

AD-A111 854

SYSTEMS AND APPLIED SCIENCES CORP RIVERDALE MD

F/G 12/1

AN INVESTIGATION OF THREE METHODS FOR SPECTRAL REPRESENTATION 0--ETC(U)

AUG 81 I M HALBERSTAM

F19628-81-C-0039

UNCLASSIFIED

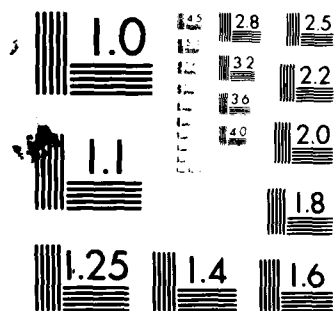
SCIENTIFIC-1

AFGL-TR-81-0234

NL

1
2
3
4
5
6
7
8
9
10
11
12

END
DATE
FILMED
4 82
DTIC



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

AFGL-TR-81-0234

ADA111854

AN INVESTIGATION OF THREE METHODS FOR
SPECTRAL REPRESENTATION OF RANDOMLY
DISTRIBUTED DATA

Isidore M. Halberstam

Systems and Applied Sciences Corporation
6811 Kenilworth Avenue
Riverdale, Maryland 20737

21 August 1981

Scientific Report No. 1

Approved for public release; distribution unlimited

AIR FORCE GEOPHYSICS LABORATORY
AIR FORCE SYSTEMS COMMAND
UNITED STATES AIR FORCE
HANSCOM AFB, MASSACHUSETTS 01731

DTIC FILE COPY

DTIC
SELECTED
MAR 10 1982
H

82 03 09 088

UNCLASSIFIED

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

REPORT DOCUMENTATION PAGE		READ INSTRUCTIONS BEFORE COMPLETING FORM
1. REPORT NUMBER AFGL-TR-81-0234	2. GOVT ACCESSION NO.	3. RECIPIENT'S CATALOG NUMBER
4. TITLE (and Subtitle) AN INVESTIGATION OF THREE METHODS FOR SPECTRAL REPRESENTATION OF RANDOMLY DISTRIBUTED DATA		5. TYPE OF REPORT & PERIOD COVERED Scientific Report No. 1
7. AUTHOR(s) Isidore M. Halberstam		6. PERFORMING ORG. REPORT NUMBER
9. PERFORMING ORGANIZATION NAME AND ADDRESS Systems and Applied Sciences Corporation 6811 Kenilworth Ave., P.O. Box 308 Riverdale, MD 20737		8. CONTRACT OR GRANT NUMBER(s) F19628-81-C-0039
11. CONTROLLING OFFICE NAME AND ADDRESS Air Force Geophysics Laboratory Hanscom AFB, MA 01731 Manager/Charles Burger/LY		10. PROGRAM ELEMENT, PROJECT, TASK AREA & WORK UNIT NUMBERS 61102F 667000AB
14. MONITORING AGENCY NAME & ADDRESS (if different from Controlling Office)		12. REPORT DATE 21 August 1981
		13. NUMBER OF PAGES 33
		15. SECURITY CLASS. (of this report) UNCLASSIFIED
		15a. DECLASSIFICATION/DOWNGRADING SCHEDULE
16. DISTRIBUTION STATEMENT (of this Report) Approved for public release; distribution unlimited		
17. DISTRIBUTION STATEMENT (of the abstract entered in Block 20, if different from Report)		
18. SUPPLEMENTARY NOTES		
19. KEY WORDS (Continue on reverse side if necessary and identify by block number) OBJECTIVE ANALYSIS INITIALIZATION NORMAL MODES ORTHONORMAL EXPANSIONS INTERPOLATION SPECTRAL MODELING DATA ASSIMILATION		
20. ABSTRACT (Continue on reverse side if necessary and identify by block number) Three methods are proposed for representing randomly distributed data by a truncated Fourier series. These methods involve the use of: 1. a set of linear equations based on a simple least-squares approach, 2. empirical orthogonal functions derived by the Gram-Schmidt process, and 3. a step-by-step approach where each coefficient is solved independently by subtracting the contribution from previously computed coefficients. The methods are tested against a known function (finite cosine series) for dif-		

DD FORM 1 JAN 73 1473

EDITION OF 1 NOV 65 IS OBSOLETE

UNCLASSIFIED

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

UNCLASSIFIED

SECURITY CLASSIFICATION OF THIS PAGE(When Data Entered)

ferent distributions of data and different truncation and simulated observation errors. Results show that if the empirical orthogonal functions are linear combinations of cosines, then Methods 1 and 2 yield identical coefficients. This offers two convenient methods for achieving the same goal, depending on the number of data points and the truncation. Results also indicate that Method 3 is not very sensitive to the number of data points or to their distributions while Methods 1 and 2 are, failing when the number of data points approaches the critical value for resolving the waves.

UNCLASSIFIED

SECURITY CLASSIFICATION OF THIS PAGE(When Data Entered)

TABLE OF CONTENTS

REPORT DOCUMENTATION PAGE	1
I. INTRODUCTION	4
II. ORTHOGONAL EXPANSIONS	5
III. THREE METHODS	6
IV. TESTS	8
V. RESULTS	9
VI. SUMMARY AND CONCLUSION	29



Accession For	
NTIS CRAAI	☒
DTIC TAB	☐
Unannounced	☐
Justification	
By _____	
Distribution/	
Availability Codes	
Dist	Avail and/or Special
A	

I. INTRODUCTION

Objective analyses of meteorological variables have always depended on interpolation schemes. Some methods center on local fits of data to specific grid points, where subsequent processing removes inconsistencies. Others attempt to fit given observations to a global network by spectral representation of the variables. But these usually require a priori interpolation to equally-spaced points in order to apply the spectral integration. Flattery (1971)¹, for example, used Hough functions as a basis for his analysis in the meridional direction, while using Fourier expansions in the zonal direction and empirical functions in the vertical. But he, too, required that the observed data first be interpolated to fixed volume elements before the spectral expansions could be determined. Once the spectral coefficients are known, the variables can be interpolated to any desired location. With finite difference models, this redundancy of interpolations may seem wasteful, but it has its advantages, especially because Hough functions are period-dependent and can be used to control initial imbalances. For spectral models, spectral coefficients are the required initial conditions, but, unless Hough functions are used in solving the model equations, they must first be transformed to the correct basis. This may be accomplished through relationships between the Hough functions and the desired basis or through numerical integration of the values represented at given grid points as with the finite difference models. In either case, current analysis techniques require interpolation of observed data to equally-spaced points, a step which may alter the spectral nature of the data, as pointed out by Yang and Shapiro (1973)².

1. Flattery, T. W., 1971: Spectral models for global analysis and forecasting, Proc. Sixth AWS Tech. Exchange Conf., U. S. Naval Academy, AWS Tech Rep 242, 42-53.
2. Yang, C.-H. and R. Shapiro, 1973: The effects of the observational system and the method of interpolation on the computation of spectra, J. Atmos. Sci., 30, 530-536.

This is mainly due to the weights imposed during interpolation. Because of distance and distribution of observation points vis-a-vis the grid points, certain data will be unevenly weighted. One would therefore like to avoid interpolation, if at all possible, and derive the spectral coefficients directly from the unevenly spaced data. This work will examine some possible procedures for accomplishing this and the problems related to each of the methods.

II. ORTHOGONAL EXPANSIONS

Before discussing the possible methods for deriving spectral coefficients from randomly distributed data points, it may be useful to review some of the properties associated with expansions in orthogonal bases. A set of functions $\{\phi_n(x)\}$ is orthonormal over $x \in [a, b]$ if $\int_a^b \phi_k(x) \phi_j(x) dx = \delta_j^k$ where $\delta_j^k = 1$ when $j = k$ and $= 0$ otherwise.

There are certain properties associated with $\{\phi_n\}$ with respect to an arbitrary function $f(x)$ in $[a, b]$ that are important:

1. If $f(x)$ is bounded and continuous at all but a finite number of points, it may be expressed as an infinite series of the orthonormal functions, i.e., $f(x) = \sum_{i=0}^{\infty} a_i \phi_i(x)$ except at the points of discontinuity.

2. The coefficients $\{a_i\}$ can be derived by implementing the orthonormal property of $\{\phi_n\}$. For any k , $a_k = \int_a^b \phi_k f(x) dx$.

3. The coefficients $\{a_i\}$ satisfy a least-squares criterion for any truncation $0 < M < \infty$. Thus, for any M , $\int_a^b \left[f(x) - \sum_{i=0}^M a_i \phi_i \right]^2 dx$

will be a minimum.

When dealing with a set of discrete data, integrals may have to be approximated in order to employ orthogonal expansions. When the data are evenly spaced, approximations to integrals can become very close, depending on the orthonormal functions in question and the highest degree required. Fourier transforms, for example,

can be approximated with a great deal of precision over the interval $[-\pi, \pi]$ if sufficient points and their spacing are suitable for the set of functions $\{\cos nx\}$, $\{\sin nx\}$, or $\{e^{inx}\}$. If the orthogonal functions are polynomials, such as Legendre polynomials over the interval $[-1, 1]$, transformations of polynomials can be made exact by introducing quadrature formulae with sufficient points to resolve the combined highest degrees of the Legendre polynomial and the polynomial in question. When the data are not equally spaced, or not spaced with regard to the particular quadrature formula employed, the degree of departure will depend on the approximation used.

III. THREE METHODS

In deriving the coefficients for an orthogonal expansion when the data points are unevenly spaced, one must first specify the criteria that one seeks to satisfy before solving for the coefficients. One may, for instance, wish that the series satisfy a least-squares fit, or one may desire that the coefficients decrease in magnitude with increasing wavenumber. For this study, three methods were selected for comparison with a known function consisting of a finite cosine series of K terms with the coefficients equal to $(-1)^n (n+1)^{-2}$ where n is the wavenumber. K is a parameter that can be varied. The generating function, $f(x)$, is, thus, equal to $\sum_{j=0}^K (-1)^j (j+1)^{-2} \cos 2\pi jx$. There are N values of x , selected between 0 and 0.5, where N is, again, a parameter of the problem. The approximating series is assumed to be truncated at M terms. The three methods are:

1. Derivation of the coefficients by a direct least-squares approach. $M + 1$ independent linear equations are derived for the coefficients, a_j , by requiring that $\sum_{i=1}^N \left[f(x_i) - \sum_{j=0}^M a_j \cos 2\pi jx_i \right]^2$ be a minimum. The coefficient matrix is symmetric with general term

$$c_{kj} = \sum_{i=1}^N \cos 2\pi jx_i \cdot \cos 2\pi kx_i. \quad (1)$$

2. Derivation of empirical orthogonal functions (EOF's). One chooses a set of functions $\{P_n\}$ such that they satisfy the conditions of orthonormality over the data points, i.e.,

$$1/N \sum_{i=1}^N P_j(x_i) P_k(x_i) = \delta_j^k. \quad (2)$$

These functions must be recalculated for each new distribution of data points. The nature of the functions is also arbitrary. Ordinarily they are chosen as polynomials with coefficients derived by the Gram-Schmidt procedure. For purposes of comparison in this study, the functions were chosen as linear combinations of cosines. That is,

$$P_k(x) = \sum_{j=0}^k b_j \cos 2\pi j x. \quad (3)$$

The Gram-Schmidt method is just as applicable here as with polynomials. $f(x)$ can then be approximated by the series

$$\sum_{n=0}^M \beta_n P_n(x), \text{ where the coefficient } \beta_k \text{ is given by virtue of the}$$

orthonormality condition, namely

$$\beta_k = 1/N \sum_{i=1}^N f(x_i) P_k(x_i). \quad (4)$$

Combinations of the β_n 's and b_j 's should yield the equivalent of the a_i 's in Method 1 for cross comparison.

