

AD-A144 284

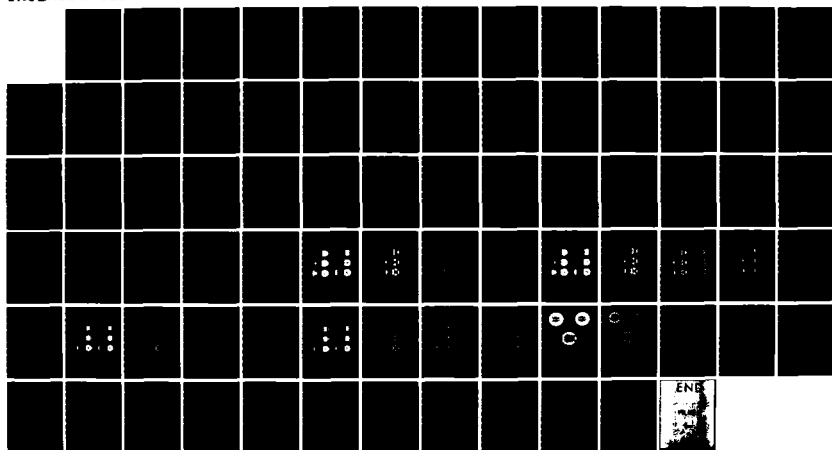
ON THE ESTIMATION OF THE FIRST-ORDER AUTOREGRESSIVE
PARAMETER(U) RHODE ISLAND UNIV KINGSTON DEPT OF
ELECTRICAL ENGINEERING F GIANNELLA ET AL. JUN 84
N00014-81-K-0144

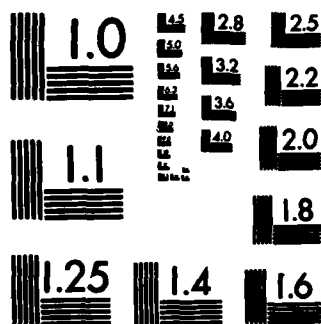
1/1

UNCLASSIFIED

F/G 12/1

NL





MICROCOPY RESOLUTION TEST CHART
NATIONAL BUREAU OF STANDARDS-1963-A

(12)

**ON THE ESTIMATION OF THE FIRST-ORDER
AUTOREGRESSIVE PARAMETER**

Fernando Giannella
Donald W. Tufts

Department of Electrical Engineering
University of Rhode Island
Kingston, Rhode Island 02881

June 1984

Prepared for

OFFICE OF NAVAL RESEARCH
Statistics and Probability Program
800 North Quincy St.
Arlington, Virginia 22217
under contract N00014-81-K-0144
D.W. Tufts, Principal Investigator

DTIC
ELECTE
AUG 8 1984
S
A

This document has been approved
for public release and sale; its
distribution is unlimited.

84 07 16 033

AD-A144 284

DTIC FILE COPY

REPORT DOCUMENTATION PAGE		READ INSTRUCTIONS BEFORE COMPLETING FORM	
1. REPORT NUMBER	2. GOVT ACCESSION NO. A144284	3. RECIPIENT'S CATALOG NUMBER April 1983 - June 1984	
4. TITLE (and Subtitle) ON THE ESTIMATION OF THE FIRST-ORDER AUTOREGRESSIVE PARAMETER		5. TYPE OF REPORT & PERIOD COVERED	
		6. PERFORMING ORG. REPORT NUMBER	
7. AUTHOR(s) F. Giannella D. W. Tufts		8. CONTRACT OR GRANT NUMBER(s) N00014-81-K-0144	
9. PERFORMING ORGANIZATION NAME AND ADDRESS Department of Electrical Engineering University of Rhode Island Kingston, Rhode Island 02881		10. PROGRAM ELEMENT, PROJECT, TASK AREA & WORK UNIT NUMBERS	
11. CONTROLLING OFFICE NAME AND ADDRESS Office of Naval Research Department of the Navy Arlington, VA 22217		12. REPORT DATE June 1984	
		13. NUMBER OF PAGES	
14. MONITORING AGENCY NAME & ADDRESS (if different from Controlling Office)		15. SECURITY CLASS. (of this report) Unclassified	
		15a. DECLASSIFICATION/DOWNGRADING SCHEDULE	
16. DISTRIBUTION STATEMENT (of this Report) This document has been approved for public release and sale; its distribution is unlimited.			
17. DISTRIBUTION STATEMENT (of the abstract entered in Block 20, if different from Report) Approved for public release; distribution unlimited			
18. SUPPLEMENTARY NOTES			
19. KEY WORDS (Continue on reverse side if necessary and identify by block number) Autoregressive process, parameter estimation, CR bounds, linear prediction, Prony's method, low-rank approximation, zero selection, modal zeros, noise zeros.			
20. ABSTRACT (Continue on reverse side if necessary and identify by block number) In this paper we present Cramer-Rao (C-R) bounds for parameter estimation of a first-order autoregressive (AR) process from a finite record of data. We use these bounds to evaluate the performance of Maximum Likelihood Estimation (MLE) and linear prediction approaches. Some estimates use low-rank approximation of an estimated covariance matrix. The latter estimates are based on the method of Tufts and Kumaresan [15]. Here we have added a zero selection technique in the last step of the procedure.			

Our low-rank, high order, linear prediction estimator performs better than the other methods which we have tested, when the pole is close to the unit circle. It is slightly biased and its variance is small and close to the variance given by the C-R bound for unbiased estimators. For a small number of samples (25 to 100) this estimator performs substantially better than the MLE.



Acquisition Rec	
THIS GRA&I	
TO TAB	
Unannounced	
Classification	
<i>Not in file</i>	
By	
Distribution	
Availability	
Dist	
A1	

1

ON THE ESTIMATION OF THE FIRST-ORDER
AUTOREGRESSIVE PARAMETER*

F. Giannella and D.W. Tufts

Department of Electrical Engineering

University of Rhode Island

Kingston, R.I. 02881

ABSTRACT

In this paper we present Cramer-Rao (C-R) bounds for parameter estimation of a first-order autoregressive (AR) process from a finite record of data. *The authors* We used these bounds to evaluate the performance of Maximum Likelihood Estimation (MLE) and linear prediction approaches. Some estimators use low-rank approximation of an estimated covariance matrix. The latter estimates are based on the method of Tufts and Kumaresan [15]. *In this document* ~~Here we have added~~ a zero selection technique in the last step of the procedure *was added.*

The Our low-rank, high order, linear prediction estimator performs better than the other methods which ~~we~~ *were* have tested, when the pole is close to the unit circle. It is slightly biased and its variance is small and close to the variance given by the C-R bound for unbiased estimators. For a small number of samples (25 to 100) this estimator performs substantially better than the MLE.

* This work was supported by the Probability and Statistics Program of the Office of Naval Research.

I. INTRODUCTION

For the past seven years we and our collaborators have been working on problems of parametric signal modeling. We have tried to find estimators of signal parameters which have the following properties : (1) they are unbiased or their biases are small compared with the standard deviation of error, and (2) their variances are close to the C-R lower-bound from high SNR to an SNR threshold that is as low as possible.

For the purpose of estimation of signal parameters we have introduced the use of low-rank approximation to data matrices or to estimated correlation or covariance matrices [1,9,10]. This technique has been applied in parameter estimation for sinusoids [2], exponentially damped sinusoids [3,4], and impulse response pole-zero identification [4], and also for direction finding and frequency finding using arrays of sensors [6,8]. Simplified calculations have also been studied [5,6,7].

During the last few years we have turned our attention to the problem of parameter estimation of random signals in the presence of noise for short records of data. Both autoregressive (AR) and autoregressive moving average (ARMA) models have been used [11-15]. We have benefited from earlier papers on system identification by Mehra [16], Astrom [17], and Parzen [18].

The modal estimation approach [11,12], which starts from the estimated correlation coefficients has been extended [15] to include the use of extra modes (poles) for modeling the effects of noise and

fluctuations. This leads to greater accuracy in estimating the signal modes. The resulting spurious modes can be suppressed using low rank approximation [15] or by subset selection techniques [5].

In this paper we more thoroughly evaluate the performance of the reduced-rank modal parameter estimator [15] for a first-order AR process.

We use a special procedure for selection of the signal zero from the set of zeros of the resulting prediction-error-filter polynomial.

The reduced-rank approximation can be carried out on the data matrix of the linear prediction equations, or equivalently on the estimated covariance matrix obtained from premultiplication by the transpose of this data matrix.

II. ANGLES BETWEEN TRUE AND ESTIMATED EIGENVECTORS

Recently, through perturbation analyses, which were motivated by the work of Wilkinson [19] and which have their roots in the work of Rayleigh, Timoshenko, and Courant, D. Tufts has traced the improvements of low rank approximation to the following fact:

The angles between the principal eigenvectors of a true covariance matrix and the corresponding principal eigenvectors of an estimated covariance tend to be very small, especially if the modal pole locations are close to the unit circle. We now present the results of measurements of some of these angles performed by I. Kirsteins using simulation of a first-order AR process.

The j 'th signal vector s_j generated by the j 'th independent realization of a first-order AR process, is described by the formulas

$$s_j = [s(1,j) \ s(2,j) \ \dots \ s(N,j)]^T \quad (1)$$

$$\begin{aligned} s(t,j) &= a \ s(t-1,j) + w(t,j) \quad ; \ t=1,2,\dots,N \quad (2) \\ &\quad ; \ j=1,2,\dots,K \end{aligned}$$

where $\{w(t,j)\}$ and $\{s(0,j)\}$ are mutually independent Gaussian random variables with zero mean and variances σ^2 and $\sigma^2/(1-a^2)$ respectively. Such vectors are segments of independent realizations (index j) of the AR process. The signal covariance matrix is estimated by

$$\hat{R} = \frac{1}{K} \sum_{j=1}^K s_j s_j^T \quad (3)$$

In our experiments we let $a = .9$, $N=10$ and $K=16$ or $K=32$.

The true correlation matrix, given by the expected value of (3), has the following eigen-decomposition :

$$R = E[\hat{R}] = V \Lambda V^T = \sum_{i=1}^N \lambda_i \underline{v}_i \underline{v}_i^T \quad (4)$$

where V contains the orthonormal eigenvectors $\{\underline{v}_i\}$ of R and Λ is a diagonal matrix containing the corresponding eigenvalues $\{\lambda_i\}$. We define the i 'th eigenvectors of R and $\hat{R}$ as $\underline{v}_i$ and $\hat{\underline{v}}_i$, respectively.

The angle in degrees between these two i 'th eigenvectors is given by

$$\phi_i = \frac{130}{\pi} \cos^{-1} | \underline{v}_i^T \hat{\underline{v}}_i | \quad (5)$$

Histograms of the absolute values of ϕ_1 and ϕ_2 for 500 simulation trials are presented in Figures 1a to 1d. Two values of K , the number of independent vector observations of the true AR sequence, are used. The results illustrate the angular stability of the estimated principal eigenvector $\hat{\underline{v}}_1$ when the value of (a) is close to unity.

III. EXACT C-R BOUNDS FOR A GAUSSIAN FIRST-ORDER AR PROCESS

In the context of autoregressive processes, as in the case of deterministic signals [1,10], we need C-R bounds for the variance of unbiased estimates of the parameters of interest, in order to have a standard of comparison for existing and proposed estimates. Until our work such bounds appear to have been available only for the asymptotic case of very long observations [20-23]. Here, we are interested in cases for which the observations are of short duration.

Suppose, as described above, that we observe K independent realizations of a set of N samples of a steady-state, first-order AR sequence. The j 'th realization $\{s(t,j)\}$ for $t=1,2,\dots,N$ is used to form the j 'th row of a matrix Y of data. We start with the joint-multivariate-normal-density function, named (f) , of the K -rows of Y , conditioned on the 2-parameter vector $\underline{\theta}$:

$$\underline{\theta} = [\theta_1, \theta_2]^T \quad (5)$$

in which θ_1 is the first-order AR parameter, denoted by (a) in formula (2), and θ_2 is σ , the standard deviation of the white-noise, $\{w(t,j)\}$, of (2).

We can then derive the elements of the Fisher Information Matrix associated with this process by the following steps [24]:

$$J_{ij} = E \left[- \frac{\partial^2}{\partial \theta_i \partial \theta_j} \ln(f) \right] \quad (7)$$

$$J_{ij} = \frac{-K}{2} \frac{\partial^2}{\partial \theta_i \partial \theta_j} [\ln(1 - \theta_1^2)] + K N \frac{\partial^2}{\partial \theta_i \partial \theta_j} [\ln(\theta_2)] + \frac{K}{2} \text{tr} \left[R \frac{\partial^2}{\partial \theta_i \partial \theta_j} (R)^{-1} \right] \quad (8)$$

After evaluation of these derivatives for $i=1,2$ and $j=1,2$ we can express the Fisher Information Matrix in the following form:

$$J = K \begin{bmatrix} \frac{1 + a^2 + (N-2)(1-a^2)}{(1-a^2)^2} & \frac{2a}{\sigma(1-a^2)} \\ \frac{2a}{\sigma(1-a^2)} & \frac{2N}{\sigma^2} \end{bmatrix} \quad (9)$$

The bounds for the variance of the estimated parameters are then determined from the diagonal of the inverse of J .

$$\text{var}(\hat{a}) \geq \begin{cases} \frac{(1-a^2)^2}{K[1+a^2+(N-2)(1-a^2)]} & ; \sigma \text{ known} \\ \frac{N(1-a^2)^2}{KN[1+a^2+(N-2)(1-a^2)] - 2Ka^2} & ; \sigma \text{ unknown} \end{cases} \quad (10)$$

$$\text{var}(\hat{\sigma}) \geq \begin{cases} \frac{\sigma^2}{2KN} & ; a \text{ known} \\ \frac{\sigma^2[1+a^2+(N-2)(1-a^2)]}{2KN[1+a^2+(N-2)(1-a^2)] - 4Ka^2} & ; a \text{ unknown} \end{cases} \quad (11)$$

We show in Figures (2a-2h) some curves of the above C-R bounds for $\text{var}(\hat{a})$.

The asymptotic bound [20] can of course be misleading for small-sample-size observations. We have added this bound in Figs. (2c) and (2d) only for a reference. One can verify from Fig.(2c) that as (a) tends to unity, the asymptotic bound is much larger than the exact bound. Fig. (2d) presents a case where the asymptotic bound is lower than the true ones (for σ known and unknown). Also it is interesting to note that as (a) tends to zero the asymptotic and the true bounds converge to $1/N-1$.

Figs. (2a,2b) show the dependence of the C-R bounds on (a) . They can help one to visualize a continuous transition between the decorrelated ($a=0$) and the highly correlated case ($a \rightarrow 1$).

In Figs (2e,2f) and (2g,2h) we show the bounds versus the number K of independent realizations for two different values of (a) for σ known and unknown.

IV. THE MAXIMUM LIKELIHOOD ESTIMATOR

For one realization ($K=1$) and N data samples of a first order AR sequence, the log-likelihood function is:

$$L = \frac{1}{2} \ln(1 - a^2) - N \ln(\sigma) - \frac{1}{2\sigma^2} (S_{00} - 2a S_{01} + a^2 S_{11}) \quad (12)$$

where:

$$S_{00} = \sum_{t=1}^N s_t^2 ; \quad S_{01} = \sum_{t=1}^{N-1} s_t s_{t+1} ; \quad S_{11} = \sum_{t=2}^{N-1} s_t^2 \quad (13)$$

Therefore, the value of (a) which maximizes (12), i.e. the Maximum Likelihood Estimator of (a), is a solution of the third order equation :

$$c^3 - \frac{S_{01}(N-2)}{S_{11}(N-1)} c^2 - \frac{(S_{00} + N S_{11})}{S_{11}(N-1)} c + \frac{N S_{01}}{S_{11}(N-1)} = 0 \quad (14)$$

The maximum likelihood estimate is the real solution of (14) of modulus less than the unity, which maximizes (12).

In Figs. 3a-3c, 4a-4c, we show the curves of estimated rms error, variance and bias of the MLE of (a) which have been obtained by simulations using 500 independent realizations. Comparison with the C-R bound curves in those figures show that for small sample sizes and values of (a) near unity, one might be able to obtain better performance than that of the MLE. In the next section we present a method for doing this.

V. AN SVD-BASED ESTIMATOR FOR THE A-R COEFFICIENT OF A FIRST-ORDER A-R PROCESS

For the special problem of estimating parameters of damped or undamped sinusoids in noise, some techniques based on low-rank approximation via Singular Value Decomposition (SVD) have been shown to possess nearly optimal properties [2],[4]. That is, the standard deviation of the estimation error, σ_e , is very close to the value given by the C-R bound. And the bias is small compared with the value of σ_e .

One of the conditions for good performance of these techniques is a rank deficiency of the signal-alone data matrix or correlation matrix.

Processes other than those of deterministic type can possess a near-to-rank-deficiency of the correlation matrix. That is, after a few large eigenvalues, the eigenvalue spectrum rapidly falls to low values (see Figs. 5a,5b,5c). We can then define subsets of random processes, possessing only a few high-variance components in their Karhunen-Loève expansions over finite data intervals.

As an example of such a process, let us consider a discrete, first-order AR process, the pole of which is located on the positive real line, near the unit circle.

Let's consider the eigenvector outer-product decomposition of the corresponding $(p \times p)$ true, population covariance matrix R_p :

$$R_p = \sum_{i=1}^p \lambda_i \underline{v}_i \underline{v}_i^T \quad (15)$$

For small values of p , a good approximant of R_p is the component carrying the most of the variance [27].

$$\lambda_1 \underline{v}_1 \underline{v}_1^T \quad (16)$$

As the size p of the covariance matrix becomes larger, more terms are required for a given level of approximation (cf. Figs. 5a,5b,5c).

This approximation is particularly interesting in the practical case of a finite observation of data in which the angle of the estimated principal eigenvector, relative to the true principal eigenvector of R_p , is small. That is, $\hat{\underline{v}}_1$ and $\underline{v}_1$ are nearly identical vectors.

The original [28,29] and low-rank linear prediction equations can be written as follows :

$$\text{(original)} \quad R_p \underline{g} = - \underline{r} \quad (17)$$

$$\text{(low-rank)} \quad (\lambda_1 \underline{v}_1 \underline{v}_1^T) \underline{b} = - \underline{r} \quad (18)$$

If we represent the correlation vector $\underline{r}$ using the eigenvectors of R_p :

$$\underline{r} = \gamma_1 \underline{v}_1 + \gamma_2 \underline{v}_2 + \dots + \gamma_p \underline{v}_p \quad (19)$$

then the minimum-norm solutions for the coefficients of the linear predictors are respectively :

$$\underline{g} = - \sum_{i=1}^p \left(\frac{\gamma_i}{\lambda_i} \right) \underline{v}_i \quad (20)$$

$$\underline{b} = - \left(\frac{\gamma_1}{\lambda_1} \right) \underline{v}_1 \quad (21)$$

Using the prediction coefficients specified by (20) and (21) we can form the associated prediction-error-filter polynomials. The set of zeros of a prediction-error-filter polynomial is denoted by $\{z_i ; i=1, \dots, p\}$. The positive real zero nearest to the unit circle is denoted by z_q .

$$|1 - z_q| = \min \{ |1 - z_i| ; z_i \text{ is positive real} \} \quad (22)$$

$$i=1, \dots, p$$

Having found z_q , if it exists, we use it to estimate the AR parameter.

$$a = \begin{cases} z_q, & |z_q| < 1 \\ 1/z_q, & |z_q| > 1 \end{cases} \quad (23)$$

If z_q does not exist, because there are no positive real zeros, then we make no estimate of (a). In practice, this could certainly be refined by projection of zeros which are very close to the positive real line.

