation results for complete sample when $X \sim GL$, $Y|X \sim \text{normal}$ - Tiku's outlier model (2.0 is added to the largest r of the X -observations)

n=20									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	-0.093	1.141	-0.005	0.809	0.501	0.996	1.001	0.569	18
mean									
LSE	-0.362	1.306	-0.005	0.852	0.501	1.013	1.019	0.720	18
variance									
MMLE	0.112	0.032	0.090	0.019	0.013	0.061	0.021	0.015	9.130
variance									
LSE	0.141	0.045	0.090	0.021	0.013	0.062	0.022	0.013	8.807
RE (LS)									
var	79.3	69.8	100.0	89.9	100.0	99.0	92.0	114.2	103.7
RE (LS)									
mse	44.3	37.1	100.0	104.9	100.0	95.0	90.6	31.5	
n=50									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	-0.127	1.151	-0.006	0.843	0.503	0.990	1.027	0.563	46
mean									
LSE	-0.384	1.318	-0.006	0.861	0.503	1.016	1.031	0.727	45
variance									
MMLE	0.042	0.012	0.035	0.007	0.005	0.025	0.008	0.005	20.769
variance									
LSE	0.055	0.019	0.035	0.007	0.005	0.026	0.008	0.005	19.641
RE (LS)									
var	76.6	62.8	100.0	95.9	100.0	98.4	90.4	114.9	105.7
RE (LS)									
mse	28.8	28.8	100.0	102.5	100.0	86.7	89.0	16.7	
n=100									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	-0.141	1.154	-0.001	0.854	0.500	0.986	1.033	0.558	93
mean									
LSE	-0.396	1.324	-0.001	0.863	0.500	1.014	1.031	0.727	90
variance									
MMLE	0.023	0.006	0.018	0.004	0.003	0.011	0.004	0.003	37.366
variance									
LSE	0.032	0.009	0.018	0.004	0.003	0.012	0.005	0.002	34.339
RE (LS)									
var	73.3	66.0	100.0	97.4	100.0	99.1	93.3	120.8	108.8
RE (LS)									
mse	22.8	26.1	100.0	100.8	100.0	77.2	96.1	11.6	

Table 3.1.5 Simulation results for complete sample when $X \sim GL$, $Y|X \sim \text{normal}$ - Tiku's outlier model (4.0 is added to the largest r of the X -observations)

n=20									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	-0.286	1.341	0.006	0.811	0.499	1.089	1.061	0.630	17
mean									
LSE	-0.870	1.694	0.006	0.855	0.499	1.120	1.144	0.803	16
variance									
MMLE	0.110	0.030	0.073	0.019	0.007	0.060	0.019	0.011	6.336
variance									
LSE	0.146	0.045	0.073	0.021	0.007	0.060	0.024	0.006	5.822
RE (LS)									
var	75.4	67.0	100.0	90.1	100.0	98.8	80.4	172.6	108.8
RE (LS)									
mse	21.2	27.9	100.0	103.7	100.0	87.9	51.5	28.1	
n=50									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	-0.308	1.338	-0.002	0.842	0.501	1.073	1.080	0.621	43
mean									
LSE	-0.884	1.700	-0.002	0.859	0.501	1.117	1.153	0.808	41
variance									
MMLE	0.044	0.012	0.029	0.008	0.003	0.024	0.008	0.004	15.067
variance									
LSE	0.059	0.018	0.029	0.008	0.003	0.024	0.010	0.002	13.250
RE (LS)									
var	75.0	64.5	100.0	96.2	100.0	99.2	81.3	173.9	113.7
RE (LS)									
mse	16.6	24.8	100.0	103.0	100.0	75.0	43.0	19.3	
n=100									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	-0.319	1.339	0.001	0.854	0.499	1.067	1.087	0.615	86
mean									
LSE	-0.891	1.704	0.001	0.863	0.499	1.116	1.154	0.806	82
variance									
MMLE	0.023	0.006	0.015	0.003	0.001	0.014	0.003	0.002	31.452
variance									
LSE	0.032	0.010	0.015	0.003	0.001	0.014	0.005	0.001	27.586
RE (LS)									
var	73.0	59.8	100.0	97.1	100.0	97.9	75.6	166.7	114.0
RE (LS)									
mse	15.1	23.9	100.0	100.8	100.0	67.8	38.4	16.0	

Similar to the previous two tables, Table 3.1.4 and 3.1.5 exhibit the effects of the outliers (in X) on the means and variances of the estimators. But, it should be noted that for these cases, outliers are created with respect to the Tiku's outlier model. As expected, the estimators become biased and they have larger variances when compared to the estimators obtained from the complete sample without outliers. However, it is easily seen that the MML estimators have superiority in the efficiencies compared to the LS estimators. Similarly, it is once again well

understood that the outliers enormously increase the cost of observing the required sample sizes. In other words, it is seen that with a fixed budget it becomes impossible to observe as many observations as we want when there are outliers in the sample.

The performance of the estimators, obtained from the censored sample without outliers (Table 3.1.6) bear a strong resemblance to the performance of the estimators, obtained from the complete sample without outliers (Table 3.1.1). The mean and variance differences observed in the aforesaid tables are unimportant, and furthermore, the estimated sample sizes are more than the actual ones, which means a high amount of money is saved by not observing the largest r observations.

Table 3.1.6 Simulation results for censored sample when $X \sim GL$, $Y|X \sim \text{normal}$ - without outliers

n=20									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.074	0.960	-0.001	0.803	0.502	0.918	0.958	0.492	26
variance									
MMLE	0.113	0.034	0.156	0.021	0.049	0.066	0.027	0.035	25.090
n=50									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.028	0.986	-0.002	0.845	0.499	0.914	0.986	0.495	65
variance									
MMLE	0.046	0.014	0.057	0.008	0.017	0.027	0.011	0.013	53.064
n=100									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.007	0.997	0.005	0.855	0.496	0.915	0.992	0.496	130
variance									
MMLE	0.023	0.007	0.030	0.004	0.009	0.014	0.006	0.007	109.641

In Tables 3.1.7 and 3.1.8, the effects of the outliers in X (created under the Dixon's outlier model) on the performance of the estimators and the sample sizes are presented. Accordingly, it is obvious that censoring improves the performance of the estimators, when there are outliers in the sample. Additionally, it will not be wrong to mention that censoring does not harm the performance of the estimators

even when there exist no outliers in the data set. Also, when the results displayed in Tables 3.1.3 and 3.1.8 are compared, it is realized that with a fixed budget, censoring increases the accessible sample size from 16, 42 and 84 to 26, 65 and 130, respectively.

Table 3.1.7 Simulation results for censored sample when $X \sim GL$, $Y|X \sim \text{normal}$ - Dixon's outlier model (2.0 is added to the r of the X -observations)

n=20									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.072	0.962	0.006	0.805	0.497	0.918	0.958	0.489	26
variance									
MMLE	0.115	0.033	0.151	0.020	0.046	0.067	0.026	0.035	24.477
n=50									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.031	0.985	-0.006	0.842	0.503	0.917	0.985	0.498	65
variance									
MMLE	0.045	0.014	0.056	0.008	0.017	0.028	0.012	0.013	54.194
n=100									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.007	0.996	-0.001	0.855	0.500	0.915	0.993	0.499	130
variance									
MMLE	0.022	0.007	0.026	0.004	0.008	0.014	0.006	0.006	106.914

Table 3.1.8 Simulation results for censored sample when $X \sim GL$, $Y|X \sim \text{normal}$ - Dixon's outlier model (4.0 is added to the r of the X -observations)

n=20									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.078	0.956	-0.001	0.805	0.501	0.916	0.957	0.490	26
variance									
MMLE	0.111	0.033	0.154	0.020	0.046	0.069	0.026	0.033	24.451
n=50									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.022	0.988	0.008	0.838	0.496	0.915	0.980	0.496	65
variance									
MMLE	0.046	0.014	0.058	0.008	0.017	0.027	0.011	0.014	52.306
n=100									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.020	0.990	0.010	0.857	0.491	0.911	0.990	0.490	130
variance									
MMLE	0.024	0.007	0.026	0.004	0.008	0.015	0.005	0.007	105.348

Tables 3.1.9 and 3.1.10 exhibit the simulation results for censored X , when it includes outliers. In fact, the performance of these estimators are almost the same as the performance of the estimators obtained from the censored sample without outliers (Table 3.1.6). The reason for this superior performance is related to the outlier creation procedure. Since, in Tiku's outlier model a constant is added to the largest observations, in the censoring procedure, they become the ones that are censored.

Table 3.1.9 Simulation results for censored sample when $X \sim GL$, $Y|X \sim \text{normal}$ - Tiku's outlier model (2.0 is added to the largest r of the X -observations)

n=20									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.069	0.961	-0.001	0.806	0.498	0.911	0.959	0.489	26
variance									
MMLE	0.117	0.034	0.148	0.020	0.046	0.070	0.027	0.034	27.027
n=50									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.026	0.986	0.007	0.842	0.498	0.920	0.983	0.495	65
variance									
MMLE	0.048	0.014	0.055	0.009	0.017	0.028	0.011	0.013	57.461
n=100									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.010	0.999	0.000	0.852	0.499	0.919	0.991	0.501	130
variance									
MMLE	0.023	0.007	0.027	0.004	0.008	0.014	0.005	0.007	96.933

Comparing the results given in Tables 3.1.5 and 3.1.10, it is seen that the accessible sample sizes increase dramatically. Indeed, from 17, 43 and 86 to 26, 65 and 131, while the actual sample sizes are also equal to 20, 50 and 100, respectively.

By combining the results of the Section 3.1, it can be said with confidence that the MML estimators perform better than the LS estimators. The performance gap between the estimators increases as sample size increases and when the sample contains outliers.

Also, the other important result is that not observing a few largest x -observations leads to advantageous results.

Table 3.1.10 Simulation results for censored sample when $X \sim GL$, $Y|X \sim \text{normal}$ - Tiku's outlier model (4.0 is added to the largest r of the X -observations)

n=20									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.073	0.964	0.000	0.802	0.501	0.922	0.957	0.493	26
variance									
MMLE	0.115	0.035	0.152	0.020	0.046	0.069	0.027	0.034	23.702
n=50									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.033	0.986	-0.001	0.843	0.502	0.923	0.986	0.499	65
variance									
MMLE	0.047	0.014	0.054	0.008	0.016	0.027	0.011	0.013	54.729
n=100									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.009	0.991	0.002	0.855	0.501	0.917	0.993	0.498	131
variance									
MMLE	0.024	0.007	0.029	0.004	0.009	0.013	0.005	0.007	104.667

3.2. Simulation Results when Both Marginal and Conditional Distributions are Generalized Logistic

In this section, the estimated sample sizes and the performance of the estimators, $\hat{\mu}_1$, $\hat{\sigma}_1$, $\hat{\mu}_{2.1}$, $\hat{\sigma}_{2.1}$, $\hat{\theta}$, $\hat{\mu}_2$, $\hat{\sigma}_2$ and $\hat{\rho}$ are evaluated for both complete and censored sample when both marginal and conditional distribution are Generalized Logistic (Section 2.4). Also, how these estimators and the estimated sample size are affected from the largest X observations is another question of interest. In the simulations b_1 , b_2 , μ_1 , σ_1 , $\mu_{2.1}$, $\sigma_{2.1}$, θ , μ_2 , σ_2 , and ρ are taken to be 4.0, 4.0, 0.0, 1.0, 0.0, 0.8660, 0.5, 0.9167, 1.0, and 0.5, respectively. The formulas for the estimated sample sizes for complete and censored samples are

$$\hat{n}_{comp} = \text{int} \left[\frac{C_0}{c + \hat{\mu}_1 + \hat{\sigma}_1 \{\psi(b_1) - \psi(1)\}} \right], \quad \hat{n}_{cens} = r + \frac{C_0}{c + \hat{\mu}_1 + \frac{\hat{\sigma}_1}{n-r} \sum_{i=1}^{n-r} t_{(i)}}. \quad (3.2.1)$$

Here, C_0 is the permissible cost and c is a threshold cost as before.

Table 3.2.1 Simulation results for complete sample when $X \sim GL$, $Y|X \sim GL$ - without outliers

n=20									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.051	0.980	0.087	0.819	0.506	1.021	0.977	0.500	20
mean									
LSE	0.032	0.981	0.044	0.843	0.499	0.998	0.979	0.493	20
variance									
MMLE	0.110	0.032	0.210	0.024	0.040	0.109	0.029	0.030	13.009
variance									
LSE	0.130	0.043	0.251	0.032	0.047	0.122	0.037	0.036	13.794
RE (LS)									
var	84.7	75.0	83.7	74.1	83.8	89.3	78.7	84.2	94.3
RE (LS)									
mse	86.0	75.3	86.1	79.6	83.8	93.1	79.2	84.1	
n=50									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.020	0.991	0.032	0.846	0.505	0.960	0.990	0.503	50
mean									
LSE	0.017	0.988	0.016	0.856	0.502	0.950	0.990	0.499	50
variance									
MMLE	0.044	0.012	0.081	0.010	0.013	0.045	0.011	0.011	28.549
variance									
LSE	0.054	0.017	0.098	0.013	0.016	0.054	0.015	0.014	29.731
RE (LS)									
var	82.3	71.5	82.5	70.9	81.0	83.7	77.2	79.0	96.0
RE (LS)									
mse	82.6	71.4	83.3	73.3	81.1	85.3	77.4	79.1	
n=100									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.008	0.999	0.012	0.856	0.504	0.938	0.997	0.504	100
mean									
LSE	0.002	1.001	0.007	0.859	0.502	0.937	0.996	0.504	100
variance									
MMLE	0.024	0.007	0.036	0.005	0.006	0.022	0.006	0.005	56.006
variance									
LSE	0.030	0.010	0.045	0.007	0.007	0.026	0.008	0.007	59.106
RE (LS)									
var	78.9	68.4	80.7	67.1	78.1	82.8	71.8	74.3	94.8
RE (LS)									
mse	79.0	68.4	80.9	67.9	78.2	83.2	71.8	74.3	

The mean and variance of the MML and LS estimators obtained from the complete sample without outliers is given in Table 3.2.1. It is seen that MML estimators have superiority compared to the LS estimators and it is seen that their efficiencies increase as sample size increases. Additionally, the estimated sample sizes are equal to the actual values.

Table 3.2.2 Simulation results for complete sample when $X \sim GL$, $Y|X \sim GL$ - Dixon's outlier model (2.0 is added to the r of the X -observations)

n=20									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.084	1.066	0.071	0.825	0.506	1.103	1.004	0.531	18
mean									
LSE	0.058	1.074	0.025	0.851	0.499	1.079	1.009	0.525	18
variance									
MMLE	0.127	0.033	0.217	0.024	0.032	0.113	0.030	0.026	8.201
variance									
LSE	0.153	0.044	0.267	0.035	0.039	0.133	0.039	0.032	8.423
RE (LS)									
var	83.1	75.5	81.2	69.9	82.6	85.4	75.9	82.4	97.4
RE (LS)									
mse	85.9	76.0	82.9	74.2	82.7	93.1	75.8	83.7	
n=50									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.054	1.082	0.030	0.848	0.503	1.056	1.015	0.535	45
mean									
LSE	0.049	1.084	0.020	0.855	0.499	1.051	1.014	0.531	45
variance									
MMLE	0.048	0.013	0.077	0.010	0.011	0.043	0.011	0.010	19.855
variance									
LSE	0.059	0.017	0.096	0.013	0.014	0.050	0.015	0.013	19.939
RE (LS)									
var	81.8	74.0	80.3	70.9	79.0	85.5	73.7	78.5	99.6
RE (LS)									
mse	83.4	80.1	80.9	72.8	79.1	91.1	74.2	81.6	
n=100									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.043	1.081	0.019	0.860	0.500	1.030	1.019	0.529	91
mean									
LSE	0.045	1.081	0.011	0.866	0.498	1.028	1.020	0.527	90
variance									
MMLE	0.025	0.007	0.038	0.005	0.005	0.020	0.006	0.005	39.732
variance									
LSE	0.031	0.010	0.049	0.007	0.007	0.024	0.008	0.006	40.614
RE (LS)									
var	80.6	72.6	77.8	71.0	75.4	82.4	74.1	74.2	97.8
RE (LS)									
mse	81.2	83.8	78.3	71.5	75.3	89.8	74.9	78.7	

From Tables 3.2.2 and 3.2.3 it is understood that outliers in X (created with respect to Dixon's outlier model) have detrimental effects on the estimators, especially on the LS estimators. Considering the estimated sample sizes, it is clear that outliers increase the cost of observing the sample. Because of those outliers, it is seen that with a fixed budget, one can only observe 84% of the required sample (Table 3.2.3).

