

where λ denotes the proportionality factor.

A number of values of λ are selected and the path of steepest descent is followed as long as $S(\theta)$ decreases. When it does not, another combination of levels of $\theta_1, \theta_2, \dots, \theta_p$ is selected and the process is continued until it converges to the value which minimises $S(\theta)$.

(iii) A method developed by Marquardt (1963), represents a compromise between the linearisation method and the steepest descent method. Marquardt's 'compromise' develops a maximum neighbourhood method, which performs an optimum interpolation between the linearisation method and the gradient method, the linearisation being based upon the maximum neighbourhood in which the truncated Taylor series gives an adequate representation of the nonlinear model.

A broad outline of the algorithm used can be described as follows: Suppose a certain point in the parameter space is considered. If the method of steepest descent is applied, then movement away from this initial point will result in a certain vector direction, δ_g , being obtained. This new direction may be the best local direction in which to move to attain smaller values of $S(\theta)$ but may not be the best overall direction, due to the attenuation in the $S(\theta)$ contours. However, the best direction must be within 90° of δ_g or else $S(\theta)$ will get larger locally. The linearisation method leads to another correction vector δ . For a number of practical problems investigated by Marquardt, he found that the angle, ϕ say, between δ_g and δ , fell in the range $80^\circ < \phi < 90^\circ$. The two directions are therefore almost at right angles and Marquardt's algorithm provides a method for interpolating between the vectors δ_g and δ , and also for obtaining a suitable step size.

4.2 Use of the Jackknife Method in Linear and Non-Linear Models

In an important paper by Hinkley (1977a), the jackknife method is applied to a general linear model of the form:

$$Y = X\beta + \epsilon \quad (4.2.1)$$

where the e_i 's are taken to be i.i.d. random variables with mean zero and variance σ^2 . He considers the quantities:

$$w_i = X_i' (X'X)^{-1} X_i, \quad i=1, \dots, n \quad (4.2.2)$$

where X_i' is the i th row of the $n \times p$ matrix X . In the balanced case where the w_i s are constant i.e. $w_i = p/n$ for $i=1, 2, \dots, n$, he shows that the standard jackknife method produces less biased estimators than the least squared method. Hinkley also considers the more interesting non-linear case where the model is unbalanced i.e. w_i is not constant. He proposes the weighted psudovalue:

$$\hat{Q}_i = \hat{\beta} + n(1-w_i) (\hat{\beta} - \hat{\beta}_{-i}), \quad (4.2.3)$$

and the weighted jackknife estimator of β :

$$\hat{Q} = n^{-1} \sum_{i=1}^n \hat{Q}_i \quad (4.2.4)$$

with covariance matrix estimator:

$$V(\hat{Q}) = (vQ_{ij})_{p \times p} = \{n(n-p)\}^{-1} \sum_{i=1}^n (\hat{Q}_i - \hat{Q})(\hat{Q}_i - \hat{Q})', \quad (4.2.5)$$

Hinkley suggests that the weighted estimator (4.2.4) may have superior performance for unbalanced data.

It is also interesting to note that Miller (1974b), in an earlier paper also investigated the performance of the jackknife in the case of unbalanced problems. Quenouille originally proposed the jackknife to eliminate a bias term of order $1/n$. This it does exactly in balanced problems (Gray et al, 1972). However, in the case of unbalanced problems such as non-linear least squares, it is not at all clear what jackknifing does to the bias of $\hat{\theta}_n = f(\hat{\beta})$, where $\hat{\beta}$ is the least squares estimator of

β Miller shows that from the expansion

$$f(\hat{\beta}) = f(\beta) + (\hat{\beta} - \beta)^T f'(\beta) + 1/2 (\hat{\beta} - \beta)^T f''(\beta) (\hat{\beta} - \beta) + \dots \quad (4.2.6)$$

one obtains

$$E(f(\hat{\beta})) = f(\beta) = \frac{\sigma^2}{2} \text{tr. } f''(\beta) (X_n^T X_n)^{-1} + \dots \quad (4.2.7)$$

After jackknifing, the first order bias term becomes

$$\frac{\sigma^2}{2} \left\{ \text{tr. } f''(\beta) (X_n^T X_n)^{-1} - \frac{(n-1)}{n} \sum_{i=1}^n \frac{x_i (X_n^T X_n)^{-1} f''(\beta) (X_n^T X_n)^{-1} x_i^T}{1 - x_i (X_n^T X_n)^{-1} x_i^T} \right\}$$

What the second term does to the absolute size of the bias is unknown. This provides an important motivating factor for the research topic relating to non-linear parameter estimation which is investigated via Monte-Carlo simulation studies in Section 4.4.

In a paper by Duncan (1978), he examined the jackknife procedure for constructing confidence regions in non-linear regression using Monte Carlo simulation. Although attempts have been made by Williams (1962), Halperin (1963) and Hartley (1964), to determine exact confidence regions for the parameters in a non-linear model, the methods proposed are only applicable to the very simplest of non-linear models and, when the number of parameters is moderate, they produce regions which are very elliptic and non-robust if the distribution of errors e is non-normal.

Duncan considered the model

$$y = f(x, \theta) + e$$

where

$$f(x, \theta) = \frac{\theta_1}{(\theta_1 - \theta_2)} \{ \exp(-\theta_2 x) - \exp(-\theta_1 x) \} \quad (4.2.8)$$

and the error e followed the standard normal or 'contaminated' normal distribution. Univariate and joint confidence intervals were determined for the parameters θ_1 and θ_2 . For the jackknife method, successive deletions of size 1, 2, 3 and 6 were considered based on a moderate sample of size $n=24$. The jackknife was found to be untrustworthy in establishing joint confidence intervals as larger blocks of observations were deleted. For example, based on nominal 95% confidence regions, only a 35.8% and 34.2% coverage were obtained for a deleted block size of 6 where the error distribution was normal and 'contaminated' normal respectively.

In an interesting paper by Fox, Hinkely and Larntz (1980), they consider the standard jackknife and two linear jackknife methods in the context of non-linear regression, and also investigate the findings of Duncan's paper. The linear jackknife methods could be based on linear approximations obtained from estimators of the influence function for θ , using methods described by Devlin, Gnanadesikan and Kettenring (1975), and Hinkely (1977b), or, alternatively, by using a Taylor series expansion, i.e. consider a Taylor expansion of the least squares estimating equation for $\hat{\theta}_{-i}$, assuming this to be a stationary point of

$$\sum_{j=1}^n \{Y_j - f(x_j, \theta)\}^2$$

and pick off the linear term. The result is

$$\hat{\theta}_{-i} = \hat{\theta} - (\hat{Z}^T \hat{Z})^{-1} \hat{Z}_i \left\{ \frac{r_i}{1 - \hat{w}_i} \right\}, \quad (i=1, 2, \dots, n), \quad (4.2.9)$$

where

$$\hat{Z}_i = \nabla f(x_i, \hat{\theta}) = \left(\frac{\partial}{\partial \theta} f(x_i, \theta), \dots, \frac{\partial}{\partial \theta} f(x_i, \theta) \right)_{\theta=\hat{\theta}}^T$$

$$\hat{Z}^T = (\hat{Z}_1^T, \dots, \hat{Z}_n^T), \quad \hat{w}_i = \hat{Z}_i^T (\hat{Z}^T \hat{Z})^{-1} \hat{Z}_i$$

and

$$r_i = y_i - f(x_i, \hat{\theta})$$

Substitution of (4.2.9) into the standard formula for jackknife pseudovalues, and replacement of $(n-1)$ by n , gives the linear pseudovalues

$$LP_i = \hat{\theta} + n(\hat{Z}^T \hat{Z})^{-1} \hat{Z}_i \left\{ \frac{r_i}{1 - \hat{w}_i} \right\} \quad (4.2.10)$$

The linear jackknife estimator of θ is:

$$\overline{LP} = n^{-1} \sum_{i=1}^n LP_i \quad (4.2.11)$$

with variance-covariance matrix estimate:

$$V(\overline{LP}) = \frac{1}{n(n-1)} \sum_{i=1}^n (LP_i - \overline{LP})(LP_i - \overline{LP})' \quad (4.2.12)$$

An alternative weighted jackknife is also suggested where the pseudovalues LP_i are replaced by

$$LQ_i = \hat{\theta} + n(\hat{Z}^T \hat{Z})^{-1} \hat{Z}_i \cdot r_i \quad (4.2.13)$$

The weighted jackknife estimator $\overline{LQ}$ and variance estimator $V(\overline{LQ})$ are defined as:

$$\overline{LQ} = n^{-1} \sum_{i=1}^n LQ_i \quad (4.2.14)$$

$$\text{and } V(\overline{LQ}) = (\hat{Z}^T \hat{Z})^{-1} \sum_{j=1}^n \hat{Z}_j \hat{Z}_j^T \cdot r_j^2 (\hat{Z}^T \hat{Z})^{-1} \quad (4.2.15)$$

Using the same example as Duncan, the linear jackknife estimators $\overline{LP}$ and $\overline{LQ}$ produced similar estimators to the standard jackknife method for the parameters θ_1 and θ_2 . In the overlap with Duncan results,

there was general agreement, suggesting that all the jackknife methods based on normal approximation are poor for joint confidence regions.