3. Deriving the coefficients one-by-one. Were the cosine functions truly orthogonal under a summation over the randomly-spaced points, the coefficients could be derived singly and would still satisfy the least-squares requirement. Because they are not orthogonal for the unequally-spaced points, however, coefficients derived singly will not satisfy the least-squares requirement for the entire series, but they can be restricted so that they will eventually decrease in magnitude with increasing wavenumber. This criterion dictates that low wavenumbers be made representative of the function's large scale, while high wavenumbers represent the unfiltered smaller waves. To achieve this, one can simply sequentially subtract the longer waves leaving only the higher wavenumbers. A

particular wavenumber is then filtered out by multiplication by the orthonormal function of the desired wavelength and adding over all points. Were the points correctly spaced or the orthonormal function truly orthonormal, the filter would work perfectly. Because of uneven spacing, however, the method captures more waves than intended. There is a guarantee, however, that the coefficients of higher order wavenumber will ultimately represent waves of decreasing length. This aspect has been termed the "finality" condition by Holmström (1963)³ who imposed it on the creation of his EOF's for analyzing geopotential heights. It, in effect, guarantees convergence of the series throughout the interval.

Accordingly, the first order coefficient, a_0 , should represent the mean of the function, $1/N \sum_{i=1}^N f(x_i)$. The second coefficient, a_1 , represents the departure from the mean filtered by $\cos 2\pi x_i$ and normalized by $\sum_{i=1}^N \cos^2 2\pi x_i$.

Thus $a_1 = \sum_{i=1}^N \left[f(x_i) - a_0 \right] \cos 2\pi x_i \left(\sum_{i=1}^N \cos^2 2\pi x_i \right)^{-1}$. This process

is continued so that any coefficient is given by

$$a_k = \sum_{i=1}^N \left[f(x_i) - \sum_{j=0}^{k-1} a_j \cos 2\pi j x_i \right] \cos 2\pi k x_i \left(\sum_{i=1}^N \cos^2 2\pi k x_i \right)^{-1}. \quad (5)$$

IV. TESTS

A series of tests was made with regards to distribution, number, and error. The tests can best be understood by a description of the label connected with any given run. A typical run was described by a

3. Holmström, I., 1963: On a method for parametric representation of the state of the atmosphere, Tellus, 15, 127-149.

label which looked like d , $N = n$, $M = m$, $K = k$, R . Capital letters are fixed, lower case letters are variable.

d - Distribution. Three types of distribution were tested. Evenly spaced data points, $d \equiv E$, implies that the N data points were placed in positions which would enable fast Fourier transforms, i.e., $x_j = (2j-1)(4N)^{-1}$, $j = 1, \dots, N$. Random placing of the data was performed uniformly, $d \equiv U$, or by staggering the data points, $d \equiv S$, so that 90 percent of the points lay in the interval $0 \leq x < 0.25$ and 10 percent in the remaining half.

N, M, K - Number. N refers to the number of data points (n), assumed known. M refers to the truncation limit (m) on the approximating series, while K refers to the upper limit (k) on the generating function. The default values for m and k are 15, when not specified on the label.

R - Random Error. When the label "R" appears, "observational" errors were added to the data before the coefficients were determined. The error, ϵ , was uniform random and of magnitude ≤ 0.01 (as compared with a maximum coefficient magnitude of 1.0).

Thus a label U , $N = 50$, R refers to a test where 50 data points were uniformly distributed between 0 and 0.5, 15 terms were used for both the generating and approximating series, and an error of at most ± 0.01 was added to each data point. A label S , $N = 50$, $M = 10$ implies a staggered distribution of the 50 points, 15 terms in the generating function and only 10 terms in the approximating series.

V. RESULTS

Although Methods 1 and 2 differ procedurally, they yield the same results. The reason for this is that the EOF's are linear combinations and of the same degree as the cosine series. Because of this, they both span the same vector space and will give the same approximation to any function. In general we can prove the following theorem:

Given two sets of functions $\{Q_n(x)\}$ and $\{\phi_n(x)\}$, integrable on $[a, b]$, where $Q_n(x) = \sum_{l=0}^N b_{ln} \phi_l(x)$, $b_{ln} \neq 0$, and also a set of

coefficients $\{\beta_n\}$ such that $\int_a^b \left[f(x) - \sum_{n=0}^M \beta_n Q_n(x) \right]^2 dx$ is a minimum

for a given integrable function $f(x)$, then the coefficients $\{\alpha_n\}$,

where $\sum_{n=0}^M \alpha_n \phi_n(x) = \sum_{n=0}^M \beta_n Q_n(x)$ are exactly those which will minimize $\int_a^b \left[f(x) - \sum_{n=0}^M \alpha_n \phi_n(x) \right]^2 dx$.

Proof: Since the coefficients β_j , $j = 0, \dots, M$, minimize the integral of the squared difference, we have by differentiation with respect to each β_k the condition $\int_a^b \left[f(x) - \sum_{j=0}^M \beta_j Q_j(x) \right] Q_k(x) dx = 0$,

$k = 0, \dots, M$. Substituting for Q_k and Q_j , we get

$$\int_a^b \left[f(x) - \sum_{j=0}^M \alpha_j \phi_j(x) \right] \sum_{\ell=0}^k b_{\ell k} \phi_\ell(x) dx = 0 \text{ for } k = 0, \dots, M. \text{ Since}$$

$b_{\ell k} \neq 0$, we can eliminate each term of the series sequentially by starting with $k = 0$ and working up to $k = M$. This leaves us with the

condition $\int_a^b \left[f(x) - \sum_{j=0}^M \alpha_j \phi_j(x) \right] \phi_k(x) dx = 0$. But this is simply the

necessary condition for proving that the coefficients α_j , $j = 0, \dots, M$, minimize the squared difference between the $f(x)$ and the series in terms of the ϕ_j 's.

Because the two methods differ computationally, however, they will yield different results when approximate solutions are either very close to the true solution or very far. Table 1 exemplifies this feature, in giving the root mean square (rms) error between coefficients of the approximating series and the true coefficients for equally spaced points. Without the addition of random error to the data, rms errors for all three methods are within the noise level of the computer algorithms. The addition of random error to the data affects all methods equally; rms errors for fewer points are only slightly greater than for $N = 30$. Truncation to $M = 10$ has no detrimental effect even in the face of random error.

TABLE 1. RMS Differences between Coefficients of the Generating Series and the Coefficients of the Approximating Series for Uniform Distribution.

<u>Run</u>	<u>Method 1</u>	<u>Method 2</u>	<u>Method 3</u>
E,N=30	2.1543×10^{-14}	2.5977×10^{-14}	1.9510×10^{-14}
E,N=16	2.0282×10^{-14}	2.5152×10^{-14}	1.4059×10^{-14}
E,N=30,M=10	2.5182×10^{-14}	2.6447×10^{-14}	2.3845×10^{-14}
E,N=16,M=10	1.7838×10^{-14}	2.1707×10^{-14}	1.6992×10^{-14}
E,N=16,M=8	1.8979×10^{-14}	2.4168×10^{-14}	1.8973×10^{-14}
E,N=30,R	1.1936×10^{-3}	1.1936×10^{-3}	1.1936×10^{-3}
E,N=30,M=10,R	1.2822×10^{-3}	1.2822×10^{-3}	1.2822×10^{-3}

Figures 1a-b show the effect of number on U and S distributions, respectively, for all three methods. Methods 1 and 2 do not coincide for these cases because they are still within the noise range of the algorithm, even for S distribution where the rms error rises rapidly with decreasing N. It is obvious from these figures that while Methods 1 and 2 are better than Method 3 for both distributions, at least for $N = 30$, they become highly unreliable for decreasing N. Method 3 is barely affected by number and maintains approximately the same estimates of coefficients even when N approaches the critical minimum of 15 points for resolving the 15 cosine waves. Note also that with U distribution, increasing the number of points beyond 30 has very little impact on error for any of the methods. For S distributions, however, there is a distinct effect caused by varying the number of points.

