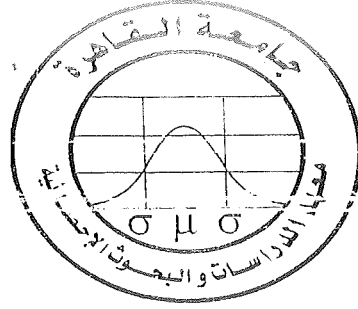




جامعة القاهرة
معهد الدراسات والبحوث الإحصائية

المؤتمر السنوى السابع والأربعون
لإحصاء

وعلوم الحاسب
وبحوث العمليات

إحصاء رياضى

٢٤-٢٧ ديسمبر ٢٠١٢

- [7] **Dunsmore, I. R.** (1983). The future occurrence of records. *Annals of the Institute Statistical Mathematics*, 17, 267-277.
- [8] **Feller, W.** (1966). *An introduction to probability theory and its applications*, 2, Second Edition, John Wiley, New York.
- [9] **Klimczak, M.** (2006). Prediction of k th records. *Statistics & Probability Letters*, 76, 117-127.
- [10] **Madi, M. T., Raqab, M. Z.** (2004). Bayesian prediction of temperature records using the Pareto model. *Environmetrics* 15, 701-710.
- [11] **Nagaraja, H. N.** (1988). Record values and related statistics. A review. *Communication in Statistics- Theory and Methods*, 17, 7, 2223-2238.
- [12] **Nagaraja, H. N.** (1984). Asymptotic linear prediction of extreme order statistics. *Annals of the Institute of statistical Mathematics*, 36, 289- 299.
- [13] **Sultan, K. S.** (2008). Bayesian Estimates Based on Record Values from the Inverse Weibull Lifetime Model .*Quality Technology and Quantitative Management*, 5, 363-374.
- [14] **Wald, A.** (1942).Setting of tolerance limits when the sample is large, *Annals of Mathematical Statistics*, 13, 389-399.

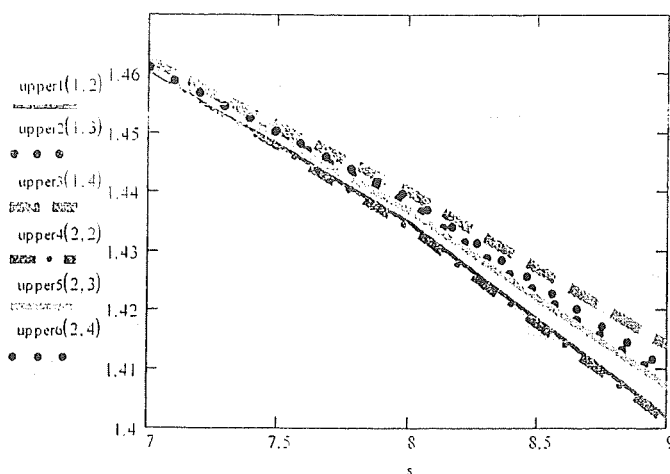


Fig. (4) relation between the record index and the upper bound of the interval for different combinations of (a, b), with $m = 6$ and β unknown.

References

- [1] Aboubecrine, M. (2010). On prediction of future record values. publié dans Advances and Applications in Statistics, 17, 79-93.
- [2] Ahsanullah, M. (1980). Linear prediction of record values for the two parameter exponential distribution. Annals Institute of Statistical Mathematics, 32, 363-368.
- [3] Ahsanullah, M. (1994). Record statistic. Nova Science Publishers, Inc.
- [4] AL-Hussaini, E. K. and Ahmed, A. A. (2003). On Bayesian interval prediction of future records. Test, 12, 1-20.
- [5] Arnold, C. B.; Balakrishnan, N. and Nagaraja, H.N. (1998). Records. John Wiley & Sons, Inc.
- [6] Chandler, K. N. (1952). The distribution and frequency of record values. Journal of Royal Statistics Society, series B, 220-228.

interval for different combinations of (a, b), with $m = 5$ and β unknown.

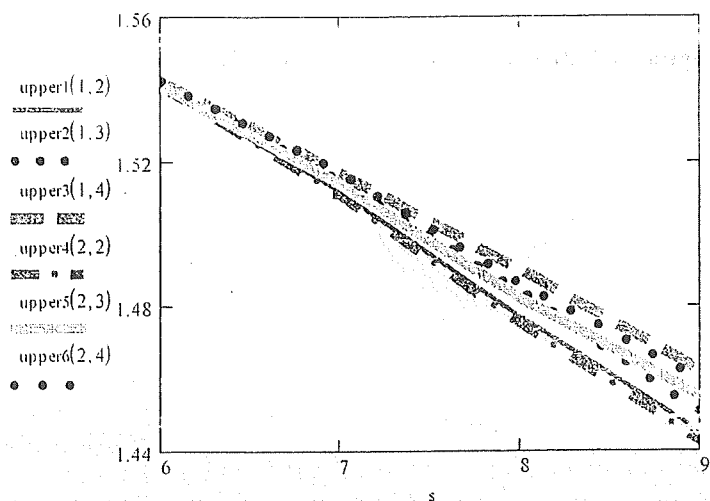


Fig. (2) relation between the record index and the upper bound of the interval for different combinations of (a, b), with $m = 5$ and β unknown.

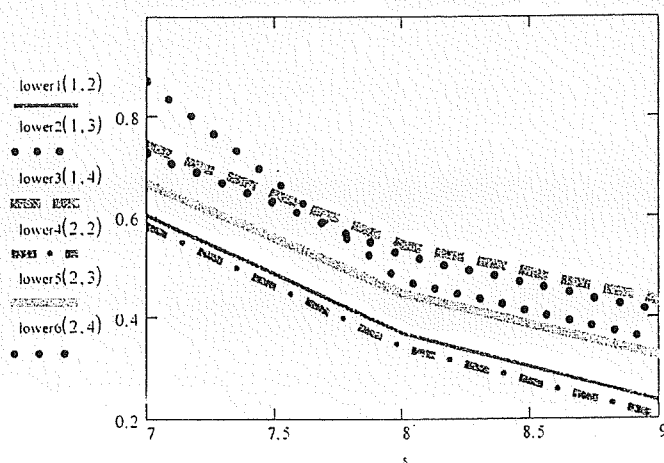


Fig. (3) relation between the record index and the lower bound of the interval for different combinations of (a, b), with $m = 6$ and β unknown.

Table (3)

Bayesian Prediction Interval of the s^{th} Future Upper Record Values for Data Generated from Inverse Weibull Distribution with $\alpha = 4$, $\beta = 1$ for Different Prior Parameters When both Parameters are Unknown and $m = 6$ ($c = 3, d = 4$)

a	B	R_m	s	lower interval	upper interval
1	2	1.453	7	0.605	1.461
		1.123	8	0.37	1.435
		0.729	9	0.234	1.402
1	3	3	7	0.87	1.461
			8	0.474	1.438
			9	0.352	1.408
1	4	4	7	0.744	1.462
			8	0.547	1.44
			9	0.433	1.414

a	b	S	lower interval	upper interval
2	2	7	0.584	1.461
		8	0.343	1.434
		9	0.205	1.401
2	3	7	0.668	1.461
		8	0.45	1.437
		9	0.325	1.407
2	4	7	0.727	1.461
		8	0.525	1.439
		9	0.408	1.41

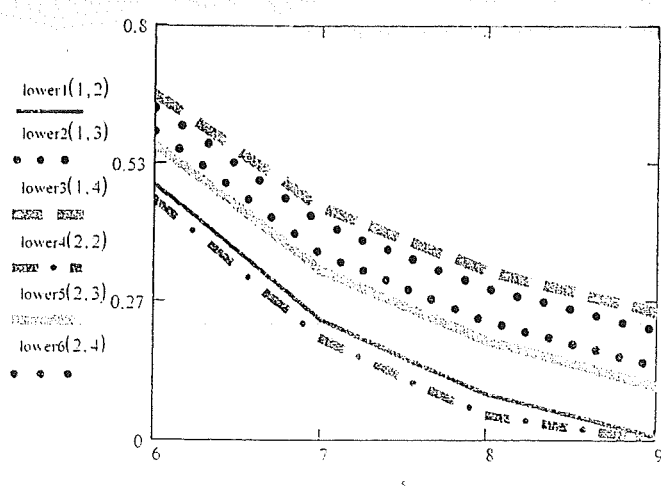


Fig. (1) relation between the record index and the lower bound of the

Table (2)

Bayesian Prediction Interval of the s^{th} Future Upper Record Values for Data Generated from Inverse Weibull Distribution with $\alpha = 4$, $\beta = 1$ for Different Prior Parameters When Both Parameters are Unknown and $m = 5$ ($c = 3, d = 4$)

a	b	R_m	s	lower interval	upper interval
1	2	1.465	6	0.491	1.541
		1.453	7	0.23	1.512
		1.123	8	0.087	1.477
		0.729	9	0.007	1.445
1	3		6	0.593	1.542
			7	0.357	1.515
			8	0.228	1.484
			9	0.141	1.45
1	4		6	0.664	1.542
			7	0.445	1.517
			8	0.326	1.49
			9	0.246	1.462
2	2		6	0.461	1.541
			7	0.194	1.511
			8	0.047	1.475
			9	0	1.442
2	3		6	0.566	1.541
			7	0.324	1.514
			8	0.191	1.482
			9	0.103	1.455
2	4		6	0.639	1.542
			7	0.415	1.517
			8	0.291	1.486
			9	0.21	1.459

(1) Using MathCad (2001) program, applying Equation (3.9) and (3.10) to obtain the 95% Bayesian prediction interval for the s^{th} future lower records for ($s = 6, 7, 8, 9$) and for several different values of the prior parameters a, b, c and d . Table (2) and Table (3) contained the results which show:

- (i) The values of the prior parameters a and b
- (ii) The 95% Bayesian prediction interval for the s^{th} records.

From the results in Figs. (1), (2), (3) and (4), one can see a monotonically decreasing lower-(upper) bound as related to s for fixed values of a irrespective of the value of b , while if a is kept fixed, increasing b increases the lower (upper) bound.

4. Numerical Illustrations

In this section, numerical results of the Bayesian and non Bayesian predictive intervals are obtained.

4.1 Numerical Results of the Prediction Interval using Wald's Prediction Interval

In this case, the calculations were carried out according to the following steps:

- (1) For given values of the Inverse Weibull parameters α and β , we generated a random variable X from the Inverse Weibull distribution (1.1) and selected the first 10 records ($R_m = (23.469 \ 6.808 \ 5.027 \ 2.221 \ 1.73 \ 1.546 \ 1.465 \ 1.453 \ 1.123 \ 0.729)$) by using the software package Mathcad (2001).
- (2) Consider the first six records as the observed lower records ($m = 5$), while the last four records as the unobserved records, which are to be predicted.
- (3) Applying Equations (2.12), to obtain the 95% equal tail Bayesian prediction interval for the s^{th} lower values for $m = 5$, ($m = 6$) and $s = 6, 7, 8, 9$, ($s = 7, 8, 9$).

Table (1) contains the 95% Bayesian prediction interval for the s^{th} records.

Form Table (1) shows that both the lower and upper confidence limits go down an either m or s (or both) increase(s).

Table (1)

Wald's Prediction Interval of the s^{th} Future Upper Record Values for Data
Generated from Inverse Weibull Distribution with $\alpha = 4$ $\beta = 1$.

m	R_m	s	lower interval	upper interval
5	1.465	6	0.491	1.541
	1.453	7	0.23	1.512
	1.123	8	0.067	1.477
	0.729	9	0.007	1.445
6		7	0.357	1.515
		8	0.228	1.484
		9	0.141	1.45

4.2 Numerical Results of the Prediction Interval when the two Parameters are Unknown

In this case, the two parameters are considered to be unknown with prior density function is given by (3.7). Many different values of prior parameters will be considered here.

After generating a random number from Inverse Weibull distribution and selecting the first 10 records, the first six records are considered as the informative values (i.e., $m = 5, m = 6$) and the next four records (three records) are considered as the non informative ($s = 6, 7, 8, 9, s = 7, 8, 9$) (as (1) and (2) in case 1) then the calculations are carried out according to the following steps:

Two sided Bayesian prediction bounds of the s^{th} function lower record values R_s ,

$s = m+1, m+2, \dots$ are obtained from (3.9) by solving the following Equations:

$$\begin{aligned} p(R_s < l(r_m) | \underline{r}) &= p(R_m - y_{m,s} < l(r_m) | \underline{r}) = p(y_{m,s} > R_m - l(r_m) | \underline{r}) \\ &= p(y_{m,s} > l'(r_m) | \underline{r}) = \int_{l'(r_m)}^{\infty} f(y_{m,s} | r_m) dy_{m,s} \\ &= A \int_{l'(r_m)}^{\infty} \int_c^d J(s, \beta, y_{m,s}) dy_{m,s} = \frac{\tau}{2}, \end{aligned} \quad (3.10)$$

and

$$\begin{aligned} p(R_s < u(r_m) | \underline{r}) &= p(R_m - y_{m,s} < u(r_m) | \underline{r}) = p(y_{m,s} > R_m - u(r_m) | \underline{r}) \\ &= p(y_{m,s} > u'(r_m) | \underline{r}) = \int_{u'(r_m)}^{\infty} f(y_{m,s} | r_m) dy_{m,s} \\ &= A \int_{u'(r_m)}^{\infty} \int_c^d J(s, \beta, y_{m,s}) dy_{m,s} = 1 - \frac{\tau}{2}. \end{aligned} \quad (3.11)$$

Numerical techniques are needed to solve Equation (3.10) and (3.11). Hence, the values of lower and upper bound for the s^{th} future record values R_s can be obtained from $l'(r_m)$ and $u'(r_m)$ as follows

$$l(r_m) = R_m - l'(r_m) \quad \text{and} \quad u(r_m) = R_m - u'(r_m).$$

$$f(y_{m,s} | r) \propto \int_0^\infty \int_0^\infty \int_0^\infty \alpha^{\nu(s)} \beta^{m+3} \left[(x_s - y_{m,s}) \prod_{i=0}^m r_i \right]^{-\beta-1} x_s^{-\beta(m+1)-1} \\ \left[(x_s - y_{m,s})^{-\beta} - x_s^{-\beta} \right]^{\varepsilon(s)} e^{-\alpha \left[(x_s - y_{m,s})^{-\beta} + (b + r_m^{-\beta}) \right]} dx_s d\alpha d\beta \\ , 0 < y_{m,s} < \infty.$$

On simplification,

$$f(y_{m,s} | r_m) \propto \int_0^\infty \int_0^\infty \beta^{m+3} \left[(x_s - y_{m,s}) \prod_{i=0}^m r_i \right]^{-\beta-1} \\ \frac{x_s^{-\beta(m+1)-1} \left[(x_s - y_{m,s})^{-\beta} - x_s^{-\beta} \right]^{\varepsilon(s)}}{\left[(x_s - y_{m,s})^{-\beta} + (b + r_m^{-\beta}) \right]^{\nu(s)+1}} dx_s d\beta \\ , 0 < y_{m,s} < \infty.$$

$$f(y_{m,s} | r_m) = A \int_0^\infty J(s, \beta, y_{m,s}) d\beta \quad , 0 < y_{m,s} < \infty, \quad (3.9)$$

where

$$J(s, \beta, y_{m,s}) = \int_0^\infty \beta^{m+3} \left[(x_s - y_{m,s}) \prod_{i=0}^m r_i \right]^{-\beta-1} \\ \frac{x_s^{-\beta(m+1)-1} \left[(x_s - y_{m,s})^{-\beta} - x_s^{-\beta} \right]^{\varepsilon(s)}}{\left[(x_s - y_{m,s})^{-\beta} + (b + r_m^{-\beta}) \right]^{\nu(s)+1}} dx_s \quad , y_{m,s} > 0,$$

and

$$A = \frac{1}{\int_0^\infty \int_0^\infty J(s, \beta, y_{m,s}) d\beta} d\beta.$$

Consider now the parameter α has a gamma prior distribution and the prior density of the shape parameter is the uniform distribution with density function given by

$$\pi(\alpha) \propto \alpha^{a-1} e^{-b\alpha}, \quad \alpha > 0, a, b > 0. \quad (3.5)$$

$$\pi(\beta) \propto \frac{1}{d-c}, \quad c < \beta < d. \quad (3.6)$$

Then the joint pdf of α and β will be obtained from (3.5) and (3.6) as follows

$$\pi(\alpha, \beta) \propto \frac{\alpha^{a-1} e^{-b\alpha}}{d-c}, \quad \alpha > 0, c < \beta < d. \quad (3.7)$$

Combining the prior pdf (3.7) with the likelihood function (2.3) then the bivariate posterior distribution of α and β is given by

$$\pi(\alpha, \beta | \underline{r}) \propto \beta^{m+1} \alpha^{m+a} \prod_{i=1}^m r_i^{-\beta-1} e^{-\alpha(b+r_m^{-\beta})}, \quad \alpha > 0, c < \beta < d. \quad (3.8)$$

Consider the probability density function of the difference between two record values (3.4). Then the Bayes predictive density function of s^{th} the future record R_s , $s = m+1, m+2, \dots$ is obtained by integrating the product of Equations (3.4) and (3.8) with respect to the parameters α and β as follows

Where $H(x) = -\ln F(x)$, $h(x) = \frac{-dH(x)}{dx}$. Now, let us consider the following transformation, $X_s = R_m$ and $Y_{m,s} = R_m - R_s$ then, the joint pdf of $X_s, Y_{m,s}$ will be

$$f(x_s, y_{m,s}) = \frac{[H(x_s)]^m}{m!} h'(x_s) f(x_s, y_{m,s}) \frac{[H(x_s - y_{m,s}) - H(x_s)]^{s-m-1}}{(s-m-1)!}, \quad 0 < y_{m,s} < x_s < \infty. \quad (3.2)$$

From (3.2) integrate out X_s , then the probability density function of the difference between two lower record values is given by

$$f(y_{m,s}) = \int_{x_s} \frac{[H(x_s)]^m}{m!} h'(x_s) \frac{[H(x_s - y_{m,s}) - H(x_s)]^{s-m-1}}{(s-m-1)!} f(x_s - y_{m,s}) dx_s, \quad 0 < y_{m,s} < \infty. \quad (3.3)$$

For the inverse Weibull distribution with probability density function (1.1), the pdf of $Y_{m,s} = R_m - R_s$ is given as follows

$$f(y_{m,s} | \alpha) = \frac{\alpha^{s+1} \beta^2}{m!(s-m-1)!} \int_{y_{m,s}}^{\infty} (x_s - y_{m,s})^{-\beta-1} x_s^{-\beta(m+1)-1} \left[(x_s - y_{m,s})^{-\beta} - x_s^{-\beta} \right]^{s-m-1} e^{-\alpha(x_s - y_{m,s})^{-\beta}} dx_s, \quad 0 < y_{m,s} < \infty. \quad (3.4)$$

Integration (3.4) needs numerical computation, i.e., there is no explicit form for it.

From (2.11), we conclude that the distribution of $R_s^{-\hat{\beta}}$ can be approximated by gamma distribution with parameters $(s+1, \hat{\alpha})$. From Wald (1942) and the fact that $2\hat{\alpha} R_s^{-\hat{\beta}}$ has a chi square distribution with $2(s+1)$ degrees of freedom, and hence a $100(1-\tau)\%$ prediction interval for R_s is $(l(R_s), u(R_s))$,

where

$$l(R_s) = \left[\frac{\chi^2_{(2(s+1)), (1-\tau/2)}}{2\hat{\alpha}} \right]^{-1/\hat{\beta}} \quad \text{and} \quad u(R_s) = \left[\frac{\chi^2_{(2(s+1)), \tau/2}}{2\hat{\alpha}} \right]^{-1/\hat{\beta}} \quad (2.12)$$

with $\hat{\alpha}$ and $\hat{\beta}$ being the maximum likelihood estimators of α and β .

3. Bayesian Prediction Interval for the sth Future Lower Record Values when both Parameters are Unknown

The prediction intervals provide bounds to contain the results of a future record, based upon the results of the previous record observed from the same distribution. Prediction problems comes up naturally in several real life situations, for example, Ahsanullah (1980) and Nagraja (1984). This section is devoted to derive Bayes predictive density function, which is necessary to obtain bounds for the predictive interval of future records.

Let $R_0, R_1, \dots, R_m$ be the first $(m+1)$ lower records, the joint pdf of R_m and R_s where $s < m$ is given as follows

$$f(r_m, r_s) = \frac{[H(r_m)]^m}{m!} h(r_m) \frac{[H(r_s) - H(r_m)]^{s-m-1}}{(s-m-1)!} f(r_s) \quad (3.1)$$

$$0 < r_s < r_m < \infty$$

$$(m+2) - \left[1 - \alpha r_s^{-\beta^*} \right] \beta^* \ln r_s - \beta^* \sum_{i=0}^m \ln r_i + \frac{\beta^* (s-m-1)}{(r_s^{-\beta^*} - r_m^{-\beta^*})} \left[r_m^{-\beta^*} \ln r_m - r_s^{-\beta^*} \ln r_s \right] = 0 \quad (2.8)$$

Numerical solution for (2.8) yields the predictive maximum likelihood estimator β^* (PMLE) of the shape parameter β . Back substitution by the value of β^* in equations (2.6) and (2.7) one can obtain the PMLE α^* of the parameter α and the maximum likelihood prediction (MLP) of the s^{th} future lower record values r_s^* .

2.2 Wald's prediction interval

Let $R_0, R_1, \dots, R_m$ be the first $(m+1)$ lower record from IWD (1.1). Let R_s be s^{th} future lower record values, the density function of R_s is

$$f(r_s) = \frac{\alpha^{s+1} \beta}{s!} r_s^{-(s+1)\beta-1} e^{-\alpha r_s^{-\beta}}, \quad r_s > 0. \quad (2.9)$$

It can be seen from (2.9) that the random variable $R_s^{-\beta}$ has a gamma distribution with parameters $(s+1, \alpha)$.

On replacing α and β by their maximum likelihood estimators $\hat{\alpha}$ and $\hat{\beta}$ given by (see, Sultan (2008))

$$\hat{\alpha} = (m+1) r_m^{\hat{\beta}}, \text{ and } \hat{\beta} = \frac{(m+1)}{\sum_{i=0}^m \ln r_i - (m+1) \ln r_m}. \quad (2.10)$$

It follows that the pdf of R_s can be approximated by

$$\hat{f}(r_s) = \frac{\hat{\alpha}^{s+1} \hat{\beta}}{s!} r_s^{-(s+1)\hat{\beta}-1} e^{-\hat{\alpha} r_s^{-\hat{\beta}}}, \quad 0 < r_s < \infty. \quad (2.11)$$

$$\text{where } L(\alpha, \beta) = \alpha^{m+1} \beta^{m+1} e^{-\alpha r_m^{-\beta}} \prod_{i=1}^m r_i^{-\beta-1} \quad (2.3)$$

is the likelihood function based on the first $(m+1)$ lower record values (see: Sultan (2008)) and $f(r_s | r_m)$ is the conditional density function of R_s given R_m , and substituting from (2.1) and (2.3) in (2.2) we have

$$L(\underline{R}, R_s; \alpha, \beta) \propto \beta^{m+2} \alpha^{s+1} (r_s \prod_{i=0}^m r_i)^{-\beta-1} (r_s^{-\beta} - r_m^{-\beta})^{s-m-1} e^{-\alpha r_s^{-\beta}} \quad (2.4)$$

Taking the logarithm for (2.4) we get

$$\ln L(\underline{R}, R_s; \alpha, \beta) \propto (s+1) \ln \alpha + (m+2) \ln \beta - (\beta+1) \left[\ln r_s + \sum_{i=0}^m \ln r_i \right] + (s-m-1) \ln \left[r_s^{-\beta} - r_m^{-\beta} \right] - \alpha r_s^{-\beta} \quad (2.5)$$

To obtain the maximum likelihood prediction for the s^{th} future lower record R_s , differentiate (2.5) with respect to α, β, r_s and equating with zero we get

$$\alpha^* = \frac{(s+1)(1-m\beta^* - \beta^*)}{(1-s\beta^*)r_m^{-\beta^*}} \quad (2.6)$$

$$r_s^* = \left[\frac{1-s\beta^*}{1-\beta^*(m+1)} \right]^{\frac{-1}{\beta^*}} r_m \quad (2.7)$$

and

Record values and associated statistics are of great importance in several real live problems involving weather, economic, and support data. The statistical study of record values started with Chandler (1952) and spread in different directions. Interested readers may refer to Feller (1966) Dunsmore (1983), Nagaraja (1988), Ahsanullah (1994), Arnold et al (1998) ,AL-Hussaini and Ahmed(2003), Madi and Raqab (2004) , Klimczak (2006) and Aboubecrine (2010) for a review of developments in this area of research.

2. Non Bayesian Prediction interval for the s^{th} future lower record values for inverse Weibull distribution

In this section we discussed two methods of non Bayesian prediction for the s^{th} future lower record values based on the first $(m+1)$ lower record values for the inverse Weibull distribution.

2.1 Maximum likelihood prediction interval

Let $R_0, R_1, \dots, R_m$ be the first $(m+1)$ lower record from IWD (1.1). Let R_s be s^{th} future lower record values. Ahsanullah (1994) obtained the following conditional density function of R_s given R_m

$$f(r_s | r_m) = \frac{\alpha^{s-m} \beta (r_s^{-\beta} - r_m^{-\beta})^{s-m-1}}{(s-m-1)!} r_s^{-\beta-1} e^{-\alpha(r_s^{-\beta} - r_m^{-\beta})} \quad 0 < r_s < r_m < \infty \quad (2.1)$$

The likelihood function $L(\underline{R}, R_s; \alpha, \beta)$ which based on the first $(m+1)$ lower record and the s^{th} future lower record values R_s , can be factored in two parts as follows

$$L(\underline{R}, R_s; \alpha, \beta) \propto L(\alpha, \beta) f(r_s | r_m) \quad (2.2)$$

Prediction Interval of Lower Record Value for the Inverse Weibull Distribution

W. Yahia¹ A.N. Ahmed²

Abstract

This article discusses Bayesian and non- Bayesian prediction interval for the s^{th} future lower record values based on the first $(m+1)$ lower record values for two parameters Inverse Weibull distribution. Non Bayesian prediction interval is based on two different approaches maximum likelihood prediction and Wald's prediction. Bayesian prediction interval is based on the distribution of the difference between two lower records. Finally, numerical illustrations for the new results are presented.

Key words: Inverse Weibull distribution, lower recode values, non-Bayesian prediction, Bayesian prediction.

1. Introduction:

The Inverse Weibull Distribution (IWD) plays an important role in many applications, including the dynamic components of diesel engines and several data set such as the times to breakdown of an insulating fluid subject to the action of a constant tension.

The IWD probability density function (pdf) with two parameters is given by

$$f(x) = \alpha \beta x^{-\beta-1} e^{-\alpha x^{-\beta}}, \quad x \geq 0, \beta > 0, \alpha > 0,$$

and cumulative distribution function (cdf)

$$F(x) = e^{-\alpha x^{-\beta}}, \quad x \geq 0, \beta, \alpha > 0. \quad (1.2)$$

where α is the scale parameter, β is the shape parameter.

¹²Department of Mathematical Statistics, ISSR, Cairo University, Cairo, Egypt

第

一

卷

上

الملخص العربي

إن الأبحاث التي تناولت مشكلة البيانات المفقودة قد تزايدت في الأونة الأخيرة حيث يعتبر الاستدلال الإحصائي في حالة البيانات التي تحتوي على قيم مفقودة من أهم التطبيقات العملية وذلك لأن القيم المفقودة سواء كانت مخطئة أو غير مخطئة تقابلنا كثيراً في الحياة العملية.

فالبيانات المفقودة تعني عدم وجود قيمة مخزنة للمتغير في البيانات الحالية للمشاهدات، فمشكلة البيانات المفقودة مشكلة متكررة لأنها تسبب التحيز أو إنها تؤدي إلى تحليلات غير كفاء. وطبقاً Roth وآخرون (1999) البيانات المفقودة لها تأثيرات سلبية ضخمة منها:-

1- أن لها أثر سلبي على القوة الأحصائية حيث إنها تؤثر على حجم العينة، وحجم العينة هو أحد العوامل الأساسية التي تؤثر على القوة الإحصائية.

2- أنها قد تؤدي إلى تقديرات متحيزة بأشكال مختلفة فقد تؤثر بالسلب على مقاييس النزعة المركزية ومقاييس التشتت ومقاييس الارتباط.

ووصولاً إلى هدف تقدير البيانات المفقودة تم عرض طرقاً كثيرة لمعالجة البيانات المفقودة حيث تم عرض تلك الطرق من جانبين جانب تناول الطرق التقليدية والآخر تناول الطرق الحديثة ومن الطرق التقليدية التي تم عرضها في البحث Complete Case Analysis ومن الطرق الحديثة طريقة (Maximum likelihood (ML) وطريقة EM algorithm والتي تتكون من خطوتين Expectation step و maximization step ويتم تكرار تلك الخطوتين في كل محاولة ولمحاولات عديدة حتى يتم التوصل إلى التقارب وذلك للوصول إلى مقدرات النموذج.

وأخيراً تعاملنا مع طريقة Bayesian وأوضحنا أنها تختلف عن الطرق التقليدية حيث أن الطريقة التقليدية تعتبر مصدر البيانات الوحيد هي المعلومات الحالية كما تعتبر المعلمة تكون ثابتة وأن البيانات المشاهدة تكون متغير عشوائي. أما في طريقة Bayesian عملية التقدير تنتج من خلال دمج البيانات الحالية مع المعلومات الناتجة من الخبرات السابقة.

ثم قمنا بعد ذلك بتطبيق طريقة Bayesian في نموذج الانحدار المتعدد في حالة احتواء البيانات على قيم مفقودة وتم عمل مقارنة بين طريقتي EM, Bayesian في نموذج الانحدار عند احتواء على بيانات مفقودة وتوصلنا أن طريقة Bayesian تعطي نتائج أفضل ولكن تحتاج إلى وقت أطول حيث يتم إجراء مجموعة من المحاولات للوصول إلى عملية التقارب للحصول على مقدرات النموذج كما تم استخدام طريقة Bayesian في نماذج السلاسل الزمنية وأوضحنا أهمية استخدام هذه الطريقة في تلك النماذج وللتبسيط قمنا باستخدام نموذج الانحدار الذاتي من الدرجة الأولى (AR(1) في حالة احتواء البيانات على قيم مفقودة.