The estimator resulting from the zero selection of formulas (22) and (23), when applied to the prediction-error-filter resulting from (20), is known to be unbiased, given the true R_p . On the other hand it may be slightly biased when using (21) instead of (20). This can be seen for two values of (a) in Tables [1] and [2] and in Fig. 6, where the prediction-error-filter zeros resulting from (21) are plotted for different predictor-lengths (p). The presence of the bias is not a serious drawback in the estimator, especially when the data record is

short, and the poles are close to the unit circle, in which case the biases may be lower than the minimum standard deviation given by the C-R bound.

In Tables [1] and [2], we compare the biases for two different values of (a) : .95 and .99. We can observe that the bias decreases as (a) tends to unity, and it increases as the predictor length gets larger.

Thus, as one would expect, if the true covariance matrix is replaced by a different, rank-1 approximating matrix, some error is introduced in the value of the AR parameter which is obtained using linear prediction. As we show below, the acceptance of this small error with exact covariance information leads to improved insensitivity to the errors in an estimated covariance matrix.

VI. SIMULATIONS

In this section we present some results of simulations, where we compare the performance of maximum likelihood estimation and least squares techniques based on linear prediction. The linear-prediction approach is here carried out in two ways, using the estimated full-rank correlation matrix and using its low-rank approximant.

Given the recursion (2), $K=500$ independent realizations were generated. We have started from independent initial conditions $\{s(0,j) ; j=1,500\}$ given by a zero-mean gaussian random number generator of variance equal to $\sigma^2/(1-a^2)$. This was done in order to guarantee the stationarity of the sequence $\{s(t,j) ; t=1,\dots,N\}$ for a given trial j . Without loss of generality σ has been normalized to unity.

We have used the Forward (F) and the Forward-Backward (FB) [25],[26] structures for the data matrix Y defined as follows:

$$\begin{aligned}
 \text{F-data matrix : } Y &= Y_f & ; & & \text{F-data vector : } \underline{h} &= \underline{h}_f \\
 \text{FB-data matrix : } Y &= \begin{bmatrix} Y_f \\ \hline Y_b \end{bmatrix} & ; & & \text{FB-data vector : } \underline{h} &= \begin{bmatrix} \underline{h}_f \\ \hline \underline{h}_b \end{bmatrix} \\
 & & & & & (24)
 \end{aligned}$$

These matrices and vectors are written out more explicitly as

$$\begin{aligned}
 Y_f &= \begin{bmatrix} s(p) & s(p-1) & \dots & s(2) & s(1) \\ s(p+1) & s(p) & \dots & s(3) & s(2) \\ \cdot & \cdot & & \cdot & \cdot \\ \cdot & \cdot & & \cdot & \cdot \\ s(N-1) & s(N-2) & \dots & s(N-p+1) & s(N-p) \end{bmatrix} & \underline{h}_f &= \begin{bmatrix} s(p+1) \\ s(p+2) \\ \cdot \\ \cdot \\ s(N) \end{bmatrix} \\
 Y_b &= \begin{bmatrix} s(2) & s(3) & \dots & s(p) & s(p+1) \\ s(3) & s(4) & \dots & s(p+1) & s(p+2) \\ \cdot & \cdot & & \cdot & \cdot \\ \cdot & \cdot & & \cdot & \cdot \\ s(N-p+1) & \dots & \dots & s(N-1) & s(N) \end{bmatrix} & \underline{h}_b &= \begin{bmatrix} s(1) \\ s(2) \\ \cdot \\ \cdot \\ s(N-p) \end{bmatrix}
 \end{aligned}
 \tag{25}$$

For all cases we have taken predictor-lengths, p , to be less than half the observed data-vector length N .

$$p < N/2 \tag{26}$$

Since our estimated correlation matrix is proportional to the transpose product of the data matrix Y , namely $Y^T Y$, the theoretical formulas for the prediction coefficients, (20) and (21), can be rewritten for the available data as:

$$\underline{g} = -(Y^T Y)^{-1} Y^T \underline{h} = -\hat{R}_p^{-1} \hat{r} \tag{27}$$

$$\underline{b} = -\left(\frac{\hat{\gamma}_1}{\hat{\lambda}_1}\right) \hat{y}_1 \tag{28}$$

As opposed to the case of (20) and (21), the scalars and vectors of (27) and (28) are determined by the observed data vector of N samples.

The prediction coefficients of (28) have been previously used in [15] where the objective was spectral peak estimation of an AR process observed in the presence of extra noise. We are now interested in estimation of pole location for the first-order AR process and we use the zero selection technique of formulas (22,23).

Here we consider the case where only one realization ($K=1$) of a finite record of data (N) is available for estimation of the AR parameter. Five hundred independent records are used in order to measure the properties of the estimation error.

The simulation results shown in this paper have been obtained from only one value of the AR parameter, i.e. $a = .95$. We have obtained similar results for values of (a) varying from .8 up to .99.

The noise-free case is presented in Tables 3 through 14 from data record lengths of $N=25$ for Tables 1 through 8, and $N=100$ for Tables 9 through 14. These results are summarized in Figs. 21, 22.

The signal-plus-noise case is presented in Figs. 23, 24 for an SNR of 5db, and in Fig. 25 for an SNR of 0db (see also Tables 15 and 16).

Associated with each table is a corresponding figure showing superpositions of zeros. The indices (a,b,c) are, respectively, the cases of:

a) superposition of the zeros of the estimated prediction-error-filter (P-E-F) polynomials for all 500 trials ;

b) superposition of the zeros of the estimated P-E-F polynomials which have at least one positive real zero ;

c) superposition of the zeros of the estimated P-E-F polynomials which have no positive real zeros.

If, for a given predictor length, we have at least one positive real zero for each of the 500 sets of P-E-F polynomials, the figures associated with the indices (b) and (c) are not plotted.

For the noise-free case we have used predictor lengths (p) from 1 to 6. For the signal-plus-noise case, (p) has been chosen up to 16.

In order to compare performance we have tabulated some information other than the standard deviation and bias of the estimates :

1) the number of zeros outside the unit circle of the 500 estimated P-E-F polynomials : column "uns" ;

2) the number of trials in which the selected positive real zeros were outside the unit circle : column "out" ;

3) the number of trials in which there was no solution i.e. , there was no positive real zero : column "Opz" ;

4) the number of trials in which there was only one positive real zero : column "1pz" ; and

5) the number of trials in which there were two positive real zeros : column "2pz".

The number of cases with more than two positive real zeros can of course be deduced from 3), 4) and 5).

Tables 1 and 2 give the asymptotic (large number of samples) bias of the rank-1 approximation method. These values have been obtained numerically from (17) and (18), given the true correlation matrix. From Tables 4 and 6 one can see that this estimator may be viewed as practically unbiased for a small number of samples ($N=25$) if one considers the bias relative to the lower-bound standard deviation (C-R bound).

VII. COMMENTS ON THE RESULTS

The Cramer-Rao bounds place lower bounds on the estimation error variances for unbiased estimators. The maximum likelihood estimator may perform badly with respect to the C-R bound, especially near threshold and for a small number of samples, N . As conventional methods of linear prediction, including all of the standard variations, are approximations to maximum likelihood, the same comments hold for linear prediction techniques.

The defects in maximum likelihood and conventional linear prediction can be corrected by tailoring the data with SVD to incorporate structural information (such as low or approximate-low rank), fitting a model to the data that is substantially higher order than the model is known to be, and separating modal zeros from noise zeros using prior information as a fitting rule.

Our summarizing results of Figs. 21 and 22 substantiate these claims. Beyond this, it is shown in Figs. 23, 24 and 25 that even more impressive results are obtained in the presence of additive noise.

We conclude with the following, more detailed comments :

a) Classical true-order ($p=1$) linear prediction

For a small number of samples ($N=25$) we can see from Tables 5 and 7 for $p=1$ that the Forward-Backward Linear Prediction (FBLP) performs better than the Forward Linear Prediction (FLP). This is true not only for rms error but also for the location of zeros inside the unit circle. For $N=25$ we had 37 out of 500 cases in which zeros were outside the unit

circle for the FLP method, while for the FBLP all the Prediction-Error Filters (P-E-F) were minimum phase ("stable").

Nevertheless, the FLP and FBLP methods tend to have the same performances as N gets large. We can see for $N=100$ in Tables 11 and 13 that for $p=1$, the performance is about the same.

From Tables 3 and 9 one verifies that the MLE is slightly better than the classical true-order linear prediction, but still the C-R bound is not achieved.

b) Full-rank larger-order ($p>1$) linear prediction

For this case where $p>1$, using the full-rank estimated covariance matrix and our particular zero selection procedure, we also conclude that for small $N(=25)$ the FBLP method is better than the FLP method. This, in terms of location of the zeros of the estimated P-E-F polynomials inside the unit circle and in terms of the rms error (for $p<5$). For $N=100$ the performances of both methods are practically the same for $p<4$.

On the other hand the number of trials in which there were no positive real zeros (Opz) was slightly bigger for FBLP than for FLP for smaller samples ($N=25$), cf. Tables (5,7) and (11,13).

As far as the comparison with the true order ($p=1$) case is concerned, we can see that for $N=25$ the full-rank high-order predictors provide better estimates, especially for the FBLP case. The optimal predictor length was $p=3$ for both FLP and FBLP.

On the other hand, for $N=100$, the true-order predictor ($p=1$) performed better than those for $p>1$.

One could infer that for a small number of samples, fluctuations of the estimated covariance are not negligible, making it useful to use extra poles to model and, hence, isolate the fluctuations. That is, the fluctuations are filtered by modeling them. But, for large-enough data records, the covariance is accurate and the true order can be used.

If we look at Fig.21, we can notice that the full-rank high-order predictors may provide better estimates than the MLE for a small number of samples ($N=25$). This is the case for FBLP with $p=3$.

c) The reduced-rank high-order predictors

For this case, the value of p is greater than one and a low-rank approximation to the estimated covariance matrix is used to obtain the prediction coefficients. This method provided more accurate estimates than the other methods which we examined.

Comparing Tables 4 and 6 for $N=25$, and Tables 10 and 12 for $N=100$, once more the FBLP method is shown to perform better than the FLP method. This can be also viewed in Figs.21 and 22.

We summarize the performances of these estimators for the noise-free case in Figs. 21 and 22 and for the signal-plus-noise case in Figs. 23, 24 and 25.

For the signal-plus-noise case, in order to filter the noise, we had to use a higher order predictor.

After completion of this report we discovered the recently published paper (ICASSP - march 1984) by S. Shon [30] in which the C-R bound for the first-order AR parameter is also presented. In the last ASSP Spectrum Estimation Workshop, in Tampa (november 1983) we have also presented the same bound [31]. Our proposed low-rank estimator provides better performance than the estimators discussed by Shon, as we have shown in this report. The C-R bounds for the estimation of AR coefficients of a general p-order AR process can also be determined exactly (for small samples) [32].

ACKNOWLEDGEMENTS

The authors wish to thank I. Kirsteins for allowing them to include his results on the angles between true and estimated eigenvalues. They also wish to thank Dr. L.L. Scharf for his many useful comments.

REFERENCES

- [1] D.W. Tufts, R. Kumaresan, "Improved Spectral Resolution II" , Proc. ICASSP IEEE(April 1980),pp.592-597.
- [2] D.W. Tufts, R. Kumaresan, "Estimation of Frequencies of Multiple Sinusoids: Making Linear Prediction Perform Like Maximum Likelihood" ,Proc. of the IEEE, vol.70, No.9, September 1982, pp.975-989.
- [3] R. Kumaresan, D.W. Tufts, "Accurate Parameter Estimation of Noisy Speech-Like Signals" ,Proc. ICASSP-IEEE, April 1982, pp.1357-1361.
- [4] R. Kumaresan, D.W. Tufts, "Estimating the Parameters of Exponentially Damped Sinusoids and Pole-Zero Modeling in Noise" ,IEEE Trans. Acoust. Speech and Signal Processing, vol.ASSP-30, No.6, December 1982, pp. 833-840.
- [5] R. Kumaresan, D.W. Tufts, L.L. Scharf, "A Prony Method for Noisy Data : Choosing the Signal Components and Selecting the Order in Exponential Signal Models" , Report No.3, December 1982, Dept. of Electrical Engineering, University of Rhode Island, prepared for the Office of Naval Research. Published also on Proc. of IEEE, vol.72, No.2, pp.230-233, february 1984.
- [6] R. Kumaresan, D.W. Tufts, "A Two-Dimensional Technique for Frequency-Wavenumber Estimation" , Proc. of the IEEE, vol.69, No.11, November 1981, pp.1515-1517.
- [7] R. Kumaresan, D.W. Tufts, "Improved Spectral Resolution III" , Proc. of the IEEE, vol.68, No.10, October 1980, pp.1354-1355.

[8] R. Kumaresan, D.W. Tufts, "Estimating the Angle of Arrival of Multiple Plane Waves", IEEE Trans. on Aerospace and Electronic Systems, vol.19, No.1, January 1983, pp.134-139.

[9] R. Kumaresan, D.W. Tufts, "Data-Adaptive Principal Component Signal Processing", Proc. IEEE, Conference on Decision and Control, Albuquerque, N.M., December 1980.

[10] R. Kumaresan, D.W. Tufts, "Singular Value Decomposition and Spectral Analysis", Proc. IEEE Conference on Decision and Control, December 1981.

[11] L.L. Scharf, A.A. Beex, T. von Reyn, "Modal Decomposition of Covariance Sequences for Parametric Spectrum Analysis", Proc. ICASSP-IEEE, 1981, pp.512-517.

[12] A.A. Beex, L.L. Scharf, "Covariance Sequence Approximation for Parametric Spectrum Modeling", IEEE Trans. on ASSP, vol. ASSP-29, No.5, October 1981, pp.1042-1052.

[13] C. Gueguen, L.L. Scharf, "Exact Maximum Likelihood Identification of ARMA Models: A Signal Processing Perspective", 1st. EUSIPCO, September 1980, Lausanne, Switzerland.

[14] C. Gueguen, F. Giannella, "Analyse Spectrale et Approximation de Formes Quadratiques", Proc. GRETSI, June 1981, Nice, France.

[15] R. Kumaresan, D.W. Tufts, "Singular Value Decomposition and Spectral Analysis", Proc. First Workshop on Spectral Estimation, Hamilton, Canada, August 1981, pp.6.4.1-6.4.12.

[16] R.K. Mehra, "On-Line Identification of Linear Dynamic Systems with Applications to Kalman Filtering", IEEE Trans. on Automatic Control, vol.AC-16, No.1, February 1971, pp.12-21 .

[17] K.J. Astrom, Eykhoff, "System Identification - a Survey", Automatica, vol.7, pp.123-162, Pergamon Press, 1971.

[18] E. Parzen, "Some Recent Advances in Time Series Modeling", IEEE Transactions on Automatic Control, vol.AC-19, No.6, December 1974, pp.723-730 .

[19] J.H. Wilkinson, "The Algebraic Eigenvalue Problem", Clarendon Press, Oxford, 1965.

[20] G.E. Box, G.M. Jenkins, "Time Series Analysis Forecasting and Control", Holden Day, 1976, (Revised Edition), p.281 .

[21] M. Pagano, "Estimation of Models of AR Signals Plus White Noise", Ann. Sta., vol.2, No.1, pp.98-108, January 1974.

[22] Y. Hosoya, "Efficient Estimation of a Model with an Autoregressive Signal with White Noise", Technical Report No.31, Dept. of Statistics, Stanford University, March 1979.

[23] B. Friedlander, "On the Computation of the Cramer-Rao Bound for ARMA Parameter Estimation, Proc. ICASSP-IEEE (march 1984), pp. 14.1.1-14.1.4 .

[24] H.L. van Trees, "Detection, Estimation and Modulation Theory", vol. I, Wiley, New York, 1968.

[25] A.H. Nuttall, "Spectral Analysis of a Univariate process with bad data points, via Maximum Entropy and Linear Predictive Technique", in NUSC (Naval Underwater Systems Center) New London Laboratory, 26 march 1976, Technical Report No.5303.

[26] T.J. Ulrych, R.W. Clayton, "Time series modelling and maximum entropy", Physics of the earth and planetary interiors, vol.12, pp.188-200, august 1976.

[27] C.R. Rao, "Linear Statistical Inference and Its Applications", 2nd. edition, John Wiley and Sons, 1973.

[28] J. Makhoul, "Linear Prediction : a Tutorial Review", Proc. IEEE, VOL.63, pp.561-580, april 1975

[29] J.D. Markel, A.H. Gray, "Linear Prediction of speech", Springer-Verlag, N.Y., 1976.

[30] S. Shon, "Performance Comparison of Autoregressive Estimation Methods", Proc. ICASSP-IEEE, march 1984, pp. 14.3.1-14.3.4 .

[31] D.W. Tufts, F. Giannella, I. Kirsteins, L.L. Scharf, "Cramer-Rao Bounds on the Accuracy of Autoregressive Parameter Estimators", Proc. IEEE-ASSP Spectrum Estimation Workshop II, november 1983, pp.12-16 .

[32] F. Giannella, "Cramer-Rao bounds for AR Parameter Estimation from Small Samples", submitted for publication.

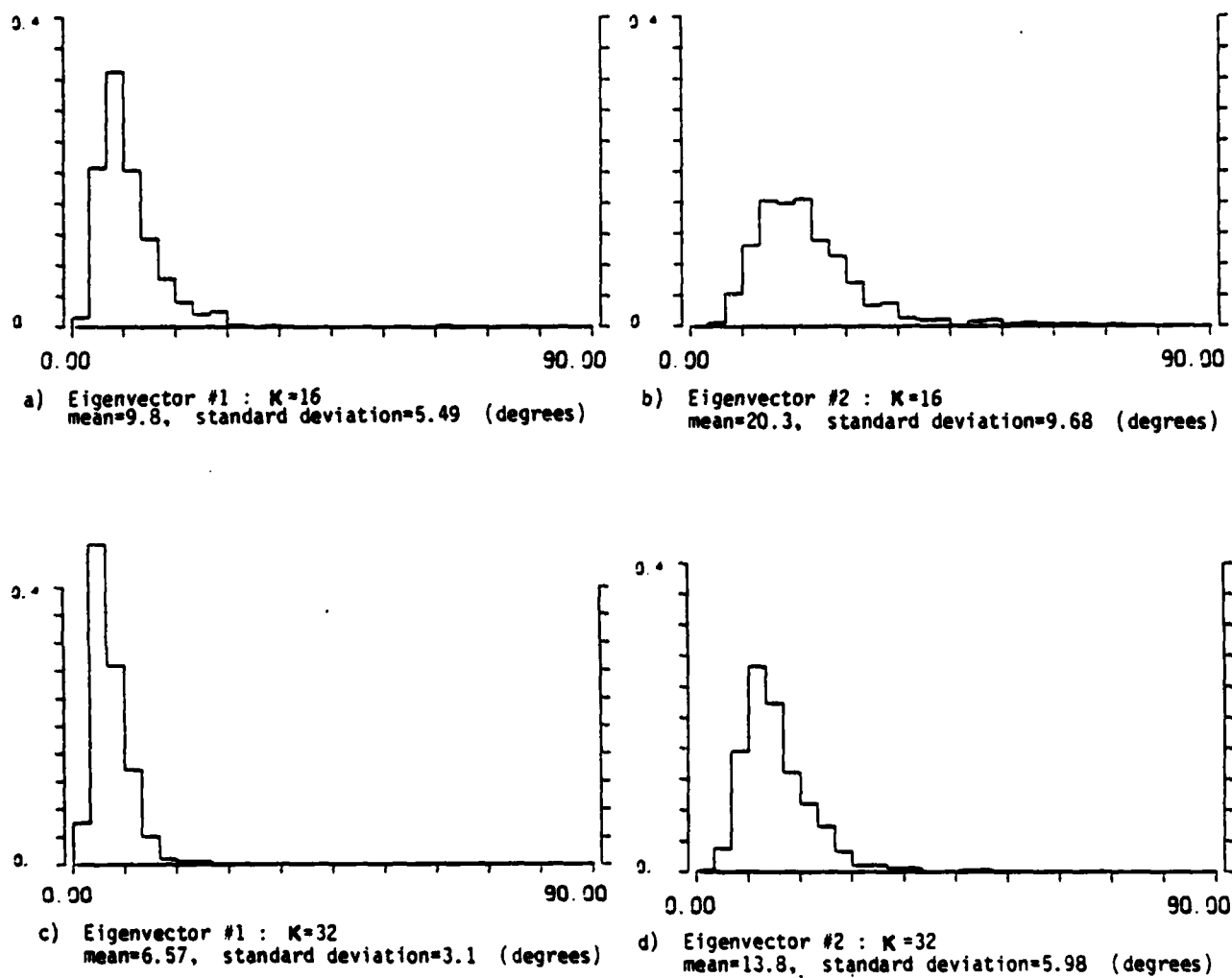


Fig. 1) Histograms of the absolute value of the angular perturbations of the principle eigenvectors of R .