In the most comprehensive simulation study to date, Frangos (1981), used seven different estimators of a two parameter non-linear model. The estimators were the standard least squares and jackknife estimators, Hinkley's weighted jackknife, the two linear jackknives discussed by Fox, Hinkley and Larntz, and two alternative weighted jackknives. The weights w_i were functions of the jackknife estimator and Hinkley's weighted jackknife, respectively, for these two alternative weighted jackknives. The non-linear model considered was of the form:

$$Y = \theta_1 \exp(\theta_2 X) + e, \quad (4.2.16)$$

where the e_j 's were assumed to be i.i.d. normal random variables with mean zero and variance σ^2 .

The results of the simulation study indicated that the standard jackknife and the three weighted jackknife procedures gave smaller biases than the least squares procedure, at the expense of a very small increase of the mean square error. The two linear jackknives did not show any improvement in terms of bias over the least squares procedures. Also, the jackknife procedures were not effective in establishing joint confidence regions, which is in agreement with Duncan, and Fox et al.

4.3 Confidence Interval Estimation for the Different Procedures in Non-Linear Models

Firstly, the non-linear least squares estimate $\hat{\theta} = (\hat{\theta}_1, \hat{\theta}_2, \dots, \hat{\theta}_p)$ of θ , is calculated using the linearisation method. From Hartley and Booker (1965), the asymptotic estimate of the covariance matrix of $\hat{\theta} = (\hat{\theta}_1, \hat{\theta}_2, \dots, \hat{\theta}_p)$ is calculated from the following $p \times p$ matrix:

$$\hat{\Sigma} = (\hat{\Sigma}_{ij}) = (n-p)^{-1} \hat{\sigma}^2 A^{-1} \quad (4.3.1)$$

where

$$\hat{\sigma}^2 = \sum_{i=1}^n (Y_i - f(x_i; \hat{\theta}))^2$$

and

$$A = (\alpha_{ij})_{p \times p} = \left\{ \sum_{k=1}^n \frac{\partial f(x_k; \theta)}{\partial \theta_i} \frac{\partial f(x_k; \theta)}{\partial \theta_j} \right\}_{\theta = \hat{\theta}}$$

The 100(1- α)% least squares confidence limit for component θ_i of θ is given by:

$$\hat{\theta}_i \pm \sqrt{\hat{\sum}_{ii} X_{1, 1-\alpha}^2} \quad (i=1, 2, \dots, p) \quad (4.3.2)$$

with

$$SS(\theta) = \sum_{i=1}^n (Y_i - f(x_i; \theta))^2$$

The standard jackknife estimate of θ is given by:

$$\hat{P} = \bar{P} - \frac{1}{n} \sum_{i=1}^n \hat{P}_i \quad (4.3.3)$$

and its estimated covariance matrix is given by the following:

$$V(\hat{P}) = (VP_{ij})_{p \times p} = \{n(n-1)\}^{-1} \sum_{i=1}^n (\hat{P}_i - \hat{P})(\hat{P}_i - \hat{P})' \quad (4.3.4)$$

The separate confidence limit for the component θ_i of θ is given by:

$$\hat{P}_i \pm \sqrt{\hat{V}P_{ii} \cdot X_{1, 1-\alpha}^2} \quad (4.3.5)$$

Similarly, if the standard bootstrap method is considered for estimation purposes, the bootstrap estimate of θ is given by:

$$\hat{\theta}^* = \frac{\sum_{k=1}^B \hat{\theta}^{*k}}{B} \quad (4.3.6)$$

with variance estimate

$$\hat{V} = \left\{ \frac{1}{B-1} \sum_{k=1}^B (\hat{e}^{*k} - \hat{\theta}^*)^2 \right\} \quad (4.3.7)$$

where $\hat{\theta}^{*k}$ are the bootstrap replications and B is the number of bootstrap replications.

The separate confidence limit for the component θ_i of θ is given by:

$$\hat{\theta} \pm \sqrt{\hat{V} \chi^2_{1, 1-\alpha}} \quad (4.3.8)$$

Other bootstrap confidence limit methods such as Efron's percentile, the bias corrected, the improved method using the 'acceleration' constant (see Section 2.9) as well as Abramovitch and Singh's (1985) studentized pivotal quantity method (see Sections 2.9 and 4.4) are also available for confidence limit estimation in non-linear models.

4.4 Monte-Carlo Simulation Studies

This section describes the simulation studies that were undertaken regarding the non-parametric estimation of parameters for various non-linear models. Four different p.d.fs were used to generate the error deviates, namely the normal, exponential, gamma and long tailed h distributions. For each simulation run, unless otherwise stated, 1000 samples were generated, each for a given size of $N=20$. The parameter values for θ_1 and θ_2 in the non-linear models under consideration were $+1.0$ and -1.0 respectively. Estimators of the value of the ratio $-\theta_2/\theta_1$, were also considered. In the case of the non-linear models considered by Duncan (1978), (see equation 4.2.8), and Frangos (1981), (see equation 4.2.16), $x_i = -4.0(0.2) - 0.2$, $i=1$ to 20. The accuracy of the parameter estimators were expressed in terms of bias, sample variance and mean square error. The 95% and 99% levels were considered for the confidence limits and their robustness was determined in terms of coverage, average length and variance of length.

The following simulation studies were undertaken:

(i) a comparison between the least squares and jackknife and bootstrap methods using Duncan's non-linear model. The results of the simulation study are summarised in Table 4.4.1 and 4.4.2. The conclusions of the study are as follows:

- a. the jackknife estimator has a smaller bias for the parameter θ_2 than the least squares estimator, although this is at the expense of a larger mean square error. In general, however, the jackknife does not consistently outperform the least squares method in terms of bias reduction or mean square error.
- b. the bootstrap produced the worst bias results although the mean square errors were similar to those obtained by the least squares method.
- c. both the least squares and jackknife methods generally produced robust confidence limits in terms of coverage. The least squares method did, however, produce only modest coverage results for the parameters at the 95% level when the errors were normally distributed.
- d. the bootstrap method produced poor coverage results generally for the different error distributions and is an unsuitable method for establishing robust confidence limits for this area of application.

(ii) a comparison between the least squares and jackknife method using Duncan's non-linear model for different error distributions where the variance of the errors is not constant. In particular, for a sample size of 20, the variance of the errors is different for each consecutive sub-sample of size 5. Evaluation of the bootstrap method was not considered due to its poor performance in the previous study. The results of the simulation study are summarised in Tables 4.4.3 and 4.4.4. The conclusions of the study are as follows:

Table 4.4.1 : Bias, Sample Variance and Mean Square Error for a Two Parameter Non-Linear Model

		Least Squares Method			Jackknife Method			Bootstrap Method		
		θ_1	θ_2	$-\theta_2/\theta_1$	θ_1	θ_2	$-\theta_2/\theta_1$	θ_1	θ_2	$-\theta_2/\theta_1$
A	Bias	-2.43×10^{-5}	-4.73×10^{-3}	3.15×10^{-3}	2.72×10^{-4}	-7.90×10^{-4}	-1.32×10^{-3}	1.06×10^{-4}	-7.92×10^{-3}	6.39×10^{-3}
	Sample Variance	2.07×10^{-4}	1.62×10^{-2}	1.27×10^{-2}	2.42×10^{-4}	1.84×10^{-2}	1.44×10^{-2}	2.09×10^{-4}	1.65×10^{-2}	1.29×10^{-2}
	Mean Square Error	2.07×10^{-4}	1.62×10^{-2}	1.27×10^{-2}	2.42×10^{-4}	1.84×10^{-2}	1.44×10^{-2}	2.09×10^{-4}	1.66×10^{-2}	1.29×10^{-2}
B	Bias	-1.51×10^{-4}	-3.67×10^{-3}	2.20×10^{-3}	2.20×10^{-5}	1.32×10^{-3}	-3.17×10^{-3}	-2.29×10^{-4}	-7.04×10^{-3}	5.61×10^{-3}
	Sample Variance	2.07×10^{-4}	1.66×10^{-2}	1.30×10^{-2}	2.37×10^{-4}	1.87×10^{-2}	1.44×10^{-2}	2.11×10^{-4}	1.70×10^{-2}	1.33×10^{-2}
	Mean Square Error	2.07×10^{-4}	1.66×10^{-2}	1.30×10^{-2}	2.37×10^{-4}	1.87×10^{-2}	1.44×10^{-2}	2.11×10^{-4}	1.71×10^{-2}	1.33×10^{-2}
C	Bias	-2.10×10^{-6}	-1.99×10^{-3}	1.26×10^{-3}	1.24×10^{-4}	8.45×10^{-5}	-1.04×10^{-3}	-5.50×10^{-5}	-3.42×10^{-3}	2.73×10^{-3}
	Sample Variance	9.60×10^{-5}	7.29×10^{-3}	5.74×10^{-3}	1.10×10^{-4}	8.23×10^{-3}	6.43×10^{-3}	9.70×10^{-6}	7.42×10^{-3}	5.84×10^{-3}
	Mean Square Error	9.60×10^{-5}	7.29×10^{-3}	5.74×10^{-3}	1.10×10^{-4}	8.23×10^{-3}	6.43×10^{-3}	9.70×10^{-6}	7.42×10^{-3}	5.84×10^{-3}
D	Bias	-4.49×10^{-4}	-1.40×10^{-2}	9.71×10^{-3}	-1.49×10^{-4}	-3.50×10^{-4}	-4.74×10^{-3}	-1.90×10^{-2}	-6.68×10^{-3}	2.08×10^{-3}
	Sample Variance	1.47×10^{-3}	4.08×10^{-2}	3.13×10^{-2}	1.54×10^{-3}	4.78×10^{-2}	3.69×10^{-2}	2.01×10^{-2}	6.05×10^{-2}	5.06×10^{-2}
	Mean Square Error	1.47×10^{-3}	4.10×10^{-2}	3.14×10^{-2}	1.54×10^{-3}	4.78×10^{-2}	3.69×10^{-2}	2.05×10^{-2}	6.06×10^{-2}	5.06×10^{-2}