وتبين أن طريقة Bayesian طريقة جيدة في تحليل السلاسل الزمنية حيث إنها تعطي نتائج أفضل خاصة عندما يكون حجم البيانات كبير وعندما توجد اختلافات وفتية (موسمية- وفتية- عشوائية) لذلك فطريقة Bayesian يمكن أن يستخدم في تحليل أحوال الطقس وأيضاً في تحليل البيانات الناتجة عن الرادارات خاصة في حالة وجود مجموعة من العوائق مثل الأمطار والفيضانات.

- Rubin, D.B., and Schafer, J.L. (1990)**, Efficiently Creating Multiple Imputation for Incomplete Multivariate Normal Data, Proceeding of the Statistical Computing Section. American Statistical Association.
- Swamy, P. A.V. B., and Mehta, J.S. (1975)**, "On Bayesian Estimation of Seemingly Unrelated Regression When some Observations are missing", Journal of Econometrics, 3, pp. 157 - 169.
- Tanner, M. (1996)**, Tools For Statistical inference, Third Edition, New York: Springer Verlag.
- Tanner, M.A. (1991)**, "Tools for Statistical inference: Observed Data and Data Augmentation Methods", New York: Springer Verlag.
- Tanner, M.A., and Wong, W.H. (1987)**, "The calculation of posterior distributions by data augmentation", Journal of the American Statistical Association, 82, pp. 528 - 550.

Reference

- Box, G.E.P., and Tiao, G.C. (1973), Bayesian Inference in Statistical Analysis. Addison-Wesley, Reading, Massachusetts.
- Chen, C.F. (1986), "A Bayesian approach to nested missing data problem", in Bayesian Inference and Decision Techniques, eds. P. Goel and A. Zellner, New York: Elsevier Press, pp. 355 - 361.
- Chen, T.,A. and Liu, J., J. M. (1996), "E-e modeling of turbulent airflow downwind of a model forest edge", Boundary-Layer Meteorology, 77:21- 44.
- Dickey, J. M. (1967), "Matric-variate generalizations of the multivariate distribution t distribution and the inverted multivariate distribution", Annals of Mathematical Statistics, 38, pp. 511-518.
- Guttman, I., and Menzefrieke, U. (1983), "Bayesian inference in multivariate regression with missing observation on the response variables," Journal of Business and Economic Statistics, 1. pp. 239 - 248.
- Little, R.J.A. (1988b), "Small sample inference about means from bivariate normal data with missing values", Computational Statistics and Data Analysis, 7, pp. 161 - 178.
- Mabdy, M. (2005), "The Extensions of the Expectation– Maximization (EM) Using a Bayesian Approach", MSC. Thesis, Zagazig University, Benha Branch.
- Press, S. J. and Scott, A. J. (1976), "Missing variables in Bayesian regression, II, Journal of the American Statistical Association, 71, pp. 366 -369.
- Rubin, D.B. (1978), Multiple Imputations in Sample Surveys – A phenomenological Bayesian Approach to Nonresponse, Imputation and Editing of Faulty on Missing Survey Data, U.S. Department of Commerce, pp. 1 - 23.
- Rubin, D.B. (1987a), Multiple Imputation for Nonresponse in Surveys,. New York: Wiley.
- Rubin, D.B. (1987b), "Comment on "The calculation of posterior distributions by data augmentation" by M.A. Tanner and W.H. Wong." Journal of the American Statistical Association, 81, pp. 366 - 374.

Applying the chain rule to $p(y_{OBS}, y_{MIS} | \phi_1, y_1)$ we obtain

$$p(y_{OBS}, y_{MIS} | \phi_1, y_1) = \prod_{t=2}^T p(y_t | y_{t-1}, y_{t-2}, \dots, y_1, \phi_1) \quad (4.1)$$

Under the AR (1) assumption, we have:

$$p(y_t | y_{t-1}, \dots, y_2, y_1, \phi_1) = p(y_t | y_{t-1}, \phi_1)$$

then the equation (4.1) implies

$$p(y_{OBS}, y_{MIS} | \phi_1, y_1) = \prod_{t=2}^T p(y_t | y_{t-1}, \phi_1)$$

Under the assumption of normality of AR (1) model then ,

$$p(y_t | y_{t-1}, \phi_1) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left[-\frac{1}{2\sigma^2} (y_t - \phi_1 y_{t-1})^2\right]$$

and then the posterior p.d.f. is as follows:-

$$p(\phi_1 | y_{OBS}, y_1) \propto$$

$$p(\phi_1) \prod_{\ell \in G} \left\{ \left(\sum_{t=0}^{\ell} \phi_1^{2t} \right)^{-\frac{\ell |D_{\ell}|}{2}} \exp \left[\frac{-1}{2\sigma^2} \left(\sum_{t=0}^{\ell} \phi_1^{2t} \right)^{-1} \begin{bmatrix} 1 \\ -\phi_1^{\ell+1} \end{bmatrix}^T S_{\ell} \begin{bmatrix} 1 \\ -\phi_1^{\ell+1} \end{bmatrix} \right] \right\}$$

Where

1. D is the set $OBS \cup \{1\}$
2. $GAP(i)$ be the length of a missing-data sequence after y_i , and $i \in D$ and
3. $y_i, y_{i+GAP(i)+1} \in OBS$ and $y_{i+1}, y_{i+2}, \dots, y_{i+GAP(i)} \in MIS$
4. G is the set of all sizes of missing-data sequences for every i.i.e. $G = \{\ell \in \mathbb{R} \exists i \in D | \ell = GAP(i)\}$
5. D_{ℓ} be the subset of D

$$D_{\ell} = \{i \in D | \ell = GAP(i)\}$$

6. The statistic S_{ℓ} is as follows :-

$$S_{\ell} = \sum_{j \in D_{\ell}} \begin{bmatrix} y_{j+GAP(j)+1}^2 & y_{j+GAP(j)+1} y_j \\ y_{j+GAP(j)+1} y_j & y_j^2 \end{bmatrix}$$

7. Finally the likelihood function is as follows:-

$$p(y_{OBS} | \phi_1, y_1) \propto$$

$$\prod_{\ell \in G} \left\{ \left(\sum_{t=0}^{\ell} \phi_1^{2t} \right)^{-\frac{\ell |D_{\ell}|}{2}} \cdot \exp \left[\frac{1}{2\sigma^2} \left(\sum_{t=0}^{\ell} \phi_1^{2t} \right)^{-1} \begin{bmatrix} 1 \\ -\phi_1^{\ell+1} \end{bmatrix}^T S_{\ell} \begin{bmatrix} 1 \\ -\phi_1^{\ell+1} \end{bmatrix} \right] \right\}$$

and $|\cdot|$ denotes the number of elements of a set.

the random shocks $\epsilon_1, \epsilon_2, \epsilon_3, \dots$ in different time period are assumed to be statistically independent of each other.

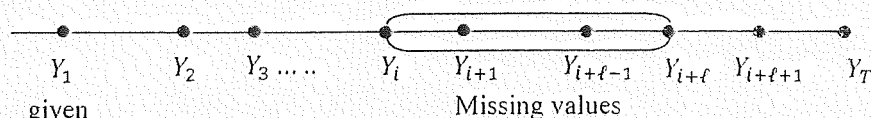
To be able to analyze estimator, it is useful to know exactly the posterior probability density function.

When the whole sample of data is available, the solution is well known. The problem has a finite sufficient statistic and it is possible to compute the posterior probability density function. But when some data are missing, and there are some gaps in data sample, very little can be said about posterior p.d.f.

In this section we try to deal with the identification of posterior p.d.f. of the parameter of AR (1) from incomplete data sample.

To do this, we consider only certain data to be observed while the other missing, such that the observed data is indicated by $OBS \subset \{2, 3, \dots, T\}$, while the missing data by $MIS \subset \{\ell, \ell + 1, \dots, T\}$ and we denote by y_{OBS} and y_{MIS} the observed and missing data respectively. The initial value y_1 is not included in data and its value is given.

We assume, without loss of generality, that the last data Y_T is observed ($T \in OBS$)



It should be noted that our problem is based on the Bayesian point of view.

Then it means that ϕ_1 is a random variable and we are interested in the posterior distribution

$$p(\phi_1 | y_{OBS}, y_1)$$

For the sake of completeness, the indices of observed data are assumed to be random and independent on value of data. This missing-data mechanism is called "missing at random".

Now, by using Bayes rule we have:

$$p(\phi_1 | y_{OBS}, y_1) \propto p(\phi_1) p(y_{OBS} | \phi_1, y_1)$$

Where $p(\phi_1)$ is prior information about ϕ_1 . Under the assumption of independence of OBS and y_1 the likelihood function $p(y_{OBS} | \phi_1, y_1)$ follows:-

$$p(y_{OBS} | \phi_1, y_1) = \int_{R[MIS]} p(y_{OBS}, y_{MIS} | \phi, y_1) dy_{MIS}$$

conditional on the observed data. The posterior distribution of the parameters is then represented by values of the parameters that are sampled after the chain has converged. This two-step iteration procedure is called data augmentation (Tanner, 1996). So in next we apply the Bayesian method on the data in example (1) and then we make comparison between this method and EM method.

Example (2): The Bayesian method is used here to be applied on each iteration to know the fitting goodness in each phase until the convergence phase in estimating the missing data.

Our results of this example is as follows

After 25 iterations, we have:-

$$\begin{aligned} \hat{\beta}_0 &= 81.59019 & \hat{\beta}_1 &= 1.577932 & \hat{\sigma}_{y,x}^2 &= 15.2722 \\ \hat{Y}_6 &= 132.10628 & \hat{Y}_9 &= 125.7042 & \hat{Y}_{13} &= 142.9751 \end{aligned}$$

The following table summarizes the estimated missing observations by using both EM algorithm and Bayesian approach.

Table(2) The estimated missing observations by using both EM algorithm and Bayesian approach.

Methods	Y_6	Y_9	Y_{13}
EM algorithm	132.908227	127.8243671	146.4651898
Bayesian approach	132.10628	125.7042	142.9751

From these comparisons between these methods we note that the EM method is faster in making convergence to estimate the missing values than the Bayesian methods but the last one achieved better estimations especially if the size of the data is large.

4. Autoregressive model AR (1) process with missing data.

To introduce autoregressive model of order one we consider, first, the non-seasonal autoregressive model of order one as

$$Y_t = \phi_1 Y_{t-1} + \epsilon_t$$

This model says that y_t is a constant multiple of Y_{t-1} which is the time series value in the previous period, plus a random shock ϵ_t that describes the effect of all factors other than Y_{t-1} on Y_t . The constant ϕ_1 relating to Y_t to Y_{t-1} is an unknown parameter that must be estimated from sample data. The random shock ϵ_t is a value that is assumed to have been randomly selected from a normal distribution that has mean zero and a variance that is the same for each and every time period. Moreover,

We note here that as $A \rightarrow 0, M \rightarrow \hat{B}$ and $W \rightarrow 0$, as may easily be verified, so that $(v + v_0)\mathcal{G} \rightarrow vS$ as $v_0 \rightarrow 0$, and $A \rightarrow 0$. Upon integration with respect to Σ^{-1} in (4.3), we have the posterior of B given there are no missing data can be written as

$$p(B|Y) \propto |I_p + \frac{1}{v+v_0} \times \mathcal{G}^{-1}(B-M)^T[C+A](B-M)|^{-(n+v_0+\delta)/2} \quad (3.6)$$

which is to say, that $p(B|Y)$ is a matrix-variate t distribution, discussed, for example, by Dickey (1967) and. Box and Tiao (1973, Sec. 8.4.3).

The predictive distribution of $Z_1, f(Z_1|Y_1, Y_2)$ is easily obtained since $f(Z_1|Y_1, Y_2) \propto \iint l(Z_1, Y_1, Y_2|X, B, \Sigma^{-1}) \times p(B, \Sigma^{-1}) d\Sigma^{-1} dB$,

and using the work that leads to (3.6), along with properties of the matrix-variate t density, we easily find

$$f(Z_1|Y_1, Y_2) \propto |(v + v_0)\mathcal{G}|^{-(n+v_0-k+\delta)/2}. \quad (3.7)$$

We recall from (3.5) that the $(p \times p)$ matrix $(v + v_0)\mathcal{G} = v_0S_0 + vS + W$.

Now suppose we partition $\tilde{\mathcal{G}} = (v + v_0)\mathcal{G}$ as

$$(v + v_0)\mathcal{G} = \tilde{\mathcal{G}} = \begin{bmatrix} \tilde{\mathcal{G}}_{11} & \tilde{\mathcal{G}}_{12} \\ \tilde{\mathcal{G}}_{21} & \tilde{\mathcal{G}}_{22} \end{bmatrix},$$

where $\tilde{\mathcal{G}}_{11}$ is $p_1 \times p_1$, $\tilde{\mathcal{G}}_{12}$ is $p_1 \times p_2$, where $p_1 + p_2 = p$, and so on where $\tilde{\mathcal{G}}_{11}$ does not involve the missing data Z_1 , so that we may write (3.7) as

$$f(Z_1|Y_1, Y_2) \propto |\tilde{\mathcal{G}}_{22} - \tilde{\mathcal{G}}_{21} \tilde{\mathcal{G}}_{11}^{-1} \tilde{\mathcal{G}}_{12}|^{-(n+v_0-k+\delta)/2}.$$

Now, to compute the posterior distribution a Markov chain Monte Carlo (MCMC) resampling algorithm sampler can be used in Bayesian analysis. MCMC techniques can also be used to adjust for data missing at random. The missing data are just regarded as more unknown parameters to be resampled. In some instances, implementation is almost seamless. The change to the program code given to a software package such as WinBUGS can sometimes be as simple as designating that an observation is missing, as opposed to having a known value, and giving it an initial value. In essence, two iterative steps are used by the resampling algorithm to compute posterior distribution in the presence of data missing at random. First, the missing data are filled in with draws from their predictive distribution conditional on last sampled values of the parameters. Then the parameter values are drawn from their posterior distribution conditional on the completed data set of observed and filled-in missing data. Iterating between draws of missing data and parameter values is a Markov-chain sequence that eventually converges to draws from the joint distribution of the parameters and the missing data

The likelihood function for (3.1) can be written as

$$l(Z_1, Y_1, Y_2 | X, B, \Sigma^{-1}) \propto |\Sigma^{-1}|^{n/2} \times \exp \left\{ -1/2 \operatorname{tr} \Sigma^{-1} \left[\nu S + (B - \hat{B})^T C (B - \hat{B}) \right] \right\} \quad (3.2)$$

where $C = X^t X$, $\hat{B} = C^{-1} X^t Y$ and $\nu = n - p - k + 1$,

with

$$\begin{aligned} \nu S &= (Y - X\hat{B})^T (Y - X\hat{B}), \\ &= (Y - XC^{-1}X^T Y)^T (Y - XC^{-1}X^T Y) \\ &= (Y - X(X^T X)^{-1}X^T Y)^T (Y - X(X^T X)^{-1}X^T Y) \\ &= Y^T Y - Y^T X C^{-1} X^T Y - Y^T X C^{-1} X^T Y + Y^T X C^{-1} X^T X C^{-1} X^T Y \\ &= Y^T Y - 2Y^T X C^{-1} X^T Y + Y^T X C^{-1} C C^{-1} X^T Y \\ &= Y^T (I - X C^{-1} X^T)^T (I - X C^{-1} X^T) Y \\ &= (I - X C^{-1} X^T)^T (I - X C^{-1} X^T) \\ &= I - X C^{-1} X^T - X C^{-1} X^T + X C^{-1} X^T X C^{-1} X^T \\ &= [I - X C^{-1} X^T - X C^{-1} X^T + X C^{-1} X^T] \\ &= (I - X C^{-1} X^T) = R(X) \end{aligned}$$

$R(X)$ should be idempotent matrix

i.e.

$$\nu S = Y^T R(X) Y$$

We assume that the prior information for (B_0, Σ^{-1}) is given by the natural conjugate prior for this situation, namely,

$$p(B_0, \Sigma^{-1}) \propto |\Sigma^{-1}|^{\delta/2} |\Sigma^{-1}|^{(\nu_0 - p - 1)/2} \times \exp \left\{ -1/2 \operatorname{tr} \Sigma^{-1} [\nu_0 S_0 + (B - B_0)^T A (B - B_0)] \right\} \quad (3.3)$$

where $\delta = \delta(A) = 0$, if $A = 0$, and $\delta = 1$, otherwise.

However, using (3.2) and (3.3) yields, on application of Bayes' theorem,

$$p(B, \Sigma^{-1} | Y) \propto |\Sigma^{-1}|^{(n + \nu_0 - p - 1 + \delta)/2} \times \exp \left\{ -1/2 \operatorname{tr} \Sigma^{-1} [(\nu + \nu_0) \varphi + (B - M)^T (A + C) (B - M)] \right\} \quad (3.4)$$

Where,

$$M = [C + A]^{-1} (C\hat{B} + AB_0),$$

and, defining W as

$$W = (\hat{B} + B_0)^T (A^{-1} + C^{-1})^{-1} (\hat{B} + B_0),$$

then

$$(\nu + \nu_0) \varphi = \nu S + \nu_0 S_0 + W. \quad (3.5)$$

$$\begin{aligned}
 Y_{n \times p} &= \begin{array}{c} \begin{array}{cc} & p_1 \end{array} & \begin{array}{cc} & p_2 \end{array} \\ \left[\begin{array}{cccc|ccc} y_{11} & y_{12} & \cdots & y_{1p_1} & y_{1(p_1+1)} & \cdots & y_{1(p_1+p_2)} \\ y_{21} & y_{22} & \cdots & y_{2p_1} & y_{2(p_1+1)} & \cdots & y_{2(p_1+p_2)} \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ y_{n_1,1} & y_{n_1,2} & \cdots & y_{n_1 p_1} & y_{n_1(p_1+1)} & \cdots & y_{n_1(p_1+p_2)} \\ y_{n_1+1,1} & & & y_{n_1+n_2,p_1} & y_{n_1+1(p_1+1)}^* & \cdots & y_{n_1+1(p_1+p_2)}^* \\ \vdots & & & \vdots & \vdots & \vdots & \vdots \\ y_{n_1+n_2,1} & \cdots & \cdots & y_{n_1+n_2,p_1} & y_{n_1+n_2,1}^* & \cdots & y_{n_1+n_2,p_1+p_2}^* \end{array} \right] \begin{array}{c} \left. \begin{array}{c} \\ \\ \\ \end{array} \right\} n_1 \\ \\ \left. \begin{array}{c} \\ \\ \\ \end{array} \right\} n_2 \end{array} \\ \\ = \begin{bmatrix} Y_1 & Z_1 \\ Y_1^* & Y_2 \end{bmatrix}
 \end{array}$$

$$Y_1 = \begin{bmatrix} y_{11} & y_{12} & \cdots & y_{1p_1} \\ y_{21} & y_{22} & \cdots & y_{2p_1} \\ \vdots & \vdots & \vdots & \vdots \\ y_{n_1,1} & y_{n_1,2} & \cdots & y_{n_1 p_1} \end{bmatrix}_{n_1 \times p_1}$$

$$Y_1^* = \begin{bmatrix} y_{11}^* & y_{12}^* & \cdots & y_{1p_2}^* \\ \cdots & \cdots & \cdots & \cdots \\ y_{n_2,1}^* & y_{n_2,2}^* & \cdots & y_{n_2 p_2}^* \end{bmatrix}_{n_2 \times p_2}$$

$$Y_2 = \begin{bmatrix} y_{(n_1+1)(p_1+1)} & \cdots & y_{(n_1+1)(p_1+p_2)} \\ \vdots & \vdots & \vdots \\ y_{(n_1+n_2)(p_1+1)} & \cdots & y_{(n_1+n_2)(p_1+p_2)} \end{bmatrix}$$

And Z_1 is an $(n_1 \times p_2)$ matrix, with $n_1 + n_2 = n$, $p_1 + p_2 = p$

$$Z_1 = \begin{bmatrix} z_{11} & z_{12} & \cdots & z_{1p_2} \\ \vdots & \vdots & \vdots & \vdots \\ z_{n_1,1} & z_{n_1,2} & \cdots & z_{n_1 p_2} \end{bmatrix} = \begin{bmatrix} y_{1(p_1+1)} & \cdots & y_{1(p_1+p_2)} \\ \vdots & \vdots & \vdots \\ y_{n_1(p_1+1)} & \cdots & y_{n_1(p_1+p_2)} \end{bmatrix}$$

and X and B conformably partitioned. We assume that only Z_1 is unobserved. To get the posterior distribution for B given the observed data, we use the fact that

$$p(B|Y_1, Y_2) = \int p(B|Y) f(Z_1|Y_1, Y_2) dZ_1,$$

where $p(B|Y)$ is the posterior distribution given there are no missing data and $f(Z_1|Y_1, Y_2)$ is the predictive distribution for the missing observations given the observed data.

3. Bayesian Regression for Missing Data

The Bayesian approach has been applied to multivariate problems with incomplete dependent variable (Chen, 1986; Guttman and Menzefricke, 1983 ; Little, 1988 ; Press and scott, 1976 ; Rubin, 1977, 1978, 1987a and Swamy and Mehta, 1975) .For general patterns of missing data, more complex simulation techniques (tanner, 1991), such as data augmentation (Tanner and wong, 1987), and the importance sampling (Rubin, 1987b), are needed to simulate the posterior distribution . Rubin and Schafer (1990) applied these ideas to the normal model. Chen and liu (1996) proposed algorithms based on random draws from predictive distributions of unknown quantities (missing values). This procedure can either be iterative, which is a special variation of Gibbs sampler, or be sequential, which is a variation of sequential imputation.

3.1 Bayesian Regression Model for Missing Data

According to Guttman and Menzefricke (1983), suppose the multivariate linear model:

$$Y = XB + \epsilon \quad (3.1)$$

Where,

Y is an $(n \times p)$ matrix of n observations on a p -dimensional response variable.

X is an $(n \times k)$ design matrix, B is an $(k \times p)$ matrix of regression coefficients.

Then,

$$\begin{aligned} Y &= \begin{pmatrix} Y_1 & Z_1 \\ Y_2 \end{pmatrix} = \begin{pmatrix} Y_1 & Z_1 \\ Y_1^* & Y_2 \end{pmatrix} \\ &= XB + \epsilon = \begin{pmatrix} X_1 \\ X_2 \end{pmatrix} (B_1 \ B_2) + \epsilon, \end{aligned}$$

with Y_1 is an $(n_1 \times p_1)$ matrix, Y_2 is an $(n_2 \times p_2)$ matrix, with

$$n_1 + n_2 = n \quad p_1 + p_2 = p$$

The algorithm can be simplified by noting that (2.1) does not involve $(\hat{\sigma}_{y,x}^t)^2$, and note that at convergence point we have:

$$(\hat{\sigma}_{y,x}^{t+1})^2 = (\hat{\sigma}_{y,x}^t)^2 = \hat{\sigma}_{y,x}^2.$$

By using

$$\hat{\sigma}_{y,x}^2 = \frac{1}{n} \sum_{i=1}^m (Y_i - X_i \hat{\beta}_{y,x}^{(t)})^2 + \frac{(n-m)}{n} (\hat{\sigma}_{y,x}^t)^2.$$

We can get,

$$\hat{\sigma}_{y,x}^2 = \frac{1}{m} \sum_{i=1}^m (Y_i - X_i \hat{\beta}_{y,x}^{(t)})^2.$$

Again as we know, we can prove that:

$$\hat{\sigma}_{\beta}^2 = \hat{\sigma}_{y,x}^2 (X^T X)^{-1}.$$

Consequently, the EM iterations can ignore the M-step estimation of $\hat{\sigma}_{y,x}^2$ and E-step estimation of $E(Y_i^2 | Y_{obs}; \theta_{y,x}^{(t)})$ and then find $\hat{\beta}_{y,x}$ by iteration.

Example(1): Consider 15 observations taken from 2 variables and containing three missing values. We find the estimates of missing values, where missing values denotes* and we also estimate the parameters of a linear model with X as an exogenous variable and Y being endogenous variable, so we have,

E – step, estimation of missing data:

$$Y_i = \begin{cases} Y_i & ; i = 1, \dots, m \\ \beta_0^{(t)} + \beta_1^{(t)} X_i & ; i = m + 1, \dots, n \end{cases}$$

Table (1): Data for Example

X	24	26	20	25	35	32	32	25	29	30	26	33	40	24	28
Y	115	125	110	124	140	*	130	119	*	122	125	138	*	127	125

Start with $\beta_0^{(0)} = 78.68045303$ and $\beta_1^{(0)} = 1.69461757$, the following table shows results were obtained.

So that, the estimated parameters are

$$\hat{\beta}_0 = 66.27477, \quad \hat{\beta}_1 = 1.267281 \text{ and } \hat{\sigma}_{y,x}^2 = 14.54004,$$

and the estimated missing values are

$$\hat{y}_6 = 132.9082278, \quad \hat{y}_9 = 127.8243671, \quad \hat{y}_{13} = 146.4651898.$$

Consequently,

$$E(Y^2 | X, Y_{obs}; \theta_{y,x}^{(t)}) = V(Y) + [E(Y | X, Y_{obs}; \theta_{y,x}^{(t)})]^2$$

$$= (\sigma_{y,x}^t)^2 + (X_i \hat{\beta}_{y,x}^{(t)})^2$$

Where X is the $(n \times (p+1))$ matrix of X values. Let $\underline{Y}$ be the $(n \times 1)$ vector of Y values and $\underline{Y}^{(t)}$ be the $(n \times 1)$ vector of Y with missing components Y_i replaced by estimates from the E-steps at iteration t .

2. M – step

The M-Steps can be used to calculate $\hat{\beta}_{y,x}^{(t)}$ and $(\hat{\sigma}_{y,x}^{(t)})^2$, we start to estimate $\hat{\beta}_{y,x}^{(t)}$

since,

$$\underline{\epsilon} = (\underline{Y} - X \hat{\beta}_{y,x}^{(t)}),$$

then,

$$\underline{\epsilon} \underline{\epsilon}^T = (\underline{Y} - X \hat{\beta}_{y,x}^{(t)})(\underline{Y} - X \hat{\beta}_{y,x}^{(t)})^T.$$

Consequently,

$$\frac{\partial \underline{\epsilon} \underline{\epsilon}^T}{\partial \hat{\beta}_{y,x}^{(t)}} = -2 (\underline{Y} - X \hat{\beta}_{y,x}^{(t)}) X^T,$$

$$\frac{\partial \underline{\epsilon} \underline{\epsilon}^T}{\partial \hat{\beta}_{y,x}^{(t)}} = 0.$$

and,

$$X^T \underline{Y} = X^T X \hat{\beta}_{y,x}^{(t)}.$$

Then,

$$\hat{\beta}_{y,x}^{(t)} = (X^T X)^{-1} X^T \underline{Y}$$

Secondly, to estimate $\hat{\sigma}_{y,x}^2$ we have,

$$(\hat{\sigma}_{y,x}^{(t)})^2 = \frac{1}{n} \left[\sum_{i=1}^m (Y_i - X_i \hat{\beta}_{y,x}^{(t)})^2 + \sum_{i=1}^n (Y_i - X_i \hat{\beta}_{y,x}^{(t)})^2 \right] \quad (2.1)$$

consequently,

$$(\hat{\sigma}_{y,x}^{t+1})^2 = n^{-1} \left[\sum_{i=1}^m (Y_i - X_i \hat{\beta}_{y,x}^{(t)})^2 + (n-m) (\hat{\sigma}_{y,x}^t)^2 \right]$$

- (i). X is $(n \times (p + 1))$ matrix of observation of the exogenous variable.
- (ii). $\underline{\beta}_{y,x}$ is $((p + 1) \times 1)$ vector of unknown parameters.
- (iii). $\underline{\epsilon}$ is $(n \times 1)$ vector of error, and $\underline{\epsilon} \sim N(0, \sigma^2 I_n)$, where σ^2 is unknown.

therefore,

$$E(\underline{Y}) = E[X\underline{\beta}_{y,x}] + E(\underline{\epsilon}).$$

and, it follows that

$$\underline{\hat{Y}} = X\underline{\hat{\beta}}_{y,x}.$$

If the EM is applied to this problem, the M-step corresponds to the least squares analysis on the original design and the E-step involves finding the expected values and expected squared of the missing Y_i 's given the current estimated parameters.

$$\theta_{y,x}^{(t)} = (\underline{\beta}_{y,x}^{(t)}, (\hat{\sigma}_{y,x}^t)^2),$$

so, steps of EM algorithm are

1. E – step

$$E(Y_i | X, Y_{obs}, \theta_{y,x}^{(t)}) = \begin{cases} Y_i & , \text{ if } Y_i \text{ is observed } (i=1, \dots, m) \\ X_i \hat{\underline{\beta}}_{y,x}^{(t)} & , \text{ if } Y_i \text{ is missing } (i=m+1, \dots, n) \\ & \text{and } X_i \text{ is } (1 \times (p+1)) \text{ vector} \end{cases}$$

and

$$E(Y_i^2 | X, Y_{obs}, \theta_{y,x}^{(t)}) = \begin{cases} Y_i^2 & , \text{ if } Y_i \text{ is observed } i=1, \dots, m. \\ \left(X_i \hat{\underline{\beta}}_{y,x}^{(t)} \right)^2 + \sigma_{y,x}^{2(t)} & , \text{ if } Y_i \text{ is missing } i=m+1, \dots, n. \end{cases}$$

where,

$$Y_i = (X\underline{\beta}_{y,x}^{(t)} + \epsilon_i),$$

also,

$$\text{Var}(Y) = \text{Var}(X\underline{\beta}_{y,x}^{(t)} + \underline{\epsilon}).$$

$$= E(Y^2 | X, Y_{obs}, \theta_{y,x}^{(t)}) - \left[E(Y | X, Y_{obs}, \theta_{y,x}^{(t)}) \right]^2 = (\sigma_{y,x}^t)^2$$

(iii). Non-randomly missing observations-this is a case in which the researcher has evidence to believe that neither (a) nor (b) is true and that some special reason exists for non-response to a particular question.

Missing data is a part of almost studies. There are a number of alternative ways of dealing with missing data. These ways can be divided into two main directions:

- (i). Traditional methods for handling missing data such as the complete case analysis (cc), Pairwise Deletion and Single Imputation (imputing unconditional means, imputing from unconditional distributions and imputing conditional mean).
- (ii). Modern methods for handling missing data such as EM and Bayesian approaches.