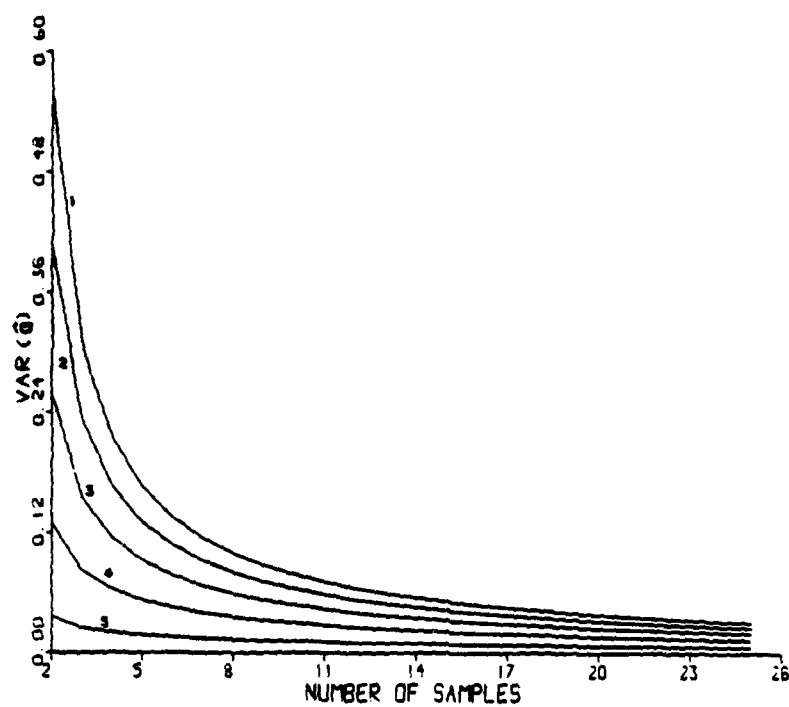


Fig.2a

C-R bounds for only one realization ($K=1$) and σ known 1) $a=.5$; 2) $a=.6$; 3) $a=.7$; 4) $a=.8$; 5) $a=.9$

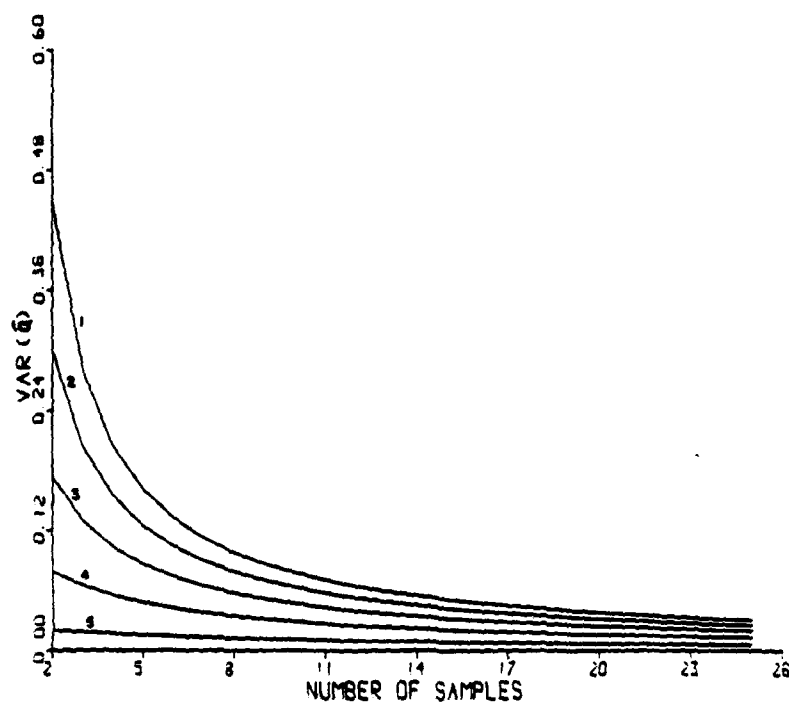


Fig.2b

C-R bounds for only one realization ($K=1$) and σ unknown 1) $a=.5$; 2) $a=.6$; 3) $a=.7$; 4) $a=.8$; 5) $a=.9$

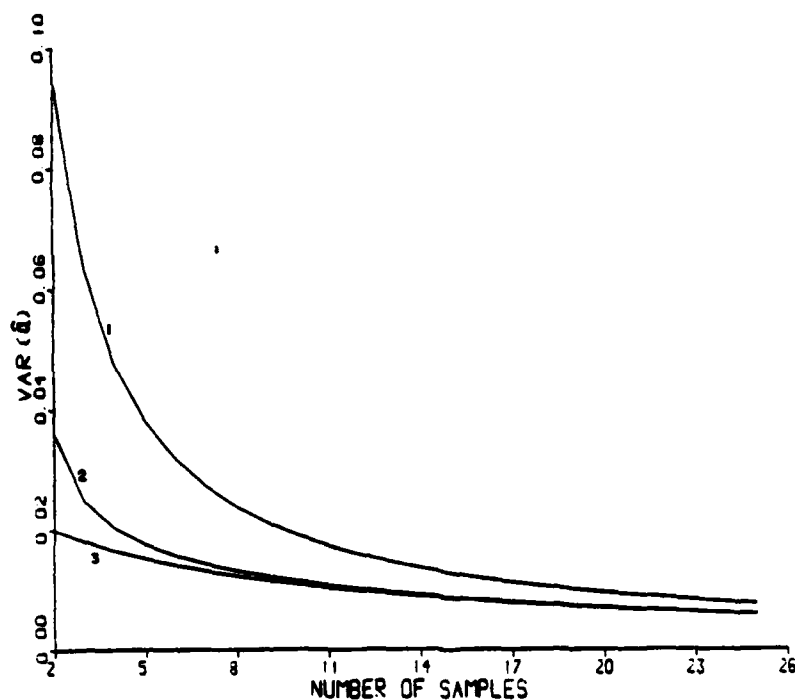


Fig.2c

C-R bounds for $a=.9$ and $K=1$ 1) asymptotic bound (plotted only as a reference) 2) true bound for σ unknown 3) true bound for σ known

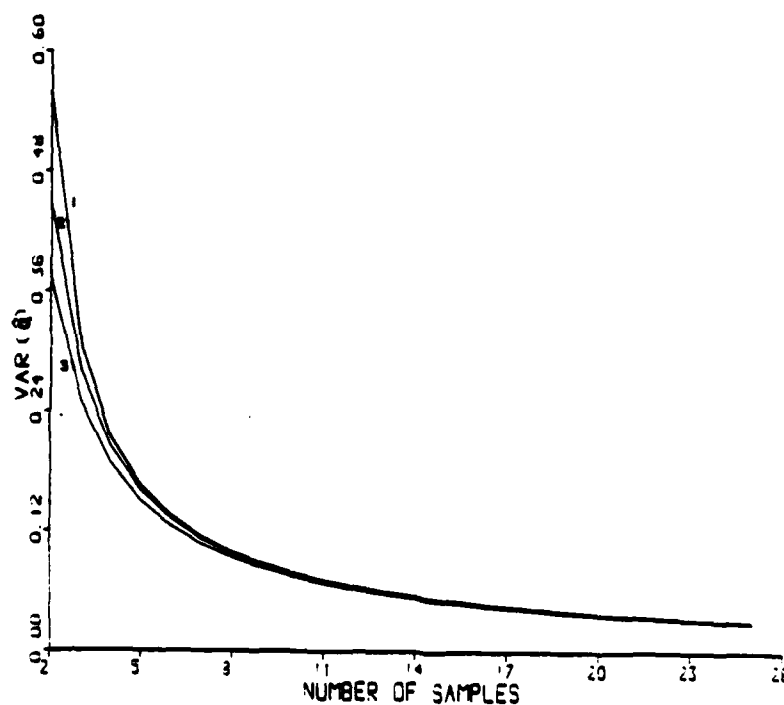


Fig.2d

C-R bounds for $a=.5$ and $K=1$ 1) true bound for σ unknown 2) true bound for σ known 3) asymptotic bound (plotted only as a reference)

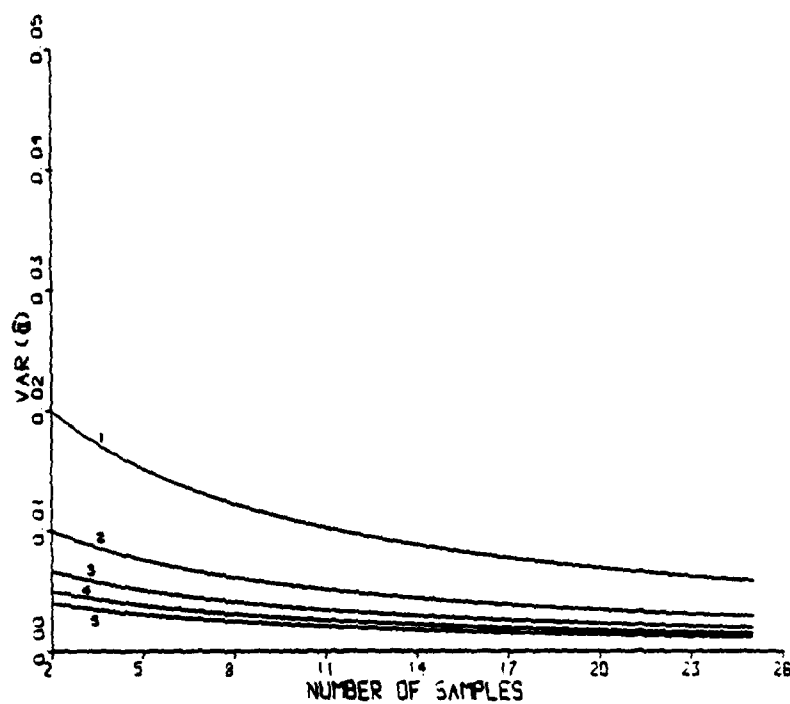


Fig.2e

C-R bounds for $a=.9$ and σ known 1)K=1 ; 2)K=2 ;
3)K=3 ; 4)K=4 ; 5)K=5

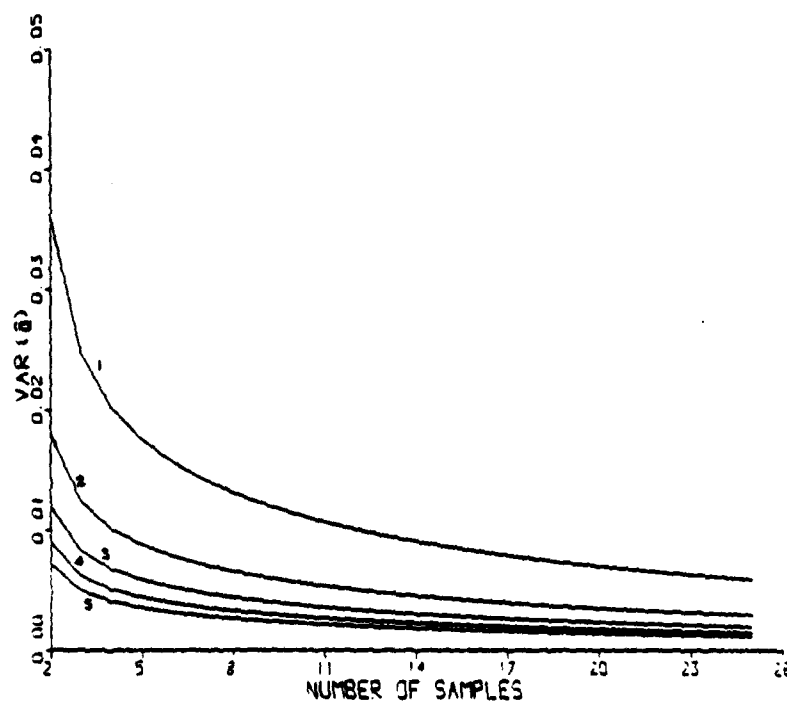


Fig.2f

C-R bounds for $a=.9$ and σ unknown 1)K=1 ;
2)K=2 ; 3)K=3 ; 4)K=4 ; 5)K=5

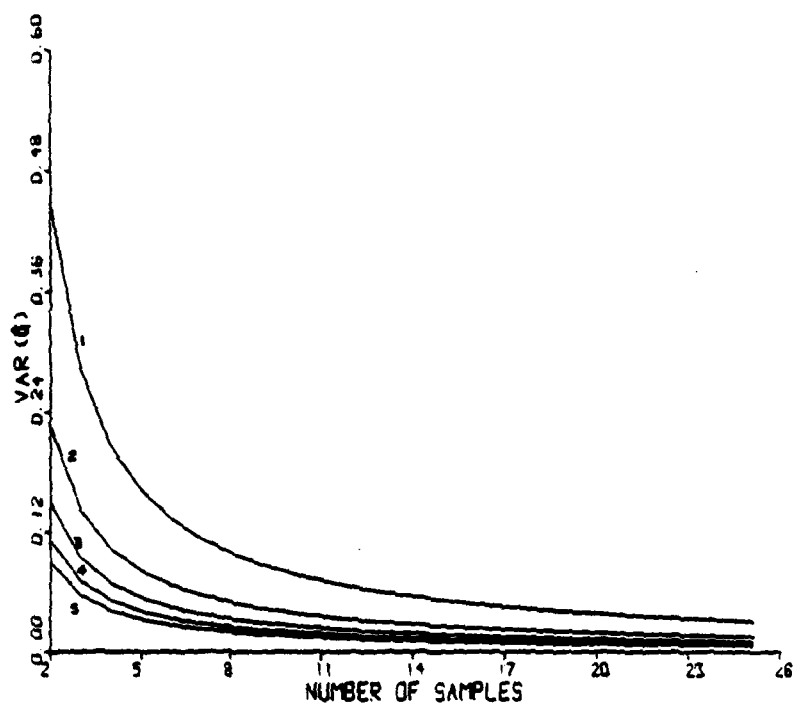


Fig.2g

C-R bounds for $a=.5$ and σ known 1)K=1 ; 2)K=2 ;
3)K=3 ; 4)K=4 ; 5)K=5

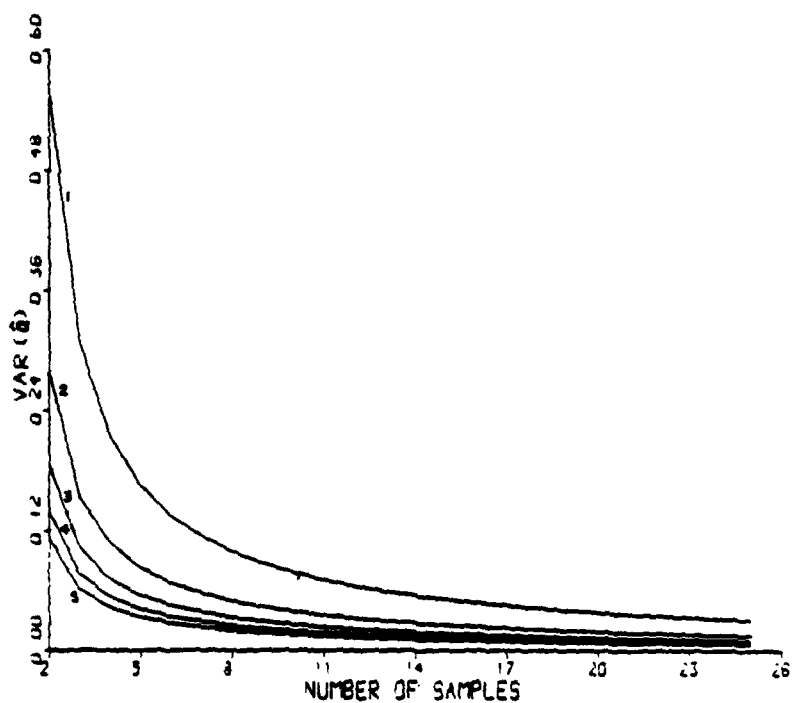


Fig.2h

C-R bounds for $a=.5$ and σ unknown 1)K=1 ;
2)K=2 ; 3)K=3 ; 4)K=4 ; 5)K=5

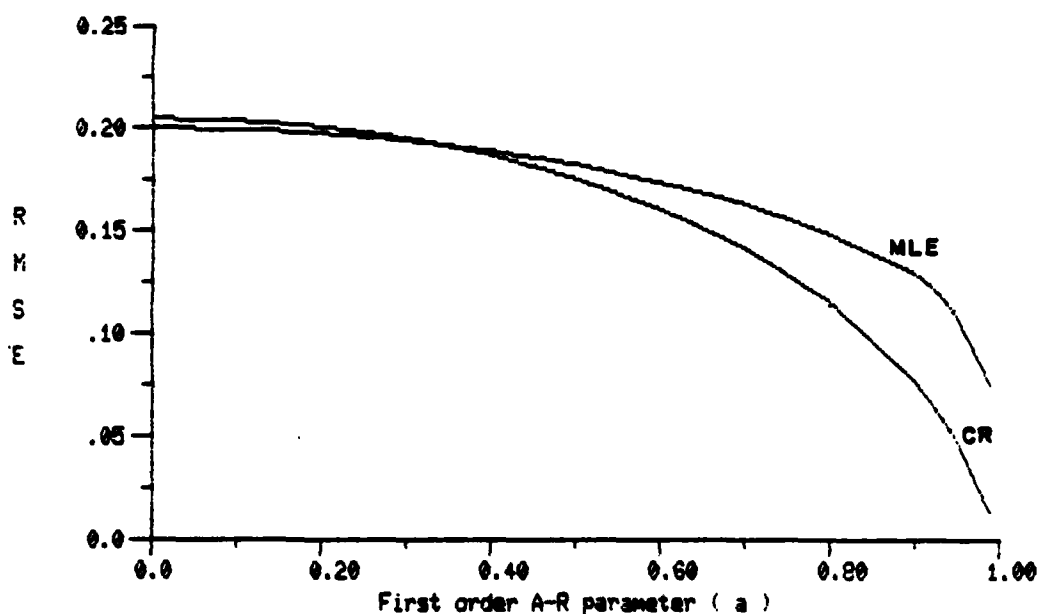


Fig. 3a

Comparison of the Cramer-Rao (CR) lower bound standard deviation for unbiased estimators of the first-order AR parameter (a) with the estimated root-mean-square error (rmse) of the Maximum Likelihood Estimator (MLE) of (a), for a sample size of $N=25$ over 500 independent realizations.