Table 3.2.3 Simulation results for complete sample when $X \sim GL$, $Y|X \sim GL$ - Dixon's outlier model (4.0 is added to the r of the X -observations)

n=20									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	-0.033	1.225	0.071	0.819	0.510	1.199	1.046	0.594	16
mean									
LSE	-0.187	1.322	0.032	0.845	0.502	1.194	1.076	0.611	16
variance									
MMLE	0.127	0.034	0.178	0.024	0.019	0.110	0.029	0.019	6.035
variance									
LSE	0.175	0.046	0.217	0.033	0.024	0.127	0.040	0.023	5.656
RE (LS)									
var	72.9	72.7	82.2	71.8	82.2	87.0	73.8	80.9	106.7
RE (LS)									
mse	61.2	56.3	84.1	77.3	82.6	93.5	69.0	77.7	
n=50									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	-0.073	1.233	0.027	0.847	0.504	1.130	1.056	0.587	42
mean									
LSE	-0.196	1.323	0.013	0.856	0.501	1.146	1.084	0.610	41
variance									
MMLE	0.050	0.013	0.070	0.010	0.007	0.047	0.011	0.007	14.928
variance									
LSE	0.068	0.018	0.087	0.014	0.009	0.054	0.016	0.009	13.632
RE (LS)									
var	73.2	69.6	80.3	67.1	79.1	87.3	66.5	76.6	109.5
RE (LS)									
mse	51.7	54.7	81.0	69.1	79.3	86.6	60.1	68.7	
n=100									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	-0.066	1.229	0.003	0.860	0.505	1.108	1.064	0.583	84
mean									
LSE	-0.184	1.318	-0.004	0.862	0.505	1.133	1.090	0.610	82
variance									
MMLE	0.026	0.007	0.033	0.005	0.004	0.022	0.006	0.004	30.391
variance									
LSE	0.036	0.009	0.044	0.007	0.005	0.026	0.008	0.005	27.053
RE (LS)									
var	71.9	72.8	74.7	70.8	75.5	83.7	66.3	79.2	112.3
RE (LS)									
mse	43.4	53.3	74.6	71.2	75.7	80.3	57.9	63.5	

Table 3.2.4 Simulation results for complete sample when $X \sim GL$, $Y|X \sim GL$ - Tiku's outlier model (2.0 is added to the largest r of the X -observations)

n=20									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	-0.103	1.147	0.069	0.822	0.511	1.089	1.024	0.568	18
mean									
LSE	-0.375	1.314	0.037	0.848	0.497	1.089	1.073	0.602	18
variance									
MMLE	0.111	0.031	0.173	0.024	0.021	0.111	0.027	0.019	8.711
variance									
LSE	0.142	0.045	0.203	0.034	0.026	0.128	0.039	0.025	8.345
RE (LS)									
var	78.4	67.6	85.3	70.8	83.5	87.1	70.1	77.6	104.4
RE (LS)									
mse	43.1	36.4	87.1	75.8	83.9	89.7	62.9	67.7	
n=50									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	-0.130	1.156	0.026	0.847	0.507	1.034	1.035	0.565	46
mean									
LSE	-0.390	1.325	0.016	0.855	0.502	1.056	1.084	0.611	45
variance									
MMLE	0.045	0.012	0.062	0.010	0.007	0.044	0.010	0.007	21.704
variance									
LSE	0.061	0.020	0.076	0.014	0.009	0.053	0.015	0.009	20.281
RE (LS)									
var	73.5	61.7	81.2	70.3	80.2	82.6	66.9	75.0	107.0
RE (LS)									
mse	29.0	29.2	81.8	72.3	80.7	79.6	51.2	51.5	
n=100									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	-0.128	1.152	0.022	0.858	0.500	1.015	1.036	0.556	92
mean									
LSE	-0.381	1.321	0.014	0.864	0.498	1.038	1.086	0.605	90
variance									
MMLE	0.022	0.006	0.033	0.005	0.003	0.023	0.005	0.003	39.956
variance									
LSE	0.030	0.010	0.041	0.007	0.004	0.028	0.008	0.005	37.035
RE (LS)									
var	74.2	63.3	79.8	67.6	78.0	84.7	63.9	72.3	107.9
RE (LS)									
mse	21.9	26.0	80.5	68.5	78.0	77.9	41.9	41.9	

Like Tables 3.2.2 and 3.2.3, Tables 3.2.4 and 3.2.5 also exhibit the effects of the X – outliers (created with respect to the Tiku's outlier model) on the means and variances of the estimators. Similar to the previous two tables, it is seen that biases and the variances of the estimators have increased when compared to the ones obtained from the complete samples without outliers. Also, it is observed that MML estimators are more efficient than LS estimators. Additionally, once again the adverse effects of outliers on the cost of observing the required sample

sizes are easily noticed. Estimated sample sizes which are obtained by using MML estimators are nearly 85% of the required sample sizes (Table 3.2.5). However, if LS estimators are used in the estimation procedure this percentage decreases to 82%.

Table 3.2.5 Simulation results for complete sample when $X \sim GL$, $Y|X \sim GL$ - Tiku's outlier model (4.0 is added to the largest r of the X -observations)

n=20									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	-0.287	1.344	0.052	0.824	0.516	1.176	1.089	0.636	17
mean									
LSE	-0.879	1.701	0.022	0.851	0.503	1.189	1.208	0.704	16
variance									
MMLE	0.113	0.030	0.144	0.024	0.012	0.108	0.026	0.013	6.164
variance									
LSE	0.150	0.046	0.169	0.033	0.014	0.125	0.040	0.015	5.662
RE (LS)									
var	75.5	64.7	85.2	72.6	85.3	86.3	64.2	88.6	108.9
RE (LS)									
mse	21.2	27.6	86.6	77.6	87.1	87.9	40.5	56.1	
n=50									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	-0.315	1.341	0.034	0.849	0.503	1.113	1.090	0.619	43
mean									
LSE	-0.890	1.703	0.027	0.856	0.498	1.155	1.206	0.701	41
variance									
MMLE	0.044	0.011	0.056	0.009	0.004	0.044	0.010	0.005	15.517
variance									
LSE	0.059	0.019	0.069	0.013	0.006	0.051	0.017	0.006	13.816
RE (LS)									
var	73.5	60.3	81.7	70.7	78.2	86.1	59.0	88.3	112.3
RE (LS)									
mse	16.8	24.9	82.5	72.2	78.3	76.4	30.2	41.8	
n=100									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	-0.322	1.341	0.008	0.857	0.504	1.087	1.094	0.618	86
mean									
LSE	-0.897	1.709	0.010	0.858	0.501	1.139	1.213	0.706	82
variance									
MMLE	0.021	0.006	0.028	0.005	0.002	0.020	0.005	0.003	29.902
variance									
LSE	0.029	0.010	0.034	0.007	0.003	0.025	0.009	0.003	26.209
RE (LS)									
var	72.3	60.4	80.6	68.7	80.8	81.5	61.2	86.2	114.1
RE (LS)									
mse	15.0	23.9	80.6	69.3	81.5	66.2	26.2	36.5	

In Table 3.2.6, the simulation results are given for censored samples without outliers. Comparing the results with the ones obtained from the complete

sample without outliers, it is easily noticed that the performances of the estimators are similar. As expected and seen in the previous sections, censoring has increased the bias and the variance, but in a small amount. Also, considering the estimated sample sizes, it is observed that they increased substantially.

Table 3.2.6 Simulation results for censored sample when $X \sim GL$, $Y|X \sim GL$ - without outliers

n=20									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.060	0.964	0.092	0.847	0.501	1.005	1.004	0.467	26
variance									
MMLE	0.116	0.033	0.266	0.029	0.072	0.126	0.038	0.047	26.353
n=50									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.030	0.981	0.033	0.876	0.497	0.943	1.015	0.476	65
variance									
MMLE	0.044	0.013	0.094	0.012	0.026	0.050	0.015	0.018	52.692
n=100									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.017	0.993	0.015	0.887	0.498	0.928	1.021	0.481	130
variance									
MMLE	0.022	0.006	0.050	0.006	0.012	0.027	0.008	0.008	98.365

Table 3.2.7 Simulation results for censored sample when $X \sim GL$, $Y|X \sim GL$ - Dixon's outlier model (2.0 is added to the r of the X -observations)

n=20									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.101	1.048	0.083	0.844	0.504	1.104	1.024	0.502	24
variance									
MMLE	0.125	0.036	0.276	0.029	0.060	0.126	0.040	0.042	16.431
n=50									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.057	1.074	0.011	0.878	0.507	1.039	1.044	0.517	59
variance									
MMLE	0.051	0.015	0.096	0.011	0.021	0.051	0.015	0.016	36.861
n=100									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.045	1.081	0.005	0.885	0.502	1.022	1.044	0.516	118
variance									
MMLE	0.027	0.007	0.051	0.006	0.011	0.026	0.008	0.008	72.491

Table 3.2.8 Simulation results for censored sample when $X \sim GL$, $Y|X \sim GL$ - Dixon's outlier model (4.0 is added to the r of the X -observations)

n=20									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.017	1.154	0.066	0.848	0.509	1.154	1.058	0.543	23
variance									
MMLE	0.131	0.043	0.252	0.029	0.047	0.129	0.042	0.036	15.295
n=50									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	-0.039	1.189	0.024	0.873	0.502	1.099	1.068	0.554	57
variance									
MMLE	0.055	0.019	0.093	0.012	0.017	0.049	0.017	0.014	31.419
n=100									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	-0.058	1.203	0.011	0.884	0.501	1.087	1.075	0.558	113
variance									
MMLE	0.029	0.012	0.045	0.006	0.008	0.026	0.008	0.007	64.748

Table 3.2.9 Simulation results for censored sample when $X \sim GL$, $Y|X \sim GL$ - Tiku's outlier model (2.0 is added to the largest r of the X -observations)

n=20									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.074	0.957	0.094	0.844	0.500	1.006	1.000	0.466	26
variance									
MMLE	0.118	0.034	0.274	0.030	0.073	0.125	0.039	0.047	25.258
n=50									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.021	0.990	0.024	0.877	0.499	0.941	1.018	0.480	65
variance									
MMLE	0.043	0.014	0.098	0.012	0.026	0.051	0.015	0.018	50.465
n=100									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.009	0.993	0.014	0.883	0.496	0.920	1.018	0.481	131
variance									
MMLE	0.024	0.007	0.051	0.006	0.014	0.025	0.008	0.010	111.474

Tables 3.2.7 - 3.2.10 show the simulation results for censored X , when the sample contains outliers. Accordingly, it is obvious that censoring improves the performance of the estimators, when there are outliers in the sample. Also, when the results obtained from complete and censored samples are compared, once

again it is observed that censoring results in saving a great amount of money. It can be understood from the corresponding tables, because in those tables estimated sample sizes are more than the actual sample sizes.

Table 3.2.10 Simulation results for censored sample when $X \sim GL$, $Y|X \sim GL$ - Tiku's outlier model (4.0 is added to the largest r of the X -observations)

n=20									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.075	0.960	0.093	0.849	0.498	1.004	1.003	0.462	26
variance									
MMLE	0.114	0.034	0.275	0.029	0.075	0.122	0.038	0.049	25.699
n=50									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.021	0.988	0.031	0.875	0.501	0.948	1.017	0.481	65
variance									
MMLE	0.046	0.014	0.099	0.012	0.025	0.052	0.015	0.018	57.996
n=100									
	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
mean									
MMLE	0.005	0.998	0.007	0.885	0.500	0.924	1.022	0.485	130
variance									
MMLE	0.024	0.007	0.048	0.006	0.012	0.026	0.008	0.009	108.679

Summarizing the results of Section 3.2, it can be said that MML estimators are more efficient than LS estimators. Also, the efficiency of MML estimators increases as sample size increases and when the sample contains outliers. Furthermore, censoring a few largest X -observations leads to advantageous results. In other words, from the simulations it is understood that censoring never increases the cost of observing the required sample sizes, on the contrary, it decreases especially when there are outliers in the sample.

3.3. Accuracy of the Fisher Information Matrix

In Tables 3.3.1 – 3.3.6, the minimum variance bounds (MVB) of the estimators are displayed for both complete and censored samples when marginal and conditional distributions are generalized logistic and normal, respectively (See Chapter 2 for the derivations). In Table 3.3.1, the MVB's are obtained by

using Fisher information matrix. However, for censored samples since it is difficult to take the expectation of concomitant terms, instead of obtaining Fisher information matrix, the sample information matrices are used. Therefore, MVB's given in Tables 3.3.2 – 3.3.6 are based on sample information matrices.

Table 3.3.1 MVB's for complete sample when $X \sim GL$, $Y|X \sim normal$

	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ
n=20	0.112	0.032	0.103	0.019	0.019
n=50	0.045	0.013	0.041	0.007	0.008
n=100	0.022	0.006	0.021	0.004	0.004

Table 3.3.2 MVB's for censored sample when $X \sim GL$, $Y|X \sim normal$ - without outliers

	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ
n=20	0.109	0.033	0.134	0.019	0.041
n=50	0.045	0.013	0.053	0.008	0.016
n=100	0.023	0.007	0.024	0.004	0.008

Table 3.3.3 MVB's for censored sample when $X \sim GL$, $Y|X \sim normal$ - Dixon's outlier model (2.0 is added to the r of the X -observations)

	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ
n=20	0.109	0.034	0.134	0.019	0.041
n=50	0.045	0.014	0.054	0.008	0.016
n=100	0.022	0.007	0.026	0.004	0.008

Table 3.3.4 MVB's for censored sample when $X \sim GL$, $Y|X \sim normal$ - Dixon's outlier model (4.0 is added to the r of the X -observations)

	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ
n=20	0.108	0.034	0.135	0.019	0.041
n=50	0.045	0.014	0.053	0.008	0.016
n=100	0.023	0.007	0.026	0.004	0.008

As it is seen from the tables, the simulated variances of the estimators are close to MVB values and also they are getting closer as sample size increases.

Table 3.3.5 MVB's for censored sample when $X \sim GL$, $Y|X \sim \text{normal}$ - Tiku's outlier model (2.0 is added to the largest r of the X -observations)

	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ
n=20	0.110	0.033	0.132	0.019	0.041
n=50	0.045	0.014	0.053	0.008	0.016
n=100	0.022	0.007	0.027	0.004	0.008

Table 3.3.6 MVB's for censored sample when $X \sim GL$, $Y|X \sim \text{normal}$ - Tiku's outlier model (4.0 is added to the largest r of the X -observations)

	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ
n=20	0.109	0.033	0.133	0.018	0.041
n=50	0.045	0.014	0.054	0.008	0.016
n=100	0.023	0.007	0.027	0.004	0.008

The MVB's of the estimators when both marginal and conditional distributions are generalized logistic are given in Tables 3.3.7 – 3.3.12. Similar to the previous case, for complete sample the MVB's are obtained from the Fisher information matrix (Table 3.3.7), while for censored samples they are obtained by using the sample information matrices (Table 3.3.8 – 3.3.12).

From the results it is once again understood that it will give no harm to use the simulated variances instead of asymptotic variances.

