

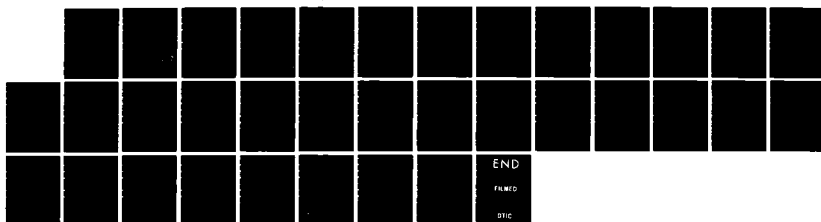
AD-A156 122

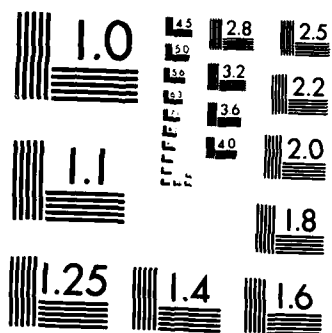
PRUNED ERROR-TRELLIS DECODING OF CERTAIN NON-SYSTEMATIC 1/1
CONVOLUTIONAL CODES(U) ADAPTIVE SENSORS INC SANTA
MONICA CA I S REED 31 DEC 84 N00014-84-C-0720

UNCLASSIFIED

F/G 9/4

NL





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

ASI

Adaptive Sensors, Incorporated

12

PRUNED ERROR-TRELLIS DECODING

OF CERTAIN NON-SYSTEMATIC CONVOLUTIONAL CODES

I. S. Reed

31 December 1984

AD-A156 122

First Quarterly Report

Submitted to

The Office of Naval Research

Arlington, VA 22219

Under Contract # N00014-84-C-0720

by

ADAPTIVE SENSORS, INCORPORATED

216 Pico Boulevard

Santa Monica, CA 90405

DTIC
ELECTE
JUN 27 1985
S B

DTIC FILE COPY

DISTRIBUTION STATEMENT A
Approved for public release
Distribution Unlimited

85 6 4 072

PRUNED ERROR-TRELLIS DECODING OF CERTAIN NON-SYSTEMATIC CONVOLUTIONAL CODES

I. S. Reed

I. INTRODUCTION

The previous works on syndrome decoding of convolutional codes (CCs), e.g., see Ref. 1, led to what is called error-trellis decoding [2]. Since the ending of the previous NAVAIR contract (Contract Number N00019-83-C-0075) in March 1984 and the beginning of the present contract, considerable progress has been made in the concept of error-trellis decoding. It was found, using the fact that CCs are capable of correcting only t errors in some multiple of the constraint length, that the new algorithm requires generally only a reduced or "pruned" trellis. In other words, it was shown that the finite error-correcting capability of the CC makes possible a pruning of sometimes a substantial number of states or paths in a constraint length of the error trellis.

Quantitative formulas for the extent of the primary process for all systematic CCs and certain non-systematic CCs were found recently [3] in collaboration with J. M. Jensen, a visiting scholar at the University of Southern California. In this report, the underlying theory and derivation of these results are given and applied to the special case, the dual-K CCs. The non-systematic dual-K CCs were first introduced by Viterbi and Odenwalder [4]. Such a code is non-binary and, for some applications, a possible substitute for a Reed-Solomon block code.

It will be shown that a CC with a pruned error trellis can be decoded using a modification of either the Viterbi or sequential decoding algorithm

with often a significant reduction in complexity. It is expected that decoding with a pruned error trellis may make it possible to decode CCs--and, in particular, the high-rate CCs--with a larger error-correcting capability than heretofore considered practical.

In order that this report may stand alone, the properties of the CCs required in the report are reviewed briefly in this introductory section. In Section II, the general techniques developed previously [2] for error-trellis decoding are summarized and a set, E , of error sequences which can be decoded by error-trellis decoding is defined. In Section III, a specific procedure is found for pruning the error trellis of a certain class of non-systematic CCs, which includes all systematic CCs and the dual-K CC. Formulas for the number of resulting states and transitions after pruning are found for any CC with free distance, d_{free} , of this sub-class of CCs. The results are applied to the dual-K CC in Section IV.