Error Distributions

A Normal : Mean=0.0, S.D=0.2;

B Exponential : Mean=0.2

C Gamma : Mean=0.2 (k=2)

D Long tailed h : Mean=0.2 (h=0.2)

Table 4.4.2 : Sample Confidence Intervals for a Two Parameter Non-Linear Model

		Least Squares Method		Jackknife Method		Bootstrap Method		
		θ_1	θ_2	θ_1	θ_2	θ_1	θ_2	
	Coverage	95%	92.1	89.1	95.2	94.5	90.5	90.1
		99%	98.2	97.7	98.8	98.8	96.8	96.2
A	Average Length	95%	4.59×10^{-2}	3.83×10^{-1}	6.84×10^{-2}	5.93×10^{-1}	5.01×10^{-2}	4.45×10^{-1}
		99%	6.27×10^{-2}	5.23×10^{-1}	9.35×10^{-2}	8.11×10^{-1}	6.59×10^{-2}	5.85×10^{-1}
	Variance of Length	95%	7.29×10^{-5}	5.71×10^{-3}	4.05×10^{-4}	2.47×10^{-2}	8.30×10^{-5}	7.70×10^{-3}
		99%	1.36×10^{-4}	1.07×10^{-2}	7.57×10^{-4}	4.61×10^{-2}	1.44×10^{-4}	1.33×10^{-2}
	Coverage	95%	95.9	92.9	95.8	95.4	90.2	90.1
		99%	99.0	98.0	99.1	99.4	96.7	96.6
B	Average Length	95%	6.27×10^{-2}	5.29×10^{-1}	6.57×10^{-2}	5.74×10^{-1}	4.92×10^{-2}	4.38×10^{-1}
		99%	8.57×10^{-2}	7.24×10^{-1}	8.98×10^{-2}	7.84×10^{-1}	6.47×10^{-2}	5.76×10^{-1}
	Variance of Length	95%	3.49×10^{-4}	3.45×10^{-2}	7.20×10^{-4}	5.43×10^{-2}	2.44×10^{-2}	2.14×10^{-2}
		99%	6.53×10^{-4}	6.45×10^{-2}	1.35×10^{-3}	1.01×10^{-1}	4.21×10^{-4}	3.70×10^{-2}
	Coverage	95%	97.9	97.0	96.1	96.3	91.0	91.0
		99%	99.5	99.1	99.3	99.4	96.9	97.2
C	Average Length	95%	5.42×10^{-2}	4.55×10^{-1}	4.63×10^{-2}	4.03×10^{-1}	3.44×10^{-2}	3.04×10^{-1}
		99%	7.40×10^{-2}	6.22×10^{-1}	6.33×10^{-2}	5.51×10^{-1}	4.52×10^{-2}	3.99×10^{-1}
	Variance of Length	95%	1.49×10^{-4}	1.38×10^{-2}	2.80×10^{-4}	1.87×10^{-2}	7.30×10^{-5}	5.78×10^{-3}
		99%	2.78×10^{-4}	2.57×10^{-2}	5.24×10^{-4}	3.50×10^{-2}	1.26×10^{-4}	9.99×10^{-3}
	Coverage	95%	94.7	92.7	97.3	97.8	89.0	89.5
		99%	99.1	98.3	99.4	99.7	94.3	95.2
D	Average Length	95%	7.48×10^{-2}	6.38×10^{-1}	8.90×10^{-2}	7.92×10^{-1}	6.74×10^{-2}	5.75×10^{-1}
		99%	1.02×10^{-1}	8.73×10^{-1}	1.22×10^{-1}	1.08×10^0	8.86×10^{-2}	7.56×10^{-1}
	Variance of Length	95%	1.87×10^{-3}	1.71×10^{-1}	4.36×10^{-3}	4.39×10^{-1}	2.42×10^{-3}	1.55×10^{-1}
		99%	3.48×10^{-3}	3.19×10^{-1}	8.14×10^{-3}	8.18×10^{-1}	4.18×10^{-3}	2.67×10^{-1}

Error Distribution

A Normal : Mean=0.0, S.D.=0.2

B Exponential : Mean=0.2

C Gamma : Mean=0.2 (k=2)

D Long tailed h : Mean=0.2 (h=0.2);

Table 4.4.3 : Bias, Sample Variance and Mean Square Error for a Two Parameter Model

- (errors are not constant)

	Least Squares Method			Jackknife Method		
	θ_1	θ_2	$-\theta_2/\theta_1$	θ_1	θ_2	$-\theta_2/\theta_1$
A	Bias	-4.11×10^{-5}	-2.26×10^{-3}	1.43×10^{-3}	-4.37×10^{-5}	1.57×10^{-4}
	Sample Variance	1.18×10^{-4}	8.49×10^{-3}	6.62×10^{-3}	1.78×10^{-4}	1.25×10^{-2}
	Mean Square Error	1.18×10^{-4}	8.50×10^{-3}	6.62×10^{-3}	1.78×10^{-4}	1.25×10^{-2}
B	Bias	-7.54×10^{-3}	7.45×10^{-2}	-6.83×10^{-2}	-7.26×10^{-3}	7.56×10^{-2}
	Sample Variance	1.16×10^{-4}	7.59×10^{-3}	6.04×10^{-3}	1.73×10^{-4}	1.14×10^{-2}
	Mean Square Error	1.73×10^{-4}	1.31×10^{-2}	1.07×10^{-2}	2.26×10^{-4}	1.71×10^{-2}
C	Bias	-7.69×10^{-3}	7.70×10^{-2}	-7.03×10^{-2}	-7.46×10^{-3}	7.64×10^{-2}
	Sample Variance	5.85×10^{-5}	3.72×10^{-3}	2.96×10^{-3}	8.49×10^{-5}	5.37×10^{-3}
	Mean Square Error	1.17×10^{-4}	9.65×10^{-3}	7.90×10^{-3}	1.40×10^{-4}	1.12×10^{-2}
D	Bias	-1.12×10^{-2}	8.12×10^{-2}	-7.48×10^{-2}	-1.04×10^{-2}	8.08×10^{-2}
	Sample Variance	4.05×10^{-3}	1.08×10^{-2}	9.63×10^{-3}	4.62×10^{-3}	1.99×10^{-2}
	Mean Square Error	4.19×10^{-3}	1.74×10^{-2}	1.52×10^{-2}	4.73×10^{-3}	2.64×10^{-2}

Error Distribution

A Normal : Mean=0.0; S.D.=0.20, S.D.=0.15, S.D.=0.10, S.D.=0.5 ;

B Exponential : Mean=0.20, 0.15, 0.10, 0.05;

C Gamma : Mean=0.20, 0.15, 0.10, 0.05; (k=2)

D Long tailed h : Mean=0.20, 0.15, 0.10, 0.05; (h=0.2)

Table 4.4.4 : Sample Confidence Intervals for a Two Parameter Non-Linear Model

- (error variances are not constant)

		Least Squares Method		Jackknife Method	
		θ_1	θ_2	θ_1	θ_2
A	Coverage	95%	81.6	80.4	92.7
		99%	92.5	92.3	97.8
	Average Length	95%	2.84×10^{-2}	2.37×10^{-1}	5.24×10^{-2}
		99%	3.88×10^{-2}	3.24×10^{-1}	7.17×10^{-2}
	Variance of Length	95%	3.72×10^{-5}	2.73×10^{-3}	4.56×10^{-4}
		99%	6.94×10^{-5}	5.10×10^{-3}	8.52×10^{-4}
B	Coverage	95%	74.2	66.7	91.6
		99%	90.1	82.9	98.1
	Average Length	95%	3.06×10^{-2}	2.43×10^{-1}	5.07×10^{-2}
		99%	4.18×10^{-2}	3.33×10^{-1}	6.92×10^{-2}
	Variance of Length	95%	9.20×10^{-5}	5.72×10^{-3}	7.43×10^{-4}
		99%	1.72×10^{-4}	1.07×10^{-2}	1.39×10^{-3}
C	Coverage	95%	66.6	55.8	83.8
		99%	85.2	74.4	96.3
	Average Length	95%	2.32×10^{-2}	1.85×10^{-1}	3.58×10^{-2}
		99%	3.18×10^{-2}	2.53×10^{-1}	4.89×10^{-2}
	Variance of Length	95%	3.00×10^{-5}	1.95×10^{-3}	2.62×10^{-4}
		99%	5.61×10^{-5}	3.64×10^{-3}	4.90×10^{-4}
D	Coverage	95%	87.1	73.8	96.1
		99%	98.9	94.4	99.2
	Average Length	95%	3.36×10^{-2}	2.59×10^{-1}	5.94×10^{-2}
		99%	4.59×10^{-2}	3.53×10^{-1}	8.12×10^{-2}
	Variance of Length	95%	1.39×10^{-3}	5.69×10^{-2}	7.27×10^{-3}
		99%	2.60×10^{-3}	1.06×10^{-1}	1.36×10^{-2}

Error Distributicas

A Normal : Mean=0.0; S.D.=0.20, SD=0.15, SD=0.10, S.D.=0.5;

B Exponential : Mean=0.20, 0.15, 0.10, 0.05;

C Gamma : Mean=0.20, 0.15, 0.10, 0.05;

D Long tailed h : Mean=0.20, 0.15, 0.10, 0.05 ; (h=0.2)

- a. the jackknife estimator, in the majority of cases, has a smaller bias for the parameters θ_1 and θ_2 than the least squares estimate, although this is again at the expense of a larger mean square error.
- b. the jackknife method produced robust confidence limits in terms of coverage except in the case where the errors were gamma distributed. However, the least squares method produced untrustworthy confidence limits in all cases, in terms of coverage, and is an unsuitable method for this area of application.

