

AD 87 2201

RO-S-68-2

A TUTORIAL DERIVATION OF RECURSIVE
WEIGHTED LEAST SQUARES STATE-VECTOR
ESTIMATION THEORY (KALMAN THEORY)

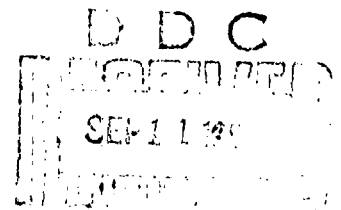
(Special Report)

by

James S. Pappas

AUGUST 1968

U. S. ARMY TEST AND EVALUATION COMMAND
ANALYSIS AND COMPUTATION DIRECTORATE
DEPUTY FOR NATIONAL RANGE OPERATIONS
WHITE SANDS MISSILE RANGE, NEW MEXICO



DISTRIBUTION OF THIS DOCUMENT IS UNLIMITED

REPORT NO.

A TUTORIAL DERIVATION OF RECURSIVE
WEIGHTED LEAST SQUARES STATE-VECTOR
ESTIMATION THEORY (KALMAN THEORY)

(Special Report)

by

James S. Pappas

Approved by



GUENTHER HINTZE
Director

August 1968

U. S. ARMY TEST AND EVALUATION COMMAND

ANALYSIS AND COMPUTATION DIRECTORATE
DEPUTY FOR NATIONAL RANGE OPERATIONS
WHITE SANDS MISSILE RANGE, NEW MEXICO

DISTRIBUTION OF THIS DOCUMENT IS UNLIMITED

CONTENTS

	Page No.
ABSTRACT -----	iv
INTRODUCTION -----	v
NOTATION -----	vi
I. DISCREET DETERMINISTIC TRAJECTORIES -----	1
a. One Parameter Family of Discreet Trajectories -----	1
b. Two Parameter Family of Discreet Trajectories -----	8
c. Three Parameter Family of Discreet Observation Trajectories -----	12
II. SCALAR ESTIMATION EQUATIONS -----	14
III. VECTOR ESTIMATION EQUATIONS -----	28
IV. EQUATION SUMMARY FOR COMPUTER APPLICATION -----	39
APPENDIX A. MATRIX TRACE PROPERTIES -----	41
APPENDIX B. GRADIENTS OF SCALARS WITH RESPECT TO MATRICES -----	44
APPENDIX C. MINIMIZATION -----	58
REFERENCES -----	60

ABSTRACT

This is one of a series of tutorial reports deriving the discrete recursive Kalman estimation equations. The series proceeds from unweighted to weighted least squares parameter estimation in a vector space setting to the stage-wise updating, to time varying parameter or state variable estimation. The derivations have been stage-wise tutorial in an attempt to make the theory accessible to the "non-specialist" in optimization theory.

INTRODUCTION

Modern state-vector recursive estimation theory (in the sense of Kalman) has been derived from many viewpoints: conditional probability, orthogonal projections, least squares, minimum variance, maximum likelihood, etc. The individual user of the theory may understand the derivation of the equations obtained by any one of the above theories or possibly by a number of approaches depending upon the individuals background.

This report is the third of a tutorial series deriving the estimation equations in a vector space weighted least squares setting. The essential areas for understanding the theory via this approach are:

1. Unweighted Least Squares Parameter-Vector Estimation and the Variance-of-the-Estimate Matrix.
2. Discrete Matrix Recursive Methods apply to (1) for Real Time (on line) Computer Processing.
3. Weighted Least Squares Parameter Estimation and Variance-of-the-Estimate Matrix for Correlated Noise.
4. Discrete Matrix Recursive Methods applied to (3) for Real Time Computer Processing.
5. Recursive Weighted Least Squares State-Vector Estimation Theory (Kalman Theory).

Items 1 and 2 are completed and published in Reference 4. Item 3 is completed and published in Reference 6. Item 4 is not completed yet. Item 5 is the contents of the current report.

This report is restricted to full-rank matrices so that the reader may understand the basic principles of the theory. Non-full rank matrices processed recursively require a knowledge and background in Pseudo Inverses, and many more sophisticated linear algebra concepts. The minimization criterion is derived by taking the partial derivative of a scalar with respect to a $p \times k$ weighting matrix- the classical gradient approach.

The derivations can be achieved completely algebraically in terms of orthogonal projections relating to spaces and sub-spaces. Many sophisticated papers exist in the literature using this approach. The algebraic approach may not be as readily accessible to the engineer trained in classical mathematics.

NOTATION

The notations used in the report is an effort to blend the notation of Friedman for inner-products and dyadic products with the current journal-literature on vector-spaces, psuedo-inverses, state-vectors, etc.

$X_{p \times k}$ Capital letters designate matrices of size p rows and k columns.

$\langle p \rangle$: when $p = 1$, the matrix is called a column vector, and we use Friedmans symbol to distinguish this matrix.

$\langle p \rangle \times$: when $k = 1$, the matrix is a row vector of dimension p .

$\langle p \rangle \times \langle p \rangle$ "inner-product" or scalar product of two vectors.

$\langle p \rangle \times \langle p \rangle$: "outer-product" or dyadic product of two vectors.

$X_{p \times k} = [\langle p \rangle_1, \dots, \langle p \rangle_k]$ Matrix X partitioned into a row k -tuple of column vectors from a p -space.

$X_{p \times k} = \begin{bmatrix} \langle 1 \rangle \times \\ \langle k \rangle \times \\ \vdots \\ \langle p \rangle \times \end{bmatrix}$ Matrix X partitioned into a p -column tuple of row vectors from a k -space

x small x is a scalar

x^i scalar from a column vector

x_j scalar from a row vector

Scalar here is a "real field" element.

Consider the system of two vector equations

$$\langle x(k+1) \rangle = \Phi(k+1, k) \langle x(k) \rangle + f(k) + u(k)$$

and

$$\langle z(k) \rangle = H_{m \times p}(k) \langle x(k) \rangle + v(k)$$

where:

$x(k)$, $x(k+1)$ are p -dimensional column vectors describing the states at stage k and stage $k+1$.

$\Phi(k+1, k)$ is a $p \times p$ state transition matrix.

$f(k)$ is a p -dimensional deterministic forcing vector for which we can write a vector function.

$u(k)$ is a p -dimensional uncertainty or noise vector, it is the composite of the random noises and the variables we fail to model.

$z(k)$ is the m -dimensional observation vector, m is less than or equal to p .

$H(k)$
 $m \times p$ is the known matrix describing how the state vector is functionally related to the observation vector (if the instruments were noise free).

$v(k)$ is an m -dimensional additive instrument noise vector.

1. ONE PARAMETER FAMILY OF DISCRETE TRAJECTORIES

This section considers the one parameter family of i trajectories generated by i different values of the initial condition parameter $x_1(1)$. Each of the i trajectories consists of a sequence of k different points or stages.

We consider

- (1) The homogeneous case
- (2) The deterministically forced case.

Homogeneous System

Consider the system of scalar equations

$$x(k+1) = \phi(k+1, k)x(k) + f(k) + u(k) \quad (1)$$

$$z(k) = h(k)x(k) + v(k) \quad (2)$$

where the following constraints apply

$$f(k) = u(k) = v(k) = 0 \quad (3)$$

$$h(k) = h_0 \text{ a constant} \quad (4)$$

$$\phi(k+1, k) = \phi_0 \text{ a constant} \quad (5)$$

then equation (1) and (2) become

$$x(k+1) = \phi_0 x(k) \quad (6)$$

$$z(k) = h_0 x(k) \quad (7)$$

If further, we consider the measurement to be exact

$$z(k) = x(k) \quad (8)$$

or

$$h_0 = 1 \quad (9)$$

then we need only consider

$$x(k+1) = \phi_0 x(k) \quad (10)$$

Case 1. The most elementary case occurs when

$$\phi_0 = 1, \quad (11)$$

then

$$x(k+1) = x(k) \quad (12)$$

and for values of k

$$x(2) = x(1) \quad (13)$$

$$x(3) = x(2) = x(1) \quad (14)$$

$$x(k) = x(1) \quad (15)$$

hence we can plot

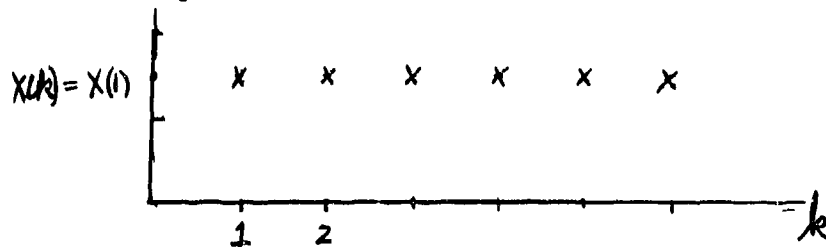


Fig. (1)

If we vary the initial condition we obtain a sequence of trajectories

$$x_i(k) = x_i(1) \quad (16)$$

for

$$x_1(1), x_2(1) \dots x_i(1), x_{i+1}(1) \dots$$

as shown in Fig. (2)

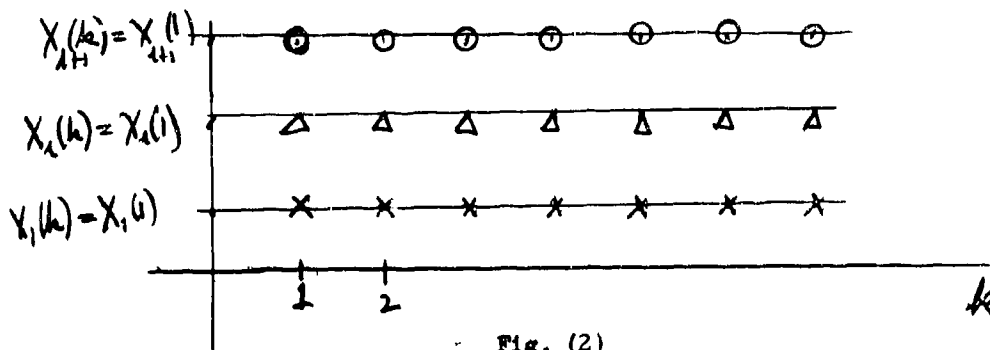


Fig. (2)

For any constant ϕ_0 , we obtain

$$\begin{aligned} x(2) &= \phi_0 x(1) \\ x(3) &= \phi_0 x(2) = \phi_0(\phi_0 x(1)) = \phi_0^2 x(1) \\ &\vdots \\ x(k) &= \phi_0^{k-1} x(1). \end{aligned} \tag{17}$$

If ϕ_0 is less than one, the trajectory of Equation (), for a given initial condition, decreases, for example

$$\begin{aligned} \phi_0 &= 1/2 \\ x(2) &= \frac{1}{2} x(1) \\ x(3) &= \left(\frac{1}{2}\right)^2 x(1) = \frac{1}{4} x(1) \\ x(4) &= \frac{1}{8} x(1) \\ &\vdots \\ x(k) &= \left(\frac{1}{2}\right)^{k-1} x(1) \end{aligned} \tag{18}$$

The sequence of k vectors are

$$\begin{pmatrix} 1 \\ x(1) \end{pmatrix} \quad \begin{pmatrix} 2 \\ \frac{1}{2} x(1) \end{pmatrix} \quad \begin{pmatrix} 3 \\ \frac{1}{4} x(1) \end{pmatrix} \quad \dots \quad \begin{pmatrix} k \\ \left(\frac{1}{2}\right)^{k-1} x(1) \end{pmatrix} \tag{19}$$

which plot as shown in Fig. (3)

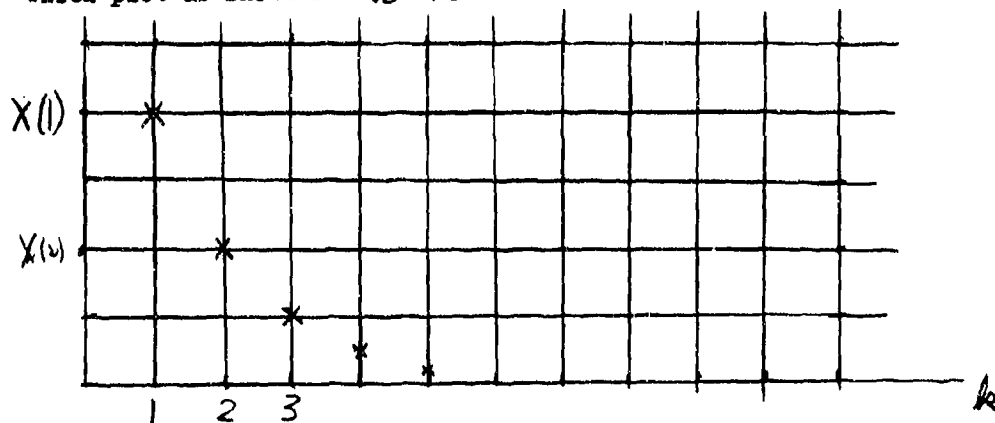


Fig. (3) Contracting or Decaying Sequence $\phi_0 = 1/2$

The first noisy observation is

$$z_{jn}(1) = h(1) x_j(1) + v_n(1)$$

Using (3) and (6) in (4)

$$\tilde{z}_{jn}(1, 0) = h(1)[x_j(1) - \hat{x}_j(1, 0)] + v_n(1). \quad (7)$$

Using (7) in (5)

$$\begin{aligned} \hat{x}_{jn}(1, 1) &= \hat{x}_j(1, 0) + w(1)h(1)[x_j(1) - \hat{x}_j(1, 0)] + w(1)v_n(1) \\ &= [1 - w(1)h(1)] \hat{x}_j(1, 0) + w(1)h(1)x_j(1) + w(1)v_n(1). \end{aligned} \quad (8)$$

We now define two new errors as differences of the previous variables.

$$x_j(1) - \hat{x}_j(1, 0) = \tilde{x}_j(1, 0) \quad (9)$$

and

$$x_j(1) - \hat{x}_{jn}(1, 1) = \tilde{x}_{jn}(1, 1). \quad (10)$$

The variable $x_j(1)$ is the actual (but unknown) process state variable that exists on trajectory j . The variable $\hat{x}_{jn}(1, 1)$ is the estimate of $x_j(1)$, having used the first noisy observation.

