

ESD-TR-69-184

ESD ACCESSION LIST

ESTI Call No. 66628

Copy No. 1 of 2 cys.

W-7550

ESD RECORD COPY

RETURN TO
SCIENTIFIC & TECHNICAL INFORMATION DIVISION
(ESTI), BUILDING 1211

NONLINEAR PREDICTION AND DIGITAL SIMULATION

J. F. A. Ormsby

JUNE 1969

Prepared for
ELECTRONIC SYSTEMS DIVISION
AIR FORCE SYSTEMS COMMAND
UNITED STATES AIR FORCE
L. G. Hanscom Field, Bedford, Massachusetts



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

Project 6151
Prepared by
THE MITRE CORPORATION
Bedford, Massachusetts
Contract AF19(628)-2390

AD0693549

When U.S. Government drawings, specifications, or other data are used for any purpose other than a definitely related government procurement operation, the government thereby incurs no responsibility nor any obligation whatsoever; and the fact that the government may have formulated, furnished, or in any way supplied the said drawings, specifications, or other data is not to be regarded by implication or otherwise, as in any manner licensing the holder or any other person or corporation, or conveying any rights or permission to manufacture, use, or sell any patented invention that may in any way be related thereto.

Do not return this copy. Retain or destroy.

NONLINEAR PREDICTION AND DIGITAL SIMULATION

J. F. A. Ormsby

JUNE 1969

Prepared for
ELECTRONIC SYSTEMS DIVISION
AIR FORCE SYSTEMS COMMAND
UNITED STATES AIR FORCE
L. G. Hanscom Field, Bedford, Massachusetts



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

Project 6151
Prepared by
THE MITRE CORPORATION
Bedford, Massachusetts
Contract AF19(628)-2390

FOREWORD

The work reported in this document was performed by The MITRE Corporation, Bedford, Massachusetts, for Electronic Systems Division, Air Force Systems Command under Contract AF 19(628)-2390. This information was originally published in Working Paper W-7550, The MITRE Corporation, February 1965.

REVIEW AND APPROVAL

Publication of this technical report does not constitute Air Force approval of the report's findings or conclusions. It is published only for the exchange and stimulation of ideas.

ANTHONY P. TRUNFIO, Technical Advisor
Development Engineering Division
Directorate of Planning and Technology

ABSTRACT

An approach to digital nonlinear prediction is proposed and analyzed. The basic relations are developed. The nonlinear operator is obtained by quantization of the data.

The model is developed in terms of occupancy of data cells in N-space. Extensions to increase occupancy and reduce error are formulated. Illustrative results are included.

A comparison with linear techniques is made and over-all conclusions on error, quantization level, length of data required, time invariance, etc., are provided.

TABLE OF CONTENTS

	<u>Page</u>
1.0 INTRODUCTION	1
2.0 THE MODEL	1
3.0 QUANTIZATION AND REALIZATION	4
3.1 Quantization Error and Data Calibration	7
4.0 METHODS OF ESTABLISHING ALL OCCUPANCY.	11
4.1 Method B	12
4.2 Method A	13
4.3 Model C	14
4.4 Adaptive Control.	15
4.5 Population of Cell Occupancies.	16
4.6 Comparison of Methods A and B	17
5.0 GENERAL LINEAR PREDICTOR	17
6.0 OBSERVATIONS, REMARKS, AND CONCLUSIONS	19
6.1 Linear	19
6.2 Nonlinear	19
6.3 Linear versus Nonlinear	20
REFERENCE	25

1.0 INTRODUCTION

Methods to effect statistical prediction have involved polynomial fitting and correlation and/or spectral analysis. Both can use a minimum mean square error criterion and result in a set of weights as an optimum linear operator. In both cases, calculations of the weights based on knowledge or computation of the pertinent statistics is made.

Our concern here is with a non linear approach which involves any intermediate determination of the statistics. It also offers a readily available means for judging error and adjusting for an improved prediction. Being a non linear method, the results should be at least as good as a corresponding linear technique.

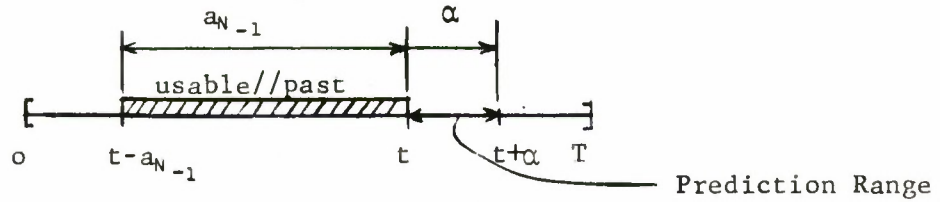
To obtain the desired non linearity, a quantization of the data is required. Of course such a quantization itself degrades the error possible with the technique. The technique as applied to digital simulation forms a variation on an approach discussed in Reference 1. We now discuss the approach.*

2.0 THE MODEL

We let t be the present time and T the total interval of data for processing. We let α be the time advance of prediction. Without loss of generality, we take the data interval as $[0, T]$ and consider the usable past at any t from t to $t - a_{N-1}$. For equally spaced sampled data $a_{N-1} = (N-1)\delta$ where δ is the sampling interval. These definitions may be summed up in the following sketch

Author's Note

*The work reported here was originally considered by the author in the summer of 1960 while at S.T.L. Other matters prevented a proper evaluation and summary at that time. The present report represents a current effort to fill this need.



Let the continuous parameter time series be $X(t)$. We then estimate $X(t+\alpha)$ by

$X^*(t+\alpha)$, the predicted value at $t+\alpha$ using $X(\sigma)$; $\sigma \leq t$

In the sampled finite past case we have

$X_N^*(t+\alpha)$, the predicted value at $t+\alpha$ using $X(t-a_i)$; $i=0, 1, \dots, N-1$

i.e. by $\{X_0, X_{-1}, X_{-2}, \dots, X_{-(N-1)}\}$.