Table 3.3.7 MVB's for complete sample when $X \sim GL$, $Y|X \sim GL$

	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ
n=20	0.112	0.032	0.182	0.024	0.029
n=50	0.045	0.013	0.073	0.010	0.012
n=100	0.022	0.006	0.036	0.005	0.006

Table 3.3.8 MVB's for censored sample when $X \sim GL$, $Y|X \sim GL$ - without outliers

	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ
n=20	0.102	0.044	0.238	0.026	0.065
n=50	0.043	0.018	0.093	0.010	0.025
n=100	0.022	0.008	0.047	0.005	0.012

Table 3.3.9 MVB's for censored sample when $X \sim GL$, $Y|X \sim GL$ - Dixon's outlier model (2.0 is added to the r of the X -observations)

	μ_1	σ_1	$\mu_{2,1}$	$\sigma_{2,1}$	θ
n=20	0.122	0.040	0.255	0.029	0.056
n=50	0.051	0.017	0.096	0.011	0.021
n=100	0.025	0.007	0.050	0.006	0.011

Table 3.3.10 MVB's for censored sample when $X \sim GL$, $Y|X \sim GL$ - Dixon's outlier model (4.0 is added to the r of the X -observations)

	μ_1	σ_1	$\mu_{2,1}$	$\sigma_{2,1}$	θ
n=20	0.141	0.052	0.238	0.029	0.046
n=50	0.058	0.025	0.092	0.012	0.017
n=100	0.032	0.014	0.045	0.006	0.008

Table 3.3.11 MVB's for censored sample when $X \sim GL$, $Y|X \sim GL$ - Tiku's outlier model (2.0 is added to the largest r of the X -observations)

	μ_1	σ_1	$\mu_{2,1}$	$\sigma_{2,1}$	θ
n=20	0.102	0.040	0.242	0.025	0.066
n=50	0.043	0.016	0.093	0.010	0.025
n=100	0.023	0.009	0.047	0.005	0.012

Table 3.3.12 MVB's for censored sample when $X \sim GL$, $Y|X \sim GL$ - Tiku's outlier model (4.0 is added to the largest r of the X -observations)

	μ_1	σ_1	$\mu_{2,1}$	$\sigma_{2,1}$	θ
n=20	0.104	0.040	0.238	0.025	0.065
n=50	0.042	0.016	0.093	0.010	0.025
n=100	0.022	0.008	0.047	0.005	0.012

3.4. Simulated Powers of the Test Statistic Z

In section 3.4, the simulated powers of the test statistic, Z , presented in (2.4.1.3.1) are given for complete samples when marginal and conditional distributions are Generalized Logistic. Also, the effects of the outliers in the marginal distribution X are evaluated. The outliers are created according to the Dixon's outlier model, by multiplying the randomly selected 10% X values by 4.0. b_1 and b_2 denote the shape parameters for marginal and conditional distributions.

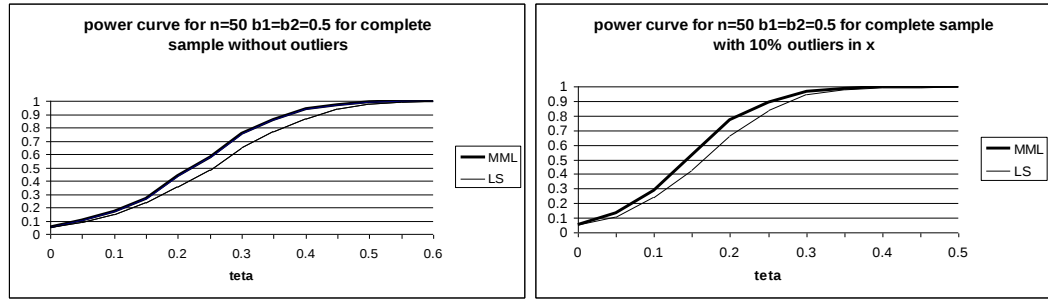


Figure 3.4.1: Power curves for complete sample, without outliers and with 10% outliers in X , for $b_1=b_2=0.5$, $n=50$.

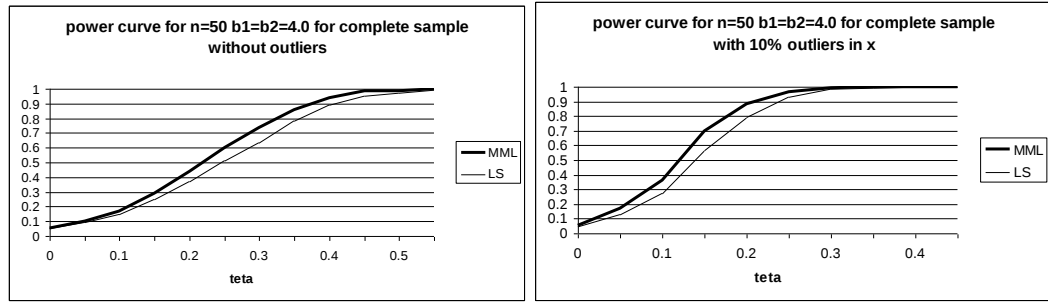


Figure 3.4.2: Power curves for complete samples, without outliers and with 10% outliers in X , for $b_1=b_2=4.0$, $n=50$.

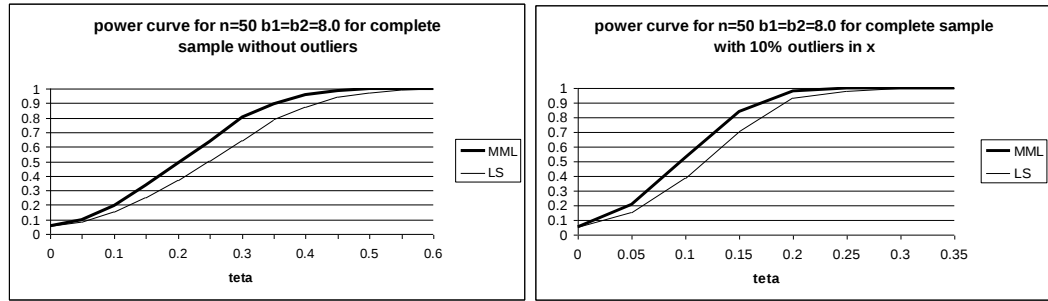


Figure 3.4.3: Power curves for complete sample, without outliers and with 10 % outliers in X , for $b_1=b_2=8.0$, $n=50$.

From the Figures 3.4.1 - 3.4.3, it is easily seen that test statistic (for testing $H_0 : \theta = 0$ against $H_1 : \theta > 0$) obtained by using MML estimators has higher power than the test statistic obtained by using LS estimators. This situation is true for both complete samples with and without outliers. This was to be expected because the MML estimators are more efficient than the LS estimators; see Sundrum (1954).

3.5. Illustrative Examples

3.5.1. Real Life Example 1

Woolson (1981) did research about a psychiatric disease and psychiatric inpatients who are admitted to the University of Iowa hospitals in the period 1935 – 1948 are taken into consideration. The information included in the data are age of the patients when they are at first admitted to the hospital, sexuality and number of years of follow-up. Then, this data is used in order to see whether that psychiatric disease has an effect on the lifetimes of the patients.

The years from admission to death (X) and the age of death (Y) are given in Table 3.5.1.1. Also, it is obvious that as the follow up time of the patients increases, the cost of the research rises.

Table 3.5.1.1. Woolson data

i	1	2	3	4	5	6	7
X	1	1	2	22	28	32	11
Y	52	59	57	50	47	57	59

i	8	9	10	11	12	13	14
X	14	25	22	26	24	35	40
Y	61	61	63	69	69	67	76

It is reasonable to consider the distribution of X as Generalized Logistic (with shape parameter equal to 0.5) and Y as normal. This is due to the fact that $b=0.5$ maximizes $(1/n)\ln \hat{L}$, $\hat{L}$ is the likelihood function evaluated at $\mu_1 = \hat{\mu}_1$ and $\sigma_1 = \hat{\sigma}_1$. The Q-Q plots of the variables are given in Figure 3.5.1.1 – 3.5.1.2.

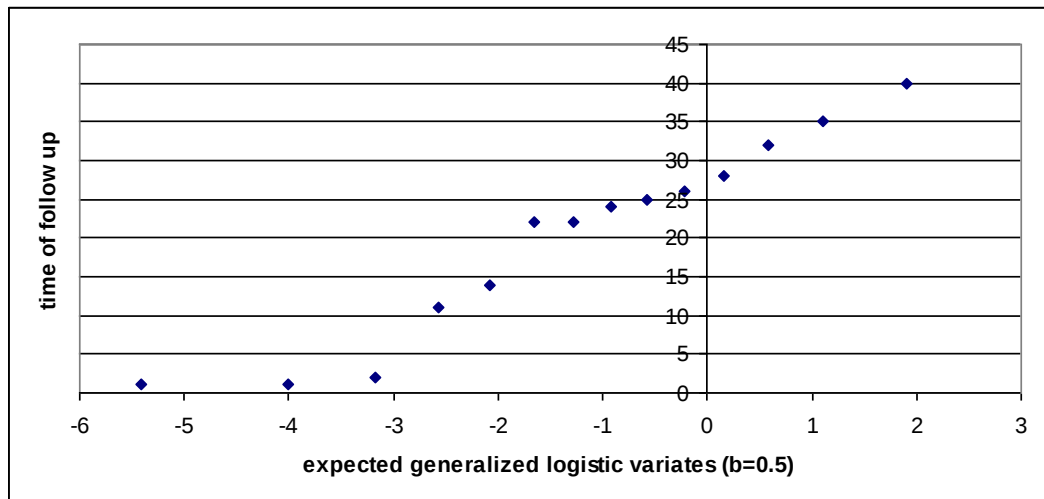


Figure 3.5.1.1. Q-Q plot of X when $b=0.5$

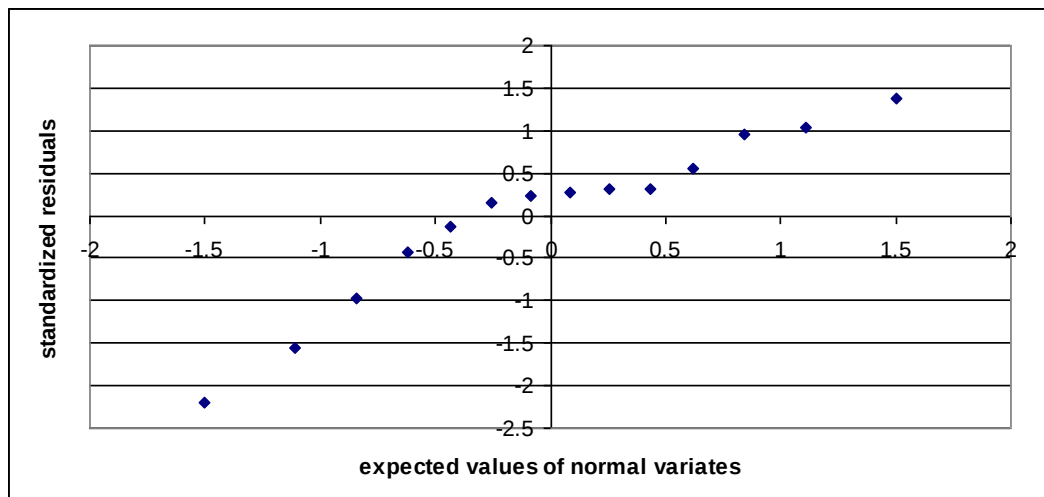


Figure 3.5.1.2. Q-Q plot of the standardized residuals for Woolson data

The MML and LS estimators and the estimated sample sizes for the complete sample are given in Table 3.5.1.2. Also, the table includes the results after censoring the highest X and the corresponding Y observations. Although X has no outliers, it is seen that censoring has minor positive effects on the available sample size.

Table 3.5.1.2. MML and LS estimators for Woolson data

	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
Complete sample									
MMLE	27.551	5.225	54.743	6.860	0.285	60.526	7.019	0.212	Co/(c+20.31)
LSE	29.883	6.975	54.743	7.409	0.285	60.500	6.587	0.452	Co/(c+20.21)
Censored sample ($n-r=13$)									
MMLE	27.310	5.201	56.353	6.369	0.158	59.530	6.422	0.128	Co/(c+19.28)

The parametric bootstrap variances of the estimators before and after censoring are given in Table 3.5.1.3. It is seen that variances of the estimators increase after censoring, but not excessively. Therefore, from Tables 3.5.1.2 and 3.5.1.3, it can be concluded that censoring does no harm. In fact, it reduces the cost substantially because the largest observation is the most expensive to observe. This coincides with the results obtained in Section 3.1 and 3.2. Also, from complete sample results, it is seen that MML estimators are much more efficient than the LS estimators.

Table 3.5.1.3. Parametric bootstrap variances of the estimators for Woolson data

	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ	n
Complete sample									
var MMLE	10.223	1.552	15.490	1.676	0.027	4.401	1.607	0.018	8.208
var LSE	12.112	3.658	15.490	1.955	0.027	4.461	1.541	0.053	9.083
RE(LS)	84.4	42.4	100.0	85.7	100.0	98.7	104.3	34.8	90.4
Censored sample ($n-r=13$)									
var MMLE	10.345	1.597	17.567	1.812	0.036	4.582	1.733	0.023	12.35

The similarity in the performance of MML and LS estimators coming from the conditional distribution is because of the fact that conditional distribution is normal.

3.5.2. Real Life Example 2

In the following data obtained from Gross and Clark (1975), U represents 100 times the white blood counts and Y represents the survival times (in weeks) of patients who died of acute myelogenous leukemia.

Table 3.5.2.1. Gross and Clark data

i	1	2	3	4	5	6	7	8
U	7.5	23	26	43	54	60	70	94
Y	156	65	134	100	39	16	143	56

i	9	10	11	12	13	14	15	16
U	100	105	170	320	350	520	1000	1000
Y	121	108	4	26	22	5	1	1

In Tiku and Vaughan (2000) it is mentioned that the seventeenth data point $(U, Y) = (1000, 65)$ is an outlier and therefore it is removed. Also, the distribution of U is considered as Weibull with $p=0.8$ and the distribution of Y is taken as normal. This is also understood from the Q-Q plots of U and the standardized residuals given in Figures 3.5.2.1 – 3.5.2.2.

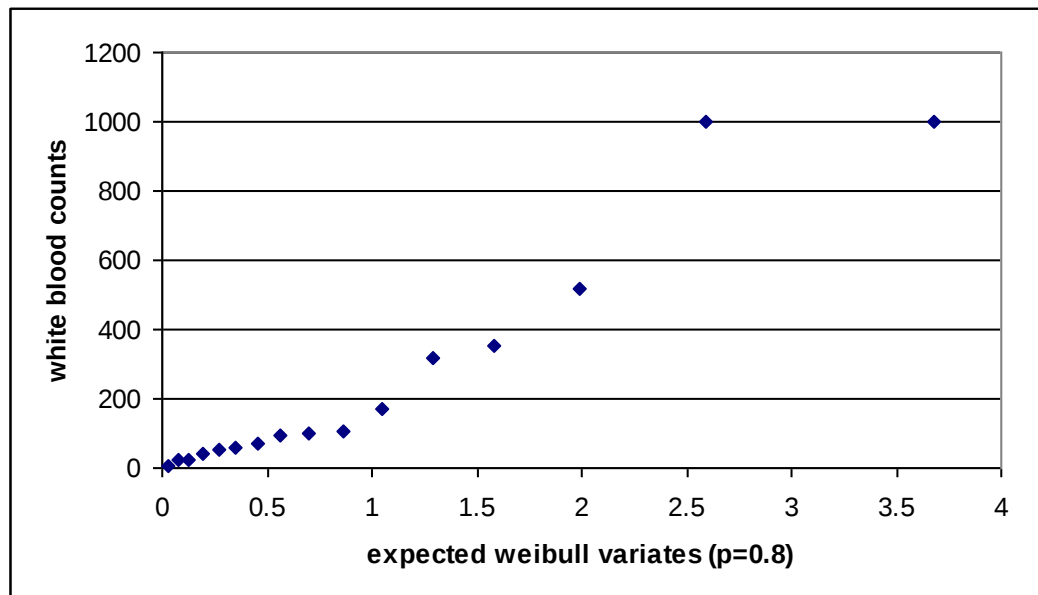


Figure 3.5.2.1. Q-Q plot of U when $p=0.8$

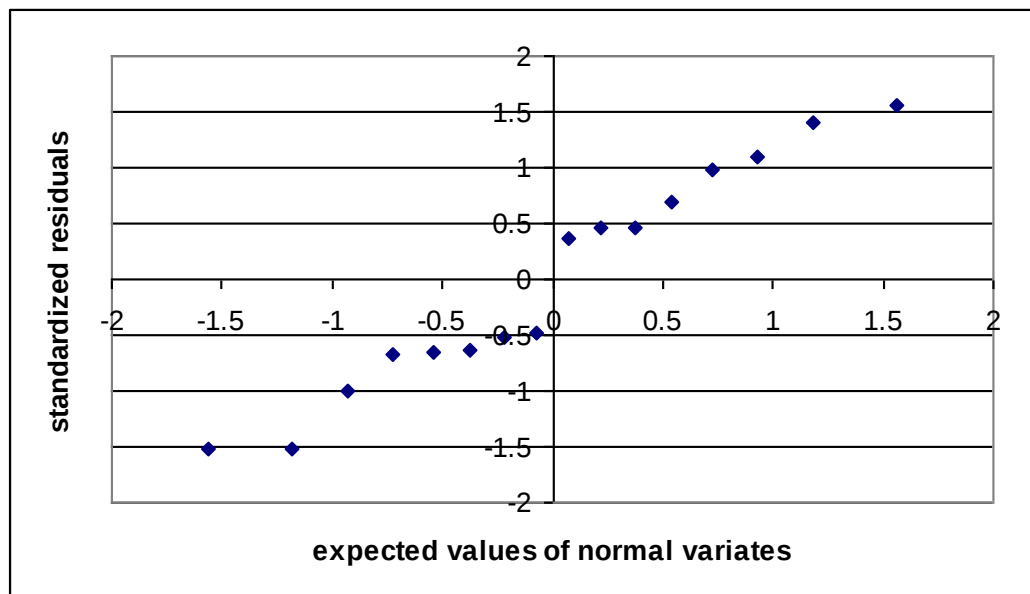


Figure 3.5.2.2. Q-Q plot of the standardized residuals for Gross and Clark data

MML and LS estimators for complete sample and MML estimators for censored sample (after censoring the highest U and the corresponding Y observations) are given in the Table 3.5.2.2. From the table it is seen that censoring has not changed the estimates so widely.