(iii) the previous study was then repeated using a different non-linear model i.e. the one parameter model $Y = \exp(\theta_1 X) + e$, to determine whether the least squares and jackknife method performed as before, in terms of coverage for the parameters. The results of the study are shown in Table 4.4.5. Although the jackknife method only produces modest coverage for the parameters, it is again far superior to that of the least squares method.

(iv) the final case study concerned a comparison between the least squares, jackknife, the standard bootstrap, the 'improved' bootstrap confidence limit methods i.e. Efron's percentile, bias corrected and 'improved' method using the 'acceleration' constant (see Section 2.9) and also Abramovitch and Singh's (1985) studentized pivotal quantity method. A description of how the studentized pivotal quantity method (see Efron, 1987) is applied in a non-linear regression model is as follows:

Consider the model

$$Y = \theta_1 e^{\theta_2 X} + e = g(\theta_1, \theta_2) + e, \text{ where } g(\theta_1, \theta_2) = \theta_1 e^{\theta_2 X} \quad (4.4.1)$$

and the e_i 's are taken to be i.i.d. random variables with mean zero and variance σ^2 .

Table 4.4.5 : Sample Confidence Intervals for a One Parameter Non-Linear Model

- (error variances are not constant)

		Least Squares Method	Jackknife Method
		θ_1	θ_2
A	Coverage	95%	80.9
		99%	91.8
	Average Length	95%	2.57×10^{-2}
		99%	3.51×10^{-2}
	Variance of Length	95%	2.86×10^{-5}
		99%	5.34×10^{-5}
B	Coverage	95%	75.1
		99%	86.9
	Average Length	95%	3.21×10^{-2}
		99%	4.38×10^{-2}
	Variance of Length	95%	8.74×10^{-5}
		99%	1.63×10^{-4}
C	Coverage	95%	71.5
		99%	85.9
	Average Length	95%	2.73×10^{-2}
		99%	3.73×10^{-2}
	Variance of Length	95%	3.37×10^{-5}
		99%	6.30×10^{-5}
D	Coverage	95%	75.4
		99%	89.5
	Average Length	95%	3.45×10^{-2}
		99%	4.71×10^{-2}
	Variance of Length	95%	9.27×10^{-4}
		99%	1.73×10^{-3}

Error Distributions

A Normal : Mean=0.0, S.D.=0.20, S.D.=0.15, S.D.=0.10, S.D.=0.5;

B Exponential : Mean=0.20, 0.15, 0.15, 0.05;

C Gamma : Mean=0.20, 0.15, 0.10, 0.05 (k=2)

D Long tailed h : Mean=0.20, 0.15, 0.10, 0.05; (h=0.2)

Step 1 : Calculate the least squares estimators $\hat{\theta}_1$ and $\hat{\theta}_2$ using the linearisation method such that

$$\hat{\theta}_i = t(X_1, \dots, X_n; Y_1, \dots, Y_n) \quad \text{for } i=1,2 \quad (4.4.2)$$

Step 2 : Calculate the bootstrap pseudo estimators $\hat{\theta}_1^{*k}, \hat{\theta}_2^{*k}, k=1, B$ ($B=120$ say), in the normal manner (see Section 2.4), from the model:

$$X_j^* = g_i(\hat{\theta}_i) + \epsilon_j^*, \quad \text{for } j=1, n; i=1,2 \quad (4.4.3)$$

where the bootstrap residuals $\epsilon^* = (\epsilon_1^*, \dots, \epsilon_n^*)$ are constructed by drawing with replacement from the observed residuals $\hat{\epsilon}_1, \dots, \hat{\epsilon}_n$.

Step 3 : For each bootstrap sample calculate

$$\hat{\theta}_{i,-j}^* = t(X_1^*, \dots, X_{j-1}^*, X_{j+1}^*, \dots, X_n^*), \quad \text{for } j=1, n; i=1,2 \quad (4.4.4)$$

then, the estimator of the influence function

$$\hat{I}_{ij}^* = (n-1) (\hat{\theta}_i^{*k} - \hat{\theta}_{i,-j}^*) \quad \text{for } k=1, B \quad (4.4.5)$$

The bootstrap estimator

$$\hat{\theta}_i^* = \frac{\sum_{k=1}^B \hat{\theta}_i^{*k}}{B} \quad (4.4.6)$$

and the estimator of the variance of $\hat{\theta}_i^*$,

$$\text{var}(\hat{\theta}_i^*) = \frac{\sum_{j=1}^n \hat{I}_{ij}^{*2}}{n^2} \quad (4.4.7)$$

The studentized statistic T_k^i can then be calculated as follows:

$$T_k^i = \frac{\hat{\theta}_i^k - \hat{\theta}_i}{\sqrt{\text{var}(\hat{\theta}_i^k)}} , \text{ for } K=1, \dots, B \quad \left. \vphantom{\frac{\hat{\theta}_i^k - \hat{\theta}_i}{\sqrt{\text{var}(\hat{\theta}_i^k)}}} \right\}_{i=1,2}$$

Step 4 : Rank the T_k^i 's in ascending order and calculate the T_k^i value at $(1-\alpha/2)\%$ and $\alpha/2\%$ levels. i.e. $T_{(1-\alpha/2)}^i$ and $T_{(\alpha/2)}^i$.

Step 5 : Calculate the variance estimator of $\hat{\theta}_i$ using estimators of the first order influence functions of the functionals $\hat{\theta}_i, i=1,2$ where $\hat{\theta}_i$ is the least squares estimator.

$$\begin{aligned} \text{i.e. } \hat{\theta}_i &= t(X_1, \dots, X_n; Y_1, \dots, Y_n) \\ \hat{\theta}_{i,-j} &= t(X_1, \dots, X_{j-1}, X_{j+1}, \dots, X_n; Y_1, \dots, Y_{j-1}, Y_{j+1}, \dots, Y_n) \end{aligned} \quad (4.4.9)$$

The estimator of the influence function

$$\hat{I}_{ij} = (n-1)(\hat{\theta}_i - \hat{\theta}_{i,-j}) \quad (4.4.10)$$

Hence,

$$\text{var}(\hat{\theta}_i) = \frac{\sum_{j=1}^n \hat{I}_{ij}^2}{n^2} \quad (4.4.11)$$

Step 6 : Calculate the confidence intervals for $\theta_i, i=1,2$

$$\text{i.e. } (\hat{\theta}_i - T_{(1-\alpha/2)}^i \sqrt{\text{var} \hat{\theta}_i} , \hat{\theta}_i - T_{(\alpha/2)}^i \sqrt{\text{var} \hat{\theta}_i})$$

The different methods were compared in terms of coverage of the parameters and the simulations performed were similar to those carried out in the first study except that the non-linear model considered by Frangos (1981), was used instead. Also, due to the large C.P.U. time required for some of these iterative methods, the number of simulation

runs was reduced to 100 in each case. The results of the simulation study are summarised in Tables 4.4.6 and 4.4.7.

The conclusions of the study are as follows:

- a. the jackknife and studentized pivotal quantity method consistently produce robust coverage results for both parameters. There are some cases where the studentized pivotal quantity method outperforms the jackknife method in terms of coverage, shorter average length and smaller variance in length, although, in general, both methods produce very similar results.
- b. the other methods, including the least squares, produce poor coverage results where the error deviates are either normal or exponentially distributed. The least squares method also produces poor coverage results for the parameter θ_1 in the case of the long tailed h distribution.

4.5 Summary of simulation results

The following conclusions can be made regarding the results presented in the previous section:

(points (i), (ii) and (iii) refer to the studies where the variance of the errors is constant).

- (i) the least squares and jackknife method produced similar results in terms of bias reduction although for the latter method this was at the expense of a larger mean square error. The standard bootstrap method did not improve on the least squares method in terms of bias or mean square error.
- (ii) for the two different two parameter non-linear models considered, the jackknife produced consistently robust coverage results. The least squares method produced poor coverage results for Frangos' (1981), model for most of the different error distributions considered. The standard bootstrap method as well as Efron's 'improved' bootstrap confidence limit methods generally only produced modest coverage results.