Values of (a) : .0, .1, .2, .3, .4, .5, .6, .7, .8, .9, .92, .95, .99

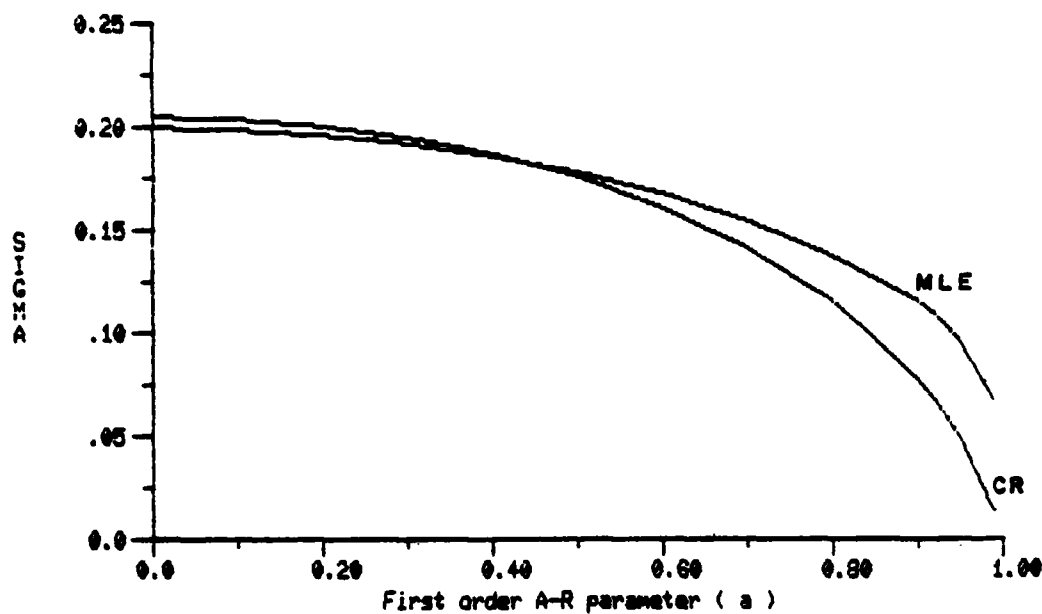


Fig. 3b

Comparison of the Cramer-Rao (CR) lower bound standard deviation for unbiased estimators of the first-order AR parameter (a) with the estimated standard deviation (sigma) of the Maximum Likelihood Estimator (MLE) of (a), for a sample size of N=25 over 500 independent realizations.

Values of (a) : .0, .1, .2, .3, .4, .5, .6, .7, .8, .9, .92, .95, .99

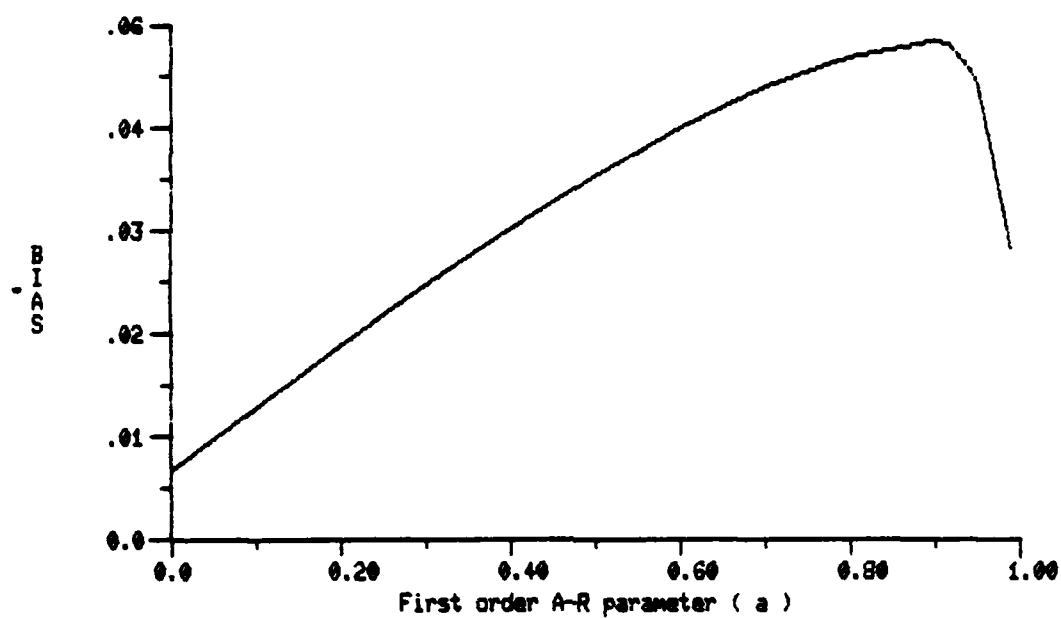


Fig. 3c

Estimated bias of the Maximum Likelihood Estimator (MLE) of (a), for a sample size of $N=25$ over 500 independent realizations.

Values of (a) : .0, .1, .2, .3, .4, .5, .6, .7, .8, .9, .92, .95, .99

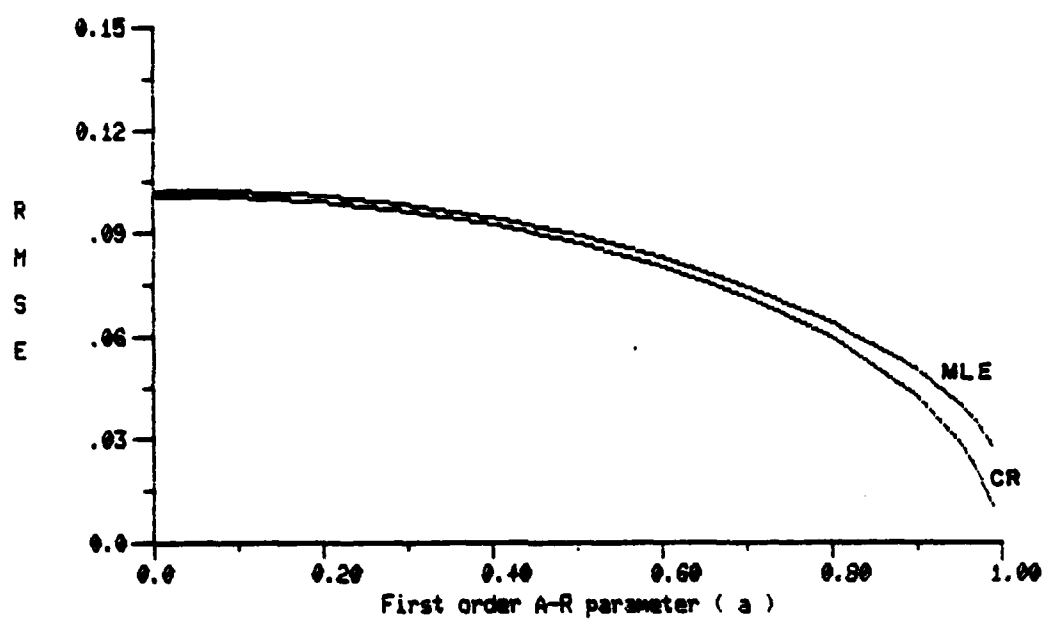


Fig. 4a

Comparison of the Cramer-Rao (CR) lower bound standard deviation for unbiased estimators of the first-order AR parameter (a) with the estimated root-mean-square error (rmse) of the Maximum Likelihood Estimator (MLE) of (a), for a sample size of $N=100$ over 500 independent realizations.

Values of (a) : .0, .1, .2, .3, .4, .5, .6, .7, .8, .9, .95, .97, .99

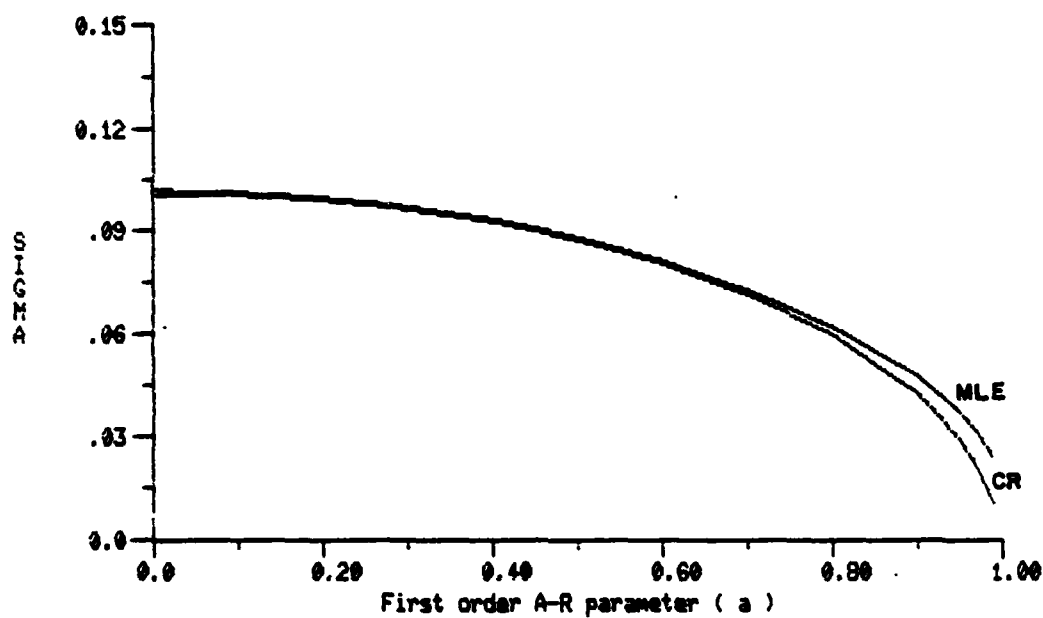


Fig. 4b

Comparison of the Cramer-Rao (CR) lower bound standard deviation for unbiased estimators of the first-order AR parameter (a) with the estimated standard deviation (sigma) of the Maximum Likelihood Estimator (MLE) of (a), for a sample size of $N=100$ over 500 independent realizations.

Values of (a) : .0, .1, .2, .3, .4, .5, .6, .7, .8, .9, .95, .97, .99

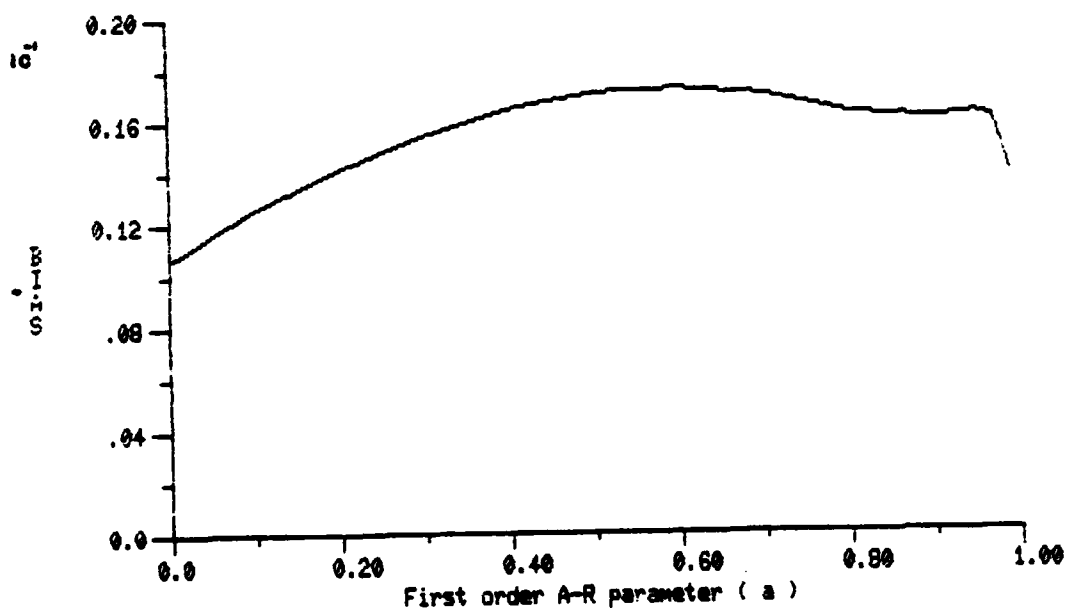


Fig. 4c

Estimated bias of the Maximum Likelihood Estimator (MLE) of (a), for a sample size of $N=100$ over 500 independent realizations.

Values of (a) : .0, .1, .2, .3, .4, .5, .6, .7, .8, .9, .95, .97, .99

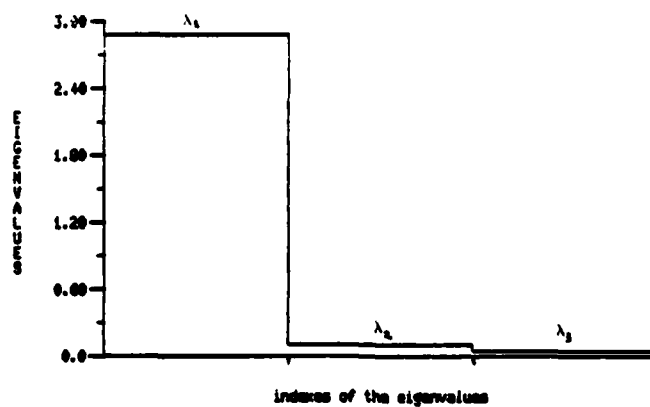


Fig. 5a (p=3)

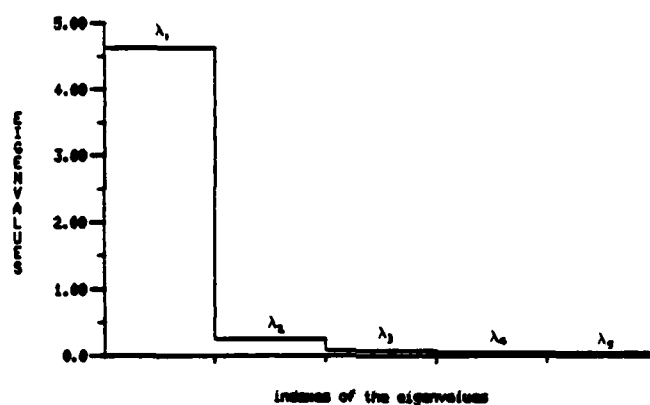


Fig. 5b (p=5)

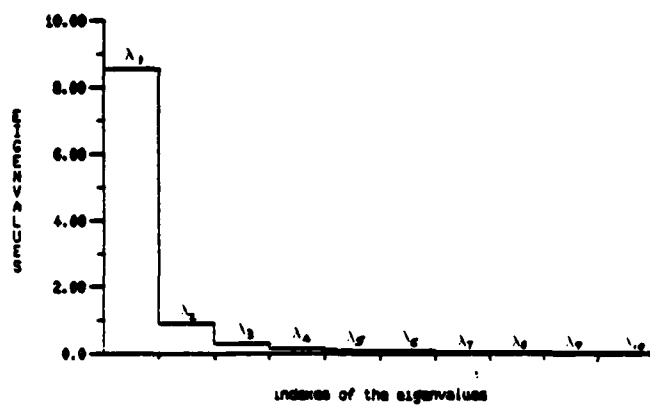


Fig. 5c (p=10)

Fig. 5

Eigenvalues of the p-dimension normalized (trace=p) correl. matrix R_p , of a first-order AR process in which $a=.95$

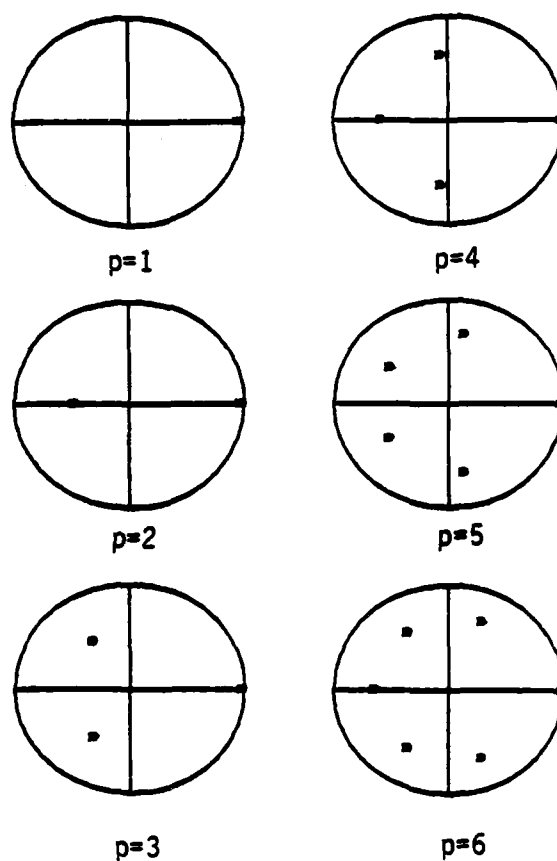


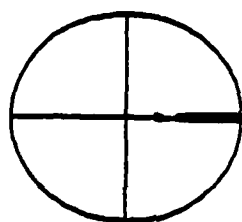
Fig. 6

$a = .95$; Noise-free case ; $N = \infty$

Zeros of the Prediction-Error-Filter Polynomial for the
 asymptotic case of large samples from the rank-1 approximation
 of the (true) correlation matrix.

[p : predictor length]

(cf. Table 1)



$p=1$

Fig. 7

$a = .95$; Noise-free case ; $N = 25$

Superposition of the Zeros estimated by the MLE over 500 trials.

(cf. Table 3)

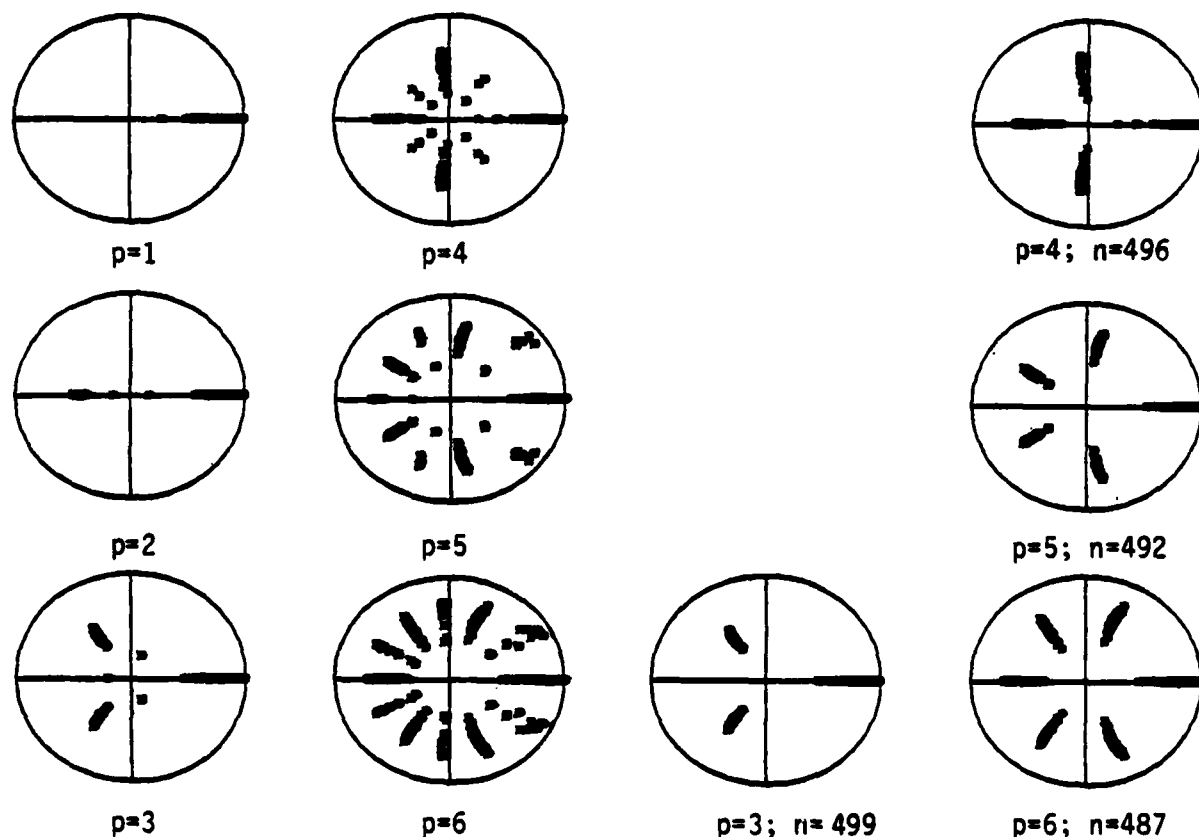


Fig. 8a

Fig. 8b

$a = .95$; Noise-free case ; $N = 25$

Superposition of the Zeros of the Estimated Prediction-Error-Filter Polynomials obtained from the rank-1 approximation of the estimated Forward-Backward covariance matrix for : a) All 500 trials. b) Cases in which each set of polynomial zeros has at least one positive real zero.