Let the information or message sequence, the input to the CC, be represented by

$$\underline{x}(D) = [x_1(D), \dots, x_k(D)] \quad , \quad (1a)$$

where

$$x_i(D) = \sum_{j=0}^{\infty} x_{ji} D^j \quad (1b)$$

for $1 \leq j \leq k$ are elements in $F[D]$, the ring of polynomials in the unit delay operator D over $F = GF(q)$, a Galois field, with q a power of a prime integer. Vector $\underline{x}(D)$ is a generating function in D of the input message sequence $\underline{x} = [\underline{x}_0, \dots, \underline{x}_j, \dots]$, where $\underline{x}_j = [x_{j1}, \dots, x_{jk}]$ is a vector belonging to $V_k(F)$, the k -dimensional vector space over F . $\underline{x}(D)$ is sometimes

☒	
☐	
☐	
PER LETTER	
ides	
or	



A-1

called a D-transform of the message or information sequence $\underline{x}$. The k component vector $\underline{x}_j$ in $\underline{x}$ is called the information frame at stage or frame time j .

In a similar manner, the output sequence is

$$\underline{y}(D) = [y_1(D), \dots, y_n(D)] , \quad (2)$$

where $y_i(D) \in F[D]$. Vector $\underline{y}(D)$ is the D-transform of output coded sequence $\underline{y} = [\underline{y}_0, \dots, \underline{y}_j, \dots]$, where $\underline{y}_j = [y_{j1}, \dots, y_{jk}]$ belongs to $V_n(F)$. The n -vector $\underline{y}_j$ is called the j -th codeword frame of code sequence $\underline{y}$.

The information and code sequences of an (n, k) convolutional code are linearly related by a $k \times n$, rank k , generator matrix $G(D)$ of polynomial elements in $F[D]$, as follows:

$$\underline{y}(D) = \underline{x}(D) G(D) . \quad (3)$$

The maximum degree m of the polynomial elements of $G(D)$ in D is called the memory, and the constraint length L is defined as $L = m + 1$.

The free distance of a CC is defined by

$$d_{\text{free}} = \min_{\underline{y}(D) \neq 0} W_H(\underline{y}(D)) , \quad (4)$$

where $W_H(\underline{y}(D))$ is the cumulative Hamming weight of the coefficients y_j of D^j for all j , $0 \leq j$, where $\underline{y}_j$ is the j -th codeword frame. Note that the computation of d_{free} requires at least L codeword frames for all codes of practical interest.

To avoid catastrophic error propagation, $G(D)$ is assumed to be a basic encoder [5]. The Smith normal form of a basic encoder [2] is

$$G(D) = A(D) \begin{bmatrix} I_k, 0 \end{bmatrix} B(D) , \quad (5)$$

where $A(D)$ and $B(D)$ are, respectively, $k \times k$ and $n \times n$ invertible matrices over $F[D]$ and I_k is a $k \times k$ identity matrix.

In Eq. (4), let matrix $B(D)$ be partitioned as

$$B(D) = \left[B_1(D)^T, B_2(D)^T \right]^T,$$

where $B_1(D)$ consists of the first k rows of $B(D)$ and "T" denotes matrix transpose. Similarly, let

$$B(D)^{-1} = \left[\bar{B}_1(D), \bar{B}_2(D) \right],$$

where $\bar{B}_1(D)$ consists of the first k columns of $B(D)^{-1}$. Since $B(D) \cdot B(D)^{-1} = I_n$, the following identities evidently hold:

$$\begin{aligned} B_1(D) \cdot \bar{B}_1(D) &= I_k, & B_1(D) \cdot \bar{B}_2(D) &= 0 \\ B_2(D) \cdot \bar{B}_1(D) &= 0, & B_2(D) \cdot \bar{B}_2(D) &= I_{n-k}. \end{aligned} \quad (6)$$

A parity-check matrix $H(D)$ is an $(n - k) \times n$ matrix of rank $(n - k)$, satisfying

$$G(D) \cdot H^T(D) = 0. \quad (7)$$

From Eqs. (5), (6), and (7), it is seen next that

$$H(D) = \bar{B}_2(D)^T \quad (8)$$

has the properties of a parity-check matrix $H(D)$ associated with $G(D)$.

By Eq. (3), the CC generated by $G(D)$ is the set

$$C = \left\{ \underline{y}(D) = [y_1(D), \dots, y_n(D)] \mid \underline{y}(D) = \underline{x}(D) G(D) \right\}. \quad (9)$$

It is now shown also that

$$C = \left\{ \underline{y}(D) = [y_1(D), \dots, y_n(D)] \mid \underline{y}(D) H^T(D) = 0 \right\}, \quad (10)$$

where $H(D)$ is given in Eq. (8). To see this, denote the right side of Eq. (10) by C_H . Clearly, an element of C , as given in Eq. (9), belongs to C_H , and hence $C \subseteq C_H$.