Table 4.4.6 : Sample Confidence Intervals for a Two Parameter Non-Linear Model

Parameter θ_1		Least Squares	Jackknife	Bootstrap	Percentile	Bias Corrected	Improved Bootstrap (using acceleration constant)	Studentised Pivotal Quantities
A	Coverage	95%	74.0	94.0	83.0	84.0	83.0	87.0
		99%	93.0	100.0	94.0	93.0	93.0	92.0
	Average Length	95%	7.57×10^{-2}	1.53×10^{-1}	1.14×10^{-1}	1.13×10^{-1}	1.13×10^{-1}	1.14×10^{-1}
		99%	1.03×10^{-1}	2.09×10^{-1}	1.50×10^{-1}	1.49×10^{-1}	1.44×10^{-1}	1.39×10^{-1}
	Variance in Length	95%	4.39×10^{-4}	1.48×10^{-3}	4.79×10^{-4}	5.06×10^{-4}	5.20×10^{-4}	5.71×10^{-4}
		99%	8.20×10^{-4}	2.76×10^{-3}	8.28×10^{-4}	1.06×10^{-3}	9.81×10^{-4}	8.75×10^{-4}
	Coverage	95%	80.0	94.0	84.0	86.0	86.0	88.0
		99%	90.0	99.0	91.0	92.0	92.0	90.0
B	Average Length	95%	9.85×10^{-2}	1.53×10^{-1}	1.14×10^{-1}	1.16×10^{-1}	1.15×10^{-1}	1.17×10^{-1}
		99%	1.35×10^{-1}	2.09×10^{-1}	1.50×10^{-1}	1.50×10^{-1}	1.45×10^{-1}	1.40×10^{-1}
	Variance in Length	95%	1.22×10^{-3}	3.27×10^{-3}	1.39×10^{-3}	1.65×10^{-3}	1.54×10^{-3}	1.73×10^{-3}
		99%	2.29×10^{-3}	6.11×10^{-3}	2.40×10^{-3}	2.71×10^{-3}	2.46×10^{-3}	2.55×10^{-3}
	Coverage	95%	93.0	98.0	94.0	93.0	95.0	93.0
		99%	96.0	100.0	97.0	97.0	97.0	97.0
C	Average Length	95%	8.52×10^{-2}	1.07×10^{-1}	7.79×10^{-2}	7.74×10^{-2}	7.77×10^{-2}	7.80×10^{-2}
		99%	1.16×10^{-1}	1.46×10^{-1}	1.02×10^{-1}	1.02×10^{-1}	9.76×10^{-2}	9.39×10^{-2}
	Variance in Length	95%	5.56×10^{-4}	1.19×10^{-3}	3.88×10^{-4}	3.98×10^{-4}	4.11×10^{-4}	3.90×10^{-4}
		99%	1.04×10^{-3}	2.22×10^{-3}	6.69×10^{-4}	8.30×10^{-4}	7.53×10^{-4}	6.27×10^{-4}
	Coverage	95%	88.0	99.0	95.0	95.0	94.0	97.0
		99%	95.0	100.0	99.0	99.0	99.0	99.0
D	Average Length	95%	1.27×10^{-1}	2.02×10^{-1}	1.72×10^{-1}	1.76×10^{-1}	1.77×10^{-1}	1.82×10^{-1}
		99%	1.73×10^{-1}	2.76×10^{-1}	2.26×10^{-1}	2.39×10^{-1}	2.29×10^{-1}	2.23×10^{-1}
	Variance Length	95%	9.16×10^{-3}	1.74×10^{-2}	2.30×10^{-2}	2.74×10^{-2}	2.94×10^{-2}	3.04×10^{-2}
		99%	1.71×10^{-2}	3.25×10^{-2}	3.97×10^{-2}	5.33×10^{-2}	4.58×10^{-2}	4.61×10^{-2}
	Coverage	95%	88.0	99.0	95.0	95.0	94.0	97.0
		99%	95.0	100.0	99.0	99.0	99.0	99.0

Error Distributions

A Normal : Mean=0.0, S.D.=0.2

B Exponential : Mean=0.2;

C Gamma : Mean=0.2 (k=2)

D Long tailed h : Mean=0.2 (h=0.2)

Table 4.4.6 : Sample Confidence Intervals for a Two Parameter Non-Linear Model

	Parameter θ_1		Least Squares	Jackknife	Bootstrap	Percentile	Bias Corrected	Improved Bootstrap (using acceleration constant)	Studentised Pivotal Quantities
	Coverage								
A	Coverage	95%	74.0	94.0	83.0	84.0	83.0	87.0	95.0
		99%	93.0	100.0	94.0	93.0	93.0	92.0	97.0
	Average Length	95%	7.57×10^{-2}	1.53×10^{-1}	1.14×10^{-1}	1.13×10^{-1}	1.13×10^{-1}	1.14×10^{-1}	1.40×10^{-1}
		99%	1.03×10^{-1}	2.09×10^{-1}	1.50×10^{-1}	1.49×10^{-1}	1.44×10^{-1}	1.39×10^{-1}	2.08×10^{-1}
	Variance in Length	95%	4.39×10^{-4}	1.48×10^{-3}	4.79×10^{-4}	5.06×10^{-4}	5.20×10^{-4}	5.71×10^{-4}	1.41×10^{-3}
		99%	8.20×10^{-4}	2.76×10^{-3}	8.28×10^{-4}	1.06×10^{-3}	9.81×10^{-4}	8.75×10^{-4}	4.34×10^{-3}
B	Coverage	95%	80.0	94.0	84.0	86.0	86.0	88.0	93.0
		99%	90.0	99.0	91.0	92.0	92.0	90.0	98.0
	Average Length	95%	9.85×10^{-2}	1.53×10^{-1}	1.14×10^{-1}	1.16×10^{-1}	1.15×10^{-1}	1.17×10^{-1}	1.37×10^{-1}
		99%	1.35×10^{-1}	2.09×10^{-1}	1.50×10^{-1}	1.50×10^{-1}	1.45×10^{-1}	1.40×10^{-1}	1.96×10^{-1}
	Variance in Length	95%	1.22×10^{-3}	3.27×10^{-3}	1.39×10^{-3}	1.65×10^{-3}	1.54×10^{-3}	1.73×10^{-3}	2.65×10^{-3}
		99%	2.29×10^{-3}	6.11×10^{-3}	2.40×10^{-3}	2.71×10^{-3}	2.46×10^{-3}	2.55×10^{-3}	8.36×10^{-3}
C	Coverage	95%	93.0	98.0	94.0	93.0	95.0	93.0	98.0
		99%	96.0	100.0	97.0	97.0	97.0	97.0	99.0
	Average Length	95%	8.52×10^{-2}	1.07×10^{-1}	7.79×10^{-2}	7.74×10^{-2}	7.77×10^{-2}	7.80×10^{-2}	9.96×10^{-2}
		99%	1.16×10^{-1}	1.46×10^{-1}	1.02×10^{-1}	1.02×10^{-1}	9.76×10^{-2}	9.39×10^{-2}	1.40×10^{-1}
	Variance in Length	95%	5.56×10^{-4}	1.19×10^{-3}	3.88×10^{-4}	3.98×10^{-4}	4.11×10^{-4}	3.90×10^{-4}	1.08×10^{-3}
		99%	1.04×10^{-3}	2.22×10^{-3}	6.69×10^{-4}	8.30×10^{-4}	7.53×10^{-4}	6.27×10^{-4}	2.38×10^{-3}
D	Coverage	95%	88.0	99.0	95.0	95.0	94.0	97.0	97.0
		99%	95.0	100.0	99.0	99.0	99.0	99.0	100.0
	Average Length	95%	1.27×10^{-1}	2.02×10^{-1}	1.72×10^{-1}	1.76×10^{-1}	1.77×10^{-1}	1.82×10^{-1}	1.78×10^{-1}
		99%	1.73×10^{-1}	2.76×10^{-1}	2.26×10^{-1}	2.39×10^{-1}	2.29×10^{-1}	2.23×10^{-1}	2.51×10^{-1}
	Variance in Length	95%	9.16×10^{-3}	1.74×10^{-2}	2.30×10^{-2}	2.74×10^{-2}	2.94×10^{-2}	3.04×10^{-2}	1.34×10^{-2}
		99%	1.71×10^{-2}	3.25×10^{-2}	3.97×10^{-2}	5.33×10^{-2}	4.58×10^{-2}	4.61×10^{-2}	2.77×10^{-2}

Error Distributions

A Normal : Mean=0.0, S.D.=0.2

B Exponential : Mean=0.2;

C Gamma : Mean=0.2 (k=2)

D Long tailed h : Mean=0.2 (h=0.2)