Hence $\tilde{x}_{jn}(1, 1)$ is the error in our estimate of the state at stage 1 based on one measurement.

From equation (3) we see that if our initial guess $\hat{x}_j(1, 0)$ is wrong, then $\tilde{z}_j(1, 0)$ will also be wrong and after our observation $z(1)$ comes in and our computed error $\tilde{z}_j(1, 0)$ of equation (4) is large, then by equation (5) we will have a correction term which is the product $(w(1)\tilde{z}_j(1, 1))$ of a large term times the weight value $w(1)$.

Continuing with the derivation of the weights, we obtain the error term by equation (8) in equation (11) as

$$\begin{aligned} \tilde{x}_{jn}(1, 1) &= x_j(1) - \hat{x}_{jn}(1, 1) = x_j(1) - w(1)h(1)x_j(1) \\ &\quad - [(1-w(1)h(1)) \hat{x}_j(1, 0) - w(1)v_n(1)]. \end{aligned} \quad (11)$$

or

$$\tilde{x}_{jn}(1, 1) = [1-w(1)h(1)][x_j(1) - \hat{x}_j(1, 0)] - w(1)v_n(1). \quad (12)$$

Using equation (9) in equation (12)

$$\tilde{x}_{jn}(1,1) = [1 - w(1)h(1)] \tilde{x}_j(1,0) - w(1) v_n(1). \quad (13)$$

If we square the error of equation (13)

$$\begin{aligned} \tilde{x}_{jn}^2(1,1) &= [1-w(1)h(1)]^2 \tilde{x}_j^2(1,0) - 2w(1)v_n(1)[1-w(1)h(1)] \hat{x}_j(1,0) \\ &\quad + w^2(1)v_n^2(1). \end{aligned} \quad (14)$$

If we take the partial derivative of equation (14) with respect to $w(1)$ and equate to zero we obtain

$$\begin{aligned} \frac{\partial \tilde{x}_{jn}^2(1,1)}{\partial w(1)} &= 2[1 - w(1)h(1)] (-h(1)) \tilde{x}_j^2(1,0) \\ &\quad - 2v_n(1)[1 - w(1)h(1)] \hat{x}_j(1,0) \\ &\quad + 2w(1) v_n^2(1) = 0 \end{aligned} \quad (15)$$

Taking the expected value over all experiments in which j and n are varying we obtain

$$\begin{aligned} E \left\{ \frac{\partial \tilde{x}_{jn}^2(1,1)}{\partial w(1)} \right\} &= -[1 - w(1)h(1)] E \{ \tilde{x}_j^2(1,0) \} h(1) \\ &\quad + w(1) E \{ v_n^2(1) \} = 0 \end{aligned} \quad (16)$$

Define the variances

$$E \{ \tilde{x}_j^2(1,0) \} = \sigma_{\tilde{x}\tilde{x}}(1,0) = p(1,0) \quad (17)$$

$$E \{ v_n^2(1) \} = \sigma_{vv}(1) \quad (18)$$

We follow a number of current Kalman papers and papers about Kalman Estimation and designate the variance in state as $p(1,0)$.

The cross term of equation (15) under the expectation operator is zero when the following conditions hold,

$$\begin{aligned} E \{ [1 - w(1)h(1)] v_n(2) \hat{x}_j(1,0) \} \\ = [1 - w(1)h(1)] E \{ v_n(1) \hat{x}_j(1,0) \} \end{aligned} \quad (19)$$

We now assume that our initial guess of $\hat{x}(1,0)$ is independent of $v_n(1)$ as n ranges over all allowable values, that is

$$E \{ v_n(1) \hat{x}_j(1,0) \} = 0. \quad (20)$$

Under these assumptions equation (16) and (18) become

$$0 = -[1 - w(1)h(1)]h(1) p(1, 0) + w(1)\sigma_{vv}(1) \quad (21)$$

or

$$-h(1)p(1, 0) + w(1)h^2(1) p(1, 0) + w(1)\sigma_{vv}(1) = 0 \quad (22)$$

or

$$w(1)[h^2(1) p(1, 0) + \sigma_{vv}(1)] = h(1) p(1, 0) \quad (23)$$

$$w(1) = h(1) p(1, 0)[h^2(1) p(1, 0) + \sigma_{vv}(1)]^{-1} \quad (24)$$

The value of $p(1, 0)$ which by equation (17) is

$$p(1, 0) = E\left\{[x_j(1) - \hat{x}_j(1, 0)]^2\right\} = \text{guess} \quad (25)$$

can be obtained by guess or by a "learning process", or by a-priori knowledge about the system.

The variance of the estimate of state at stage 1 based on the observation at stage one can also be obtained by equation (14) as

$$E\{\hat{x}_{jn}^2(1, 1)\} = [1 - w(1)h(1)]^2 E\{\hat{x}_j^2(1, 0)\} + w^2(1) E\{v_n^2(1)\} \quad (26)$$

when the cross terms are zero, that is

$$E\{v_n(1) \hat{x}_j(1, 0)\} = 2w(1)[1 - w(1)h(1)] = 0. \quad (27)$$

The variance terms of equation () will be denoted as

$$\sigma_{\hat{x}\hat{x}}(1, 1) = E\{\hat{x}_{jn}^2(1, 1)\} \equiv p(1, 1) \quad (28)$$

$$E\{v_n^2(1)\} = \sigma_{vv}(1) \quad (29)$$

The $p(1, 1)$ designation is in keeping with the now near classical notation.

Expanding the terms of equation (26)

$$\begin{aligned} p(1, 1) &= [1 - w(1)h(1)]^2 p(1, 0) + w^2(1)\sigma_{vv}(1) \\ &= [1 - 2w(1)h(1) + w^2(1)h^2(1)] p(1, 0) + w^2(1)\sigma_{vv}(1) \\ &= p(1, 0) - 2w(1)h(1) p(1, 0) \\ &\quad + w^2(1)[h^2(1) p(1, 0) + \sigma_{vv}(1)]. \end{aligned} \quad (30)$$

Consider the last term in equation (30) and use equation (24) for $w^2(1)$, then

$$\begin{aligned}
 & w^2(1)[h^2(1) p(1, 0) + \sigma_{vv}(1)] \\
 &= h^2(1) p^2(1, 0)[h^2(1) p(1, 0) + \sigma_{vv}(1)]^{-2} [h^2(1) p(1, 0) + \sigma_{vv}(1)] \\
 &= h^2(1) p^2(1, 0)[h^2(1) p(1, 0) + \sigma_{vv}(1)]^{-1} \\
 &= h(1) p(1, 0) \left\{ h(1) p(1, 0) [h^2(1) p(1, 0) + \sigma_{vv}(1)]^{-1} \right.
 \end{aligned} \tag{31}$$

Replacing the bracket term of equation (31) by $w(1)$ in equation (24) we obtain

$$w^2(1)[h^2(1) p(1, 0) + \sigma_{vv}(1)] = h(1) p(1, 0) w(1) \tag{32}$$

Using (32) in equation (30) we obtain

$$\begin{aligned}
 p(1, 1) &= p(1, 0) - 2w(1) h(1) p(1, 0) + h(1) p(1, 0) w(1) \\
 &= p(1, 0) - w(1) h(1) p(1, 0).
 \end{aligned} \tag{33}$$

$$p(1, 1) = p(1, 0)[1 - w(1)h(1)]$$

also using equation (24) in equation (33) we can write $p(1, 1)$ as

$$p(1, 1) = p(1, 0) - p(1, 0)h(1)[h(1)p(1, 0)h(1) + \sigma_{vv}(1)]^{-1}p(1, 0)h(1) \tag{34}$$

Equation (34) is the scalar variance (one dimensional uncertainty ellipsoid) of the estimate of the state using all of the current observations.

We can now predict what the next state variable value should be by propagating forward with the known dynamics as

$$\hat{x}(2, 1) = \Phi(2, 1) \hat{x}(1, 1) + f(1). \tag{35}$$

The prediction of the next observation is

$$\hat{z}(2, 1) = h(2) \hat{x}(2, 1). \tag{36}$$

The ellipsoid of uncertainty of the predicted state is

$$p(2, 1) = E\{\tilde{x}^2(2, 1)\} \tag{37}$$

where

$$\tilde{x}(2, 1) = x(2) - \hat{x}(2, 1) \tag{38}$$

Using equation (1) and equation (35) in equation (38)

$$\hat{x}(2,1) = \phi(2,1) \tilde{x}(1,1) + u(1) \quad (39)$$

$$\tilde{x}^2(2,1) = \phi(2,1) \tilde{x}^2(1,1) \phi(2,1) + 2\phi(2,1) \tilde{x}(1,1) u(1) + u^2(1) \quad (40)$$

The expected value over all experiments of equation (40) is

$$E\{\tilde{x}^2(2,1)\} = p(2,1) = \phi(2,1) p(1,1) \phi(2,1) + \sigma_{uu}(1) \quad (41)$$

where

$$\sigma_{uu}(1) = E_n\{u_n^2(1)\} \quad (42)$$

Cross term statistical independence assumptions are

$$E\{\tilde{x}(1,1) u(1)\} = 0. \quad (43)$$

etc.

In conclusion, at stage one, we do:

Given

$$\sigma_{vv}(1)$$

guess or estimate

$$p(1,0) \text{ and } \hat{x}(1,0)$$

compute

$$\hat{z}(1,0) = h(1) \hat{x}(1,0) \quad (44)$$

Measure and compute error

$$\tilde{z}(1,0) = z(1) - \hat{z}(1,0) \quad (45)$$

Compute weight $w(1)$ equation ()

$$w(1) = h(1) p(1,0) [h(1) p(1,0) h(1) + \sigma_{vv}(1)]^{-1}. \quad (46)$$

Update or correct state-estimate

$$\hat{x}(1,1) = \hat{x}(1,0) + w(1) \tilde{z}(1,0) \quad (47)$$

Compute ellipsoid of uncertainty of state estimation

$$p(1,1) = p(1,0) [1 - w(1) h(1)] \quad (48)$$

Predict or propagate one stage into future with known dynamics the following three variables:

Predict next stage.

$$\hat{x}(2, 1) = \phi(2, 1) \hat{x}(1, 1) + f(1) \quad (49)$$

Predict next observation.

$$\hat{z}(2, 1) = h(2) \hat{x}(2, 1) \quad (50)$$

Predict ellipsoid of uncertainty in state

$$p(2, 1) = \phi(2, 1) p(1, 1) \phi(2, 1) + \sigma_{uu}(1). \quad (51)$$

Wait for stage $k = 2$ and second observation to come in.

Stage $k = 2$.

The second observation is taken at stage 2 and is

$$z_{jn}(2) = h(2) x_j(2) + v_n(2) \quad (52)$$

where, as before, $z_{jn}(2)$ is known, but $x_j(2)$ and $v_n(2)$ are unknown.

We compute the error signal

$$\tilde{z}_{jn}(2, 1) = z_{jn}(2) - \hat{z}_{jn}(2, 1), \quad (53)$$

The sub-scripts j, n can now be dropped for simplicity of notation, where it is henceforth understood that $\hat{x}(2, 1)$ for example means the best estimate of the true j th trajectory based on the sequence of instrument noises $\underset{n}{\lt} v$.

Using equation (50) and equation (52) in equation (53)

$$\tilde{z}(2, 1) = h(2)[x(2) - \hat{x}(2, 1)] + v(2) \quad (54)$$

we now have computed the error in the observation, hence we can correct our estimate of what the state variable should be based on the use of the second observation, that is

$$\hat{x}(2, 2) = \hat{x}(2, 1) + w(2) \tilde{z}(2, 1). \quad (55)$$

As before, we now seek the "feed-back" or correction weight $w(2)$ at stage 2. Before deriving the $w(2)$ expression, use (54) in (55)

$$\hat{x}(2, 2) = \hat{x}(2, 1) + w(2)h(2)[x(2) - \hat{x}(2, 1)] + w(2) v(2) \quad (56)$$

$$\hat{x}(2, 2) = [1 - w(2)h(2)] \hat{x}(2, 1) + w(2) h(2) x(2) + w(2) v(2). \quad (57)$$

From the error in state estimation at stage 2 using the second observation as before

$$\hat{x}(2,2) = x(2) - \hat{x}(2,2) = [1-w(2)h(2)]\tilde{x}(2,1) + w(2)v(2) \quad (58)$$

Multiply equation (58) by itself to obtain the square of the error term as

$$\begin{aligned} \hat{x}^2(2,2) &= [1-w(2)h(2)]^2 \tilde{x}^2(2,1) \\ &+ 2w(2)v(2)[1-w(2)h(2)]\tilde{x}(2,1) \\ &+ w^2(2)v^2(2). \end{aligned} \quad (59)$$

Observe that the error and the square of the error in equation (59) is a function of $w(2)$.