Table 3.5.2.2. MML and LS estimators for Gross and Clark data

	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ
Complete sample							
MMLE	214.392	88.870	42.359	-0.108	62.690	48.252	-0.479
LSE	228.479	88.870	45.284	-0.108	62.313	47.303	-0.627
Censored sample ($n-r=15$)							
MMLE	232.404	90.710	43.226	-0.124	58.079	51.942	-0.555

The parametric bootstrap variances of the estimators for complete and censored samples are given in Table 3.5.2.3. From the results it is seen that for complete sample MML estimators are more efficient than LS estimators. Also, it is observed that censoring increases the variances of the estimators in some amount.

Table 3.5.2.3. Parametric bootstrap variances of the estimators for Gross and Clark data

	σ_1	$\mu_{2,1}$	$\sigma_{2,1}$	θ	μ_2	σ_2	ρ
Complete sample							
var MMLE	4569.54	218.11	54.82	0.003	178.37	77.10	0.04
var LSE	6945.60	218.11	62.65	0.003	177.99	71.91	0.05
RE (LS)	65.8	100.0	87.5	100.0	100.2	107.2	81.6
Censored sample ($n-r=15$)							
var MMLE	4853.14	271.36	60.07	0.006	201.05	106.65	0.07

Note that since the conditional distribution is normal the MML and LS estimators obtained from that distribution perform similarly.

3.5.3. Real Life Example 3

In the literature there exist so many studies that investigate whether the dividend payment of a company has an effect on its market capitalization. In the table below, data of the 24 companies which are included in the index ISE100 (Istanbul Stock Exchange) and which paid dividend for the year 2005 is given. X denotes the dividend ratio (total amount of dividend paid/total profit) while Y denotes the annual change in the market capitalization of the firms between the years 2005 – 2006. Data is obtained from ISE website. Another related study is done by Topcu (2008). In her study, she investigated whether the dividend payment of a company has an effect on the price of the stocks.

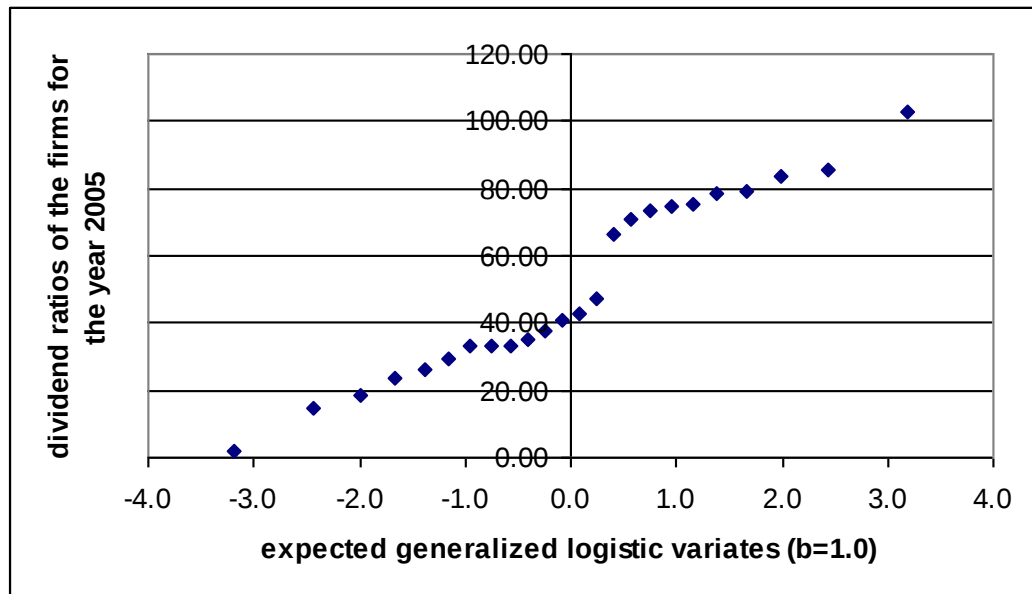
Table 3.5.3.1. Firm data

i	1	2	3	4	5	6	7	8
X	75.39	33.02	37.54	79.01	1.87	29.17	83.67	40.98
Y	11.91	15.89	-4.44	19.21	-30.91	-22.24	62.31	-4.71
i	9	10	11	12	13	14	15	16
X	33.19	43.04	47.50	14.82	26.47	33.42	102.61	66.50
Y	23.50	20.88	0.56	-4.49	-19.46	-29.48	72.73	17.39

Table 3.5.3.1. Firm data (continued)

<i>i</i>	17	18	19	20	21	22	23	24
<i>X</i>	70.90	18.27	74.94	23.81	85.32	35.25	78.22	73.64
<i>Y</i>	-22.14	8.82	3.42	6.74	59.02	72.54	-4.29	-2.42

It is appropriate to consider the distributions of X and Y as generalized logistic with shape parameter equal to 1.0 and normal, respectively. Note that the value $b=1.0$ maximizes $(1/n)\ln \hat{L}$. Also, see Figure 3.5.3.1 – 3.5.3.2 for the Q-Q plots.

**Figure 3.5.3.1.** Q-Q plot of X when $b=1.0$

From its Q-Q plot, X seems to have no outliers in it. The MML and LS estimators for both complete and censored samples are given in Table 3.5.3.2. In the censoring procedure the highest two X - observations and the corresponding Y values are removed.

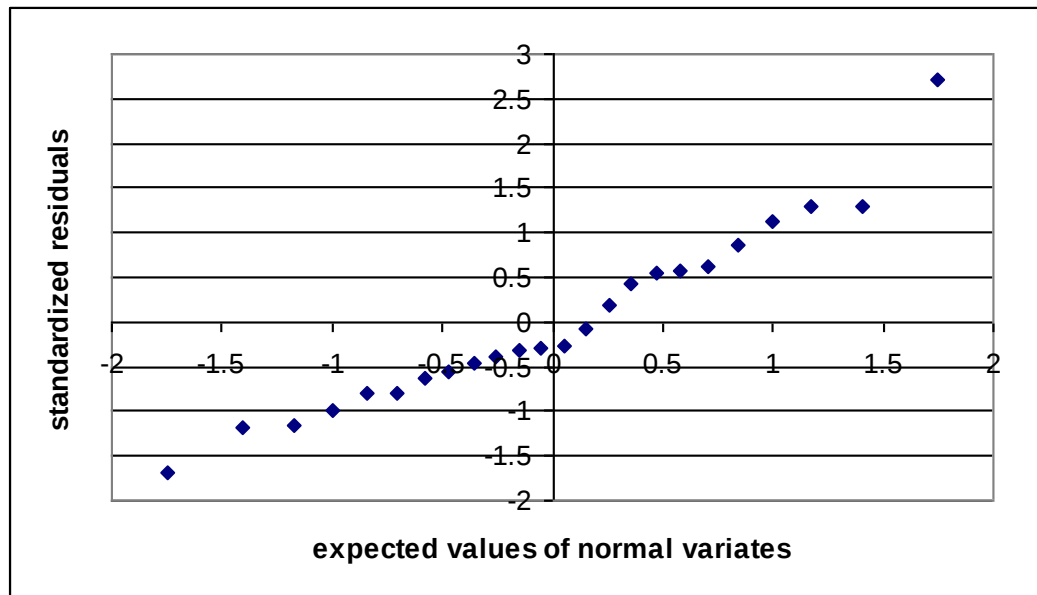


Figure 3.5.3.2. Q-Q plot of the standardized residuals for firm data

Table 3.5.3.2. MML and LS estimators for firm data

	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ
Complete sample								
MMLE	49.803	16.337	-17.345	25.516	0.552	10.126	27.061	0.333
LSE	50.356	14.948	-17.345	26.651	0.552	10.431	24.005	0.498
Censored sample ($n-r=22$)								
MMLE	49.898	16.768	-9.078	24.067	0.312	6.484	24.628	0.212

In Table 3.5.3.3 the parametric bootstrap variances of the estimators are presented. Similar to the first real life example, it is seen that censoring increases the variances of the estimators (other than $\hat{\mu}_1$), because there are no outliers in X . Also, MML estimators are more efficient than LS estimators.

Note that since the conditional distribution is normal, the MML and LS estimators obtained from that distribution perform similarly.

Table 3.5.3.3. Parametric bootstrap variances of the estimators for the firm data

	μ_1	σ_1	$\mu_{2.1}$	$\sigma_{2.1}$	θ	μ_2	σ_2	ρ
Complete sample								
Var MMLE	33.549	7.948	116.345	13.244	0.036	37.154	12.765	0.015
Var LSE	35.972	8.513	116.345	14.448	0.036	37.853	11.444	0.025
RE (LS)	93.3	93.4	100.0	91.7	100.0	98.2	111.5	60.0
Censored sample ($n-r=22$)								
Var MMLE	33.129	8.538	145.918	14.500	0.055	40.772	14.421	0.021

CHAPTER 4

CONCLUSION

Let X and Y be random variables and Y depends on $X=x$, which is a very common situation in many real life applications and in these situations non-normal situations occur frequently. However, when the marginal or conditional distributions are non-normal, ML estimation technique leads to problems. Therefore, statistical methods which give efficient results under non-normal distributions are needed. For these cases, using Tiku's MML estimation technique, which linearizes the intractable terms in the likelihood equations, is very useful. Note that the MML estimators are easy to compute and are explicit functions of sample observations. Also, they have the same forms irrespective of the underlying marginal and conditional distributions.

The aim of this thesis is to evaluate the effect of the largest order statistics $x_{(i)}$ ($i \geq n-r$) in a random sample of size n

- (i) on the mean $E(X)$ and variance $V(X)$ of X ,
- (ii) on the cost of observing the x -observations, and
- (iii) on the conditional mean $E(Y|X=x)$ and variance $V(Y|X=x)$,
- (iv) on the regression coefficient

for the distributions coming from p -family, Weibull and Generalized Logistic. It is shown that unduly large x -observations have detrimental effects on (i)-(iv). The advantages of not observing a few largest observations are also evaluated.

In the simulations, the case when marginal distribution is Generalized Logistic and conditional distribution is normal, also the case when both marginal and conditional distributions are Generalized Logistic are taken into consideration for both complete and censored samples. But, the simulations can also be done for other location – scale distributions. In the censoring situation, samples from the assumed marginal distribution are censored and r (censoring number) is to be $\text{int}[0.10*n]$. In the simulation study, the effects of the largest observations (in the marginal distribution) on the sample size and on the estimators are evaluated. Outliers are created according to Tiku's and Dixon's outlier model.

A cost constraint for observing $X=x$ is defined and it is assumed that the larger x is the more expensive it becomes. In the simulations C_0 , which is the total permissible cost, is taken as $nc + n\mu_1 + \sigma_1 \sum_{i=1}^n t_{(i)}$. It is assumed that $c + \mu_1 + \sigma_1 t_{(1)} > 0$, as expected. This is essentially saying that the cost of observing the i^{th} ordered observation is directly related to its expected value.

In the complete sample simulations, the performance of the MML estimators are compared with the bias corrected LS estimators by looking at their relative efficiencies. Note that no bias correction is applied to MML estimators since the sample sizes are large enough and MML estimators are self bias correcting. Also, it is seen that simulated variances of the MML estimators are close to the MVB's even for small n .

From the simulations it is understood that the outliers in X increase the cost of observing the required sample sizes, in other words they make it impossible to observe the required sample sizes subject to fixed cost. Besides, the estimators become biased and they have larger variances when compared to the estimators obtained from the complete sample without outliers. Also, it can be said with confidence that the MML estimators perform better than the LS

estimators. The performance gap between the estimators increases as sample size increases and when the sample contains outliers.

The other important result is that not observing a few largest x -observations leads to advantageous results, especially when the sample contains outliers. It is seen that the performance of the estimators, obtained from the censored sample without outliers bear a strong resemblance to the performance of the estimators, obtained from the complete sample without outliers. The increase in the bias and the variance is in a small amount. When the results obtained from complete and censored samples are compared, it is seen that estimated sample sizes increase substantially. In other words it can be said that censoring results in saving a great amount of money.

Finally from the simulations it is easily seen that test statistic obtained by using MML estimators has higher power than the test statistic obtained by using LS estimators.

REFERENCES

- Abramowitz, M. and Stegun, I. A. (1985). *Handbook of Mathematical Functions*. Dover: New York.
- Agresti, A. (1996). *Categorical Data Analysis*. John Wiley: New York.
- Aitken, M., Anderson, D., Francis, B. and Hinde, J. (1989). *Statistical Modelling In Glim*. Oxford Science: New York.
- Akkaya, A. D. and Tiku, M.L. (2001). Corrigendum: Time series models with asymmetric innovations. *Commun. Stat. – Theory Meth.*, 30, 2227 – 2230.
- Akkaya, A. D. and Tiku, M.L. (2008). Robust estimation in multiple linear regression model with non-Gaussian noise. *Automatica*, 44, 407 – 417.
- Balakrishnan, N. and Leung, W. Y. (1988). Means, variances and covariances of order statistics, BLUE for the Type I generalized logistic distribution and some applications. *Commun. Stat. – Simula.*, 17, 51 – 84.
- Barnett, V. D. (1966). Evaluation of the maximum likelihood estimator when the likelihood equation has multiple roots. *Biometrika*, 53, 151 – 165 .
- Berkson, J. (1951). Why I prefer logits to probits. *Biometrics*, 7, 327 – 339.
- Cohen, A. C. (1991). *Truncated and Censored Samples, Theory and Applications*. Marcel Dekker: New York.
- Cohen, A. C. (1957). On the solution of estimating equations from truncated and censored samples from normal populations. *Biometrika*, 44, 225- 236.