Although $\Delta_i = a_{i+1} - a_i$ may not be equal, nothing is lost by assuming so, since the average value of the Δa is constrained by the sampling theorem.

We take

$X^*(t+\alpha) = F[X(t), X(t-\delta), \dots, X(t-a_{N-1})]$ where

F is a time dependent non linear function. In this case we can always make the instantaneous square error $\epsilon = (X(t+\alpha) - F[X(t), X(t-a_1), \dots, X(t-a_{N-1})])^2 = 0$.

However since F is to be used when $X(t+\alpha)$ is not available, it is better not to have F time dependent. Hence we consider a range of times over which a time invariant operator F must minimize the mean squared error given by

$$\overline{\epsilon}_\alpha = \frac{1}{\beta_\alpha} \int_{a_{N-1}}^{T-\alpha} (X(t+\alpha) - F[X(t), X(t-a_1), \dots, X(t-a_{N-1})])^2 dt$$

with $[a_{N-1}, T - \alpha]$ the maximum range over which to consider error or in other words over which we can distinguish a "present" value. In the above $\beta_\alpha = T - \alpha - a_{N-1}$. As seen the value of $X^*(t+\alpha)$ depends on the

number, N , of data points used in the memory.

We replace $^\dagger F$ by an infinite series of terms whose orthogonality is invariant to X . Thus,

$$F = \sum_{-\infty}^{\infty} A_n \phi_n [X(t), X(t-a_1), \dots, X(t-a_{N-1})]$$

where, independent of X values,

$$\frac{1}{T-a_{N-1}} = \int_{a_{N-1}}^{T-a} \phi_n [] \phi_m [] dt = 1; n = m$$

$$= 0; n \neq m$$

Then minimum $\overline{\epsilon_a^2}$ gives, using $\frac{\partial \overline{\epsilon_a^2}}{\partial A_n} = 0$ and the orthogonality of the

$$\{A_n\}$$

$\{\phi\}$,

$$A_n = \frac{1}{T-a_{N-1}} \int_{a_{N-1}}^{T-a} X(t+a) \phi_n [] dt$$

$$\frac{\overline{X\phi_n}}{\overline{\phi_n^2}} = \frac{\langle X, \phi_n \rangle}{\langle \phi_n, \phi_n \rangle}$$

It is impossible, however, to have completeness using such a representation on any $X(t)$ (i.e., for continuous valued X). However, quantization of X provides a realization of the desired orthogonal set for all $X(t)$. In other words quantizing X allows construction of a $\{\phi_n\}$ set for all X , the approximation being dependent on the fineness (degree) of quantization. As X is more finely quantized, the

$$F[X(t), \dots, X(t-a_{N-1})] \rightarrow \sum_n A_n \phi_n [X(t), \dots, X(t-a_{N-1})]$$

for any $X(t)$ with the $\{\phi_n\}$ orthogonal and independent of X .

[†] See Reference 1

The orthogonality depends however on the range β_α , chosen. We then have

$$\begin{aligned}\overline{\epsilon_\alpha^2} &= \frac{1}{\beta_\alpha} \int_{a_{N-1}}^{T-\alpha} (x(t+\alpha) - \sum_0^\infty A_n \phi_n [])^2 dt \\ &= \frac{1}{\beta_\alpha} \left(\int_{a_{N-1}}^{T-\alpha} x^2(t+\alpha) dt - \sum_0^\infty \beta_n A_n^2 \right) \\ &= \frac{1}{\beta_\alpha} \left(\int_{T-\alpha-\beta_\alpha}^{T-\alpha} x^2(t+\alpha) dt - \sum_0^\infty \frac{\left(\int_{T-\alpha-\beta_\alpha}^{T-\alpha} x(t+\alpha) \phi_n dt \right)^2}{\beta_n} \right)\end{aligned}$$

with
$$\beta_n = \int_{T-\alpha-\beta_\alpha}^{T-\alpha} \phi_n^2 dt$$

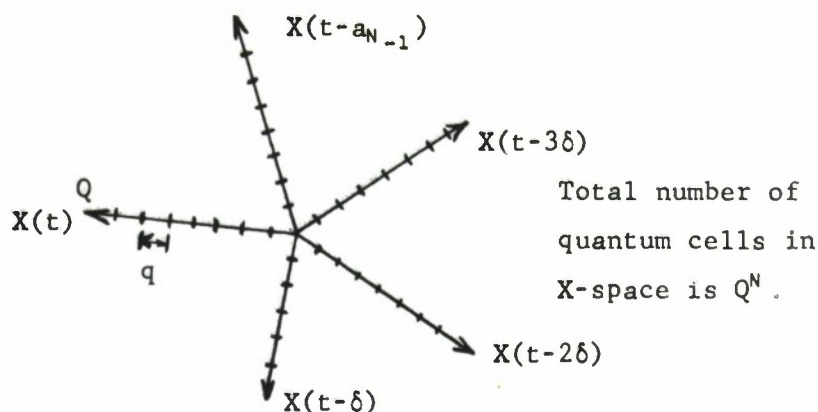
We see that A_n represents the projection of predicted X values

in the ϕ_n direction of the $\{\phi_n\}$ space. Thus $\overline{\epsilon_\alpha^2}$ compares the true value $\frac{\int_{T-\alpha-\beta_\alpha}^{T-\alpha} x^2(t+\alpha) dt}{\beta_\alpha}$ with $\sum_0^\infty \left(\frac{\beta_n}{\beta_\alpha} \right) A_n^2$ where $\frac{\beta_n}{\beta_\alpha}$ represents the

probability of getting A_n over β_α . For a fixed t , $X(t+\alpha)$ is estimated by a single A_n for some ϕ_n .