In order to select a $w(2)$ which will minimize the square of the error we take the partial derivative of equation (59) with respect to $w(2)$ and equate this "gradient" term to zero, hence

$$\begin{aligned} \frac{\partial \hat{x}^2(2,2)}{\partial w(2)} &= -2h(2)[1-w(2)h(2)]\tilde{x}^2(2,1) \\ &+ 2v(2)[1-w(2)h(2)]\tilde{x}(2,1) \\ &+ 2w(2)v(2)(-h(2))\tilde{x}(2,1) \\ &+ 2w(2)v^2(2) = 0 \end{aligned} \quad (60)$$

Since we want $w(2)$ to be the same regardless of how many times we repeat the experiment; or, from a "single-test" stand-point, $w(2)$ should be selected to hold regardless of the unknown sequence driving the two variables at this stage $u_j(2)$ and $v_n(2)$. Consequently we take the expected value over all admissible vectors $\langle u \rangle$ and $\langle v \rangle$ and obtain

$$E \left\{ \frac{\partial \hat{x}^2(2,2)}{\partial w(2)} \right\} = [1-w(2)h(2)](-h(2))p(2,1) + w(2)\sigma_{vv}(2) = 0. \quad (61)$$

The assumption of statistical independence of the variables $\tilde{x}(2,1)$ and $v(2)$ is again assumed

$$E \{ \tilde{x}(2,1) v(2) \} = 0 \quad (62)$$

also the assumption about $\langle v \rangle$

$$\begin{aligned} E_n \left\{ \begin{pmatrix} v_n(1) \\ v_n(2) \end{pmatrix} [v_n(1), v_n(2)] \right\} \\ = E_n \left\{ \begin{bmatrix} u_n(1) v_n(1) & v_n(1) v_n(2) \\ v_n(2) v_n(1) & v_n(2) v_n(2) \end{bmatrix} \right\} = \begin{bmatrix} \sigma_{vv}(1) & 0 \\ 0 & \sigma_{vv}(2) \end{bmatrix} \end{aligned} \quad (63)$$

The assumptions of equation (63) can be relaxed but then one needs a deeper knowledge of multi-linear algebras, matrix-packaging and partitioning, matrix psuedo-inverse, etc. Hence such "highly-colored" or correlated noise cases will not be discussed in this paper. The majority of published papers on Kalman theory make the implied Gaussian assumptions implied by equation (63) for arbitrary large k. Computer storage problems for correlated noise for large k are also a problem.

Solving equation (60) for $w(2)$

$$w(2)[h(2) p(2, 1) h(2) + \sigma_{vv}(2)] = h(2) p(2, 1) \quad (64)$$

or

$$w(2) = h(2) p(2, 1) [h(2) p(2, 1) h(2) + \sigma_{vv}(2)]^{-1} \quad (65)$$

The second weight can be computed since $h(2)$, $\sigma_{vv}(2)$ are assumed known and $p(2, 1)$ was computed at stage 1.

We can now derive the expression for $p(2, 2)$, using equation (65) and taking the expected value over all experiments

$$p(2, 2) = E \{ \tilde{x}^2(2, 2) \} = [1 - w(2)h(2)]^2 p(2, 1) + w^2(2)\sigma_{vv}(2), \quad (66)$$

with assumptions

$$E \{ v(2) \tilde{x}(2, 1) \} = 0. \quad (67)$$

Expanding equation (66)

$$p(2, 2) = [1 - 2w(2) h(2) + w^2(2) h^2(2)] p(2, 1) + w^2(2)\sigma_{vv}(2) \quad (68)$$

$$p(2, 2) = p(2, 1) - 2w(2) h(2) p(2, 1) + w^2(2) [h(2) p(2, 1) h(2) + \sigma_{vv}(2)] \quad (69)$$

Consider the last term of equation (69) using equation (65) for $w(2)$

$$\begin{aligned} & w^2(2) [h(2) p(2, 1) h(2) + \sigma_{vv}(2)] \\ &= h^2(2) p^2(2, 1) [h^2(2) p(2, 1) + \sigma_{vv}(2)]^2 [h^2(2)p(2,1) + \sigma_{vv}(2)] \\ &= h^2(2) p^2(2, 1) [h^2(2) p(2, 1) + \sigma_{vv}(2)]^{-1} \end{aligned} \quad (70)$$

and by equation (65) for $w(2)$ we obtain

$$w^2(2)[h(2) p(2, 1) h(2) + \bar{\sigma}_{vv}(2)] = h(2) p(2, 1) w(2) \quad (71)$$

Using equation (71) in equation (69)

$$p(2,2) = p(2, 1) - 2w(2) h(2) p(2, 1) + h(2) p(2, 1) w(2) \quad (72)$$

or

$$p(2,2) = p(2, 1) - p(2,1) h(2) w(2) \quad (73)$$

$$p(2,2) = p(2, 1) [1 - w(2) h(2)] \quad (74)$$

Equation (73) can also be written as by equation (65) in equation (73)

$$p(2, 2) = p(2, 1) - p(2, 1) [h^2(2)p(2,1) + \bar{\sigma}_{vv}(2)]^{-1} h(2)p(2, 1) \quad (75)$$

We can now predict at stage $k = 3$ the state

$$\hat{x}(3, 2) = \phi(3, 2) \hat{x}(2,2) + f(2) \quad (76)$$

and the observation

$$\hat{z}(3, 2) = h(3) \hat{x}(3, 2) \quad (77)$$

and the state-variance term

$$p(3, 2) = E \{\tilde{x}^2(3, 2)\}. \quad (78)$$

By equation (1) and equation (76)

$$x(3) = \phi(3,2) x(2) + f(2) + u(2) \quad (79)$$

$$\hat{x}(3,2) = \phi(3, 2) \hat{x}(2, 2) + f(2) \quad (80)$$

and subtracting

$$x(3) - \hat{x}(3, 2) = \tilde{x}(3, 2) = \phi(3, 2) \tilde{x}(2, 2) + u(2) \quad (81)$$

Squaring equation (81)

$$\begin{aligned} \tilde{x}^2(3, 2) &= \phi(3, 2) \tilde{x}^2(2, 2) \phi(3, 2) \\ &+ 2\phi(3, 2) \tilde{x}(2, 2) u(2) \\ &+ u^2(2). \end{aligned} \quad (82)$$

Taking the expected value

$$E\{\hat{x}^2(3, 2)\} = \phi(3, 2) E\{\hat{x}^2(2, 2)\} \phi(3, 2) + E\{u^2(2)\} \quad (83)$$

or

$$p(3, 2) = \phi(3, 2) p(2, 2) \phi(3, 2) + \sigma_{uu}(2) \quad (84)$$

Summarizing the steps at stage $k = 2$, then we:

measure and compute error

$$\tilde{z}(2, 1) = z(2) - \hat{z}(2, 1) \quad (85)$$

compute weight $w(2)$ by equation (65)

$$w(2) = h(2) p(2, 1) [h(2) p(2, 1) h(2) + \sigma_{vv}(2)]^{-1} \quad (86)$$

compute corrected estimate of state

$$\hat{x}(2, 2) = \hat{x}(2, 1) + w(2) \tilde{z}(2, 1) \quad (87)$$

compute variance of state equation (74)

$$p(2, 2) = p(2, 1) - w(2) h(2) p(2, 1) \quad (88)$$

Predict (update) state via dynamics equation (80)

$$\hat{x}(3, 2) = \phi(3, 2) \hat{x}(2, 2) + f(2) \quad (89)$$

Predict observation at next stage

$$\hat{z}(3, 2) = h(3) \hat{x}(3, 2) \quad (90)$$

Predict next stage state-variance

$$p(3, 2) = \phi(3, 2) p(2, 2) \phi(3, 2) + \sigma_{uu}(2) \quad (91)$$

wait for next stage or third observation to arrive.

Stage k .

The derivations of the equations will not be repeated for stage k , the relations will be based on the mathematical process of reasoning by analogy. The treatment of the multi-variable or vector case will derive the relations at stage k , but will not develop the stage-wise logic at $k = 1$ and $k = 2$.

We have available from previous stage predictions

$$\begin{aligned} \hat{x}(k, k-1) \\ \hat{z}(k, k-1) \\ p(k, k-1) \end{aligned}$$

and stored $\sigma_{vv}(k)$, $\sigma_{uu}(k)$

Measure and compute error

$$\tilde{z}(k, k-1) = z(k) - \hat{z}(k, k-1) \quad (92)$$

Compute $w(k)$

$$w(k) = h(k) p(k, k-1) [h^2(k) p(k, k-1) + \sigma_{vv}(k)]^{-1} \quad (93)$$

Compute corrected state estimate

$$\hat{x}(k, k) = \hat{x}(k, k-1) + w(k) \tilde{z}(k, k-1) \quad (94)$$

Compute variance

$$p(k, k) = p(k, k-1) [1 - w(k) h(k)] \quad (95)$$

Predict next state

$$\hat{x}(k+1, k) = \phi(k+1, k) \hat{x}(k, k) + f(k) \quad (96)$$

Predict next observation

$$\hat{z}(k+1, k) = h(k+1) \hat{x}(k+1, k) \quad (97)$$

Predict variance of state

$$p(k+1, k) = \phi(k+1, k) p(k, k) \phi(k+1, k) + \sigma_{uu}(k). \quad (98)$$

We define the noise variances in the notation of the many Kalman oriented papers, that is

$$\sigma_{uu}(k) = q(k) \quad (99)$$

$$\sigma_{vv}(k) = r(k) \quad (100)$$

The three familiar equations can be written as

$$\begin{aligned} \hat{x}(k+1, k+1) &= \phi(k+1, k) \hat{x}(k, k) \\ &+ p(k+1, k) h(k+1) [h(k+1) p(k+1, k) h(k+1) + r(k)]^{-1} \\ &[z(k+1) - h(k+1) \phi(k+1, k) \hat{x}(k, k)] \end{aligned} \quad (101)$$

$$p(k+1, k) = \phi(k+1, k) p(k, k) \phi(k+1, k) + q(k) \quad (102)$$

$$p(k+1, k+1) = p(k+1, k) - p(k+1, k) h(k+1) \quad (103)$$

$$[h(k+1) p(k+1, k) h(k+1) + r(k)] h(k+1) p(k+1, k)$$

We shall now define in words the meanings at stage k of the variables and rewrite the equations using the distinguishing j and n .

$$x_j(k+1) = \phi(k+1, k) x_j(k) + f(k) + u_j(k) \quad (104)$$

$$z_{jn}(k) = h(k) x_j(k) + v_n(k) \quad (105)$$

$$\begin{aligned} \hat{x}_{jn}(k+1, k+1) &= \phi(k+1, k) \hat{x}_{jn}(k, k) \\ &\quad + p(k+1, k) h(k) [h^2(k+1) p(k+1) + r(k)]^{-1} \\ &\quad [z_{jn}(k+1) - h(k+1) \phi(k+1, k) \hat{x}_{jn}(k, k)] \end{aligned} \quad (106)$$

$$p(k+1, k) = \phi(k+1, k) p(k, k) \phi(k+1, k) \quad (107)$$

$$p(k+1, k+1) = p(k+1, k) - p(k+1, k) h(k+1) \quad (108)$$

$$[h^2(k+1) p(k+1, k) + r(k)]^{-1} h(k+1) p(k+1, k)$$

$x(k) = x_j(k)$ is the true (unknown) value of the process state at stage k as a result of the unknowns $u_j(1), \dots, u_j(k)$ forcing the system.

$z_j(k) = z_{jn}(k)$ is measurement of the true noise process $x_j(k)$ with additive unknown measurement noise $v_n(k)$.

$\hat{x}_{jn}(k, k) = \hat{x}(k, k)$ is the best estimate of the state at stage k of the j th trajectory based on past observations up to stage k , that is recursively we have used noisy

$$z_{jn}(1), z_{jn}(2) \dots z_{jn}(k)$$

made noisy by $v_n(1), \dots, v_n(k)$.

$\hat{x}_{jn}(k+1, k) = \hat{x}(k+1, k)$ is the best estimate of the state of the j th trajectory at stage $k+1$, based on observations only up to k . Also interpreted as the prediction of the state at next stage $k+1$, based on current stage k and past measurements.

III. VECTOR ESTIMATION EQUATIONS

This section derives the Kalman Estimation Equations for the multivariable case using matrix analysis methods. The derivation techniques are the same as used in the previous scalar case. The essential difference lies in the minimization methods. The variance of the estimate in state for the scalar case is a scalar valued function of a scalar argument $w(k)$. For the vector case, the trace of the variance matrix of the estimate of state is likewise a scalar-valued function, but a function of a matrix of p rows and m columns $W(k)$. The minimization of a scalar-valued function with respect to this matrix.

One can arrive at the equations via strictly algebraic concepts of orthogonal projection matrices etc., in which one does not have to enter into discussions of partial derivatives, continuity of continuous variables and gradients. Since the majority of expected readers are assumed to be more familiar with the least-squares criterion via gradients, this report will stick strictly with this method.

The general linearized vector equations are

$$x(k+1) \begin{matrix} \text{p} \\ \text{pxg} \end{matrix} = \bar{f}(k+1, k) x(k) \begin{matrix} \text{p} \\ \text{pxg} \end{matrix} + B(k) f(k) \begin{matrix} \text{g} \\ \text{pxg} \end{matrix} + N(k) u(k) \begin{matrix} \text{g} \\ \text{pxq} \end{matrix} \quad (1)$$

$$z(k) \begin{matrix} \text{m} \\ \text{mxp} \end{matrix} = H(k) x(k) \begin{matrix} \text{p} \\ \text{pxg} \end{matrix} + v(k) \begin{matrix} \text{m} \\ \text{mxp} \end{matrix} \quad (2)$$

The deterministic k -varying vector $f(k) \begin{matrix} \text{g} \\ \text{pxg} \end{matrix}$ is of dimension g less than or equal to p , and gets distributed or cross-coupled into all p state variables $x(k+1) \begin{matrix} \text{p} \\ \text{pxg} \end{matrix}$ via the functional relations of $B(k)$.

The same statements apply to the noise input vector $u(k) \begin{matrix} \text{g} \\ \text{pxq} \end{matrix}$.