Next, suppose $y_1(D)$ is an element of C_H , i.e., by Eqs. (8) and (10),

$$y_1(D) H^T(D) = y_1(D) \bar{B}_2(D) = 0 .$$

But, by definition, $\bar{B}_2(D)$ consists of the last $(n - k)$ columns of $B(D)^{-1}$, so that

$$\bar{B}_2(D) = B(D)^{-1} \begin{bmatrix} 0 \\ I_{n-k} \end{bmatrix} . \quad (11)$$

where "0" denotes a block of k rows of zeros and I_{n-k} is the $(n - k)$ row identity matrix. Thus, $y_1(D)$ satisfies the equation

$$y_1(D) B^{-1}(D) \begin{bmatrix} 0 \\ I_{n-k} \end{bmatrix} = 0 .$$

The most general solution of this equation for $y_1(D) B^{-1}(D)$ is

$$y_1(D) B^{-1}(D) = [\tau_1(D), \dots, \tau_k(D), 0, \dots, 0] = [\underline{\tau}(D), 0] ,$$

where $\tau_j(D)$ for $1 \leq j \leq k$ can be chosen to be any arbitrary elements of $F[D]$.

Solving for $y_1(D)$ yields, finally, by Eq. (5),

$$y_1(D) = \underline{\tau}(D) [I_k, 0] B(D) = \underline{\tau}(D) A^{-1}(D) G(D) ,$$

which belongs to C , as given in Eq. (9). Thus, $C_H \subseteq C$ and Eq. (10) is proved.

The fact that the CC given by set C in Eq. (9) can be characterized by Eq. (10) is used in the following section to find the coset of solutions

to the syndrome equation. It is also used to find another representation of this coset needed in error-trellis decoding.

II. ERROR-TRELLIS DECODING

Let $\underline{y}(D)$ in Eq. (3) be transmitted and $\underline{z}(D)$ be received. Then,

$$\underline{z}(D) = \underline{y}(D) + \underline{e}(D) , \quad (12)$$

where $\underline{e}(D)$ is the D-transform of the error sequence. By Eqs. (12) and (7), the syndrome of the received sequence is

$$\begin{aligned} \underline{s}(D) &= \underline{z}(D) \cdot H^T(D) = [\underline{y}(D) + \underline{e}(D)] \cdot H^T(D) \\ &= \underline{e}(D) \cdot H^T(D) . \end{aligned} \quad (13)$$

The problem of syndrome decoding is, given $\underline{s}(D) = \underline{z}(D) \cdot H^T(D)$, to solve the syndrome equation

$$\underline{s}(D) = \underline{z}(D) H^T(D) = \underline{e}(D) H^T(D) , \quad (14a)$$

or its equivalent,

$$(\underline{e}(D) - \underline{z}(D)) H^T(D) = 0 , \quad (14b)$$

for all solutions $\underline{e}(D)$.

By Eqs. (10) and (9), the term $(\underline{e}(D) - \underline{z}(D))$ in Eq. (14b) must be some code sequence $\underline{v}(D) G(D)$. Hence, the most general solution of the syndrome equation, Eq. (14a), is

$$\underline{e}(D) = \underline{z}(D) + \underline{v}(D) G(D) , \quad (15)$$

where $\underline{v}(D)$ is the D-transform of an arbitrary message-like sequence $\underline{v} = [\underline{v}_0, \dots, \underline{v}_j, \dots]$ of k-vectors $\underline{v}_j \in V_k(F)$.

Equation (15) shows that the most general solution of the syndrome equation, Eq. (14a), for $\underline{e}(D)$ is the coset

$$C_z = \left\{ \underline{e}(D) = \underline{z}(D) + v(D) G(D) \mid \underline{v}(D) = [v_1(D), \dots, v_k(D)] \right\} \quad (16)$$

of code C , defined by either Eq. (9) or Eq. (10). A minimization of the Hamming weights over all elements of coset C_z yields the standard minimum-error solution for message $\underline{v}(D)$. Efficient methods for achieving this minimization include the Viterbi algorithm and all sequential decoding methods for convolutional codes.

The difficulty with the standard decoding methods of CCs, i.e., Viterbi or sequential decoding, is the need to consider a sometimes prohibitively large number of states and paths in the decoding trellis. Such minimum-weight, path-finding decoding methods do not take advantage of the limited error-correcting capability which one might expect could reduce the number of paths in the trellis over which this minimum is taken. One method which allows for such a reduction in the number of paths in the trellis is to use another equivalent solution of the syndrome equation, Eq. (14), which is independent of the transmitted codeword.