[p : predictor length ; n : number of cases out of 500 trials]

(cf. Table 4)

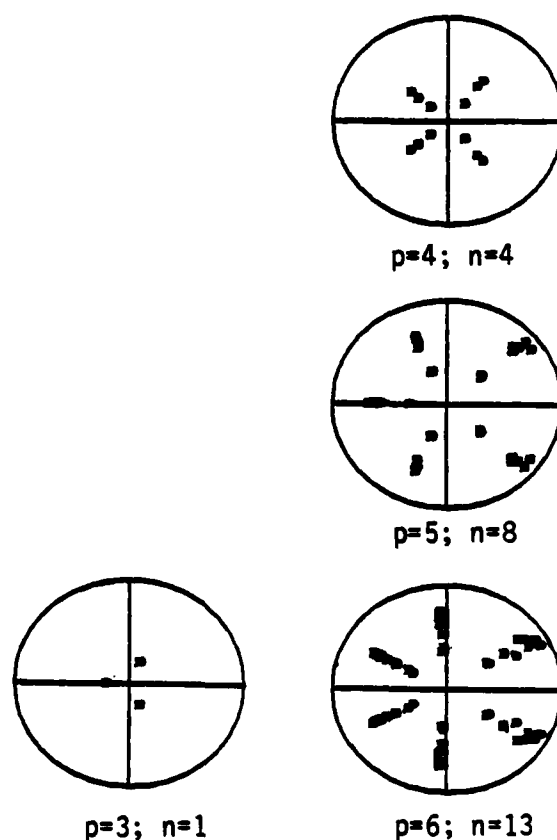


Fig. 8c

$a = .95$; Noise-free ; $N = 25$

Superposition of the Zeros of the Estimated Prediction-Error-Filter Polynomials obtained from the rank-1 approximation of the estimated Forward-Backward covariance matrix. c) Cases in which each set of polynomial zeros has no positive real zero.

[p : predictor length ; n : number of cases out of 500 trials]

(cf. Table 4)

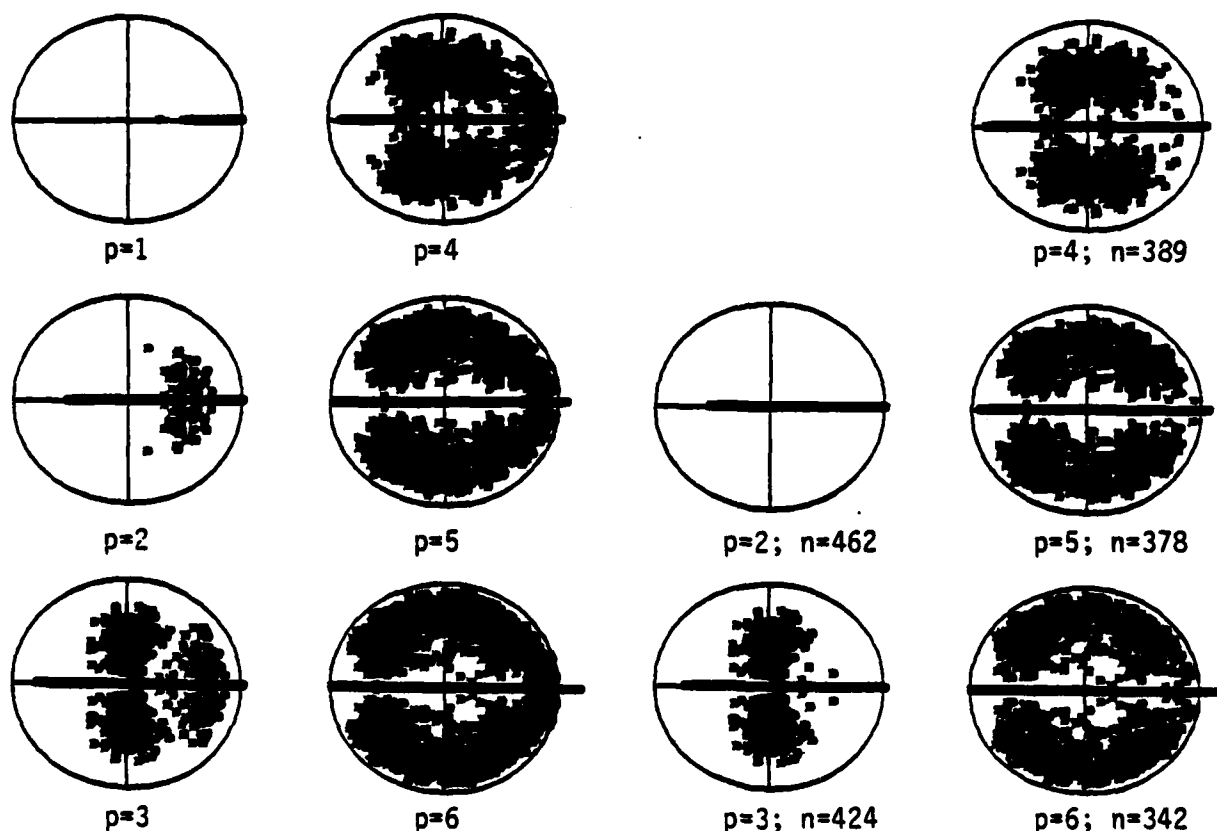


Fig. 9a

Fig. 9b

$a = .95$; Noise-free case ; $M = 25$

Superposition of the Zeros of the Estimated Prediction-Error-Filter Polynomials obtained from the estimated full-rank Forward-Backward covariance matrix for :

a) All 500 trials. b) Cases in which each set of polynomial zeros has at least one positive real zero.

[p : predictor length ; n : number of cases out of 500 trials]

(cf. Table 5)

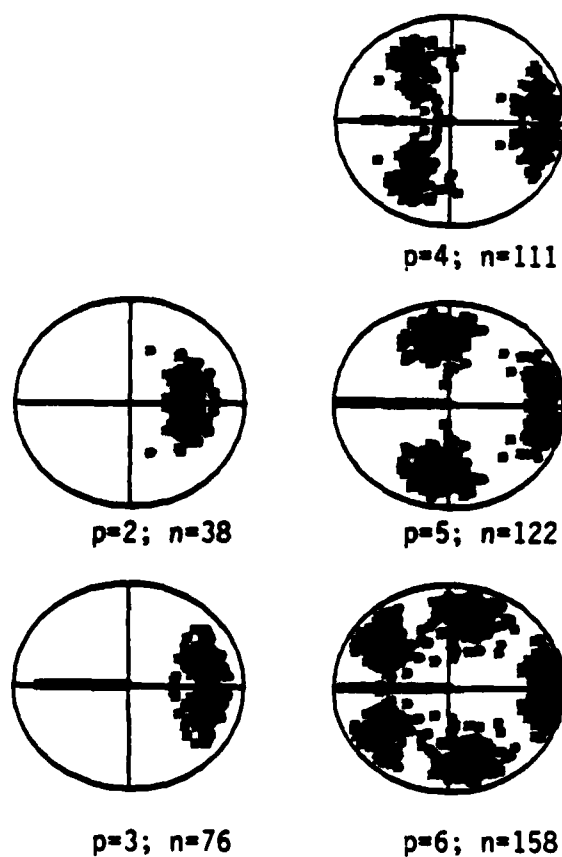


Fig. 9c

$a = .95$; Noise-free ; $N = 25$

Superposition of the Zeros of the Estimated Prediction-Error-Filter Polynomials obtained from the estimated full-rank Forward-Backward covariance matrix. c) Cases in which each set of polynomial zeros has no positive real zero.

[p : predictor length ; n : number of cases out of 500 trials]

(cf. Table 5)

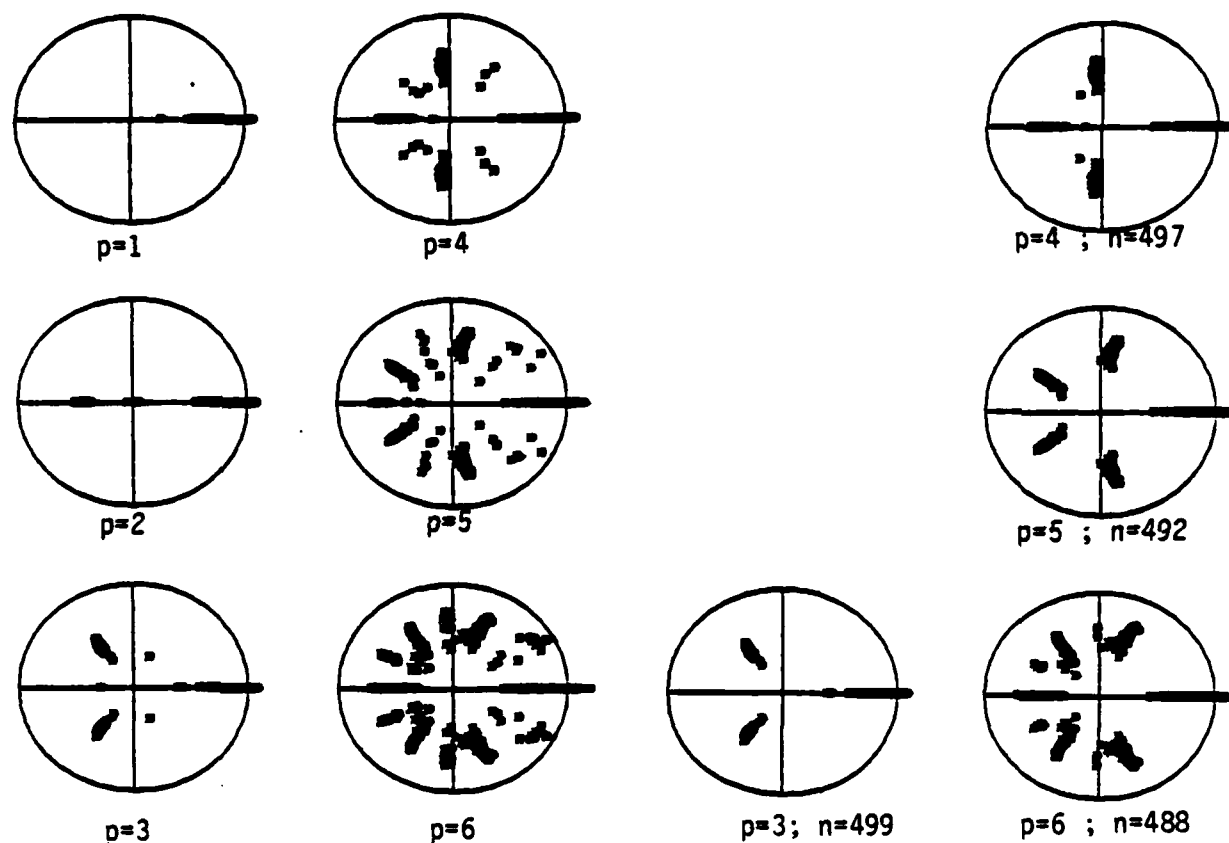


Fig. 10a

Fig. 10b

$a = .95$; Noise-free case ; $N = 25$

Superposition of the Zeros of the Estimated Prediction-Error-Filter Polynomials obtained from the rank-1 approximation of the estimated Forward covariance matrix for :

a) All 500 trials. b) Cases in which each set of polynomial zeros has at least one positive real zero.

[p : predictor length ; n : number of cases out of 500 trials]

(cf. Table 6)

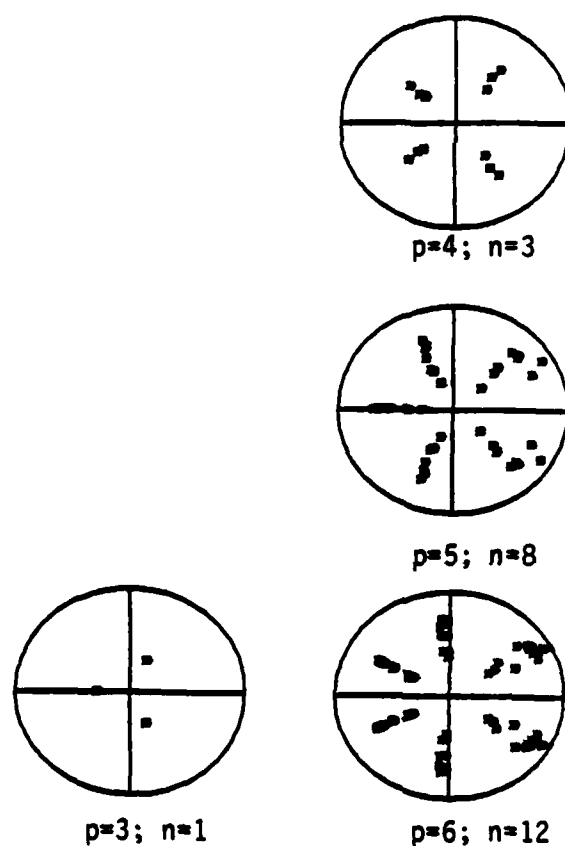


Fig. 10c

$a = .95$; Noise-free case ; $N = 25$

Superposition of the Zeros of the Estimated Prediction-Error-Filter Polynomials obtained from the rank-1 approximation of the estimated Forward covariance matrix. c) Cases in which each set of polynomial zeros has no positive real zero.

[p : predictor length ; n : number of cases out of 500 trials]

(cf. Table 6)

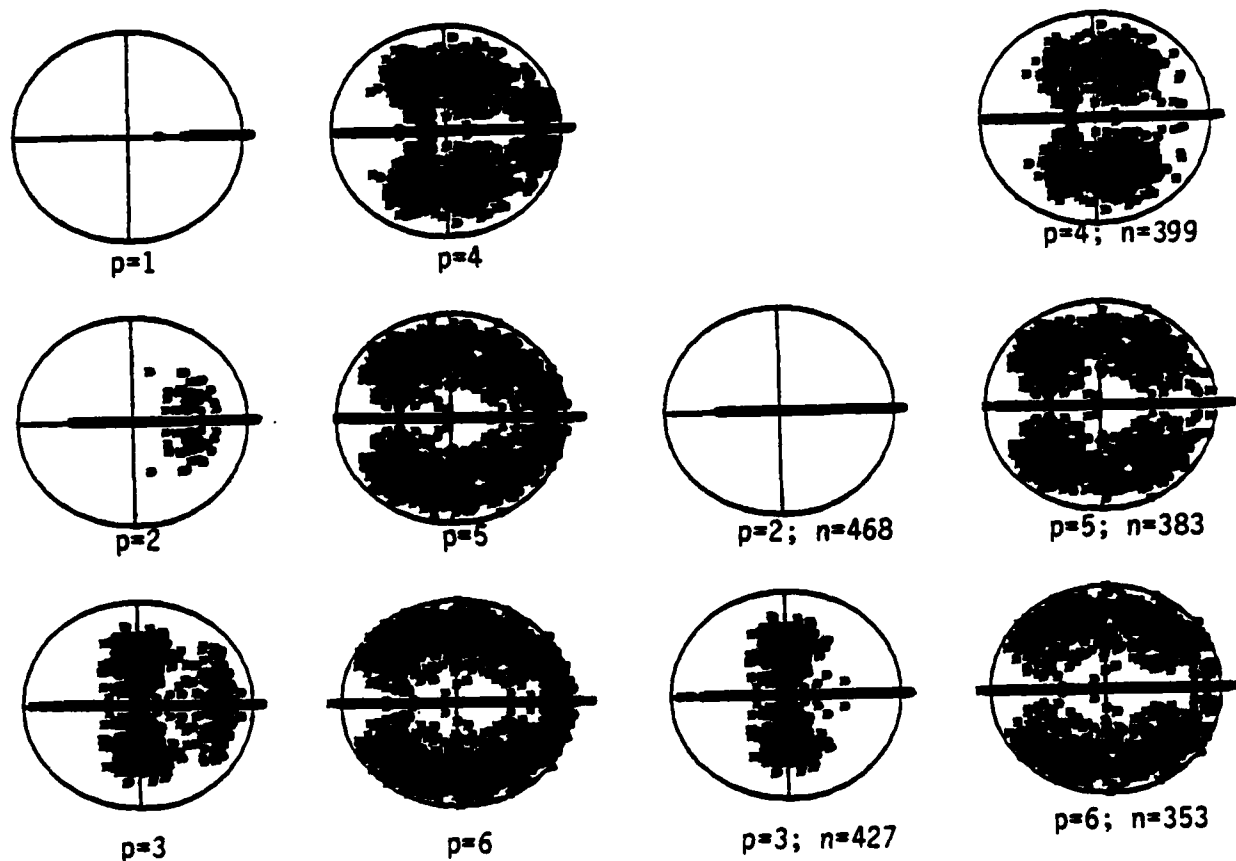


Fig. 11a

Fig. 11b

$a = .95$; Noise-free case ; $N = 25$

Superposition of the Zeros of the Estimated Prediction-Error-Filter Polynomials obtained from the estimated full-rank Forward covariance matrix for :

a) All 500 trials. b) Cases in which each set of polynomial zeros has at least one positive real zero.

[p : predictor length ; n : number of cases out of 500 trials]

(cf. Table 7)

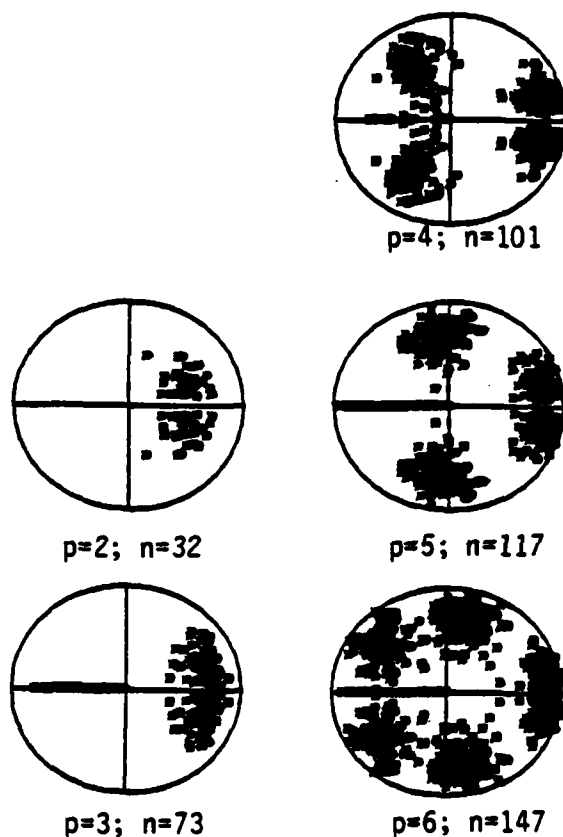


Fig. 11c

$a = .95$; Noise-free case ; $N = 25$

Superposition of the Zeros of the Estimated Prediction-Error-Filter Polynomials obtained from the estimated full-rank Forward covariance matrix. c) Cases in which each set of polynomial zeros has no positive real zero.