The reader should keep in mind the families of trajectories accurately described by the j and n indices, that is

$$x(k+1) \begin{matrix} \text{p} \\ j \end{matrix} = \bar{f}(k+1, k) x(k) \begin{matrix} \text{p} \\ j \end{matrix} + B(k) f(k) \begin{matrix} \text{g} \\ j \end{matrix} + N(k) u(k) \begin{matrix} \text{g} \\ j \end{matrix} \quad (3)$$

$$z_{jn}(k) \begin{matrix} \text{m} \\ j \end{matrix} = H(k) x(k) \begin{matrix} \text{p} \\ j \end{matrix} + v(k) \begin{matrix} \text{m} \\ n \end{matrix} \quad (4)$$

As before, the accurate descriptions designated by j and n will be dropped for simplicity of representation.

The equations are developed as a "recursive process" or an "on-line" processor; that is, as the observations "role in" the mechanized computer-estimator utilizes the data, and discards it or stores it on tape or what have you. All past data is sequentially accumulated in the "memory of the

III. VECTOR ESTIMATION EQUATIONS

This section derives the Kalman Estimation Equations for the multivariable case using matrix analysis methods. The derivation techniques are the same as used in the previous scalar case. The essential difference lies in the minimization methods. The variance of the estimate in state for the scalar case is a scalar valued function of a scalar argument $v(k)$. For the vector case, the trace of the variance matrix of the estimate of state is likewise a scalar-valued function, but a function of a matrix of p rows and p columns $W(k)$. The minimization of a scalar-valued function with respect to this matrix.

One can arrive at the equations via strictly algebraic concepts of orthogonal projection matrices etc., in which one does not have to enter into discussions of partial derivatives, continuity of continuous variables and gradients. Since the majority of expected readers are assumed to be more familiar with the least-squares criterion via gradients, this report will stick strictly with this method.

The general linearized vector equations are

$$x(k+1) \begin{matrix} \text{p} \\ \text{pxg} \end{matrix} = \bar{f}(k+1, k) x(k) \begin{matrix} \text{p} \\ \text{pxg} \end{matrix} + B(k) f(k) \begin{matrix} \text{g} \\ \text{pxg} \end{matrix} + N(k) u(k) \begin{matrix} \text{q} \\ \text{pxq} \end{matrix} \quad (1)$$

$$z(k) \begin{matrix} \text{m} \\ \text{mxp} \end{matrix} = H(k) x(k) \begin{matrix} \text{p} \\ \text{mxp} \end{matrix} + v(k) \begin{matrix} \text{m} \\ \text{mxp} \end{matrix} \quad (2)$$

The deterministic k -varying vector $f(k) \begin{matrix} \text{g} \\ \text{pxg} \end{matrix}$ is of dimension g less than or equal to p , and gets distributed or cross-coupled into all p state variables $x(k+1) \begin{matrix} \text{p} \\ \text{pxg} \end{matrix}$ via the functional relations of $B(k)$.

The same statements apply to the noise input vector $u(k) \begin{matrix} \text{q} \\ \text{pxq} \end{matrix}$.

The reader should keep in mind the families of trajectories accurately described by the j and n indices, that is

$$x(k+1) \begin{matrix} \text{p} \\ j \end{matrix} = \bar{f}(k+1, k) x(k) \begin{matrix} \text{p} \\ j \end{matrix} + B(k) f(k) \begin{matrix} \text{g} \\ j \end{matrix} + N(k) u(k) \begin{matrix} \text{q} \\ j \end{matrix} \quad (3)$$

$$z_{jn}(k) \begin{matrix} \text{m} \\ j \end{matrix} = H(k) x(k) \begin{matrix} \text{p} \\ j \end{matrix} + v(k) \begin{matrix} \text{m} \\ n \end{matrix} \quad (4)$$

As before, the accurate descriptions designated by j and n will be dropped for simplicity of representation.

The equations are developed as a "recursive process" or an "on-line" processor; that is, as the observations "role in" the mechanized computer-estimator utilizes the data, and discards it or stores it on tape or what have you. All past data is sequentially accumulated in the "memory of the

math-ware" via up dated estimates and variance matrices etc.

Stage k.

Suppose we are at stage k and have computed during stage k-1, the following

$$\begin{aligned} \hat{x}(k-1, k-1) \\ P(k-1, k-1) \end{aligned}$$

and predicted via dynamics

$$\hat{x}(k, k-1) = \Phi(k, k-1) \hat{x}(k-1, k-1) + B(k-1) f(k-1) \quad (5)$$

$$\hat{z}(k, k-1) = H(k) \hat{x}(k, k-1) \quad (6)$$

$$\hat{z}(k, k-1) = H(k) \Phi(k, k-1) \hat{x}(k-1, k-1) + H(k) B(k-1) f(k-1) \quad (7)$$

$$\begin{aligned} P(k, k-1) &= \Phi(k, k-1) P(k-1, k-1) \Phi^T(k, k-1) \\ &+ N(k-1) Q(k-1) N^T(k-1). \end{aligned} \quad (8)$$

We now receive the kth observation

$$z(k) = H(k) x(k) + v(k) \quad (9)$$

where $z(k)$ is known but $x(k)$ and $v(k)$ are unknown. We can compute an observation error vector by equation (7) and (9) as

$$\tilde{z}(k, k-1) = H(k) [x(k) - \hat{x}(k, k-1)] + v(k) \quad (10)$$

we now can correct the estimate in the state vector based on the observable and computable estimate in the observation vector as

$$\hat{x}(k, k) = \hat{x}(k, k-1) + W(k) \tilde{z}(k, k-1) \quad (11)$$

where the weighting matrix $W(k)$ at stage k has p rows and m columns.

We next seek a procedure for selecting at each stage a pxm weighting matrix $W(k)$.

Using equation (10) in equation (11)

$$\begin{aligned} \hat{x}(k, k) &= \hat{x}(k, k-1) + W(k) \{H(k)[x(k) - \hat{x}(k, k-1)] + v(k)\} \\ &= \begin{bmatrix} I & -W(k)H(k) \end{bmatrix} \hat{x}(k, k-1) + W(k)H(k)x(k) + W(k)v(k) \end{aligned} \quad (12)$$

$\begin{matrix} p \times p & p \times m & m \times p \end{matrix}$

If we now define the "unknown" error vectors

$$\hat{x}(k) - \hat{x}(k, k) = \tilde{x}(k, k) \quad (13)$$

and

$$\hat{x}(k) - \hat{x}(k, k-1) = \tilde{x}(k, k-1) \quad (14)$$

then theoretically equation (12) in (13) yields

$$\tilde{x}(k, k) = [I - W(k)H(k)]\hat{x}(k) - [I - W(k)H(k)]\hat{x}(k, k-1) - W(k)v(k) \quad (15)$$

$$\tilde{x}(k, k) = [I - W(k)H(k)]\tilde{x}(k, k-1) - W(k)v(k) \quad (16)$$

Transposing (16) we obtain

$$\langle \tilde{x}(k, k) = \langle \tilde{x}(k, k-1) [I - H^T(k) W^T(k)] - \langle v(k) W^T(k) \quad (17)$$

The dyadic product of equation (16) and (17) yields

$$\begin{aligned} \tilde{x}(k, k) \tilde{x}(k, k) &= \{ [I - W(k)H(k)] \tilde{x}(k, k-1) - W(k)v(k) \} \\ &\quad \{ \langle \tilde{x}(k, k-1) [I - H^T(k) W^T(k)] - \langle v(k) W^T(k) \} \\ &= [I - W(k)H(k)] \tilde{x}(k, k-1) \tilde{x}(k, k-1) [I - H^T(k) W^T(k)] \\ &\quad - [I - W(k)H(k)] \tilde{x}(k, k-1) v(k) W^T(k) \\ &\quad - W(k)v(k) \tilde{x}(k, k-1) [I - H^T(k) W^T(k)] \\ &\quad + W(k)v(k) v(k) W^T(k). \end{aligned} \quad (18)$$

The square of the magnitude of the error vector $\tilde{x}(k, k)$ is given as the inner-product of equation (16) and equation (17) or, as the trace of the outer-product of equation (18) as

$$\begin{aligned} \langle \tilde{x}(k, k) \tilde{x}(k, k) &= \langle \tilde{x}(k, k-1) \tilde{x}(k, k-1) - 2 \langle \tilde{x}(k, k-1) W(k) H \tilde{x}(k, k-1) \\ &\quad + \langle \tilde{x}(k, k-1) H^T(k) W^T(k) W H \tilde{x}(k, k-1) \\ &\quad - 2 \langle \tilde{x}(k, k-1) W(k) v(k) \\ &\quad + 2 \langle \tilde{x}(k, k-1) H^T(k) W^T(k) W(k) v(k) \\ &\quad + \langle v(k) W^T(k) W(k) v(k) \end{aligned} \quad (19)$$

Equation (19) is a scalar valued function of a matrix argument $W(k)$ of size $p \times m$. We shall take the partial derivative of the scalar with respect to the matrix $W(k)$,

$$\frac{\partial}{\partial W} \langle \tilde{x}(k, k) \tilde{x}(k, k) \rangle \quad \text{term by term.}$$

By equation (19), there are six additive terms, each term will be handled via the "gradient" methods in Appendix B.

The first term is not a function of

$W(k)$, hence

$$\frac{\partial}{\partial W(k)} \langle \tilde{x}(k, k-1) \tilde{x}(k, k-1) \rangle = [0] \quad (20)$$

The second term is by equation (b-61)

$$\frac{\partial}{\partial W(k)} \left\{ -2 \langle \tilde{x}(k, k-1) W(k) H(k) \tilde{x}(k, k-1) \rangle \right\} = -2H(k) \tilde{x}(k, k-1) \times \tilde{x}(k, k-1) \quad (21)$$

The third term is

$$\begin{aligned} & \frac{\partial}{\partial W} \{ \langle \tilde{x}(k, k-1) H^T(k) W^T(k) W(k) H(k) \tilde{x}(k, k-1) \rangle \} \\ & = 2H(k) \tilde{x}(k, k-1) \times \tilde{x}(k, k-1) H^T(k) W^T(k) \end{aligned} \quad (22)$$

The latter derivation is based on equation (b-80) and setting

$$\begin{aligned} \langle c &= \tilde{x}(k, k-1) H^T(k) \\ b(m) &= H(k) \tilde{x}(k, k-1) \end{aligned} \quad (23)$$

The fourth term is by equation (b-56)

$$\frac{\partial}{\partial W} \{ -2 \langle \tilde{x}(k, k-1) W(k) v(k) \rangle \} = -2v(k) \times \tilde{x}(k, k-1) \quad (24)$$

The fifth term by equation (b-80) is

$$\begin{aligned} & \frac{\partial}{\partial W} \{ 2 \langle \tilde{x}(k, k-1) H^T(k) W^T(k) W(k) v(k) \rangle \} \\ & = \{ H(k) \tilde{x}(k, k-1) \times v(k) + v(k) \times \tilde{x}(k, k-1) H^T \} W^T \end{aligned} \quad (26)$$

based on setting

$$\langle c = \langle \tilde{x} | I^T, v(k) \rangle = b \rangle \quad (27)$$

in equation (b-80)

The sixth term is

$$\frac{\partial}{\partial W} \{ \langle v(k) W^T(k) W(k) v(k) \rangle \} = 2v(k) \times v(k) W^T(k) \quad (28)$$

Utilizing the above six expressions in equation (19) after the partial derivative has been taken

$$\begin{aligned} & \frac{\partial}{\partial W} \{ \langle \tilde{x}(k, k) \tilde{x}(k, k) \rangle \} \\ &= -2H(k) \tilde{x}(k, k-1) \times \tilde{x}(k, k-1) \\ &+ 2H(k) \tilde{x}(k, k-1) \times \tilde{x}(k, k-1) H^T(k) W^T(k) \\ &- 2v(k) \times \tilde{x}(k, k-1) \\ &+ \{ H(k) \tilde{x}(k, k-1) \times v(k) + v(k) \times \tilde{x}(k, k-1) H^T(k) \} W^T(k) \\ &+ 2v(k) \times v(k) W^T(k) = [0]_{\text{exp}} \end{aligned} \quad (29)$$

The expected value over all experiments and allowable values of j and n yields

$$\begin{aligned} E \left\{ \frac{\partial}{\partial W} \langle \tilde{x}(k, k) \tilde{x}(k, k) \rangle \right\} &= -2H(k) E \{ \tilde{x}(k, k-1) \times \tilde{x}(k, k-1) \} \\ &+ 2H(k) E \{ \tilde{x}(k, k-1) \times \tilde{x}(k, k-1) \} H^T(k) W^T(k) \\ &- 2E \{ v(k) \times \tilde{x}(k, k-1) \} \\ &+ [H(k) E \{ \tilde{x}(k, k-1) \times v(k) + v(k) \times \tilde{x}(k, k-1) H^T(k) \} W^T(k) \\ &+ 2E \{ v(k) \times v(k) \} W^T(k) = [0] \end{aligned} \quad (30)$$

If we use the notation

$$P(k, k) = \cancel{\mathbb{E}_{\tilde{x}\tilde{x}}}(k, k) = E \{ \tilde{x}(k, k) \cancel{\tilde{x}(k, k)} \} \quad (31)$$

$$P(k, k-1) = \cancel{\mathbb{E}_{\tilde{x}\tilde{x}}}(k, k-1) = E \{ \tilde{x}(k, k-1) \cancel{\tilde{x}(k, k-1)} \} \quad (32)$$

$$Q(k) = E \{ u(k) \cancel{u(k)} \} \quad (33)$$

$$R(k) = E \{ v(k) \cancel{v(k)} \} \quad (34)$$

and assume that the conventional statistical independence assumptions hold,

$$E \{ \tilde{x}(k, k-1) \cancel{u(k)} \} = [0] \quad (35)$$

$$E \{ v(k) \cancel{v(k-1)} \} = [0] \quad (36)$$

$$E \{ v(k) \cancel{u(k)} \} = [0], \text{ etc.} \quad (37)$$

By equation (31) through (37) in equation (30)

$$\begin{aligned} \frac{\partial (\text{tr } P(k, k))}{\partial W(k)} &= -2H(k) P(k, k-1) \\ &+ 2H(k) P(k, k-1) H^T(k) W^T(k) \\ &+ 2\cancel{\mathbb{E}_{vv}}(k) W^T(k) = [0] \end{aligned} \quad (38)$$

$$[H(k) P(k, k-1) H^T(k) + \cancel{\mathbb{E}_{vv}}(k)] W^T(k) = H(k) P(k, k-1) \quad (39)$$