The rest of the article is organized as follows. In Section 2, we provide the definition of linear regression with missing values confined to the endogenous variable and different properties of it. Bayesian regression for missing data is discussed in Section 3; also we introduce Bayesian analysis of time series with missing data as special cases of Bayesian regression in Section 4.

2. Linear Regression with Missing Values Confined to the Endogenous Variable

When a scalar outcome variable Y is regressed on P predictor variable $X_1, X_2, \dots, X_P$ and missing values are confined to Y , the incomplete observation do not contain information about the regression parameters $\theta_{Y,X} = (\beta_{Y,X}, \sigma_{Y,X}^2)$, if $\theta_{Y,X}$ is distinct, and the X 's are regarded as fixed constants, nevertheless, the EM algorithm can be applies to all observations and will obtain iteratively the same ML estimates as would have been obtain one-iteratively using only the complete observation. In some cases it may be easier to find these ML estimates iteratively by EM than non-iteratively. In designed experiments the set of values of $X_1, \dots, X_P$ is chosen to make least squares estimation of parameters computationally simple. When Y given $X_1, \dots, X_n$ is normal, least squares computation yield ML estimates. When values of Y , say $Y_i, i = m + 1..n$ are missing at random.

We assume the regression equation is:

$$\underline{Y} = X \underline{\beta}_{Y,X} + \underline{\epsilon},$$

then,

$$\underline{\hat{Y}} = X \underline{\hat{\beta}}_{Y,X},$$

where,

On Some Results of Bayesian Regression with Missing Data

Abd El-Fattah Kandil^{*}, Mervat Mahdy^{**} and Dina El-Telbany^{***}

Abstract

Missing data are a recurring problem that can cause bias or lead to inefficient analyses. In this study we make comparison between EM algorithm and Bayesian method and conclude that EM method is faster in making convergence to estimate the missing values than the Bayesian methods but the last one achieved better estimations especially if the size of the data is large. Finally we apply the Bayesian method on Time Series data, and to simplify we use AR (1) model which the data have some missing values.

Key words: missing data, pattern and mechanisms of missing data, Bayesian method.

1. Introduction

The literatures which deal with missing data problem have increased recently such that the statistical inference with missing data is considered one of the most important practical applications. Missing data means that no data value is stored for the variable in the current observation.

In fact, we usually face a lot of problems that caused by the missing values either planned or not planned, that it can cause bias or lead to inefficient analyses, that missing data have two major negative effects:

First, they have a negative impact on statistical power.

Second, missing data may result in biased estimates.

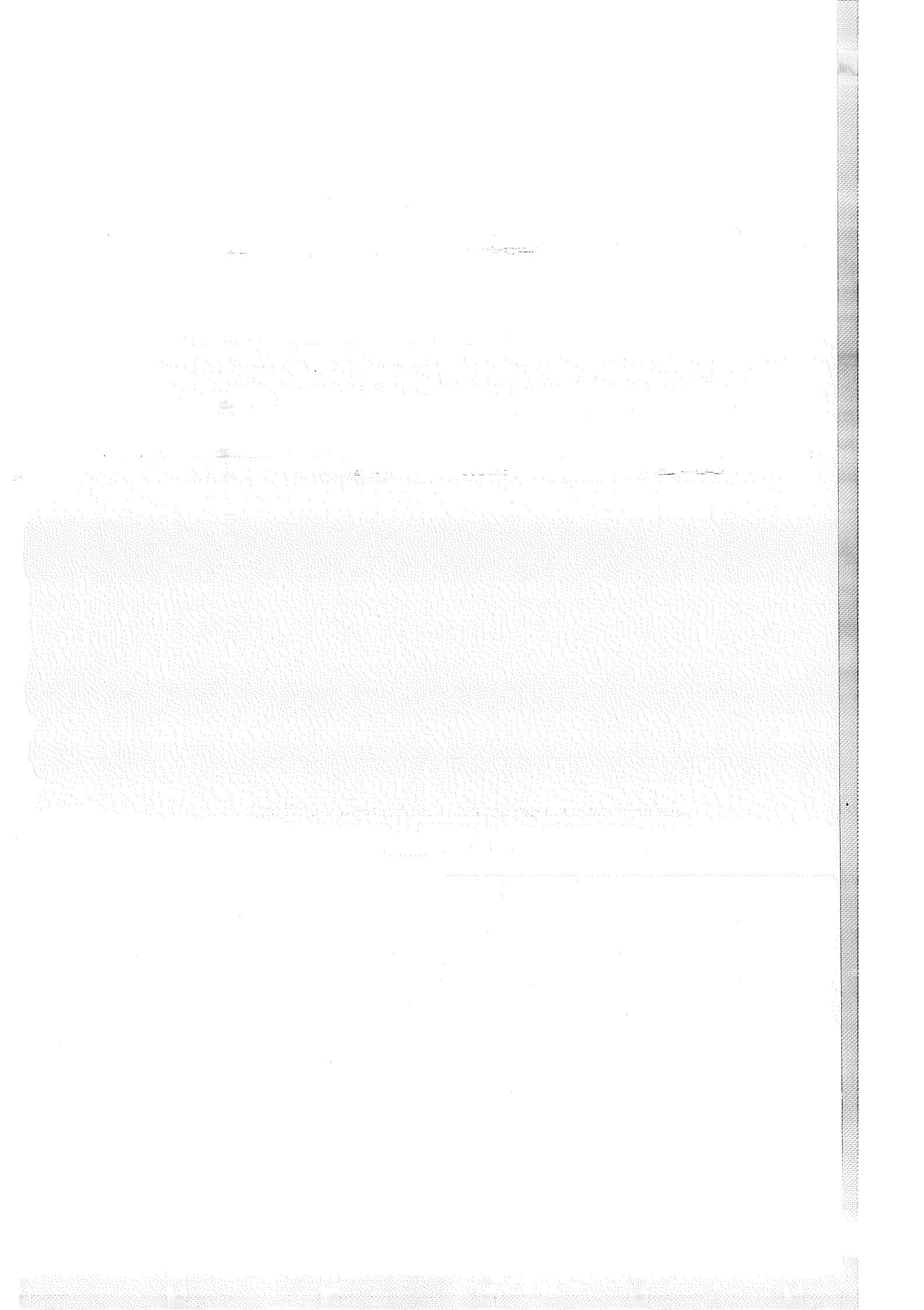
Several ways-both pragmatic and statistically sophisticated-have been devised to cope with the missing data problem depending on their nature and proportion. Three general cases are distinguishable:

- (i). Randomly missing observations;
- (ii). Missing categories, meaning that no answer is available owing to the fact that the question refers to some non-existing category in the responding unit;
- (iii). Non-randomly missing observations-this is a case in which the researcher has evidence to believe that neither (a) nor (b) is true and

* Present address: Department of Statistics, Mathematics and Insurance College of Commerce, Benha University, Egypt.E-mail address: abdel fattah.kandil@Gmail.com

** Present address: Department of Statistics, Mathematics and Insurance College of Commerce, Benha University, Egypt.E-mail address: drmervat.mahdy@fcom.bu.edu.eg

*** Present address: Department of Statistics, Mathematics and Insurance College of Commerce, Benha University, Egypt.E-mail address: d_eltelbany@yahoo.com



- 12) **Nagaraja, H. N. and Abo-Eleneen, Z. A. (2008).** Fisher Information in Order Statistics and their Concomitants in Bivariate Censored Samples. *Metrika*, 67, 327-347.
- 13) **Sarhan, A. M. and Kundu, D. (2009).** Generalized Linear Failure Rate Distribution. *Communications in Statistics – Theory and Methods*, 38(5), 642-660.
- 14) **Sarhan, A. M., Hamilton, D. C., Smith, B. and Kundu, D. (2011).** The Bivariate Generalized Linear Failure Rate Distribution and Its Multivariate Extension *Computational Statistics and Data Analysis*, 55, 644 – 654.
- 15) **Qinying, He. (2007).** Inference on Correlation from Incomplete Bivariate Models. Ph. D. Thesis, The Ohio State University, United States.

We have considered the maximum likelihood estimators of the unknown parameters of the BGLFR distribution in case of type II censored samples based on concomitants of order statistics and hybrid random censored samples in case of both exponential and Weibull censored times. The associated asymptotic variance-covariance matrix for the unknown parameters was derived in each case. To illustrate the results, numerical examples were given where the MLEs, their approximate variance-covariance matrix and asymptotic confidence intervals were determined.

References

- 1) **Aly, H. M. and Muhammed, H. Z. (2012).** Estimation of the Bivariate Generalized Linear Failure Rate Distribution Parameters under Random Censoring. *International Journal of Contemporary Mathematical Sciences*. 7 (17), 821 - 837
- 2) **Bhattacharya, P.K. (1974).** Convergence of Sample Paths of Normalized Sums of Induced Order Statistics. *The Annals of Statistics*, 2, 1034-1039.
- 3) **Cohen, A. C. (1965).** Maximum Likelihood Estimation in The Weibull Distribution Based On Complete and Censored Samples. *Technometrics*, 7, 579-588.
- 4) **David, H. A. (1973).** Concomitants of Order Statistics. *Bulletin of International Statistical Institute*, 45, 295-300.
- 5) **Epstein, B. (1954).** Truncated Lifetests in The Exponential Case. *Annals of Mathematical Statistics*. 25, 555-604.
- 6) **Hanagal, D.D. (1997).** Inference Procedures In Some Bivariate Exponential Models under Hybrid Random Censoring. *Statistical Papers*. 38, 169-189.
- 7) **Hougaard, P., Harvald, B. and Holm, N. V. (1992).** Measuring the similarities between the lifetimes of adult Danish twins born between 1881-1930. *Journal of the American Statistical Association*. 87, 17-24.
- 8) **Kundu, D. and Gupta, R. D., (2009).** Bivariate Generalized Exponential Distribution. *Journal of Multivariate Analysis*. 100(4), 581 – 593.
- 9) **Lin, D. Y., Sun, W. and Ying, Z. (1999).** Nonparametric Estimation of The Gap Time Distribution for Serial Events with Censored Data. *Biometrika*. 86, 59- 70.
- 10) **Marshall, A. W. and Olkin, I. (1967).** A Multivariate Exponential Distribution. *Journal of the American Statistical Association*, 62, 30 -44.
- 11) **Nagaraja, H. N. and David, H. A. (1994).** Distribution of the Maximum of Concomitants of Selected Order Statistics. *The Annals of Statistics*. 22, 478-494.

Table (2): The estimates, relative bias, MSE, the confidence interval and variance-covariance matrix under type II censored samples

Parameter	MLE	Confidence Intervals	The Variance-Covariance matrix		
α_1	1.027	(1.291, 1.309)	7.55E-4	9.29E-6	-5.4E-5
α_2	0.386	(0.275, 0.325)	9.29E-6	6.09E-3	-1.3E-3
α_3	0.463	(0.372, 0.428)	-5.4E-5	-1.3E-3	7.6E-3

Table (3): The estimates, Relative Bias, MSE, Confidence interval and variance- covariance matrix under hybrid random censoring in case of exponential censoring time

Parameter	MLE	Confidence Intervals	The Variance-Covariance matrix			
α_1	1.327	(1.249, 1.404)	0.156	3.36E-4	0.013	0
α_2	0.29	(0.23, 0.349)	3.36E-4	0.054	-0.038	0
α_3	0.84	(0.753, 0.935)	-0.013	-0.038	0.092	0
θ	2.5	(2.489, 2.8)	0	0	0	0.233

Table (4): The estimates, Relative Bias, MSE, Confidence interval and variance- covariance matrix under hybrid random censoring in case of Weibull censoring time

Parameter	MLE	Confidence Intervals	The Variance-Covariance matrix				
α_1	1.027	(0.948, 1.106)	0.06	1.13E-3	-5.8E-3	0	0
α_2	0.419	(0.358, 0.48)	1.13E-3	0.036	-0.017	0	0
α_3	0.572	(0.512, 0.631)	-5.8E-3	-0.017	0.034	0	0
λ	1.393	(1.331, 1.455)	0	0	0	0.037	-4.2E-3
η	1.351	(1.281, 1.422)	0	0	0	-4.2E-3	0.048

6. Conclusion

In this paper, we have provided a review of the new bivariate generalized linear failure rate distribution whose marginals are distributed as the generalized linear failure rate distributions. This new distribution is a singular distribution which can be used quit effectively instead of the Marshall-Olkin bivariate exponential model or the bivariate generalized exponential model when there are ties in the data.

Upon examining Table (3) and Table (4), we find also that the MIEs of the estimators $\hat{\alpha}_1$, $\hat{\alpha}_2$ and $\hat{\alpha}_3$ fall in their correspondence confidence intervals and the lengths of the confidence intervals are relatively short. In terms of the variance-covariance matrix we find that Weibull censoring time is better than the exponential censoring time.

Table (1): Samples observations from BGLFR distribution, the exponential distribution with $\theta = 3.25$ and Weibull distribution with $\lambda = 1.5$ and $\eta = 1.5$

No	x_1	x_2	t_i		t_i^v	
			Exponential distribution	Weibull distribution	Exponential distribution	Weibull distribution
1	0.26	0.20	2.052	1.393	0.24	0.24
2	0.63	0.18	0.506	0.660	0.246	0.24
3	0.19	0.19	0.165	1.032	0.165	0.24
4	0.66	0.85	0.323	0.336	0.24	0.24
5	0.4	0.4	0.06	1.451	0.06	0.24
6	0.49	0.49	0.538	0.489	0.24	0.24
7	0.08	0.08	0.105	1.123	0.105	0.24
8	0.69	0.71	0.366	1.789	0.24	0.24
9	0.39	0.39	0.736	1.542	0.24	0.24
10	0.82	0.48	0.48	0.051	0.24	0.051
11	0.72	0.72	0.00356	1.654	0.00356	0.24
12	0.66	0.62	0.655	2.813	0.24	0.24
13	0.25	0.09	1.452	0.736	0.24	0.24
14	0.41	0.03	0.194	0.637	0.194	0.24
15	0.16	0.75	0.156	1.477	0.156	0.24
16	0.18	0.18	0.552	0.859	0.24	0.24
17	0.22	0.14	0.245	2.017	0.24	0.24
18	0.42	0.42	0.881	0.391	0.24	0.24
19	0.36	0.52	0.075	0.754	0.075	0.24
20	0.34	0.34	0.201	0.260	0.201	0.24
21	0.53	0.39	0.041	0.127	0.041	0.24
22	0.54	0.07	0.014	0.725	0.014	0.24
23	0.51	0.28	0.19	0.842	0.19	0.24
24	0.76	0.64	0.238	0.280	0.238	0.24
25	0.64	0.15	0.046	0.396	0.046	0.24
26	0.26	0.48	0.077	0.022	0.077	0.022
27	0.16	0.16	0.00099	0.623	0.00099	0.24
28	0.44	0.6	0.151	1.205	0.151	0.24
29	0.25	0.14	0.407	0.312	0.24	0.24
30	0.55	0.11	0.054	0.986	0.054	0.24
31	0.49	0.49	0.301	0.534	0.24	0.24
32	0.24	0.24	0.12	2.818	0.12	0.24
33	0.44	0.30	1.456	1.184	0.24	0.24
34	0.42	0.03	0.396	0.656	0.24	0.24
35	0.27	0.47	0.163	0.315	0.163	0.24
36	0.28	0.28	0.055	0.806	0.055	0.24
37	0.02	0.02	0.223	0.444	0.223	0.24

$$a_{45} = -\frac{\partial^2 \ln L}{\partial \lambda \partial \eta} \bigg|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3, \hat{\lambda}, \hat{\eta}} = -\frac{r_2}{\eta} - \frac{\hat{\lambda}}{\hat{\eta}} \left\{ \sum_{i=1}^n \left(\frac{\max(x_{1i}, x_{2i})}{\hat{\eta}} \right)^{\hat{\lambda}} \ln \left(\frac{\max(x_{1i}, x_{2i})}{\hat{\eta}} \right) + \sum_{i=1}^{r_2} \left(\frac{t_i}{\hat{\eta}} \right)^{\hat{\lambda}} [\ln(t_i)] \right\} \\ + \frac{1}{\hat{\eta}} \left\{ \sum_{i=1}^n \left(\frac{\max(x_{1i}, x_{2i})}{\hat{\eta}} \right)^{\hat{\lambda}} + \sum_{i=1}^{r_2} \left(\frac{t_i}{\hat{\eta}} \right)^{\hat{\lambda}} \right\},$$

$$a_{55} = -\frac{\partial^2 \ln L}{\partial \eta^2} \bigg|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3, \hat{\lambda}, \hat{\eta}} = \frac{\hat{\lambda}}{\hat{\eta}^2} \left\{ \sum_{i=1}^n \left(\frac{\max(x_{1i}, x_{2i})}{\hat{\eta}} \right)^{\hat{\lambda}} + \sum_{i=1}^{r_2} \left(\frac{t_i}{\hat{\eta}} \right)^{\hat{\lambda}} \right\} \\ + \left(\frac{\hat{\lambda}}{\hat{\eta}} \right)^2 \left\{ \sum_{i=1}^n \left(\frac{\max(x_{1i}, x_{2i})}{\hat{\eta}} \right)^{\hat{\lambda}} + \sum_{i=1}^{r_2} \left(\frac{t_i}{\hat{\eta}} \right)^{\hat{\lambda}} \right\} - \frac{r_2}{\hat{\eta}^2},$$

$$a_{51} = -\frac{\partial^2 \ln L}{\partial \eta \partial \alpha_1} \bigg|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3, \hat{\lambda}, \hat{\eta}} = 0,$$

$$a_{52} = -\frac{\partial^2 \ln L}{\partial \eta \partial \alpha_2} \bigg|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3, \hat{\lambda}, \hat{\eta}} = 0 \text{ and}$$

$$a_{53} = -\frac{\partial^2 \ln L}{\partial \eta \partial \alpha_3} \bigg|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3, \hat{\lambda}, \hat{\eta}} = 0.$$

5. Numerical Results

To illustrate the results, a numerical example in each case is given. Thirty seven observations from BGLFR distribution (X_1, X_2) have been obtained from Sarhan et al. (2011) and are listed in Table (1). Two samples of thirty seven observations are selected at random from the exponential distribution with parameter $\theta = 3.25$ where $t_i \sim \exp(\theta)$ and Weibull distribution with $\lambda = 1.5$ and $\eta = 1.5$ respectively, and are listed in Table (1). The maximum likelihood estimators (MLEs) are obtained by Newton-Raphson technique using Mathcad. Their approximate variance-covariance matrix and a 95% asymptotic confidence intervals for both type II and hybrid random censoring are listed in Table (2), Table (3) and Table (4) respectively.

Table (2) provides the MLEs, variance-covariance matrix and 95% confidence intervals of the estimators $\hat{\alpha}_1$, $\hat{\alpha}_2$ and $\hat{\alpha}_3$ for $r = 21$ and $n = 37$. From Table (2), it is found that the MLEs fall in their correspondence confidence intervals and the lengths of the confidence intervals are relatively short. It is evident the variance of α_1 is the smallest one in the variance-covariance matrix.

The log of the likelihood function (4.1) is given as

$$\ln L(\alpha_1, \alpha_2, \alpha_3, \lambda, \eta) \propto \delta(\alpha_1, \alpha_2, \alpha_3) + \phi(\alpha_1, \alpha_2, \alpha_3) + \psi(\alpha_1, \alpha_2, \alpha_3) + \tau(\lambda, \eta), \quad (4.9)$$

where

$$\tau(\lambda, \eta) = r_2 \ln \lambda - r_2 \lambda \ln \eta + (\lambda - 1) \sum_{i=1}^{r_2} \ln(t_i) - \sum_{i=1}^n \left(\frac{\max(x_{1i}, x_{2i})}{\eta} \right)^\lambda - \sum_{i=1}^{r_2} \left(\frac{t_i}{\eta} \right)^\lambda.$$

On differentiating (4.9) with respect to $\alpha_1, \alpha_2, \alpha_3, \lambda$ and η in turn and equating to zero we obtain the same three equations (4.3) to (4.5) and the following two likelihood equations

$$\frac{r_2}{\hat{\lambda}} - r_2 \ln \hat{\eta} + \sum_{i=1}^{r_2} \ln t_i - \sum_{i=1}^n \left(\frac{\max(x_{1i}, x_{2i})}{\hat{\eta}} \right)^\lambda \ln \left(\frac{\max(x_{1i}, x_{2i})}{\hat{\eta}} \right) - \sum_{i=1}^{r_2} \left(\frac{t_i}{\hat{\eta}} \right)^\lambda \ln \left(\frac{t_i}{\hat{\eta}} \right) = 0 \quad \text{and}$$

$$\frac{\hat{\lambda}}{\hat{\eta}} \cdot \left[\sum_{i=1}^n \left(\frac{\max(x_{1i}, x_{2i})}{\hat{\eta}} \right)^\lambda - \sum_{i=1}^{r_2} \left(\frac{t_i}{\hat{\eta}} \right)^\lambda \right] - \frac{r_2}{\hat{\eta}} = 0.$$

The asymptotic variance-covariance matrix in this case is given by

$$I^{-1}(\hat{\underline{\alpha}}) = \begin{bmatrix} a_{11} & a_{12} & a_{13} & a_{14} & a_{15} \\ a_{21} & a_{22} & a_{23} & a_{24} & a_{25} \\ a_{31} & a_{32} & a_{33} & a_{34} & a_{35} \\ a_{41} & a_{42} & a_{43} & a_{44} & a_{45} \\ a_{51} & a_{52} & a_{53} & a_{54} & a_{55} \end{bmatrix}^{-1},$$

where

$a_{11}, a_{12}, a_{13}, a_{14}, a_{22}, a_{23}, a_{24}, a_{33}$ and a_{34} are the same as in (4.7). But, the elements $a_{44}, a_{45}, a_{55}, a_{15}, a_{25}$ and a_{35} have the following forms

$$a_{44} = - \frac{\partial^2 \ln L}{\partial \lambda^2} \bigg|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3, \hat{\lambda}, \hat{\eta}} = \frac{r_2}{\hat{\lambda}^2} + \sum_{i=1}^n \left(\frac{\max(x_{1i}, x_{2i})}{\hat{\eta}} \right)^\lambda \left[\ln \left(\frac{\max(x_{1i}, x_{2i})}{\hat{\eta}} \right) \right]^2 + \sum_{i=1}^{r_2} \left(\frac{t_i}{\hat{\eta}} \right)^\lambda \left[\ln \left(\frac{t_i}{\hat{\eta}} \right) \right]^2.$$

$$\begin{aligned}
 a_{41} &= - \frac{\partial^2 \ln L}{\partial \theta \partial \alpha_1} \bigg|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3, \hat{\theta}} = 0, \\
 a_{42} &= - \frac{\partial^2 \ln L}{\partial \theta \partial \alpha_2} \bigg|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3, \hat{\theta}} = 0, \\
 a_{43} &= - \frac{\partial^2 \ln L}{\partial \theta \partial \alpha_3} \bigg|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3, \hat{\theta}} = 0, \text{ and} \\
 a_{44} &= - \frac{\partial^2 \ln L}{\partial \theta^2} \bigg|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3, \hat{\theta}} = \frac{r_2}{\hat{\theta}^2}.
 \end{aligned} \tag{4.7}$$

The asymptotic normality results to obtain the asymptotic confidence intervals of $\alpha_1, \alpha_2, \alpha_3$ and $\hat{\theta}$ can be stated as follows

$$[\sqrt{n}(\hat{\underline{\alpha}} - \underline{\alpha})] \rightarrow N_4(\underline{0}, I^{-1}(\underline{\alpha})) \text{ as } n \rightarrow \infty \tag{4.8}$$

where $I^{-1}(\underline{\alpha})$ is the variance-covariance matrix,

$$\hat{\underline{\alpha}} = (\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3, \hat{\theta}) \text{ and } \underline{\alpha} = (\alpha_1, \alpha_2, \alpha_3, \theta).$$

Since $\underline{\alpha}$ is unknown in (4.8), $I^{-1}(\underline{\alpha})$ is estimated by $I^{-1}(\hat{\underline{\alpha}})$; the asymptotic variance-covariance matrix that defined above and this can be used to obtain the asymptotic confidence intervals of $\alpha_1, \alpha_2, \alpha_3$ and θ .

4.2 Assuming Weibull Censoring Time

when assuming the censoring time T_i distributed as two parameter Weibull distribution with scale parameter η and shape parameter λ , the pdf and the survivor function of T_i defined in (4.1), will be respectively

$$g(t_i) = \frac{\lambda}{\eta} \left(\frac{t_i}{\eta}\right)^{\lambda-1} e^{-\left(\frac{t_i}{\eta}\right)^\lambda}; \quad t_i \geq 0 \text{ and}$$

$$\bar{G}(t_i) = P[T_i > \max(x_{1i}, x_{2i})] = e^{-\frac{\max\{x_{1i}, x_{2i}\}^\lambda}{\eta}}$$

$$M_2 = \sum_{i=1}^{n_1} u_2(x_{2i}, \beta, \gamma) + \sum_{i=1}^{n_2} u_2(x_{2i}, \beta, \gamma) + \sum_{i=1}^{n_3} u_2(x_{2i}, \beta, \gamma) + \sum_{i=1}^{n_4} u_2(x_{2i}, \beta, \gamma),$$

and

$$M_3 = \sum_{i=1}^{n_1} u_2(x_{1i}, \beta, \gamma) + \sum_{i=1}^{n_2} u_2(x_{2i}, \beta, \gamma) + \sum_{i=1}^{n_3} u_2(x_{1i}, \beta, \gamma) + \sum_{i=1}^{n_4} u_2(x_{1i}, \beta, \gamma) + \sum_{i=1}^{n_5} u_2(x_{2i}, \beta, \gamma).$$

The three equations (4.3) to (4.5) have not explicit form; therefore, their solutions are numerically obtained using Newton-Raphson method as will be seen in Section 5. They are solved simultaneously to obtain $\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3$ and $\hat{\theta}$. For simplification, we assumed that β and γ are fixed known.