3.0 QUANTIZATION and REALIZATION

If for each n , the $\phi_n [X(t), X(t-\alpha_1), \dots, X(t-\alpha_{N-1})] = 1$ or 0 as a function of the N - dimensional argument of X values, then $A_n = \frac{[X, \phi_n]}{\beta_n}$ is a conditional expectation. That is the average of $X(t+\alpha)$ conditioned on the occurrence of $\phi_n(t) = 1$ as t varies over β_α and β_n counts the number of such occurrences. Let us now consider for equally spaced data the N - dimensional space of $\{X(t-k\delta)\}_0^{N-1}$ values. Let each $X(t-k\delta)$ range be divided into Q quantum each of width q as shown.



Let $n = 1, 2, 3, \dots, Q^N$

$$\begin{aligned} \text{We take } \phi_n [\{X(t-i\delta)\}_0^{N-1}] &= 1 && \text{if } \{X\}_0^{N-1} \in \text{cell } n \\ &= 0 && \text{if } \{X\}_0^{N-1} \notin \text{cell } n \end{aligned}$$

Thus if ϕ_n occurs at t then we take A_n as $X^*(t+\alpha)$ and A_n is calculated as average of $X(t'+\alpha)$ values at all $t' < t$ times when ϕ_n occurs. The $\{\phi_n\}$ set remains fixed only if the cell structure in N -space does. We note that the orthogonality of the $\{\phi_n\}$ is independent of the t range of integration. Also since $\langle \phi_n, \phi_n \rangle = 0$ only if $\phi_n = 0$, the

$$A_n = \frac{\langle X, \phi_n \rangle}{\langle \phi_n, \phi_n \rangle} \text{ are bounded}$$

The normality condition

$$\frac{1}{\beta\alpha} \int_{T-\alpha-\beta\alpha}^{T-\alpha} \phi_n^* () dt = \frac{1}{\beta\alpha} \int_{T-\alpha-\beta\alpha}^{T-\alpha} \phi_n dt$$

depends on the X values that is which member of the ensemble is chosen and the span of calculation $\beta\alpha$. We see that $\left\{ \frac{\phi_n}{\beta_n^{1/2}} \right\}$ produces an orthonormal

set; $\beta_n = [\phi_n, \phi_n]$.

It is useful to rewrite the formulae for A_n and $\overline{\epsilon_\alpha^2}$ in a style which reflects the cell matching imposed by the quantization as reflected by

the $\{\phi_n\}$ set.

$$A_n = \frac{\int_{\tau_n} X(t+\alpha) dt}{\int_{\tau_n} dt} = \frac{\int_{\tau_n} X(t_{i,n}+\alpha) dt}{\tau_n}$$

where $\tau_n = \{t_i: \phi_n[t] = 1, t \in [a_{N-1}, T-\alpha]\}$ so that

$$A_n = \overline{\chi(t+\alpha)} \quad \text{conditioned an occurrence of cell } n.$$

$$= \frac{\sum_{i=1}^{N_n} X(t_i+\alpha)}{N_n} \quad \text{for sampled data}$$

and $\sum_{n=1}^{Q^N} N_n = \beta_\alpha$ for a fixed cell structure over $t_i \in [a_{N-1}, T-\alpha]$

so that any $\{X\}_0^{N-1}$ sequence can be in no more than one cell.

For non-multiple use of data sets $\{X\}_0^{N-1}$ but allowing overlapping cells the total number of matches is still β_α . However the orthogonality of the ϕ_n set and the expressions for A_n and $\overline{\epsilon_\alpha^2}$ previously developed rely on the cells in N-space being non-overlapping.

We also have

$$X^*(t+\alpha) = \sum_{1 \leq n \leq Q^N} A_n \phi_n = \{A_n : \phi_n(t) = 1, t \in [a_{N-1}, T-\alpha]\}$$

and

$$\overline{\epsilon_\alpha^2} = \frac{\sum_{i=1}^{\beta_\alpha} [X(t_i + \alpha) - A_{n_i}]^2}{\beta_\alpha}$$

Incidentally we may view the quantization as a means to set up a special set of orthogonal functions for any X by writing $\hat{\Phi}$ as a product given by

$$\hat{\Phi}_{n,m,k}, \dots, \ell[X(t), X(t-a_1); \dots, X(t-a_{N-1})] = \hat{\Phi}_n[X(t)] \hat{\Phi}_m[X(t-a_1)] \dots$$

$$\text{Finally } e_{\alpha}^2 = (X(t+\alpha) - F[\])^2 = (X(t+\alpha) - \sum A_n \hat{\Phi}_n)^2$$

$$\text{which can be written } (X(t+\alpha) - \sum A_n \hat{\Phi}_n - (F[\] - \sum A_n \hat{\Phi}_n))^2$$

$$\text{or } (X(t+\alpha) - F - (\sum A_n \hat{\Phi}_n - F))^2$$

These forms indicate two kinds of errors.

- (1) Optimum F is not sufficient, i.e. future not predictable from the past (not usual).
- (2) $\sum A_n \hat{\Phi}_n$ is not a good enough approximation, i.e. orthogonal set not complete.

For $\hat{\Phi}_n$ derived via quantization, the level q plays a roll in (2) and indeed contributes a quantization error as well.

Finally we remark that when $F = F(t)$ that is the $\hat{\Phi}_n$ are changing with t , the number of possible A_n = number of possible cells $> Q^N$.

3.1 Quantization Error and Data Calibration

Given a joint distribution over data, we can always construct an orthonormal function set relative to the distribution. For example, if X is normally distributed, we choose the $\hat{\Phi}_n$ as Hermite polynomials. In dealing with time series we interchange time and ensemble averages with time averages taken over sufficiently long data stretches providing the process is ergodic and so stationary. With X normal, the optimum predictor is a linear one. This result can serve as a

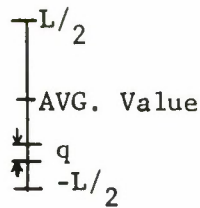
reference to the non-linear models.