Transposing

$$W(k) [H(k) P(k, k-1) H^T(k) + \cancel{\mathbb{E}_{vv}}(k)] = P(k, k-1) H^T(k) \quad (40)$$

Inverting

$$W(k) = P(k, k-1) H^T(k) [H(k) P(k, k-1) H(k) + \cancel{\mathbb{E}_{vv}}(k)]^{-1} \quad (41)$$

or

$$W(k) = P(k, k-1) H^T(k) [H(k) P(k, k-1) H^T(k) + R(k)]^{-1} \quad (42)$$

$$W^T(k) = [H(k) P(k, k-1) H^T(k) + R(k)]^{-1} H(k) P(k, k-1) \quad (43)$$

The $p \times p$ matrix variance of the estimate of state (the p -space ellipsoid of uncertainty) can be obtained by taking the expected value over all experiments of the dyadic product of equation (18),

$$\begin{aligned} P(k, k) &= E\{\tilde{x}(k, k) \tilde{x}^T(k, k)\} \\ &= [I - W(k) H(k)] P(k, k-1) [I - H^T(k) W^T(k)] \\ &\quad + W(k) R(k) W^T(k) \end{aligned} \quad (44)$$

Multiplying out the terms of equation (44) we obtain

$$\begin{aligned} P(k, k) &= P(k, k-1) - P(k, k-1) H^T(k) W^T(k) \\ &\quad - W(k) H(k) P(k, k-1) + W(k) H(k) P(k, k-1) H^T(k) W^T(k) \\ &\quad + W(k) R(k) W^T(k) \\ &= P(k, k-1) - P(k, k-1) H^T(k) W^T(k) \\ &\quad - W(k) H(k) P(k, k-1) \\ &\quad + W(k) [H(k) P(k, k-1) H^T(k) + R(k)] W^T(k) \end{aligned} \quad (45)$$

Consider the last term of the above equation and equation (42) for $W(k)$ with the transpose (43), then the last term becomes

$$\begin{aligned} &W(k) [H(k) P(k, k-1) H^T(k) + R(k)] W^T(k) \\ &= P(k, k-1) H^T(k) [H(k) P(k, k-1) H^T(k) + R(k)]^{-1} [H(k) P(k, k-1) H^T(k) + R(k)] \\ &\quad \times [H(k) P(k, k-1) H^T(k) + R(k)]^{-1} H(k) P(k, k-1) \\ &= P(k, k-1) H^T(k) [H(k) P(k, k-1) + R(k)]^{-1} H(k) P(k, k-1) \\ &= W(k) H(k) P(k, k-1). \end{aligned} \quad (46)$$

Using the above expression for the last term in equation (45) we obtain

$$\begin{aligned} P(k, k) &= P(k, k-1) - P(k, k-1) H^T(k) W^T(k) \\ &\quad - W(k) H(k) P(k, k-1) + W(k) H(k) P(k, k-1) \end{aligned} \quad (47)$$

or

$$P(k, k) = P(k, k-1) - P(k, k-1) H^T(k) W^T(k) \quad (48)$$

$$P(k, k) = P(k, k-1) \underset{p \times p}{[I - H^T(k) W^T(k)]} \underset{p \times p}{\quad} \quad (49)$$

The matrix $P(k, k-1)$ was predicted and computed during stage $k-1$.

We can now predict the next stage (k+1) state vector via the deterministic process-dynamics

$$\hat{x}(k+1, k) = \Phi(k+1, k) \hat{x}(k, k) + B(k) f(k). \quad (50)$$

The next stage prediction of the observation vector is

$$\hat{z}(k+1, k) = H(k+1) \hat{x}(k+1, k). \quad (51)$$

The error in the state vector at stage k+1 based on the prediction of equation (50) is

$$\tilde{x}(k+1, k) = x(k+1) - \hat{x}(k+1, k) \quad (52)$$

where the unknown state vector is

$$x(k+1) = \Phi(k+1, k)x(k) + B(k)f(k) + N(k)u(k) \quad (53)$$

and the unknown error is by equation (50) and equation (53) in equation (52)

$$\begin{aligned} \tilde{x}(k+1, k) &= \Phi(k+1, k)x(k) + B(k)f(k) + N(k)u(k) \\ &\quad - \Phi(k+1, k)\hat{x}(k, k) - B(k)f(k) \\ &= \Phi(k+1, k)[x(k) - \hat{x}(k, k)] + N(k)u(k) \end{aligned} \quad (54)$$

Using equation (54) in equation (52)

$$\tilde{x}(k+1, k) = \Phi(k+1, k) \tilde{x}(k, k) + N(k)u(k) \quad (55)$$

The transpose of equation (55) is

$$\tilde{x}(k+1, k)^T = \tilde{x}(k, k)^T \Phi^T(k+1, k) + u(k)^T N^T(k) \quad (56)$$

The dyadic product of equation (55) and equation (56) is

$$\begin{aligned} \tilde{x}(k+1, k) \tilde{x}(k+1, k)^T &= \Phi(k+1, k) \tilde{x}(k, k) \tilde{x}(k, k)^T \Phi^T(k+1, k) \\ &\quad + \Phi(k+1, k) \tilde{x}(k, k) u(k)^T N^T(k) \\ &\quad + N(k) u(k) \tilde{x}(k, k)^T \Phi^T(k+1, k) \\ &\quad + N(k) u(k) u(k)^T N^T(k) \end{aligned} \quad (57)$$

The expectation over all experiments of equation (57) is

$$\begin{aligned} E\{\tilde{x}(k+1, k) \tilde{x}(k+1, k)\} &= P(k+1, k) \\ &= \Phi(k+1, k) E\{\tilde{x}(k, k) \tilde{x}(k, k)\} \Phi^T(k+1, k) \\ &+ N(k) E\{u(k) u(k)\} N^T(k). \end{aligned} \quad (58)$$

The statistical independence assumption was invoked:

$$E\{\tilde{x}(k, k) u(k)\} = [0].$$

Define the process noise variance matrix

$$E\{u(k) u(k)\} = Q(k) \quad (59)$$

and equation (58) becomes

$$P(k+1, k) = \Phi(k+1, k) P(k, k) \Phi^T(k+1, k) + N(k) Q(k) N^T(k). \quad (60)$$

We may now summarize the equations and the computations to be performed at the k th stage as the k th stage summary.

We have available from stage $k-1$:

$$\begin{aligned} \hat{x}(k, k-1) \\ \hat{z}(k, k-1) &= H(k) \hat{x}(k, k-1) \\ P(k, k-1) \end{aligned}$$

Receive $z(k)$

Compute error in observation

$$\tilde{z}(k, k-1) = z(k) - \hat{z}(k, k-1) \quad (61)$$

Compute weight matrix $W(k)$ by equation (42)

$$W(k) = P(k, k-1) H^T(k) [H(k) P(k, k-1) H^T(k) + R(k)]^{-1} \quad (62)$$

Correct state estimate by equation (11)

$$\hat{x}(k, k) = \hat{x}(k, k-1) + W(k) \tilde{z}(k, k-1) \quad (63)$$

Compute new state variance by equation (49)

$$P(k, k) = P(k, k-1)[I - H^T(k) W^T(k)] \quad (64)$$

Predict to stage k+1, by equation (50)

$$\hat{x}(k+1, k) = \phi(k+1, k) \hat{x}(k, k) + B(k) f(k) \quad (65)$$

Predict observation by equation (51)

$$\hat{z}(k+1, k) = H(k+1) \hat{x}(k+1, k) \quad (66)$$

Predict state variance by equation (60)

$$P(k+1, k) = \phi(k+1, k) P(k, k) \phi^T(k+1, k) + N(k) Q(k) N^T(k) \quad (67)$$

wait for stage k+1 and new measurement vector.

The equations can be substituted and juggled around to obtain alternate expressions, for example using equation (42) and (61) in (63) we obtain

$$\begin{aligned} \hat{x}(k, k) &= \hat{x}(k, k-1) \\ &+ P(k, k-1) H^T(k) [H(k) P(k, k-1) H^T(k) + R(k)]^{-1} \\ &\times \{ z(k) - H(k) \phi(k, k-1) \hat{x}(k-1, k-1) - H(k) B(k-1) f(k-1) \} \end{aligned} \quad (68)$$

Many similar variations of the above systems of equations occur in the literature.

Stage k = 1.

The vector starting system of equations can be derived from equation for k = 1,

$$\hat{x}(1, 0) = \text{intelligent guess} \quad (69)$$

$$\hat{z}(1, 0) = H(1) \hat{x}(1, 0) \quad (70)$$

and

$$P(1, 0) = \text{intelligent guess based on experience about the process.} \quad (71)$$

Receive $z(1)$

Compute error

$$\tilde{z}(1, 0) = z(1) - \hat{z}(1, 0) \quad (72)$$

Compute first weight

$$W(1) = P(1, 0) [H^T(1) P(1, 0) H^T(1) + R(1)]^{-1} \quad (73)$$

Correct state

$$\hat{x}(1, 1) = \hat{x}(1, 0) + W(1) \tilde{z}(1, 0) \quad (74)$$

Compute state variance matrix by equation 49

$$P(1, 1) = P(1, 0) [I - H^T(1) W^T(1)] \quad (75)$$

Predict stage 2 by equation (50)

$$\hat{x}(2, 1) = \Phi(2, 1) \hat{x}(1, 1) + B(1) r(1) \quad (76)$$

Predict stage 2 observation

$$\hat{z}(2, 1) = H(2) \hat{x}(2, 1) \quad (77)$$

Predict state variance matrix by equation 67

$$P(2, 1) = \Phi(2, 1) P(1, 1) \Phi^T(2, 1) + N(1)Q(1)N^T(1) \quad (78)$$

ETC.

IV. EQUATION SUMMARY FOR COMPUTER APPLICATION

This section summarizes the equations of the previous sections and points out how to compute mechanize the estimation equations to recursively estimate the state vector as the observations "roll into the computer". Precomputation of the sequence of weighting matrices for large dynamical systems is a necessity.

The dynamical process is described by the state vector equation

$$X(k+1) = \Phi(k+1, k)X(k) + Bf(k) + Bf(k) + N(k)U(k) \quad (1)$$

and a system of noisy instruments whose outputs are functionally related to the states by the observation equation

$$Z(k) = H(k)X(k) + V(k) \quad (2)$$

The system of estimation equations are:

The state vector prediction equation

$$X(k+1, K) = \Phi(k+1, k)\hat{X}(k, k) + B(k)f(k) \quad (3)$$

The observation prediction

$$\hat{Z}(k+1, k) = H(k+1)\hat{X}(k+1, k) \quad (4)$$

The Observation error

$$\tilde{Z}(k+1, k) = Z(k+1) - \hat{Z}(k+1, k) \quad (5)$$

and the correction to the predicted states at $k+1$ after the $(k+1)^{th}$ observation is available

$$\hat{X}(k+1, k+1) = \hat{X}(k+1, k) + W(k+1)\tilde{Z}(k+1, k) \quad (6)$$

The sequence of weighting matrices $W(k)$ can be precomputed and stored in memory. The weights are:

$$W(k+1) = P(k+1, k)H(k+1) \left[H(k+1)P(k+1, k)H^T(k+1) + R(k+1) \right]^{-1} \quad (7)$$

where

$$P(k+1, k) = \phi(k+1, k)P(k, k)\phi^T(k+1, k) + N(k+1)Q(k+1)N^T(k+1). \quad (8)$$

The block diagram is shown in Figure (1) as the conventional feedback (discrete) system.

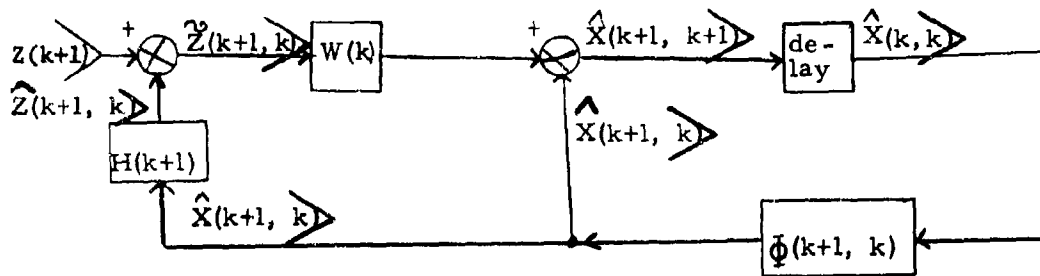


FIGURE (1) - DISCREET FEEDBACK BLOCK

By re-arranging the positions of the feedback blocks one can obtain a flow-block which looks more familiar to a digital programmer as shown in Figure (2).

Tests or applications in which one can plan or design the experiment and the times $k, k+1$, etc. at which instrument-data will be used to estimate appear to admit of pre-computing the weights. If the estimation times are not pre-designed one must compute the weights on-line.