[p : predictor length ; n : number of cases out of 5000 trials]

(cf. Table 7)

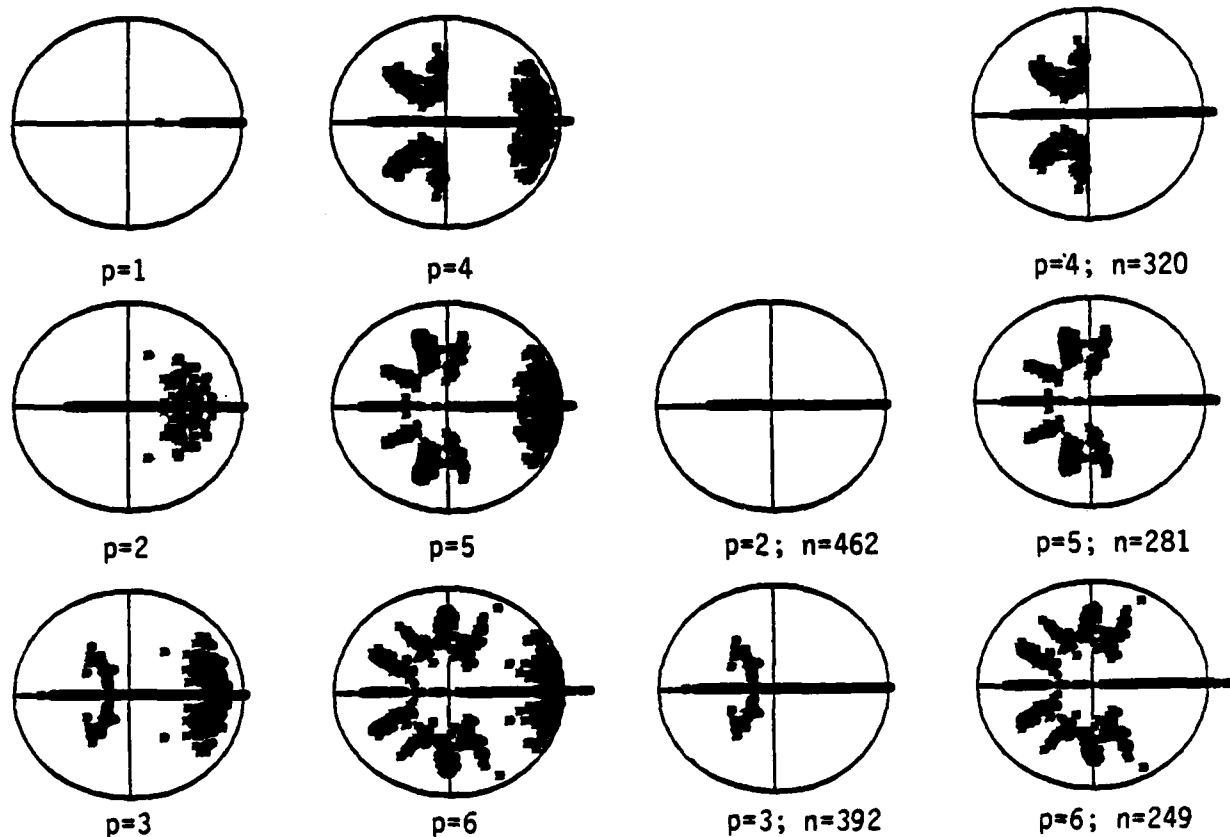


Fig. 12a

Fig. 12b

$a = .95$; Noise-free case ; $N = 25$

Superposition of the Zeros of the Estimated Prediction-Error-Filter Polynomials obtained from the rank-2 approximation of the estimated Forward-Backward covariance matrix for : a) All 500 trials. b) Cases in which each set of polynomial zeros has at least one positive real zero.

[p : predictor length ; n : number of cases out of 500 trials]

(cf. Table 3)

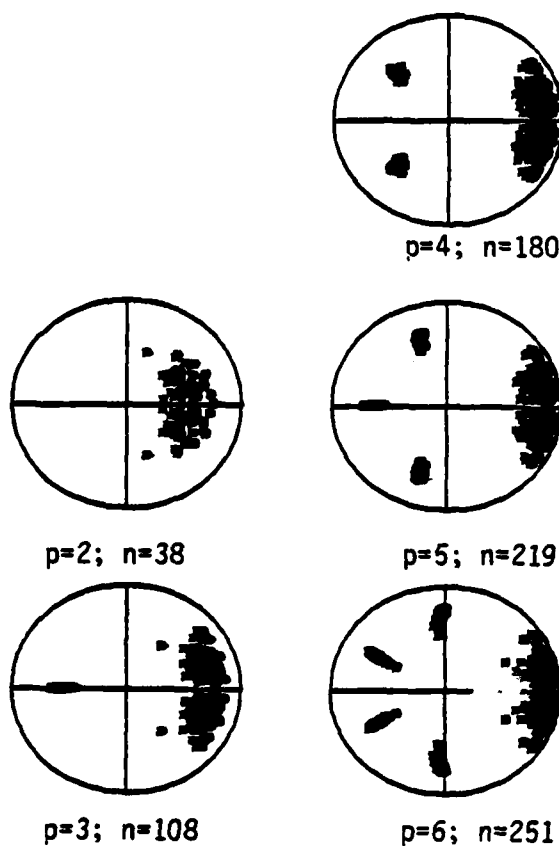


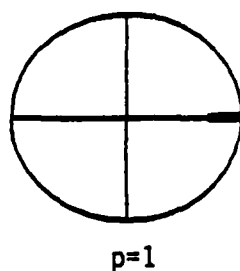
Fig. 12c

$a = .95$; Noise-free case ; $N = 25$

Superposition of the Zeros of the Estimated Prediction-Error-Filter Polynomials obtained from the rank-2 approximation of the estimated Forward-Backward covariance matrix. c) Cases in which each set of polynomial zeros has no positive real zero.

[p : predictor length ; n : number of cases out of 500 trials]

(cf. Table 3)



$p=1$

Fig. 13

$a = .95$: Noise-free case ; $N = 100$

Superposition of the Zeros estimated by the MLE over 500 trials.

(cf. Table 9)

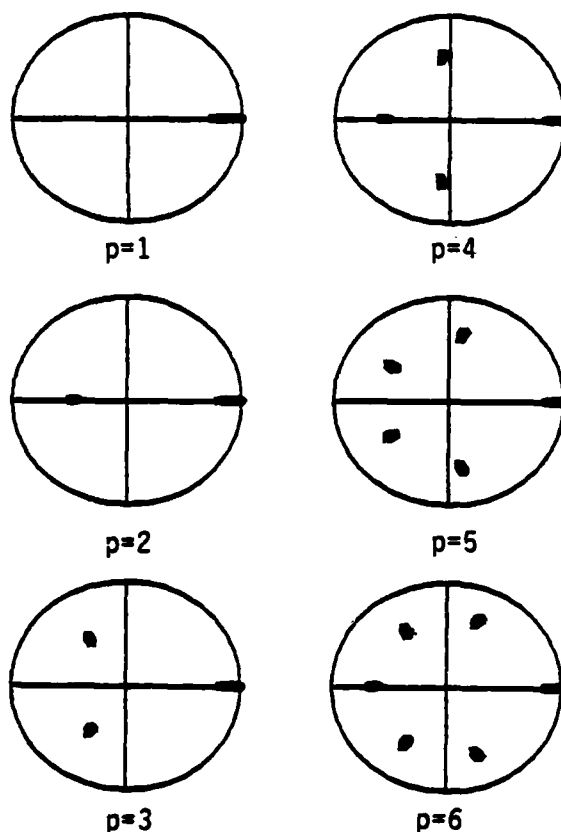


Fig. 14

$a = .95$; Noise-free case ; $N = 100$

Superposition of the Zeros of the Estimated Prediction-Error-Filter Polynomials obtained from the rank-1 approximation of the estimated Forward-Backward covariance matrix. Note : For all 500 trials, each set of polynomial zeros has only one positive real zero.

[p : predictor length]

(cf. Table 10)

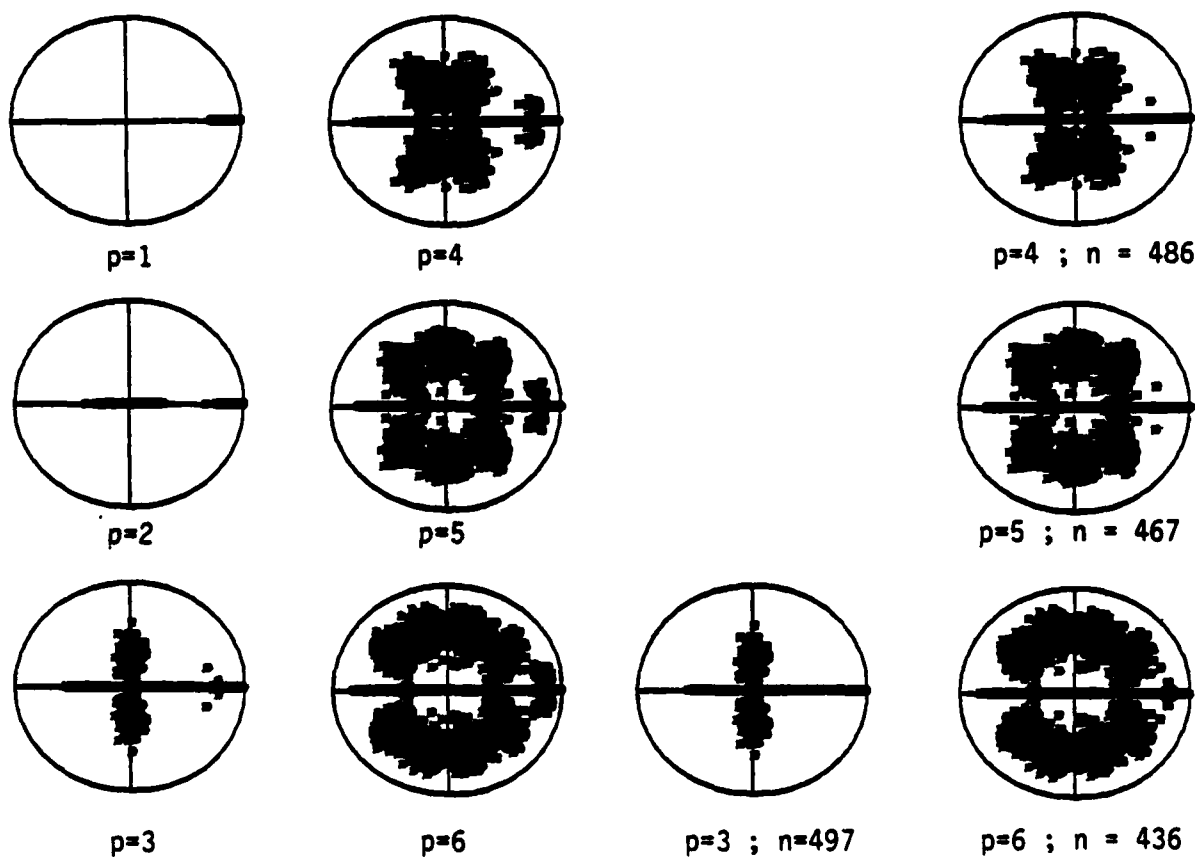


Fig. 15a

Fig. 15b

$a = .95$; Noise-free case ; $N = 100$

Superposition of the Zeros of the Estimated Prediction-Error-Filter Polynomials obtained from the estimated full-rank Forward-Backward covariance matrix for :

a) All 5000 trials. b) Cases in which each set of polynomial zeros has at least one positive real zero.

[p : predictor length ; n : number of cases out of 500 trials]

(cf. Table 11)

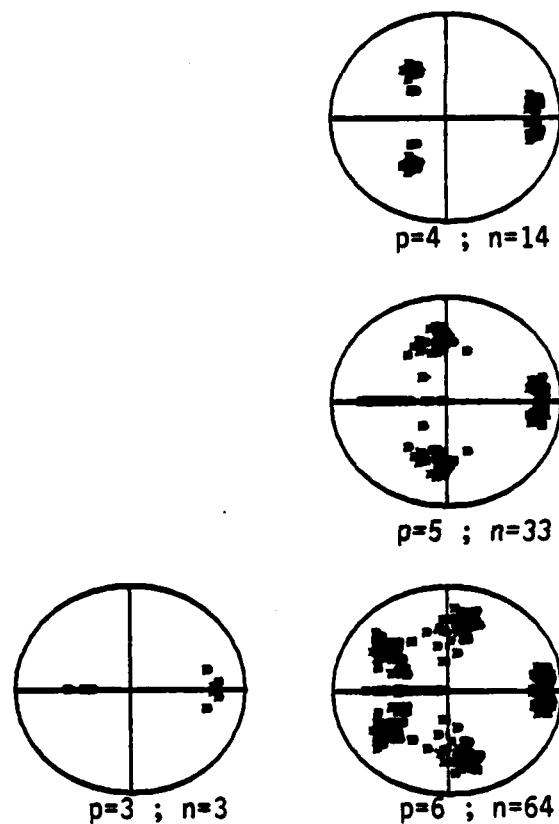


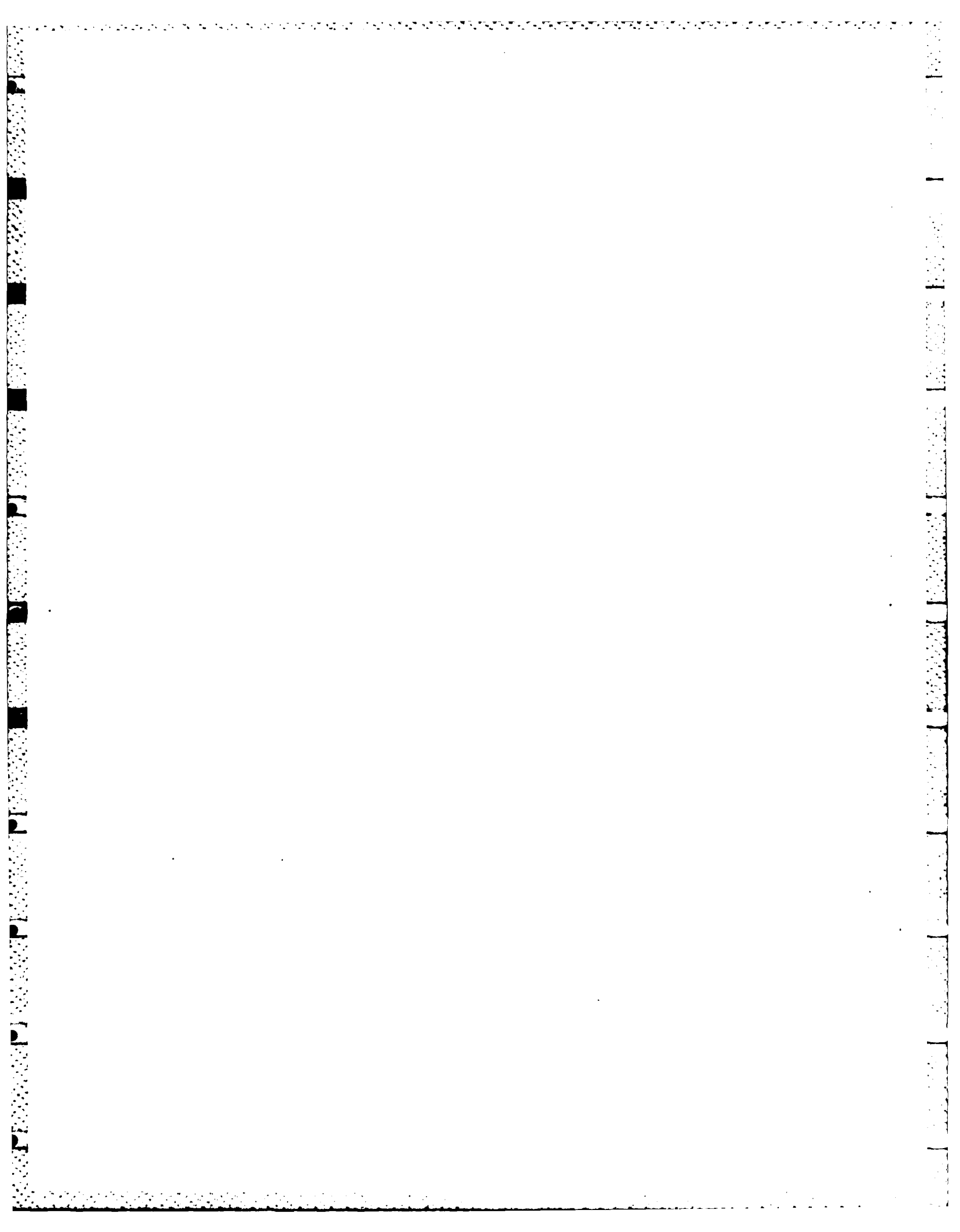
Fig. 15c

$a = .95$; Noise-free case ; $N = 100$

Superposition of the Zeros of the Estimated Prediction-Error-Filter Polynomials obtained from the estimated full-rank Forward-Backward covariance matrix. c) Cases in which each set of polynomial zeros has no positive real zero.

[p : predictor length ; n : number of case out of 500 trials]

(cf. Table 11)



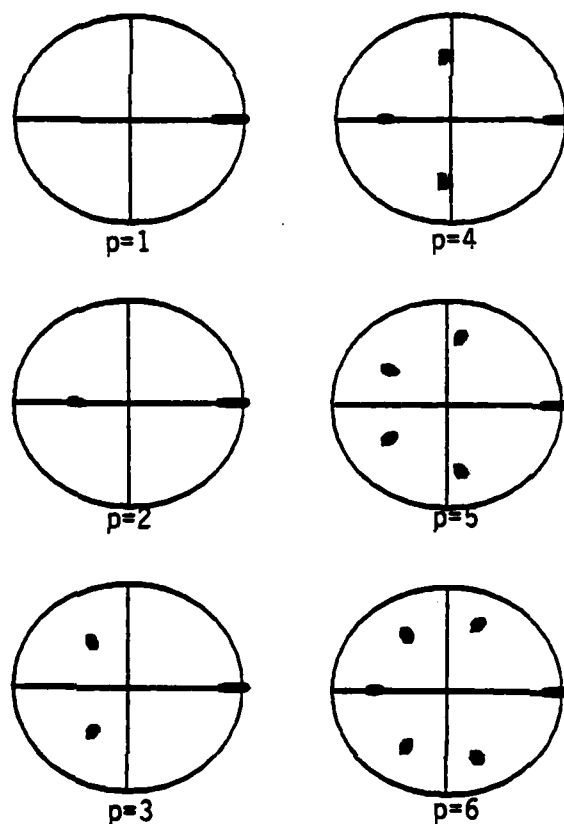


Fig. 16

$a = .95$; Noise-free case ; $N = 100$

Superposition of the Zeros of the Estimated Prediction-Error-Filter Polynomials obtained from the rank-1 approximation of the estimated Forward covariance matrix. Note : for all 500 trials, each set of polynomial zeros has only one positive real zero.