- Cohen, A. C. and Whitten, B. (1982). Modified maximum likelihood and modified moment estimators for the three-parameter Weibull distribution. *Commun. Stat. – Theory Meth.*, 11, 2631 – 2656.
- Cox, D. R. (1972). Regression models and life-tables. *J. R. Statist. Soc. B*, 34, 187 – 202 .
- Cox, D. R. (1975). Partial likelihood. *Biometrika*, 62, 269 – 276.
- Dixon, W. J. (1950). Analysis of extreme values. *Ann. Math. Stat.*, 21, 488 – 506.
- Dixon, W. J. (1953). Processing data for outliers. *Biometrics*, 22, 74 -89.
- Gross, A. J. and Clark, V. A. (1975). *Survival Distributions*. John Wiley: New York.
- Gupta, A. K. (1952). Estimation of the mean and standard deviation of a normal population from a censored sample. *Biometrika*, 39, 260 – 273.
- Hand, D. J., Daly, F., Lunn, A. D., McConway, K. J. and Ostrowski, E. (1994). *Small Data Sets*. Chapman & Hall: New York.
- Harter, H. L. (1964). Exact confidence bounds, based on one order statistic, for the parameters of an exponential population. *Technometrics*, 6, 301 – 317.
- Hawkins, D. M. (1977). Comment on ‘A new statistic for testing suspected outliers’. *Commun. Stat.*, A6 (5), 435 – 438.
- Hosmer, D. W. and Lemeshow, S. (1989). *Applied Logistic Regression*. John Wiley: New York.
- Hotelling, H. (1931). The generalization of Student's ratio. *Ann. Math. Statist.*, 2, 360 – 378.
- Islam M. Q. and Tiku, M. L. (2004). Multiple linear regression model under non-normality. *Commun. Stat. – Theory Methods*, 33, 2443 – 2467.

- Islam M. Q., Tiku M. L. and Yildirim, F. (2001). Nonnormal regression I. skew distributions. *Commun. Stat.-Theory Meth.*, 30(6), 993 – 1020.
- Johnson, R. C. and Johnson, N. L. (1979). *Survival Models and Data Analysis*. John Wiley: New York.
- Lawless, J. F. (1982). *Statistical Models and Methods for Life Data*. John Wiley: New York.
- Lee, K. R., Kapadia, C. H. and Dwight, B.B. (1980). On estimating the scale parameter of Rayleigh distribution from censored samples. *Statist. Hefte*, 21, 14 – 20.
- Menon, M. V. (1963). Estimation of the shape and scale parameters of the Weibull distribution. *Technometrics*, 5, 175 – 182.
- Puthenpura, S. and Sinha, N. K. (1986). Modified maximum likelihood method for the robust estimation of system parameters from very noisy data. *Automatica*, 22, 231 – 235.
- Qumsiyeh, S. B. (2007). Non-normal bivariate distributions: estimation and hypothesis testing. *Ph. D. Thesis, Faculty of Art and Science, Department of Statistics, Middle East Technical University, Turkey*.
- Sarhan, A. E. and Greenberg, B. G., eds. (1962). *Contributions to Order Statistics*, John Wiley: New York.
- Sazak, H. S. (2003). Estimation and hypothesis testing in stochastic regression. *Ph. D. Thesis, Faculty of Art and Science, Department of Statistics, Middle East Technical University, Turkey*.
- Schneider (1986). *Truncated and Censored Samples from Normal Populations*. Marcel Dekker: New York.
- Sazak, H. S., Tiku, M. L. and Islam, M. Q. (2006). Regression Analysis with a Stochastic Design Variable. *International Statistical Review*, 74, 1, 77 - 88.

- Smith, R. L. (1985). Maximum likelihood estimation in a class of nonregular cases. *Biometrika*, 72, 67 -90.
- Sundrum, R. M. (1954). On the relation between estimating efficiency and the power of the test. *Biometrika*, 41, 542 – 548.
- Şenoglu, B. and Tiku, M. L. (2001). Analysis of variance in experimental design with nonnormal error distributions. *Commun. Statist.-Theor. Meth.*, 30(7), 1335 – 1352.
- Şenoglu, B. and Tiku, M. L. (2002). Linear contrasts in experimental design with nonidentical error distributions. *Biometrical Journal*, 44(3), 359 – 374.
- Şenoğlu, B. and Tiku, M. L. (2004). Censored and truncated samples in experimental design under non-normality. *Statistical Methods*, 6(2), 173 – 199.
- Thode, H. C. (2002). *Testing for Normality*. Marcel Dekker: New York.
- Tiku, M. L. (1967). Estimating the mean and the standard deviation from a censored normal sample. *Biometrika*, 54, 155 – 165.
- Tiku, M. L. (1968). Estimating the parameters of normal and logistic distributions from censored samples. *Austral. J. Statist.*, 10, 64-74.
- Tiku, M. L. (1975). A new statistic for testing suspected outliers, *Commun. Stat.*, 4 (8), 737 – 752.
- Tiku, M. L. (1977). Rejoinder: “Comment on ‘A new statistic for testing suspected outliers’”, *Commun. Stat – Theory and Methods*, A6 (14), 1417 – 1422.
- Tiku, M. L. (1980). Robustness of MML estimators based on censored samples and robust test statistics. *J. Stat. Plann. Inf.*, 4, 123 – 143.
- Tiku, M. L. (1981). Testing equality of location parameters of two exponential distributions. *Aligarh J. Statist.*, 1, 1 – 7 (invited paper).

- Tiku, M. L., Tan, W. Y. and Balakrishnan, N. (1986). *Robust Inference*, Marcel Dekker, New York.
- Tiku, M. L. and Suresh, R. P. (1992). A new method of estimation for location and scale parameters. *J. Stat. Plann. Inf.*, 30, 281 – 292.
- Tiku, M. L. and Vaughan, D. C. (1997). Logistic and nonlogistic density functions in binary regression with nonstochastic covariates. *Biometrical J.*, 39, 883 – 898.
- Tiku, M. L., Islam, M. Q. and Selcuk, A. S. (2001). Nonnormal regression II. symmetric distributions, *Commun. Stat.-Theory Meth.*, 30(6), 1021 – 1045.
- Tiku, M. L. and Akkaya, A. D. (2004). *Robust Estimation and Hypothesis Testing*. New Age International Publishers (Wiley Eastern): New Delhi.
- Topcu, A. (2008). Temettü Ödemelerinin Hisse Senedi Fiyatlarına Etkisi, Sermaye Piyasası Kurulu Araştırma Raporu No. AT 2008/2.
- Vaughan, D. C. (1992). On the Tiku-Suresh method of estimation. *Commun. Stat.-Theory Meth.*, 21 (2), 451-469.
- Vaughan, D.C. (2002). The generalized secant hyperbolic distributions and its properties. *Commun. Stat.-Theory Meth.*, 31, 219 – 238.
- Vaughan, D. C. and Tiku, M. L. (2000). Estimation and hypothesis testing for a non-normal bivariate distribution with applications. *J. Mathematical and Computer Modelling*, 32, 53 – 67.
- Woolson, R. F. (1981). Rank tests and a one-sample log rank test for comparing observed survival data to a standard population. *Biometrics*, 37, 687 – 696.

APPENDIX

PROGRAM

WHEN BOTH MARGINAL AND CONDITIONAL DISTRIBUTIONS ARE GENERALIZED LOGISTIC

```
*****  
Written by Oya CAN MUTAN, 2009, Ankara  
*****
```

```
use numerical_libraries
```

```
real b1,mu1,sigma1,b2,mu21,sigma2,rho,teta,sigma21,mu2,u1(100)  
real x(100),u2(100),y(100),sumx,sumy,xbar,ybar,sxy,xxx,syy  
real LSEsigma1,LSEmu1,LSEteta,LSErho,wls(100),sumwls,wlsbar  
real swwls,LSEsigma21,LSEmu21,LSEsigma2,LSEmu2,orderx,concoy  
real q(100),t1(100),alfa1(100),bbeta1(100),m1,delta1(100),delta11  
real dd1,kk1,bb1,cc1,MMLEsigma1,MMLEmu1,MMLEteta(3),w(2,100)  
real wo,cx,cy,t2(100),alfa2(100),bbeta2(100),m2,delta2(100)  
real sumcy,sumcx,cybar,cxbar,k2num,k2den,d2num,kk2,dd2,bb2,cc2  
real MMLEsigma21,MMLEmu21,MMLEmu2,ss,MMLEsigma2,MMLERho,delta22  
real avMMLEmu1(10000),avMMLEsigma1(10000),avMMLEmu21(10000)  
real avMMLEsigma21(10000),avMMLEteta(10000),avLSEmu1(10000)  
real avLSEsigma1(10000),avLSEmu21(10000),avLSEsigma21(10000)  
real avLSEteta(10000),t_teta(10000),t_teta_LSE(10000),tetah(10000)  
real mu21h(10000),sigma21h(10000),mulh(10000),c0  
real sigma1h(10000),mu2h(10000),sigma2h(10000),rhoh(10000)  
real sumtetah,summu21h,sumsigma21h,summulh,sumsigma1h,summu2h  
real sumsigma2h,sumrhoh,simtetahbar,simmu21hbar,simsigma21hbar  
real simmulhbar,simsigma1hbar,simmu2hbar,simsigma2hbar,simrhohbar  
real varnumtetah,varnummu21h,varnumsigma21h,varnummulh  
real varnumsigma1h,varnummu2h,varnumsigma2h,varnumrhoh,vartetah  
real varmu21h,varsigma21h,varmulh,varsigma1h,varmu2h,varsigma2h  
real varrhoh,t_teta_count,t_teta_count_LSE,tetaLSE(10000)  
real mu21LSE(10000),sigma21LSE(10000),mulLSE(10000)  
real sigma1LSE(10000),mu2LSE(10000),sigma2LSE(10000),rhoLSE(10000)  
real sumtetaLSE,summu21LSE,sumsigma21LSE,summulLSE,sumsigma1LSE  
real summu2LSE,sumsigma2LSE,sumrhoLSE,power  
real power_LSE,simtetaLSEbar,simmu21LSEbar  
real simsigma21LSEbar,simmulLSEbar,simsigma1LSEbar,simmu2LSEbar  
real simsigma2LSEbar,simrhoLSEbar,varnumtetahLSE,varnummu21LSE  
real varnumsigma21LSE,varnummulLSE,varnumsigma1LSE,varnummu2LSE  
real varnumsigma2LSE,varnumrhoLSE,vartetahLSE,varmu21LSE  
real varsigma21LSE,varmulLSE,varsigma1LSE,varmu2LSE,varsigma2LSE  
real varrhoLSE,nhead,nh(10000),sumnh,simnhbar,varnumnh,varnh  
real psid_1,psid_2,psid_b1,psid_b1plus1,psid_b2,psid_b2plus1  
real e(100),nheadLSE,nhLSE(10000),sumnhLSE,simnhbarLSE  
real simnhbarLSE2,varnumnhLSE,varnhLSE,remul,resigmal,remu2
```

```

real resigma2, rerho, ren, MVBmul, MVBsigma1, MVBmu21, MVBsigma21
real MVBteta, vartetaMMLHo, vartetaLSHo

real xc(100), yc(100), orderxc, concoyc, sumxc, sumyc, xbarc
real ybarc, sxyc, sxxc, LSEtetac, qc(100), tc(100), alfac(100), sumalfac
real bbetac(100), fc, fpc, Fcdfc, beta2c, alfa2c, sumbxc, mmc
real delta1c(100), knumc, kdenc, kkc, dnumc, ddc, bblc, bbc, cclc
real ccc, MMLEsigma1c, MMLEmulc, nheadc, MMLEtetac(3), wc(2,100), woc
real cxc, cyc, q2c(100), t2c(100), alfa22c(100), bbeta22c(100), m22c
real delta22c(100), sdelta22c, sumcyc, sumcxc, cybarc, cxbarc, k22numc
real k22denc, d22numc, kk22c, dd22c, bb22c, cc22c, MMLEsigma21c
real MMLEmu21c, MMLEmu2c, ssc, MMLEsigma2c, MMLErhoc, tetahc(10000)
real mu21hc(10000), sigma21hc(10000), mulhc(10000), sigma1hc(10000)
real mu2hc(10000), sigma2hc(10000), rhohc(10000), nhc(10000)
real sumtetahc, summu21hc, sumsigma21hc, summulhc, sumsigma1hc
real summu2hc, sumsigma2hc, sumrhohc, sumnhc, simtetahbarc
real simmu21hbarc, simsigma21hbarc, simmulhbarc, simsigma1hbarc
real simmu2hbarc, simsigma2hbarc, simrhohbarc, simnhbarc, simnhbar2c
real varnumtetahc, varnummu21hc, varnumsigma21hc, varnummulhc
real varnumsigma1hc, varnummu2hc, varnumsigma2hc, varnumrhohc
real varnumnhc, vartetahc, varmu21hc, varsigma21hc, varmulhc
real varsigma1hc, varmu2hc, varsigma2hc, varrhohc, varnhc, eec(100)
real ac(100), eac(100), ea2c(100), z2c(100), glzc(100), fzc, fzpc, Fzcdfc
real g2zc, g2pc, slc, s2c, s3c, sz2c, sglzc, szglc, sumac, sumeac, sumea2c
real sumaeac, sumaea2c, suma2ea2c, sumxeac, sumxea2c, sumx2ea2c
real sumxaea2c, MVBmulhc, MVBsigma1hc, MVBmu21hc, MVBsigma21hc
real MVBtetahc, MVBmulc(10000), MVBsigma1c(10000), MVBmu21c(10000)
real MVBsigma21c(10000), MVBtetac(10000), sumMVBmulc, sumMVBsigma1c
real sumMVBmu21c, sumMVBsigma21c, sumMVBtetac, simMVBmulbarc
real simMVBsigma1barc, simMVBmu21barc, simMVBsigma21barc
real simMVBtetabarc, sumtc, tbarc

integer n, nn, scount, s, order, resul, no, no2, r, n2, orderc, resultc

parameter (LDA=5, LDAINV=5, NA=5)
real A(LDA, LDA), AINV(LDAINV, LDAINV)

parameter (LDB=5, LDBINV=5, NB=5)
real B(LDB, LDB), BINV(LDBINV, LDBINV)

parameter (LDC=5, LDCINV=5, NC=5)
real C(LDC, LDC), CINV(LDCINV, LDCINV)

parameter (LDD=5, LDDINV=5, ND=5)
real D(LDD, LDD), DINV(LDDINV, LDDINV)

open(unit=1, file='C:\oya\PHD\program\XglYXgl\all\all.txt')

print*, 'enter n'
read*, n

r=int(0.1*n)
n2=n-r

nn=int(100000.0/(n*1.0))

c-----
c      entering the parameter values
c-----

b1=4.0
mul=0.0
sigma1=1.0
b2=4.0

```

```

mu21=0.0
sigma2=1.0
rho=0.5
teta=rho*(sigma2/sigma1)
sigma21=sigma2*sqrt(1.0-rho**2.0)
mu2=mu21+teta*(mu1+sigma1*(psi(b1)-psi(1.0)))

```

```

c-----
c      specifying the values of the Trigamma functions
c-----

```

```

psid_1=1.6449
psid_2=0.6449

if(b1.eq.0.1) then
psid_b1=101.4316
psid_b1plus1=1.4333
else if(b1.eq.0.2) then
psid_b1=26.2674
psid_b1plus1=1.2672
else if(b1.eq.0.5) then
psid_b1=4.9348
psid_b1plus1=0.9348
else if(b1.eq.1.0) then
psid_b1=1.6449
psid_b1plus1=0.6449
else if(b1.eq.2.0) then
psid_b1=0.6449
psid_b1plus1=0.3949
else if(b1.eq.3.0) then
psid_b1=0.3949
psid_b1plus1=0.2838
else if(b1.eq.4.0) then
psid_b1=0.2838
psid_b1plus1=0.2213
else if(b1.eq.6.0) then
psid_b1=0.1813
psid_b1plus1=0.1535
else if(b1.eq.8.0) then
psid_b1=0.1331
psid_b1plus1=0.1175
endif