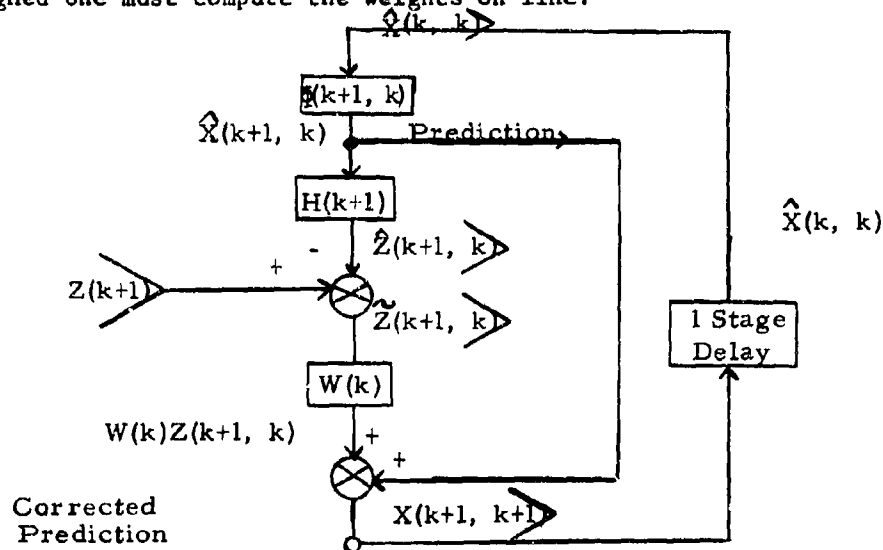


Figure (3) - Flow Block of Estimator

APPENDIX A - MATRIX TRACE PROPERTIES

The trace of a matrix, the trace of the product of two matrices, and the trace of a matrix-sum are useful notions to aid the development of the topics of Appendix B.

Consider a matrix A of p rows and m columns where $m < p$ and a matrix B , then the product

$$Q_1 = \begin{matrix} & & B \\ A & & \\ \text{pxp} & \text{pxm} & \text{mxp} \end{matrix} \quad (1)$$

is a pxp matrix.

The matrices A and B can be partitioned into their row and column spaces as shown

$$A = \left[\begin{matrix} \text{a}(p) \\ \text{a}(m) \\ \text{a}(1) \end{matrix} \right], \dots, \left[\begin{matrix} \text{a}(p) \\ \text{a}(m) \\ \text{a}(1) \end{matrix} \right] = \left[\begin{matrix} \text{a}(p) \\ \text{a}(m) \\ \text{a}(1) \end{matrix} \right] \quad (2)$$

$$B = \left[\begin{matrix} \text{b}(p) \\ \text{b}(m) \\ \text{b}(1) \end{matrix} \right] = \left[\begin{matrix} \text{b}(m) \\ \text{b}(1) \end{matrix} \right], \dots, \left[\begin{matrix} \text{b}(m) \\ \text{b}(1) \end{matrix} \right] \quad (3)$$

The product Q can be written as a matrix of inner-products

$$Q_1 = AB = \left[\begin{matrix} \text{a}(p) \\ \text{a}(m) \\ \text{a}(1) \end{matrix} \right] \left[\begin{matrix} \text{b}(m) \\ \text{b}(1) \end{matrix} \right], \dots, \left[\begin{matrix} \text{a}(p) \\ \text{a}(m) \\ \text{a}(1) \end{matrix} \right] \left[\begin{matrix} \text{b}(m) \\ \text{b}(1) \end{matrix} \right] \quad (4)$$

$$= \left[\begin{matrix} \text{a}(p) \text{b}(m) & \dots & \text{a}(p) \text{b}(1) \\ \text{a}(m) \text{b}(m) & \dots & \text{a}(m) \text{b}(1) \\ \text{a}(1) \text{b}(m) & \dots & \text{a}(1) \text{b}(1) \end{matrix} \right] \quad (5)$$

or as a sum of dyads (outer products)

$$Q_1 = AB = \begin{bmatrix} \langle a(p) |_1 & \dots & \langle a(p) |_m \end{bmatrix} \begin{bmatrix} |1\rangle b \\ \vdots \\ |m\rangle b \end{bmatrix} \quad (6)$$

$$Q_1 = \langle a(p) |_1 |1\rangle b + \dots + \langle a(p) |_m |m\rangle b \quad (7)$$

Equation (7) expresses Q_1 as a sum of m rank-one matrices.

If we commute the product we obtain a square $m \times m$ matrix

$$Q_2 = \underset{m \times m}{B} \underset{m \times p}{A} \underset{p \times m}{A}$$

and as before Q_2 can be written as a matrix of inner-products

$$Q_2 = \begin{bmatrix} \langle 1 | b \rangle \\ \vdots \\ \langle m | b \rangle \end{bmatrix} \begin{bmatrix} \langle a(p) |_1 & \dots & \langle a(p) |_m \end{bmatrix} \quad (8)$$

$$= \begin{bmatrix} \langle 1 | b \rangle \langle a(p) |_1 & \dots & \langle 1 | b \rangle \langle a(p) |_m \\ \vdots & & \vdots \\ \langle m | b \rangle \langle a(p) |_1 & \dots & \langle m | b \rangle \langle a(p) |_m \end{bmatrix} \quad (9)$$

or as a sum of dyadic products

$$Q_2 = \begin{bmatrix} |b(m)\rangle_1 & \dots & |b(m)\rangle_p \end{bmatrix} \begin{bmatrix} \langle 1 | a \rangle \\ \vdots \\ \langle p | a \rangle \end{bmatrix} \quad (10)$$

$$= |b(m)\rangle_1 \langle 1 | a \rangle + \dots + |b(m)\rangle_p \langle p | a \rangle \quad (11)$$

Clearly matrix multiplication is not commutative, that is

$$\begin{matrix} AB \\ p \times p \end{matrix} \neq \begin{matrix} BA \\ m \times m \end{matrix} \quad (12)$$

in fact the matrices are not even of the same size.

However the trace of both products are equal, that is

$$\text{tr} \begin{matrix} (AB) \\ p \times p \end{matrix} = \text{tr} \begin{matrix} (BA) \\ m \times m \end{matrix} \quad (13)$$

The following will clarify the above relation.

If we have a column vector $x \begin{matrix} p \\ \times \end{matrix}$ and a row vector $\begin{matrix} p \\ \times \end{matrix} y$ of the same dimension p then the dyadic product is the square, rank-one, matrix D of p rows and p columns

$$\begin{matrix} D \\ p \times p \end{matrix} = \begin{matrix} p \\ \times \end{matrix} \begin{matrix} p \\ \times \end{matrix} y = \begin{pmatrix} x^1 y_1 & \dots & x^1 y_p \\ x^p y_1 & \dots & x^p y_p \end{pmatrix} \quad (14)$$

If we commute the product of Equation (14) we obtain

$$\begin{matrix} d \\ 1 \times 1 \end{matrix} = \begin{matrix} p \\ \times \end{matrix} y x \begin{matrix} p \\ \times \end{matrix} = y_1 x^1 + y_2 x^2 + \dots + y_p x^p \quad (15)$$

a scalar.

When the elements y_i and x^i are real field elements the products commute, hence

$$y_i x^i = x^i y_i \quad (16)$$

and Equation (15) [the inner product] can be written as the sum of the main diagonal terms of $\begin{matrix} p \\ \times \end{matrix} \begin{matrix} p \\ \times \end{matrix}$, which the conventional definition of the trace (tr) of a matrix, hence

$$\text{tr} \begin{bmatrix} \begin{matrix} p \\ \times \end{matrix} \begin{matrix} p \\ \times \end{matrix} \end{bmatrix} = \begin{matrix} p \\ \times \end{matrix} \begin{matrix} p \\ \times \end{matrix} \quad (17)$$

The dyadic product is not as mysterious as many novices might imagine; in fact, if we write Equation (14) as

$$\begin{matrix} D \\ p \times p \end{matrix} = \begin{matrix} p \\ \times \end{matrix} \begin{matrix} p \\ \times \end{matrix} y = \begin{matrix} p \\ \times \end{matrix} \begin{pmatrix} y_1 & y_2 & \dots & y_p \end{pmatrix} \\ = \begin{bmatrix} \begin{matrix} p \\ \times \end{matrix} y_1 & \begin{matrix} p \\ \times \end{matrix} y_2 & \dots & \begin{matrix} p \\ \times \end{matrix} y_p \end{bmatrix} \quad (18)$$

we see that the matrix D when partitioned into its column space is a row of p parallel column vectors - all p of the vectors lie on a line, hence $\begin{matrix} p \\ \times \end{matrix} \begin{matrix} p \\ \times \end{matrix}$ is said to have rank one - that is, there is only one linearly independent vector in the row "package" of column vectors.

APPENDIX B
GRADIENTS OF SCALARS WITH RESPECT TO MATRICES

This appendix develops the gradient of a scalar-valued function with respect to a vector variable and also with respect to a matrix variable.

Case 1. $q = \langle p \rangle a x \langle p \rangle$. Consider the scalar q which is the inner-product

$$q = \langle p \rangle a x \langle p \rangle \quad (b-1)$$

where $\langle a$ is a fixed p dimensional row vector and $x \rangle$ is a variable column vector, or q is said to be a scalar-valued variable which is a function of the vector variable $x \rangle$.

In equation (b-1) q may be considered to have vector factors $\langle a$ and $x \rangle$.

If we have a dyad

$$Q = x \rangle \langle a \quad (b-2)$$

then it was shown in appendix A that

$$\text{tr } Q = q \quad (b-3)$$

or

$$\text{tr } [x \rangle \langle a] = \langle a x \rangle = q \quad (b-4)$$

The differential of equation (b-2) is

$$dQ = dx \rangle \langle a \quad (b-5)$$

and the trace of (b-5) is

$$\text{tr } dQ = \text{tr } [dx \rangle \langle a] = \langle a dx \rangle = dq \quad (b-6)$$

We may now ask to express the differential matrix dQ in terms of vector factors $dx \rangle$ and a gradient vector, that is

$$dQ = dx \rangle \left\langle \frac{\partial q}{\partial x} \right. \quad (b-7)$$

such that the trace of equation (b-7) is

$$dq = \text{tr } dQ = \text{tr } \left[dx \rangle \left\langle \frac{\partial q}{\partial x} \right. \right] = \left\langle \frac{\partial q}{\partial x} dx \right\rangle \quad (b-8)$$

If we take the trace of AB by Equation (5) as the sum of diagonals we obtain

$$\text{tr}(AB) = \langle \overset{1}{\cancel{m}} \rangle a \langle \overset{1}{\cancel{m}} \rangle b + \dots + \langle \overset{p}{\cancel{m}} \rangle a \langle \overset{p}{\cancel{m}} \rangle b \quad (12)$$

If we take the trace of dyadic sum decomposition of AB given by Equation (7) we obtain

$$\text{tr}(AB) = \text{tr} \left[a \langle \overset{1}{\cancel{p}} \rangle \langle \overset{1}{\cancel{p}} \rangle b + \dots + a \langle \overset{m}{\cancel{p}} \rangle \langle \overset{m}{\cancel{p}} \rangle b \right] \quad (13)$$

The trace of a sum of matrices is the sum of the traces, hence by Equation (17)

$$\text{tr}(AB) = \text{tr} a \langle \overset{1}{\cancel{p}} \rangle \langle \overset{1}{\cancel{p}} \rangle b + \text{tr} a \langle \overset{2}{\cancel{p}} \rangle \langle \overset{2}{\cancel{p}} \rangle b + \dots + \text{tr} a \langle \overset{m}{\cancel{p}} \rangle \langle \overset{m}{\cancel{p}} \rangle b \quad (14)$$

$$\text{tr}(AB) = \langle \overset{1}{\cancel{p}} \rangle b a \langle \overset{1}{\cancel{p}} \rangle + \langle \overset{2}{\cancel{p}} \rangle b a \langle \overset{2}{\cancel{p}} \rangle + \dots + \langle \overset{m}{\cancel{p}} \rangle b a \langle \overset{m}{\cancel{p}} \rangle \quad (15)$$

Equation (12) is a sum of p inner-products of m-dimensional vectors and Equation (15) is a sum of m inner-products of p-dimensional vector.

The sum of the main diagonal terms of Equation (9) is

$$\text{tr}(BA) = \langle \overset{1}{\cancel{p}} \rangle b a \langle \overset{1}{\cancel{p}} \rangle + \dots + \langle \overset{m}{\cancel{p}} \rangle b a \langle \overset{m}{\cancel{p}} \rangle \quad (16)$$

which by Equation (15) and Equation (16)

$$\text{tr}_{pp}(AB) = \text{tr}_{m \times m}(BA) \quad (17)$$

By equation (b-7) and (b-5) we can state

$$\left\langle \frac{\partial q}{\partial x} \right\rangle = \left\langle a \right\rangle. \quad (b-9)$$

We arrive at the result of equation (b-9) directly from (1)

$$dq = \left\langle a \right\rangle dx = \left\langle \frac{\partial q}{\partial x} \right\rangle dx \quad (b-10)$$

hence

$$\left\langle a \right\rangle = \left\langle \frac{\partial q}{\partial x} \right\rangle. \quad (b-11)$$

Also one can consider the gradient as an operator $\left\langle \frac{\partial}{\partial p} \right\rangle$

$$q \left\langle \frac{\partial}{\partial q} \right\rangle = \left\langle a \right\rangle x \left\langle \frac{\partial}{\partial q} \right\rangle = \left\langle a \right\rangle \left[x \left\langle \frac{\partial}{\partial q} \right\rangle \right] \quad (b-12)$$

The dyadic-type operator

$$x \left\langle \frac{\partial}{\partial x} \right\rangle = \begin{pmatrix} x^1 \\ x^2 \\ \vdots \\ x^p \end{pmatrix} \left(\frac{\partial}{\partial x_1}, \dots, \frac{\partial}{\partial x_p} \right) \quad (b-13)$$

$$= \begin{bmatrix} \frac{\partial x^1}{\partial x_1} & \frac{\partial x^1}{\partial x_2} & \dots & \frac{\partial x^1}{\partial x_p} \\ \frac{\partial x^p}{\partial x_1} & \frac{\partial x^p}{\partial x_2} & \dots & \frac{\partial x^p}{\partial x_p} \end{bmatrix} \quad (b-14)$$

when the coordinates are independent of each other, then

$$x \left\langle \frac{\partial}{\partial x} \right\rangle = I. \quad (b-15)$$

Hence

$$q \left\langle \frac{\partial}{\partial x} \right\rangle = \left\langle \frac{\partial q}{\partial x} \right\rangle = \langle a \rangle \quad (b-16)$$

In conclusion:

$$\boxed{\begin{array}{l} \text{if } q = \langle a \rangle \\ \text{then} \end{array}} \quad (b-17)$$

$$\left\langle \frac{\partial q}{\partial x} \right\rangle = \langle a \rangle \quad (b-18)$$

Case 2. $q = \langle x \rangle$.