Approximating F via quantization gives

$$A_n = \langle X, \phi_n; L \rangle$$

where L is the peak to peak signal variation in the interval used.

We choose the quantum step q so as to produce the desired e^2 and divide the X range into $Q = L/q$ quantum steps.



Quantization thus converts $X(t)$ into a step function $S(t)$ where any step level S_i results from X values given by $S_i - q/2 \leq X < S_i + q/2$. The quantization error, e, is a random variable with the X distribution and

$-q/2 \geq e < q/2$. With q small enough we assume that e is approximately uniformly distributed in any one interval so that the phase (ensemble) average of e^2 is given by

$$\begin{aligned} \text{AVG.}(e^2) &= \frac{1}{q} \int_{-q/2}^{q/2} e^2 de = \frac{2}{q} \frac{q^3}{3 \times 8} = q^2/12 \\ &= \frac{1}{12} \left(\frac{L}{Q} \right)^2 \end{aligned}$$

If we assume 1024 divisions for + or - values (a representative value for some A-D converters) then $Q = 2048$ and

$$\text{AVG.}(e^2) \sim \left(\frac{L}{2^{10}} \right)^2$$

Thus writing the mean square error of quantization as $K q^2$, K, a constant,

the mean square error in the prediction model $\overline{\epsilon^2} \geq K q^2$

From

$$\overline{\epsilon_\alpha^2} = \frac{1}{\beta_\alpha} \int_{T-\alpha-\beta_\alpha}^{T-\alpha} x^2(t+\alpha) dt - \sum_{n=0}^q \beta_n A_n^2$$

if quantization level q is too large, $A_{n_0} \rightarrow \bar{X}$ for some single $n=n_0$

and $A_n \rightarrow 0$, $n \leq n_0$, $\beta_{n_0} \rightarrow \beta_\alpha$,

$$\overline{\epsilon^2} \sim 0 \left(\frac{1}{\beta_\alpha} \int_{T-\alpha-\beta_\alpha}^{T-\alpha} (X(t+\alpha) - \bar{X})^2 dt \right) \sim 0(K' q^2), \quad K' > K.$$

Of course $\overline{\epsilon^2}$ may be large with q small for operator F not near optimum.

It is well known that in general the non-linear operator extends bandwidth due to intermodulation of components. Also, if we consider quantization error in N -space as a distance d , then

$$d_{\max} = \sqrt{\sum_{i=1}^N q_i^2} \quad \text{with } q_i \text{ the quantum level in the } i^{\text{th}} \text{ dimension.}$$

If $q_i = q$ for all i ,

$$d_{\max} = (Nq^2)^{\frac{1}{2}} = N^{\frac{1}{2}}q$$

For preserving constant distance (error), we have

$$q \sim \frac{1}{N^{\frac{1}{2}}}$$

so that the more past samples used to form a data set $\{X\}^{N-1}$, the smaller q required to maintain the quantization error.

With $\alpha, \beta_\alpha, N, q$ variable there is a range of possibilities with various influences on $\overline{\epsilon_\alpha^2}$.

As N increases, q should decrease. With decreasing q and increasing

N the β_n , number of matches over β_α , should drop. This means that the number of matches of data sets $\{X(t')\}_o^{N-1}$ to some reference set $\{X(t)\}_o^{N-1}$, $t' < t$ decreases. Thus although quantization error may be fixed, variability of the estimate of an A_n (associated with $\{X(t)\}_o^{N-1}$) may rise due to the reduced number of matches. In general we desire as many cell matches as possible at each t .

For example, unless q is large (and so larger prediction errors) or a trend sensitive parameter, such as conditional expectation function, is used, no reasonable prediction can be made on trend type data since obtaining cell matches becomes difficult or impossible.

One method to increase number of cell occupancies involves using multiple use of data and/or overlapping cell structure with time. These procedures allow for time varying $\{\phi_n\}$ sets and form a degression from the theory presented above. Details of implementing such procedures are discussed in the next section.

Another way to improve cell matching is to deal with essentially trend free data using polynomial interpolation techniques. Simple linear calibration may be sufficient.

The calibration of the input data can be given as

$$X(t) = ay(t)+b = L(y)$$

where L represents a linear operation $a(\neq 0)$ and b are known factors which can be time varying.

In the prediction model we operate an $X(\sigma)$, $\sigma \leq t$ with F to give $X^*(t+\alpha)$ so that in a 1:1 manner we have

$$y^*(t+\alpha) = \frac{X^*(t+\alpha)-b}{a}$$

4.0 METHODS OF ESTABLISHING CELL OCCUPANCY

As noted above we desire models which increase cell occupancy.

This involves basically

1. defining the cell structure
2. establishing occurrence of occupancy

We now discuss three variations on dealing with these two factors.

1) This definition establishes from the next new data set left from the $\{X\}_0^{N-1}$ data sets remaining after having filled all previous cells set up. This allows no multiple use of data and reduced amount of cell overlap. A data set defining a cell is taken as the center of the cell. In order to maximize the number of cell prediction the data sets are used starting first with the data set associated with the current time, t . The total number of sets equals β_α .

2) This is a variation on (1) which allows multiple use of any set $\{X\}_0^{N-1}$ for matching except those previously used as cell centers.

Thus both multiple use of data and cell overlap occurs. We have

$K(\text{number of cells setup}) \leq \beta_\alpha \leq \sum_{j=1}^K d_{nj} = \text{total number of sets used}$
in the K setup cells where $d_{nj} = \text{number of matching } \{X\}_0^{N-1} \text{ sets}$
into cell setup at n_j .