```

```

if(b2.eq.0.1) then
psid_b2=101.4316
psid_b2plus1=1.4333
else if(b2.eq.0.2) then
psid_b2=26.2674
psid_b2plus1=1.2672
else if(b2.eq.0.5) then
psid_b2=4.9348
psid_b2plus1=0.9348
else if(b2.eq.1.0) then
psid_b2=1.6449
psid_b2plus1=0.6449
else if(b2.eq.2.0) then
psid_b2=0.6449
psid_b2plus1=0.3949
else if(b2.eq.3.0) then
psid_b2=0.3949
psid_b2plus1=0.2838
else if(b2.eq.4.0) then
psid_b2=0.2838
psid_b2plus1=0.2213
else if(b2.eq.6.0) then

```

```

        psid_b2=0.1813
        psid_b2plus1=0.1535
        else if(b2.eq.8.0) then
        psid_b2=0.1331
        psid_b2plus1=0.1175
        endif

c-----
c      obtaining MVB's - complete sample -
c-----

        C(1,1)=(1.0*n*b1)/((sigma1**2.0)*(b1+2.0))
        C(1,2)=C(1,1)*(psi(b1+1.0)-psi(2.0))
        C(1,3)=0.0
        C(1,4)=0.0
        C(1,5)=0.0

        C(2,1)=C(1,2)
        C(2,2)=(1.0*n)/(sigma1**2.0)+C(1,1)*(psid_b1plus1+psid_2+
&(psi(b1+1.0)-psi(2.0))**2.0)
        C(2,3)=0.0
        C(2,4)=0.0
        C(2,5)=0.0

        C(3,1)=C(1,3)
        C(3,2)=C(2,3)
        C(3,3)=(n*1.0*b2)/((b2+2.0)*(sigma21**2.0))
        C(3,4)=C(3,3)*(psi(b2+1.0)-psi(2.0))
        C(3,5)=C(3,3)*(mul+sigma1*(psi(b1)-psi(1.0)))

        C(4,1)=C(1,4)
        C(4,2)=C(2,4)
        C(4,3)=C(3,4)
        C(4,4)=(n*1.0/(sigma21**2.0))+C(3,3)*(psid_b2plus1+psid_2
&+(psi(b2+1.0)-psi(2.0))**2.0)
        C(4,5)=C(3,3)*(psi(b2+1.0)-psi(2.0))*(mul+sigma1*
&(psi(b1)-psi(1.0)))

        C(5,1)=C(1,5)
        C(5,2)=C(2,5)
        C(5,3)=C(3,5)
        C(5,4)=C(4,5)
        C(5,5)=C(3,3)*((mul**2.0)+2.0*mul*sigma1*(psi(b1)-psi(1.0))+(sigma
&1**2.0)*(psid_b1+psid_1+((psi(b1)-psi(1.0))**2.0)))

        CALL LINRG (NC,C,LDC,CINV,LDCINV)

        MVBmul=CINV(1,1)
        MVBsigma1=CINV(2,2)
        MVBmu21=CINV(3,3)
        MVBsigma21=CINV(4,4)
        MVBteta=CINV(5,5)

c-----
c      c0=total permissible cost
c-----

        c0=1.0*n*(mul+sigma1*(psi(b1)-psi(1.0)))

c-----
c      starting simulation
c-----

        do 1000 s=1,nn
        scount=s

```

```

c-----
c      generating x(i) from generalized logistic distribution
c-----

      call rnun(n,u1)

      do i=1,n
      x(i)=mul-sigma1*alog((u1(i)**(-1.0/b1))-1.0)
      enddo

c-----
c      generating outliers _x+2.0_ x+4.0 Dixon's model
c-----

c-----
c      no:number of outliers
c-----

c      no=int(0.5+0.10*n)

c      do i=1,no
c      x(i)=x(i)+2.0
c      x(i)=x(i)+4.0
c      enddo

c-----
c      generating outliers for Tiku's outlier model
c-----

c-----
c      ordering x(i)
c-----

c      order=1
c5     if (order.eq.1) then
c      order=0

c      do 8 i=1,n-1

c      if (x(i).gt.x(i+1)) then

c      orderx=x(i)
c      x(i)=x(i+1)
c      x(i+1)=orderx

c      order=1

c      endif

c8     continue

c      go to 5

c      endif

c-----
c      no=number of outliers
c-----

c      no=0.1*n
c      no2=n-no+1

c      do i=no2,n
c      x(i)=x(i)+4.0

```

```

c      enddo
c-----
c      generating y(i)
c-----

      call rnun(n,u2)

      do i=1,n
      e(i)=-alog((u2(i)**(-1.0/b2))-1.0)
      enddo

      do i=1,n
      y(i)=(mu21+teta*x(i))+sigma21*e(i)
      enddo

      do i=1,n
      xc(i)=x(i)
      yc(i)=y(i)
      enddo

c-----
c      obtaining estimators for complete sample
c-----

c-----
c      estimating least squares estimators
c-----

      sumx=0.0
      sumy=0.0

      do i=1,n
      sumx=sumx+x(i)
      sumy=sumy+y(i)
      enddo

      xbar=sumx/(n*1.0)
      ybar=sumy/(n*1.0)

      sxy=0.0
      sxx=0.0
      syy=0.0

      do i=1,n
      sxy=sxy+(x(i)-xbar)*(y(i)-ybar)
      sxx=sxx+(x(i)-xbar)**2.0
      syy=syy+(y(i)-ybar)**2.0
      enddo

      LSEsigma1=sqrt((sxx/(1.0*n-1.0))/(psid_b1+psid_1))
      LSEmu1=xbar-LSEsigma1*(psi(b1)-psi(1.0))
      LSEteta=sxy/sxx
      LSErho=sxy/(sqrt(sxx*syy))

      do i=1,n
      wls(i)=y(i)-LSEteta*x(i)
      enddo

      sumwls=0.0

      do i=1,n
      sumwls=sumwls+wls(i)
      enddo

      wlsbar=sumwls/(1.0*n)

```

```

swwls=0.0

do i=1,n
swwls=swwls+(wls(i)-wlsbar)**2.0
enddo

swls=sqrt(swwls/(1.0*n-2.0))

LSEsigma21=swls/(sqrt(psid_b2+psid_1))
LSEmu21=wlsbar-LSEsigma21*(psi(b2)-psi(1.0))
LSEsigma2=sqrt((syy/(1.0*n-1.0))/((LSErho**2.0)*(psid_b1+psid_1)+
&(1.0-LSErho**2.0)*(psid_b2+psid_1)))
LSEmu2=ybar-(psi(b2)-psi(1.0))*sqrt(1.0-LSErho**2.0)*LSEsigma2

c-----
c      ordering x(i),obtaining the corresponding concomitant y[i]
c-----

      order=1
51      if (order.eq.1) then
          order=0

          do 8 i=1,n-1

              if (x(i).gt.x(i+1)) then

                  orderx=x(i)
                  x(i)=x(i+1)
                  x(i+1)=orderx

                  concoy=y(i)
                  y(i)=y(i+1)
                  y(i+1)=concoy

                  order=1

              endif

81      continue

          go to 51

      endif

c-----
c      obtaining MMLE of mul and signal
c-----

      do i=1,n
          q(i)=(1.0*i)/(1.0*n+1.0)
          enddo

      do i=1,n
          t1(i)=-alog(((q(i))**(-1.0/b1))-1.0)
          enddo

      do i=1,n
          alfa1(i)=(1.0+exp(t1(i))+t1(i)*exp(t1(i)))/(1.0+exp(t1(i)))**2.0
          enddo

      do i=1,n
          bbeta1(i)=exp(t1(i))/(1.0+exp(t1(i)))**2.0
          enddo

```

```

m1=0.0

do i=1,n
m1=m1+bbeta1(i)
enddo

do i=1,n
delta1(i)=(1.0/(b1+1.0))-alfa1(i)
enddo

delta11=0.0

do i=1,n
delta11=delta11+delta1(i)
enddo

ddl=delta11/m1

kk1=0.0

do i=1,n
kk1=kk1+(bbeta1(i)*x(i))/m1
enddo

bb1=0.0

do i=1,n
bb1=bb1+(b1+1.0)*((delta1(i))*(x(i)-kk1))
enddo

cc1=0.0

do i=1,n
cc1=cc1+(b1+1.0)*((bbeta1(i))*((x(i)-kk1)**2.0))
enddo

MMLEsigma1=(bb1+sqrt((bb1**2.0)+4.0*n*cc1))/(2.0*n)
MMLEmu1=kk1+ddl*MMLEsigma1

c-----
c   estimating the sample size with respect to total cost
c-----

nhead=c0/(MMLEmu1+MMLEsigma1*(psi(b1)-psi(1.0)))
nheadLSE=c0/(LSEmu1+LSEsigma1*(psi(b1)-psi(1.0)))

c-----
c   obtaining MMLE of mu2.1,sigma2.1 & teta
c-----

c-----
c   ordering w(i)'s and finding the concomitants x[i],y[i]
c-----

MMLEteta(1)=LSEteta

do 11 j=1,2

do i=1,n
w(j,i)=y(i)-MMLEteta(j)*x(i)
enddo

resul=1
55 if (resul.eq.1) then
resul=0

```

```

do 22 i=1,n-1

  if (w(j,i).gt.w(j,i+1)) then

    wo=w(j,i)
    w(j,i)=w(j,i+1)
    w(j,i+1)=wo

    cx=x(i)
    x(i)=x(i+1)
    x(i+1)=cx

    cy=y(i)
    y(i)=y(i+1)
    y(i+1)=cy

    resul=1

  endif

22  continue

  go to 55

endif

c-----
c      obtaining t2(i),alfa2(i),beta2(i),m2,delta2(i)
c-----

do i=1,n
  t2(i)=-alog(((q(i))**(-1.0/b2))-1.0)
enddo

do i=1,n
  alfa2(i)=(1.0+exp(t2(i))+t2(i)*exp(t2(i)))/(1.0+exp(t2(i)))**2.0
enddo

do i=1,n
  bbeta2(i)=exp(t2(i))/(1.0+exp(t2(i)))**2.0
enddo

m2=0.0

do i=1,n
  m2=m2+bbeta2(i)
enddo

do i=1,n
  delta2(i)=alfa2(i)-1.0/(1.0*b2+1.0)
enddo

delta22=0.0

do i=1,n
  delta22=delta22+delta2(i)
enddo

c-----
c      calculating ybar[.],xbar[.],K,D,B,C
c-----

sumcy=0.0
sumcx=0.0

```

```

do i=1,n
sumcy=sumcy+bbeta2(i)*y(i)
sumcx=sumcx+bbeta2(i)*x(i)
enddo

cybar=sumcy/m2
cxbar=sumcx/m2

k2num=0.0
k2den=0.0
d2num=0.0

do i=1,n
k2num=k2num+bbeta2(i)*(x(i)-cxbar)*y(i)
k2den=k2den+bbeta2(i)*((x(i)-cxbar)**2.0)
d2num=d2num+delta2(i)*(x(i)-cxbar)
enddo

kk2=k2num/k2den
dd2=d2num/k2den

bb2=0.0
cc2=0.0

do i=1,n
bb2=bb2+(b2+1.0)*(delta2(i)*((y(i)-cybar)-kk2*(x(i)-cxbar)))
cc2=cc2+(b2+1.0)*((bbeta2(i)*((y(i)-cybar)**2.0))-kk2*(bbeta2(i)*
&(x(i)-cxbar)*y(i)))
enddo

MMLEsigma21=(-bb2+sqrt((bb2**2.0)+4.0*n*cc2))/(2.0*n)
MMLEteta(j+1)=kk2-dd2*MMLEsigma21
MMLEmu21=cybar-cxbar*MMLEteta(j+1)-(delta22/m2)*MMLEsigma21
11 continue

MMLEmu2=MMLEmu21+(MMLEteta(3))*(MMLEmu1+MMLEsigma1*(psi(b1)-psi(1
&.0)))
ss=(MMLEsigma21**2.0)+((MMLEteta(3))**2.0)*(MMLEsigma1**2.0)
MMLEsigma2=sqrt(ss*1.0)
MMLErho=(MMLEteta(3))*MMLEsigma1/MMLEsigma2

tetah(scount)=MMLEteta(3)
mu21h(scount)=MMLEmu21
sigma21h(scount)=MMLEsigma21

mulh(scount)=MMLEmu1
sigma1h(scount)=MMLEsigma1
mu2h(scount)=MMLEmu2
sigma2h(scount)=MMLEsigma2
rhoh(scount)=MMLErho

tetaLSE(scount)=LSEteta
mu21LSE(scount)=LSEmu21
sigma21LSE(scount)=LSEsigma21

mulLSE(scount)=LSEmu1
sigma1LSE(scount)=LSEsigma1
mu2LSE(scount)=LSEmu2
sigma2LSE(scount)=LSEsigma2
rhoLSE(scount)=LSErho

nh(scount)=nhead
nhLSE(scount)=nheadLSE

```

```

C-----
c      the estimated Fisher Information Matrix with MMLE's
c      (mul,sigma1,mu21,sigma21,teta) for b=4.0
C-----

      A(1,1)=(1.0*n*b1)/((sigma1h(scount)**2.0)*(b1+2.0))
      A(1,2)=A(1,1)*(psi(b1+1.0)-psi(2.0))
      A(1,3)=0.0
      A(1,4)=0.0
      A(1,5)=0.0

      A(2,1)=A(1,2)
      A(2,2)=(1.0*n)/(sigma1h(scount)**2.0)+A(1,1)*(psid_b1plus1+psid_2+
&(psi(b1+1.0)-psi(2.0))**2.0)
      A(2,3)=0.0
      A(2,4)=0.0
      A(2,5)=0.0

      A(3,1)=A(1,3)
      A(3,2)=A(2,3)
      A(3,3)=(n*1.0*b2)/((b2+2.0)*(sigma21h(scount)**2.0))
      A(3,4)=A(3,3)*(psi(b2+1.0)-psi(2.0))
      A(3,5)=A(3,3)*(mulh(scount)+sigma1h(scount)*(psi(b1)-psi(1.0)))

      A(4,1)=A(1,4)
      A(4,2)=A(2,4)
      A(4,3)=A(3,4)
      A(4,4)=(n*1.0/(sigma21h(scount)**2.0))+A(3,3)*(psid_b2plus1+psid_2
&+(psi(b2+1.0)-psi(2.0))**2.0)
      A(4,5)=A(3,3)*(psi(b2+1.0)-psi(2.0))*(mulh(scount)+sigma1h(scount)
&*(psi(b1)-psi(1.0)))

      A(5,1)=A(1,5)
      A(5,2)=A(2,5)
      A(5,3)=A(3,5)
      A(5,4)=A(4,5)
      A(5,5)=A(3,3)*((mulh(scount)**2.0)+2.0*mulh(scount)*sigma1h(scount
&)*(psi(b1)-psi(1.0))+(sigma1h(scount)**2.0)*(psid_b1+psid_1+((psi(
&b1)-psi(1.0))**2.0)))

      CALL LINRG (NA,A,LDA,AINV,LDAINV)

      avMMLEmul(scount)=AINV(1,1)
      avMMLEsigma1(scount)=AINV(2,2)
      avMMLEmu21(scount)=AINV(3,3)
      avMMLEsigma21(scount)=AINV(4,4)
      avMMLEteta(scount)=AINV(5,5)

C-----
c      the estimated Fisher Information Matrix with
c      LSE's(mul,sigma1,mu21,sigma21,teta) for b=4.0
C-----

      B(1,1)=(1.0*n*b1)/((sigma1LSE(scount)**2.0)*(b1+2.0))
      B(1,2)=B(1,1)*(psi(b1+1.0)-psi(2.0))
      B(1,3)=0.0
      B(1,4)=0.0
      B(1,5)=0.0

      B(2,1)=B(1,2)
      B(2,2)=(1.0*n)/(sigma1LSE(scount)**2.0)+B(1,1)*(psid_b1plus1+psid_
&2+(psi(b1+1.0)-psi(2.0))**2.0)
      B(2,3)=0.0
      B(2,4)=0.0