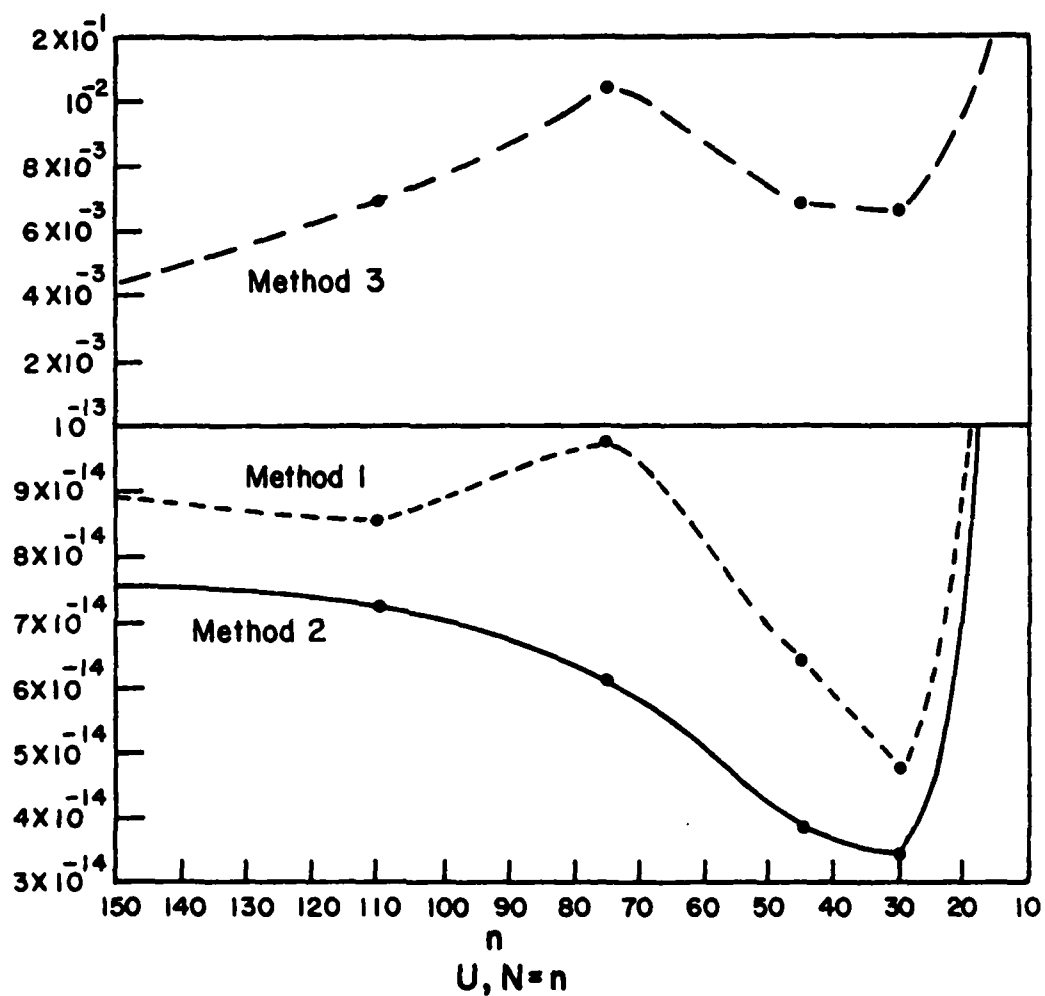


Figure 1a. RMS difference between coefficients of the approximating series and the generating series as a function of number of data points for all three methods and for U distribution.

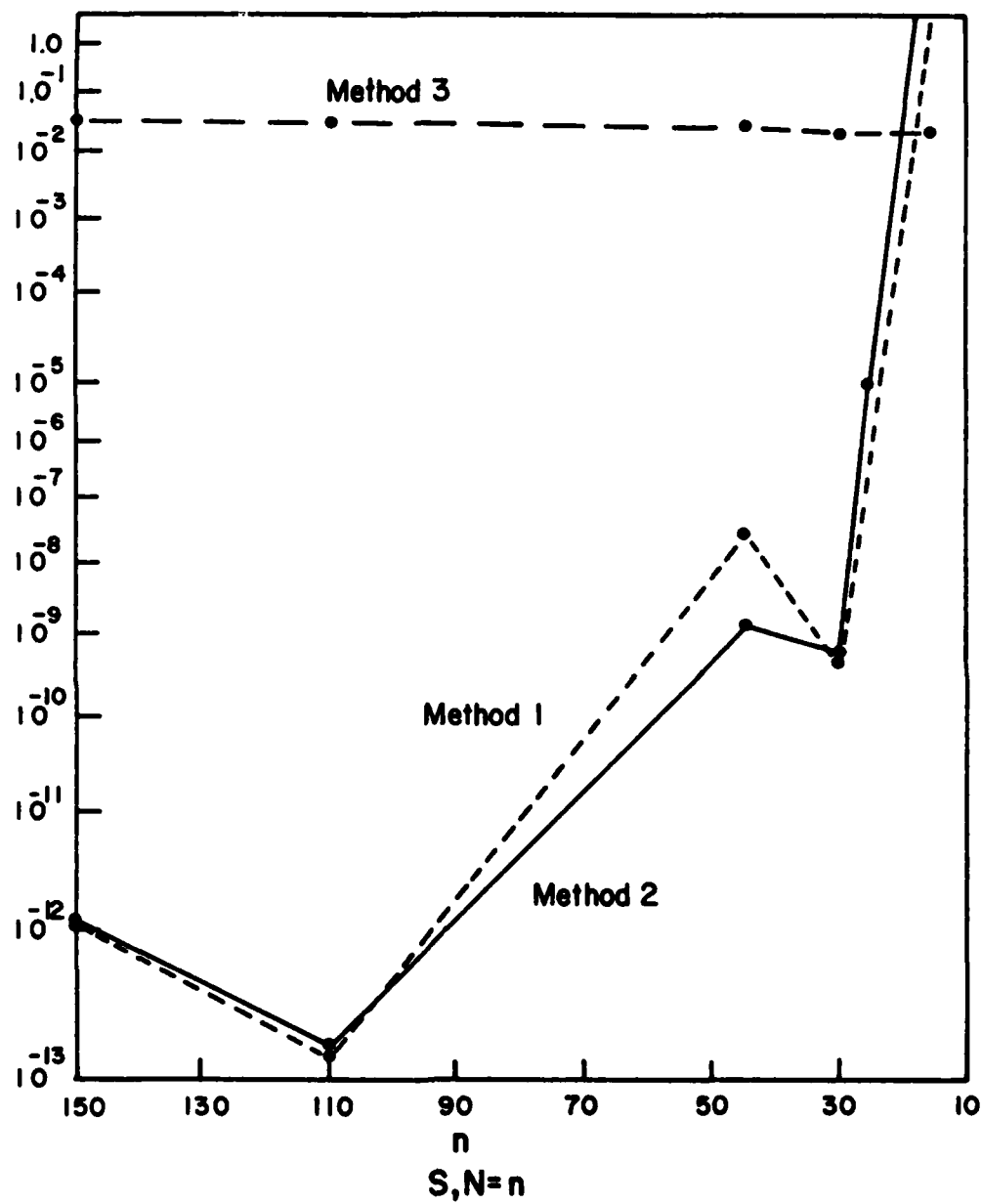


Figure 1b. RMS difference between coefficients of the approximating series and the generating series as a function of number of data points for all three methods and for S distribution.

Figures 2a-b depict the effects of truncation ($M = 10$) on the rms error. Here, error levels are substantially higher than for $M = 15$. Oddly, the rms error for all methods increases drastically for $N \rightarrow 15$ for U distribution, while for S distribution, only Methods 1 and 2 cause a precipitous rise in rms error as N drops.

Table 2 lists the rms errors for runs including the addition of random error to the data. Method 3 fares much better here than in previous comparisons. In general, Methods 1 and 2 seem superior in cases when the ratio of N/M is high, while Method 3 prevails when the ratio is low, or when the S distribution is encountered. In fact, for S distribution and $N = 16$, Methods 1 and 2 fail to achieve credible approximations of the coefficients.

TABLE 2. RMS Differences between Coefficients of the Generating Series and the Approximating Series for Runs Involving the Addition of Random Error to the Data (Other than E Distribution Portrayed in Table 1).

<u>Run</u>	<u>Method 1</u>	<u>Method 2</u>	<u>Method 3</u>
U,N=45,R	1.9280×10^{-3}	1.9280×10^{-3}	6.9638×10^{-3}
U,N=16,R	5.4372	5.4372	2.0119×10^{-2}
S,N=45,R	7.1964×10^{-1}	7.1964×10^{-1}	4.5070×10^{-2}
S,N=16,R	Failed	Failed	3.8026×10^{-2}
U,N=30,M=10,R	3.1441×10^{-3}	3.1441×10^{-3}	8.7045×10^{-3}
S,N=30,M=10,R	8.0520×10^{-3}	8.0520×10^{-3}	4.2164×10^{-2}
S,N=16,K=20,R	Failed	Failed	3.8103×10^{-2}

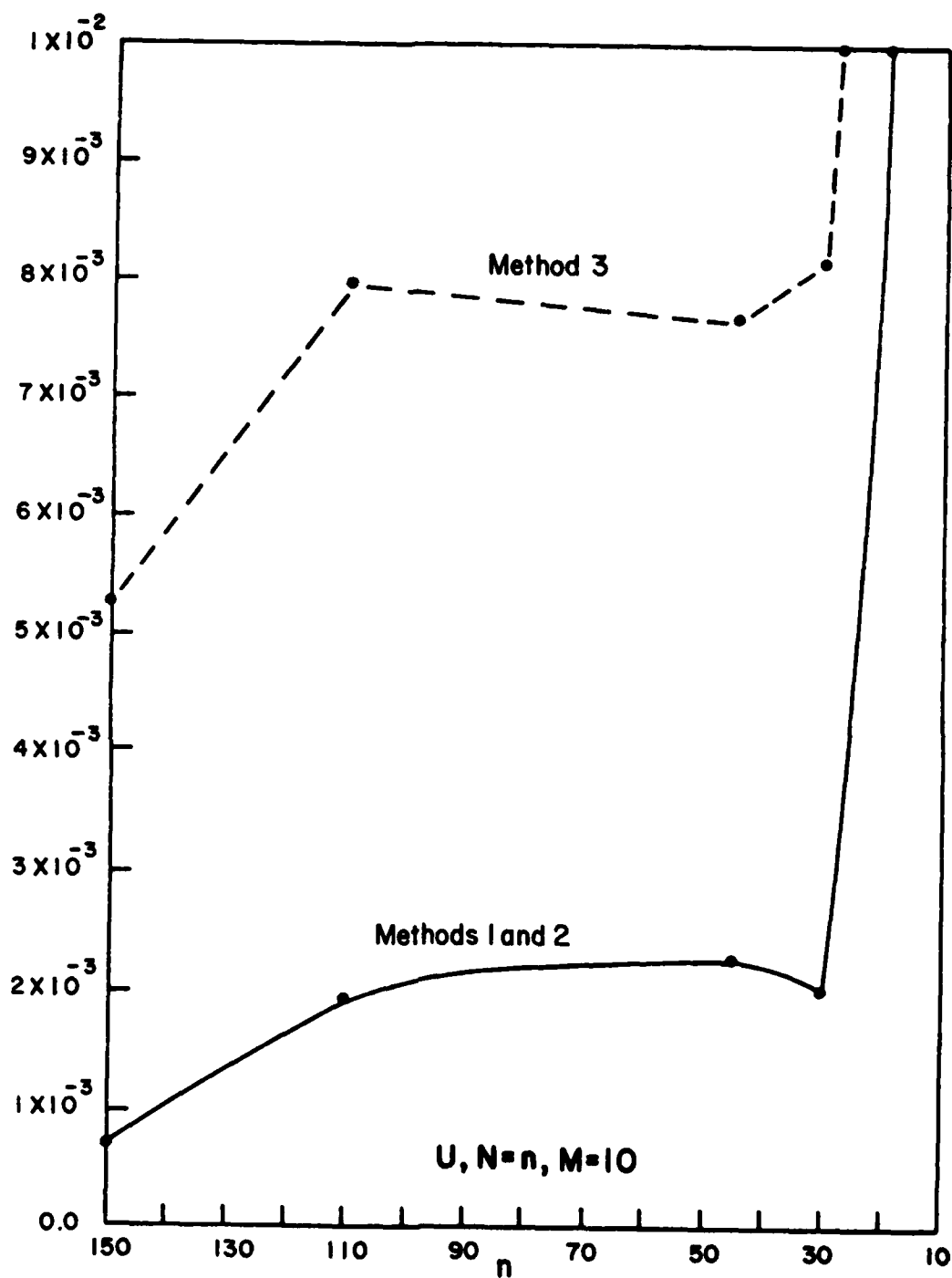


Figure 2a. RMS difference between coefficients of the approximating series and the generating series as a function of number of data points for all three methods and for U distribution. M is truncated at wavenumber 10.