Table 4.4.7 : Sample Confidence Intervals for a Two Parameter Non-Linear Model

	Parameter- θ_2		Least Squares	Jackknife	Bootstrap	Percentile	Bias Corrected	Improved Bootstrap (Using acceleration constant)	Studentised Pivotal Quantities
A	Coverage	95%	83.0	93.0	84.0	84.0	83.0	87.0	94.0
		99%	95.0	99.0	93.0	91.0	91.0	96.0	96.0
	Average Length	95%	2.08×10^{-2}	3.90×10^{-2}	2.85×10^{-2}	2.81×10^{-1}	2.80×10^{-2}	2.84×10^{-2}	3.61×10^{-2}
		99%	2.85×10^{-2}	5.33×10^{-2}	3.74×10^{-2}	3.69×10^{-2}	3.56×10^{-2}	3.43×10^{-2}	5.33×10^{-2}
	Variance in Length	95%	3.16×10^{-5}	1.20×10^{-4}	2.80×10^{-5}	2.90×10^{-5}	3.00×10^{-5}	3.50×10^{-5}	1.20×10^{-2}
		99%	5.90×10^{-5}	2.24×10^{-4}	4.80×10^{-5}	6.20×10^{-5}	5.60×10^{-5}	5.20×10^{-5}	3.55×10^{-4}
B	Coverage	95%	83.0	94.0	84.0	87.0	87.0	87.0	93.0
		99%	91.0	99.0	91.0	91.0	91.0	89.0	98.0
	Average Length	95%	2.74×10^{-2}	3.88×10^{-2}	2.84×10^{-2}	2.89×10^{-2}	2.89×10^{-2}	2.90×10^{-2}	3.48×10^{-2}
		99%	3.74×10^{-2}	5.31×10^{-2}	3.74×10^{-2}	3.74×10^{-2}	3.61×10^{-2}	3.47×10^{-2}	5.05×10^{-2}
	Variance in Length	95%	1.11×10^{-4}	2.39×10^{-4}	8.50×10^{-5}	1.00×10^{-4}	9.70×10^{-5}	9.90×10^{-5}	1.96×10^{-4}
		99%	2.07×10^{-4}	4.47×10^{-4}	1.47×10^{-4}	1.72×10^{-4}	1.45×10^{-4}	1.52×10^{-4}	6.71×10^{-4}
C	Coverage	95%	94.0	98.0	94.0	93.0	95.0	95.0	97.0
		99%	97.0	100.0	97.0	97.0	97.0	97.0	99.0
	Average Length	95%	2.36×10^{-2}	2.73×10^{-2}	1.96×10^{-2}	1.92×10^{-2}	1.94×10^{-2}	1.94×10^{-2}	2.56×10^{-2}
		99%	3.22×10^{-2}	3.73×10^{-2}	2.57×10^{-2}	2.54×10^{-2}	2.46×10^{-2}	2.34×10^{-2}	3.66×10^{-2}
	Variance in Length	95%	4.77×10^{-5}	9.34×10^{-5}	2.40×10^{-5}	2.50×10^{-5}	2.50×10^{-5}	2.50×10^{-5}	8.60×10^{-5}
		99%	8.91×10^{-5}	1.75×10^{-4}	4.10×10^{-5}	5.20×10^{-5}	5.20×10^{-5}	3.70×10^{-5}	1.78×10^{-4}
D	Coverage	95%	92.0	99.0	95.0	95.0	92.0	96.0	97.0
		99%	97.0	100.0	99.0	99.0	99.0	99.0	100.0
	Average Length	95%	3.50×10^{-2}	5.11×10^{-2}	4.28×10^{-2}	4.38×10^{-2}	4.41×10^{-2}	4.41×10^{-2}	4.51×10^{-2}
		99%	4.79×10^{-2}	6.99×10^{-2}	5.62×10^{-2}	5.95×10^{-2}	5.77×10^{-2}	5.54×10^{-2}	6.38×10^{-2}
	Variance in Length	95%	6.62×10^{-4}	1.22×10^{-3}	1.29×10^{-3}	1.52×10^{-3}	1.53×10^{-3}	1.53×10^{-3}	8.53×10^{-4}
		99%	1.24×10^{-3}	2.27×10^{-3}	2.22×10^{-3}	2.96×10^{-3}	2.93×10^{-3}	2.91×10^{-3}	1.51×10^{-3}

Error Distributions

A Normal : Mean=0.0, S.D.=0.2

B Exponential : Mean=0.2;

C Gamma : Mean=0.2 (k=2)

D Long tailed h : Mean=0.2 (h=0.2)

- (iii) the studentized pivotal quantity method produced very promising coverage results and the method was comparable to that of the jackknife. In this case, although the simulation study was rather limited due to the excessive computer C.P.U. time required for the method, there was sufficient evidence to suggest that the studentized pivotal quantity method might be a realistic competitor to the jackknife method.
(This is a very important result and it suggests that further application of the studentized pivotal quantity method in non-linear confidence interval estimation would be a very worthwhile research topic - see Section 2.11).
- (iv) for the one and two parameter non-linear models considered, when the variance of the errors is not constant, the jackknife outperformed the least squares method in terms of coverage. The jackknife also produced a smaller bias for the parameters although this was at the expense of a larger mean square error.

CHAPTER 5

NON-PARAMETRIC ESTIMATION IN TIME-SERIES MODELS

5.1 Introduction

A discrete univariate time-series is a sequence of observations ordered and equally spaced in time. Any continuous time-series may be converted into a discrete one by sampling at equal time intervals. There are two major approaches to the study of time-series processes, namely time-domain and frequency-domain analysis. The former may be considered as an approach generally called Box-Jenkins modelling and the latter as an approach involving two techniques - periodogram analysis and spectrum analysis.

Generally speaking, time-series analysis has an important role to play in many practical situations which usually involves some type of forecasting activity. In Government as well as industry, time series analysis has a major contribution to make, in terms of planning and decision taking. An example of the consequences of poor forecasting involved the spot price for chartering oil tankers for the four year period 1968-71.

Choosing the correct time-series model and deriving accurate estimators of the parameter value is thus essential in producing accurate forecasted results. This latter area will be discussed in some detail in section 5.5 of this chapter.

5.2 Models for Time-Series Data

Time-series models which were formalised by Box and Jenkins (1970, 1976) can be categorised into three classes.

- (i) autoregressive (AR) processes
- (ii) moving average (MA) processes
- (iii) mixed AR-MA processes (ARMA)

Consider an equation of the form

$$Y_t = a + b_1 Y_{t-1} + \dots + b_k Y_{t-k} + e_t, \quad (5.2.1)$$

where Y_t = value of dependent variables at time t

and e_t = error term at time t

The variables $Y_{t-1}, \dots, Y_{t-k}$ are simply time-lagged values of the dependent variable, and therefore the name autoregression (AR) is used to describe equations of the form (5.2.1).

In a similar manner, equation (5.2.1) can be written in terms of past error terms, as follows :

$$Y_t = a + b_1 e_{t-1} + \dots + b_k e_{t-k} + e_t \quad (5.2.2)$$

Here, explicitly, a dependence relationship is set up among the successive error terms, and the equation is called a moving average (MA) model.

Autoregressive (AR) models can be effectively coupled with moving average (MA) models to form a very general and useful class of time-series models called autoregressive/moving average (ARMA) schemes or processes. ARMA models can also be integrated which refers to the 'differencing' of the data series. For example, consider the AR process

$$Y_t = Y_{t-1} + e_t \quad (5.2.3)$$

where the coefficient of Y_{t-1} is unity.

The equation can be rewritten as :

$$Y_t - Y_{t-1} = e_t \quad (5.2.4)$$

to show that the first differences of the Y_t series form a random model. It is often convenient to redefine $(Y_t - Y_{t-1})$ as W_t , the first-difference series, so that W_t can be considered as a stationary series, whereas Y_t is nonstationary. Consequently, ARMA models which include a differencing of the data series, are referred to as autoregressive integrated moving average (ARIMA) models. The three numbers following an ARIMA model i.e. ARIMA(p,d,q) refer to the degree of the AR process, the degree of differencing, and the degree of the MA process.

5.3 Estimation of Parameters

Having a tentative model identification, the AR and MA parameters have to be determined in the best possible manner. For a simple AR model only, the standard least squares algorithm can be used to derive the parameter estimators. This method cannot be used for an MA model since the values of the e_i , $i=1, \dots, n$ terms are unknown. Generally speaking, for the more complicated ARMA models, the first step is to determine preliminary estimates and then to optimize these values using a non-linear, iterative search technique.

Preliminary estimators for an AR model can be derived by using the Yule-Walker equations, which are defined as follows :

$$\left. \begin{aligned} \rho_1 &= \phi_1 + \phi_2 \rho_1 + \dots + \phi_p \rho_{p-1} \\ &\vdots \\ &\vdots \\ &\vdots \\ \rho_p &= \phi_1 \rho_{p-1} + \phi_2 \rho_{p-2} + \dots + \phi_p \end{aligned} \right\}, \quad (5.3.1)$$

where $\rho_1, \dots, \rho_p$ are the theoretical autocorrelations for lags 1, 2, ..., p respectively, and $\phi_1, \phi_2, \dots, \phi_p$ are the p AR coefficients of the AR(p) process. Since the theoretical values of ρ are not known, they can be replaced by their empirical counterparts, and then solved for the ϕ_i , $i=1, \dots, p$ values. For example, considere an AR(2) process. Rewriting the Yule-Walker equations for $p=2$, yields:

$$\begin{aligned}\rho_1 &= \phi_1 + \phi_2 \rho_1 \\ \rho_2 &= \phi_1 \rho_1 + \phi_2\end{aligned}\tag{5.3.2}$$

Replacing ρ_1 and ρ_2 with the autocorrelation coefficients r_1 and r_2 and solving for ϕ_1 and ϕ_2 gives preliminary estimates:

$$\begin{aligned}\hat{\phi}_1 &= \frac{r_1(1-r_2)}{1-r_1^2}, \\ \text{and } \hat{\phi}_2 &= \frac{r_2 - r_1^2}{1-r_1^2}\end{aligned}\tag{5.3.3}$$

For an MA model, the theoretical autocorrelations for the process can be expressed in terms of the MA coefficients, as follows :

$$\rho_k = \begin{cases} \frac{-\theta_k + \theta_1\theta_{k+1} + \dots + \theta_{q-k}\theta_q}{1 - \theta_1^2 - \theta_2^2 - \dots - \theta_q^2}, & k=1, \dots, q \\ 0 & k > q \end{cases}\tag{5.3.4}$$

Since the theoretical values, ρ_k , are unknown, preliminary estimators of the coefficients $\theta_1, \theta_2, \dots, \theta_q$ can be obtained by substituting empirical autocorrelations, r_k , in equation (5.3.4), and solving. However, these equations are by no means easy to solve. In the case of a MA(2) model, Box and Jenkins (1976), offer tables and charts to handle preliminary estimators for θ_1 and θ_2 .