When q is quadratic we can write q as the trace of the dyad

$$Q = \rangle \rangle \langle \langle \quad (b-19)$$

for

$$\text{tr } Q = \text{tr } (\rangle \rangle \langle \langle) = \langle x \rangle = q. \quad (b-20)$$

The differential of the dyad

$$dQ = dx \rangle \langle x + x \rangle dx \quad (b-21)$$

$$dq = \text{tr } dQ = \langle x \rangle dx + \langle dx \rangle x = 2 \langle x \rangle dx = \left\langle \frac{\partial q}{\partial x} \right\rangle dx \quad (b-22)$$

hence

$$\left\langle \frac{\partial q}{\partial x} \right\rangle = 2 \langle x \rangle \quad (b-23)$$

Case 3. $q = \langle x B x \rangle$ (b-24)

For this case we have two different matrices

$$Q_1 = B \rangle \rangle \langle \langle x \quad (b-25)$$

and

$$Q_2 = \rangle \rangle \langle \langle x B \quad (b-26)$$

which under the trace operation map down to the same scalar

$$q = \text{tr } Q_1 = \text{tr } Q_2 = \langle x B x \rangle \quad (\text{b-27})$$

The differential of $Q_2 = Q$ is

$$dQ = dx \langle x B + x \rangle dx B \quad (\text{b-28})$$

The trace of (b-28) is

$$\text{tr } dQ = \langle x B dx \rangle + \langle dx B x \rangle \quad (\text{b-29})$$

The differential of (b-24) is

$$dq = \langle dx B x \rangle + \langle x B dx \rangle = \text{tr } Q \quad (\text{b-30})$$

$$dq = \langle x B^T dx \rangle + \langle x B dx \rangle \quad (\text{b-31})$$

$$dq = \langle x [B + B^T] dx \rangle \quad (\text{b-32})$$

we have

$$dq = \left\langle \frac{\partial q}{\partial x} dx \right\rangle \quad (\text{b-33})$$

and by (b-32) and (b-33)

$$\left\langle \frac{\partial q}{\partial x} \right\rangle = \langle x [B + B^T] \rangle \quad (\text{b-34})$$

and for symmetric B

$$B = B^T \quad (\text{b-35})$$

then

$$\left\langle \frac{\partial q}{\partial x} \right\rangle = 2 \langle x B \rangle \quad (\text{b-36})$$

Case 4.

$$q = \langle p \rangle_a X b(m) \quad (\text{b-37})$$

The scalar q is a function of the matrix X of p -rows and m columns.

The scalar q can be written as the trace of the matrix

$$Q = b(m) \langle p \rangle_a X \quad (\text{b-38})$$

The differential of Q is

$$dQ = \langle b | a dX \rangle \quad (b-39)$$

By equation (b-37), differentiating

$$dq = \langle a dX b \rangle = \text{tr } dQ. \quad (b-40)$$

We seek a gradient matrix $\frac{\partial q}{\partial X}$ of m rows and p columns as one of the factors of dQ that is

$$dQ = \begin{matrix} \text{mxm} & \frac{\partial q}{\partial X} & \text{dX} \\ & \text{mxp} & \text{pxm} \end{matrix} \quad (b-41)$$

such that

$$\text{trd}Q = dq = \langle a dX b \rangle \quad (b-42)$$

Clearly by equation (b-39) and (b-41) if

$$\begin{matrix} \frac{\partial q}{\partial X} \\ \text{mxp} \end{matrix} = \begin{matrix} b(m) \\ \langle b \rangle \end{matrix} \begin{matrix} \langle a \rangle \\ p \end{matrix} \quad (b-43)$$

then (b-42) is satisfied.

An alternate, more direct, approach is given below. Partition X into a row of column vectors (all "contravariant" vectors), then

$$\begin{aligned} q &= \langle p | a \left[x^{(p)}_1, \dots, x^{(p)}_m \right] b^{(m)} \rangle \quad (b-44) \\ &= \left[\langle a x_1 \rangle, \langle a x_2 \rangle, \dots, \langle a x_m \rangle \right] b^{(m)} \\ &= \left[\langle a x_1 \rangle, \langle a x_2 \rangle, \dots, \langle a x_m \rangle \right] \begin{bmatrix} 1 \\ b \\ b^2 \\ \vdots \\ b^m \\ b \end{bmatrix} \end{aligned}$$

$$q = \langle a | x \rangle_1 b^1 + \langle a | x \rangle_2 b^2 + \dots + \langle a | x \rangle_m b^m$$

$$= q_1 \langle x \rangle_1 + \dots + q_m \langle x \rangle_m$$
(b-45)

where each q_i is a function of a single column vector $\langle x \rangle_1$.

The scalar differential of q is

$$dq = \left\langle \frac{\partial q}{\partial x} \right|_1 dx_1 + \left\langle \frac{\partial q}{\partial x} \right|_2 dx_2 + \dots + \left\langle \frac{\partial q}{\partial x} \right|_m dx_m$$
(b-46)

$$dq = \left[\left\langle \frac{\partial q}{\partial x} \right|_1, \left\langle \frac{\partial q}{\partial x} \right|_2, \dots, \left\langle \frac{\partial q}{\partial x} \right|_m \right] \begin{bmatrix} dx_1 \\ dx_2 \\ \vdots \\ dx_m \end{bmatrix}$$
(b-47)

Equation (b-47) can be written as

$$dq = \text{tr} \left\{ \begin{bmatrix} \left\langle \frac{\partial q}{\partial x} \right|_1 \\ \vdots \\ \left\langle \frac{\partial q}{\partial x} \right|_m \end{bmatrix} \begin{bmatrix} dx_1 & \dots & dx_m \end{bmatrix} \right\}$$
(b-48)

$$= \text{tr} \left\{ \begin{bmatrix} \left\langle \frac{\partial q}{\partial x} \right|_1 dx_1 & \dots & \left\langle \frac{\partial q}{\partial x} \right|_m dx_m \\ \vdots & & \vdots \\ \left\langle \frac{\partial q}{\partial x} \right|_1 dx_1 & \dots & \left\langle \frac{\partial q}{\partial x} \right|_m dx_m \end{bmatrix} \right\}$$
(b-49)

The differential of X is a row of column vectors

$$dX_{p \times m} = \begin{bmatrix} dx^{(p)}_1 & \dots & dx^{(p)}_m \end{bmatrix} \quad (b-50)$$

and the gradient matrix is a column of row gradient-vectors.

$$\frac{\partial q}{\partial X} = \begin{pmatrix} \langle p | \frac{\partial q}{\partial x} \\ \vdots \\ \langle p | \frac{\partial q}{\partial x} \end{pmatrix} \quad (b-51)$$

From the foregoing we write

$$dQ = \frac{\partial q}{\partial X} dX \quad (b-52)$$

$m \times p$

and

$$dq = \text{tr } dQ = \text{tr} \begin{bmatrix} \frac{\partial q}{\partial X} & dX \end{bmatrix} \quad (b-52)$$

By equation (b-45), (b-46) and (b-16)

$$\begin{aligned} \langle p | \frac{\partial q}{\partial x} &= \langle \frac{\partial q}{\partial x} | 1 \rangle = b^1 \langle a \\ &\vdots \\ \langle p | \frac{\partial q}{\partial x} &= \langle \frac{\partial q}{\partial x} | m \rangle = b^m \langle a \end{aligned} \quad (b-53)$$

Packaging the row vector gradients of (b-53) into the column of (b-51) we obtain

$$\frac{\partial q}{\partial X} = \begin{pmatrix} b^1 \langle p | a \\ b^2 \langle p | a \\ \vdots \\ b^m \langle p | a \end{pmatrix} = \begin{pmatrix} b^1 \\ b^2 \\ \vdots \\ b^m \end{pmatrix} \langle p | a \quad (b-54)$$

5_1

or

$$\frac{\partial q}{\partial X_{m \times p}} = b(m) \langle p \rangle a \quad (b-55)$$

hence in conclusion

if $q = \langle p \rangle a \underset{p \times m}{X} b(m)$

then $\frac{\partial q}{\partial X_{m \times p}} = b(m) \langle p \rangle a.$

(b-56)

Case 5. $q = \langle p \rangle a \underset{p \times m}{X} \underset{m \times p}{B} a(p)$ (b-57)

For this case we set

$$\underset{m \times p}{B} a(p) = b(m) \quad (b-58)$$

as in equation (b-56), then

$$q = \langle a \underset{p \times m}{X} b(m) \rangle \quad (b-59)$$

and we obtain the case 4, hence

$$\frac{\partial q}{\partial X_{m \times p}} = \underset{m \times p}{B} a(p) \langle p \rangle a \quad (b-60)$$

or

if $q = \langle p \rangle a \underset{p \times m}{X} \underset{m \times p}{B} a(p)$

then $\frac{\partial q}{\partial X_{m \times p}} = \underset{m \times p}{B} a(p) \langle p \rangle a$

(b-61)

Case 6. $q = \langle \underline{p} \rangle_a X_{p \times m} X_{m \times p}^T b(\underline{p}) \rangle$ (b-62)

This case is the matrix analog of the quadratic vector case of equation (b-24).

We can partition X into its column space and X^T into its row space and obtain

$$q = \langle \underline{p} \rangle_a \left[x(\underline{p})_1, \dots, x(\underline{p})_m \right] \begin{bmatrix} \langle \underline{p} \rangle_1 x \\ \vdots \\ \langle \underline{p} \rangle_m x \end{bmatrix} b(\underline{p}) \quad (b-63)$$

and

$$q = \langle a \left[x(\underline{p})_1 \langle \underline{p} \rangle_1 x + \dots + x(\underline{p})_m \langle \underline{p} \rangle_m x \right] b(\underline{p}) \rangle \quad (b-64)$$

Distributing the two end vectors over the dyadic-sum decomposition of XX^T we obtain

$$q = \langle a \rangle_1 \langle \underline{p} \rangle_1 b + \langle a \rangle_2 \langle \underline{p} \rangle_2 b + \dots + \langle a \rangle_m \langle \underline{p} \rangle_m b \quad (b-65)$$

Because of inner-product commutativity

$$\langle \underline{x} \rangle_1 b = \langle b \rangle_1 \underline{x} \quad (b-66)$$

hence

$$\begin{aligned} q &= \langle a \rangle_1 \langle b \rangle_1 + \dots + \langle a \rangle_m \langle b \rangle_m \\ &= p_1(\underline{x})_1 q_1(\underline{x})_1 + \dots + p_m(\underline{x})_m q_m(\underline{x})_m \end{aligned} \quad (b-67)$$

hence the scalar q is a sum of products of scalars $p_i q_i$.

We have as before

$$dq = \left[\frac{\partial q}{\partial x}, \frac{\partial q}{\partial x}, \dots, \frac{\partial q}{\partial x} \right] \begin{bmatrix} d\underline{x} \\ \vdots \\ d\underline{x} \end{bmatrix} = \text{tr } dQ \quad (b-68)$$

where dQ is as in equation (b-41).

$$\left\langle \frac{\partial Q}{\partial x} \right\rangle = \left\langle \frac{\partial (p_1 q_1)}{\partial x} \right\rangle = q_1 \left\langle \frac{\partial p_1}{\partial x} \right\rangle + p_1 \left\langle \frac{\partial q_1}{\partial x} \right\rangle \quad (b-69)$$

and

$$p_1 = \langle a | x | 1 \rangle \quad (b-70)$$

$$\left\langle \frac{\partial p_1}{\partial x} \right\rangle = \langle a | \quad (b-71)$$

$$q_1 = \langle 1 | b | x \rangle \quad (b-72)$$

$$\left\langle \frac{\partial q_1}{\partial x} \right\rangle = \langle 1 | b \quad (b-73)$$

Using (b-70), (b-71), (b-72), (b-73) in (b-69)

$$\begin{aligned} \left\langle \frac{\partial Q}{\partial x} \right\rangle &= q_1 \langle a + p_1 b \rangle \\ &= \langle x | b | 1 \rangle \langle a + p_1 b \rangle \end{aligned} \quad (b-74)$$

$$\left\langle \frac{\partial Q}{\partial x} \right\rangle = \langle x | [b(p) \langle p | a + a(p) \langle p | b] \quad (b-75)$$

Packaging (b-75) into the gradient matrix of equation (b-51)

$$\frac{\partial Q}{\partial \mathbf{x}}_{m \times p} = \begin{bmatrix} \langle x | [b \langle p | a + a \langle p | b] \\ \vdots \\ \langle x | [b \langle p | a + a \langle p | b] \end{bmatrix} \quad (b-76)$$

or

$$\frac{\partial q}{\partial X} = \begin{bmatrix} \langle p \rangle x \\ \vdots \\ \langle m \rangle x \end{bmatrix} \left[\begin{array}{c} b \langle a + a \rangle b \end{array} \right] \quad (b-77)$$

$$\frac{\partial q}{\partial X} = X^T \left[\begin{array}{c} b \langle a + a \rangle b \end{array} \right] \quad (b-78)$$

In conclusion

$$\begin{array}{l} \text{if } q = \langle p \rangle a X_{pxm} X_{m \times p}^T b \langle p \rangle \\ \text{then } \frac{\partial q}{\partial X}_{m \times p} = X^T_{m \times p} \left[\begin{array}{c} b \langle a + a \rangle b \end{array} \right] \end{array} \quad (b-79)$$

In a similar fashion it can be shown that

$$\begin{array}{l} \text{if } q = \langle m \rangle c X_{m \times p}^T X_{p \times m} b \langle m \rangle \\ \text{then } \frac{\partial q}{\partial X}_{m \times p} = \left[\begin{array}{c} c \langle m \rangle \langle m \rangle b + b \langle m \rangle c \end{array} \right] X_{m \times p}^T \end{array} \quad (b-80)$$

Consider the $p \times p$ matrix L which has factors as shown

$$L = B \cdot X \quad (b-81)$$

$p \times p \quad p \times k \quad k \times p$

where X is a variable matrix.