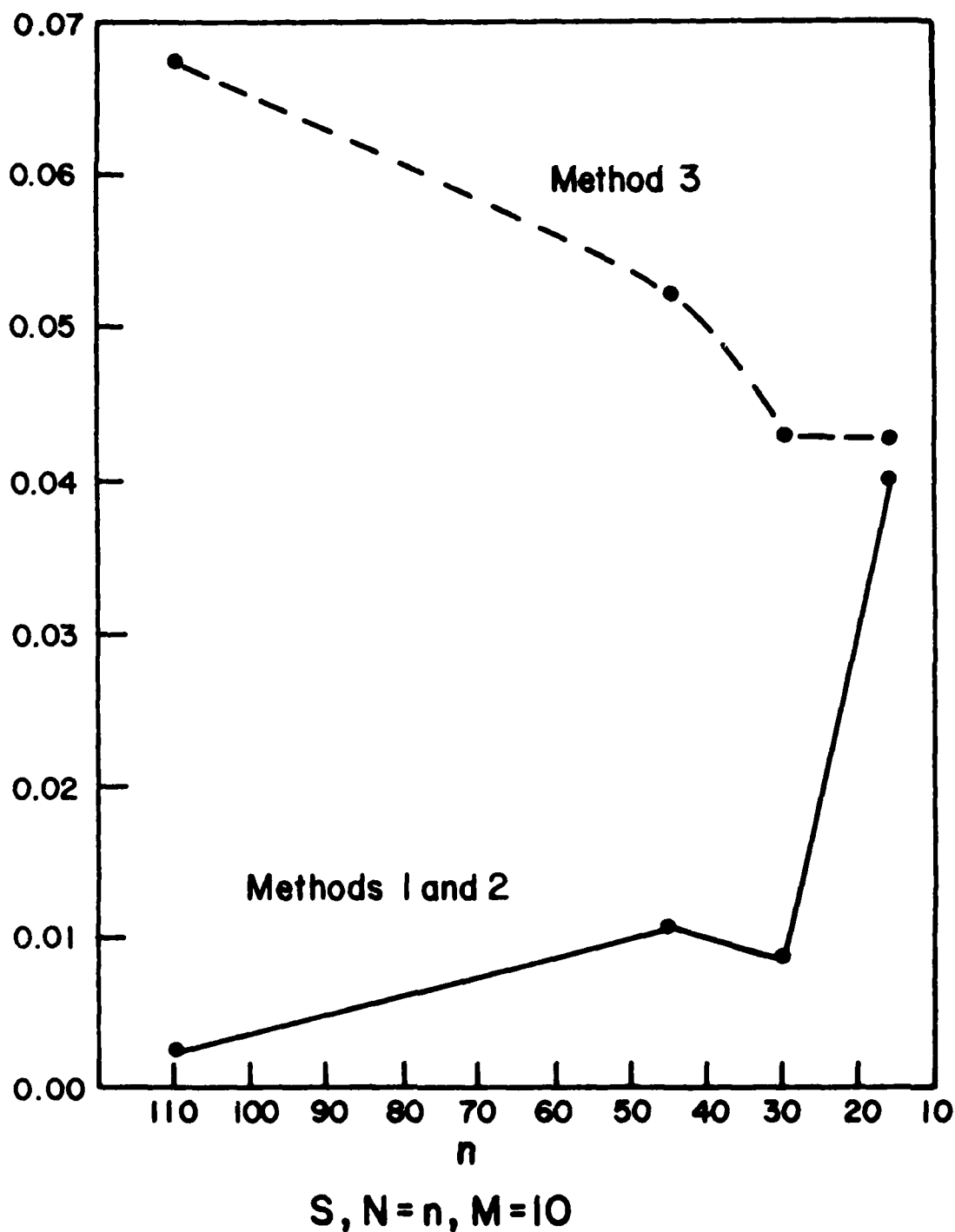


Figure 2b. RMS difference between coefficients of the approximating series and the generating series as a function of number of data points for all three methods and for S distribution. M is truncated at wavenumber 10.

The spectra shown in Figures 3 - 5 demonstrate the effect, by wavenumber, of the various errors. Figure 3 depicts the spectra of the generating functions along with the approximating function for E , $N = 30$, $M = 10$, R . By Table 1, all methods coincide for this case and so only one approximating spectrum is depicted. Notice that the effect of truncation and the addition of random error seem only to perturb the high wavenumbers ($k \geq 7$). Figures 4a-c show the effects of random distribution in comparison with Figure 3. In Figure 4a, differences in the methods appear for lower wavenumbers, but, despite the presence of truncation at $M = 10$, departures from true values are not systematic or severe. The addition of random error to the data does not result in any major spectral changes as seen in Figure 4b. The reduction of data points, however, as depicted by Figure 4c, has a pronounced effect on the spectra, resulting in a substantial underestimate of values for wavenumbers higher than 5. Figures 5a-c are similar to Figures 4a-c, except that they are for S distribution. Here the departures are more severe than for U distribution, while, again, the addition of random error has very little impact. The reduction of data points, however, does not result in a systematic underestimate of the spectra as with U distribution. If anything, there seems to be an overestimation of the spectral coefficients in this case. It is also important to notice that the departure of Method 3 is more pronounced for higher wavenumbers than Methods 1 and 2, but that they are both as accurate for lower wavenumbers.

Depiction of the function itself is found in Figures 6 and 7. Figures 6a-b are for Methods 1 and 2 and Method 3, respectively, for $N = 30$, $M = 10$. Both U and S distributions are shown on each graph. Here, one can clearly see the superiority of Methods 1 and 2 over Method 3 as far as accuracy is concerned. When N drops to 16 as in Figures 7a-b, this superiority, especially with regards to the S distribution, is non-existent. Notice that Methods 1 and 2 still guarantee a least-squares fit to the data. Even for S , $N = 16$, $M = 10$ (Figure 7a) where very large discrepancies appear, the approx-

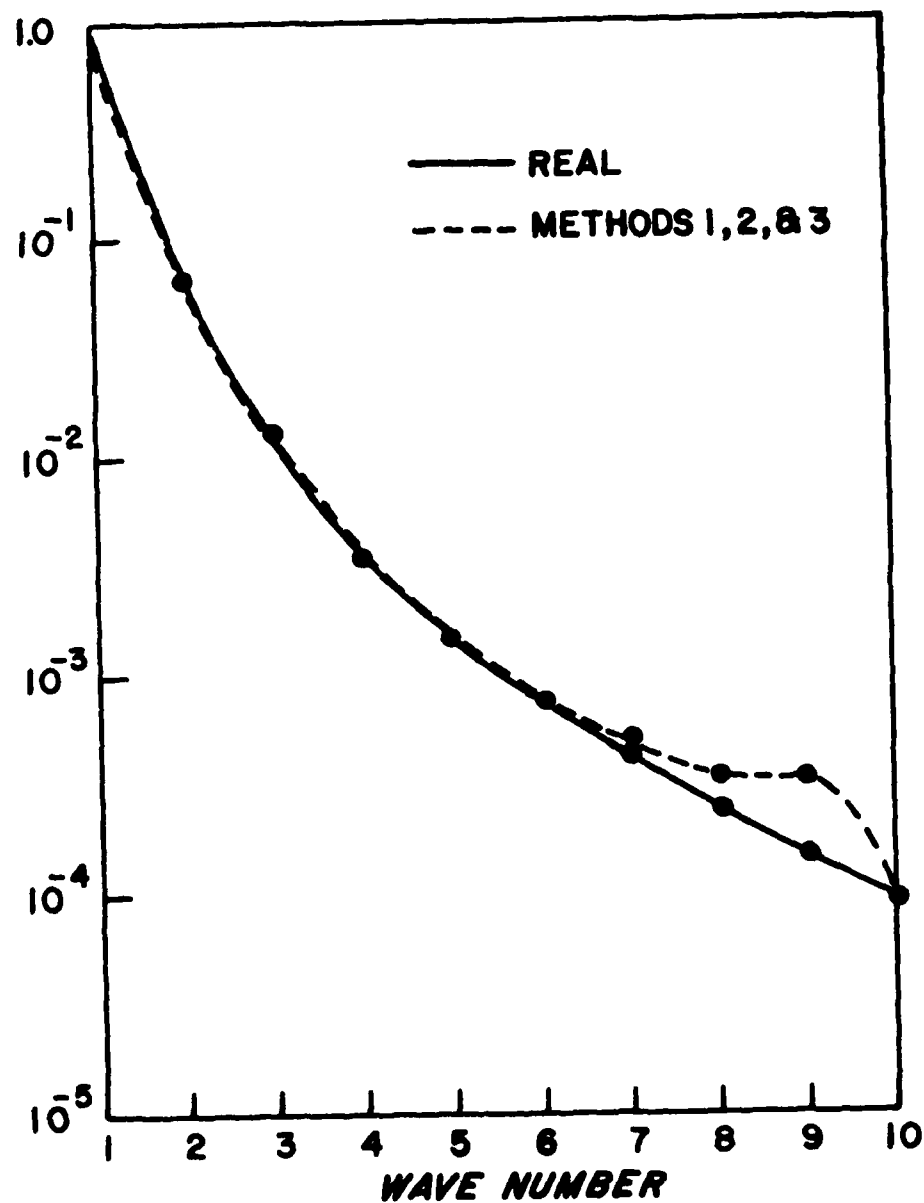


Figure 3. Spectra of the approximating series for all three methods and for the generating series for E, $N = 30$, $M = 10$, R .