Having obtained preliminary estimators for the AR and MA parameters, the next stage is to use an iterative search method such as the Marquardt compromise (1963), to select new coefficients which produce a smaller 'sum of squared residuals' (SSR). In choosing new coefficients, this method not only moves in the right direction towards a smaller SSR, but also chooses a good coefficient correction size, to ensure rapid convergence to a minimum SSR.

5.4 Jackknifing in Time-Series

The application of the jackknife does cause some difficulty in time-series, due to the inter-dependence between successive values. The first application of the jackknife method in time-series was the jackknife estimator which was obtained by dividing a sequence of observations into two groups. This was considered by Quenouille (1949), and was the original application of the jackknife method. In particular, Quenouille jackknifed the partial correlation coefficient of a series of observations. He stated that, if r is the partial correlation coefficient calculated from a series, and r_1 and r_2 are the corresponding coefficients from the first and second halves of the series, then

$$R = 2r - \frac{1}{2}(r_1 + r_2) \quad (5.4.1)$$

will be free from bias to the order $1/n^2$. This result was validated using some simple numerical examples.

An interesting paper by Davis (1977), considers various methods for setting confidence intervals for the error variance in an ARMA model, which are robust against non-normality. As suggested by Box (1953), the fourth moment must be studentized for inferences about variances just as the second moment is studentized when making inferences about means. Davis considered a standard error procedure which was a moment estimator of the kurtosis, an analogue of Box's (1953) simple split data and the jackknife technique, and showed that all these methods for studentizing the kurtosis in the i.i.d. case were valid for a non-normal ARMA time-series model.

A further paper by Davis (1979), studies robust methods for detecting a shift in variance at a specified point of a time series. The likelihood ratio test of equality of variance is the F test, but the test statistic is dependent on the assumption that the data are normally distributed. Shorack (1969), for the i.i.d. case, compared several robust criteria on the basis of Pitman asymptotic relative efficiency (a.r.e.) and Monte

Carlo studies of power functions. The Box-Andersen permutation test (1955), and the jackknife procedure were found to be the most satisfactory.

Davis considered robust procedures for detecting a variance shift when the observations were correlated. He found that analogs of the Box-Anderson procedure and the jackknife had identical asymptotic properties to the i.i.d. case. Uniform, normal and double exponential distributions were considered in the Monte-Carlo studies, and while the Box-Anderson procedure did a better job maintaining the stated significance level, it lagged slightly behind the jackknife in power.

5.5 Monte Carlo Simulation Studies

This section describes the simulation studies that were undertaken in new areas of research regarding the non-parametric estimation of parameters for stationary autoregressive models of order one and two respectively.

$$\text{i.e. } Y_t = \phi_1 Y_{t-1} + \mu + e_t \quad \sim \quad \text{ARMA}(1,0,0) \quad (5.5.1)$$

$$\text{and } Y_t = \phi_1 Y_{t-1} + \phi_2 Y_{t-2} + \mu + e_t \quad \sim \quad \text{ARMA}(2,0,0) \quad (5.5.2)$$

where

Y_t = value of the dependent variable at time t

e_t = error term at time t

μ = mean value

and ϕ_1, ϕ_2 are the parameters

Four different p.d.f.s. were used to generate the error deviates for each simulation run, namely the normal, exponential, gamma and long tailed h distribution. In these cases, the e 's satisfy the assumption of serial independence. Other simulation runs were also carried out where the error term was serially correlated.

i.e. consider the simple model

$$Y_t = \phi_1 Y_{t-1} + \mu + e_t \quad (5.5.3)$$

where $u_t = \alpha \mu_{t-1} + e_t$

and $|\alpha| < 1$

(The e 's are defined in the same manner as before)

For each simulation run, 1000 samples were generated. The accuracy of the parameter estimators was expressed in terms of bias, sample variance and mean square error. The 95% and 99% levels were considered for the confidence limits and their robustness was determined in terms of coverage, average length and variance in length. It should be noted that, for the jackknife method, the number of groups $g=2$, due to the inter-dependence between successive values.

The following simulation studies were undertaken:

(i) a comparison between the least squares method and the jackknife and bootstrap methods in the one and two parameter linear model using either normal, exponential, gamma or long tailed h distributions for the error deviates. A sample size of 30 and 60 was considered for the one and two parameter linear models respectively. The results of the simulation study are summarised in Tables 5.5.1-5.5.4.

The conclusions of the study are as follows:

- a. the jackknife estimator has the smallest bias in all cases although the least square estimator is better than either the jackknife or bootstrap estimator in terms of mean square error
- b. the jackknife method produces untrustworthy confidence limits in terms of coverage. This result is expected due to the restriction of using only two different groups i.e. $g=2$, in the method
- c. the least squares method outperforms the bootstrap method based on the criteria of better coverage, shorter average lengths and smaller variances of length,

(the bootstrap method was not considered in the two parameter linear models due to the excessive C.P.U. time required and also the generally disappointing results which were obtained in the one parameter linear model).

**Table 5.5.1 : Bias, Sample Variance and Mean Square Error for a One
Parameter Linear Model**

		Least Squares Method	Jackknife Method	Bootstrap Method
A	Bias	-0.0796	0.0062	-0.1555
	Sample Variance	2.89×10^{-2}	4.58×10^{-2}	2.34×10^{-2}
	Mean Square Error	3.52×10^{-2}	4.58×10^{-2}	4.75×10^{-2}
B	Bias	-0.0656	0.0088	-0.1292
	Sample Variance	2.19×10^{-2}	3.63×10^{-2}	1.97×10^{-2}
	Mean Square Error	2.62×10^{-2}	3.63×10^{-2}	3.64×10^{-2}
C	Bias	-0.0714	0.0039	-0.1257
	Sample Variance	2.22×10^{-2}	3.43×10^{-2}	2.12×10^{-2}
	Mean Square Error	2.70×10^{-2}	3.43×10^{-2}	3.69×10^{-2}
D	Bias	-0.0629	0.0017	-0.1297
	Sample Variance	1.72×10^{-2}	5.92×10^{-2}	1.51×10^{-2}
	Mean Square Error	2.12×10^{-2}	5.92×10^{-2}	3.19×10^{-2}

Error Distributions

A Normal : Mean = 0.0, S.D = 0.5

B Exponential : Mean = 0.5

C Gamma : Mean = 0.5 (k=2)

D Long tailed h : Mean = 0.5 (h=0.3)

Starting Value

$Y_0 = 0.1$

Sample Size

N=30

Mean

$\mu = 0.7$

Table 5.5.2 : Sample Confidence Intervals for a One Parameter Linear

Model

		Least Squares Method	Jackknife Method	Bootstrap Method
A	Coverage	95%	93.7	88.3
		99%	97.6	96.2
	Average Length	95%	6.57×10^{-1}	5.54×10^{-1}
		99%	8.63×10^{-1}	7.28×10^{-1}
	Variance of Length	95%	4.49×10^{-3}	1.71×10^{-1}
		99%	7.75×10^{-3}	2.96×10^{-1}
B	Coverage	95%	95.1	92.5
		99%	98.1	97.7
	Average of Lenth	95%	6.19×10^{-1}	5.03×10^{-1}
		99%	8.13×10^{-1}	6.61×10^{-1}
	Variance of Length	95%	5.50×10^{-3}	1.59×10^{-1}
		99%	9.51×10^{-3}	2.74×10^{-1}
C	Coverage	95%	93.6	90.7
		99%	98.3	97.4
	Average Length	95%	5.76×10^{-1}	5.01×10^{-1}
		99%	7.57×10^{-1}	6.59×10^{-1}
	Variance of Length	95%	5.10×10^{-3}	1.61×10^{-1}
		99%	8.81×10^{-3}	2.79×10^{-1}
D	Coverage	95%	97.2	94.6
		99%	99.1	98.5
	Average Length	95%	6.48×10^{-1}	5.03×10^{-1}
		99%	8.52×10^{-1}	6.61×10^{-1}
	Variance of Length	95%	2.55×10^{-2}	5.01×10^{-1}
		99%	4.40×10^{-2}	8.64×10^{-1}

Error Distributions

A Normal : Mean=0.0, S.D=0.5

B Exponential : Mean=0.5

C Gamma : Mean=0.5 (k=2)

D Long tailed h : Mean=0.5 (h=0.3)

Starting Value

$Y_0=0.1$

Sample Size

N=30

Mean

$\mu=0.7$

**Table 5.5.3 : Bias, Sample Variance and Mean Square Error for a Two
Parameter Linear Model**

		Least Squares Method		Jackknife Method	
		ϕ_1	ϕ_2	ϕ_1	ϕ_2
A	Bias	-0.0313	-0.0469	0.0092	-0.0014
	Sample Variance	1.74×10^{-2}	1.52×10^{-2}	2.20×10^{-2}	2.02×10^{-2}
	Mean Square Error	1.84×10^{-2}	1.74×10^{-2}	2.20×10^{-2}	2.02×10^{-2}
B	Bias	-0.0342	-0.0549	0.0101	-0.0048
	Sample Variance	1.74×10^{-2}	1.57×10^{-2}	2.39×10^{-2}	2.24×10^{-2}
	Mean Square Error	1.85×10^{-2}	1.87×10^{-2}	2.40×10^{-2}	2.24×10^{-2}
C	Bias	-0.0313	-0.0501	0.0120	-0.0037
	Sample Variance	1.49×10^{-2}	1.40×10^{-2}	2.59×10^{-2}	2.31×10^{-2}
	Mean Square Error	1.59×10^{-2}	1.65×10^{-2}	2.60×10^{-2}	2.31×10^{-2}