The asymptotic variance-covariance matrix is given by

$$I^{-1}(\hat{\alpha}) = \begin{bmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ a_{21} & a_{22} & a_{23} & a_{24} \\ a_{31} & a_{32} & a_{33} & a_{34} \\ a_{41} & a_{42} & a_{43} & a_{44} \end{bmatrix}^{-1},$$

where

$$\begin{aligned} a_{11} &= -\frac{\partial^2 \ln L}{\partial \alpha_1^2} \Big|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3} = \left\{ \frac{n_1 + n_4}{(\hat{\alpha}_1 + \hat{\alpha}_3)^2} + \frac{n_2}{\hat{\alpha}_1^2} + v_2(\hat{\alpha}_1, 0, 0, u_1, u_2, t_i^*, n_5) \right. \\ &\quad \left. + v_2(\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3, u_1, u_2, t_i^*, n_6) \right\}, \\ a_{12} &= -\frac{\partial^2 \ln L}{\partial \alpha_1 \partial \alpha_2} \Big|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3} = v_2(\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3, u_1, u_2, t_i^*, n_6), \\ a_{13} &= -\frac{\partial^2 \ln L}{\partial \alpha_1 \partial \alpha_3} \Big|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3} = \frac{n_1 + n_4}{(\hat{\alpha}_1 + \hat{\alpha}_3)^2} + v_2(\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3, u_1, u_2, t_i^*, n_6), \\ a_{22} &= -\frac{\partial^2 \ln L}{\partial \alpha_2^2} \Big|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3} = \frac{n_2 + n_5}{(\hat{\alpha}_2 + \hat{\alpha}_3)^2} + \frac{n_1}{\hat{\alpha}_2^2} + v_2(0, \hat{\alpha}_2, 0, u_1, u_2, t_i^*, n_5) \\ &\quad + v_2(\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3, u_1, u_2, t_i^*, n_6), \\ a_{23} &= -\frac{\partial^2 \ln L}{\partial \alpha_2 \partial \alpha_3} \Big|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3} = \frac{n_2 + n_5}{(\hat{\alpha}_1 + \hat{\alpha}_3)^2} + v_2(\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3, u_1, u_2, t_i^*, n_6), \\ a_{33} &= -\frac{\partial^2 \ln L}{\partial \alpha_3^2} \Big|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3} = \frac{n_1 + n_4}{(\hat{\alpha}_1 + \hat{\alpha}_3)^2} + \frac{n_2}{\hat{\alpha}_3^2} + \frac{n_2 + n_5}{(\hat{\alpha}_1 + \hat{\alpha}_3)^2} \\ &\quad + v_2(\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3, u_1, u_2, t_i^*, n_6). \end{aligned}$$

$$\begin{aligned}
 & + (\alpha_1 + \alpha_2 + \alpha_3 - 1) \sum_{i=1}^{n_2} \ln(1 - e^{-(\beta x_{2i} + \frac{\gamma}{2} x_{2i}^2)}) + (\alpha_1 + \alpha_3 - 1) \sum_{i=1}^{n_4} \ln(1 - e^{-(\beta x_{4i} + \frac{\gamma}{2} x_{4i}^2)}) \\
 & + (\alpha_2 + \alpha_3 - 1) \sum_{i=1}^{n_5} \ln(1 - e^{-(\beta x_{5i} + \frac{\gamma}{2} x_{5i}^2)}), \\
 \psi(\alpha_1, \alpha_2, \alpha_3) = & \sum_{i=1}^{n_4} \ln[1 - (1 - e^{-(\beta t_i^* + \frac{\gamma}{2} t_i^{*2})})^{\alpha_2}] + \sum_{i=1}^{n_5} \ln[1 - (1 - e^{-(\beta t_i^* + \frac{\gamma}{2} t_i^{*2})})^{\alpha_1}] \\
 & + \sum_{i=1}^{n_6} \ln[1 - (1 - e^{-(\beta t_i^* + \frac{\gamma}{2} t_i^{*2})})^{\alpha_1 + \alpha_2 + \alpha_3}]
 \end{aligned}$$

and

$$\zeta(\theta) = r_2 \ln \theta - \theta \left[\sum_{i=1}^{r_1} t_i + \sum_{i=1}^{r_2} \max(x_{1i}, x_{2i}) \right].$$

On differentiating (4.2) with respect to $\alpha_1, \alpha_2, \alpha_3$ and θ in turn and equating to zero, we obtain the following likelihood equations

$$\frac{n_1 + n_4}{\hat{\alpha}_1 + \hat{\alpha}_3} + \frac{n_2}{\hat{\alpha}_1} + M_1 - v_1(\hat{\alpha}_1, 0, 0, u_1, u_2, t_i^*, n_5) - v_1(\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3, u_1, u_2, t_i^*, n_6) = 0, \quad (4.3)$$

$$\frac{n_2 + n_5}{\hat{\alpha}_2 + \hat{\alpha}_3} + \frac{n_1}{\hat{\alpha}_2} + M_2 - v_1(0, \hat{\alpha}_2, 0, u_1, u_2, t_i^*, n_4) - v_1(\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3, u_1, u_2, t_i^*, n_6) = 0, \quad (4.4)$$

$$\frac{n_1 + n_4}{\hat{\alpha}_1 + \hat{\alpha}_3} + \frac{n_2 + n_5}{\hat{\alpha}_2 + \hat{\alpha}_3} + \frac{n_3}{\hat{\alpha}_3} + M_3 - v_1(\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3, u_1, u_2, t_i^*, n_6) = 0, \quad (4.5)$$

and

$$\frac{r_2}{\hat{\theta}} - \left[\sum_{i=1}^{r_2} t_i + \sum_{i=1}^{r_1} \max(x_{1i}, x_{2i}) \right] = 0. \quad (4.6)$$

where

$$v_1(\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3, u_1, u_2, t_i, n) = \sum_{i=1}^n \frac{[u_1(t_i, \beta, \gamma)]^{\hat{\alpha}_1 + \hat{\alpha}_2 + \hat{\alpha}_3} u_2(t_i, \beta, \gamma)}{1 - [u_1(t_i, \beta, \gamma)]^{\hat{\alpha}_1 + \hat{\alpha}_2 + \hat{\alpha}_3}},$$

$$u_1(t_i, \beta, \gamma) = 1 - e^{-(\beta t_i + \frac{\gamma}{2} t_i^2)}, \quad u_2(t_i, \beta, \gamma) = \ln u_1(t_i, \beta, \gamma),$$

$$M_1 = \sum_{i=1}^{n_1} u_2(x_{1i}, \beta, \gamma) + \sum_{i=1}^{n_2} u_2(x_{1i}, \beta, \gamma) + \sum_{i=1}^{n_3} u_2(x_{1i}, \beta, \gamma) + \sum_{i=1}^{n_4} u_2(x_{1i}, \beta, \gamma),$$

$$f_3(x_i) = \alpha_3(\beta + \gamma x_i) e^{-(\beta x_i + \frac{\gamma}{2} x_i^2)} [1 - e^{-(\beta x_i + \frac{\gamma}{2} x_i^2)}]^{a_{123}-1}; \quad 0 < (x_{1i} = x_{2i}) < t_i^*,$$

$$f_4(x_{1i}, t_i^*) = \lim_{\delta x_i \rightarrow 0} \frac{P(x_{1i} < X_{1i} < x_{1i} + \delta x_{1i} / X_{2i} > t_i^*) P(X_{2i} > t_i^*)}{\delta x_i}$$

$$= (\alpha_1 + \alpha_3)(\beta + \gamma x_{1i}) e^{-(\beta x_{1i} + \frac{\gamma}{2} x_{1i}^2)} [1 - e^{-(\beta x_{1i} + \frac{\gamma}{2} x_{1i}^2)}]^{a_1 + \alpha_3 - 1}$$

$$[1 - (1 - e^{-(\beta t_{1i} + \frac{\gamma}{2} t_{1i}^2)})^{a_2}]^{a_3}; \quad 0 < x_{1i} < t_i^*,$$

$$f_3(t_i^*, x_{2i}) = \lim_{\delta x_i \rightarrow 0} \frac{P(x_{2i} < X_{2i} < x_{2i} + \delta x_{2i} / X_{1i} > t_i^*) P(X_{1i} > t_i^*)}{\delta x_i}$$

$$= (\alpha_2 + \alpha_3)(\beta + \gamma x_{2i}) e^{-(\beta x_{2i} + \frac{\gamma}{2} x_{2i}^2)} [1 - e^{-(\beta x_{2i} + \frac{\gamma}{2} x_{2i}^2)}]^{a_1 + \alpha_3 - 1}$$

$$[1 - (1 - e^{-(\beta t_{2i} + \frac{\gamma}{2} t_{2i}^2)})^{a_1}]^{a_3}; \quad 0 < x_{2i} < t_i^*,$$

$$\bar{F}(t_i^*, t_i^*) = P(X_{1i} > t_i^*, X_{2i} > t_i^*) = 1 - (1 - e^{-(\beta t_{1i} + \frac{\gamma}{2} t_{1i}^2)})^{a_1 + \alpha_2 + \alpha_3 - 1},$$

$$g(t_i) = \theta \cdot e^{-\theta t_i}; \quad 0 < t_i < \infty,$$

$$\text{and } \bar{G}(t_i) = P[T_i > \max(x_{1i}, x_{2i})] = \exp[-\theta \cdot \max(x_{1i}, x_{2i})],$$

The counts $n_i, i = 1, \dots, 6$ are the number of observations falling in the range corresponds to $f_i, i = 1, \dots, 5$ and $\bar{F}$ such that $n_1 + \dots + n_5 = D^*$ and $n_6 = n - D^*$.

r_1 is the number of observations with $\max(x_{1i}, x_{2i}) < t_i$, $r_2 = n - r_1$,

f_1 and f_2 are with respect to lebesgue measure in R_2 and f_3, f_4, f_5 are with respect to lebesgue measure in R_1 .

After substituting from (2.2) in (4.1) and taking the logarithm we get

$$\ln L(\alpha_1, \alpha_2, \alpha_3, \theta) \propto \delta(\alpha_1, \alpha_2, \alpha_3) + \phi(\alpha_1, \alpha_2, \alpha_3) + \psi(\alpha_1, \alpha_2, \alpha_3) + \zeta(\theta), \quad (4.2)$$

where

$$\delta(\alpha_1, \alpha_2, \alpha_3) = n_1 \ln(\alpha_1 + \alpha_3) + n_1 \ln \alpha_2 + n_2 \ln(\alpha_2 + \alpha_3) + n_2 \ln \alpha_1 + n_3 \ln \alpha_3 + n_4 \ln(\alpha_1 + \alpha_3) + n_5 \ln(\alpha_2 + \alpha_3),$$

$$\begin{aligned} \phi(\alpha_1, \alpha_2, \alpha_3) = & (\alpha_1 + \alpha_3 - 1) \sum_{i=1}^{n_1} \ln(1 - e^{-(\beta x_{1i} + \frac{\gamma}{2} x_{1i}^2)}) + (\alpha_2 - 1) \sum_{i=1}^{n_1} \ln(1 - e^{-(\beta x_{2i} + \frac{\gamma}{2} x_{2i}^2)}) \\ & + (\alpha_1 - 1) \sum_{i=1}^{n_2} \ln(1 - e^{-(\beta x_{1i} + \frac{\gamma}{2} x_{1i}^2)}) + (\alpha_2 + \alpha_3 - 1) \sum_{i=1}^{n_2} \ln(1 - e^{-(\beta x_{2i} + \frac{\gamma}{2} x_{2i}^2)}) \end{aligned}$$

rate θ and T is independent of the lifetimes (X_1, X_2) and $Y_{(r)} = \min(X_1, X_2)$. Let D and D^* be the number of observations that would fail before T^* and T , respectively

$$\text{and } D^* = \begin{cases} r & \text{if } Y_{(r)} \leq T_i, i = 1, \dots, n \\ D < r & \text{if } Y_{(r)} > T_i, i = 1, \dots, n, \end{cases}$$

such that $D \leq D^* \leq r$ as $T_i^* \leq T_i, i = 1, \dots, n$.

The lifetimes associated with i^{th} pair of the two components are given by

$$(x_{1i}, x_{2i}) = \begin{cases} (x_{1i}, x_{2i}) & \text{if } \max(x_{1i}, x_{2i}) < t_i^* \\ (x_{1i}, t_i^*) & \text{if } x_{1i} < t_i^* < x_{2i} \\ (t_i^*, x_{2i}) & \text{if } x_{2i} < t_i^* < x_{1i} \\ (t_i^*, t_i^*) & \text{if } \min(x_{1i}, x_{2i}) > t_i^* \end{cases}$$

where $t_i^* = \min(Y_{(r)}, t_i)$.

Note that $f(x_{1i}, x_{2i})$ and $g(t_i)$ are the joint pdf of (x_{1i}, x_{2i}) and the pdf of T respectively and $\bar{F}(x_{1i}, x_{2i})$ and $\bar{G}(t_i)$ are the survivor functions of (x_{1i}, x_{2i}) and T respectively.

Now, according to Hanagal (1997) the likelihood function of the sample of size n is given by

$$L(\alpha_1, \alpha_2, \alpha_3, \theta) = \prod_{i=1}^{n_1} f_1(x_{1i}, x_{2i}) \prod_{i=1}^{n_2} f_2(x_{1i}, x_{2i}) \prod_{i=1}^{n_3} f_3(x_i) \prod_{i=1}^{n_4} f_4(x_{1i}, t_i^*) \prod_{i=1}^{n_5} f_5(t_i^*, x_{2i}) \\ \cdot \prod_{i=1}^{n_6} \bar{F}(t_i^*, t_i^*) \prod_{i=1}^{r_1} \bar{G}(t_i) \prod_{i=1}^{r_2} g(t_i). \quad (4.1)$$

Upon using (2.1), we get

$$f_1(x_{1i}, x_{2i}) = \alpha_{13} \alpha_2 (\beta + \gamma x_{1i}) (\beta + \gamma x_{2i}) e^{-(\beta \gamma_{1i} + \frac{\gamma}{2} x_{1i}^2)} e^{-(\beta \gamma_{2i} + \frac{\gamma}{2} x_{2i}^2)} \\ [1 - e^{-(\beta \gamma_{1i} + \frac{\gamma}{2} x_{1i}^2)}]^{\alpha_{13}-1} [1 - e^{-(\beta \gamma_{2i} + \frac{\gamma}{2} x_{2i}^2)}]^{\alpha_2-1}, \\ ; 0 < x_{1i} < x_{2i} < t_i^*,$$

$$f_2(x_{1i}, x_{2i}) = \alpha_1 \alpha_{23} (\beta + \gamma x_{1i}) (\beta + \gamma x_{2i}) e^{-(\beta \gamma_{1i} + \frac{\gamma}{2} x_{1i}^2)} e^{-(\beta \gamma_{2i} + \frac{\gamma}{2} x_{2i}^2)} \\ [1 - e^{-(\beta \gamma_{1i} + \frac{\gamma}{2} x_{1i}^2)}]^{\alpha_1-1} [1 - e^{-(\beta \gamma_{2i} + \frac{\gamma}{2} x_{2i}^2)}]^{\alpha_{23}-1}, \\ ; 0 < x_{2i} < x_{1i} < t_i^*.$$

$$a_{33} = -\frac{\partial^2 \ln L}{\partial \alpha_3^2} \bigg|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3} = \frac{r}{\hat{\alpha}_3^2} + \frac{r}{(\hat{\alpha}_1 + \hat{\alpha}_3)^2} + \frac{r}{(\hat{\alpha}_2 + \hat{\alpha}_3)^2},$$

$$a_{12} = -\frac{\partial^2 \ln L}{\partial \alpha_1 \partial \alpha_2} \bigg|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3} = 0,$$

$$a_{13} = -\frac{\partial^2 \ln L}{\partial \alpha_1 \partial \alpha_3} \bigg|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3} = \frac{r}{(\hat{\alpha}_1 + \hat{\alpha}_3)^2} \text{ and}$$

$$a_{23} = -\frac{\partial^2 \ln L}{\partial \alpha_2 \partial \alpha_3} \bigg|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3} = \frac{r}{(\hat{\alpha}_2 + \hat{\alpha}_3)^2}.$$

Now, we state the asymptotic normality results to obtain the asymptotic confidence intervals of α_1 , α_2 and α_3 . It can be stated as follows:

$$[\sqrt{n}(\hat{\underline{\alpha}} - \underline{\alpha})] \rightarrow N_3(0, I^{-1}(\underline{\alpha})) \text{ as } n \rightarrow \infty, \quad (3.5)$$

where $I^{-1}(\underline{\alpha})$ is the variance-covariance matrix, $\hat{\underline{\alpha}} = (\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3)$ and $\underline{\alpha} = (\alpha_1, \alpha_2, \alpha_3)$.

Since $\underline{\alpha}$ in (3.5) is unknown, $I^{-1}(\underline{\alpha})$ is estimated by $I^{-1}(\hat{\underline{\alpha}})$; the asymptotic variance-covariance matrix that defined above and this can be used to obtain the asymptotic confidence intervals of α_1 , α_2 and α_3 .

4. Maximum Likelihood Estimation under Hybrid Random Censoring

In this Section we introduce the maximum likelihood estimators for the unknown parameters of BGLFR distribution and their asymptotic variance-covariance matrix under hybrid random censoring. We assume that the censoring times follow either the exponential or Weibull distribution.

4.1 Assuming Exponential Censoring Time

We consider a sampling plan under which the experiment is terminated at the random time $T_i^* = \min(Y_{(r)}, T_i)$. This censoring scheme is called hybrid random censoring. We assume that censoring time T_i is exponential distribution with failure

$$\frac{r}{\hat{\alpha}_1 + \hat{\alpha}_3} + \frac{r}{\hat{\alpha}_1} - 3(n-r) \frac{[u_{1r}(\beta, \gamma)]^{\hat{\alpha}_1} \ln u_{1r}(\beta, \gamma)}{1 - [u_{1r}(\beta, \gamma)]^{\hat{\alpha}_1}} - 3 \sum_{i=1}^r (n-i) \frac{[u_{1i}(\beta, \gamma)]^{\hat{\alpha}_1} \ln u_{1i}(\beta, \gamma)}{1 - [u_{1i}(\beta, \gamma)]^{\hat{\alpha}_1}} + 3 \sum_{i=1}^r i \ln u_{1i}(\beta, \gamma) = 0,$$

$$\frac{r}{\hat{\alpha}_2 + \hat{\alpha}_3} + \frac{r}{\hat{\alpha}_2} + \sum_{i=1}^r \ln u_{1i}(\beta, \gamma) + 2 \sum_{i=1}^r \ln u_{2i}(\beta, \gamma) = 0 \quad \text{and}$$

$$\frac{r}{\hat{\alpha}_1 + \hat{\alpha}_3} + \frac{r}{\hat{\alpha}_3} + \frac{r}{\hat{\alpha}_2 + \hat{\alpha}_3} + 2 \sum_{i=1}^r \ln u_{1i}(\beta, \gamma) + \sum_{i=1}^r \ln u_{2i}(\beta, \gamma) = 0.$$

These three equations have not explicit form; therefore, their solutions are numerically obtained using Newton-Raphson method as will be seen in Section 5. They are solved simultaneously to obtain $\hat{\alpha}_1, \hat{\alpha}_2$ and $\hat{\alpha}_3$. For simplification, we assumed that β and γ are fixed known.

The asymptotic variance-covariance matrix of $(\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3)$ is obtained by inverting the Fisher information matrix with elements that are negatives of the expected values of the second order derivatives of logarithms of the likelihood function. In the present situation, it seems appropriate to approximate the expected values by their maximum likelihood estimates [see Cohen (1965)]. Accordingly, we have the following approximate variance-covariance matrix

$$I^{-1}(\hat{\alpha}) = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{bmatrix}^{-1}, \quad (3.4)$$

where

$$a_{11} = - \frac{\partial^2 \ln L}{\partial \alpha_1^2} \bigg|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3} = \frac{r}{\hat{\alpha}_1^2} + \frac{r}{(\hat{\alpha}_1 + \hat{\alpha}_3)^2} + 3(n-r) \frac{[u_{1r}(\beta, \gamma)]^{\hat{\alpha}_1} [\ln u_{1r}(\beta, \gamma)]^2}{(1 - [u_{1r}(\beta, \gamma)]^{\hat{\alpha}_1})^2} + 3 \sum_{i=1}^r (n-i) \frac{[u_{1i}(\beta, \gamma)]^{\hat{\alpha}_1} [\ln u_{1i}(\beta, \gamma)]^2}{(1 - [u_{1i}(\beta, \gamma)]^{\hat{\alpha}_1})^2},$$

$$a_{22} = - \frac{\partial^2 \ln L}{\partial \alpha_2^2} \bigg|_{\hat{\alpha}_1, \hat{\alpha}_2, \hat{\alpha}_3} = \frac{r}{\hat{\alpha}_2^2} + \frac{r}{(\hat{\alpha}_2 + \hat{\alpha}_3)^2},$$

3. Maximum Likelihood Estimation under Type II Censored Samples based on Concomitants of Order Statistics

In this Section, we will introduce the maximum likelihood estimates for the unknown parameters of the BGLFR distribution under type II censoring scheme. The detailed procedure is as follows

Suppose $\{(x_{1i}, x_{2i}), \dots, (x_{1n}, x_{2n})\}$ be a random sample from BGLFR distribution, this random sample is subject to type II censoring, and the experiment is terminated at time $X_{1(r)}$. Let $X_{1(i)}$ be the i^{th} order statistic of X_1 and $X_{2(i)}$ be its concomitant of X_2 . According to Nagaraja and Abo-Eleneen (2008) the likelihood function for a type II right censored sample is given as

$$L(\underline{\theta}) = \frac{n!}{(n-r)!} [1 - F_{X_1}(x_{1(r)})]^{n-r} \prod_{i=1}^r \frac{n!}{(i-1)!(n-i)!} f(x_{1i}, x_{2i}) [F_{X_1}(x_{1i})]^{i-1} [1 - F_{X_1}(x_{1i})]^{n-i}. \quad 1 \leq i \leq r < n. \quad (3.1)$$

In case of BGLFR, the likelihood function for a type II right censored sample $\{(x_{1(i)}, x_{2(i)}), 1 \leq i \leq r < n\}$ is given as

$$L(\underline{\alpha}) = \left(\frac{n!}{(n-r)!} [1 - F_{X_1}(x_{1(r)})]^{n-r} \right)^3 \prod_{i=1}^r \left(\frac{n!}{(i-1)!(n-i)!} [F_{X_1}(x_{1i})]^{i-1} [1 - F_{X_1}(x_{1i})]^{n-i} \right)^3 f_1(x_{1i}, x_{2i}) f_2(x_{1i}, x_{2i}) f_3(x_{1i}) \quad (3.2)$$

After substituting from (2.2), in (3.2) and taking the logarithm we get

$$\begin{aligned} \ln L(\underline{\alpha}) \propto & \{r \ln \alpha_1 + r \ln \alpha_2 + r \ln \alpha_3 + r \ln(\alpha_1 + \alpha_3) + r \ln(\alpha_2 + \alpha_3) \\ & + 3(n-r) \ln[1 - (u_{1r}(\beta, \gamma))^{\alpha_1}] + 3 \sum_{i=1}^r (n-i) \ln[1 - (u_{1i}(\beta, \gamma))^{\alpha_1}] \\ & + \sum_{i=1}^r (3i \alpha_1 + \alpha_2 + 2\alpha_3 - 3) \ln u_{1i}(\beta, \gamma) + (2\alpha_2 + \alpha_3 - 2) \sum_{i=1}^r \ln u_{2i}(\beta, \gamma)\}. \end{aligned} \quad (3.3)$$

where

$$u_{ji}(\beta, \gamma) = 1 - e^{-(\beta x_{ji} + \frac{\gamma}{2} x_{ji}^2)}, j = 1, 2 \text{ and}$$

$$\underline{\alpha} = (\alpha_1, \alpha_2, \alpha_3).$$

On differentiating (3.3) with respect to $\alpha_1, \alpha_2, \alpha_3$ in turn and equating to zero, we obtain the following likelihood equations

$$\begin{aligned} F_1(x_1, x_2) &= F_{GLFR}(x_1; \alpha_{13}, \beta, \gamma) F_{GLFR}(x_2; \alpha_2, \beta, \gamma) \\ &= [1 - e^{-(\beta x_1 + \frac{\gamma}{2} x_1^2)}]_{\alpha_{13}} [1 - e^{-(\beta x_2 + \frac{\gamma}{2} x_2^2)}]_{\alpha_2}, \\ F_2(x_1, x_2) &= F_{GLFR}(x_1; \alpha_1, \beta, \gamma) F_{GLFR}(x_2; \alpha_{23}, \beta, \gamma) \\ &= [1 - e^{-(\beta x_1 + \frac{\gamma}{2} x_1^2)}]_{\alpha_1} [1 - e^{-(\beta x_2 + \frac{\gamma}{2} x_2^2)}]_{\alpha_{23}}, \end{aligned}$$

$$\begin{aligned} F_3(x) &= F_{GLFR}(x; \alpha_{123}, \beta, \gamma) \\ &= [1 - e^{-(\beta x + \frac{\gamma}{2} x^2)}]_{\alpha_{123}}. \end{aligned}$$

Then joint pdf of (X_1, X_2) is given as

$$f_{BGLFR}(x_1, x_2) = \begin{cases} f_1(x_1, x_2) & \text{if } 0 < x_1 < x_2 < \infty \\ f_2(x_1, x_2) & \text{if } 0 < x_2 < x_1 < \infty \\ f_3(x) & \text{if } 0 < x_1 = x_2 = x < \infty. \end{cases} \quad (2.2)$$

where

$$\begin{aligned} f_1(x_1, x_2) &= f_{GLFR}(x_1; \alpha_{13}, \beta, \gamma) f_{GLFR}(x_2; \alpha_2, \beta, \gamma) \\ &= \alpha_{13} \alpha_2 (\beta + \gamma x_1) (\beta + \gamma x_2) e^{-(\beta x_1 + \frac{\gamma}{2} x_1^2)} e^{-(\beta x_2 + \frac{\gamma}{2} x_2^2)} \\ &\quad [1 - e^{-(\beta x_1 + \frac{\gamma}{2} x_1^2)}]_{\alpha_{13}-1} [1 - e^{-(\beta x_2 + \frac{\gamma}{2} x_2^2)}]_{\alpha_2-1}, \\ f_2(x_1, x_2) &= f_{GLFR}(x_1; \alpha_1, \beta, \gamma) f_{GLFR}(x_2; \alpha_{23}, \beta, \gamma) \\ &= \alpha_1 \alpha_{23} (\beta + \gamma x_1) (\beta + \gamma x_2) e^{-(\beta x_1 + \frac{\gamma}{2} x_1^2)} e^{-(\beta x_2 + \frac{\gamma}{2} x_2^2)} \\ &\quad [1 - e^{-(\beta x_1 + \frac{\gamma}{2} x_1^2)}]_{\alpha_1-1} [1 - e^{-(\beta x_2 + \frac{\gamma}{2} x_2^2)}]_{\alpha_{23}-1}, \end{aligned}$$

$$\text{and } f_3(x) = \alpha_3 (\beta + \gamma x) e^{-(\beta x + \frac{\gamma}{2} x^2)} [1 - e^{-(\beta x + \frac{\gamma}{2} x^2)}]_{\alpha_{123}-1}.$$

The marginal distributions of both X_1 and X_2 have GLFR $(\alpha_{13}, \beta, \gamma)$ and GLFR $(\alpha_{23}, \beta, \gamma)$, respectively. The bivariate generalized exponential distribution that defined by Kundu and Gupta (2009) can be considered as a special case of BGLFR distribution when $\gamma = 0$.

Epstein (1954) combined both type I and type II censoring schemes and called hybrid censoring. For this censoring scheme, the test is terminated either at the time T or at the time of the r th failure, say $x_{(r)}$ whichever is reached first when n independent and identical observations are put on test i.e., hybrid life test experiment terminates at $\text{Min}(x_{(r)}, T)$.

The maximum likelihood estimation, estimated variance-covariance matrix and asymptotic confidence intervals for BGLFR distribution under both type II and hybrid random censoring are provided in Section 3 and Section 4, respectively. In Section 5, we introduce numerical examples to illustrate the results. Finally, we write the conclusion of the paper in Section 6.

2. The Bivariate Generalized Linear Failure Rate Distribution

Sarhan et al. (2011) defined the BGLFR distribution as follows: Suppose U_1, U_2 and U_3 are three independent random variables such that $U_i \sim \text{GLFR}(\alpha_i, \beta, \gamma)$ for $i = 1, 2$ and 3 . Define $X_1 = \max(U_1, U_3)$ and $X_2 = \max(U_2, U_3)$, then it is said that the bivariate vector (X_1, X_2) has BGLFR distribution with parameters $(\alpha_1, \alpha_2, \alpha_3, \beta, \gamma)$, denoted by $\text{BGLFR}(\alpha_1, \alpha_2, \alpha_3, \beta, \gamma)$. Then, the joint cdf of (X_1, X_2) is given as follows

$$F_{\text{BGLFR}}(x_1, x_2) = \prod_{i=1}^3 F_{\text{GLFR}}(x_i; \alpha_i, \beta, \gamma)$$

where $x_3 = \min(x_1, x_2)$, $\alpha_{13} = \alpha_1 + \alpha_3$, $\alpha_{23} = \alpha_2 + \alpha_3$ and $\alpha_{123} = \alpha_1 + \alpha_2 + \alpha_3$.

the joint cdf of (X_1, X_2) can be written as

$$F_{\text{BGLFR}}(x_1, x_2) = \begin{cases} F_1(x_1, x_2) & \text{if } 0 < x_1 < x_2 < \infty \\ F_2(x_1, x_2) & \text{if } 0 < x_2 < x_1 < \infty \\ F_3(x) & \text{if } 0 < x_1 = x_2 = x < \infty \end{cases} \quad (2.1)$$

where

Suppose $(X_{1i}, X_{2i}), i = 1 \dots n$, is a random sample from an absolutely continuous bivariate population (X_1, X_2) . If we order the sample by the X_1 -variate, and obtain the order statistics $X_{1(1)} < X_{1(2)} < \dots < X_{1(n)}$ for the X_1 sample, then the X_2 -variate associated with the r^{th} order statistic $X_{1(r)}$ is called concomitant of the r^{th} order statistic, and is denoted by $X_{2[r]}$. The term concomitant order statistics was introduced by David (1973). Independently, Bhattacharya (1974) used the term induced order statistic for concomitant order statistics. One important application of concomitants arises in selection procedures when items or subjects may be chosen on the basis of their X_1 characteristic, for example, the score on a screening test for job applicants. While an associated characteristic X_2 that is hard to measure or can be observed only later may be of interest, such as the future performance for the job applicants. See Nagaraja and David (1994) for other examples and applications for concomitants.

There are several mechanisms that can lead to censored data. So-called type II censoring that describes the experiment would be continued until a fixed proportion of items have failed. A type II bivariate censored sample consists of $X_1 = (X_{1(1)}, X_{1(2)}, \dots, X_{1(r)})$ and $X_2 = (X_{2[1]}, X_{2[2]}, \dots, X_{2[r]})$ for some $r, 1 \leq r \leq n$. For many life testing problems, a type II censoring scheme is adopted, i.e. the experiment is run until the first r failures out of n occur, so that failure times $X_{1(1)}, X_{1(2)}, \dots, X_{1(r)}$ are observed. An experimenter may wish to study the relationship between the failure time and a set of concomitant variates. This procedure assumes that though $X_{1(1)}, X_{1(2)}, \dots, X_{1(r)}$ are only observable, the entire set $(X_{21}, X_{22}, \dots, X_{2n})$ is given prior to experimentation. In many life testing problems, $X_{2[i]}$ is observable only when $X_{1(i)}$ is so. This is particularly true if recording of X_{2i} necessitates the failure of the i^{th} unit. So that for the surviving $(n - r)$ units, the X_{2i} 's are not observable. For example, in an experiment designed to study the relationship between arteriosclerosis and length of life, rhesus monkeys are fed an atherogenic diet. For the first r monkeys to die an autopsy is performed to measure the amount of fatty plaque in the aorta. The experimenter may wish to quantify the strength of the relationship between this measurements and life length.

$$F_{GLFR}(x; \alpha, \beta, \gamma) = [1 - e^{-(\beta x + \frac{\gamma}{2} x^2)}]^\alpha, \quad x > 0, \alpha > 0, \beta > 0, \gamma > 0.$$

Its corresponding pdf has the form

$$f_{GLFR}(x; \alpha, \beta, \gamma) = \alpha(\beta + \gamma x) e^{-(\beta x + \frac{\gamma}{2} x^2)} [1 - e^{-(\beta x + \frac{\gamma}{2} x^2)}]^{\alpha-1}, \quad x > 0.$$

Several properties of this distribution were established by the author's. They observed that several known distributions like the exponential, the Rayleigh and the linear failure rate distribution distributions can be obtained as special cases of the GLFR distribution.

Many times the life/ failure data of interest is bivariate in nature. Any study on twins or on failure data recoded twice on the same system naturally leads to bivariate data. Paired data could also consist of blindness in the left / right eye, failure time of the left / right kidney or age at death of parent / child in a genetic study. For example, Houggard et al. (1992) studied data on life length of Danish twins and Lin et al. (1999) considered data of colon cancer and the time from treatment to death. Recently, Sarhan et al. (2011) introduced a new bivariate generalized linear failure rate (BGLFR) distribution, whose marginals were GLFR distribution. This new five parameters BGLFR distribution was obtained using a method similar to that used to obtain the Marshall – Olkin bivariate exponential model [Marshall and Olkin (1967)]. The proposed BGLFR was constructed from three independent GLFR distributions using a maximization process. This new distribution is a singular distribution, and it can be used quit conveniently if there are ties in the data, and it often used to fit paired data in survival studies where there is a possibility of simultaneous occurrence of both the events. Sarhan et al. (2011) provided many interpretations for the BGLFR model as competing risks model, Shock model, Stress Model and Maintenance Model. They discussed several properties of this distribution, and used the EM algorithm to compute the maximum likelihood estimators of the unknown parameters in the case of complete sampling. Aly and Muhammed (2012) introduced the moment generating function, the maximum likelihood estimators and asymptotic confidence intervals under random left censoring for BGLFR distribution.