If we could essentially assume no cell overlap that is a fixed

$\{\phi_n\}$ set (fixed cell structure) then for (1) and (2) we would have

$$\overline{\epsilon_\alpha^2} = \frac{1}{K} \sum_{j=1}^K (X(t_{ij} + \alpha) - A_{ij})^2$$

$$A_{ij} = \frac{1}{d_{ij}(\alpha)} \sum_{i_j=1}^{d_{ij}(\alpha)} X(t_{ij} + \alpha)$$

3) Another variation takes all past sets as cell centers so that $K = \beta_\alpha$. This method gives more cell overlap and multiple use of data but maximizes cell occupancy at each t in $[a_{N-1}, T-\alpha]$ and places the center of the cell for t at $\{X(t)\}_0^{N-1}$ at each t . Briefly in summary (1) tends to minimum occupancy (overlap) with the maximum of these at current time t while (3) produces maximum occupancy (overlap). We call Method A, Approach 3) and Method B, Approach 1) together with cells held non-overlapping. We first discuss Method B for $\overline{\epsilon_\alpha^2}$ and A_N . We also include another variation called Model C.

4.1 Method B

With fixed cell structure that is $\{\phi_n\}$ set fixed, we have from before,

$$\overline{\epsilon_\alpha^2} = \frac{1}{\beta_\alpha} \left(\int_{T-\alpha-\beta_\alpha}^{T-\alpha} X^2(t+\alpha) dt - \frac{\sum_0^\infty \left(\int_{T-\alpha-\beta_\alpha}^{T-\alpha} X(t+\alpha) \phi_n [] dt \right)^2}{\beta_n} \right)$$

with sampled (discrete time series) data

$$\overline{\epsilon_\alpha^2} = \frac{1}{\beta_\alpha} \sum_{k=1}^K \left(\sum_{j=1}^{n_k(\alpha)} (X(t_{jk}+\alpha) - A_{jk})^2 \right); \quad \tilde{\beta}_\alpha = \sum_{k=1}^K n_k(\alpha)$$

with

$$A_{n_k} = \frac{\frac{1}{\beta_\alpha} \sum_{j=1}^{n_k} X(t_{jk}+\alpha) \phi_{jk}}{\frac{1}{\tilde{\beta}_\alpha} \sum_{j=1}^{n_k} \phi_{jk}^2} = \frac{\sum_{j=1}^{n_k} X(t_{jk}+\alpha)}{\sum_{j=1}^{n_k} 1}$$

so that

$$\overline{\epsilon_\alpha^2} = \frac{1}{\beta_\alpha} \left(\sum_{n=1}^{\tilde{\beta}_\alpha} X^2(t_n+\alpha) - \sum_{k=1}^K n_k A_{jk}^2 \right)$$

$$= \frac{1}{\tilde{\beta}_\alpha} \left(\sum_{n=1}^{\tilde{\beta}_\alpha} X^2(t_n + \alpha) - \sum_{k=1}^K \left[\frac{\sum_{j=1}^{n_k} X(t_{jk} + \alpha)}{n_k} \right]^2 \right)$$

The matchings (for each α) to the cell setup by past set $\{X(T)\}_0^{N-1}$ taken at present time T (i.e., n_T) produce A_{n_T} as $X^*(T+\alpha)$ for each α .

4.2 Method A

As noted, the method sets up each of the $\tilde{\beta}_\alpha$ data sets as centers for cells. The approach thus allows multiple use of any data set $\{X\}_0^{N-1}$ and overlapping cell structures.

Now,

$$\overline{\epsilon_\alpha^2} = \frac{1}{\tilde{\beta}_\alpha} \sum_{n=1}^{\tilde{\beta}_\alpha} [X(t_n + \alpha) - A_n]^2 = \frac{1}{\tilde{\beta}_\alpha} \sum_{n=1}^{\tilde{\beta}_\alpha} [X^2(t_n + \alpha) - 2X(t_n + \alpha)A_n + A_n^2]$$

for each of the $\tilde{\beta}_\alpha$, $n=1, 2, \dots, \tilde{\beta}_\alpha$ where we take

$$A_n = \frac{1}{l_n(\alpha)} \sum_{j=1}^{l_n(\alpha)} X(t_{jn} + \alpha) ; \quad \sum_{n=1}^{\tilde{\beta}_\alpha} l_n \geq \tilde{\beta}_\alpha$$

Then, without orthogonality effects allowed since the $\{\phi_n\}$ sets correspond to overlapping cells,

$$\overline{\epsilon_\alpha^2} = \frac{1}{\tilde{\beta}_\alpha} \sum_{n=1}^{\tilde{\beta}_\alpha} \left[X^2(t_n + \alpha) - 2X(t_n + \alpha) \frac{\sum_{j=1}^{l_n} X(t_{jn} + \alpha)}{l_n} + \left(\frac{\sum_{j=1}^{l_n} X(t_{jn} + \alpha)}{l_n} \right)^2 \right]$$

4.3 Model C

In order to increase further the number of occupancies and so hopefully reduce $\overline{\epsilon_\alpha^2}$ averaged over α for $\alpha_{\min} < \alpha < \alpha_{\max}$, for the cell specified at the current T value, we consider sequentially $\{X(T-j\delta)\}_0^{N-1}$ as potentially equivalent $\{X(T)\}_0^{N-1}$ data sets with adjusted α values[†]. This is valuable when $\{X(t)\}_0^{N-1}$ has none or few matches. For α the data sets then used for equivalent prediction points are those at $t \leq T-\alpha$. We designate the corresponding data sample numbers as n_T and n_{Tj} where $t_j = T-\alpha-j\delta$ where δ is the data spacing.

Thus, at α using the set at n_{Tj} , we use $\alpha_j = \alpha+j\delta$, $j = 0, 1, \dots, R$ and look for occupancy of the cell determined by $\{X(T-j\delta)\}_0^{N-1}$ by the sets $\{X(t)\}_0^{N-1}$, $t < T - j\delta$ that is for $n < n_{Tj}$.

At j the maximum α value possible as in previous methods is determined by $\alpha_{\max_j} = n_{Tj} \cdot n$ where n are the matching data sets $\{X\}_0^{N-1}$ to $\{X(T-j\delta)\}_0^{N-1}$. The minimum α in $\alpha_{\min_j} = j\delta$.