If we factor B into its column space and X into its row space

$$L = [b(p)_1, \dots, b(p)_k] \begin{bmatrix} 1 \\ \langle p \rangle x \\ \vdots \\ k \\ \langle p \rangle x \end{bmatrix} \quad (b-82)$$

$$= b(p)_1 \langle p \rangle x + \dots + b(p)_k \langle p \rangle x \quad (b-83)$$

The trace of L is

$$\text{tr } L = \ell = \langle x \rangle b_1 + \dots + \langle x \rangle b_k \quad (b-84)$$

The differential of (b-84) is

$$dL = B \, dX. \quad (b-85)$$

The factors of dL can also be expressed as

$$dL = \frac{\partial(\text{tr} L)}{\partial x_{p \times k}} \frac{\partial x}{k \times p} \quad (b-86)$$

where the $p \times k$ gradient matrix is

$$\frac{\partial \ell}{\partial x_{p \times k}} = \left[\frac{\partial \ell}{\partial x} \langle p \rangle_1, \frac{\partial \ell}{\partial x} \langle p \rangle_2, \dots, \frac{\partial \ell}{\partial x} \langle p \rangle_k \right] \quad (b-87)$$

The differential of Equation (b-84) is

$$d(\text{tr } L) = dl = dl_1 + \dots + dl_k \quad (\text{b-88})$$

$$= \left\langle \frac{\partial l}{\partial x} \right\rangle_1 + \dots + \left\langle \frac{\partial l}{\partial x} \right\rangle_k \quad (\text{b-89})$$

where

$$\left\langle \frac{\partial l}{\partial x} \right\rangle_1 = \frac{\partial l_1}{\partial x} = \frac{\partial}{\partial x} \left\langle x \right\rangle_1 b_1 = b_1 \quad (\text{b-90})$$

...

$$\left\langle \frac{\partial l}{\partial x} \right\rangle_k = \frac{\partial l_k}{\partial x} = \frac{\partial}{\partial x} \left\langle x \right\rangle_k b_k = b_k$$

and

$$\frac{\partial l}{\partial X} = [b(p)_1, \dots, b(p)_k] = B_{pxk} \quad (\text{b-91})$$

In summary,

If

$$L = B_{pxp} X_{pxk} (kxp) \quad (\text{b-92})$$

then

$$\frac{\partial(\text{tr } L)}{\partial X_{pxk}} = B_{pxk} \quad (\text{b-93})$$

APPENDIX C

MINIMIZATION

Consider the linear surface

$$l = \langle b \ x \rangle \quad (1)$$

and the quadratic surface

$$q = \langle x Q x \rangle \quad (2)$$

and the difference

$$q - l = \phi. \quad (3)$$

If l is a constant, $l = l_0$, then we seek a vector x that lies on the linear surface and on the quadratic surface such that difference in the linear surface and the quadratic surface is a minimum.

Differentiating

$$d\phi = dq - dl \quad (4)$$

and

$$d\phi = \left\langle \frac{\partial \phi}{\partial x} dx \right\rangle \quad (5)$$

$$= \left\langle \frac{\partial q}{\partial x} dx \right\rangle - \left\langle \frac{\partial l}{\partial x} dx \right\rangle \quad (6)$$

$$= \left\langle \left[\frac{\partial q}{\partial x} - \frac{\partial l}{\partial x} \right] dx \right\rangle \quad (7)$$

or

$$\left\langle \frac{\partial \phi}{\partial x} \right\rangle = \left\langle \frac{\partial q}{\partial x} - \frac{\partial l}{\partial x} \right\rangle \quad (8)$$

If we equate the gradient vector to zero

$$\left\langle \frac{\partial q}{\partial x} \right\rangle = \left\langle \frac{\partial l}{\partial x} \right\rangle. \quad (9)$$

By equation p^{-3b} and equation b^{-9}

$$2\langle x \rangle = \langle b \rangle \quad (10)$$

and solving for $\langle x \rangle$

$$\langle x \rangle = \frac{\langle b \rangle}{2} \quad (11)$$

Multiplying equation (11) by $\langle b \rangle$ and using equation (C-1)

$$\langle x \rangle \langle b \rangle = \frac{\langle b \rangle^2}{2} = \ell_0 \quad (12)$$

or

$$\frac{1}{2} = \frac{\ell_0}{\langle b \rangle} \quad (13)$$

Using (13) in (11)

$$\langle x \rangle = \ell_0 \frac{\langle b \rangle}{\langle b \rangle^2} \quad (14)$$

If

$$\ell_0 = 1$$

then

$$\langle x \rangle = \frac{1}{\langle b \rangle} \quad (15)$$

REFERENCES

1. Dwyer, Paul S., "Some Applications of Matrix Derivatives on Multi-Variate Analysis", American Statistical Association Journal, June 1967.
2. Kalman, R. E., "New Methods in Weiner Filtering Theory", Proceedings of the First Symposium in Engineering Applications of Random Function Theory and Probability. John Wiley and Sons, 1963.
3. Lee, C. K. Robert, "Optimal Estimation, Identification and Control", The M.I.T. Press, Cambridge, Mass., 1964.
4. Pappas, James S., "Application of the Kalman Filter To Sequential Optimal Parameter Estimation Via Householder's Matrix Inversion Method", R-O-S-67-1, Special Report, USATECOM, Analysis and Computation Directorate, Deputy for NRO, White Sands Missile Range, New Mexico, June 1967.
5. Pappas, J. S., Diaz, A., "A Math Model for Computing Noise Variance Matrices for a System of Radar Trackers", R-O-S-67-2, Special Report USATECOM, Analysis and Computation Directorate, Deputy for NRO, White Sands Missile Range, New Mexico, July 1967.
6. Pappas, J. S., "A Vector Space Derivation Using Dyads-of-Weighted Least Squares For Correlated Noise, RO-S-68-1, Special Report, USATECOM, Analysis and Computation Directorate, Deputy for NRO, White Sands Missile Range, New Mexico, June 1967.
7. Scheffe, Henry, "The Analysis of Variance", John Wiley, 1959.
8. Scott, Maceo T., Rede, Edmundo, and Wygant, Marthe, "Discrete Recursive Estimation: An Optimum Automated Data Processing Technique for Multi-Radar Tracking Systems", Technical Report Nr. 5, Analysis and Computation Directorate, White Sands Missile Range, New Mexico, April 1967.
9. Friedman, B., "Principles and Techniques of Applied Mathematics", John Wiley, 1956.

DISTRIBUTION

Special Report Nr. RO-S-68-2, "A Tutorial Deviation of Recursive Weighted Least Squares State-Vector Estimation Theory (Kalman Theory)", UNCLASSIFIED, National Range Operations, White Sands Missile Range, New Mexico 88002, August 1968.

White Sands Missile Range
New Mexico (STEWS)

Number
of
Copies

National Range Operations

Analysis and Computation Directorate (Record Copy)
(Ref. Copy)

1
1

Math Services Division, ATTN: Mr. Pappas

69

Technical Library

10

Defense Documentation Center
Cameron Station
Alexandria, Virginia 22314

20

UNCLASSIFIED

Security Classification

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) ANALYSIS AND COMPUTATION DIRECTORATE DEPUTY FOR NATIONAL RANGE OPERATIONS WHITE SANDS MISSILE RANGE, NEW MEXICO 88002		2a. REPORT SECURITY CLASSIFICATION UNCLASSIFIED 2b. GROUP
3. REPORT TITLE A TUTORIAL DERIVATION OF RECURSIVE WEIGHTED LEAST SQUARES STATE-VECTOR ESTIMATION THEORY (KALMAN THEORY)		
4. DESCRIPTIVE NOTES (Type of report and inclusive dates) SPECIAL REPORT		
5. AUTHOR(S) (Last name, first name, initial) PAPPAS, JAMES S.		
6. REPORT DATE AUGUST 1968	7a. TOTAL NO. OF PAGES 67	7b. NO. OF REFS 9
8a. CONTRACT OR GRANT NO.	8a. ORIGINATOR'S REPORT NUMBER(S)	
b. PROJECT NO.		
c.	8b. OTHER REPORT NO(S) (Any other numbers that may be assigned this report)	
d.		
10. AVAILABILITY/LIMITATION NOTICES DISTRIBUTION OF THIS DOCUMENT IS UNLIMITED		
11. SUPPLEMENTARY NOTES		12. SPONSORING MILITARY ACTIVITY
13. ABSTRACT This is one of a series of tutorial reports deriving the discrete recursive Kalman estimation equations. The series proceeds from unweighted to weighted least squares parameter estimation in a vector space setting to the stage-wise updating, to time varying parameter or state variable estimation. The derivations have been stage-wise tutorial in an attempt to make the theory accessible to the "non-specialist" in optimization theory.		

DD FORM 1473
1 JAN 64

UNCLASSIFIED

Security Classification

14. KEY WORDS	LINK A		LINK B		LINK C	
	ROLE	WT	ROLE	WT	ROLE	WT
1. Matrices						
2. Weighted least squares						
3. Correlated noise						
4. Kalman Filtering						
5. Optimal Estimator						
6. State vector						

INSTRUCTIONS

1. ORIGINATING ACTIVITY: Enter the name and address of the contractor, subcontractor, grantee, Department of Defense activity or other organization (*corporate author*) issuing the report.

2a. REPORT SECURITY CLASSIFICATION: Enter the overall security classification of the report. Indicate whether "Restricted Data" is included. Marking is to be in accordance with appropriate security regulations.

2b. GROUP: Automatic downgrading is specified in DoD Directive 5200.10 and Armed Forces Industrial Manual. Enter the group number. Also, when applicable, show that optional markings have been used for Group 3 and Group 4 as authorized.

3. REPORT TITLE: Enter the complete report title in all capital letters. Titles in all cases should be unclassified. If a meaningful title cannot be selected without classification, show title classification in all capitals in parenthesis immediately following the title.

4. DESCRIPTIVE NOTES: If appropriate, enter the type of report, e.g., interim, progress, summary, annual, or final. Give the inclusive dates when a specific reporting period is covered.

5. AUTHOR(S): Enter the name(s) of author(s) as shown on or in the report. Enter last name, first name, middle initial. If military, show rank and branch of service. The name of the principal author is an absolute minimum requirement.

6. REPORT DATE: Enter the date of the report as day, month, year; or month, year. If more than one date appears on the report, use date of publication.

7a. TOTAL NUMBER OF PAGES: The total page count should follow normal pagination procedures, i.e., enter the number of pages containing information.

7b. NUMBER OF REFERENCES: Enter the total number of references cited in the report.

8a. CONTRACT OR GRANT NUMBER: If appropriate, enter the applicable number of the contract or grant under which the report was written.

8b, 8c, & 8d. PROJECT NUMBER: Enter the appropriate military department identification, such as project number, subproject number, system numbers, task number, etc.

9a. ORIGINATOR'S REPORT NUMBER(S): Enter the official report number by which the document will be identified and controlled by the originating activity. This number must be unique to this report.

9b. OTHER REPORT NUMBER(S): If the report has been assigned any other report numbers (*either by the originator or by the sponsor*), also enter this number(s).

10. AVAILABILITY/LIMITATION NOTICES: Enter any limitations on further dissemination of the report, other than those imposed by security classification, using standard statements such as:

- (1) "Qualified requesters may obtain copies of this report from DDC."
- (2) "Foreign announcement and dissemination of this report by DDC is not authorized."
- (3) "U. S. Government agencies may obtain copies of this report directly from DDC. Other qualified DDC users shall request through _____."
- (4) "U. S. military agencies may obtain copies of this report directly from DDC. Other qualified users shall request through _____."
- (5) "All distribution of this report is controlled. Qualified DDC users shall request through _____."

If the report has been furnished to the Office of Technical Services, Department of Commerce, for sale to the public, indicate this fact and enter the price, if known.

11. SUPPLEMENTARY NOTES: Use for additional explanatory notes.

12. SPONSORING MILITARY ACTIVITY: Enter the name of the departmental project office or laboratory sponsoring (*paying for*) the research and development. Include address.

13. ABSTRACT: Enter an abstract giving a brief and factual summary of the document indicative of the report, even though it may also appear elsewhere in the body of the technical report. If additional space is required, a continuation sheet shall be attached.

It is highly desirable that the abstract of classified reports be unclassified. Each paragraph of the abstract shall end with an indication of the military security classification of the information in the paragraph, represented as (TS), (S), (C), or (U).

There is no limitation on the length of the abstract. However, the suggested length is from 150 to 225 words.

14. KEY WORDS: Key words are technically meaningful terms or short phrases that characterize a report and may be used as index entries for cataloging the report. Key words must be selected so that no security classification is required. Identifiers, such as equipment model designation, trade name, military project code name, geographic location, may be used as key words but will be followed by an indication of technical context. The assignment of links, rules, and weights is optional.

UNCLASSIFIED
Security Classification