Estimation of the Bivariate Generalized Linear Failure Rate Distribution under Different Censoring Schemes

Ahmed F. Attia¹, Hanan M. Aly² and Hiba Z. Muhammed³

Abstract

Recently a new distribution, named a bivariate generalized linear failure rate distribution has been introduced by Sarhan et al. (2011). In this paper, we get the maximum likelihood estimators for the unknown parameters of bivariate generalized linear failure rate distribution and their approximate variance-covariance matrix in case of type II censored samples based on concomitants of order statistics. Moreover, these estimators are obtained assuming hybrid random censored samples when censored times follow either exponential or Weibull distribution. Numerical examples are carried out to check the properties of the estimators.

Keywords: *Type II censoring, Hybrid censoring, Generalized linear failure rate distribution, Bivariate generalized linear failure rate distribution, Concomitants of order statistics.*

1. Introduction

In analyzing lifetime data, one often uses the exponential, Rayleigh, linear failure rate or generalized exponential distributions. It is well known that the exponential distribution can has constant hazard function whereas Rayleigh, linear failure rate and the generalized exponential distribution can have monotone (increasing in the case of Rayleigh or linear failure rate and increasing / decreasing in case of the generalized exponential distribution) hazard functions. Unfortunately, in practice often one needs to consider non-monotonic functions such as bathtub shaped hazard function. Recently Sarhan and Kundu (2009) presented a new distribution: called the generalized linear failure rate distribution (GLFR) which may has bathtub shaped hazard function. The GLFR has a cdf in the form

¹ Department of Mathematical Statistics, Institute of Statistical Studies and Research, Cairo University, email: Attia16152@yahoo.com.

² Department of Statistics, Faculty of Economics & Political Science, Cairo University, email: han_m_hasan@yahoo.com.

³ Department of Mathematical Statistics, Institute of Statistical Studies and Research, Cairo University, email: heba_stat@yahoo.com.

References

- 1- Basu, A.P. (1981), The estimation of $P(X < Y)$ for distributions useful in life testing, Naval Research Logistics Quarterly, Vol. 28, Issue 3, 383-392.
- 2- Barnett. V. and Lewis, T. (1994), Outliers in Statistical Data. John Wiley & Sons, 3rd edition.
- 3- Birnbaum, Z. W. (1956), On a use of the Mann-Whitney statistic, proceedings of the third barkeley symposium on mathematical statistics and probability, vol. I.
- 4- Birnbaum, Z. W. and McCarty (1956), R. C., A distribution free upper confidence bound for $P(Y < X)$, based on independent samples of X and Y, Annals of Mathematical Statistics, 29, 557-562.
- 5- Church, J.D. and Harris, B. (1970), The estimation of reliability from stress strength relationships, Technometrics, vol. 12, 49 - 54.
- 6- Constantine, K. & Karson, M. (1986), The estimation of $P(Y < X)$ in gamma case. Communications in Statistics - Computations and Simulations, Vol. 15, 365-388.
- 7- Davarzani, Nasser (2009), Estimation of $P(X < Y)$ for a Bivariate Weibull Distribution, Journal of Applied Probability & Statistics Vol. 4, No. 2, 227-238.
- 8- Deiri, E., (2011a), Estimation of $P(Y < X)$ for Exponential Distribution in the Presence of Two Outliers. International Journal of Academic Research, Vol. 3, No. 2, 508-514.
- 9- Deiri, E., (2011b), Estimation of $P(Y < X)$ for Generalized Exponential Distribution in the Presence of Two Outliers when Scale Parameter is Known. International Journal of Academic Research, Vol. 3, No. 2, 1179-1185.
- 10- Deiri, E., (2011c), Estimation of Parameters of the Gamma Distribution in the Presence of Two Outliers. Australian Journal of Basic and Applied Science, Vol. 5, No. 6, p 1496-1502.
- 11- Deiri, E., (2011d), Estimation of Parameters of the Minimax Distribution in the Presence of Two Outliers. Australian Journal of Basic and Applied Science, Vol. 5, No. 6, p 1510-1516.
- 12- Deiri, E., (2011e), Estimation of $P(Y < X)$ in the Rayleigh Distribution in the Presence of Two Outliers, Australian Journal of Basic & Applied Sciences, Vol. 5, No. 6, p1503-1509.
- 13- Dixit, U. J. Moor, K. L. and Barnett, V. (1996), on the estimation of the power of the scale parameter of the exponential distribution in the presence of outliers generated from uniform distribution. Metron, 54, 201-211.
- 14- Dixit, U. J. and Nasiri, P.F. (2001), Estimation of parameters of the exponential distribution in the presence of outliers generated from uniform distribution. Metron, 49(3 - 4), p 187- 198.
- 15- Ghanizadeh, A., (2011), Estimation of $P(Y < X)$ in the rayleigh distribution in the presence of k outliers, Journal of Applied Sciences, Vol. 11, No. 17, p3192-3197.
- 16- Jafari, A., (2011), Estimation of $P(Y < X)$ in the Rayleigh Distribution with Presence of Outlier. International Journal of Academic Research, Vol. 3, No. 2, 528-534.
- 17- Kotz, S., Lumelskii, Y., Pensky, M., (2003), the stress- strength model and its generalizations, "Theory and Applications".
- 18- Krishnamoorthy, K., Shubhabrata, M., Guo, H., (2007) Inference on reliability in two-parameter exponential "stress-strength model", Metrika, Vol. 65, issue 3, pages 261-273.
- 19- Kundu, D. and Gupta, R.D. (2006), "Estimation of $R = P(Y < X)$ for Weibull distribution", IEEE Transactions on Reliability, vol. 55, 270 - 280.
- 20- Kundu, D. and Raqab. M.Z. (2009), "Estimation of $R = P(Y < X)$ for three parameter Weibull distribution", Statistics and Probability Letters, vol. 79, 1839 - 1846.
- 21- McCool, J.I. (1991), Inference on $P(Y < X)$ in the Weibull case. Communications in Statistics - Simulation and Computations, Vol. 20, 129 - 148.
- 22- Nasiri, P.F. and H. Pazira, (2010), Bayesian and non-Bayesian estimations on the Generalized Exponential distribution in the presence outliers, Journal of Statistical Theory and Practice, 4(3): 453-475.
- 23- Read, R. R. (1981), Representation of certain covariance matrices with application to asymptotic efficiency. Journal of the American Statistical Associate, 76, 148-154.

Table (5)

Estimates of R, Biases and Mean Squared Errors MSE's of the point estimates from Weibull Distribution, when $k = 1$, $\beta = 1.5$, $\alpha = 0.5$, $\gamma = 2$ and $p = 1$

n, m	$\hat{R}$			BIAS			MSE		
	MLE	Mom	Mix	MLE	Mom	Mix	MLE	Mom	Mix
15,15	0.231	0.287	0.312	0.077	0.057	0.121	0.092	0.065	0.081
20,20	0.249	0.276	0.341	0.065	-0.043	-0.107	0.083	0.068	0.087
25,25	0.247	0.303	0.362	0.055	-0.037	0.109	0.094	0.079	0.065
15,20	0.256	0.292	0.382	-0.057	0.035	0.113	0.088	0.088	0.077
20,15	0.251	0.298	0.354	0.068	0.049	0.126	0.087	0.067	0.083
15,25	0.234	0.308	0.322	0.086	0.057	-0.125	0.092	0.076	0.083
25,15	0.256	0.287	0.387	-0.092	-0.069	0.131	0.071	0.059	0.071
20,25	0.248	0.294	0.353	0.079	0.067	-0.117	0.094	0.067	0.090
25,20	0.262	0.314	0.361	0.108	-0.087	0.132	0.092	0.085	0.057

Table (6)

Estimates of R, Biases and Mean Squared Errors MSE's of the point estimates from Weibull Distribution, when $k = 2$, $\beta = 1.5$, $\alpha = 0.5$, $\gamma = 2$ and $p = 1$

n, m	$\hat{R}$			BIAS			MSE		
	MLE	Mom	Mix	MLE	Mom	Mix	MLE	Mom	Mix
15,15	0.306	0.331	0.375	0.105	0.114	-0.091	0.083	0.156	0.085
20,20	0.311	0.342	0.352	-0.098	-0.123	0.086	0.069	0.134	0.093
25,25	0.321	0.352	0.348	0.076	0.108	-0.065	0.066	0.144	0.103
15,20	0.335	0.367	0.394	0.089	0.099	0.076	0.087	0.161	0.075
20,15	0.347	0.356	0.371	-0.114	-0.093	0.077	0.093	0.187	0.093
15,25	0.351	0.385	0.388	0.086	-0.104	0.079	0.067	0.154	0.077
25,15	0.337	0.382	0.403	0.109	0.103	0.089	0.059	0.134	0.083
20,25	0.363	0.397	0.396	-0.071	0.107	-0.091	0.058	0.162	0.096
25,20	0.342	0.403	0.408	0.096	-0.097	0.076	0.087	0.167	0.109

Table (3)

Estimates of R, Biases and Mean Squared Errors MSE's of the point estimates from Weibull Distribution, when $k = 1$, $\beta = 2$, $\alpha = 0.5$, $\gamma = 2$ and $p = 1$

n, m	$\hat{R}$			BIAS			MSE		
	MLE	Mom	Mix	MLE	Mom	Mix	MLE	Mom	Mix
15,15	0.321	0.374	0.402	0.108	0.059	0.142	0.152	0.072	0.087
20,20	0.319	0.358	0.425	-0.098	-0.048	0.123	0.144	0.065	0.092
25,25	0.321	0.362	0.392	0.087	-0.027	0.115	0.134	0.084	0.081
15,20	0.331	0.371	0.418	-0.084	0.048	0.126	0.178	0.091	0.087
20,15	0.342	0.382	0.457	-0.099	0.059	0.134	0.155	0.069	0.093
15,25	0.341	0.369	0.445	0.105	-0.068	0.152	0.192	0.081	0.104
25,15	0.346	0.384	0.427	-0.112	0.079	0.143	0.171	0.064	0.084
20,25	0.365	0.374	0.432	-0.107	0.096	0.137	0.182	0.074	0.113
25,20	0.367	0.398	0.465	0.123	-0.107	0.141	0.148	0.097	0.109

Table (4)

Estimates of R, Biases and Mean Squared Errors MSE's of the point estimates from Weibull Distribution, when $k = 2$, $\beta = 2$, $\alpha = 0.5$, $\gamma = 2$ and $p = 1$

n, m	$\hat{R}$			BIAS			MSE		
	MLE	Mom	Mix	MLE	Mom	Mix	MLE	Mom	Mix
15,15	0.337	0.353	0.377	0.077	0.108	0.088	0.057	0.213	0.066
20,20	0.325	0.382	0.412	-0.078	-0.098	0.067	0.049	0.192	0.089
25,25	0.343	0.394	0.418	-0.059	-0.075	0.081	0.052	0.184	0.086
15,20	0.352	0.397	0.414	-0.064	-0.069	0.055	0.078	0.211	0.072
20,15	0.361	0.377	0.413	0.076	0.087	0.072	0.086	0.231	0.061
15,25	0.365	0.386	0.423	0.086	-0.096	0.078	0.092	0.207	0.087
25,15	0.359	0.378	0.433	-0.097	0.103	0.086	0.077	0.198	0.055
20,25	0.381	0.386	0.425	0.081	-0.109	0.091	0.087	0.215	0.067
25,20	0.373	0.397	0.423	0.092	0.113	0.086	0.069	0.197	0.093

Table (1)

Estimates of R, Biases and Mean Squared Errors MSE's of the point estimates from Weibull Distribution, when $k = 1$, $\beta = 1$, $\alpha = 0.5$, $\gamma = 2$ and $p = 1$

n, m	$\hat{R}$			Bias			MSE		
	MLE	Mom	Mix	MLE	Mom	Mix	MLE	Mom	Mix
15,15	0.207	0.260	0.317	0.059	0.053	-0.110	0.209	0.025	0.018
20,20	0.203	0.253	0.288	0.043	0.032	-0.085	0.196	0.012	0.010
25,25	0.201	0.244	0.257	0.039	-0.011	-0.071	0.155	0.009	0.007
15,20	0.205	0.254	0.315	0.042	0.037	-0.092	0.214	0.031	0.013
20,15	0.213	0.261	0.332	0.054	0.046	-0.088	0.238	0.053	0.027
15,25	0.215	0.243	0.298	0.047	-0.048	0.098	0.241	0.037	0.031
25,15	0.221	0.253	0.302	0.053	-0.061	-0.077	0.247	0.067	0.026
20,25	0.241	0.282	0.317	0.044	-0.077	-0.103	0.211	0.045	0.022
25,20	0.238	0.279	0.322	0.048	0.081	-0.113	0.216	0.057	0.053

Table (2)

Estimates of R, Biases and Mean Squared Errors MSE's of the point estimates from Weibull Distribution, when $k = 2$, $\beta = 1$, $\alpha = 0.5$, $\gamma = 2$ and $p = 1$

n, m	$\hat{R}$			BIAS			MSE		
	MLE	Mom	Mix	MLE	Mom	Mix	MLE	Mom	Mix
15,15	0.220	0.274	0.322	0.091	0.063	0.160	0.254	0.045	0.027
20,20	0.217	0.268	0.303	0.075	-0.044	-0.091	0.234	0.033	0.019
25,25	0.216	0.272	0.272	0.078	0.023	-0.083	0.201	0.034	0.016
15,20	0.221	0.282	0.323	0.065	0.051	0.094	0.264	0.053	0.021
20,15	0.226	0.291	0.352	0.084	-0.062	-0.092	0.286	0.066	0.039
15,25	0.231	0.269	0.314	0.085	0.065	0.106	0.297	0.051	0.038
25,15	0.236	0.283	0.317	0.098	-0.083	0.095	0.298	0.076	0.034
20,25	0.255	0.312	0.336	0.087	-0.107	-0.114	0.252	0.054	0.035
25,20	0.257	0.306	0.341	0.104	-0.113	0.121	0.262	0.068	0.061

$$\frac{\partial l(p, \gamma)}{\partial p} = -n + \left(\frac{1}{p}\right)^\beta \sum_{i=1}^n x_i^\beta + \frac{b}{b} \left(\frac{p}{\gamma}\right)^\beta \sum_{i=1}^n \frac{\left(\left(\frac{x_i}{p}\right)^\beta - \left(\frac{x_i}{\gamma}\right)^\beta\right) e^{-\left(\frac{1}{\gamma^\beta} - \frac{1}{p^\beta}\right)x_i^\beta}}{\psi(x_i; p, \gamma)} = 0 \quad (C)$$

There is no closed-form solution to this system of equations (A, B), so we will solve for $\hat{p}, \hat{\gamma}$ iteratively. In case of no outlier presence, $R, \hat{p}$ and $\hat{\gamma}$ were proposed by Kundu, D., Gupta, R.D. (2006).

4. Mixture of method of Moment and Maximum likelihood

Read (1981) proposed the methods, which avoid the difficulty of complicated equations. According to Read (1981), replacement of some, but not all, of the equations in the system of likelihood may make it more manageable. From moment estimation, we have (21)

$$\hat{\gamma} = \frac{-\xi_2 + \sqrt{\xi_2^2 - 4\xi_1\xi_3}}{2\xi_1}$$

and, from equation (C)

$$= -n + \left(\frac{1}{p}\right)^\beta \sum_{i=1}^n x_i^\beta + \frac{b}{b} \left(\frac{p}{\gamma}\right)^\beta \sum_{i=1}^n \frac{\left(\left(\frac{x_i}{p}\right)^\beta - \left(\frac{x_i}{\gamma}\right)^\beta\right) e^{-\left(\frac{1}{\gamma^\beta} - \frac{1}{p^\beta}\right)x_i^\beta}}{\psi(x_i; p, \gamma)} = 0$$

We obtain $\hat{p}$.

5. Conclusion

In this paper, the researcher has addressed the problem of estimating $P(Y < X)$ for the Weibull distribution with presence of k outliers. The moment, maximum likelihood and mixture estimators of R are derived using MathCAD Package. All the results are base on 10000 replications and are given in Tables 1 to 6.

From Tables, the researcher infers that the moment estimator of R is asymptotically unbiased and as expected when $m=n$ and m, n increase then the average biases and the MSEs decrease, and the MSEs of any three estimators are tending to zero and when $m \rightarrow \infty$.

Tables show the estimates of $R=P(Y < X)$ with presence of one and two outliers under different values of β, γ, p and α , on the other hand they show some of the previous results in the literatures such as Ghanizadeh (2011) and Deiri (2011e) can be achieved as special case of our results.

Tables show that almost the mixture estimator has the smallest estimated MSEs as compared with the moment and maximum likelihood estimators, so the researcher strongly feels that mixture estimator is better and easy to calculate than the maximum likelihood and moment estimations. From the previous observations, the researcher suggests to use mixture method for estimating $R = P(Y < X)$ in the Weibull distribution in the presence of k outliers because it is easy to calculate than the rest.

$$= \frac{b\beta}{b\gamma} \left(\frac{p}{\gamma}\right)^{\beta} e^{-\left(\frac{1}{\gamma^{\beta}} - \frac{1}{p^{\beta}}\right)x_i^{\beta}} \left[1 - \left(\frac{x_i}{p}\right)^{\beta}\right]$$

$$= \frac{\frac{b\beta}{b\gamma} \left(\frac{p}{\gamma}\right)^{\beta} e^{-\left(\frac{1}{\gamma^{\beta}} - \frac{1}{p^{\beta}}\right)x_i^{\beta}} \left[1 - \left(\frac{x_i}{p}\right)^{\beta}\right]}{1 + \frac{b}{b} \left(\frac{p}{\gamma}\right)^{\beta} e^{-\left(\frac{1}{\gamma^{\beta}} - \frac{1}{p^{\beta}}\right)x_i^{\beta}}}$$

$$\frac{\partial \psi}{\partial \gamma} = \frac{b}{b} \left[\left(\frac{p}{\gamma}\right)^{\beta} e^{-\left(\frac{1}{\gamma^{\beta}} - \frac{1}{p^{\beta}}\right)x_i^{\beta}} \beta \gamma^{-\beta-1} x_i^{\beta} - \beta e^{-\left(\frac{1}{\gamma^{\beta}} - \frac{1}{p^{\beta}}\right)x_i^{\beta}} - \beta e^{-\left(\frac{1}{\gamma^{\beta}} - \frac{1}{p^{\beta}}\right)x_i^{\beta}} \left(\frac{p}{\gamma}\right)^{\beta-1} p \gamma^{-2} \right]$$

$$= \frac{b\beta}{b\gamma} \left(\frac{p}{\gamma}\right)^{\beta} e^{-\left(\frac{1}{\gamma^{\beta}} - \frac{1}{p^{\beta}}\right)x_i^{\beta}} \left[\left(\frac{x_i}{\gamma}\right)^{\beta} - 1 \right]$$

$$\frac{\partial \ln \psi}{\partial \gamma} = \frac{1}{\psi} \frac{\partial \psi}{\partial \gamma}$$

$$= \frac{\frac{b\beta}{b\gamma} \left(\frac{p}{\gamma}\right)^{\beta} e^{-\left(\frac{1}{\gamma^{\beta}} - \frac{1}{p^{\beta}}\right)x_i^{\beta}} \left[\left(\frac{x_i}{\gamma}\right)^{\beta} - 1 \right]}{1 + \frac{b}{b} \left(\frac{p}{\gamma}\right)^{\beta} e^{-\left(\frac{1}{\gamma^{\beta}} - \frac{1}{p^{\beta}}\right)x_i^{\beta}}}$$

$$\frac{\partial l(p, \gamma)}{\partial \gamma} = \frac{b\beta}{b\gamma} \left(\frac{p}{\gamma}\right)^{\beta} \sum_{i=1}^n \frac{\left[\left(\frac{x_i}{\gamma}\right)^{\beta} - 1 \right] e^{-\left(\frac{1}{\gamma^{\beta}} - \frac{1}{p^{\beta}}\right)x_i^{\beta}}}{\psi(x_i; p, \gamma)} = 0$$

$$\frac{\partial l(p, \gamma)}{\partial \gamma} = \sum_{i=1}^n \frac{\left[\left(\frac{x_i}{\gamma}\right)^{\beta} - 1 \right] e^{-\left(\frac{1}{\gamma^{\beta}} - \frac{1}{p^{\beta}}\right)x_i^{\beta}}}{\psi(x_i; p, \gamma)} = 0$$

$$\sum_{i=1}^n \frac{\left(\frac{x_i}{\gamma}\right)^{\beta} e^{-\left(\frac{1}{\gamma^{\beta}} - \frac{1}{p^{\beta}}\right)x_i^{\beta}}}{\psi(x_i; p, \gamma)} - \sum_{i=1}^n \frac{e^{-\left(\frac{1}{\gamma^{\beta}} - \frac{1}{p^{\beta}}\right)x_i^{\beta}}}{\psi(x_i; p, \gamma)} = 0$$

$$\sum_{i=1}^n \frac{\left(\frac{x_i}{\gamma}\right)^{\beta} e^{-\left(\frac{1}{\gamma^{\beta}} - \frac{1}{p^{\beta}}\right)x_i^{\beta}}}{\psi(x_i; p, \gamma)} = \sum_{i=1}^n \frac{e^{-\left(\frac{1}{\gamma^{\beta}} - \frac{1}{p^{\beta}}\right)x_i^{\beta}}}{\psi(x_i; p, \gamma)} \quad (B)$$

From (B) in (A)

$$-n + \left(\frac{1}{p}\right)^{\beta} \sum_{i=1}^n x_i^{\beta} + \frac{b}{b} \left(\frac{p}{\gamma}\right)^{\beta} \left[\sum_{i=1}^n \frac{\left(\frac{x_i}{p}\right)^{\beta} e^{-\left(\frac{1}{\gamma^{\beta}} - \frac{1}{p^{\beta}}\right)x_i^{\beta}}}{\psi(x_i; p, \gamma)} - \sum_{i=1}^n \frac{\left(\frac{x_i}{\gamma}\right)^{\beta} e^{-\left(\frac{1}{\gamma^{\beta}} - \frac{1}{p^{\beta}}\right)x_i^{\beta}}}{\psi(x_i; p, \gamma)} \right] = 0$$

$$= \prod_{i=1}^n \left[\frac{k}{n} \frac{\beta}{p^\beta} (x_i)^{\beta-1} e^{-\left(\frac{x_i}{p}\right)^\beta} \left[1 + \frac{(n-k)}{k} \left(\frac{p}{\gamma}\right)^\beta e^{-\left(\left(\frac{x_i}{\gamma}\right)^\beta - \left(\frac{x_i}{p}\right)^\beta\right)} \right] \right]$$

$$= \prod_{i=1}^n \left[\frac{k}{n} \frac{\beta}{p^\beta} (x_i)^{\beta-1} e^{-\left(\frac{x_i}{p}\right)^\beta} (\psi(x_i; p, \gamma)) \right]$$

where:

$$\psi(x_i; p, \gamma) = 1 + \frac{b}{b} \left(\frac{p}{\gamma}\right)^\beta e^{-\left(\frac{1}{\gamma^\beta} - \frac{1}{p^\beta}\right)x_i^\beta}$$

$$L(x, p, \gamma) = \left(\frac{b\beta}{p^\beta}\right)^n e^{-\sum_{i=1}^n \left(\frac{x_i}{p}\right)^\beta} \left[\prod_{i=1}^n x_i^{\beta-1} \right] \left[\prod_{i=1}^n \psi(x_i; p, \gamma) \right]$$

$$l(p, \gamma) = \ln L(x, p, \gamma)$$

$$l(p, \gamma) = n \ln(b\beta) - n\beta \ln p - \frac{1}{p^\beta} \sum_{i=1}^n x_i^\beta + (\beta - 1) \sum_{i=1}^n \ln x_i + \sum_{i=1}^n \ln \psi(x_i; p, \gamma)$$

The solve for our MLEs of p and γ we take the derivative of the log likelihood ($l(p, \gamma)$) with respect to each parameter set the partial derivatives equal to zero and solve for $\hat{p}$ and $\hat{\gamma}$:

$$\frac{\partial l(p, \gamma)}{\partial p} = \frac{-n\beta}{p} + \beta p^{-\beta-1} \sum_{i=1}^n x_i^\beta + \frac{b\beta}{p^\beta} \left(\frac{p}{\gamma}\right)^\beta \sum_{i=1}^n \frac{\left(1 - \left(\frac{x_i}{p}\right)^\beta\right) e^{-\left(\frac{1}{\gamma^\beta} - \frac{1}{p^\beta}\right)x_i^\beta}}{\psi(x_i; p, \gamma)} = 0$$

$$= -n + \left(\frac{1}{p}\right)^\beta \sum_{i=1}^n x_i^\beta + \frac{b}{b} \left(\frac{p}{\gamma}\right)^\beta \sum_{i=1}^n \frac{\left(1 - \left(\frac{x_i}{p}\right)^\beta\right) e^{-\left(\frac{1}{\gamma^\beta} - \frac{1}{p^\beta}\right)x_i^\beta}}{\psi(x_i; p, \gamma)} = 0$$

$$= -n + \left(\frac{1}{p}\right)^\beta \sum_{i=1}^n x_i^\beta + \frac{b}{b} \left(\frac{p}{\gamma}\right)^\beta \left[\sum_{i=1}^n \frac{\left(\frac{x_i}{p}\right)^\beta e^{-\left(\frac{1}{\gamma^\beta} - \frac{1}{p^\beta}\right)x_i^\beta}}{\psi(x_i; p, \gamma)} - \sum_{i=1}^n \frac{e^{-\left(\frac{1}{\gamma^\beta} - \frac{1}{p^\beta}\right)x_i^\beta}}{\psi(x_i; p, \gamma)} \right] = 0 \quad (A)$$

where

$$\frac{\partial \ln \psi}{\partial p} = \frac{1}{\psi} \frac{\partial \psi}{\partial p}$$

$$\frac{\partial \psi}{\partial p} = \frac{b}{b} \left[\left(\frac{p}{\gamma}\right)^\beta e^{-\left(\frac{1}{\gamma^\beta} - \frac{1}{p^\beta}\right)x_i^\beta} (-\beta) p^{-\beta-1} x_i^\beta + \frac{\beta}{\gamma} e^{-\left(\frac{1}{\gamma^\beta} - \frac{1}{p^\beta}\right)x_i^\beta} \left(\frac{p}{\gamma}\right)^{\beta-1} \right]$$

Let $y_1, y_2, \dots, y_m$ be a random sample and let $g(y_i, \alpha, \beta)$ denotes the Weibull distribution function

$$g(y_i, \alpha, \beta) = \frac{\beta}{\alpha} \left(\frac{y_i}{\alpha}\right)^{\beta-1} e^{-\left(\frac{y_i}{\alpha}\right)^\beta}$$

The likelihood function for weibull distribution given by

$$L = \prod_{i=1}^m g(y_i, \alpha, \beta) = \prod_{i=1}^m \left[\frac{\beta}{\alpha} \left(\frac{y_i}{\alpha}\right)^{\beta-1} e^{-\left(\frac{y_i}{\alpha}\right)^\beta} \right]$$

$$L = \left(\frac{\beta}{\alpha}\right)^m \left(\frac{1}{\alpha}\right)^{m(\beta-1)} e^{-\sum_{i=1}^m \left(\frac{y_i}{\alpha}\right)^\beta} \prod_{i=1}^m [(y_i)^{\beta-1}]$$

The log-likelihood given by

$$l = \ln(L) = m \ln(\beta) - m \ln(\alpha) - m(\beta - 1) \ln(\alpha) - \alpha^{-\beta} \sum_{i=1}^m (y_i)^\beta + (\beta - 1) \sum_{i=1}^m \ln y_i$$

The partial derivatives of l with respect to α is

$$\frac{\partial l}{\partial \alpha} = \frac{-m}{\alpha} - \frac{m(\beta - 1)}{\alpha} + \beta \alpha^{-\beta} \sum_{i=1}^m (y_i)^\beta = 0$$

$$\frac{\partial l}{\partial \alpha} = -\frac{m\beta}{\alpha} + \beta \alpha^{-\beta-1} \sum_{i=1}^m (y_i)^\beta = 0$$

$$\frac{1}{\alpha^\beta} \sum_{i=1}^m y_i^\beta = m$$

$$\hat{\alpha} = \left(\frac{1}{m} \sum_{i=1}^m y_i^\beta \right)^{\frac{1}{\beta}}$$

And to be the maximum Likelihood estimator of p and γ , we consider the joint distribution of X with presence of k outliers:

$$L(x, p, \gamma) = \prod_{i=1}^n [f(x_i; \gamma, p, \beta)]$$

$$= \prod_{i=1}^n \left[\frac{k}{n} \frac{\beta}{p} \left(\frac{x_i}{p}\right)^{\beta-1} e^{-\left(\frac{x_i}{p}\right)^\beta} + \frac{(n-k)}{n} \frac{\beta}{\gamma} \left(\frac{x_i}{\gamma}\right)^{\beta-1} e^{-\left(\frac{x_i}{\gamma}\right)^\beta} \right]$$

$$b \left(\frac{m'_1}{b \Gamma \frac{1}{\beta} + 1} - \frac{\bar{b}\gamma}{b} \right)^2 + \bar{b}\gamma^2 = \frac{m'_2}{\Gamma \frac{2}{\beta} + 1}$$

$$b \left[\frac{m_1'^2}{b^2 \left(\Gamma \frac{1}{\beta} + 1 \right)^2} - 2 \frac{\bar{b}m'_1}{b^2 \Gamma \frac{1}{\beta} + 1} \gamma + \frac{\bar{b}^2}{b^2} \gamma^2 \right] + \bar{b}\gamma^2 = \frac{m'_2}{\Gamma \frac{2}{\beta} + 1}$$

$$\frac{m_1'^2}{b \left(\Gamma \frac{1}{\beta} + 1 \right)^2} - 2 \frac{\bar{b}m'_1}{b \Gamma \frac{1}{\beta} + 1} \gamma + \bar{b} \left(\frac{\bar{b}}{b} + 1 \right) \gamma^2 = \frac{m'_2}{\Gamma \frac{2}{\beta} + 1}$$

$$\bar{b} \left(\frac{\bar{b}}{b} + 1 \right) \gamma^2 - 2 \frac{\bar{b}m'_1}{b \Gamma \frac{1}{\beta} + 1} \gamma + \left[\frac{m_1'^2}{b \left(\Gamma \frac{1}{\beta} + 1 \right)^2} - \frac{m'_2}{\Gamma \frac{2}{\beta} + 1} \right] = 0$$