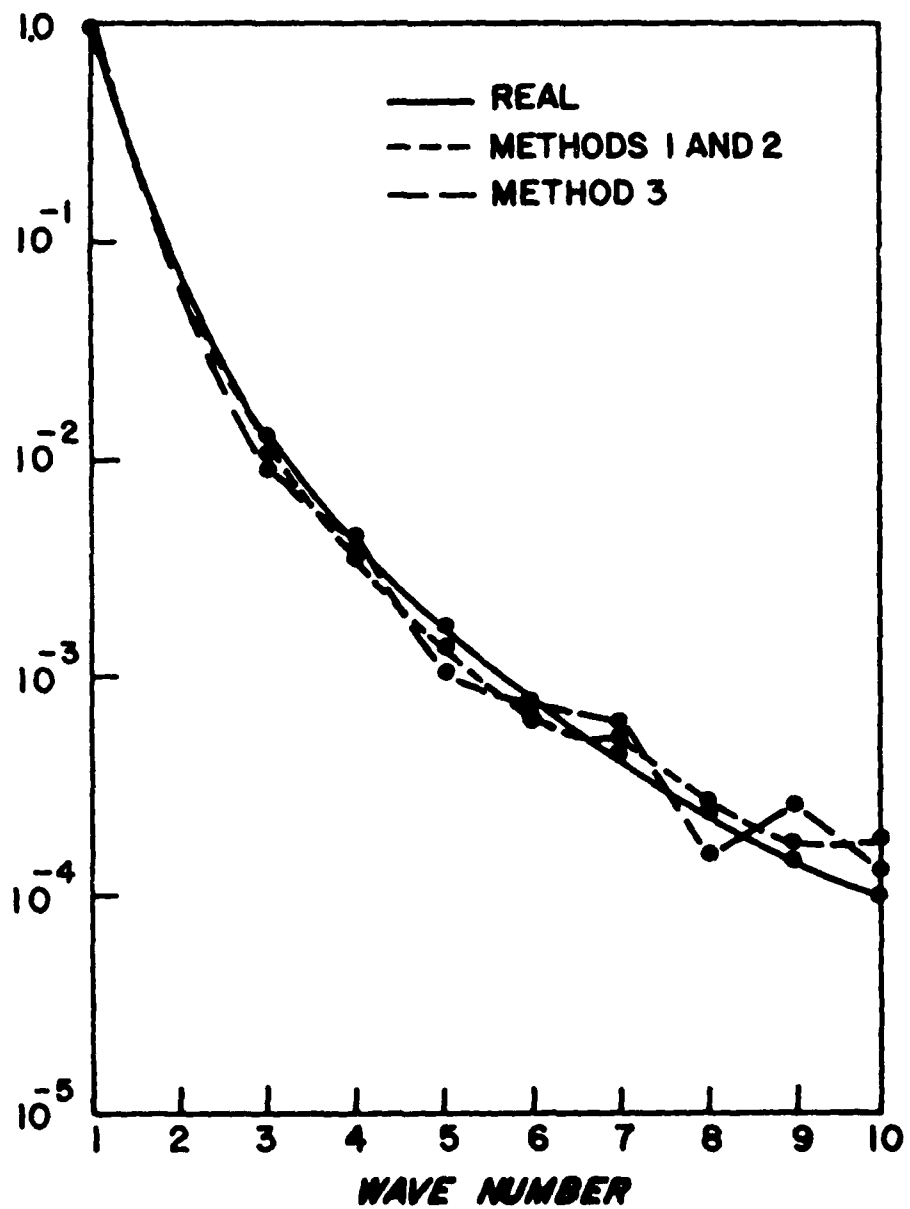


Figure 4a. Spectra for the generating series, Methods 1 and 2, and Method 3 for $U, N = 30, M = 10$.

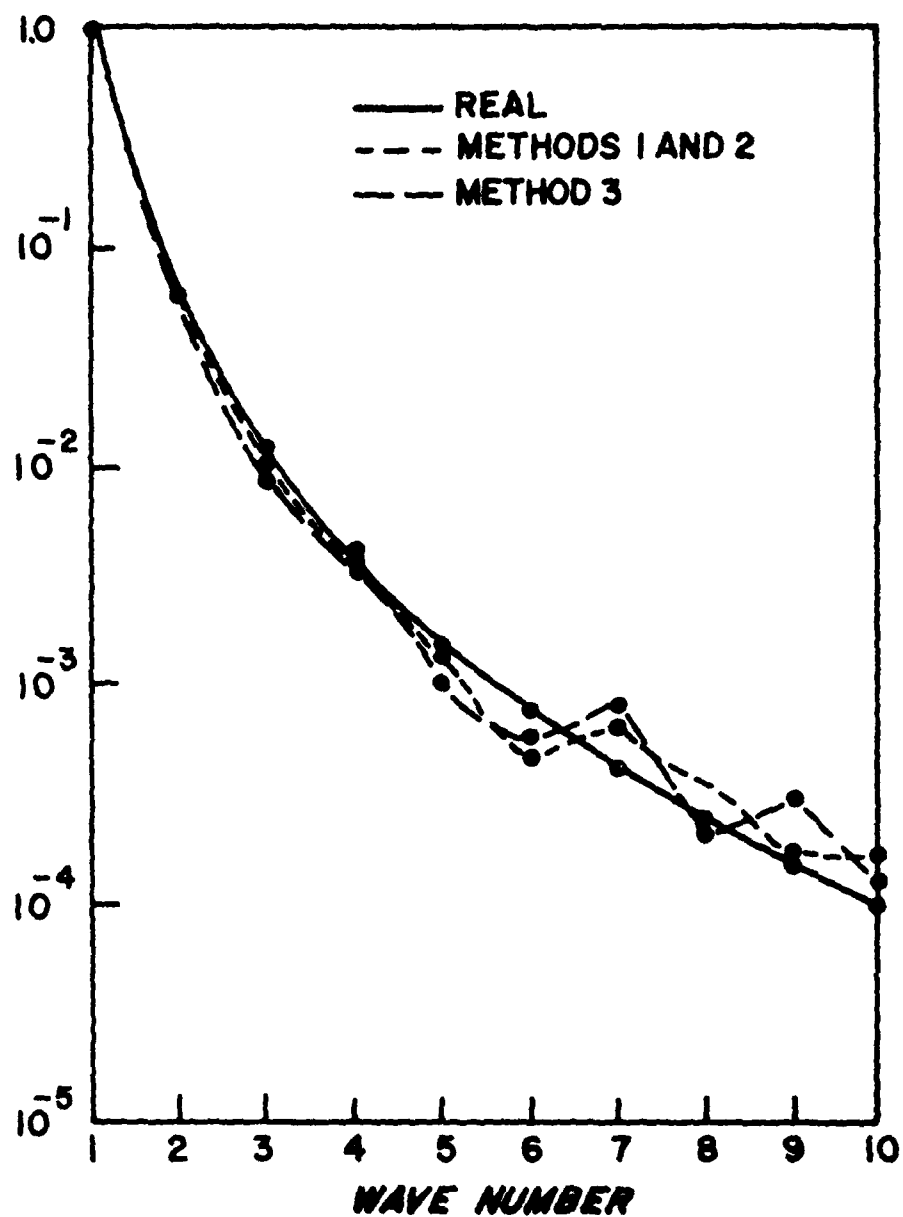


Figure 4b. Spectra for the generating series, Methods 1 and 2, and Method 3 for $U, N = 30, M = 10, R$.

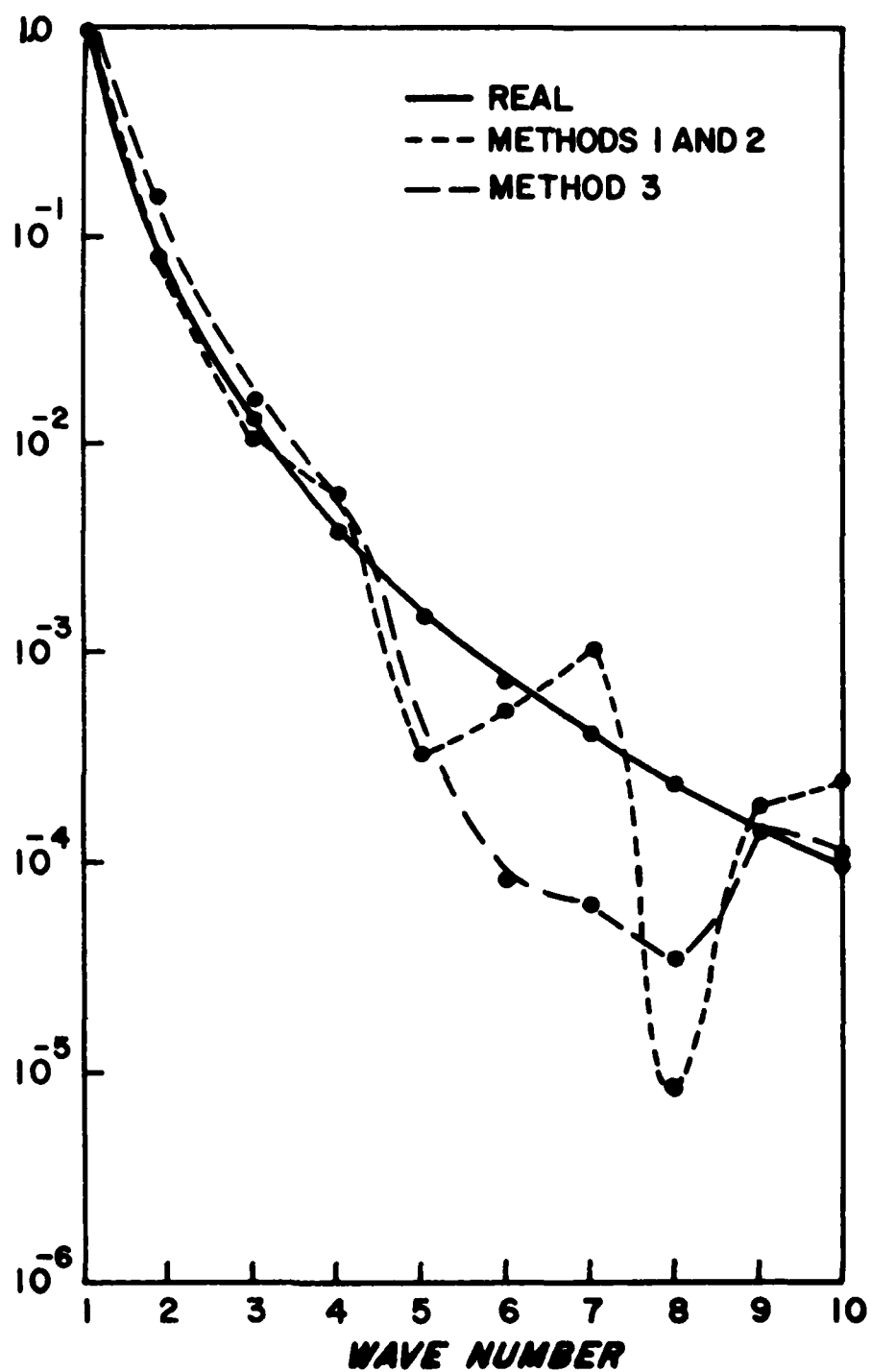


Figure 4c. Spectra for the generating series, Methods 1 and 2, and Method 3 for $U, N = 16, M = 10$.