```

```

B(2,5)=0.0

B(3,1)=B(1,3)
B(3,2)=B(2,3)
B(3,3)=(n*1.0*b2)/((b2+2.0)*(sigma21LSE(scount)**2.0))
B(3,4)=B(3,3)*(psi(b2+1.0)-psi(2.0))
B(3,5)=B(3,3)*(mulLSE(scount)+sigma1LSE(scount)*(psi(b1)-psi(1.0))
&)

B(4,1)=B(1,4)
B(4,2)=B(2,4)
B(4,3)=B(3,4)
B(4,4)=(n*1.0/(sigma21LSE(scount)**2.0))+B(3,3)*(psid_b2plus1+psid
&_2+(psi(b2+1.0)-psi(2.0))**2.0)
B(4,5)=B(3,3)*(psi(b2+1.0)-psi(2.0))*(mulLSE(scount)+sigma1LSE(sco
&unt)*(psi(b1)-psi(1.0)))

B(5,1)=B(1,5)
B(5,2)=B(2,5)
B(5,3)=B(3,5)
B(5,4)=B(4,5)
B(5,5)=B(3,3)*((mulLSE(scount)**2.0)+2.0*mulLSE(scount)*sigma1LSE
&(scount)*(psi(b1)-psi(1.0))+(sigma1LSE(scount)**2.0)*(psid_b1+psid
&_1+((psi(b1)-psi(1.0))**2.0)))

CALL LINRG (NB,B,LDB,BINV,LDBINV)

avLSEmul(scount)=BINV(1,1)
avLSEsigma1(scount)=BINV(2,2)
avLSEmu21(scount)=BINV(3,3)
avLSEsigma21(scount)=BINV(4,4)
avLSEteta(scount)=BINV(5,5)

c-----
c      obtaining estimators for censored sample
c-----

c-----
c      ordering xc(i),obtaining the corresponding concomitant yc[i]
c-----

      orderc=1
550      if (orderc.eq.1) then
          orderc=0

          do 88 i=1,n-1

              if (xc(i).gt.xc(i+1)) then

                  orderxc=xc(i)
                  xc(i)=xc(i+1)
                  xc(i+1)=orderxc

                  concoyc=yc(i)
                  yc(i)=yc(i+1)
                  yc(i+1)=concoyc

                  orderc=1

              endif

88      continue

      go to 550

```

```

endif

sumxc=0.0
sumyc=0.0
sxyz=0.0
sxxc=0.0
do i=1,n2
sumxc=sumxc+xc(i)
sumyc=sumyc+yc(i)
enddo

xbarc=sumxc/(n2*1.0)
ybarc=sumyc/(n2*1.0)
do i=1,n2
sxyz=sxyz+(xc(i)-xbarc)*(yc(i)-ybarc)
sxxc=sxxc+(xc(i)-xbarc)**2.0
enddo
LSEtetac=sxyz/sxxc

c-----
c      obtaining MMLE of mul and sigma1
c-----

do i=1,n2
qc(i)=(1.0*i)/(1.0*n+1.0)
enddo

do i=1,n2
tc(i)=-alog(((qc(i))**(-1.0/b1))-1.0)
enddo

sumtc=0.0
do i=1,n2
sumtc=sumtc+tc(i)
enddo

tbarc=sumtc/(1.0*n2)

do i=1,n2
alfac(i)=(1.0+exp(tc(i))+tc(i)*exp(tc(i)))/(1.0+exp(tc(i)))**2.0
enddo

sumalfac=0.0
do i=1,n2
sumalfac=sumalfac+alfac(i)
enddo

do i=1,n2
bbetac(i)=exp(tc(i))/(1.0+exp(tc(i)))**2.0
enddo

c-----
c      z~GL(b,sigma=1)
c      obtaining f(z), f'(z),F(z) at the point t(n-r)
c-----

fc=(b1*exp(-tc(n2)))/((1.0+exp(-tc(n2)))**(b1+1.0))
fpc=(b1*exp(-tc(n2)))*(b1*exp(-tc(n2))-1.0)/(1.0+exp(-tc(n2)))
Fcdfc=(1.0+exp(-tc(n2)))**(-1.0*b1)

c-----
c      obtaining beta2 & alfa2
c-----

beta2c=-1.0*(fpc*(1.0-Fcdfc)+fc**2.0)/((1.0-Fcdfc)**2.0)

```

```

alfa2c=(fc/(1.0-Fcdfc))+beta2c*tc(n2)

sumbxc=0.0

do i=1,n2
sumbxc=sumbxc+bbetac(i)*xc(i)
enddo

mmc=0.0

do i=1,n2
mmc=mmc+bbetac(i)
enddo

do i=1,n2
delta1c(i)=(1.0/(b1+1.0))-alfac(i)
enddo

knumc=(b1+1.0)*sumbxc-1.0*r*beta2c*xc(n2)
kdenc=(b1+1.0)*mmc-1.0*r*beta2c
kkc=knmc/kdenc

dnumc=1.0*n2-(b1+1.0)*sumalfac+1.0*r*alfa2c
ddc=dnumc/kdenc

bb1c=0.0

do i=1,n2
bb1c=bb1c+(b1+1.0)*((delta1c(i))*(xc(i)-kkc))
enddo

bbc=bb1c+1.0*r*alfa2c*(xc(n2)-kkc)

cc1c=0.0

do i=1,n2
cc1c=cc1c+(b1+1.0)*((bbetac(i))*((xc(i)-kkc)**2.0))
enddo

ccc=cc1c-1.0*r*beta2c*((xc(n2)-kkc)**2.0)

MMLEsigmalc=(bbc+sqrt((bbc**2.0)+4.0*n2*ccc))/(2.0*n2)
MMLEmulc=kkc+ddc*MMLEsigmalc

c-----
c   estimating the sample size with respect to total cost
c-----

nheadc=1.0*r+c0/(MMLEmulc+MMLEsigmalc*tbarc)

c-----
c   obtaining MMLE of mu2.1,sigma2.1 & teta
c-----

c-----
c   ordering w(i)'s and finding the concomitants x[i],y[i]
c-----

MMLEtetac(1)=LSEtetac

do 111 j=1,2

do i=1,n2
wc(j,i)=yc(i)-MMLEtetac(j)*xc(i)

```

```

        enddo

        resulc=1
555      if (resulc.eq.1) then
        resulc=0

        do 222 i=1,n2-1

          if (wc(j,i).gt.wc(j,i+1)) then

            woc=wc(j,i)
            wc(j,i)=wc(j,i+1)
            wc(j,i+1)=woc

            cxc=xc(i)
            xc(i)=xc(i+1)
            xc(i+1)=cxc

            cyc=yc(i)
            yc(i)=yc(i+1)
            yc(i+1)=cyc

            resulc=1

          endif

222      continue

        go to 555

      endif

c-----
c      obtaining q2(i),t2(i),alfa22(i),beta22(i),m22,delta22(i),sdelta22
c-----

      do i=1,n2
        q2c(i)=1.0*i/(1.0*n+1)
      enddo

      do i=1,n2
        t2c(i)=-alog(((q2c(i))**(-1.0/b2))-1.0)
      enddo

      do i=1,n2
        alfa22c(i)=(1.0+exp(t2c(i))+t2c(i)*exp(t2c(i)))/(1.0+exp(t2c(i)))
& **2.0
      enddo

      do i=1,n2
        bbeta22c(i)=exp(t2c(i))/(1.0+exp(t2c(i)))**2.0
      enddo

      m22c=0.0

      do i=1,n2
        m22c=m22c+bbeta22c(i)
      enddo

      do i=1,n2
        delta22c(i)=alfa22c(i)-1.0/(1.0*b2+1.0)
      enddo

      sdelta22c=0.0

```

```

do i=1,n2
sdelta22c=sdelta22c+delta22c(i)
enddo

c-----
c      calculating ybar[.],xbar[.],K,D,B,C
c-----

sumcyc=0.0
sumcxc=0.0

do i=1,n2
sumcyc=sumcyc+bbeta22c(i)*yc(i)
sumcxc=sumcxc+bbeta22c(i)*xc(i)
enddo

cybarc=sumcyc/m22c
cxbarc=sumcxc/m22c

k22numc=0.0
k22denc=0.0
d22numc=0.0

do i=1,n2
k22numc=k22numc+bbeta22c(i)*(xc(i)-cxbarc)*yc(i)
k22denc=k22denc+bbeta22c(i)*((xc(i)-cxbarc)**2.0)
d22numc=d22numc+delta22c(i)*(xc(i)-cxbarc)
enddo

kk22c=k22numc/k22denc
dd22c=d22numc/k22denc

bb22c=0.0
cc22c=0.0

do i=1,n2
bb22c=bb22c+(b2+1.0)*(delta22c(i)*((yc(i)-cybarc)-kk22c*(xc(i)-
&cxbarc)))
cc22c=cc22c+(b2+1.0)*((bbeta22c(i)*((yc(i)-cybarc)**2.0))-kk22c*
&(bbeta22c(i)*(xc(i)-cxbarc)*yc(i)))
enddo

MMLEsigma21c=(-bb22c+sqrt((bb22c**2.0)+4.0*n2*cc22c))/(2.0*n2)
MMLEtetac(j+1)=kk22c-dd22c*MMLEsigma21c
MMLEmu21c=cybarc-cxbarc*MMLEtetac(j+1)-(sdelta22c/m22c)*
&MMLEsigma21c

111 continue

MMLEmu2c=MMLEmu21c+(MMLEtetac(3))*(MMLEmulc+MMLEsigmalc*(psi(b1)
&-psi(1.0)))
ssc=(MMLEsigma21c**2.0)+((MMLEtetac(3))**2.0)*(MMLEsigmalc**2.0)
MMLEsigma2c=sqrt(ssc*1.0)
MMLErhoc=(MMLEtetac(3))*MMLEsigmalc/MMLEsigma2c

c-----
c      obtaining sample information matrix
c-----

do i=1,n2
eec(i)=yc(i)-MMLEmu21c-MMLEtetac(3)*xc(i)
ac(i)=eec(i)/MMLEsigma21c
eac(i)=(exp(-ac(i)))/(1.0+exp(-ac(i)))
ea2c(i)=(exp(-ac(i)))/((1.0+exp(-ac(i)))**2.0)
z2c(i)=(xc(i)-MMLEmulc)/MMLEsigmalc

```

```

glzc(i)=exp(-z2c(i))/(1.0+exp(-z2c(i)))
enddo

fzc=(b1*exp(-z2c(n2)))/((1.0+exp(-z2c(n2)))**2.0)
fzpc=(b1*exp(-z2c(n2)))*(b1*exp(-z2c(n2))-1.0)/(1.0+exp(-z2c(n2)))
Fzcdfc=(1.0+exp(-z2c(n2)))**(-1.0*b1)

g2zc=fzc/(1.0-Fzcdfc)
g2pc=(fzpc*(1.0-Fzcdfc)+fzc**2.0)/((1.0-Fzcdfc)**2.0)

s1c=0.0
s2c=0.0
s3c=0.0
sz2c=0.0
sglzc=0.0
szglc=0.0

sumac=0.0
sumeac=0.0
sumea2c=0.0
sumaeac=0.0
sumaea2c=0.0
suma2ea2c=0.0
sumxeac=0.0
sumxea2c=0.0
sumx2ea2c=0.0
sumxaea2c=0.0

do i=1,n2
s1c=s1c+exp(-z2c(i))/((1.0+exp(-z2c(i)))**2.0)
s2c=s2c+z2c(i)*exp(-z2c(i))/((1.0+exp(-z2c(i)))**2.0)
s3c=s3c+((z2c(i))**2.0)*exp(-z2c(i))/((1.0+exp(-z2c(i)))**2.0)
sz2c=sz2c+z2c(i)
sglzc=sglzc+glzc(i)
szglc=szglc+z2c(i)*glzc(i)

sumac=sumac+ac(i)
sumeac=sumeac+eac(i)
sumaeac=sumaeac+ac(i)*eac(i)
sumaea2c=sumaea2c+ac(i)*ea2c(i)
suma2ea2c=suma2ea2c+((ac(i))**2.0)*ea2c(i)
sumeac2c=sumea2c+ea2c(i)
sumxeac=sumxeac+xc(i)*eac(i)
sumxea2c=sumxea2c+xc(i)*ea2c(i)
sumx2ea2c=sumx2ea2c+((xc(i))**2.0)*ea2c(i)
sumxaea2c=sumxaea2c+xc(i)*ac(i)*ea2c(i)

enddo

D(1,1)=(b1+1.0)/(MMLEsigmalc**2.0)*s1c+((1.0*r)/(MMLEsigmalc**
&2.0))*g2pc
D(1,2)=((1.0*n2)/(MMLEsigmalc**2.0))-((b1+1.0)/(MMLEsigmalc**2.0))
&*sglzc+((b1+1.0)/(MMLEsigmalc**2.0))*s2c+((1.0*r)/(MMLEsigmalc**
&2.0))*g2zc+((1.0*r)/(MMLEsigmalc**2.0))*(z2c(n2))*g2pc
D(1,3)=0.0
D(1,4)=0.0
D(1,5)=0.0

D(2,1)=D(1,2)
D(2,2)=-((1.0*n2)/(MMLEsigmalc**2.0))+(2.0/(MMLEsigmalc**2.0))*
&sz2c-2.0*((b1+1.0)/(MMLEsigmalc**2.0))*szglc+((b1+1.0)/(MMLEsigmal
&c**2.0))*s3c+((2.0*r)/(MMLEsigmalc**2.0))*z2c(n2)*g2zc+((1.0*r)/
&(MMLEsigmalc**2.0))*((z2c(n2))**2.0)*g2pc
D(2,3)=0.0

```

```

D(2,4)=0.0
D(2,5)=0.0

D(3,1)=D(1,3)
D(3,2)=D(2,3)
D(3,3)=(b2+1.0)/(MMLEsigma21c**2.0)*sume2c
D(3,4)=(1.0*n2)/(MMLEsigma21c**2.0)-(b2+1.0)/(MMLEsigma21c**2.0)
&*sumeac+(b2+1.0)/(MMLEsigma21c**2.0)*sumaea2c
D(3,5)=(b2+1.0)/(MMLEsigma21c**2.0)*sumxea2c

D(4,1)=D(1,4)
D(4,2)=D(2,4)
D(4,3)=D(3,4)
D(4,4)=-((1.0*n2)/(MMLEsigma21c**2.0))+(2.0/(MMLEsigma21c**2.0))*
&sumac-(2.0*(b2+1.0)/(MMLEsigma21c**2.0))*sumaeac+(b2+1.0)/
&(MMLEsigma21c**2.0)*suma2ea2c
D(4,5)=(1.0/(MMLEsigma21c**2.0))*sumxc-(b2+1.0)/(MMLEsigma21c**
&2.0)*sumxeac+(b2+1.0)/(MMLEsigma21c**2.0)*sumxaea2c

D(5,1)=D(1,5)
D(5,2)=D(2,5)
D(5,3)=D(3,5)
D(5,4)=D(4,5)
D(5,5)=(b2+1.0)/(MMLEsigma21c**2.0)*sumx2ea2c

CALL LINRG (ND,D,LDD,DINV,LDDINV)

MVBmulhc=DINV(1,1)
MVBsigma1hc=DINV(2,2)
MVBmu21hc=DINV(3,3)
MVBsigma21hc=DINV(4,4)
MVBtetahc=DINV(5,5)

tetahc(scount)=MMLEtetac(3)
mu21hc(scount)=MMLEmu21c
sigma21hc(scount)=MMLEsigma21c

mulhc(scount)=MMLEmulc
sigma1hc(scount)=MMLEsigma1c
mu2hc(scount)=MMLEmu2c
sigma2hc(scount)=MMLEsigma2c
rhohc(scount)=MMLErhoc

nhc(scount)=nheadc

MVBmulc(scount)=MVBmulhc
MVBsigma1c(scount)=MVBsigma1hc
MVBmu21c(scount)=MVBmu21hc
MVBsigma21c(scount)=MVBsigma21hc
MVBtetac(scount)=MVBtetahc

1000   enddo

c-----
c      calculating simulated means
c-----

sumtetah=0.0
summu21h=0.0
sumsigma21h=0.0

summulh=0.0
sumsigma1h=0.0
summu2h=0.0
sumsigma2h=0.0