Let

$$\xi_1 = \bar{b} \left(\frac{\bar{b}}{b} + 1 \right), \quad \xi_2 = -2 \frac{\bar{b}m'_1}{b \Gamma \frac{1}{\beta} + 1}, \quad \text{and} \quad \xi_3 = \frac{m_1'^2}{b \left(\Gamma \frac{1}{\beta} + 1 \right)^2} - \frac{m'_2}{\Gamma \frac{2}{\beta} + 1}$$

$$\xi_1 \gamma^2 + \xi_2 \gamma + \xi_3 = 0$$

if $\Delta = \xi_2^2 - 4\xi_1\xi_3$ is non-negative then the roots are real. Therefore

$$\hat{\gamma} = \frac{-\xi_2 + \sqrt{\xi_2^2 - 4\xi_1\xi_3}}{2\xi_1}$$

and, from (2)

$$\hat{p} = \frac{m'_1}{b \Gamma \frac{1}{\beta} + 1} - \frac{\bar{b}}{b} \hat{\gamma}$$

3. Method of Maximum Likelihood

The maximum Likelihood estimator of R is shown to be

$$\hat{R} = \frac{b\hat{\alpha}^\beta}{\hat{p}^\beta + \hat{\alpha}^\beta} + \frac{\bar{b}\hat{\alpha}^\beta}{\hat{\alpha}^\beta + \hat{\gamma}^\beta}$$

Where $\hat{\alpha} = \left(\frac{1}{m} \sum_{i=1}^m y_i^\beta \right)^{\frac{1}{\beta}}$

$\hat{\alpha}$, $\hat{\gamma}$ and $\hat{p}$ can be obtained as;

$$\begin{aligned} E(y^r) &= \int_0^{\infty} y^r g(y) dy \\ &= \int_0^{\infty} \frac{\beta}{\alpha^{\beta}} y^{\beta+r-1} e^{-\left(\frac{y}{\alpha}\right)^{\beta}} dy \\ &= \alpha^r \Gamma \frac{r}{\beta} + 1 \end{aligned}$$

at $r=1$ then $E(y) = \alpha \Gamma \frac{1}{\beta} + 1 = \frac{\sum_{i=1}^m y_i}{m}$

$$\hat{\alpha} = \frac{\sum_{i=1}^m y_i}{m \Gamma \frac{1}{\beta} + 1}$$

$$\begin{aligned} E(x^r) &= \int_0^{\infty} x^r f(x) dx \\ &= b \int_0^{\infty} \frac{\beta}{p^{\beta}} (x)^{\beta+r-1} e^{-\left(\frac{x}{p}\right)^{\beta}} + \bar{b} \int_0^{\infty} \frac{\beta}{\gamma^{\beta}} (x)^{\beta+r-1} e^{-\left(\frac{x}{\gamma}\right)^{\beta}} dx \\ &= (bp^r + \bar{b}\gamma^r) \Gamma \frac{r}{\beta} + 1 \end{aligned}$$

at $r=1$

$$(bp + \bar{b}\gamma) \Gamma \frac{1}{\beta} + 1 = \frac{\sum_{i=1}^n x_i}{n} = m'_1 \quad (1)$$

at $r=2$

$$(bp^2 + \bar{b}\gamma^2) \Gamma \frac{2}{\beta} + 1 = \frac{\sum_{i=1}^n x_i^2}{n} = m'_2 \quad (2)$$

From (1)

$$bp \Gamma \frac{1}{\beta} + 1 = m'_1 - \bar{b}\gamma \Gamma \frac{1}{\beta} + 1$$

$$p = \frac{m'_1}{b \Gamma \frac{1}{\beta} + 1} - \frac{\bar{b}\gamma}{b} \quad (3)$$

From (3) in (2)

And by the same way, we can obtain I_2

$$I_2 = \frac{(n-k)\alpha^\beta}{n[\alpha^\beta + \gamma^\beta]}$$

Let

$$\frac{k}{n} = b, \quad \frac{(n-k)}{n} = \bar{b}$$

where $b + \bar{b} = 1$

$$I = \left[\frac{b\alpha^\beta}{p^\beta + \alpha^\beta} + \frac{\bar{b}\alpha^\beta}{\alpha^\beta + \gamma^\beta} \right]$$

$$P(y < x) = 1 - \left[\frac{b\alpha^\beta}{p^\beta + \alpha^\beta} + \frac{\bar{b}\alpha^\beta}{\alpha^\beta + \gamma^\beta} \right]$$

Since $P(y < x) + P(y > x) = 1$, we consider $P(y > x)$,

$$R = P(y > x) = 1 - P(y < x)$$

$$R = \frac{b\alpha^\beta}{p^\beta + \alpha^\beta} + \frac{\bar{b}\alpha^\beta}{\alpha^\beta + \gamma^\beta}$$

Special cases

a. Let $\beta = 2$, $p^2 = \frac{\theta}{\beta}$, $\alpha^2 = \lambda$, $\gamma^2 = \theta$

Then

$$R = \frac{b\lambda\beta}{\lambda\beta + \theta} + \frac{\bar{b}\lambda}{\lambda + \theta}$$

see Ghanizadeh(2011)

Let $\beta = 1$, $p = \frac{\theta}{\beta}$, $\alpha = \lambda$, $\gamma = \theta$

Then

$$R = \frac{b\lambda\beta}{\lambda\beta + \theta} + \frac{\bar{b}\lambda}{\lambda + \theta}$$

see Deiri (2011e)

In section 2, 3 and 4, we deal with the methods of moment, maximum likelihood and mixture of methods of moment and maximum likelihood in estimating R . We compare these methods numerically, and finally the researcher represents conclusion in section 5.

2. Method of Moment

The moment estimator of R can be shown to be

$$\hat{R} = \frac{b\hat{\alpha}^\beta}{\hat{p}^\beta + \hat{\alpha}^\beta} + \frac{\bar{b}\hat{\alpha}^\beta}{\hat{\alpha}^\beta + \hat{\gamma}^\beta}$$

$$f_2(x; \gamma, \beta) = \frac{\beta}{\gamma} \left(\frac{x}{\gamma}\right)^{\beta-1} e^{-\left(\frac{x}{\gamma}\right)^\beta} \quad x > 0, \beta > 0 \text{ and } \gamma > 0$$

then the marginal distribution of X is

$$f(x; \gamma, p, \beta) = \frac{k\beta}{n p} \left(\frac{x}{p}\right)^{\beta-1} e^{-\left(\frac{x}{p}\right)^\beta} + \frac{(n-k)\beta}{n \gamma} \left(\frac{x}{\gamma}\right)^{\beta-1} e^{-\left(\frac{x}{\gamma}\right)^\beta} \quad x > 0, \beta > 0, \gamma > 0 \text{ and } p > 0$$

Let $(Y_1, Y_2, \dots, Y_m)$ be a random sample for Y with pdf,

$$g(y; \alpha; \beta) = \frac{\beta}{\alpha} \left(\frac{y}{\alpha}\right)^{\beta-1} e^{-\left(\frac{y}{\alpha}\right)^\beta}$$

where X, Y are independent.

With presence of k outliers, the parameter R we want to estimate is

$$\begin{aligned} R &= P(y < x) \\ P(y < x) &= \int_0^\infty \int_0^x g(y) f(x) dy dx \\ &= \int_0^\infty f(x) \left[\int_0^x g(y) dy \right] dx \\ \int_0^x g(y) dy &= 1 - e^{-\left(\frac{y}{\alpha}\right)^\beta} \\ P(y < x) &= \int_0^\infty f(x) \left[1 - e^{-\left(\frac{y}{\alpha}\right)^\beta} \right] dx \\ P(y < x) &= 1 - \int_0^\infty f(x) e^{-\left(\frac{y}{\alpha}\right)^\beta} dx = 1 - I \\ I &= \int_0^\infty f(x) e^{-\left(\frac{y}{\alpha}\right)^\beta} dx \\ I &= \int_0^\infty \frac{k\beta}{n p} \left(\frac{x}{p}\right)^{\beta-1} e^{-\left[\left(\frac{x}{p}\right)^\beta + \left(\frac{y}{\alpha}\right)^\beta\right]} dx + \int_0^\infty \frac{(n-k)\beta}{n \gamma} \left(\frac{x}{\gamma}\right)^{\beta-1} e^{-\left[\left(\frac{x}{\gamma}\right)^\beta + \left(\frac{y}{\alpha}\right)^\beta\right]} dx \\ I &= I_1 + I_2 \\ I_1 &= \frac{k\beta}{n p^\beta} \int_0^\infty x^{\beta-1} e^{-[p^{-\beta} + \alpha^{-\beta}]x^\beta} dx \\ &= \frac{k}{n p^\beta [p^{-\beta} + \alpha^{-\beta}]} \int_0^\infty \beta [p^{-\beta} + \alpha^{-\beta}] x^{\beta-1} e^{-[p^{-\beta} + \alpha^{-\beta}]x^\beta} dx \\ &= \frac{k}{n \left[1 + \left(\frac{p}{\alpha}\right)^\beta\right]} = \frac{k\alpha^\beta}{n[p^\beta + \alpha^\beta]} \end{aligned}$$

Dixit, et al.(1996), assume that a set of random variables $(X_1, X_2, \dots, X_n)$ represent the distance of an infected sampled plant from a plant from a plot of plants inoculated with a virus. Some of the observations are derived from the airborne dispersal of the spores and are distributed according to the exponential distribution. The other observations out of n random variables (say k) are present because aphids which are known to be carriers of barley yellow mosaic dwarf virus (BYMDV) have passed the virus into the plants when the aphids feed on the sap.

Dixit and Nasiri (2001) considered estimation of parameters of the exponential distribution in the presence of outliers generated from uniform distribution. Deiri (2011c) presented the estimation of parameters of the gamma distribution with presence of two outliers, and presented (2011d) estimation of the parameters of minimax distribution with presence of two outliers.

Recently Nasiri and Pazira (2010) considered the problem of estimating of $R = P(Y < X)$, where Y and X have exponential distribution with presence of k outliers, Deiri (2011a) presented the same estimation of $P(Y < X)$ of the same distribution but in the presence of two outliers, then in the same year (2011e), he presented the estimation of $P(Y < X)$ but where Y and X has the rayleigh distribution with presence of two outliers, and he presented the estimation of $P(Y < X)$ for generalized exponential distribution in the presence of two outliers when scale parameter is known, also, Jafari (2011) presented the estimation of $P(Y < X)$ in the rayleigh distribution but with presence of one outlier, and Ghanizadeh (2011) presented the estimation of $P(Y < X)$ in the rayleigh distribution with presence of k outliers.

An outlier is a specific value that stands far away from the mean or from the rest of the values in a given set. When graphed, the outlier is visually far away from all the rest of the points, Barnett and Lewis (1994).

In this paper, we obtain the moment, maximum likelihood and mixture estimators of R in Weibull distribution with presence of k outliers generated from the same distribution and we compare them together.

The Weibull distributions with shape parameter β and scale parameter p will be denoted by WE $(\beta; p)$ and the corresponding density function $f(x; \beta, p)$, for $\beta > 0$ and $p > 0$, is as follows;

$$f(x; \beta, p) = \frac{\beta}{p} \left(\frac{x}{p}\right)^{\beta-1} e^{-\left(\frac{x}{p}\right)^{\beta}} \quad x > 0$$

with c.d.f $F(x; \beta, p)$

$$F(x; \beta, p) = 1 - e^{-\left(\frac{x}{p}\right)^{\beta}} \quad \text{for } x \geq 0, \text{ and } F(x; k; \lambda) = 0 \text{ for } x < 0$$

Thus, we assume that the random variables $(X_1, X_2, \dots, X_n)$ are such that k of them are distributed with p.d.f $f_1(x; \beta, p)$,

$$f_1(x; p, \beta) = \frac{\beta}{p} \left(\frac{x}{p}\right)^{\beta-1} e^{-\left(\frac{x}{p}\right)^{\beta}} \quad x > 0, \beta > 0 \text{ and } p > 0$$

and the remaining $(n-k)$ random variables are distributed with p.d.f $f_2(x; \gamma; p)$,

Estimation of $P(Y < X)$ for Weibull Distribution In the Presence of k Outliers

Elsayed, A. Elsherpieny*

Rania. M. Shalaby**

Abstract: In this paper, we consider the problem of estimating $R=P(Y<X)$, where Y has Weibull distribution with parameters β and α and X has Weibull distribution with parameters β and p in presence of k outliers, such that X and Y are independent. The moment, maximum likelihood and mixture estimators of R are derived with help of numerical technique. At the end, we conclude that mixture estimators are better than the maximum likelihood and moment estimators. Some of the previous results in the literatures can be achieved as special case of our results.

Key words: Weibull distribution, the maximum likelihood estimator, moment estimator, mixture estimator and outliers.

1. Introduction

There has been continuous interest in the problem of estimating the probability that one random variable exceeds another, that is the quantity $R=P(Y<X)$ in reliability context, which is known as stress-strength model reliability, it describes life of a component which has a random strength X , and is subjected to random stress Y . This problem arises in the classical stress-strength reliability where one is interested in estimating the proportion of the times the random strength X of a component exceeds the random stress Y to which the component is subjected. If $X \leq Y$, then either the component fails or the system that uses the component may malfunction, and there is no failure when $Y < X$.

This problem also arises in situations where X and Y represent lifetimes of two devices and one wants to estimate the probability that one fails before the other.

The germ of this idea was introduced by Birnbaum (1956) and developed by Birnbaum and McCarty (1958). The latter paper does for the first time include $P(Y < X)$ in its title [A distribution-free upper confidence bound for $P\{Y < X\}$, based on independent samples of X and Y], but the formal term "stress - strength" appeared in the title [The estimation of reliability from stress-strength relationships] of Church and Harris (1970). The component fails at the instant that stress applied to it exceeds the strength and the component will function satisfactorily whenever $Y > X$. The estimation of R is very common in the statistical literature, for example in mechanical reliability of a system if X is the strength of a component which is subject to stress Y , where X and Y independent distribution, then R is a measure of system performance. The system fails, if at any time the applied stress is greater than its strength Kotz, et al.(2003).

McCool (1991) presented inference on $P(Y < X)$ in the Weibull case, Kundu and Gupta (2006) presented the estimation of $R = P(Y < X)$ for Weibull distribution. Davarzani (2009) presented estimation of $P(Y < X)$ where X and Y are random variables from a bivariate Weibull distribution and X is censored at Y .

*Institute of Statistical Studies & Research, Cairo University.
Ahmedc55@yahoo.com

**The High Institute of Managerial Science, Culture and Science City
ronashalabi@yahoo.com

Table 2-c: Confidence bounds of the estimates at confidence level based on type II censoring data 95% and 99% via EM-algorithm.

(n)	(r)	θ	Case 1 (0.7, 2.5, 0.6)			Case 2 (0.8, 2.2, 0.5)		
			(L,U)bound 1	(L,U)bound 2	SD	(L,U)bound 1	(L,U)bound 2	SD
(15)	60% <i>n</i>	α	(0.5853,0.7471)	(0.5597,0.7727)	0.0413	(0.5686,0.9526)	(0.5079,1.0133)	0.0980
		β	(1.7471,3.0836)	(1.5358,3.2950)	0.3409	(1.4459,2.6789)	(1.2509,2.8739)	0.3145
		k	(0.4703,0.6288)	(0.4453,0.6538)	0.0404	(0.4310,0.6147)	(0.4019,0.6438)	0.0469
		θ						
(25)		α	(0.4995,0.8166)	(0.4493,0.8667)	0.0809	(0.6334,1.0214)	(0.5720,1.0828)	0.0990
		β	(1.9201,3.0832)	(1.7362,3.2672)	0.2967	(1.6695,2.6125)	(1.5203,2.7617)	0.2406
		k	(0.4897,0.8045)	(0.4399,0.8543)	0.0803	(0.3555,0.6701)	(0.3057,0.7199)	0.0803
		θ						
(35)		α	(0.4627,0.8785)	(0.3969,0.9442)	0.1061	(0.6777,0.9575)	(0.6335,1.0017)	0.0714
		β	(1.7481,2.9456)	(1.5587,3.1350)	0.3055	(1.6238,2.2831)	(1.5196,2.3873)	0.1682
		k	(0.4697,0.7451)	(0.4261,0.7887)	0.0703	(0.3835,0.6130)	(0.3472,0.6493)	0.0586
		θ						
(40)		α	(0.5045,0.6737)	(0.4778,0.7004)	0.0432	(0.6585,0.9393)	(0.6141,0.9837)	0.0716
		β	(2.0163,3.1075)	(1.8437,3.2801)	0.2784	(1.5045,2.4383)	(1.3420,2.6439)	0.2509
		k	(0.5044,0.7404)	(0.4671,0.7777)	0.0602	(0.3686,0.5810)	(0.3350,0.6146)	0.0542
		θ						
(15)	70% <i>n</i>	α	(0.5143,0.7418)	(0.4784,0.7778)	0.0580	(0.5594,0.9414)	(0.4990,1.0018)	0.0974
		β	(1.8189,3.2025)	(1.6001,3.4213)	0.3530	(1.4791,2.7707)	(1.2748,2.9750)	0.3295
		k	(0.4670,0.7659)	(0.4197,0.8132)	0.0763	(0.3267,0.6447)	(0.2764,0.6950)	0.0811
		θ						
(25)		α	(0.5037,0.8416)	(0.4502,0.8951)	0.0862	(0.5871,1.0090)	(0.5204,1.0758)	0.1076
		β	(2.0644,3.0816)	(1.9035,3.2425)	0.2595	(1.7383,2.5980)	(1.6023,2.7340)	0.2193
		k	(0.4573,0.7094)	(0.4175,0.7493)	0.0643	(0.4342,0.6002)	(0.4080,0.6264)	0.0423
		θ						
(35)		α	(0.6325,0.8952)	(0.5909,0.9368)	0.0670	(0.6502,0.9440)	(0.6037,0.9904)	0.0749
		β	(1.8727,3.4851)	(1.6176,3.7401)	0.4113	(1.8142,2.5551)	(1.6970,2.6723)	0.1890
		k	(0.4419,0.7675)	(0.3904,0.8190)	0.0831	(0.4245,0.5609)	(0.4029,0.5824)	0.0348
		θ						
(40)		α	(0.5377,0.8390)	(0.4900,0.8866)	0.0769	(0.6915,0.9792)	(0.6460,1.0247)	0.0734
		β	(2.0451,2.9264)	(1.9057,3.0658)	0.2248	(1.5363,2.6765)	(1.3560,2.8568)	0.2909
		k	(0.4874,0.7342)	(0.4483,0.7732)	0.0630	(0.3744,0.6191)	(0.3357,0.6579)	0.0624
		θ						
(15)	90% <i>n</i>	α	(0.5005,0.7861)	(0.4553,0.8312)	0.0729	(0.5407,0.9276)	(0.4795,0.9888)	0.0987
		β	(2.2267,2.8206)	(2.1328,2.9145)	0.1515	(1.6347,2.7174)	(1.4635,2.8886)	0.2762
		k	(0.5295,0.7741)	(0.4908,0.8128)	0.0624	(0.3614,0.6442)	(0.3167,0.6890)	0.0721
		θ						
(25)		α	(0.5535,0.8189)	(0.5115,0.8608)	0.0677	(0.5577,0.9790)	(0.4911,1.0456)	0.1075
		β	(1.7803,2.7833)	(1.6217,2.9419)	0.2559	(1.6148,2.6090)	(1.4576,2.7662)	0.2536
		k	(0.4384,0.7653)	(0.3867,0.8170)	0.0834	(0.4960,0.5911)	(0.4810,0.6061)	0.0243
		θ						
(35)		α	(0.5845,0.8682)	(0.5396,0.9131)	0.0724	(0.5000,0.9882)	(0.4227,1.0654)	0.1245
		β	(2.4929,2.7724)	(2.4487,2.8166)	0.0713	(1.4502,2.6318)	(1.2634,2.8187)	0.3014
		k	(0.5115,0.6395)	(0.4913,0.6597)	0.0326	(0.5255,0.6273)	(0.5094,0.6433)	0.0260
		θ						
(40)		α	(0.5060,0.7890)	(0.4613,0.8338)	0.0722	(0.7114,0.9719)	(0.6702,1.0131)	0.0665
		β	(1.6615,2.9045)	(1.4649,3.1011)	0.3171	(1.6915,3.0861)	(1.4709,3.3067)	0.3558
		k	(0.4826,0.6943)	(0.4492,0.7277)	0.0540	(0.3828,0.6046)	(0.3477,0.6397)	0.0566
		θ						

Table 1-c: Confidence bounds of the estimates at confidence level based on type II censoring data 95% and 99% via ML method.

(n)	(r)	θ	Case 1 (0.7, 2.5, 0.6)			Case 2 (0.8, 2.2, 0.5)		
			(L,U)bound 1	(L,U)bound 2	SD	(L,U)bound 1	(L,U)bound 2	SD
(15)	60% <i>n</i>	α	(0.5671,0.8372)	(0.5244,0.8799)	0.0689	(0.5625,0.9877)	(0.4952,1.0550)	0.1085
		β	(1.7139,3.3722)	(1.4516,3.6345)	0.4231	(1.9629,2.5071)	(1.8768,2.5932)	0.1388
		k	(0.4002,0.6003)	(0.3686,0.6319)	0.0510	(0.3613,0.6117)	(0.3216,0.6513)	0.0639
(25)		α	(0.5570,0.8800)	(0.5060,0.9311)	0.0824	(0.4907,0.9466)	(0.4186,1.0187)	0.1163
		β	(1.8025,3.0192)	(1.6100,3.2117)	0.3104	(1.6060,2.5858)	(1.4510,2.7407)	0.2500
		k	(0.4131,0.6074)	(0.3824,0.6381)	0.0496	(0.4851,0.6311)	(0.4620,0.6542)	0.0373
(35)		α	(0.4264,0.9494)	(0.3437,1.0321)	0.1334	(0.6477,0.8129)	(0.6216,0.8391)	0.0422
		β	(1.9340,3.2725)	(1.7223,3.4842)	0.3415	(1.5470,2.8318)	(1.3438,3.0350)	0.3278
		k	(0.4572,0.7877)	(0.4049,0.8399)	0.0843	(0.3861,0.6233)	(0.3485,0.6609)	0.0605
(40)		α	(0.4434,0.9258)	(0.3671,1.0021)	0.1231	(0.5339,0.9794)	(0.4635,1.0499)	0.1137
		β	(2.0437,3.0208)	(1.8892,3.1754)	0.2493	(1.7598,2.4798)	(1.6459,2.5937)	0.1837
		k	(0.4219,0.7631)	(0.3679,0.8171)	0.0871	(0.4930,0.5877)	(0.4780,0.6027)	0.0242
(15)	70% <i>n</i>	α	(0.5284,0.9494)	(0.4618,1.0160)	0.1074	(0.5374,0.9809)	(0.4673,1.0511)	0.1131
		β	(2.3250,3.3115)	(2.1690,3.4675)	0.2517	(1.5817,2.8284)	(1.3845,3.0256)	0.3181
		k	(0.3521,0.7376)	(0.2911,0.7986)	0.0984	(0.3574,0.5665)	(0.3243,0.5995)	0.0533
(25)		α	(0.5130,0.8887)	(0.4536,0.9481)	0.0958	(0.5985,1.0079)	(0.5338,1.0727)	0.1045
		β	(1.6600,3.1834)	(1.4191,3.4243)	0.3886	(1.4741,2.8467)	(1.2570,3.0638)	0.3502
		k	(0.4686,0.7043)	(0.4313,0.7416)	0.0601	(0.3360,0.6564)	(0.2853,0.7070)	0.0817
(35)		α	(0.5111,0.7814)	(0.4683,0.8241)	0.0690	(0.6763,0.9248)	(0.6370,0.9641)	0.0634
		β	(1.6877,3.0635)	(1.4701,3.2811)	0.3510	(1.8927,2.6932)	(1.7661,2.8198)	0.2042
		k	(0.4708,0.7802)	(0.4219,0.8292)	0.0789	(0.4216,0.5882)	(0.3953,0.6146)	0.0425
(40)		α	(0.4644,0.9407)	(0.3891,1.0160)	0.1215	(0.5773,1.0483)	(0.5028,1.1228)	0.1202
		β	(2.1289,3.2230)	(1.9559,3.3961)	0.2791	(2.0161,2.2254)	(1.9830,2.2585)	0.0534
		k	(0.4478,0.7927)	(0.3932,0.8473)	0.0880	(0.3654,0.6634)	(0.3182,0.7106)	0.0760
(15)	90% <i>n</i>	α	(0.4318,0.9583)	(0.3485,1.0415)	0.1343	(0.6244,1.0188)	(0.5620,1.0812)	0.1006
		β	(1.9066,3.0175)	(1.7308,3.1932)	0.2834	(1.7665,2.3192)	(1.6791,2.4066)	0.1410
		k	(0.4473,0.5177)	(0.4362,0.5288)	0.0180	(0.3183,0.6096)	(0.2722,0.6557)	0.0743
(25)		α	(0.4912,0.8334)	(0.4370,0.8876)	0.0873	(0.5584,0.9641)	(0.4942,1.0283)	0.1035
		β	(1.9227,3.1067)	(1.7353,3.2940)	0.3021	(1.4723,2.8302)	(1.2576,3.0450)	0.3464
		k	(0.4991,0.7225)	(0.4638,0.7578)	0.0570	(0.3413,0.6595)	(0.2910,0.7098)	0.0812
(35)		α	(0.4339,0.8737)	(0.3643,0.9433)	0.1122	(0.5983,0.9280)	(0.5461,0.9801)	0.0841
		β	(2.0971,2.9794)	(1.9575,3.1189)	0.2251	(1.5684,2.8427)	(1.3668,3.0442)	0.3251
		k	(0.4167,0.7364)	(0.3662,0.7869)	0.0815	(0.4423,0.5674)	(0.4225,0.5872)	0.0319
(40)		α	(0.4348,0.9456)	(0.3540,1.0264)	0.1303	(0.7046,1.0299)	(0.6531,1.0814)	0.0830
		β	(1.4087,3.1467)	(1.1338,3.4215)	0.4434	(1.6146,2.7889)	(1.4288,2.9747)	0.2996
		k	(0.4385,0.7920)	(0.3826,0.8480)	0.0902	(0.3822,0.6787)	(0.3353,0.7256)	0.0756

Table 2-b: Asymptotic variance and covariance of estimates based on type II censoring data via EM-algorithm.

(n)	(r)		Case 1	(0.7, 2.5, 0.6)		Case 2	(0.8, 2.2, 0.5)	
			$\hat{\alpha}$	$\hat{\beta}$	$\hat{k}$	$\hat{\alpha}$	$\hat{\beta}$	$\hat{k}$
(15)	60% <i>n</i>	$\hat{\alpha}$	0.0475	-0.0115	-0.1568	0.0440	-0.0152	-0.1283
		$\hat{\beta}$	-0.0115	0.0501	0.0491	-0.0152	0.0672	0.0620
		$\hat{k}$	-0.1568	0.0491	1.9230	-0.1283	0.0620	1.3332
(25)		$\hat{\alpha}$	0.0357	-0.0069	-0.1137	0.0270	-0.0118	-0.0888
		$\hat{\beta}$	-0.0069	0.0273	0.0266	-0.0118	0.0489	0.0520
		$\hat{k}$	-0.1137	0.0266	1.9100	-0.0888	0.0520	0.9268
(35)		$\hat{\alpha}$	0.0249	-0.0051	-0.0812	0.0180	-0.0080	-0.0530
		$\hat{\beta}$	-0.0051	0.0205	0.0211	-0.0080	0.0347	0.0337
		$\hat{k}$	-0.0812	0.0211	0.9812	-0.0530	0.0337	0.5037
(40)		$\hat{\alpha}$	0.0208	-0.0034	-0.0658	0.0152	-0.0071	-0.0490
		$\hat{\beta}$	-0.0034	0.0111	0.0134	-0.0071	0.0306	0.0320
		$\hat{k}$	-0.0658	0.0134	1.0252	-0.0490	0.0320	0.4505
(15)	70% <i>n</i>	$\hat{\alpha}$	0.0614	-0.0120	-0.2074	0.0420	-0.0159	-0.1382
		$\hat{\beta}$	-0.0120	0.0430	0.0505	-0.0159	0.0721	0.0716
		$\hat{k}$	-0.2074	1.6473	0.0505	-0.1382	0.0716	0.9509
(25)		$\hat{\alpha}$	0.0320	-0.0079	-0.1113	0.0257	-0.0106	-0.0801
		$\hat{\beta}$	-0.0079	0.0304	0.0356	-0.0106	0.0450	0.0457
		$\hat{k}$	-0.1113	0.0356	0.9834	-0.0801	0.0457	0.6207
(35)		$\hat{\alpha}$	0.0228	-0.0065	-0.0770	0.0173	-0.0073	-0.0563
		$\hat{\beta}$	-0.0065	0.0260	0.0281	-0.0073	0.0335	0.0338
		$\hat{k}$	-0.0770	0.0281	0.7823	-0.0563	0.0338	0.4353
(40)		$\hat{\alpha}$	0.0199	-0.0048	-0.0611	0.0156	-0.0074	-0.0494
		$\hat{\beta}$	-0.0048	0.0189	0.0195	-0.0074	0.0316	0.0325
		$\hat{k}$	-0.0611	0.0195	0.6080	-0.0494	0.0325	0.3810
(15)	90% <i>n</i>	$\hat{\alpha}$	0.0597	-0.0103	-0.1788	0.0451	-0.0159	-0.1570
		$\hat{\beta}$	-0.0103	0.0429	0.0411	-0.0159	0.0653	0.0718
		$\hat{k}$	-0.1788	0.0411	1.1011	-0.1570	0.0718	0.9174
(25)		$\hat{\alpha}$	0.0328	-0.0079	-0.0975	0.0273	-0.0099	-0.0816
		$\hat{\beta}$	-0.0079	0.0311	0.0322	-0.0099	0.0405	0.0403
		$\hat{k}$	-0.0975	0.0322	0.5628	-0.0816	0.0403	0.4670
(35)		$\hat{\alpha}$	0.0210	-0.0058	-0.0737	0.0225	-0.0071	-0.0632
		$\hat{\beta}$	-0.0058	0.0252	0.0272	-0.0071	0.0265	0.0274
		$\hat{k}$	-0.0737	0.0272	0.5181	-0.0632	0.0274	0.3334
(40)		$\hat{\alpha}$	0.0199	-0.0043	-0.0603	0.0157	-0.0076	-0.0605
		$\hat{\beta}$	-0.0043	0.0173	0.0179	-0.0076	0.0322	0.0382
		$\hat{k}$	-0.0603	0.0179	0.3561	-0.0605	0.0382	0.4153

Table I-b: Asymptotic variance and covariance of estimates based on type II censoring data via ML method.