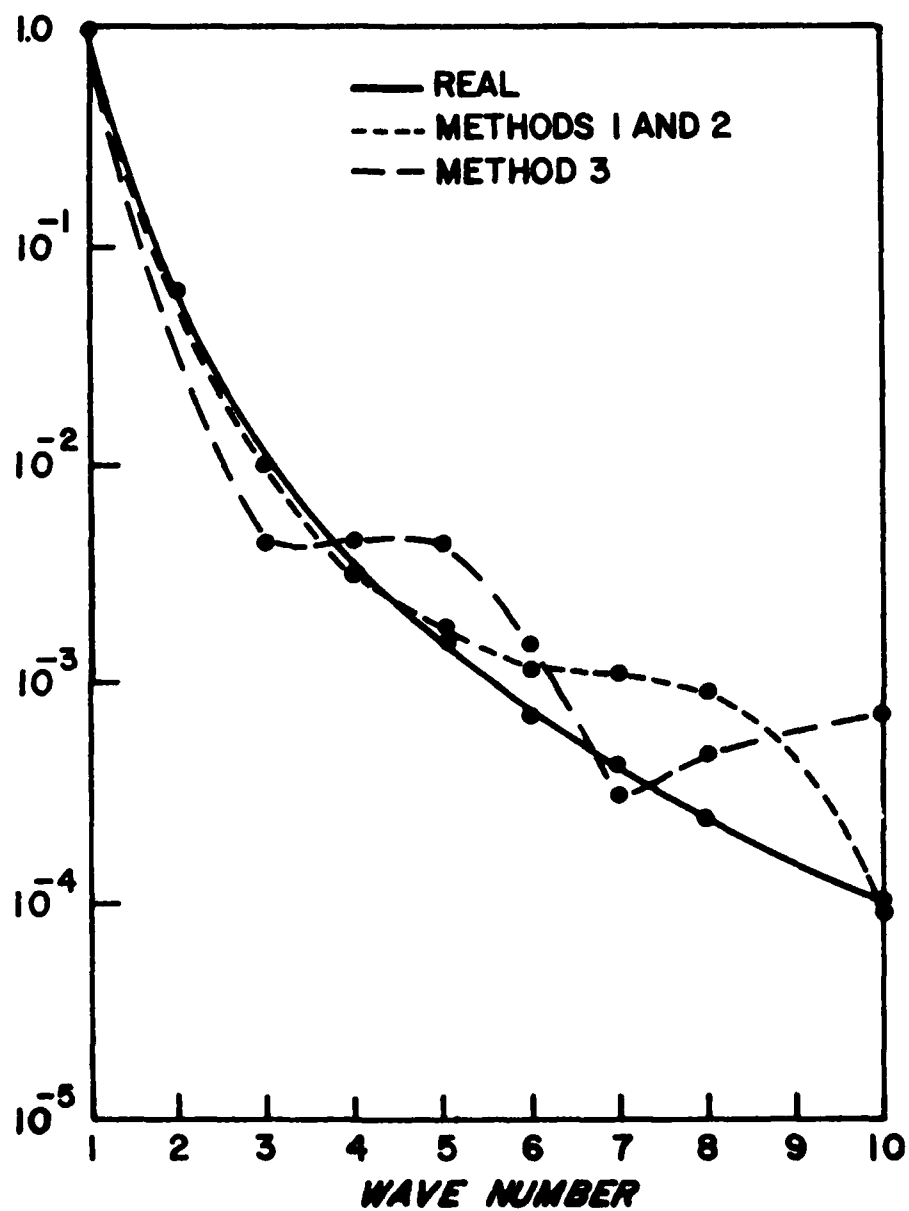


Figure 5a. Spectra for the generating series, Methods 1 and 2, and Method 3 for $S, N = 30, M = 10$.

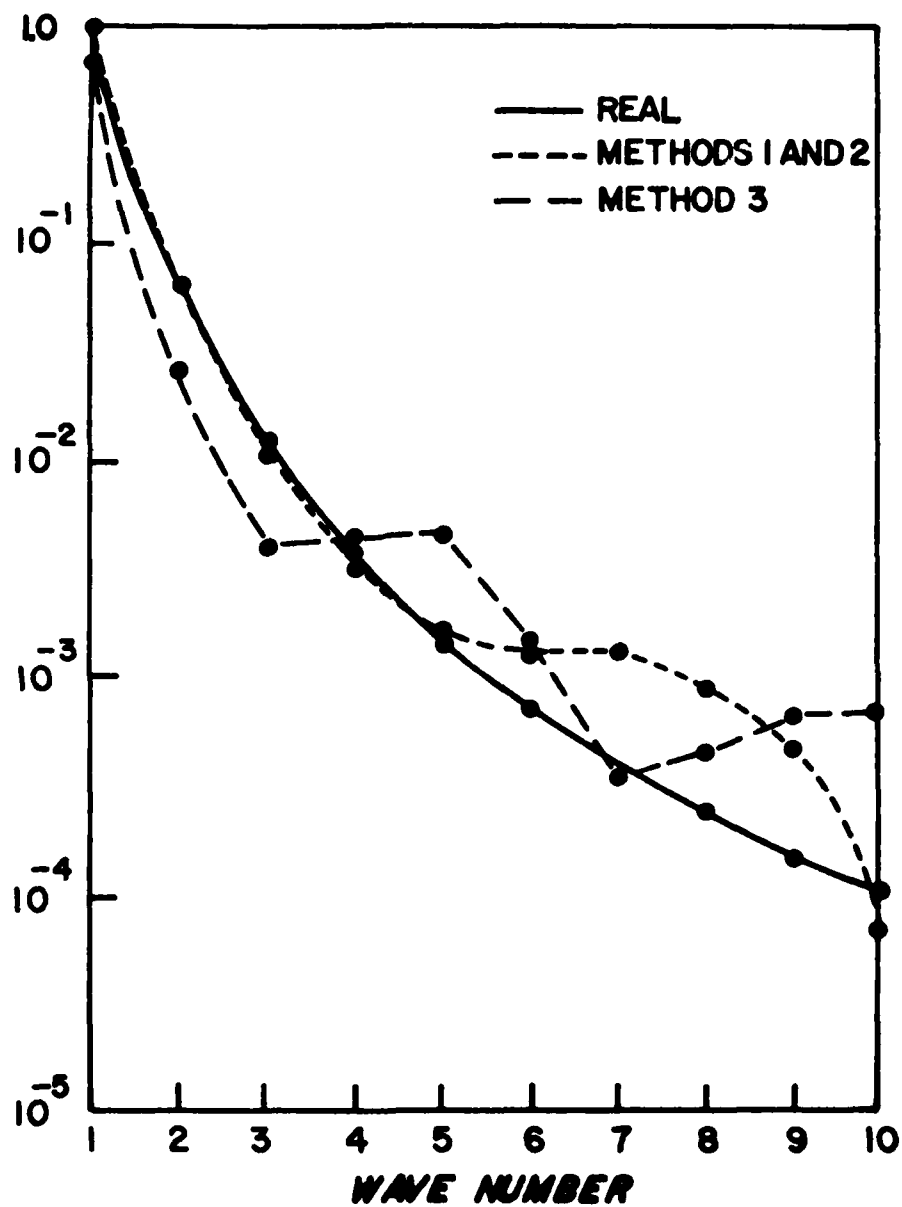


Figure 5b. Spectra for the generating series, Methods 1 and 2, and Method 3 for $S, N = 30, M = 10, R$.

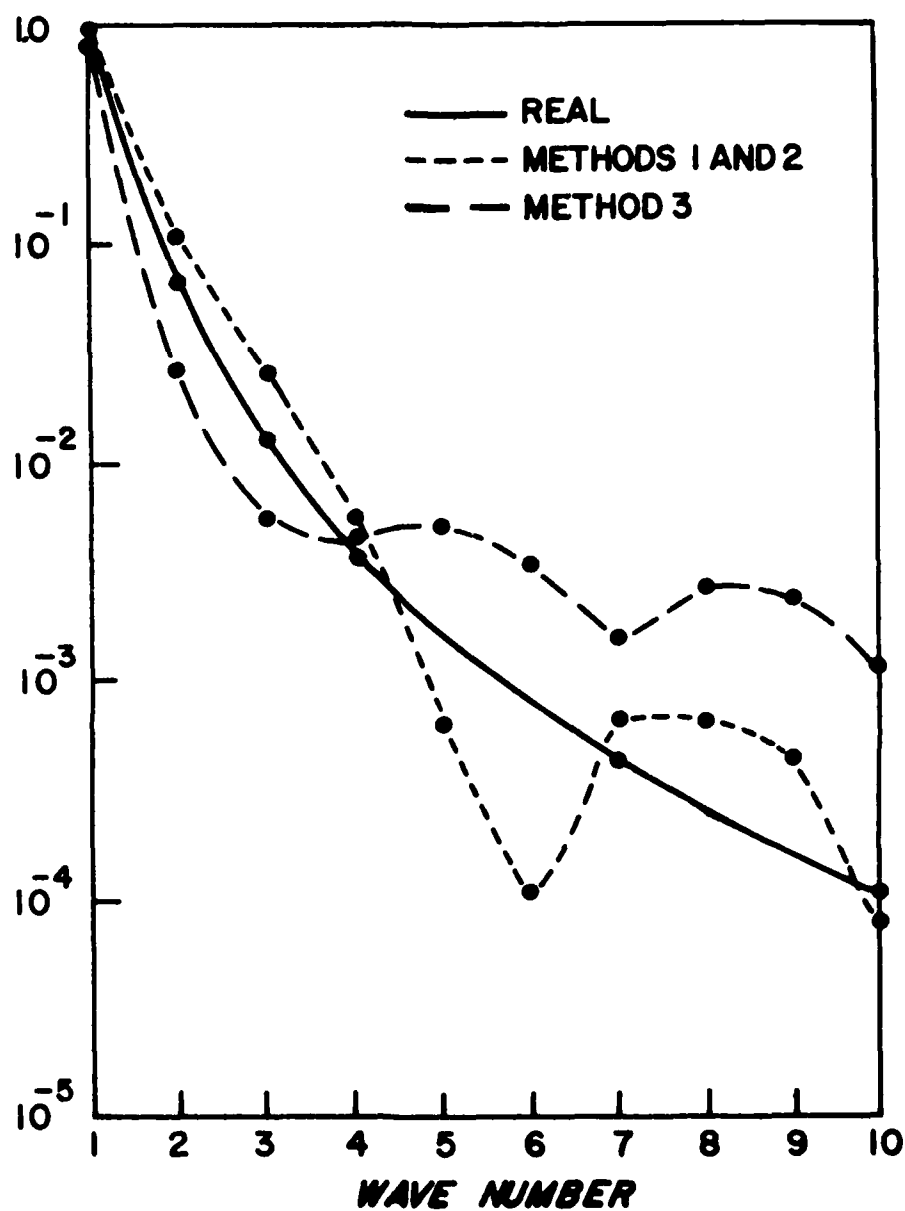


Figure 5c. Spectra for the generating series, Methods 1 and 2, and Method 3 for $S, N = 16, M = 10$.