```

```

sumrhoh=0.0

sumtetaLSE=0.0
summu21LSE=0.0
sumsigma21LSE=0.0

summu1LSE=0.0
sumsigma1LSE=0.0
summu2LSE=0.0
sumsigma2LSE=0.0
sumrhoLSE=0.0

sumnh=0.0
sumnhLSE=0.0

sumtetahc=0.0
summu21hc=0.0
sumsigma21hc=0.0

summu1hc=0.0
sumsigma1hc=0.0
summu2hc=0.0
sumsigma2hc=0.0
sumrhohc=0.0

sumnhc=0.0

sumMVBmu1c=0.0
sumMVBsigma1c=0.0
sumMVBmu21=0.0
sumMVBsigma21c=0.0
sumMVBtetac=0.0

do i=1,nn

sumtetah=sumtetah+tetah(i)
summu21h=summu21h+mu21h(i)
sumsigma21h=sumsigma21h+sigma21h(i)

summu1h=summu1h+mu1h(i)
sumsigma1h=sumsigma1h+sigma1h(i)
summu2h=summu2h+mu2h(i)
sumsigma2h=sumsigma2h+sigma2h(i)
sumrhoh=sumrhoh+rhoh(i)

sumtetaLSE=sumtetaLSE+tetaLSE(i)
summu21LSE=summu21LSE+mu21LSE(i)
sumsigma21LSE=sumsigma21LSE+sigma21LSE(i)

summu1LSE=summu1LSE+mu1LSE(i)
sumsigma1LSE=sumsigma1LSE+sigma1LSE(i)
summu2LSE=summu2LSE+mu2LSE(i)
sumsigma2LSE=sumsigma2LSE+sigma2LSE(i)
sumrhoLSE=sumrhoLSE+rhoLSE(i)

sumnh=sumnh+nh(i)
sumnhLSE=sumnhLSE+nhLSE(i)

sumtetahc=sumtetahc+tetahc(i)
summu21hc=summu21hc+mu21hc(i)
sumsigma21hc=sumsigma21hc+sigma21hc(i)

summu1hc=summu1hc+mu1hc(i)
sumsigma1hc=sumsigma1hc+sigma1hc(i)
summu2hc=summu2hc+mu2hc(i)

```

```

sumsigma2hc=sumsigma2hc+sigma2hc(i)
sumrhohc=sumrhohc+rhohc(i)

sumnhc=sumnhc+nhc(i)

sumMVBmulc=sumMVBmulc+MVBmulc(i)
sumMVBsigma1c=sumMVBsigma1c+MVBsigma1c(i)
sumMVBmu21c=sumMVBmu21c+MVBmu21c(i)
sumMVBsigma21c=sumMVBsigma21c+MVBsigma21c(i)
sumMVBtetac=sumMVBtetac+MVBtetac(i)

enddo

simtetahbar=sumtetah/(1.0*nn)
simmu21hbar=summu21h/(1.0*nn)
simsigma21hbar=sumsigma21h/(1.0*nn)

simmulhbar=summulh/(1.0*nn)
simsigma1hbar=sumsigma1h/(1.0*nn)
simmu2hbar=summu2h/(1.0*nn)
simsigma2hbar=sumsigma2h/(1.0*nn)
simrhohbar=sumrhoh/(1.0*nn)

simtetaLSEbar=sumtetaLSE/(1.0*nn)
simmu21LSEbar=summu21LSE/(1.0*nn)
simsigma21LSEbar=sumsigma21LSE/(1.0*nn)

simmulLSEbar=summulLSE/(1.0*nn)
simsigma1LSEbar=sumsigma1LSE/(1.0*nn)
simmu2LSEbar=summu2LSE/(1.0*nn)
simsigma2LSEbar=sumsigma2LSE/(1.0*nn)
simrhoLSEbar=sumrhoLSE/(1.0*nn)

simnhbar=sumnh/(1.0*nn)
simnhbar2=int(simnhbar)

simnhbarLSE=sumnhLSE/(1.0*nn)
simnhbarLSE2=int(simnhbarLSE)

simtetahbarc=sumtetahc/(1.0*nn)
simmu21hbarc=summu21hc/(1.0*nn)
simsigma21hbarc=sumsigma21hc/(1.0*nn)

simmulhbarc=summulhc/(1.0*nn)
simsigma1hbarc=sumsigma1hc/(1.0*nn)
simmu2hbarc=summu2hc/(1.0*nn)
simsigma2hbarc=sumsigma2hc/(1.0*nn)
simrhohbarc=sumrhohc/(1.0*nn)

simnhbarc=sumnhc/(1.0*nn)
simnhbar2c=int(simnhbarc)

simMVBmulbarc=sumMVBmulc/(1.0*nn)
simMVBsigma1barc=sumMVBsigma1c/(1.0*nn)
simMVBmu21barc=sumMVBmu21c/(1.0*nn)
simMVBsigma21barc=sumMVBsigma21c/(1.0*nn)
simMVBtetabarc=sumMVBtetac/(1.0*nn)

```

```

c-----
c      calculating simulated variances
c-----

```

```

varnumtetah=0.0
varnummu21h=0.0
varnumsigma21h=0.0

```

```

varnummulh=0.0
varnumsigma1h=0.0
varnummu2h=0.0
varnumsigma2h=0.0
varnumrhoh=0.0

varnumtetaLSE=0.0
varnummu21LSE=0.0
varnumsigma21LSE=0.0

varnummulLSE=0.0
varnumsigma1LSE=0.0
varnummu2LSE=0.0
varnumsigma2LSE=0.0
varnumrhoLSE=0.0

varnumnh=0.0
varnumnhLSE=0.0

varnumtetahc=0.0
varnummu21hc=0.0
varnumsigma21hc=0.0

varnummulhc=0.0
varnumsigma1hc=0.0
varnummu2hc=0.0
varnumsigma2hc=0.0
varnumrhohc=0.0

varnumnhc=0.0

do i=1,nn

varnumtetah=varnumtetah+(tetah(i)-simeetahbar)**2.0
varnummu21h=varnummu21h+(mu21h(i)-simmu21hbar)**2.0
varnumsigma21h=varnumsigma21h+(sigma21h(i)-simsigma21hbar)**2.0

varnummulh=varnummulh+(mulh(i)-simmulhbar)**2.0
varnumsigma1h=varnumsigma1h+(sigma1h(i)-simsigma1hbar)**2.0
varnummu2h=varnummu2h+(mu2h(i)-simmu2hbar)**2.0
varnumsigma2h=varnumsigma2h+(sigma2h(i)-simsigma2hbar)**2.0
varnumrhoh=varnumrhoh+(rhoh(i)-simrhohbar)**2.0

varnumtetaLSE=varnumtetaLSE+(tetaLSE(i)-simeetetaLSEbar)**2.0
varnummu21LSE=varnummu21LSE+(mu21LSE(i)-simmu21LSEbar)**2.0
varnumsigma21LSE=varnumsigma21LSE+(sigma21LSE(i)-simsigma21LSEbar)
&**2.0

varnummulLSE=varnummulLSE+(mulLSE(i)-simmulLSEbar)**2.0
varnumsigma1LSE=varnumsigma1LSE+(sigma1LSE(i)-simsigma1LSEbar)**
&2.0
varnummu2LSE=varnummu2LSE+(mu2LSE(i)-simmu2LSEbar)**2.0
varnumsigma2LSE=varnumsigma2LSE+(sigma2LSE(i)-simsigma2LSEbar)**
&2.0
varnumrhoLSE=varnumrhoLSE+(rhoLSE(i)-simrhoLSEbar)**2.0

varnumnh=varnumnh+(nh(i)-simnhbar)**2.0
varnumnhLSE=varnumnhLSE+(nhLSE(i)-simnhbarLSE)**2.0

varnumtetahc=varnumtetahc+(tetahc(i)-simeetahbarc)**2.0
varnummu21hc=varnummu21hc+(mu21hc(i)-simmu21hbarc)**2.0
varnumsigma21hc=varnumsigma21hc+(sigma21hc(i)-simsigma21hbarc)
&**2.0

```

```

varnummulhc=varnummulhc+(mulhc(i)-simmulhbarc)**2.0
varnumsigma1hc=varnumsigma1hc+(sigma1hc(i)-simsigma1hbarc)**2.0
varnummu2hc=varnummu2hc+(mu2hc(i)-simmu2hbarc)**2.0
varnumsigma2hc=varnumsigma2hc+(sigma2hc(i)-simsigma2hbarc)**2.0
varnumrhohc=varnumrhohc+(rhohc(i)-simrhohbarc)**2.0

```

```

varnumnhc=varnumnhc+(nhc(i)-simnhbarc)**2.0

```

```

enddo

```

```

vartetah=varnumtetah/(1.0*nn-1.0)
varmu21h=varnummu21h/(1.0*nn-1.0)
varsigma21h=varnumsigma21h/(1.0*nn-1.0)

```

```

varmulh=varnummulh/(1.0*nn-1.0)
varsigma1h=varnumsigma1h/(1.0*nn-1.0)
varmu2h=varnummu2h/(1.0*nn-1.0)
varsigma2h=varnumsigma2h/(1.0*nn-1.0)
varrho=varnumrho/(1.0*nn-1.0)

```

```

vartetaLSE=varnumtetaLSE/(1.0*nn-1.0)
varmu21LSE=varnummu21LSE/(1.0*nn-1.0)
varsigma21LSE=varnumsigma21LSE/(1.0*nn-1.0)

```

```

varmulLSE=varnummulLSE/(1.0*nn-1.0)
varsigma1LSE=varnumsigma1LSE/(1.0*nn-1.0)
varmu2LSE=varnummu2LSE/(1.0*nn-1.0)
varsigma2LSE=varnumsigma2LSE/(1.0*nn-1.0)
varrhoLSE=varnumrhoLSE/(1.0*nn-1.0)

```

```

varnh=varnumnh/(1.0*nn-1.0)
varnhLSE=varnumnhLSE/(1.0*nn-1.0)

```

```

vartetahc=varnumtetahc/(1.0*nn-1.0)
varmu21hc=varnummu21hc/(1.0*nn-1.0)
varsigma21hc=varnumsigma21hc/(1.0*nn-1.0)

```

```

varmulhc=varnummulhc/(1.0*nn-1.0)
varsigma1hc=varnumsigma1hc/(1.0*nn-1.0)
varmu2hc=varnummu2hc/(1.0*nn-1.0)
varsigma2hc=varnumsigma2hc/(1.0*nn-1.0)
varrhohc=varnumrhohc/(1.0*nn-1.0)

```

```

varnhc=varnumnhc/(1.0*nn-1.0)

```

```

c-----
c      calculating relative efficiency
c-----

```

```

remu1=(varmulh/varmulLSE)*100
resigma1=(varsigma1h/varsigma1LSE)*100
remu2=(varmu2h/varmu2LSE)*100
resigma2=(varsigma2h/varsigma2LSE)*100
rerho=(varrho/varrhoLSE)*100
ren=(varnh/varnhLSE)*100

```

```

c-----
c      testing Ho:teta=0.0, Ha:teta>0.0 by using MMLE's & LSE's simulated
c      variances
c-----

```

```

c      vartetaMMLHo=0.0184
c      vartetaLSHo=0.0233

```

```

c      t_teta_count=0.0

```

```

c      do i=1,nn
c      if(tetah(i)>=(1.645*sqrt(vartetah))) then
c      if(tetah(i)>=(1.645*sqrt(vartetamMLHo))) then
c      t_teta_count=t_teta_count+1.0
c      endif
c      enddo

c      power=t_teta_count/(nn*1.0)

c      t_teta_count_LSE=0.0

c      do i=1,nn
c      if(tetaLSE(i)>=(1.645*sqrt(vartetatLSE))) then
c      if(tetaLSE(i)>=(1.645*sqrt(vartetatLSho))) then
c      t_teta_count_LSE=t_teta_count_LSE+1.0
c      endif
c      enddo

c      power_LSE=t_teta_count_LSE/(nn*1.0)

c-----
c      printing means,variances & MVB
c-----

102      format(15x,a3,3x,a6,6x,a4,6x,a7,5x,a4,8x,a3,6x,a6,7x,a3,6x,a1)
         write(1,102)'MU1','SIGMA1','MU21','SIGMA21','TETA','MU2','SIGMA2',
         &'RHO','n'

103      format(a4,9x,f7.4,3x,f7.4,4x,f7.4,4x,f7.4,4x,f7.4,4x,f7.4,4x,f7.4,
         &4x,f7.4,1x,f10.4)
         write(1,103)'mean',simmulhbar,simsigma1hbar,simmu2lhbar,simsigma21
         &hbar,simtetahbar,simmu2hbar,simsigma2hbar,simrhohbar,simnhbar2

105      format(a7,6x,f7.4,3x,f7.4,4x,f7.4,4x,f7.4,4x,f7.4,4x,f7.4,4x,f7.4,
         &4x,f7.4,1x,f10.4)
         write(1,105)'meanLSE',simmulLSEbar,simsigma1LSEbar,simmu21LSEbar,
         &simsigma21LSEbar,simtetaLSEbar,simmu2LSEbar,simsigma2LSEbar,simrho
         &LSEbar,simnhbarLSE2

104      format(a8,5x,f7.4,3x,f7.4,4x,f7.4,4x,f7.4,4x,f7.4,4x,f7.4,4x,f7.4,
         &4x,f7.4,4x,f7.4)
         write(1,104)'variance',varmulh,varsigma1h,varmu21h,varsigma21h,var
         &tetah,varmu2h,varsigma2h,varrho,h,varnh

106      format(a11,2x,f7.4,3x,f7.4,4x,f7.4,4x,f7.4,4x,f7.4,4x,f7.4,4x,
         &f7.4,4x,f7.4,4x,f7.4)
         write(1,106)'varianceLSE',varmulLSE,varsigma1LSE,varmu21LSE,varsig
         &ma21LSE,vartetatLSE,varmu2LSE,varsigma2LSE,varrhoLSE,varnhLSE

113      format(a5,9x,f7.4,3x,f7.4,4x,f7.4,4x,f7.4,4x,f7.4,4x,f7.4,4x,f7.4,
         &4x,f7.4,1x,f10.4)
         write(1,113)'meanc',simmulhbarc,simsigma1hbarc,simmu2lhbarc,simsig
         &ma21hbarc,simtetahbarc,simmu2hbarc,simsigma2hbarc,simrhohbarc,
         &simnhbar2c

114      format(a9,5x,f7.4,3x,f7.4,4x,f7.4,4x,f7.4,4x,f7.4,4x,f7.4,4x,f7.4,
         &4x,f7.4,4x,f7.4)
         write(1,114)'variancec',varmulhc,varsigma1hc,varmu21hc,varsigma21
         &hc,vartetahc,varmu2hc,varsigma2hc,varrhohc,varnhc

c      write(1,*)'MVBmulc',simMVBmulbarc
c      write(1,*)'MVBsigma1c',simMVBsigma1barc
c      write(1,*)'MVBmu21c',simMVBmu21barc
c      write(1,*)'MVBsigma21c',simMVBsigma21barc

```

```

c      write(1,*) 'MVBtetac',simMVBtetabarc

209    format(a6,4x,f7.2,4x,f7.2,4x,f7.2,4x,f7.2,4x,f7.2,4x,f7.2)
      write(1,209) 'releff',remu1,resigma1,remu2,resigma2,rerho,ren

c      write(1,*) 'powerMML=',power
c      write(1,*) 'powerLSE=',power_LSE

c      write(1,*) 'MVBmu1=',MVBmu1
c      write(1,*) 'MVBsigma1=',MVBsigma1
c      write(1,*) 'MVBmu21=',MVBmu21
c      write(1,*) 'MVBsigma21=',MVBsigma21
c      write(1,*) 'MVBteta=',MVBteta

      stop
      end

```

VITA

Oya CAN MUTAN was born in Ankara in 1978. She graduated from Ankara Atatürk Anatolian High School in 1996. After receiving her B.S. degree in Statistics and minor degree in Economics from Middle East Technical University in 2001, she became a research assistant in the Department of Statistics. In the years 2004 and 2005, from Middle East Technical University, she received her M.S. degrees in Statistics and in Economics, respectively. Since 2006 May, she has been working as a statistician in Capital Markets Board of Turkey. She is married and has a daughter.