(n)	(r)		Case 1 (0.7, 2.5, 0.5)			Case 2 (0.8, 2.2, 0.5)		
		$\hat{\theta}$	$\hat{\alpha}$	$\hat{\beta}$	$\hat{k}$	$\hat{\alpha}$	$\hat{\beta}$	$\hat{k}$
(15)	60% <i>n</i>	$\hat{\alpha}$	0.0416	-0.0139	-0.1598	0.0412	-0.0167	-0.1436
		$\hat{\beta}$	-0.0139	0.0607	0.0665	-0.0167	0.0748	0.0753
		$\hat{k}$	-0.1598	0.0665	2.0334	-0.1436	0.0753	1.5193
		$\hat{\alpha}$	0.0252	-0.0091	-0.0890	0.0279	-0.0082	-0.0764
(25)		$\hat{\beta}$	-0.0091	0.0377	0.0416	-0.0082	0.0349	0.0307
		$\hat{k}$	-0.0890	0.0416	1.1519	-0.0764	0.0307	0.9405
		$\hat{\alpha}$	0.0243	-0.0055	-0.0834	0.0182	-0.0065	-0.0596
(35)		$\hat{\beta}$	-0.0055	0.0211	0.0231	-0.0065	0.0280	0.0290
		$\hat{k}$	-0.0834	0.0231	1.3609	-0.0596	0.0290	0.6375
		$\hat{\alpha}$	0.0201	-0.0049	-0.0692	0.0168	-0.0059	-0.0489
(40)		$\hat{\beta}$	-0.0049	0.0193	0.0216	-0.0059	0.0244	0.0244
		$\hat{k}$	-0.0692	0.0216	0.9787	-0.0489	0.0244	0.5987
		$\hat{\alpha}$	0.0513	-0.0166	-0.2156	0.0385	-0.0160	-0.1365
(15)	70% <i>n</i>	$\hat{\beta}$	-0.0166	0.0642	0.0862	-0.0160	0.0759	0.0757
		$\hat{k}$	-0.2156	0.0862	1.9801	-0.1365	0.0757	1.0004
		$\hat{\alpha}$	0.0318	-0.0084	-0.1021	0.0276	-0.0125	-0.0968
(25)		$\hat{\beta}$	-0.0084	0.0321	0.0365	-0.0125	0.0490	0.0575
		$\hat{k}$	-0.1021	0.0365	0.8602	-0.0968	0.0575	0.6932
		$\hat{\alpha}$	0.0263	-0.0051	-0.0812	0.0182	-0.0079	-0.0626
(35)		$\hat{\beta}$	-0.0051	0.0189	0.0214	-0.0079	0.0330	0.0370
		$\hat{k}$	-0.0812	0.0214	0.6360	-0.0626	0.0370	0.5167
		$\hat{\alpha}$	0.0208	-0.0051	-0.0694	0.0165	-0.0073	-0.0523
(40)		$\hat{\beta}$	-0.0051	0.0196	0.0219	-0.0073	0.0293	0.0315
		$\hat{k}$	-0.0694	0.0219	0.7373	-0.0523	0.0315	0.4089
		$\hat{\alpha}$	0.0418	-0.0143	-0.1713	0.0402	-0.0205	-0.1419
(15)	90% <i>n</i>	$\hat{\beta}$	-0.0143	0.0633	0.0761	-0.0205	0.0887	0.0953
		$\hat{k}$	-0.1713	0.0761	1.1734	-0.1419	0.0953	0.8164
		$\hat{\alpha}$	0.0325	-0.0075	-0.1050	0.0268	-0.0108	-0.0920
(25)		$\hat{\beta}$	-0.0075	0.0284	0.0324	-0.0108	0.0440	0.0485
		$\hat{k}$	-0.1050	0.0324	0.6851	-0.0920	0.0485	0.5356
		$\hat{\alpha}$	0.0224	-0.0052	-0.0779	0.0201	-0.0077	-0.0711
(35)		$\hat{\beta}$	-0.0052	0.0207	0.0240	-0.0077	0.0308	0.0356
		$\hat{k}$	-0.0779	0.0240	0.5104	-0.0711	0.0356	0.4287
		$\hat{\alpha}$	0.0214	-0.0051	-0.0625	0.0178	-0.0085	-0.0591
(40)		$\hat{\beta}$	-0.0051	0.0190	0.0206	-0.0085	0.0321	0.0373
		$\hat{k}$	-0.0625	0.0206	0.3640	-0.0591	0.0373	0.3682

Table 2-a: The RABias, MSE and RE of the parameter $\theta = (\alpha, \beta, k)$ based on type II censoring from Burr XII distribution via EM-algorithm.

(n)	(r)		Case 1 (0.7, 2.5, 0.6)				Case 2 (0.8, 2.2, 0.5)		
		$\hat{\theta}$	RAbias	MSE	RE	$\hat{\theta}$	RAbias	MSE	RE
(15)	60% <i>n</i>	0.6662	0.0483	0.0029	0.0762	0.7606	0.0493	0.0112	0.1320
		2.4154	0.0339	0.1234	0.1405	2.0624	0.0626	0.1179	0.1561
		0.5495	0.0841	0.0042	0.1078	0.5229	0.0457	0.0027	0.1043
(25)		0.6580	0.0600	0.0083	0.1302	0.8274	0.0343	0.0106	0.1284
		2.5017	0.0007	0.0880	0.1187	2.1410	0.0268	0.0614	0.1126
		0.6471	0.0785	0.0087	0.1552	0.5128	0.0256	0.0066	0.1626
(35)		0.6706	0.0421	0.0121	0.1573	0.8176	0.0220	0.0054	0.0919
		2.3469	0.0613	0.1168	0.1367	1.9535	0.1121	0.0891	0.1357
		0.6074	0.0124	0.0050	0.1178	0.4983	0.0035	0.0034	0.1172
(40)		0.5891	0.1584	0.0142	0.1700	0.7989	0.0014	0.0051	0.0896
		2.5619	0.0248	0.0813	0.1141	1.9935	0.0925	0.1044	0.1469
		0.6224	0.0374	0.0041	0.1071	0.4748	0.0504	0.0036	0.1195
(15)	70% <i>n</i>	0.6281	0.1028	0.0085	0.1320	0.7504	0.0620	0.0120	0.1367
		2.5107	0.0043	0.1247	0.1413	2.1249	0.0341	0.1142	0.1536
		0.6165	0.0274	0.0061	0.1300	0.4857	0.0286	0.0068	0.1647
(25)		0.6727	0.0391	0.0082	0.1292	0.7981	0.0024	0.0116	0.1346
		2.5730	0.0292	0.0727	0.1078	2.1681	0.0145	0.0491	0.1007
		0.5834	0.0277	0.0044	0.1107	0.5172	0.0344	0.0021	0.0914
(35)		0.7638	0.0912	0.0086	0.1322	0.7971	0.0037	0.0056	0.0938
		2.6789	0.0716	0.2012	0.1794	2.1846	0.0070	0.0350	0.0862
		0.6047	0.0078	0.0069	0.1386	0.4927	0.0147	0.0013	0.0711
(40)		0.6883	0.0167	0.0060	0.1111	0.8353	0.0442	0.0066	0.1018
		2.4858	0.0057	0.0507	0.0901	2.1064	0.0426	0.0934	0.1389
		0.6108	0.0180	0.0041	0.1065	0.4968	0.0065	0.0039	0.1250
(15)	90% <i>n</i>	0.6433	0.0811	0.0085	0.1319	0.7342	0.0823	0.0141	0.1483
		2.5237	0.0095	0.0235	0.0613	2.1761	0.0109	0.0769	0.1260
		0.6518	0.0864	0.0066	0.1352	0.5029	0.0057	0.0052	0.1444
(25)		0.686159	0.0198	0.0048	0.0987	0.7683	0.0396	0.0126	0.1400
		2.2818	0.0873	0.1131	0.1345	2.1119	0.0401	0.0721	0.1220
		0.6018	0.0030	0.007	0.1390	0.5436	0.0871	0.0025	0.0997
(35)		0.7263	0.0376	0.0059	0.1100	0.7441	0.0699	0.0186	0.1707
		2.6326	0.0531	0.0227	0.0602	2.0410	0.0723	0.1161	0.1549
		0.5755	0.0409	0.0017	0.0680	0.5764	0.1528	0.0065	0.1613
(40)		0.6475	0.0750	0.0080	0.1275	0.8416	0.0520	0.0062	0.0980
		2.2830	0.0868	0.1477	0.1537	2.3886	0.0858	0.1622	0.1831
		0.5884	0.0193	0.0031	0.0920	0.4937	0.0126	0.0032	0.1139

Table 1-a: The RABias, MSE and RE of the parameter $\theta = (\alpha, \beta, k)$ based on type II censoring from Burr XII distribution via ML method.

(n)	(r)	Case 1 (0.7, 2.5, 0.6)					Case 2 (0.8, 2.2, 0.5)			
		$\hat{\theta}$	$\hat{\theta}$	RABias	MSE	RE	$\hat{\theta}$	RABias	MSE	RE
(15)	60% <i>n</i>	α	0.7021	0.0030	0.0048	0.0985	0.7751	0.0311	0.0124	0.1391
		β	2.5435	0.0172	0.1808	0.1701	2.2350	0.0159	0.0205	0.0651
		k	0.5005	0.1663	0.0126	0.1868	0.4865	0.0271	0.0043	0.1306
(25)		α	0.7185	0.0265	0.0071	0.1206	0.7187	0.1017	0.0201	0.1774
		β	2.4109	0.0357	0.1043	0.1292	2.0959	0.0473	0.0733	0.1231
		k	0.5102	0.1496	0.0105	0.1709	0.5581	0.1162	0.0048	0.1381
(35)		α	0.6879	0.0173	0.0180	0.1914	0.7303	0.0871	0.0066	0.1018
		β	2.6033	0.0413	0.1273	0.1427	2.1894	0.0048	0.1075	0.1491
		k	0.6224	0.0374	0.0076	0.1454	0.5047	0.0094	0.0037	0.1214
(40)		α	0.6846	0.0220	0.0154	0.1772	0.7567	0.0542	0.0148	0.1520
		β	2.5323	0.0129	0.0632	0.1005	2.1198	0.0365	0.0402	0.0911
		k	0.5925	0.0126	0.0076	0.1456	0.5403	0.0807	0.0022	0.0940
(15)	70% <i>n</i>	α	0.7389	0.0556	0.0131	0.1632	0.7592	0.0510	0.0145	0.1503
		β	2.8183	0.1273	0.1646	0.1623	2.2051	0.0023	0.1012	0.1446
		k	0.5448	0.0919	0.0127	0.1879	0.4619	0.0762	0.0043	0.1311
(25)		α	0.7008	0.0012	0.0092	0.1369	0.8032	0.0040	0.0109	0.1306
		β	2.4217	0.0313	0.1572	0.1586	2.1604	0.0180	0.1242	0.1602
		k	0.5864	0.0226	0.0038	0.1027	0.4962	0.0077	0.0067	0.1637
(35)		α	0.6462	0.0768	0.0077	0.1249	0.8006	0.0007	0.0040	0.0792
		β	2.3756	0.0498	0.1387	0.1490	2.2930	0.0423	0.0503	0.1020
		k	0.6255	0.0426	0.0069	0.1383	0.5049	0.0098	0.0018	0.0856
(40)		α	0.7026	0.0036	0.0148	0.1736	0.8128	0.0160	0.0146	0.1510
		β	2.6760	0.0704	0.1089	0.1320	2.1208	0.0360	0.0091	0.0434
		k	0.6202	0.0337	0.0082	0.1505	0.5144	0.0288	0.0060	0.1548
(15)	90% <i>n</i>	α	0.6950	0.0071	0.0181	0.1920	0.8216	0.0270	0.0106	0.1286
		β	2.4620	0.0152	0.0818	0.1144	2.0428	0.0715	0.0446	0.0960
		k	0.4825	0.1958	0.0141	0.1981	0.4640	0.0721	0.0068	0.1652
(25)		α	0.6623	0.0539	0.0090	0.1359	0.7612	0.0485	0.0122	0.1382
		β	2.5146	0.0059	0.0915	0.1210	2.1513	0.0222	0.1224	0.1590
		k	0.6108	0.0180	0.0034	0.0967	0.5004	0.0008	0.0066	0.1624
(35)		α	0.6538	0.0660	0.0147	0.1734	0.7631	0.0461	0.0084	0.1148
		β	2.5382	0.0153	0.0521	0.0913	2.2055	0.0025	0.1057	0.1478
		k	0.5765	0.0391	0.0072	0.1414	0.5048	0.0097	0.0010	0.0646
(40)		α	0.6902	0.0140	0.0171	0.1867	0.8672	0.0840	0.0114	0.1335
		β	2.2777	0.0889	0.2460	0.1984	2.2017	0.0008	0.0898	0.1362
		k	0.6153	0.0254	0.0084	0.1525	0.5305	0.0610	0.0067	0.1631

- [13] Park, S. (2004), optimal design of step-stress degradation tests in the case of destructive measurement, *Quality Technology and Quantitative Management*, 1, 105-124.
- [14] Wang, F.K., Cheng, Y.F. (2010), EM algorithm for estimating the Burr XII parameters with multiple censored data, *Quality and Reliability Engineering International* 26, 615-630.
- [15] Watkins, A. J. (2001), Commentary: inference in simple step-stress models, *IEEE Trans. Reliab.*, 50, 36-27.
- [16] Xiong, C. (1998), Inference on simple step-stress model with type-II censored exponential data, *IEEE Trans. Reliab.*, 47, 142-146.

(0.8,2.2,0.5) respectively, for deferent (n, r) , we take $r = 60\%n, 70\%n, 90\%n$.

References

- [1] Al-Masri, A. (2009), Optimum times for step-stress cumulaive exposure model using Log-Logistic distribution with Known scale parameter, Austrian Journal of Statistics, 38, 59-66.
- [2] Bai, D. S., Kim, M. S. and Lee, S. H. (1989), Optimum simple step-stress accelerated life tests with censoring, IEEE Trans. Reliab., 38, 528-532.
- [3] Balakrishnan, N. (1996), A bayes model for step-stress accelerated life testing, IEEE Trans. Reliab., 45, 491-498.
- [4] Balakrishnan, N. (2007), Point and interval estimation for a simple step-stress model with type-II censoring, IEEE Trans. Reliab., 47, 142-146.
- [5] Balakrishnan, N., Kim J.A. (2004), EM algorithm for Type-II right censored bivariate normal data, in: M.S. Nikulin, N. Balakrishnan, M. Mesbah, N. Limnios (Eds.). Advances in Parametric and Semi-Parametric Inference with Applications in Reliability, Survival Analysis Quality of Life. Birkhauser, Boston, pp. 177-210.
- [6] Cohen, A. C. (1965), Maximum likelihood estimation in the Weibull distribution based on complete and on censored samples, Technometrics, 5, 579-588.
- [7] Dempster, A.P, Laird N.M , Rubin, D.B (1977), Maximum likelihood from incomplete data via the EM algorithm, Journal of Royal Statistics Society - Serial B Methodology 39, 1-38.
- [8] Jalali, A. (2009), On maximum likelihood estimation for the two parameter Burr XII distribution, Communication in Statistics-Theory and Methods, 38, 1916-1926.
- [9] Louis, T.A. (1982), Finding the observed information matrix when using the EM Algorithm, Journal of Royal Statistics Society - Serial B Methodology 44, 226-233.
- [10] McCool, J. I. (1980), Confidance limits for weibull regression with censored data, IEEE Trans. Reliab., R-29, 145-150.
- [11] Miller, R and Nelson, W. (1983), Optimum simple step-stress plans for accelerated life testing, IEEE Trans. Reliab., 32, 59-65.
- [12] Nelson, W. (1980), Accelerated life testing step-stress model and data analysis, IEEE Trans. Reliab., 29, 103-108.

The confidence interval depend at $r \cong 70\%$ as follows, via ML method is

$$[L_{\alpha}, U_{\alpha}] \text{ bound } 95\% = [0.4644, 0.9407], [L_{\alpha}, U_{\alpha}] \text{ bound } 99\% = [0.3891, 1.0160].$$

$$[L_{\beta}, U_{\beta}] \text{ bound } 95\% = [2.1289, 3.2230], [L_{\beta}, U_{\beta}] \text{ bound } 99\% = [1.9559, 3.3961].$$

$$[L_k, U_k] \text{ bound } 95\% = [0.4478, 0.7927], [L_k, U_k] \text{ bound } 99\% = [0.3932, 0.8473],$$

via EM method is

$$[L_{\alpha}, U_{\alpha}] \text{ bound } 95\% = [0.5377, 0.8390], [L_{\alpha}, U_{\alpha}] \text{ bound } 99\% = [0.4900, 0.8866],$$

$$[L_{\beta}, U_{\beta}] \text{ bound } 95\% = [2.0451, 2.9264], [L_{\beta}, U_{\beta}] \text{ bound } 99\% = [1.9057, 3.0658],$$

$$[L_k, U_k] \text{ bound } 95\% = [0.4874, 0.7342], [L_k, U_k] \text{ bound } 99\% = [0.4483, 0.7732].$$

The confidence interval depend at $r \cong 90\%$ as follows, via ML method is

$$[L_{\alpha}, U_{\alpha}] \text{ bound } 95\% = [0.4348, 0.9456], [L_{\alpha}, U_{\alpha}] \text{ bound } 99\% = [0.3540, 1.0264],$$

$$[L_{\beta}, U_{\beta}] \text{ bound } 95\% = [1.4087, 3.1467], [L_{\beta}, U_{\beta}] \text{ bound } 99\% = [1.1538, 3.4215],$$

$$[L_k, U_k] \text{ bound } 95\% = [0.4385, 0.7920], [L_k, U_k] \text{ bound } 99\% = [0.3826, 0.8480].$$

via EM method is

$$[L_{\alpha}, U_{\alpha}] \text{ bound } 95\% = [0.5060, 0.7890], [L_{\alpha}, U_{\alpha}] \text{ bound } 99\% = [0.4613, 0.8338].$$

$$[L_{\beta}, U_{\beta}] \text{ bound } 95\% = [1.6615, 2.9045], [L_{\beta}, U_{\beta}] \text{ bound } 99\% = [1.4649, 3.1011].$$

$$[L_k, U_k] \text{ bound } 95\% = [0.4826, 0.6943], [L_k, U_k] \text{ bound } 99\% = [0.4492, 0.7277].$$

7 Conclusion

In this paper, we have considered a simple step-stress model with two stress levels from Burr-XII distribution when the data are type II censoring. We have derived two methods (ML, EM) of the unknown parameter $\theta = (\alpha, \beta, k)$. We have also proposed procedure for constructing the estimator $\hat{\theta}$ beside RABies, MSE and RE of it, see tables (1-a, 2-a). Moreover, the variance and covariance matrix, see tables (1-b, 2-b), and the confidence interval (CIs) for θ at 95% and 99%, see tables (1-c, 2-c). Hence, we present the results of a Monte Carlo simulation study carried out in order to compare the performance of the method of inference described in section 6. And more than, we observe the Fisher matrix of the variance and covariance from using EM-algorithm method much better than others to estimate the estimators, such that the MSE will be very small rather than others. see examples (1-2). Since, we chose the values of parameter θ to be (0.7,2.5,0.6) and

Example 2: We consider the following data when $n = 40$ with $\tau_1 = 3$, $\alpha = 0.7$, $\beta = 2.5$ and $k = 0.6$. Therefore, the failure times of units as follows,

$0 < t \leq \tau_1$	0.0003	0.0069	0.0153	0.0713	0.0719	0.1276	0.1706	0.2113	0.4706
	0.4721	0.5009	0.6539	1.0755	1.1365	1.3150	1.3233	1.9780	2.6145
	2.8013	2.9079							
$\tau_1 < t \leq \infty$	3.0907	3.2131	3.2662	3.4186	3.4304	3.5804	3.6448	3.6552	3.9792
	4.4885	4.6030	4.6283	4.8869	5.1161	5.3421	6.3845	9.2052	13.6004
	16.0477	19.6542							

In this case, the life-testing experiment is terminated as soon as the r^{th} failure occurs. We take $r = 70\%n$, and $90\%n$, we obtain the estimator of $\hat{\theta} = (\hat{\alpha}, \hat{\beta}, \hat{k})$ as follows,

Via ML method, $\hat{\alpha} = 0.7626$ $\hat{\beta} = 2.6760$ $\hat{k} = 0.6202$ when $r \cong 70\%n$,

$\hat{\alpha} = 0.6902$ $\hat{\beta} = 2.2777$ $\hat{k} = 0.6153$ when $r \cong 90\%n$.

Via EM method,

$\hat{\alpha} = 0.6883$ $\hat{\beta} = 2.4858$ $\hat{k} = 0.6108$ when $r \cong 70\%n$,

$\hat{\alpha} = 0.6475$ $\hat{\beta} = 2.2830$ $\hat{k} = 0.5884$ when $r \cong 90\%n$.

Note that, when $r \cong 70\%n$, the time termination $t_{r,40} = 3.6552$ and when $r \cong 90\%n$, the time termination $t_{r,40} = 6.3845$. An asymptotic variance covariance matrix at $r \cong 70\%n$ as follows, via ML method is

$$\begin{pmatrix} 0.0208 & -0.0051 & -0.0694 \\ -0.0051 & 0.0196 & 0.0219 \\ -0.0694 & 0.0219 & 0.7373 \end{pmatrix}.$$

Via EM method is

$$\begin{pmatrix} 0.0199 & -0.0048 & -0.0611 \\ -0.0048 & 0.0189 & 0.0195 \\ -0.0611 & 0.0195 & 0.6080 \end{pmatrix}.$$

An asymptotic variance covariance matrix at $r \cong 90\%n$ as follows, via ML method is

$$\begin{pmatrix} 0.0214 & -0.0051 & -0.0625 \\ -0.0051 & 0.0190 & 0.0206 \\ -0.0625 & 0.0206 & 0.3640 \end{pmatrix}.$$

Via EM method is

$$\begin{pmatrix} 0.0199 & -0.0043 & -0.0603 \\ -0.0043 & 0.0173 & 0.0179 \\ -0.0603 & 0.0179 & 0.3561 \end{pmatrix}.$$

$$\begin{pmatrix} 0.0276 & -0.0125 & -0.0968 \\ -0.0125 & 0.0490 & 0.0575 \\ -0.0968 & 0.0575 & 0.6932 \end{pmatrix}$$

via EM method is given,

$$\begin{pmatrix} 0.0257 & -0.0106 & -0.0801 \\ -0.0106 & 0.0450 & 0.0457 \\ -0.0801 & 0.0457 & 0.6207 \end{pmatrix}$$

An asymptotic variance covariance matrix at $r \cong 90\%$ as follows, via ML method is given,

$$\begin{pmatrix} 0.0268 & -0.0108 & -0.0920 \\ -0.0108 & 0.0440 & 0.0485 \\ -0.0920 & 0.0485 & 0.5356 \end{pmatrix}$$

via EM method is given,

$$\begin{pmatrix} 0.0273 & -0.0099 & -0.0816 \\ -0.0099 & 0.0405 & 0.0403 \\ -0.0816 & 0.0403 & 0.4670 \end{pmatrix}$$

The confidence interval depend at $r \cong 70\%$ as follows, via ML method is

$$[L_\alpha, U_\alpha] \text{ bound } 95\% = [0.5985, 1.0079], [L_\alpha, U_\alpha] \text{ bound } 99\% = [0.5337, 1.0727],$$

$$[L_\beta, U_\beta] \text{ bound } 95\% = [1.4741, 2.8467], [L_\beta, U_\beta] \text{ bound } 99\% = [1.2570, 3.0638],$$

$$[L_k, U_k] \text{ bound } 95\% = [0.3360, 0.6564], [L_k, U_k] \text{ bound } 99\% = [0.2853, 0.7070].$$

via EM method is

$$[L_\alpha, U_\alpha] \text{ bound } 95\% = [0.5871, 1.0090], [L_\alpha, U_\alpha] \text{ bound } 99\% = [0.5204, 1.0758],$$

$$[L_\beta, U_\beta] \text{ bound } 95\% = [1.7383, 2.5980], [L_\beta, U_\beta] \text{ bound } 99\% = [1.6023, 2.7340],$$

$$[L_k, U_k] \text{ bound } 95\% = [0.4342, 0.6002], [L_k, U_k] \text{ bound } 99\% = [0.4080, 0.6264].$$

The confidence interval depend at $r \cong 90\%$ as follows, via ML method is

$$[L_\alpha, U_\alpha] \text{ bound } 95\% = [0.5584, 0.9641], [L_\alpha, U_\alpha] \text{ bound } 99\% = [0.4942, 1.0283],$$

$$[L_\beta, U_\beta] \text{ bound } 95\% = [1.4723, 2.8302], [L_\beta, U_\beta] \text{ bound } 99\% = [1.2576, 3.0450],$$

$$[L_k, U_k] \text{ bound } 95\% = [0.3413, 0.6595], [L_k, U_k] \text{ bound } 99\% = [0.2910, 0.7098].$$

via EM method is

$$[L_\alpha, U_\alpha] \text{ bound } 95\% = [0.5577, 0.9790], [L_\alpha, U_\alpha] \text{ bound } 99\% = [0.4911, 1.0456].$$

$$[L_\beta, U_\beta] \text{ bound } 95\% = [1.6148, 2.6090], [L_\beta, U_\beta] \text{ bound } 99\% = [1.4576, 2.7662].$$

$$[L_k, U_k] \text{ bound } 95\% = [0.4960, 0.5911], [L_k, U_k] \text{ bound } 99\% = [0.4810, 0.6061].$$

and for $1 \leq j \leq r - N_1$, we set $t_{N_1+j:n} = \left[\left(1 - U_{N_1+j:n} \right)^{\frac{1}{k}} - 1 \right]^{\frac{1}{\beta}} + \tau_1 - \tau_1^{\frac{\alpha}{\beta}}$.

Step 4: Based on n, r, N_1, τ_1 , and order observations $\{t_{1:n}, \dots, t_{N_1:n}, t_{N_1+1:n}, \dots, t_{r:n}\}$

We solve the system of nonlinear equations get by differentiating Eq. (9) and Eq. (14) with respect to α, β and k respectively and equating to zero to obtain the unknown parameters $\hat{\alpha}$ and $\hat{k}$ (ML and EM methods) and using the values which gotten of estimators to get the absolute relative bias (RABias), mean square error (MSE) and relative error (RE). Furthermore, Fisher information matrix and two-sided confidence intervals of the estimators are obtained. We give illustrate examples by using small and moderately large samples to illustrate all the methods of inference described in this paper.

Example 1: We consider the following data when $n = 25$ with $\tau_1 = 3, \alpha = 0.8, \beta = 2.2$ and $k = 0.5$. Therefore, the failure times of units as follows,

$0 < t \leq \tau_1$	0.0177	0.0495	0.1267	0.2195	0.4428	0.6312	0.7750	0.8385
	1.3069	1.5092	1.5901	2.1000	2.8311			
$\tau_1 < t \leq \infty$	3.1157	3.1718	3.5581	3.7557	3.9660	4.0858	4.1505	5.9851
	8.1848	21.7891	28.1333	29.8994				

In this case, the life-testing experiment is terminated as soon as the r^{th} failure occurs. We take $r = 70\%n, 90\%n$, we obtain the estimator of $\hat{\theta} = (\hat{\alpha}, \hat{\beta}, \hat{k})$ as follows,

Via ML method, $\hat{\alpha} = 0.8032$ $\hat{\beta} = 2.1604$ $\hat{k} = 0.4962$ when $r \cong 70\%n$,

$\hat{\alpha} = 0.7612$ $\hat{\beta} = 2.1513$ $\hat{k} = 0.5004$ when $r \cong 90\%n$.

Via EM method, $\hat{\alpha} = 0.7981$ $\hat{\beta} = 2.168$ $\hat{k} = 0.5172$ when $r \cong 70\%n$,

$\hat{\alpha} = 0.7683$ $\hat{\beta} = 2.1119$ $\hat{k} = 0.5436$ when $r \cong 90\%n$.

Note that, when $r \cong 70\%n$, the time termination $t_{r,25} = 3.9660$ and when $r \cong 90\%n$, the time termination $t_{r,25} = 21.7891$. An asymptotic variance covariance matrix at $r \cong 70\%n$ as follows, via ML method is

probability δ . z is the $\left[\frac{100(1-\delta)}{2} \right]$ standard normal percentile.

Then, the approximate confidence limits for α, β and k will be constructed using Eq. (17) with confidence levels 95% and 99%, if $\underline{\theta} = (0.8, 2.2, 0.5)$ at $n=40$ and $r=28$ then we get the estimators $\hat{\underline{\theta}} = (0.8353, 2.1064, 0.4968)$. Therefore, the confidence interval depend on ML method is,

$$\begin{aligned} [L_{\alpha}, U_{\alpha}] \text{ bound } 95\% &= [0.5773, 1.0483], [L_{\alpha}, U_{\alpha}] \text{ bound } 99\% = [0.5028, 1.1228], \\ [L_{\beta}, U_{\beta}] \text{ bound } 95\% &= [2.0161, 2.2254], [L_{\beta}, U_{\beta}] \text{ bound } 99\% = [1.9830, 2.2585], \\ [L_k, U_k] \text{ bound } 95\% &= [0.3654, 0.6634], [L_k, U_k] \text{ bound } 99\% = [0.3182, 0.7106]. \end{aligned}$$

And, the confidence interval depend on EM method is

$$\begin{aligned} [L_{\alpha}, U_{\alpha}] \text{ bound } 95\% &= [0.5915, 0.9792], [L_{\alpha}, U_{\alpha}] \text{ bound } 99\% = [0.6460, 1.0247], \\ [L_{\beta}, U_{\beta}] \text{ bound } 95\% &= [1.5363, 2.6765], [L_{\beta}, U_{\beta}] \text{ bound } 99\% = [1.3560, 2.8668], \\ [L_k, U_k] \text{ bound } 95\% &= [0.3744, 0.6191], [L_k, U_k] \text{ bound } 99\% = [0.3357, 0.6579]. \end{aligned}$$

Then, the approximate confidence limits for α, β and k will be constructed using Eq. (17) with confidence levels 95% and 99%, see Tables (1-c, 2-c).