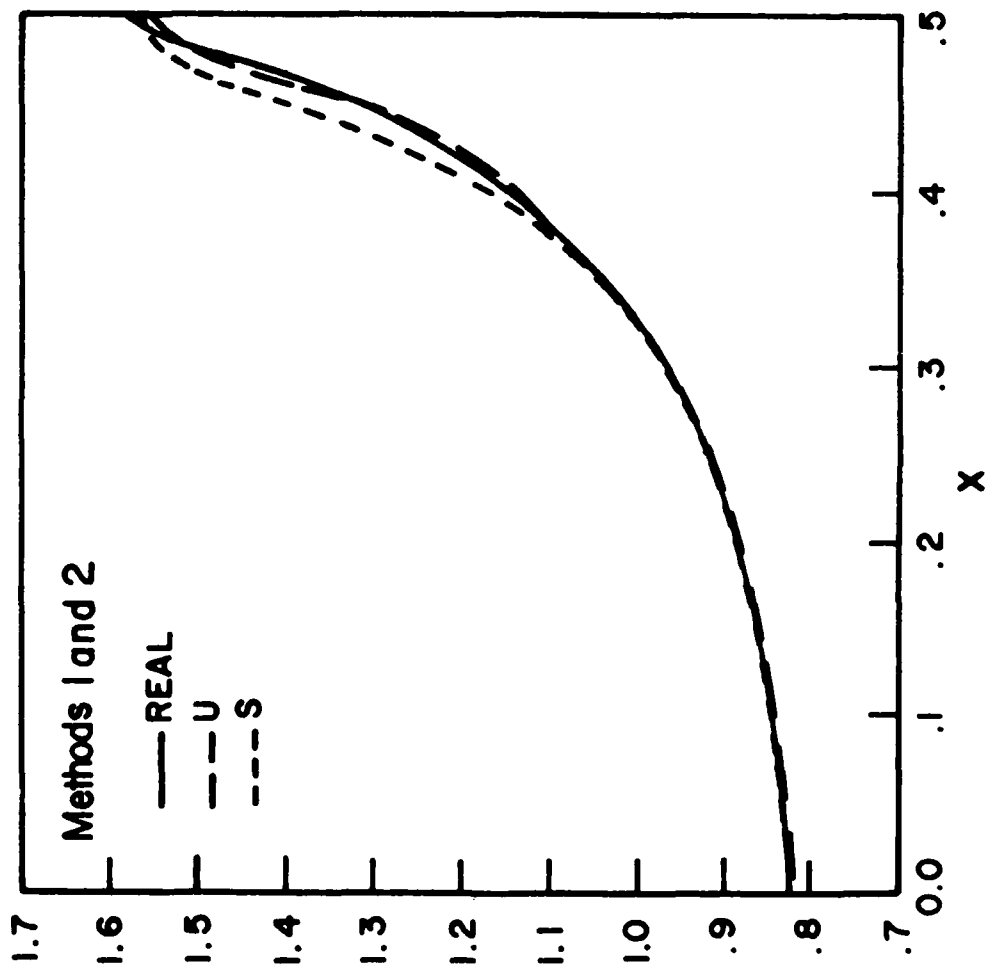


Figure 6a. Values of the generating series and the approximating series for both U and S distribution, $N = 30$, $M = 10$, as a function of x for Methods 1 and 2.

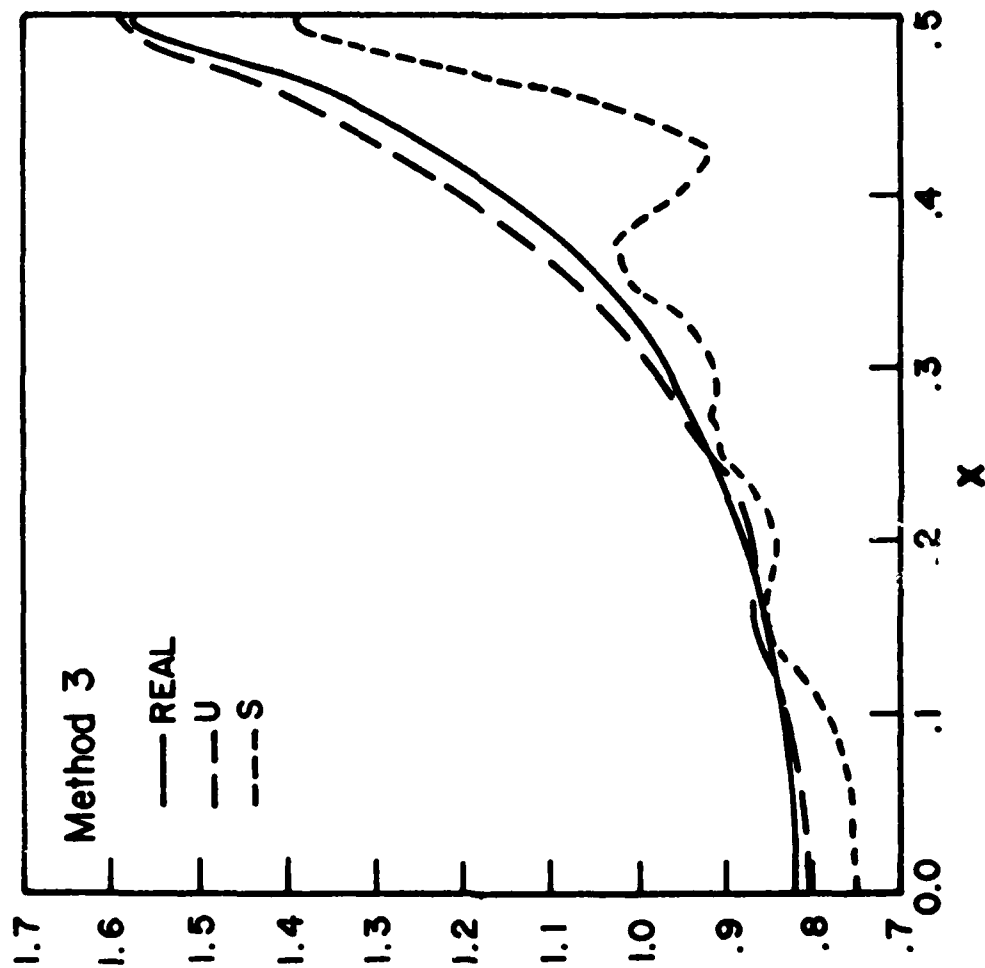


Figure 6b. Values of the generating series and the approximating series for both U and S distribution, $N = 30$, $M = 10$, as a function of x for Method 3.

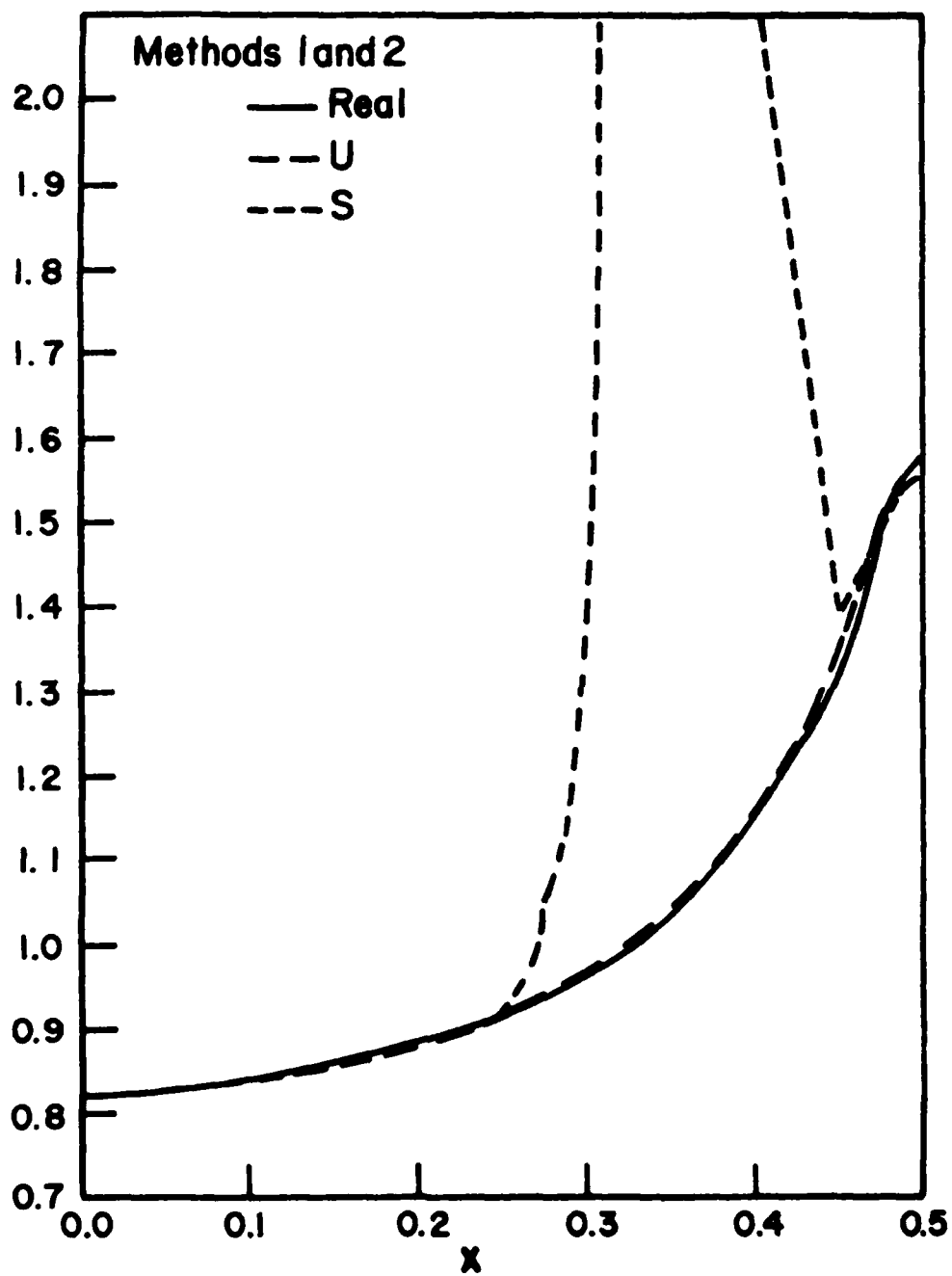


Figure 7a. Values of the generating series and the approximating series for both U and S distribution, $N = 16$, $M = 10$, as a function of x for Methods 1 and 2.