We take

$$\epsilon_j \triangleq \frac{1}{\alpha_{\max_j}} \sum_{\alpha_j=\alpha_{\min_j}}^{\alpha_{\max_j}} \overline{\epsilon_{\alpha_j}^2} \quad \text{for each } j.$$

Then we choose the n_{Tj} and associated matching sets which give $\min_j \epsilon_j$ and calculate the A_n set for the $\alpha_{\min_j} < \alpha_j < \alpha_{\max_j}$ corresponding to the α . If $\alpha_{\max_j} < \alpha_{\max}$, then for $n_{Tj} < n_T$ we choose the next lowest ϵ_j which has $\alpha_{\max_j} > \alpha_{\max_j}$. The process is repeated until $\alpha_{\max_j} = \alpha_{\max}$ or the j range is depleted.

[†] Of course increased occupancy must be weighed against basically higher $\overline{\epsilon_\alpha^2}$ as α increases.

4.4 Adaptive Control

As we have seen, parameters N , q , β_α influence $\overline{\epsilon_\alpha^2}$. It is thus useful to have criterion, even if only crude, by which to judge such effects.

We have noted the error amplitude resolution in the sampling process is $q^2/12$ where q_0 is the amplitude resolution or imposed quantization range q . For example, with binary data and $q = 2^k$, the error is $4^{k-1}/3$.

Referenced to the original (unquantized) data before sampling and excluding other error sources, the total error $\overline{\epsilon_{TOT}^2}$ is approximately $\overline{\epsilon_\alpha^2} + q^2/12$ with

$$0 \leq \overline{\epsilon_\alpha^2} \leq \frac{1}{\beta_\alpha} \int_{T-\alpha-\beta_\alpha}^{T-\alpha} X^2(t+\alpha) dt \sim \frac{1}{\tilde{\beta}_\alpha} \sum_{n=1}^{\tilde{\beta}_\alpha} X^2(t_n+\alpha)$$

Thus, we may usually expect

$$q^2/12 \leq \overline{\epsilon_{TOT}^2} \leq q^2/12 + \frac{1}{\tilde{\beta}_\alpha} \sum_{n=1}^{\tilde{\beta}_\alpha} X^2(t_n+\alpha)$$

With $\overline{\epsilon_\alpha^2} = \frac{1}{\tilde{\beta}_\alpha} \sum_{n=1}^{\tilde{\beta}_\alpha} X^2(t_n+\alpha) - \sum_{k=0}^q \beta_k A_k^2$ and $q^2/12$ not truly independent

errors since $\overline{\epsilon_\alpha^2}$ depends on q , we might have $\overline{\epsilon_{TOT}^2} \leq \overline{\epsilon_\alpha^2}$ for example.

From

$$\overline{\epsilon_\alpha^2} = \overline{X^2(t+\alpha)} - \sum_{k=1}^K \frac{[\sum_{j=1}^n X(t_{jk}+\alpha)]^2}{n_k}$$

as $n_k \rightarrow 0$ (no cell matching) $\overline{\epsilon_\alpha^2} \rightarrow \overline{X^2(t+\alpha)}$.

We may then consider using the values $q^2/2$ and $\overline{X^2(t+\alpha)}$ as bounds for testing $\overline{\epsilon_\alpha^2}$ when changing q , N , β_α to determine acceptability or for improving $\overline{\epsilon_\alpha^2}$ to obtain $X^*(T+\alpha)$.

4.5 Population of Cell Occupancies

With Q steps of magnitude q and N past point data sets there are Q^N possible cells (i.e., A_n 's using fixed structure). The percentage actually used is much smaller, say less than 1% on finite sections of data. In digital simulation it is not required to store all the possible or even expected $\{A_n\}$ set.

For example, with $N=2$ and independent normally identically distributed data at X_n , X_{n-1} for all n with a standard deviation of σ the number of A_n drops from Q^2 to $3\sigma(2Qq-3\sigma)$, $qQ/2 > 3\sigma$. An estimate of 3σ can be taken as the largest amplitude increment between successive samples.

Runs have been made on data to determine and remove trends using data types such as

- (1) monotonic increasing data of a component of velocity sampled at 10 times per second.
- (2) periodic high frequency type data from accelerometer residuals sampled at 10 times per second.
- (3) radar trajectory position data at 10 times per second.

A definite $X^*(t+\alpha)$ versus q relationship was difficult to establish for various ranges of N and β_α .

4.6 Comparison of Methods A and B

It is convenient to display errors for comparison by

$$\% \text{ rms error} = \% \left(\overline{\epsilon_{\alpha}^2} \right)^{\frac{1}{2}} = \sqrt{\frac{\overline{\epsilon_{\alpha}^2}}{\overline{x^2}}}$$

$$\text{and } \% \text{ actual error} = \% \epsilon_{\alpha}^* = \frac{X^* - X}{X}$$

Due to the flexibility of centering cells, it is expected that A gives better X^* (low $\% \epsilon_{\alpha}^*$) than B so that the possibility of cell occupancy is more likely.

It is recalled that ϵ_{α}^* was taken with respect to prediction ($x^* = A_{n_T}$) point sample n_T only while $\overline{\epsilon_{\alpha}^2}$ in Method A was calculated over matches not only to data set for n_T but matches to data sets for all n used.

It is suggested a revised method A, say A^1 , be examined to remedy these situations by (1) allowing actual predictions to be calculated over the data range and not just at n_T and (2) calculating $\overline{\epsilon_{\alpha}^2}$ at each prediction point, n_T' , from matches of past data sets to the data set at n_T' only. At each n_T' , $\{X(T')\}_0^{N-1}$ is considered fixed allowing

$$A_n = \begin{matrix} \langle X, \bar{\Phi}_n \rangle \\ \langle \bar{\Phi}_n, \bar{\Phi}_n \rangle \end{matrix} \text{ to be used with cell center fixed by } \{X(T')\}_0^{N-1} \text{ which}$$

shifts with T' to make $\{A_n\}$ set time varying with n_T' .