6 Numerical Illustration

In this section, we describe the algorithm to obtain the Type-II censored sample, and the estimators $\hat{\alpha}, \hat{\beta}, \hat{k}$. In addition, the absolute relative bias (RABias), mean square error (MSE) and relative error (RE).

Step 1: Given τ_1 and initial values for the parameters α, β and k .

Step 2: Based on $n, r, \tau_1, \alpha, \beta$ and k , we generated a random sample of size n from Uniform(0,1) distribution, and obtain the order statistics $U_{1:n}, \dots, U_{n:n}$.

Step 3: Find N_1 such that $U_{N_1:n} \leq 1 - (1 + \tau_1^{\alpha})^{-k} < U_{N_1+1:n}$.

We solve the above inequality to get the failure times of units from unite 1 to unite N_1 .

The failure times under Burr XII distribution of units as follows,

$$\text{for } 1 \leq i \leq N_1 \text{ we set } t_{i:n} = \left[(1 - U_{i:n})^{-\frac{1}{k}} - 1 \right]^{\frac{1}{\alpha}},$$

$$F = \begin{pmatrix} -\frac{\partial^2 \Psi}{\partial k^2} & -\frac{\partial^2 \Psi}{\partial k \partial \alpha} & -\frac{\partial^2 \Psi}{\partial k \partial \beta} \\ -\frac{\partial^2 \Psi}{\partial \alpha \partial k} & -\frac{\partial^2 \Psi}{\partial \alpha^2} & -\frac{\partial^2 \Psi}{\partial \alpha \partial \beta} \\ -\frac{\partial^2 \Psi}{\partial \beta \partial k} & -\frac{\partial^2 \Psi}{\partial \beta \partial \alpha} & -\frac{\partial^2 \Psi}{\partial \beta^2} \end{pmatrix}. \quad (16)$$

Where Ψ is ML function or Q-function. Consequently, the estimators $\hat{\alpha}$, $\hat{\beta}$ and $\hat{k}$ for α , β and k have an asymptotic variance covariance matrix defined by inverting the Fisher information matrix F and substituting $\hat{\theta} = (\hat{\alpha}, \hat{\beta}, \hat{k})$ for $\theta = (\alpha, \beta, k)$, if $\theta = (0.8, 2.2, 0.5)$ at $n=40$ and $r=28$ then we get the estimators $\hat{\theta} = (0.8353, 2.1064, 0.4968)$. Therefore, the asymptotic variance covariance matrix depend on ML method is given as follows, see tables (1-b),

$$\begin{pmatrix} 0.0165 & -0.0073 & -0.0523 \\ -0.0073 & 0.0293 & 0.0315 \\ -0.0523 & 0.0315 & 0.4089 \end{pmatrix}.$$

Therefore, the asymptotic variance covariance matrix depend on EM method is given as follows, see tables (2-b).

$$\begin{pmatrix} 0.0156 & -0.0074 & -0.0494 \\ -0.0074 & 0.0316 & 0.0325 \\ -0.0494 & 0.0325 & 0.3810 \end{pmatrix}.$$

5 Confidence Interval for Burr-XII distribution under Data Type-II Censoring

In this section, we construct confidence intervals (CIs) for the unknown parameters α , β and k for larger sample sizes under data type II- censored. The two-sided approximate confidence limits for α , β and k will be respectively as follows,

$$\begin{aligned} L_{\alpha} &= \hat{\alpha} - z\sigma(\hat{\alpha}), & U_{\alpha} &= \hat{\alpha} + z\sigma(\hat{\alpha}), \\ L_{\beta} &= \hat{\beta} - z\sigma(\hat{\beta}), & U_{\beta} &= \hat{\beta} + z\sigma(\hat{\beta}), \\ L_k &= \hat{k} - z\sigma(\hat{k}), & U_k &= \hat{k} + z\sigma(\hat{k}). \end{aligned} \quad (17)$$

Where L_{θ} and U_{θ} are lower and upper confidence limits which enclosed θ with

$$\begin{aligned}
 & -(k+1) \sum_{i=N_1+1}^r \mathbb{E} \left(\log \left[1 + \left(t_{i:n} - \tau_1 + \tau_1 \frac{\alpha}{\beta} \right)^\beta \right] \middle| t_{i:n} > \delta_{i:n}, \hat{\theta}_m \right) \\
 & + k(n-r) \log \left[1 + \left(\delta_{r:n} - \tau_1 + \tau_1 \frac{\alpha}{\beta} \right)^\beta \right].
 \end{aligned} \tag{14}$$

(ii) The Maximization step(M-Step)

We note that the maximization can be done by finding the solution of

$$\mathbb{E} \left(\frac{\partial \log L_\delta(t_{i:n} | t_{i:n} > \delta_{i:n})}{\partial \theta}, \theta_m \right) = 0, \tag{15}$$

therefore, $\frac{\partial \Phi(\theta, \hat{\theta}_m)}{\partial \theta} = 0$.

Then, we obtain the estimators of α , β and k by differentiating Eq. (14) with respect to α , β and k respectively and equating to zero. The results from the differential system of nonlinear equations it will be difficult to solve. The (E-step) and (M-step) are repeatedly until the estimates of parameters converges. Therefore, we use the MATLAB program to solve the nonlinear equations simultaneously to obtain $\hat{\alpha}$, $\hat{\beta}$ and $\hat{k}$, see tables (2-a).

4 Asymptotic Variances and Covariance Matrix

The asymptotic variance and covariance matrix are given by the elements of the inverse of Fisher information matrix as follows, $I_{ij}(\underline{\theta}) \cong E \left(-\frac{\partial^2 \Psi}{\partial \theta_i \partial \theta_j} \right)$. Unfortunately, the exact mathematical expressions for the previous expectation are very difficult to obtain.

Therefore, Fisher information matrix is given by $I_{ij}(\underline{\theta}) \cong \left(-\frac{\partial^2 \Psi}{\partial \theta_i \partial \theta_j} \right)$, which is obtained

by dropping the expectation on operation E and replaced α, β and k with $\hat{\alpha}, \hat{\beta}$ and $\hat{k}$ respectively, Cohen (1965). The Fisher information matrix F can be written as follows.

$$* \prod_{i=N_1+1}^r \left[\frac{k\beta(t_{i:n}-\tau_1+\tau_1^{\frac{\alpha}{\beta}})^{\beta-1} \left(1 + \left(t_{i:n}-\tau_1+\tau_1^{\frac{\alpha}{\beta}}\right)^{\beta}\right)^{-(k+1)}}{\left(1 + \left(\delta_{i:n}-\tau_1+\tau_1^{\frac{\alpha}{\beta}}\right)^{\beta}\right)^{-k}} \right] \quad (12)$$

The log-likelihood function of the Burr XII distribution via EM- algorithm as follows,

$$\begin{aligned} \log L_{\delta} = & \log \frac{n!}{(n-r)!} + N_1 \log k\alpha + (\alpha-1) \sum_{i=1}^{N_1} \log t_{i:n} - (k+1) \sum_{i=1}^{N_1} \log(1+t_{i:n}^{\alpha}) \\ & + k \sum_{i=1}^{N_1} \log(1+\delta_{i:n}^{\alpha}) + (r-N_1) \log k\beta + (\beta-1) \sum_{i=N_1+1}^r \log \left(t_{i:n}-\tau_1+\tau_1^{\frac{\alpha}{\beta}}\right) \\ & - (k+1) \sum_{i=N_1+1}^r \log \left[1 + \left(t_{i:n}-\tau_1+\tau_1^{\frac{\alpha}{\beta}}\right)^{\beta}\right] + k \sum_{i=N_1+1}^r \log \left[1 + \left(\delta_{i:n}-\tau_1+\tau_1^{\frac{\alpha}{\beta}}\right)^{\beta}\right] \\ & + k(n-r) \left\{ \log \left[1 + \left(t_{r:n}-\tau_1+\tau_1^{\frac{\alpha}{\beta}}\right)^{\beta}\right] + \log \left[1 + \left(\delta_{r:n}-\tau_1+\tau_1^{\frac{\alpha}{\beta}}\right)^{\beta}\right] \right\}. \quad (13) \end{aligned}$$

(i) The Expectation step (E-Step)

For the E-step, we calculate $\Phi(\theta, \hat{\theta}_m)$, where $\theta = (\alpha, \beta, k)$ and $\hat{\theta}_m = (\hat{\alpha}_m, \hat{\beta}_m, \hat{k}_m)$ in m^{th} iteration. Thus, the Q-function of the Burr XII distribution can be obtained as follows,

$$\begin{aligned} \Phi(\theta, \hat{\theta}_m) &= \mathbb{E}[\log L_{\delta}(t_{i:n}|t_{i:n} > \delta_{i:n}), \hat{\theta}_m] \\ &= \int \log L_{\delta}(t_{i:n}|t_{i:n} > \delta_{i:n}) f(t_{i:n}|t_{i:n} > \delta_{i:n}) dt_{i:n}. \end{aligned}$$

Therefore, the Q-function becomes,

$$\begin{aligned} \Phi(\theta, \hat{\theta}_m) = & \log \frac{n!}{(n-r)!} + N_1 \log k\alpha + k \sum_{i=1}^{N_1} \log(1+\delta_{i:n}^{\alpha}) + (r-N_1) \log k\beta \\ & + (\alpha-1) \sum_{i=1}^{N_1} \mathbb{E}(\log t_{i:n} | t_{i:n} > \delta_{i:n}, \hat{\theta}_m) - (k+1) \sum_{i=1}^{N_1} \mathbb{E}(\log(1+t_{i:n}^{\alpha}) | t_{i:n} > \delta_{i:n}, \hat{\theta}_m) \\ & - k(n-r) \log \left[1 + \left(t_{r:n}-\tau_1+\tau_1^{\frac{\alpha}{\beta}}\right)^{\beta}\right] + k \sum_{i=N_1+1}^r \log \left[1 + \left(\delta_{i:n}-\tau_1+\tau_1^{\frac{\alpha}{\beta}}\right)^{\beta}\right] \\ & + (\beta-1) \sum_{i=N_1+1}^r \mathbb{E} \left(\log \left(t_{i:n}-\tau_1+\tau_1^{\frac{\alpha}{\beta}}\right) | t_{i:n} > \delta_{i:n}, \hat{\theta}_m \right) \end{aligned}$$

complete-data problem. For which Maximum Likelihood (MLE) estimation is computationally more tractable. For instance, the complete-data problem for chosen may yield a closed form solution to the (MLE) or may be amenable to MLE computation with a standard computer package. The methodology of the EM-algorithm consists of reformulating the problem in terms of this more easily solved complete-data problem, establishing a relationship between the likelihoods of these two problems and exploiting the simpler (MLE) computation of the complete-data problem in the M-step of the iteration computing algorithm. The pdf of the Burr XII distribution under simple step stress given $t_{i:n} > \delta_{i:n}$ will be applied in Q-function of EM-algorithm. Hence, the pdf of Burr XII distribution given $t_{i:n} > \delta_{i:n}$ is calculated as follows,

$$f(t_{i:n}|t_{i:n} > \delta_{i:n}) = \begin{cases} f_1(t_{i:n}|t_{i:n} > \delta_{i:n}) = \frac{k\alpha t_{i:n}^{\alpha-1}(1+t_{i:n}^\alpha)^{-(k+1)}}{(1+\delta_{i:n}^\alpha)^{-k}}, & \delta_{i:n} \leq t_{i:n} < \tau_1, \\ f_2(t_{i:n}|t_{i:n} > \delta_{i:n}) = \frac{k\beta \left(t_{i:n}-\tau_1+\tau_1\frac{\alpha}{\beta}\right)^{\beta-1} \left(1+\left(t_{i:n}-\tau_1+\tau_1\frac{\alpha}{\beta}\right)^\beta\right)^{-(k+1)}}{\left(1+\left(\delta_{i:n}-\tau_1+\tau_1\frac{\alpha}{\beta}\right)^\beta\right)^{-k}}, & \tau_1 \leq \delta_{i:n} \leq t_{i:n} < \infty. \end{cases} \quad (10)$$

The cdf of Burr XII distribution given $t_{i:n} > \delta_{i:n}$ as follows,

$$F(t_{i:n}|t_{i:n} > \delta_{i:n}) = \begin{cases} F_1(t_{i:n}|t_{i:n} > \delta_{i:n}) = 1 - \left[\frac{(1+t_{i:n}^\alpha)^{-k}}{(1+\delta_{i:n}^\alpha)^{-k}} \right], & \delta_{i:n} \leq t_{i:n} < \tau_1, \\ F_2(t_{i:n}|t_{i:n} > \delta_{i:n}) = 1 - \left[\frac{\left(1+\left(t_{i:n}-\tau_1+\tau_1\frac{\alpha}{\beta}\right)^\beta\right)^{-k}}{\left(1+\left(\delta_{i:n}-\tau_1+\tau_1\frac{\alpha}{\beta}\right)^\beta\right)^{-k}} \right], & \tau_1 \leq \delta_{i:n} \leq t_{i:n} < \infty. \end{cases} \quad (11)$$

The likelihood function of the Burr XII distribution via EM- algorithm as follows,

$$L_\delta(t_{i:n}|t_{i:n} > \delta_i) = \frac{n!}{(n-r)!} \prod_{i=1}^{N_1} \left[\frac{k\alpha t_{i:n}^{\alpha-1}(1+t_{i:n}^\alpha)^{-(k+1)}}{(1+\delta_{i:n}^\alpha)^{-k}} \right] * \left\{ \frac{\left[1+\left(t_{i:n}-\tau_1+\tau_1\frac{\alpha}{\beta}\right)^\beta\right]^{-k}}{\left[1+\left(\delta_{i:n}-\tau_1+\tau_1\frac{\alpha}{\beta}\right)^\beta\right]^{-k}} \right\}^{n-r}$$

$$L(c, k | t) = \frac{n!}{(n-r)!} \left[\prod_{i=1}^r g_2(t_{i:n}) \right] * [1 - G_2(t_{r:n})]^{n-r}, \quad \tau_1 \leq t_{1:n} < \dots < t_{r:n} < \infty. \quad (7)$$

3. If $1 \leq N_1 \leq r-1$, the likelihood function becomes

$$\begin{aligned} L(c, k | t) &= \frac{n!}{(n-r)!} \left[\prod_{i=1}^{N_1} g_1(t_{i:n}) \right] \left[\prod_{i=N_1+1}^r g_2(t_{i:n}) \right] * [1 - G_2(t_{r:n})]^{n-r}, \\ &= \frac{n!}{(n-r)!} \left[1 + \left(t_{r:n} - \tau_1 + \tau_1^{\frac{\alpha}{\beta}} \right)^{\beta} \right]^{-k(n-r)} * \prod_{i=1}^{N_1} k \alpha t_i^{\alpha-1} (1 + t_{i:n}^{\alpha})^{-(k+1)} * \prod_{i=N_1+1}^r k \beta \left[t_{i:n} - \tau_1 + \tau_1^{\frac{\alpha}{\beta}} \right]^{\beta-1} \\ &\quad * \left[1 + \left[t_{i:n} - \tau_1 + \tau_1^{\frac{\alpha}{\beta}} \right]^{\beta} \right]^{-(k+1)}, \quad t_{1:n} < \dots < t_{N_1:n} < \tau_1 \leq t_{N_1+1:n} < \dots < t_{r:n} < \infty. \end{aligned} \quad (8)$$

Then, the log-likelihood function may be written as,

$$\begin{aligned} \ell &= \log \frac{n!}{(n-r)!} + N_1 \log(k\alpha) + (\alpha-1) \sum_{i=1}^{N_1} \log t_{i:n} - (k+1) \sum_{i=1}^{N_1} \log(1 + t_{i:n}^{\alpha}) + N_2 \log(k\beta) \\ &\quad + (\beta-1) \sum_{i=N_1+1}^r \log \left[t_{i:n} - \tau_1 + \tau_1^{\frac{\alpha}{\beta}} \right] - (k+1) \sum_{i=N_1+1}^r \log \left[1 + \left(t_{i:n} - \tau_1 + \tau_1^{\frac{\alpha}{\beta}} \right)^{\beta} \right] \\ &\quad - k(n-r) \log \left[1 + \left(t_{r:n} - \tau_1 + \tau_1^{\frac{\alpha}{\beta}} \right)^{\beta} \right]. \end{aligned} \quad (9)$$

Then, we obtain the estimators of α , β and k by differentiating Eq. (9) with respect to α , β and k respectively and equating to zero. The results from the differential system of nonlinear equations it will be difficult to solve. Therefore, we use the MATLAB program to solve the nonlinear equations simultaneously to obtain $\hat{\alpha}$, $\hat{\beta}$ and $\hat{k}$, see table (1-a).

4 EM-Algorithm under Cumulative Exposure Model

In this section, EM-algorithm is broadly applicable approach to the iterative computation of maximum Likelihood estimates (MLE). On each iteration of the EM algorithm, there are two steps- called the Expectation step or the E-step and the maximization step or the M-step. Therefore, the algorithm is called the EM algorithm. The basic idea of the EM-algorithm is associate with the given incomplete-data, a

and

$$g(t) = \begin{cases} g_1(t) = kc_1 t^{c_1-1} (1+t^{c_1})^{-(k+1)} & , 0 \leq t < \tau_1, \\ g_2(t) = kc_2 \left[t - \tau_1 + \tau_1^{\frac{c_2}{c_1}} \right]^{c_2-1} * \left[1 + \left(t - \tau_1 + \tau_1^{\frac{c_2}{c_1}} \right)^{\frac{c_2}{c_1}} \right]^{-(k+1)} & , \tau_1 \leq t < \infty. \end{cases} \quad (3)$$

Based on the type-II censored, we have n identical units under an initial stress level s_0 . The stress level is changed to s_1 at time τ_1 , and the life-testing experiment is terminated as soon as the r^{th} failure occurs, where r and $0 < \tau_1 < \infty$ are fixed in advance. Under the assumption of the cumulative exposure model, we have derived the corresponding cumulative exposure distribution $G(t)$ and probability exposure density function $g(t)$.

Let us now denote,

- N_1 = Number of units that fail before time τ_1 at stress level s_0 .
- N_2 = Number of units that fail after time τ_1 at stress level s_1 .

With these notation, we will observe the following ordered observations.

$$(t_{1:n} < t_{2:n} < \dots < t_{N_1:n} < \tau_1 \leq t_{N_1+1:n} < t_{N_1+2:n} < \dots < t_{r:n}), \quad (4)$$

where, $r = N_1 + N_2$ and t is the vector of observed Type-II censored data.

3 Maximum Likelihood Estimation under Cumulative Exposure Model

From the cumulative exposure distribution function (2) and the probability distribution function (3), we obtain the likelihood function of c_1 and c_2 based on the type-II hybrid censored sample as follows,

$$L(c_i, k | t) = \frac{n!}{(n-r)!} \left[\prod_{i=1}^r g(t_{i:n}) \right] * [1 - G(t_{r:n})]^{n-r}, \quad 0 < \dots < t_{N_1:n} < \tau_1 < \dots < t_{r:n} < \infty. \quad (5)$$

1. If $N_1 = r$ and $N_2 = 0$, the likelihood function becomes

$$L(c_i, k | t) = \frac{n!}{(n-r)!} \left[\prod_{i=1}^r g_1(t_{i:n}) \right] * [1 - G_1(t_{r:n})]^{n-r}, \quad 0 < t_{1:n} < \dots < t_{r:n} < \tau_1 < \infty. \quad (6)$$

2. If $N_1 = 0$ and $N_2 = r$, the likelihood function becomes

c_i, k The shape parameters of the Burr type XII distribution.

$f(t_{i:n}|t_{i:n} > \delta_{i:n})$ The pdf of Burr XII distribution via EM- algorithm.

$F(t_{i:n}|t_{i:n} > \delta_{i:n})$ The cdf of Burr XII distribution via EM- algorithm.

$\Phi(\theta, \hat{\theta}_m)$ The expectation of Logarithm of likelihood function via EM- algorithm.

$L(.)$ Likelihood function.

$\log L(.)$ Logarithm of likelihood function.

L_{δ_i} Likelihood function via EM- algorithm.

$\log L_{\delta_i}$ Logarithm of likelihood function via EM- algorithm.

$t_{i:n} > \delta_{i:n}$ $t_{i:n}$ denotes censored type II data at $\delta_{i:n}$.

F Fisher information matrix.

MSE Mean square error.

ARBias Absolute relative bias.

RE Relative error.

2.2 Assumptions to obtain the Cumulative Exposure Model

Suppose that the data come from Burr type XII distribution under a cumulative exposure model, and we consider a simple step-stress model based on type-II censored with two stress level s_0 and s_1 . The probability density function and cumulative distribution function for Burr type XII distribution are given by, respectively

$$f_i(t; c_i) = c_i k t^{c_i-1} (1 + t^{c_i})^{-(k+1)}, \quad t > 0 \quad k, c_i > 0 \quad i = 1, 2,$$

$$\text{and} \quad F_i(t; c_i) = 1 - (1 + t^{c_i})^{-k}, \quad t > 0 \quad k, c_i > 0 \quad i = 1, 2. \quad (1)$$

The cumulative exposure distribution (CED) $G(t)$ and probability exposure density (PED) function $G(t)$, respectively,

$$G(t) = \begin{cases} G_1(t) = 1 - (1 + t^{c_1})^{-k} & , 0 \leq t < \tau_1, \\ G_2(t) = 1 - \left[1 + \left(t - \tau_1 + \tau_1^{\frac{c_1}{c_2}} \right)^{c_2} \right]^{-k} & , \tau_1 \leq t < \infty. \end{cases} \quad (2)$$

Louis (1982) finding the observed information matrix by using the EM- algorithm. Miller (1983) present optimum simple step-stress plans for accelerated life testing, and Bai (1989) present optimum simple step-stress accelerated life tests With Censoring. Balakrishnan (1996) developed a byes model for step-stress accelerated life testing. Xiong (1998) studied inference on simple step-stress model with type-II censored exponential data. Watkins (2001) studies the inference in simple step-stress models. Balakrishnan (2004) studied the EM-algorithm for type-II right censored bivariate normal data. Park (2004) studied optimal design of step-stress degradation tests in the case of destructive measurement. Balakrishnan (2007) studied point and interval estimation for a simple step-stress model with type-II censoring. Al-Masri (2009) studied optimum times for step-stress cumulative exposure model using Log-Logistic distribution with known scale parameter. Jalali (2009) considered three related aspects of maximum likelihood estimation of parameters in the two parameter Burr XII distribution. Wang (2010) studied the estimation for Burr XII parameters with multiple censored data by using EM-algorithm.

2 Model Description

2.1 Notation

ALT Acceerated Life Testing.

SST Step Stess Testing.

s_0, s_1 Stress levels.

$G(t)$ Cumulative exposure distribution (CED) function.

$g(t)$ Probability exposure density (PED) function.

n Identical units under an initial stress level s_0 .

$r \leq n$ and $0 < \tau_1 < \tau_2 < \infty$ are fixed in advance.

$t_{1:n} < t_{2:n} < \dots < t_{n:n}$ The ordered failure times of the n unit under test.

τ_1 A fixed time before which the stress level is changed from s_0 to s_1 .

$t_{r:n}$ The time when the r^{th} failure occurs the experiment is terminated.

N_1 Number of units that fail before time τ_1 at stress level s_0 .

N_2 Number of units that fail before time τ_2 at stress level s_1 .

Point and Interval Estimation for the Parameters of Burr Type-XII Distribution based on Cumulative Exposure Model

Abd-Elfattah, A. M. ¹ , Essawy, A. ² and Mubarak, Sh. ³

Abstract: In this article, we consider the estimation problem for a simple step-stress model under type II censoring data from Burr type XII distribution. We use the maximum likelihood method and EM algorithm method to obtain the estimators of the unknown parameters with respect to a cumulative exposure model. In addition, asymptotic variance and covariance matrix of the estimators are given. The performance of the findings in the paper is showed by demonstrating some numerical illustration through Monte Carlo simulation. Finally, we present examples to illustrate all the methods of inference.

Keywords: Burr XII distribution; Cumulative exposure model; EM algorithm; Fisher information matrix; Type-II censoring data; Maximum likelihood estimation.

1 Introduction

In reliability and life-testing experiments, the researcher is often interested in the effects of extreme or varying stress factors such a temperature, voltage, and load on the lifetimes of experimental units. Step-stress testing (SST), which is a special class of accelerated life-testing (ALT), allows the experimenter to increase the stress levels at fixed times during the experiment in order to obtain information on the parameters of the life distribution more quickly than under normal operating conditions. Dempster (1977) considered the maximum likelihood from incomplete data via EM- algorithm. McCool (1980) studied the Confidence limits for Weibull regression with censored data. Nelson (1980) studies accelerated life testing step stress model and data analysis.

¹ Department of Mathematical Statistics, Institute of Statistical studies & Research, Cairo University, Egypt

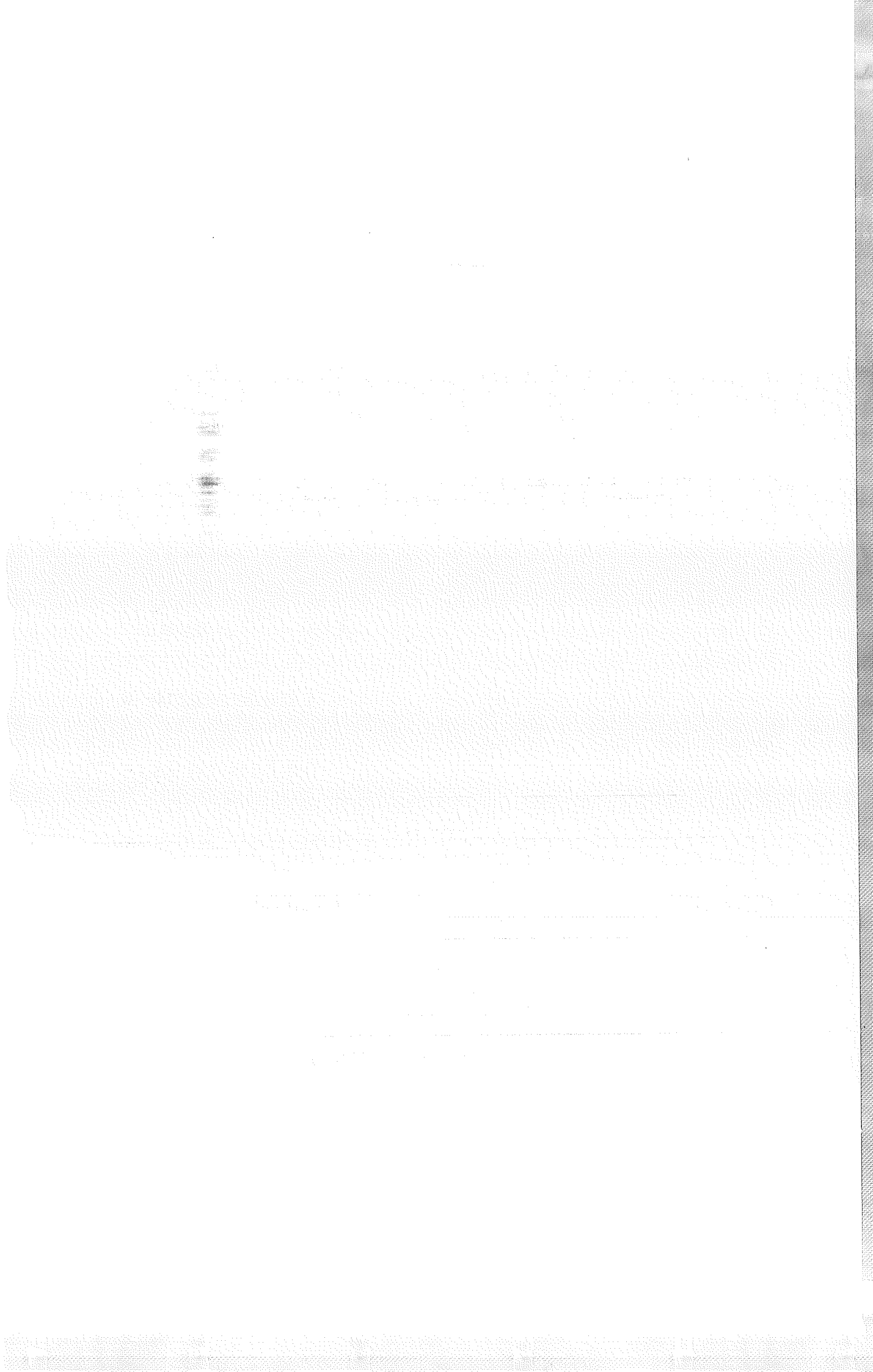
² Department of Mathematics, Faculty of Science, Minia University, Egypt.

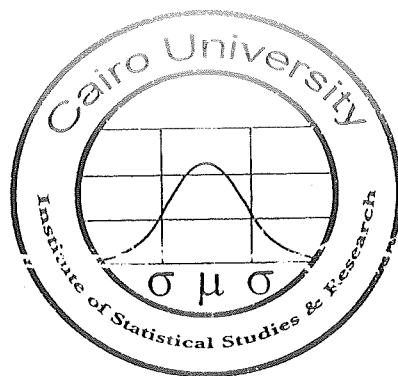
³ The Higher Technological Institute, El-Minia, Egypt.

Index

MATHEMATICAL STATISTICS

1	Point and Interval Estimation for the Parameters of Burr Type-XII Distribution on Cumulative Exposure Model Abd-Elfattah, A.M. , Essawy, A. , Mubarak, Sh	1-22
2	Estimation of $(Y < X)$ for Weibull Distribution in the Presence of k Outliers Elsayed, A.Elsherpieny , Rania, M.Shalaby	23-36
3	Estimation of the Bivariate Generalized Linear Failure Rate Distribution Under Different Censoring Schemes Ahmed F. Attia, Hanan M.Aly , Hiba Z. Muhammed	37-55
4	On Some Results of Bayesian Regression With Missing Data Abd El- Fattah Kandil , Mervat Mahdy , Dina El-Telbany	56-70
5	Prediction Interval Of Lower Record Value For The Inverse Weibull Distribution W.yahia , A.N.Ahmed	71-87





Cairo University
Institute of Statistical Studies and Research

The 47th Annual
Conference
on Statistics,
Computer Sciences,
and Operation Research

MATHEMATICAL STATISTICS

24-27, Dec . 2012