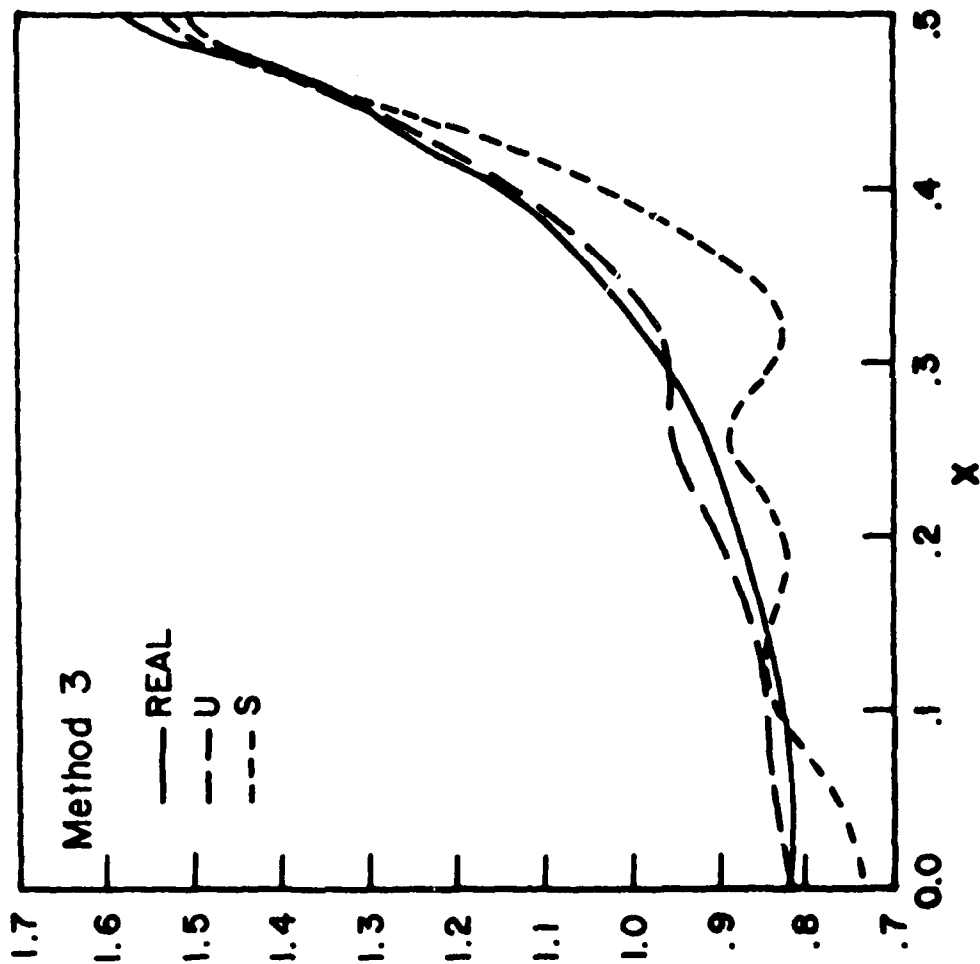


Figure 7b. Values of the generating series and the approximating series for both U and S distribution, $N = 16$, $M = 10$, as a function of x for Method 3.

imating function still minimizes the squared error with the given data. Apparently, however, the distribution of the data points left out the interval between, roughly, 0.25 and 0.45, allowing the function to roam freely between these values. Method 3 may not fit the data very well in places but, because of its finality condition, prevents runaway errors.

VI. SUMMARY AND CONCLUSION

Three methods have been studied which yield coefficients of a finite cosine series in an interval, given a sampling of data in that interval. The first method directly solved the least-squares condition by matrix inversion. The second involved the calculation of EOF's, while the third imposed the condition of finality and the coefficients were derived singly. It was shown that since the EOF's were specified as linear combinations of cosines, the first two methods would yield the same results.

The following conclusions can be drawn from examining the rms errors and the spectra of the various approximating series produced by the methods:

1. Number is less of a factor than distribution. Even when the number of points approaches the critical value, if they are equally spaced, they will result in perfect approximation. Uniformly distributed data will generally result in closer approximation than non-uniformly distributed data.

2. Method 3 does not provide as close a fit as Methods 1 and 2 but does not result in divergent behavior for poorly distributed data. The divergent behavior of Methods 1 and 2 is due to the fit of the series with the data given, which allows significant degrees of freedom for the intervals not covered by data. This is also related to the "ill-conditioned" nature of the coefficient matrix which must be inverted in Method 1, as will be discussed below.

3. The addition of random error to the data has very little impact on the results. At least for the magnitude of error tested,

the deviations from unperturbed cases are minimal. Since the errors are uniformly random, the averaging process over the data points acts as a filter and diminishes the effect of these observation errors.

4. Truncation error has a nominal effect on the coefficients especially in the high wavenumber regime. In many cases, when data points are few, it is preferable to truncate rather than to go to higher wavenumbers. This is because solution of Methods 1 and 2 makes the wavenumbers interdependent, so that the truncation is "felt" in all coefficients. In Method 3, however, there is no danger of this happening and the truncation limit is arbitrary.

From analysis of the spectra in Figures 4 and 5, one notes that Method 3, although designed to depict less of the function for increasing wavenumber, does not result in monotonically decreasing coefficients with increasing wavenumber. The reason for this, as stated above, is that, because of non-uniform distribution, the approximation of individual coefficients is not precise. Thus when subtraction of these waves from the original data is performed, the remainder will not be accurate, leading to more impreciseness. The method assures, however, that the approximation will eventually catch up and adjust itself to the data, without diverging.

One may question, however, whether numerical integration is perhaps superior to Method 3 for approximating the coefficients. That is, one may attempt to derive a_k by numerically integrating

$$2 \int_0^{.5} f(x) \cos 2\pi kx \, dx \text{ by quadrature formulae. Unless the data points}$$

are correctly spaced for the relevant quadrature formula, numerical integration will involve unwittingly weighting certain data points. This has the same effect as interpolation of data points to fixed grid points before integration. As pointed out previously, such a procedure could affect the spectral nature of the data. As an example, Table 3 indicates the rms differences between the coefficients of the generating function and coefficients derived from the trapezoidal and midpoint rule of integration. These results may be compared with those found in Tables 1 and 2 and with Figures 1 and 2. The trapezoi-

dal rule, being a closed Newton-Cotes quadrature formula, required inserting the end values of the generating function for each run. Note that for E distribution the midpoint quadrature is exact, while the trapezoidal rule is approximate. In general, rms differences for numerical integration are significantly higher than for Method 3 (in some cases by an order of magnitude). The sensitivity to number and distribution of data points is also obvious, when one notes the significant increase in error for S distribution as opposed to U distribution and the jump in error caused by a reduction in number. Method 3, which works on wave space rather than physical space, is better suited to represent the spectral nature of the particular function for these distributions, while numerical integration by concentrating on physical space is not as efficient.

TABLE 3. RMS Differences between Coefficients of the Generating Series and the Approximating Series for Runs Involving Trapezoidal and Mid-Point Quadrature Integrations.

<u>Run</u>	<u>Trapezoidal</u>	<u>Mid-point</u>
E, N=30	3.39793×10^{-2}	2.152034×10^{-14}
E, N=30, M=10	3.27108×10^{-2}	2.551610×10^{-14}
U, N=30	8.3108×10^{-2}	7.64284×10^{-2}
U, N=100, M=10	1.13156×10^{-2}	3.97745×10^{-3}
U, N=30, M=10	5.55853×10^{-2}	3.44025×10^{-2}
U, N=30, M=10, R	5.57984×10^{-2}	3.44947×10^{-2}
U, N=16, M=10	1.25008×10^{-1}	1.07500×10^{-1}
S, N=100, M=10	2.96793×10^{-1}	2.96183×10^{-1}
S, N=30, M=10	3.92814×10^{-1}	4.07799×10^{-1}
S, N=16	4.91530×10^{-1}	5.23083×10^{-1}

Methods 1 and 2 are apparently not based on physical spacing either but on finding a function which approaches most closely the available data. As such, the spacing affects the choice of coefficients by prescribing where along the x-axis the approximating function must approach the true solution. But the function is free to roam through other values of x where data are not available. The matrix of Method 1 can sometimes serve as an indicator as to when Methods 1 and 2 will fail. By examining the ratio of the eigenvalues of the coefficient matrix in the least-squares problem, one can ascertain the amount of distortion involved in the mapping of the solution vector to the given vector containing the data points. If the ratio of eigenvalues is very large, it is an indication that a great deal of magnification and diminution will be present in the mapping, leading to misrepresentation. Unfortunately, this is not a necessary condition for the solution to diverge from the true solution. In the problem considered here, the matrix is real symmetric with entries c_{kj} given by (1). The eigenvalues are therefore always real. The ideal ratio of extreme eigenvalues γ (i.e., for E distribution) is 2, because the elements of the matrix are simply $c_{kj} = \delta_k^j N/2$, except when $k = j = 0$, where $c_{00} = N$. For S, $N = 16$, $\gamma = -2.126 \times 10^{14}$, a clear indication as to why the method failed for that run. Yet for U, $N = 30$, where the fit was very close (according to rms differences in the coefficients), $\gamma = 1.681 \times 10^2$, but for U, $N = 100$, $M = 10$, where the fit was close, but not as close as U, $N = 30$, $\gamma = 3.343$. For U, $N = 30$, $M = 10$, R, where the fit is only slightly worse than U, $N = 100$, $M = 10$, $\gamma = 5.442$. S, $N = 30$, $M = 10$, R yields results which are only slightly worse than for the U distribution, but its $\gamma = 9.463 \times 10^2$. Thus, other criteria must be sought to enable evaluation of Methods 1 and 2 and their applicability.

In meteorological analysis one often encounters areas of sparse data similar to the S distribution in this study. Results here indicate that attempts to analyze data according to Methods 1 and 2 may produce runaway errors in regions of no coverage. Method 3, on the other hand, while being a safe method, could lead to serious inaccuracies. It may be best to combine methods by allowing Method 3 to provide a first guess in sparse areas which can then be used in

conjunction with the given data in data-dense areas to provide analyses according to Methods 1 and 2. The fact that Methods 1 and 2 will both provide similar coefficients will allow a pragmatic evaluation of the employability of each method for any given problem. For numerous coefficients the inversion of large matrices may be alleviated by employing Method 2. In other cases a simple matrix inversion may prove easier than the two-step EOF method.

This study was limited to one dimension. The introduction of a second dimension complicates the application of the three methods. If points are randomly distributed in two-dimensional space, it is not obvious how one can represent a sum of functions to approximate the data. With evenly distributed data, each dimension is considered separately, but this is not feasible if points are not aligned in any specific direction. Future research will attempt to solve these fundamental problems and extend the results obtained here to higher dimensions.