5.0 GENERAL LINEAR PREDICTOR

"General" is taken to mean no assumption of stationarity. The use of the general linear predictor provides a basis of comparison for the non linear predictor.

For continuous (nonquantized) data the form of the basic system of equations for determining the coefficients of the optimum linear predictor (operator) are

$$\int_0^{T-\alpha} X(t-\tau_n) X(t+\alpha) dt = \sum_{m=0}^{N-1} \left(\int_0^{T-\alpha} X(t-\tau_n) X(t-\tau_m) dt \right) a_m$$

$$n=0,1,\dots,N-1$$

which for sampled data becomes (with a slight abuse of notation)

$$\sum_{k=0}^{\tilde{\beta}_\alpha} X_{k-n} X_{k+\alpha} = \sum_{m=0}^{N-1} \left(\sum_{k=0}^{\tilde{\beta}_\alpha} X_{k-n} X_{k-m} \right) a_m; n=0,1,\dots,N-1$$

It can be noted that an analogous form of equation appropriate to the nonlinear method is given by

$$\sum_{k=0}^{\tilde{\beta}_\alpha} \Phi_n X_{k+\alpha} = \sum_{m=0}^M \left[\sum_{k=0}^{\tilde{\beta}_\alpha} \Phi_n(X_k, \dots) \Phi_m(X_k, \dots) \right] A_m; n=0,1,\dots,M$$

If we define $R(\xi, \eta) \triangleq \sum_{k=0}^{\tilde{\beta}_\alpha} X_{k-\xi} X_{k-\eta}$, we have

$$\{R(n, -\alpha)\} = ((R(n, m))) \{a_m\}$$

$$(nX1) \quad (nXm) \quad (mX1)$$

so that

$$\{a_m\} = ((R(n, m)))^{-1} \{R(n, -\alpha)\}$$

with $((R(n, m)))$ and $\{R(n, -\alpha)\}$ both functions of α and the point from which prediction is made.

With allowance for variable α , a comparison with the nonlinear method indicates that for the same amount of information, linear processing time is higher by a factor as high as 10.

We remark that calculations at data rates higher than that provided can be afforded in all processing modes by a sliding Lagrangian fit such as a cubic.

Finally comparisons are for all prediction points having one or more matches in the nonlinear processing. Also, the rms values of actual errors taken over these points are also found useful.

6.0 OBSERVATIONS, REMARKS AND CONCLUSIONS

It was found useful to use $\overline{\epsilon}_0$ = r.m.s. value of observed error,

$$\overline{\epsilon_\alpha^2}^{\frac{1}{2}} \text{ calculated over matches}$$

associated with a given point of prediction

and $\hat{\epsilon}_A$ = r.m.s. value of actual error,

$$\left(\overline{(\epsilon_A^*)^2} \right)^{\frac{1}{2}} \text{ averaged over all } x: x$$

available as point of prediction is changed.

6.1 Linear

For small α , the average $\hat{\epsilon}_A$ increases with N as should be expected since the linear operator becomes potentially more limited as the extent of the past (memory) increases. The error $\overline{\epsilon}_0$ increases by about a 1.6 factor as α goes from 1 to 2.

6.2 Nonlinear

For probability of occupancy as low as 5%, the $\hat{\epsilon}_A$ remain reasonably fixed. For q low enough, $\overline{\epsilon}_0$ decreases with increased N as the occupancy probability drops. In other words, at each prediction point

the $\overline{\epsilon}_0$ decreasing measures the degree of correct selectivity. As N increases with q high enough, large variation in actual error $\hat{X}^* - X$ is allowed. Indeed $\hat{\epsilon}_A$ decreases with increasing N while $\hat{\epsilon}_A$ increases by a factor of 1.5 as α goes from 1 to 2. As q decreases the $\hat{\epsilon}_A$ and $\overline{\epsilon}_0$ agree. As α increases the q value necessary for this agreement increases.

For low N and α (say 1), $\hat{\epsilon}_A$ can decrease with β_α while $\overline{\epsilon}_0$ increases for the same q. Both $\hat{\epsilon}_\alpha$ and $\overline{\epsilon}_0$ reduce with increased fs (sampling rate) for the same effective β_α , N, α . Also, $\hat{\epsilon}_A$ and $\overline{\epsilon}_0$ appear relatively insensitive to the interval between prediction points or to the over-all span of these points, the latter effect relying on the degree of stationarity present.

6.3 Linear Versus Nonlinear (for a given data sampling rate)

Although the linear procedure requires no quantization, comparisons with the nonlinear method are more equitable with the nonlinear results when this discrepancy is taken into account. Basically unless a sufficiently low q can be used the $\hat{\epsilon}_A$ values from linear processing are lower as was the case for $\alpha = 1, 2$ and all the N, β_α and data used. As q lowers to favor the nonlinear error, the probability of a prediction occurring falls sharply for the β_α range used.

The sample conditional expectation obtained with finite length records and the lack of completeness with q values > 0 result in gaps in the determination of the conditional expectation on which, indeed, the prediction is based.

A first and crude step toward overcoming this deficiency can be taken via a two point interpolation as follows. (More exactly as N-dimensional interpolation on the data sets $\{X(t)\}_0^{N-1}$ of nearest matches could be made to specify $X^*(t+\alpha)$ associated with $\{X(T)\}_0^{N-1}$).