[p : predictor length]

(cf. Table 12)

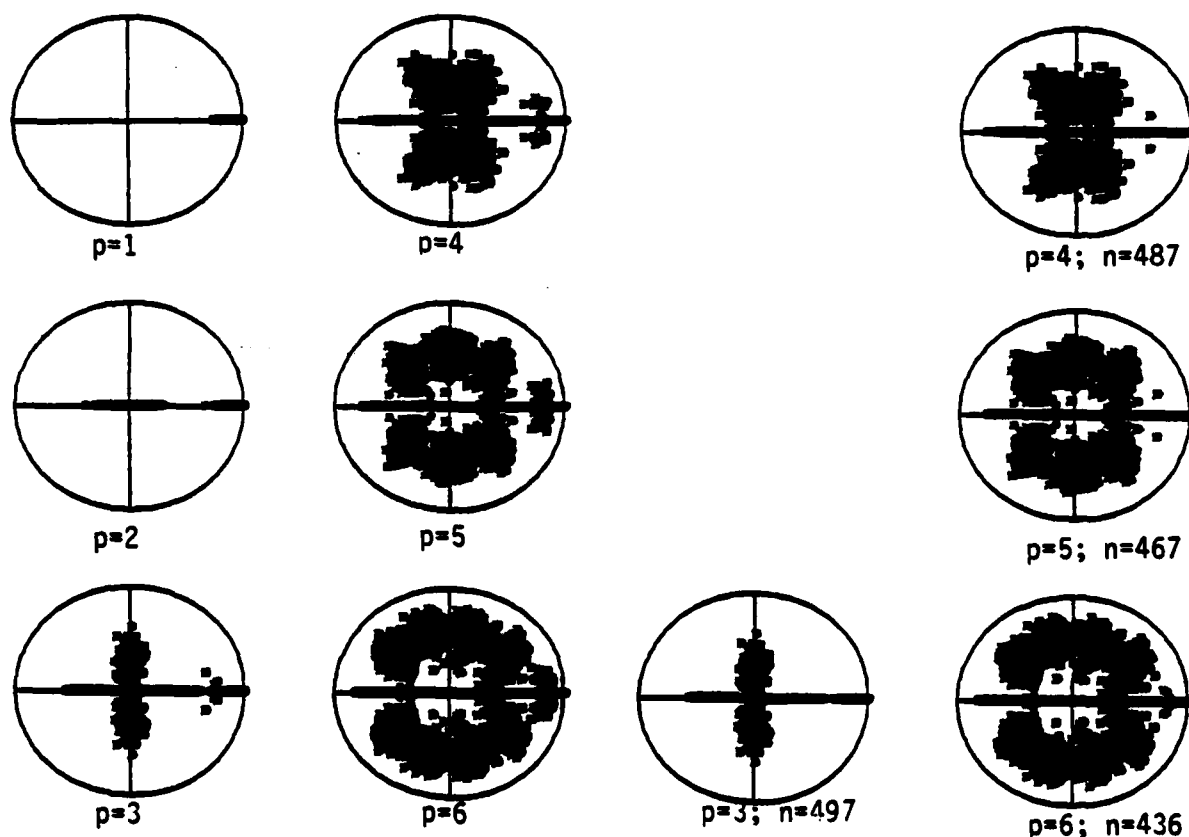


Fig. 17a

Fig. 17b

$a = .95$; Noise-free case ; $N = 100$

Superposition of the Zeros of the Estimated Prediction-Error-Filter Polynomials obtained from the estimated full-rank Forward covariance matrix for :

a) All 500 trials. b) Cases in which each set of polynomial zeros has at least one positive real zero.

[p : predictor length ; n : number of cases out of 500 trials]

(cf. Table 13)

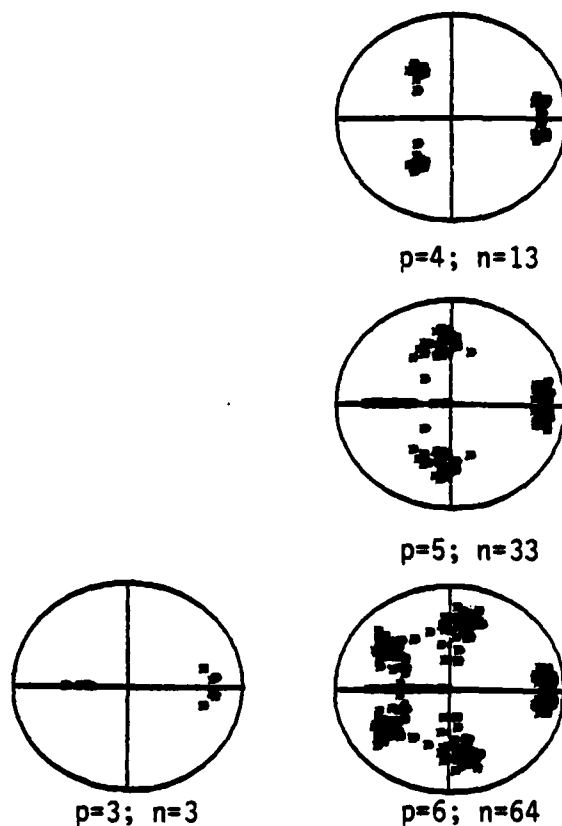


Fig. 17c

$a = .95$; Noise-free case ; $N = 100$

Superposition of the Zeros of the Estimated Prediction-Error-Filter Polynomials obtained from the estimated full-rank Forward covariance matrix. c) Cases in which each set of polynomial zeros has no positive real zero.

[p : predictor length ; n : number of cases out of 500 trials]

(cf. Table 13)

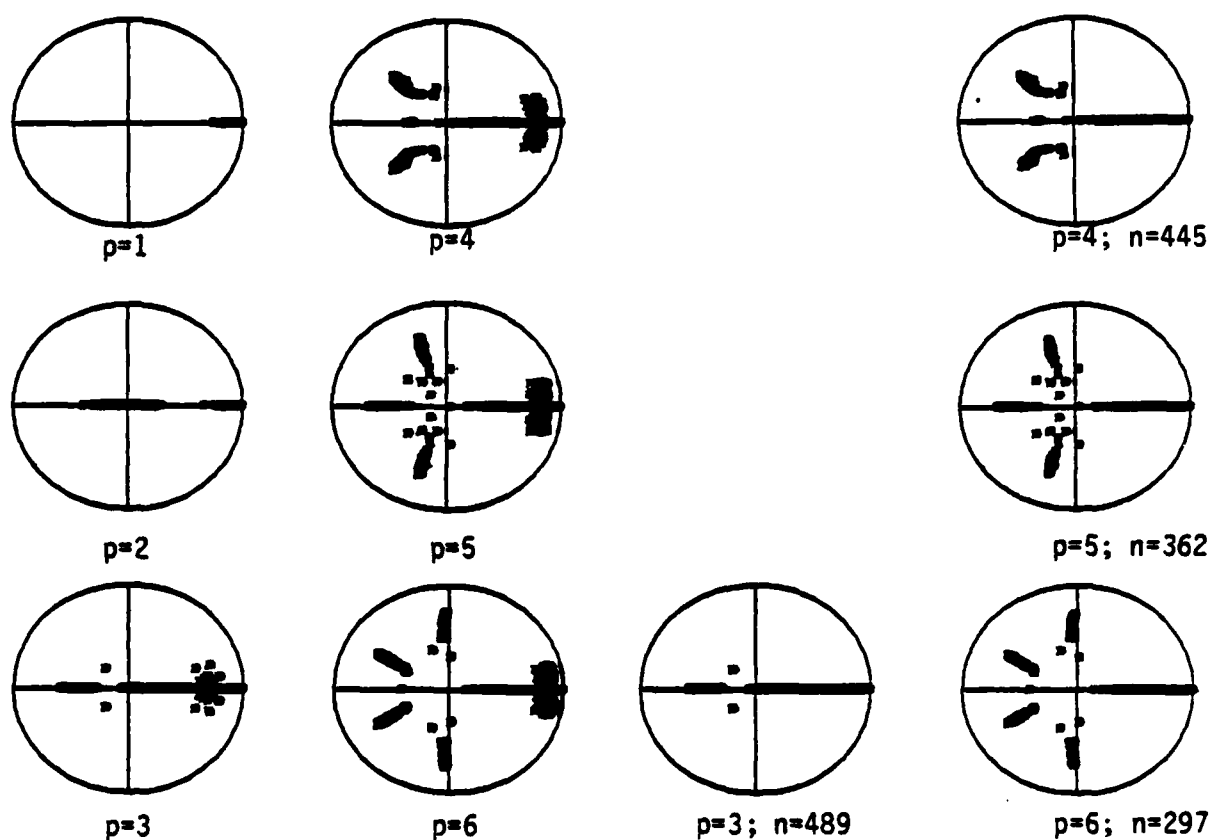


Fig. 13a

Fig. 13b

$a = .95$; Noise-free case ; $N = 100$

Superposition of the Zeros of the Estimated Prediction-Error-Filter Polynomials obtained from the rank-2 approximation of the estimated Forward-Backward covariance matrix for : a) All 500 trials. b) Cases in which each set of polynomial zeros has at least one positive real zero.

[p : predictor length ; n : number of cases out of 500 trials]

(cf. Table 14)

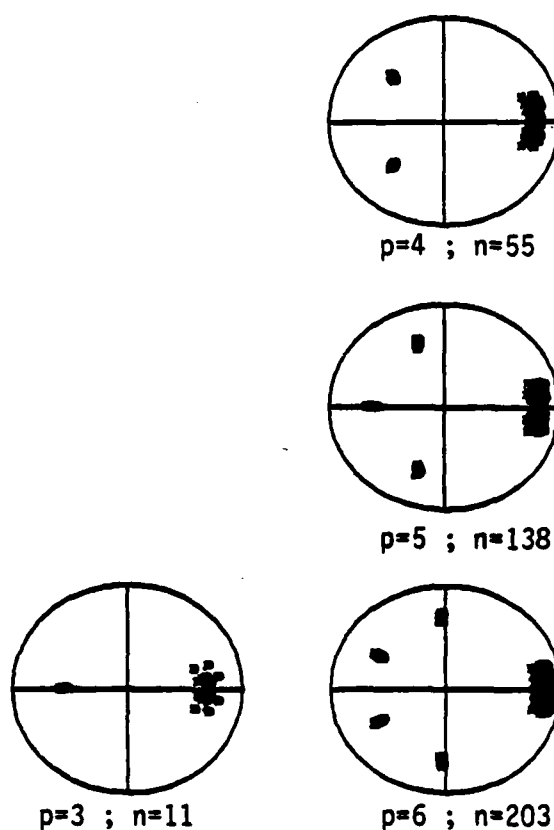


Fig 18c

$a = .95$; Noise-free case ; $N = 100$

Superposition of the Zeros of the Estimated Prediction-Error-Filter Polynomials obtained from the rank-2 approximation of the estimated Forward-Backward covariance matrix. c) Cases in which each set of polynomial zeros has no positive real zero.

[p : predictor length ; n : number of cases out of 500 trials]

(cf. Table 14)

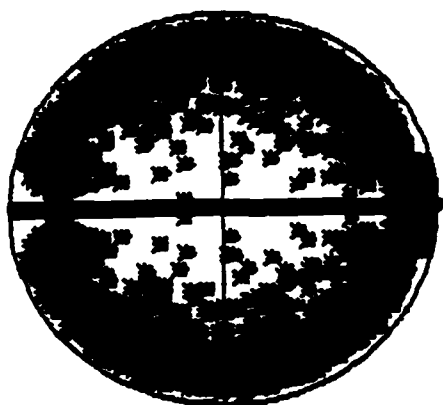


Fig. 19a

p=13

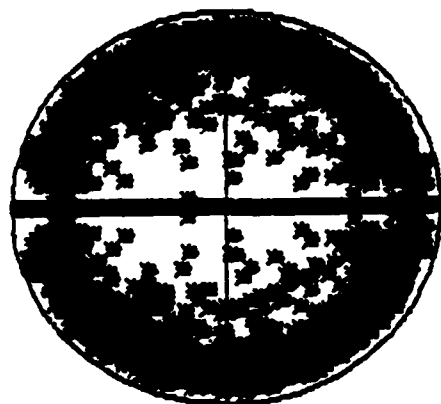


Fig. 19b

p=13 ; n=408

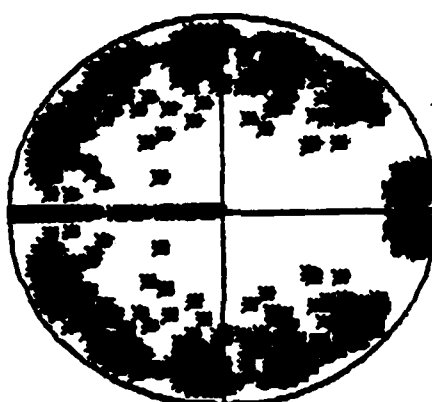


Fig. 19c

p=13 ; n=92

a = .95 ; SNR = 0 db ; N = 100

Superposition of the Zeros of the Estimated Prediction-Error-Filter Polynomials obtained from the estimated full-rank Forward-Backward covariance matrix for :

- a) All 500 trials. b) Cases in which each set of polynomial zeros has at least one positive real zero. c) Cases in which each set of polynomial zeros has no positive real zero.

[p : predictor length ; n : number of cases out of 500 trials]

(cf. Table 15)

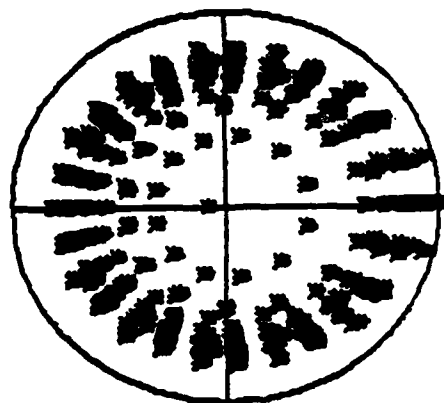


Fig. 20a

p=13

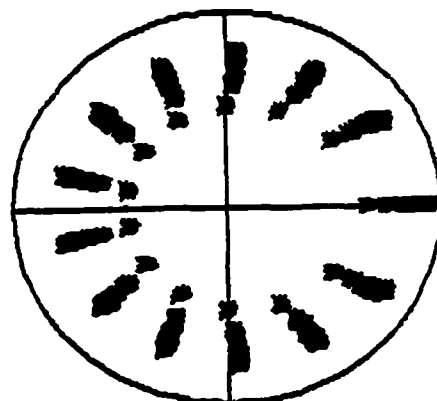


Fig. 20b

p=13 ; n=437

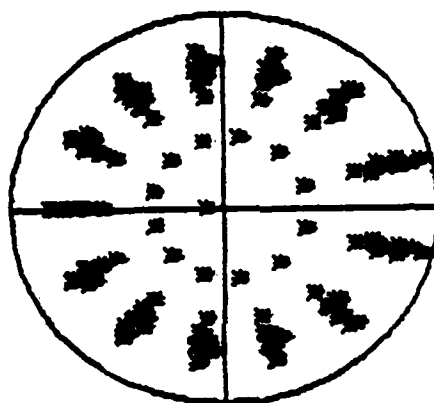


Fig. 20c

p=13 ; n=13

$a = .95$; SNR = 0 db ; $N = 100$

Superposition of the Zeros of the Estimated Prediction-Error-Filter Polynomials obtained from the rank-1 approximation of the estimated Forward-Backward covariance matrix for : a) All 500 trials. b) Cases in which each set of polynomial zeros has at least one positive real zero. c) Cases in which each set of polynomial zeros has no positive real zero.

[p : predictor length ; n : number of cases out of 500 trials]

(cf. Table 16)

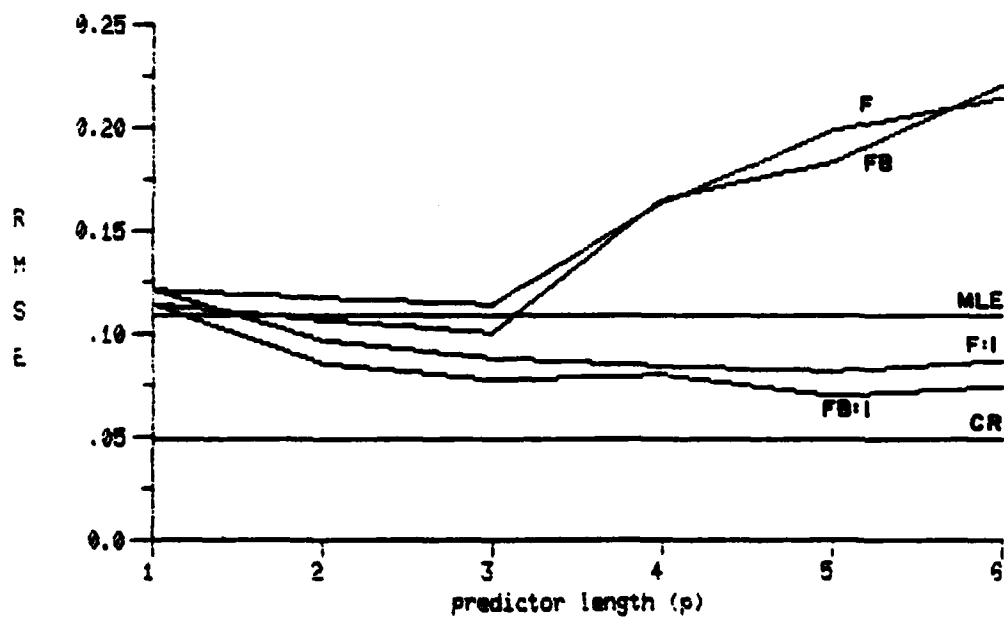


Fig. 21

Performance (in rmse) of estimators of the first-order AR parameter ($a=.95$) from segments of $N=25$ noise-free samples of data.

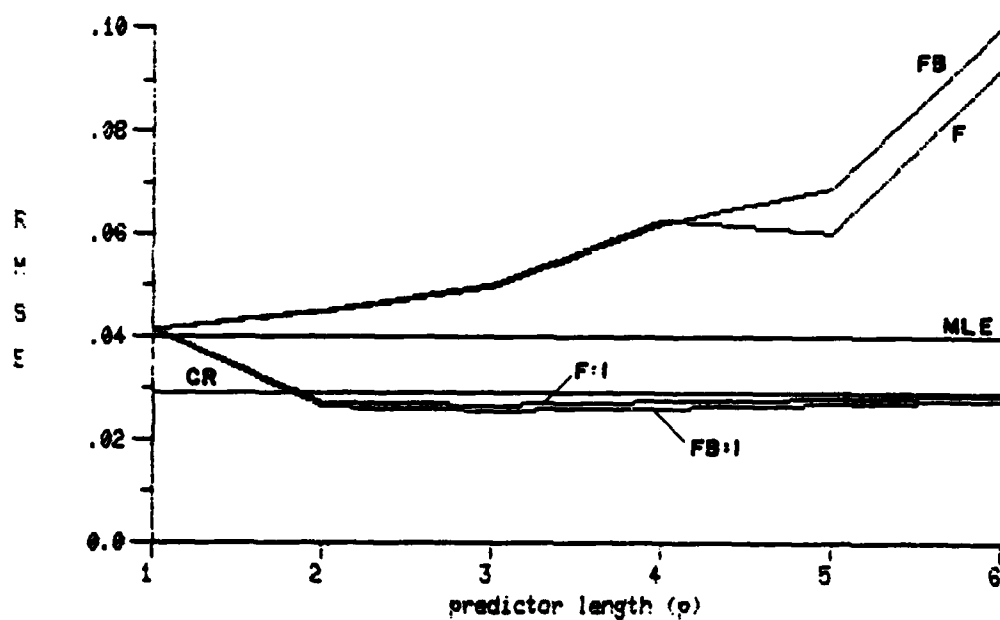


Fig. 22

Performance (in rmse) of estimators of a first-order AR parameter ($a=.95$) from segments of $N=100$ noise-free samples of data.

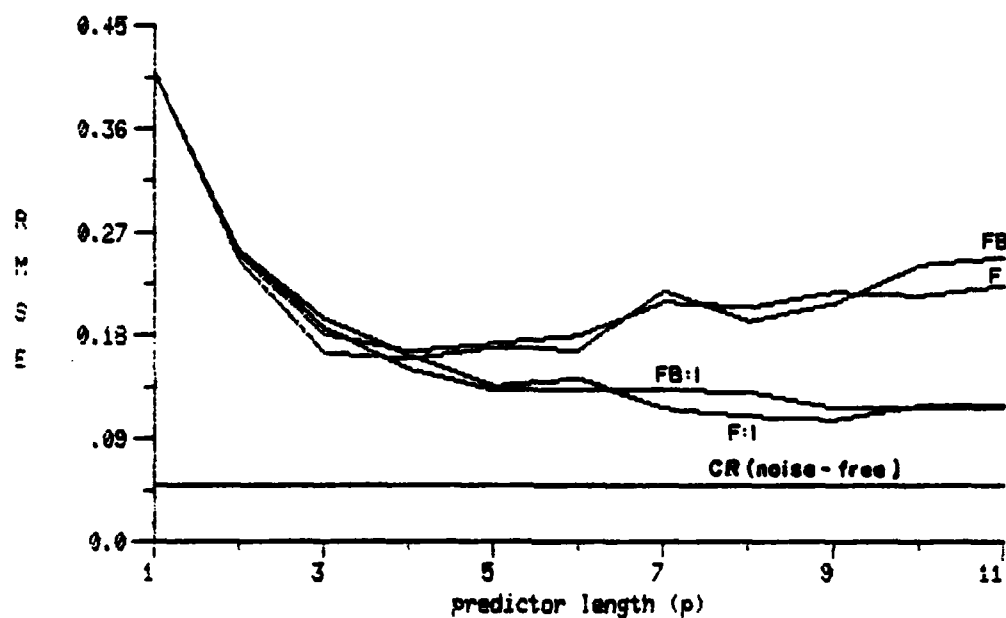


Fig. 23

Performance (in rmse) of estimators of a first-order AR parameter ($a=.95$) from segments of $N=25$ samples of data corrupted by additive white noise, such that the $SNR=5\text{db}$.