Error Distributions

A Normal : Mean=0.0, SD=0.5

B Exponential : Mean=0.5

C Long tailed h : Mean=0.5 (h=0.3)

Starting Values

$Y_0=0.35$

$Y_{-1}=-0.15$

Sample Size

N=60

Mean

$\mu=-0.4$

Table 5.5.4 : Sample Confidence Intervals for a Two Parameter Linear

Model

		Least Squares Method		Jackknife Method	
		ϕ_1	ϕ_2	ϕ_1	ϕ_2
A	Coverage	95%	93.4	94.1	71.5
		99%	98.8	99.2	78.2
	Average Length	95%	5.11×10^{-1}	4.98×10^{-1}	4.55×10^{-1}
		99%	6.71×10^{-1}	6.55×10^{-1}	5.98×10^{-1}
	Variance of Length	95%	3.53×10^{-4}	5.14×10^{-4}	1.23×10^{-1}
		99%	6.10×10^{-4}	8.88×10^{-4}	2.12×10^{-1}
B	Coverage	95%	94.2	93.9	68.6
		99%	98.6	98.7	75.0
	Average Length	95%	5.15×10^{-1}	5.17×10^{-1}	4.50×10^{-1}
		99%	6.77×10^{-1}	6.79×10^{-1}	5.92×10^{-1}
	Variance of Length	95%	4.98×10^{-4}	6.93×10^{-4}	1.27×10^{-1}
		99%	8.60×10^{-4}	1.20×10^{-3}	2.20×10^{-1}
C	Coverage	95%	95.2	95.3	68.3
		99%	99.2	98.8	76.7
	Average Length	95%	5.16×10^{-1}	5.20×10^{-1}	4.16×10^{-1}
		99%	6.78×10^{-1}	6.83×10^{-1}	5.47×10^{-1}
	Variance of Length	95%	1.71×10^{-3}	3.49×10^{-3}	1.26×10^{-1}
		99%	2.96×10^{-3}	6.02×10^{-3}	2.17×10^{-1}

Error Distributions

A Normal : Mean=0.0, S.D=0.5

B Exponential : Mean=0.5

C Long tailed h : Mean=0.5 (h=0.3)

Starting Values

Sample Size

Mean

$Y_0=0.35$

N=60

$\mu=-0.4$

$Y_{-1}=0.15$

(ii) a comparison between the least squares method and the jackknife and bootstrap methods in the one parameter linear model where the error terms were serially correlated. For the bootstrap method, the number of bootstrap repetitions was varied from 20 up to 120 in steps of 20. The results of the simulation study are summarised in Tables 5.5.5-5.5.6 for the least squares and jackknife method and Tables 5.5.7-5.5.14 for the bootstrap method.

The conclusions of the study are as follows :

- a. the bootstrap method outperforms the least squares and jackknife methods in terms of bias and mean square error. Also the values of the bias and mean square error are relatively insensitive to the number of bootstrap repetitions that were considered
- b. the least squares and jackknife methods product untrustworthy confidence limits in terms of coverage
- c. the bootstrap methods produces good confidence limits in terms of coverage where the exponential and long tailed h distributions were considered for the error distributions, at the expense of slightly longer average lengths. The bootstrap produces only modest coverage in the case of the normal and gamma distributions but this coverage is still far better than that obtained by the other methods
- d. A large number of bootstrap iterations was not required to obtain the best results in terms of coverage. Only marginal, if any, improvement was attained by increasing the number of bootstrap iterations over 100.

(iii) the final simulation study was based on the previous study and concerned two non-parametric methods for estimating standard errors and confidence limits, related to the bootstrap method, namely the percentile method and the bias corrected percentile method. A description of these methods is given in Chapter 2. The bias corrected percentile method produced very disappointing coverage results and are not presented.

Table 5.5.5 : Bias, Sample Variance and Mean Square Error for a One

Parameter Linear Model

(Errors are serially correlated $\alpha=0.5$)

		Least Squares Method	Jackknife Method
	Bias	0.222	0.3115
	Sample Variance	1.42×10^{-2}	2.81×10^{-2}
	Mean Square Error	6.36×10^{-2}	1.25×10^{-1}
B	Bias	0.1951	0.2552
	Sample Variance	1.13×10^{-2}	2.30×10^{-2}
	Mean Square Error	4.94×10^{-2}	8.81×10^{-2}
C	Bias	0.1643	0.2114
	Sample Variance	1.10×10^{-2}	2.27×10^{-2}
	Mean Square Error	3.80×10^{-2}	6.74×10^{-2}
D	Bias	0.2200	0.2893
	Sample Variance	1.00×10^{-2}	3.46×10^{-2}
	Mean Square Error	5.84×10^{-2}	1.18×10^{-1}

Error Distributions

A Normal : Mean=0.0, S.D=0.5

B Exponential : Mean=0.5

C Gamma : Mean=0.5 (k=2)

D Long tailed h : Mean=0.5 (h=0.3)

Starting Value

$Y_0=0.1$

Sample Size

N=30

Mean

$\mu=0.7$

Table 5.5.6 : Sample Confidence Intervals for a One Parameter Linear

Model

(Errors are serially correlated - $\alpha=0.5$)

		Least Squares Method	Jackknife Method
A	Coverage	95%	54.0
		99%	72.0
	Average Length	95%	5.10×10^{-1}
		99%	6.70×10^{-1}
	Variance of Length	95%	7.22×10^{-3}
		99%	1.25×10^{-2}
B	Coverage	95%	52.6
		99%	71.9
	Average Length	95%	4.36×10^{-1}
		99%	5.72×10^{-1}
	Variance of Length	95%	5.63×10^{-3}
		99%	9.72×10^{-3}
C	Coverage	95%	53.2
		99%	72.8
	Average Length	95%	3.85×10^{-1}
		99%	5.06×10^{-1}
	Variance of Length	95%	4.39×10^{-3}
		99%	7.59×10^{-3}
D	Coverage	95%	50.9
		99%	75.3
	Average Length	95%	4.72×10^{-1}
		99%	6.21×10^{-1}
	Variance of Length	95%	1.40×10^{-2}
		99%	2.41×10^{-2}

Error Distributions

- A Normal : Mean=0.0, S.D=0.5
- B Exponential : Mean=0.5
- C Gamma : Mean=0.5 (k=2)
- D Long tailed h = Mean=0.5 (h=0.3)

Starting Value

$Y_0=0.1$

Sample Size

N=30

Mean

$\mu=0.7$

Table 5.5.7 : Bias, Sample Variance and Mean Square Error for a One Parameter Linear Model using the Bootstrap

Method

(Errors are serially correlated - $\alpha=0.5$)

	Number of Bootstrap Replications					
	20	40	60	80	100	120
Bias	0.1363	0.1361	0.1355	0.1351	0.1353	0.1351
Sample Variance	1.84×10^{-2}	1.79×10^{-2}	1.77×10^{-2}	1.76×10^{-2}	1.76×10^{-2}	1.74×10^{-2}
Mean Square Error	3.69×10^{-2}	3.64×10^{-2}	3.60×10^{-2}	3.58×10^{-2}	3.59×10^{-2}	3.56×10^{-2}

Error Distribution

Normal : Mean=0.0, S.D=0.5

Starting Value

$Y_0=0.1$

Sample Size

N=30

Mean

$\mu=0.7$

Table 5.5.8 : Sample Confidence Intervals for a One Paramter Linear Model using the Bootstrap Method

(Errors are serially correlated - $\alpha=0.5$)

		Number of Bootstrap Replications					
		20	40	60	80	100	120
Coverage	95%	72.2	74.1	74.3	74.5	74.4	73.6
	99%	79.0	79.5	80.0	80.2	80.0	80.6
Average Length	95%	5.17×10^{-1}	5.21×10^{-1}	5.21×10^{-1}	5.23×10^{-1}	5.25×10^{-1}	5.24×10^{-1}
	99%	6.80×10^{-1}	6.84×10^{-1}	6.85×10^{-1}	6.87×10^{-1}	6.89×10^{-1}	6.89×10^{-1}
Variance of Length	95%	2.99×10^{-1}	2.47×10^{-2}	2.35×10^{-2}	2.29×10^{-2}	2.26×10^{-2}	2.17×10^{-2}
	99%	5.16×10^{-2}	4.26×10^{-2}	4.05×10^{-2}	3.95×10^{-2}	3.91×10^{-2}	3.76×10^{-2}

Error Distribution

Normal : Mean=0.0, S.D.=0.5

Starting Value

$Y_0=0.1$

Sample Size

$N=30$

Mean

$\mu=0.7$

Author Angus Stuart Maxwell

Name of thesis A Comparison Of The Jackknife And Bootstrap Estimators In Linear Models With Reference To Production Models Used By Sasol. 1988

PUBLISHER:

University of the Witwatersrand, Johannesburg

©2013

LEGAL NOTICES:

Copyright Notice: All materials on the University of the Witwatersrand, Johannesburg Library website are protected by South African copyright law and may not be distributed, transmitted, displayed, or otherwise published in any format, without the prior written permission of the copyright owner.

Disclaimer and Terms of Use: Provided that you maintain all copyright and other notices contained therein, you may download material (one machine readable copy and one print copy per page) for your personal and/or educational non-commercial use only.

The University of the Witwatersrand, Johannesburg, is not responsible for any errors or omissions and excludes any and all liability for any errors in or omissions from the information on the Library website.