We let n be the count corresponding to time t as before and $\{X\}_{n_1}$ and $\{X\}_{n_2}$ be the two N-dimensional data sets having the two smallest distances d_1 and d_2 ($\leq d_1$) respectively in N-space from $\{X\}_{n_1}$. We take

$$X^*(T+\alpha) = X(t_1+\alpha) + \frac{d_1}{d_1+d_2} \Delta X;$$

with

$$\Delta X = X(t_2+\alpha) - X(t_1+\alpha)$$

The minimum mean square error criterion results in the predicted value $X^*(t+\alpha)$ being given by a quantized version (for determining cell occupancy) of the sample conditional expectation $E(X(t+\alpha) | C)$ where C is the cell whose center is determined by $\{X(T)\}_0^{N-1}$ rather than the non-quantized version $E(X(t+\alpha) | \{X(T)\}_0^{N-1})$. We have

$$X^*(t+\alpha) = E(X(t+\alpha) | C) = \int_X X(t+\alpha) P(X(t+\alpha) | C) dx$$

and for sampled data

$$X^*(t+\alpha) = \sum X P_\alpha(X|C)$$

With a finite stretch of data we must consider

$$P_\alpha(X_k \leq X < X_{k+1}) | C = \frac{K_{\alpha,k,c}}{l_{\alpha,c}}$$

where $K_{\alpha,k,c}$ = number of times an $X(t+\alpha)$ value from the $\tilde{\beta}_\alpha$ sequence of X values lies within the (X_k, X_{k+1}) range given that its corresponding past data set $\{X(t)\}_0^{N-1}$ enters cell C .

$l_{\alpha,c}$ = number of past data sets which enter cell C over the $\tilde{\beta}_\alpha$ sequence of data points.

We may note here the various first order, joint and conditional probabilities as follows:

$$P_\alpha [X_k \leq X < X_{k+1} \mid C] = \frac{K_{\alpha,k,c}}{l_{\alpha,c}} = \overset{\text{Shortened Notation}}{P_\alpha(X|C)}$$

$$P_\alpha(C) = \frac{l_{\alpha,c}}{\tilde{\beta}_\alpha}$$

$$P[X_k \leq X < X_{k+1}, C] = \frac{K_{\alpha,k,c}}{\tilde{\beta}_\alpha} = P_\alpha(C) P_\alpha(X|C) = P_\alpha(X,C)$$

$$P[X_k \leq X < X_{k+1}] = \frac{K_{\alpha,k}}{\tilde{\beta}_\alpha} = P_\alpha(X)$$

$$P[C|X_k \leq X < X_{k+1}] = \frac{K_{\alpha,k,c}}{K_{\alpha,k}} = \frac{P_\alpha(X,C)}{P_\alpha(X)} = P_\alpha(C|X)$$

where $K_{\alpha,k,c}$ and $l_{\alpha,c}$ are defined above and $K_{\alpha,k}$ is the number of times an X value in the $\tilde{\beta}_\alpha$ sequence of data points lies within the range (X_k, X_{k+1}) .

In Method A', we have over a range of prediction points

$$\sum_c P_\alpha [X_k \leq X < X_{k+1}, C] \neq P_\alpha(X_k \leq X < X_{k+1})$$

since the cells overlap. Indeed

$$\left(\sum_c P_{\alpha} [X_k \leq X < X_{k+1}, C] - P_{\alpha} (X_k \leq X < X_{k+1}) \right) = \frac{\left(\sum_c K_{\alpha, k, c} \right) - K_{\alpha, k}}{\tilde{\beta}_{\alpha}}$$

However,

$$\sum_k P_{\alpha} (X_k \leq X < X_{k+1}, C) = P_{\alpha} (C)$$

Another approach, of course, for determining $X^*(t+\alpha)$ could be based on the maximum likelihood estimator, this is, on $\text{MAX}_{\alpha} P_{\alpha}(X(t+\alpha) \mid C)$ where C is the cell determined by $\{X(T)\}_0^{N-1}$, then $X^*(t+\alpha)$ can be taken as $X_{k_{m+1/2}}$ where k_m produces the maximum value of $K_{\alpha, k, c}$.

REFERENCE

Bose, A. G., "A Theory of Nonlinear Systems," M.I.T. Research
Laboratory of Electronics Technical Report 309, May 15, 1956

DOCUMENT CONTROL DATA - R & D

(Security classification of title, body of abstract and indexing annotation must be entered when the overall report is classified)

1. ORIGINATING ACTIVITY (Corporate author) The MITRE Corporation Bedford, Massachusetts		2a. REPORT SECURITY CLASSIFICATION UNCLASSIFIED	
		2b. GROUP	
3. REPORT TITLE NONLINEAR PREDICTION AND DIGITAL SIMULATION			
4. DESCRIPTIVE NOTES (Type of report and inclusive dates) N/A			
5. AUTHOR(S) (First name, middle initial, last name) J. F. A. Ormsby			
6. REPORT DATE June 1969		7a. TOTAL NO. OF PAGES 28	7b. NO. OF REFS 1
8a. CONTRACT OR GRANT NO. AF 19(628)-2390		9a. ORIGINATOR'S REPORT NUMBER(S) ESD-TR-69-184	
b. PROJECT NO. 6151		9b. OTHER REPORT NO(S) (Any other numbers that may be assigned this report) W-7550	
c.			
d.			
10. DISTRIBUTION STATEMENT This document has been approved for public release and sale; its distribution is unlimited.			
11. SUPPLEMENTARY NOTES N/A		12. SPONSORING MILITARY ACTIVITY Electronic Systems Division, Air Force Systems Command, L. G. Hanscom Field, Bedford, Massachusetts	
13. ABSTRACT An approach to digital nonlinear prediction is proposed and analyzed. The basic relations are developed. The nonlinear operator is obtained by quantization of the data. The model is developed in terms of occupancy of data cells in N-space. Extensions to increase occupancy and reduce error are formulated. Illustrative results are included. A comparison with linear techniques is made and over-all conclusions on error, quantization level, length of data required, time invariance, etc., are provided.			

14.	KEY WORDS	LINK A		LINK B		LINK C	
		ROLE	WT	ROLE	WT	ROLE	WT
	Nonlinear Prediction Digital Simulation Data Quantization Sampling Effects						