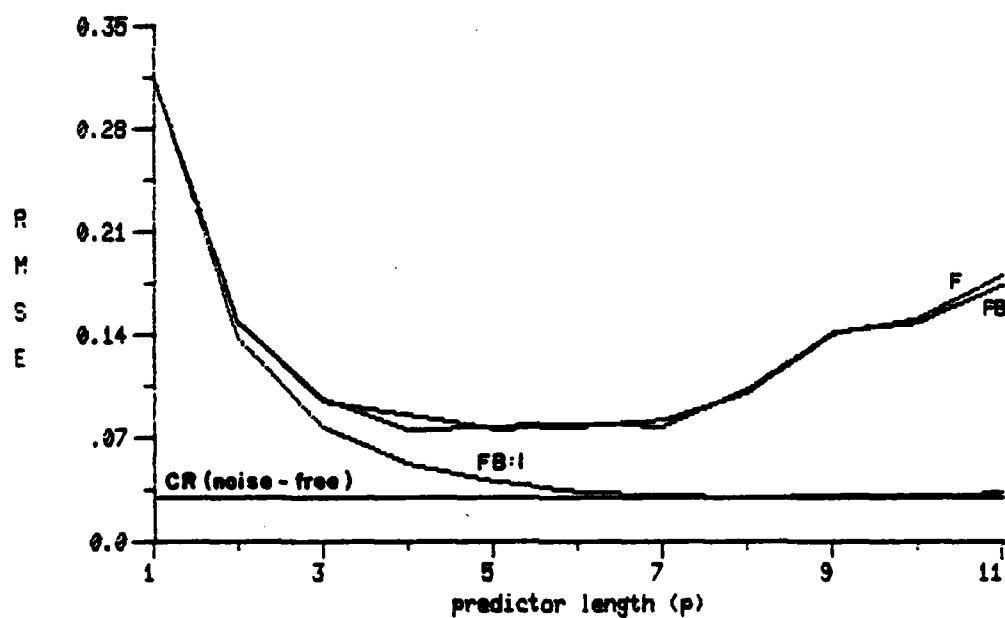


Fig. 24

Performance (in rmse) of estimators of a first-order AR parameter ($a=.95$) from segments of $N=100$ samples of data corrupted by additive white noise, such that the $SNR=5\text{db}$.

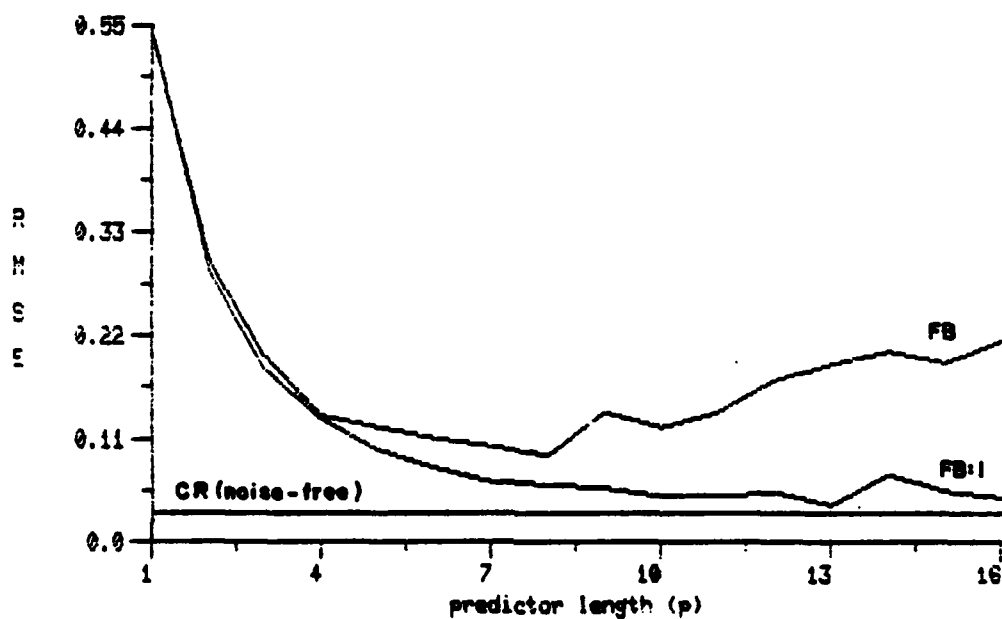


Fig. 25

Performance (in rmse) of estimators of a first-order AR parameter ($a=.95$) from segments of $N=100$ samples of data corrupted by additive white noise, such that the $SNR=0db$.

p	est.pole	bias
1	0.950	0.000
2	0.966	1.647 E-02
3	0.972	2.205 E-02
4	0.974	2.487 E-02
5	0.976	2.658 E-02
6	0.977	2.773 E-02

TABLE 1

$a=.95$; Noise-free ; $N = \infty$
 Zero selection after rank-1 approximation
 on the true correlation matrix
 (cf. Fig. 6)

p	est.pole	bias
1	0.990	0.000
2	0.993	3.326 E-03
3	0.994	4.438 E-03
4	0.995	4.995 E-03
5	0.995	5.330 E-03
6	0.996	5.554 E-03

TABLE 2

$a=.99$; Noise-free ; $N = \infty$
 Zero selection after rank-1 approximation
 on the true correlation matrix

sigma = 9.463 E-02
mean = 0.897
bias = -5.315 E-02
rms = 1.085 E-01

TABLE 3

a=.95 ; Noise-free ; N=25

Note: C-R (sigma) = 4.831 E-02

MAXIMUM LIKELIHOOD ESTIMATION

(of. Fig. 7)

p	sigma	mean	bias	rms	uns	out	Ops	lps	2ps
1	9.791 E-02	0.892	-5.741 E-02	1.135 E-01	0	0	0	500	0
2	8.027 E-02	0.922	-2.737 E-02	8.481 E-02	0	0	0	500	0
3	7.416 E-02	0.931	-1.833 E-02	7.640 E-02	0	0	1	499	0
4	7.806 E-02	0.936	-1.377 E-02	7.926 E-02	0	0	4	496	0
5	6.823 E-02	0.940	-9.422 E-03	6.888 E-02	6	6	9	492	0
6	7.281 E-02	0.941	-8.693 E-03	7.333 E-02	6	6	13	487	0

TABLE 4

a=.95 ; Noise-free case ; N=25

[Note : C-R min standard deviation (sigma) = 4.831 E-02]

Zero selection after rank-1 approximation on the estimated

Forward-Backward covariance matrix

(of. Figs. 5a, 5b, 5c)

p	sigma	mean	bias	rms	uns	out	Ops	lps	2ps
1	3.791 E-02	0.892	-5.741 E-02	1.135 E-01	0	0	0	500	0
2	9.268 E-02	0.898	-5.134 E-02	1.061 E-01	1	1	38	228	234
3	9.224 E-02	0.911	-3.820 E-02	9.984 E-02	1	1	76	254	169
4	1.566 E-01	0.899	-5.057 E-02	1.645 E-01	7	7	111	210	168
5	1.752 E-01	0.896	-5.344 E-02	1.831 E-01	15	15	122	242	126
6	2.078 E-01	0.879	-7.102 E-02	2.196 E-01	17	17	158	206	128

TABLE 5

$\alpha=0.35$; Noise-free case ; $N=25$

[Note : C-R min standard deviation (sigma) = 4.331 E-02]

Zero selection from the estimated full-rank

Forward-Backward covariance matrix

(cf. Figs. 9a,9b,9c)

p	sigma	mean	bias	rms	uns	out	Ops	lps	2ps
1	1.035 E-01	0.888	-6.111 E-02	1.202 E-01	37	37	0	500	0
2	8.940 E-02	0.915	-3.429 E-02	9.576 E-02	71	71	0	500	0
3	8.277 E-02	0.922	-2.737 E-02	8.718 E-02	85	85	1	499	0
4	7.962 E-02	0.925	-2.475 E-02	8.338 E-02	98	98	3	497	0
5	7.774 E-02	0.927	-2.297 E-02	8.107 E-02	107	107	8	492	0
6	8.218 E-02	0.924	-2.512 E-02	8.594 E-02	115	115	12	484	4

TABLE 6

$\alpha=0.95$; Noise-free case ; $N=25$

[Note : C-R min standard deviation (sigma) = 4.331 E-02]

Zero selection after rank-1 approximation on the estimated

Forward covariance matrix

(cf. Figs. 10a,10b,10c)

p	sigma	mean	bias	rms	uns	out	Ops	1ps	2ps
1	1.035 E-01	0.888	-6.111 E-02	1.202 E-01	37	37	0	500	0
2	1.003 E-01	0.890	-5.917 E-02	1.164 E-01	47	47	32	239	229
3	1.034 E-01	0.903	-4.708 E-02	1.136 E-01	52	52	73	271	156
4	1.515 E-01	0.890	-5.946 E-02	1.628 E-01	63	63	101	229	157
5	1.855 E-01	0.881	-6.823 E-02	1.977 E-01	74	78	117	256	119
6	1.984 E-01	0.871	-7.884 E-02	2.135 E-01	71	90	147	231	115

TABLE 7

a=.95 ; Noise-free case ; N=25

[Note : C-R min standard deviation (sigma) = 4.331 E-02]

Zero selection from the estimated full-rank

Forward covariance matrix

(cf. Figs. 11a,11b,11c)

p	sigma	mean	bias	rms	uns	out	Ops	1ps	2ps
1	9.791 E-02	0.892	-5.741 E-02	1.135 E-01	0	0	0	500	0
2	9.268 E-02	0.898	-5.184 E-02	1.061 E-01	1	1	38	229	234
3	8.259 E-02	0.910	-3.996 E-02	9.175 E-02	1	1	108	78	314
4	7.220 E-02	0.929	-2.030 E-02	7.500 E-02	3	3	180	60	260
5	5.969 E-02	0.941	-8.988 E-03	6.036 E-02	18	18	219	50	231
6	5.046 E-02	0.950	3.600 E-03	5.046 E-02	26	26	251	50	199

TABLE 8

a=.95 ; Noise-free case ; N=25

[Note : C-R min standard deviation (sigma) = 4.331 E-02

Zero selection after rank-2 approximation on the estimated

Forward-Backward covariance matrix

(cf. Figs. 12a,12b,12c)

sigma = 3.652 E-02

mean = 0.934

bias = -1.615 E-02

rms = 3.993 E-02

TABLE 9

$\alpha = .95$; Noise-free ; $N=100$

Note: C-R (sigma) = 2.882 E-02

MAXIMUM LIKELIHOOD ESTIMATION

(cf. Fig. 13)

p	sigma	mean	bias	rms	uns	out	Ops	lps	2ps
1	3.746 E-02	0.933	-1.680 E-02	4.106 E-02	0	0	0	500	0
2	2.567 E-02	0.955	5.318 E-03	2.622 E-02	0	0	0	500	0
3	2.195 E-02	0.962	1.264 E-02	2.533 E-02	0	0	0	500	0
4	2.028 E-02	0.966	1.616 E-02	2.594 E-02	0	0	0	500	0
5	1.952 E-02	0.968	1.814 E-02	2.665 E-02	0	0	0	500	0
6	1.917 E-02	0.969	1.941 E-02	2.729 E-02	0	0	0	500	0

TABLE 10

$\alpha = .95$; Noise-free case ; $N=100$

[Note : C-R min standard deviation (sigma) = 2.882 E-02]

Zero selection after rank-1 approximation on the estimated
Forward-Backward covariance matrix

(cf. Fig. 14)

p	sigma	mean	bias	rms	uns	out	Ops	1ps	2ps
1	3.746 E-02	0.933	-1.680 E-02	4.106 E-02	0	0	0	500	0
2	4.396 E-02	0.932	-1.751 E-02	4.455 E-02	0	0	0	260	240
3	4.586 E-02	0.931	-1.859 E-02	4.948 E-02	0	0	3	240	251
4	5.802 E-02	0.929	-2.075 E-02	6.162 E-02	0	0	14	224	246
5	6.548 E-02	0.929	-2.024 E-02	6.854 E-02	0	0	33	222	221
6	1.004 E-01	0.925	-2.464 E-02	1.034 E-01	0	0	64	214	201

TABLE 11

a=.95 ; Noise-free case ; N=100

[Note : C-R min standard deviation (sigma) = 2.882 E-02]

Zero selection from the estimated full-rank

Forward-Backward covariance matrix

(cf. Figs. 15a,15b,15c)

p	sigma	mean	bias	rms	uns	out	Ops	1ps	2ps
1	3.771 E-02	0.933	-1.651 E-02	4.117 E-02	0	0	0	500	0
2	2.641 E-02	0.955	5.482 E-03	2.697 E-02	4	4	0	500	0
3	2.330 E-02	0.962	1.255 E-02	2.647 E-02	9	9	0	500	0
4	2.186 E-02	0.966	1.597 E-02	2.707 E-02	12	12	0	500	0
5	2.117 E-02	0.968	1.789 E-02	2.772 E-02	13	13	0	500	0
6	2.082 E-02	0.969	1.913 E-02	2.828 E-02	15	15	0	500	0

TABLE 12

a=.95 ; Noise-free case ; N=100

[Note : C-R min standard deviation (sigma) = 2.382 E-02]

Zero selection after rank-1 approximation on the estimated

Forward covariance matrix

(cf. Fig. 16)

p	sigma	mean	bias	rms	uns	out	Ops	1ps	2ps
1	3.771 E-02	0.933	-1.651 E-02	4.117 E-02	0	0	0	500	0
2	4.137 E-02	0.932	-1.726 E-02	4.483 E-02	2	2	0	261	239
3	4.648 E-02	0.931	-1.842 E-02	4.999 E-02	1	1	3	243	249
4	5.871 E-02	0.929	-2.073 E-02	6.226 E-02	2	2	13	226	246
5	5.707 E-02	0.931	-1.839 E-02	5.996 E-02	1	1	33	220	221
6	8.951 E-02	0.928	-2.172 E-02	9.211 E-02	3	3	64	215	199

TABLE 13

a=.95 ; Noise-free case ; N=100

[Note : C-R min standard deviation (sigma) = 2.382 E-02]

Zero selection from the estimated full-rank

Forward covariance matrix

(cf. Figs. 17a,17b,17c)

p	sigma	mean	bias	rms	uns	out	Ops	1ps	2ps
1	3.746 E-02	0.933	-1.680 E-02	4.106 E-02	0	0	0	500	0
2	4.096 E-02	0.932	-1.751 E-02	4.455 E-02	0	0	0	260	240
3	4.762 E-02	0.924	-2.589 E-02	5.421 E-02	0	0	11	6	483
4	4.730 E-02	0.922	-2.797 E-02	5.496 E-02	0	0	55	4	441
5	3.940 E-02	0.929	-2.020 E-02	4.428 E-02	0	0	138	3	359
6	3.506 E-02	0.934	-1.534 E-02	3.827 E-02	0	0	203	1	296

TABLE 14

a=.95 ; Noise-free case ; N=100

[Note : C-R min standard deviation (sigma) = 2.382 E-02]

Zero selection after rank-2 approximation on the estimated

Forward-Backward covariance matrix

(cf. Figs. 18a,18b,18c)

p	sigma	mean	bias	rms	uns	out	Opz	1pz	2pz
1	1.642 E-01	0.431	-5.191 E-01	5.444 E-01	0	0	1	499	0
2	1.416 E-01	0.683	-2.666 E-01	3.019 E-01	0	0	1	495	4
3	1.141 E-01	0.790	-1.596 E-01	1.962 E-01	0	0	3	475	22
4	8.379 E-02	0.847	-1.030 E-01	1.328 E-01	0	0	14	446	40
5	9.195 E-02	0.872	-7.752 E-02	1.203 E-01	0	0	16	419	65
6	8.957 E-02	0.887	-6.265 E-02	1.093 E-01	0	0	19	367	112
7	8.743 E-02	0.900	-4.965 E-02	1.005 E-01	0	0	23	353	119
8	9.118 E-02	0.911	-3.894 E-02	9.915 E-02	0	0	51	296	144
9	1.286 E-01	0.904	-4.555 E-02	1.364 E-01	0	0	55	284	153
10	1.143 E-01	0.911	-3.870 E-02	1.206 E-01	0	0	80	242	162
11	1.311 E-01	0.912	-3.760 E-02	1.364 E-01	0	0	100	248	132
12	1.623 E-01	0.900	-4.969 E-02	1.697 E-01	0	0	105	228	150
13	1.768 E-01	0.889	-6.147 E-02	1.872 E-01	0	0	92	234	143
14	1.905 E-01	0.885	-6.462 E-02	2.012 E-01	0	0	123	205	153
15	1.780 E-01	0.887	-6.337 E-02	1.889 E-01	0	0	135	216	121
16	2.003 E-01	0.878	-7.207 E-02	2.128 E-01	0	0	142	220	119

Table 15

$\alpha = .95$; $N = 100$; $SNR = 0$ db

Zero selection from the estimated full-rank

Forward-Backward covariance matrix

(cf. Figs. 19a, 19b, 19c for $p=13$)

p	sigma	mean	bias	rms	uns	out	Ops	1pz	2pz
1	1.641 E-01	0.431	-5.191 E-01	5.444 E-01	0	0	1	499	0
2	1.350 E-01	0.691	-2.589 E-01	2.920 E-01	0	0	0	500	0
3	1.030 E-01	0.799	-1.504 E-01	1.823 E-01	0	0	1	499	0
4	8.500 E-02	0.854	-9.588 E-02	1.281 E-01	0	0	0	500	0
5	7.233 E-02	0.885	-6.440 E-02	9.684 E-02	0	0	1	499	0
6	6.405 E-02	0.906	-4.385 E-02	7.762 E-02	0	0	0	500	0
7	5.556 E-02	0.920	-2.955 E-02	6.293 E-02	0	0	2	498	0
8	5.601 E-02	0.928	-2.134 E-02	5.994 E-02	0	0	0	500	0
9	5.399 E-02	0.936	-1.414 E-02	5.581 E-02	0	0	1	499	0
10	4.648 E-02	0.942	-7.789 E-03	4.713 E-02	0	0	4	496	0
11	4.682 E-02	0.946	-3.997 E-03	4.699 E-02	0	0	6	494	0
12	5.169 E-02	0.947	-2.802 E-03	5.177 E-02	0	0	5	494	1
13	3.812 E-02	0.953	3.640 E-03	3.829 E-02	0	0	13	487	0
14	7.052 E-02	0.950	-2.086 E-04	7.052 E-02	0	0	8	492	0
15	5.315 E-02	0.953	3.273 E-03	5.325 E-02	0	0	10	490	0
16	4.604 E-02	0.954	3.712 E-03	4.619 E-02	0	0	20	480	0

Table 16

$\alpha = .95$; $N = 100$; $SNR = 0. \text{ db}$

Zero selection after rank-1 approximation on the estimated
Forward-Backward covariance matrix

(cf. Figs. 20a,20b,20c for $p=13$)

END

FILMED

9-84

DTIC