Another solution of Eq. (14) is given by

$$\underline{e}(D) = \underline{u}(D) G(D) + \underline{z}(D) R(D) , \quad (17a)$$

where

$$R(D) = \overline{B}_2(D) B_2(D) \quad (17b)$$

is an $n \times n$ matrix of rank $(n - k)$ and where matrices $\overline{B}_2(D)$, $B_2(D)$ are defined in Eq. (6). To prove Eq. (17) is a general solution of the syndrome equation, note, by Eqs. (17b) and (6), that

$$\begin{aligned} \left(\underline{z}(D) - \underline{z}(D) R(D) \right) H^T &= \underline{z}(D) \overline{B}_2(D) - \underline{z}(D) \overline{B}_2(D) B_2(D) \overline{B}_2(D) \\ &= \underline{z}(D) \overline{B}_2(D) - \underline{z}(D) \overline{B}_2(D) I_{n-k} = 0 . \end{aligned}$$

Hence, by Eq. (10), $\underline{z}(D) - \underline{z}(D) R(D)$ is a possible codeword, say $\underline{x}_1(D) G(D)$, where $\underline{x}_1(D)$ is some k-vector of elements in $F[D]$. That is,

$$\underline{z}(D) - \underline{z}(D) R(D) = \underline{x}_1(D) G(D)$$

or

$$\underline{z}(D) = \underline{x}_1(D) G(D) + \underline{z}(D) R(D)$$

for some information vector $\underline{x}_1(D)$. A substitution of $\underline{z}(D)$ in the above equation into Eq. (15) yields

$$\begin{aligned} \underline{e} &= (\underline{x}_1(D) + \underline{v}(D)) G(D) + \underline{z}(D) R(D) \\ &= \underline{u}(D) G(D) + \underline{z}(D) R(D) \end{aligned}$$

as another general solution of the syndrome equation, where $\underline{u}(D)$ is an arbitrary k-vector of elements in $F[D]$. Hence, Eq. (17) is established.

The solution, Eq. (17), of the syndrome equation, Eq. (15), has the desirable property that it is independent of the transmitted codeword $\underline{y}(D)$. To see this, let $\underline{e}_a(D)$ replace $\underline{e}(D)$ in Eq. (17) as the actual error sequence. Then, a substitution of Eq. (12) into Eq. (17) yields, by Eq. (10),

$$\begin{aligned} \underline{e}(D) &= \underline{u}(D) G(D) + (\underline{y}(D) + \underline{e}_a(D)) H^T(D) B_2(D) \\ &= \underline{u}(D) G(D) + \underline{e}_a(D) R(D) , \end{aligned}$$

which is independent of the transmitted message $\underline{y}(D)$.

By the maximum likelihood principle, the most likely error sequence $\hat{\underline{e}}(D)$ is the one with minimum Hamming weight. That is,

$$W_H(\hat{\underline{e}}(D)) = \min_{\underline{u}(D)} W_H(\underline{u}(D) G(D) + \underline{z}(D) R(D)) \quad (18a)$$

with

$$\hat{\underline{e}}(D) = \hat{\underline{u}}(D) G(D) + \underline{z}(D) R(D) , \quad (18b)$$

where the minimization is taken over all k -vectors $\underline{u}(D)$ over $F[D]$ and $\hat{\underline{u}}(D)$ is some k -vector which achieves this minimization.

It is seen, from Eq. (17), that all possible error sequences are obtained by adding the sequence $\underline{z}(D) R(D)$ to the code tree or trellis C in Eq. (9), presented graphically. This new trellis is called an error trellis. Thus, the most likely error sequence $\hat{\underline{e}}(D)$ in Eq. (18) equals also the minimum weight path in the error trellis. Hence, by Eq. (12), a subtraction of $\hat{\underline{e}}(D)$ from $\underline{z}(D)$ produces the best estimate $\hat{\underline{y}}(D) = \underline{z}(D) - \hat{\underline{e}}(D)$ of the original transmitted codeword $\underline{y}(D)$. Therefore, the best estimate $\hat{\underline{x}}(D)$ of the original information sequence $\underline{x}(D)$ is

$$\begin{aligned} \hat{\underline{x}}(D) &= \hat{\underline{y}}(D) G(D)^{-1} \\ &= [\underline{z}(D) - \hat{\underline{u}}(D) G(D) - \underline{z}(D) R(D)] G(D)^{-1} \\ &= \underline{z}(D) G(D)^{-1} - \hat{\underline{u}}(D) . \end{aligned} \quad (19)$$

The last inequality follows, using Eqs. (5), (17b), and (6), from

$$\begin{aligned} R(D) G(D)^{-1} &= \bar{B}_2(D) B_2(D) B^{-1}(D) \begin{bmatrix} I_k \\ 0 \end{bmatrix} A^{-1}(D) \\ &= \bar{B}_2(D) B_2(D) [\bar{B}_1(D), \bar{B}_2(D)] \begin{bmatrix} I_k \\ 0 \end{bmatrix} A^{-1}(D) \\ &= \bar{B}_2(D) [0, I_{n-k}] \begin{bmatrix} I_k \\ 0 \end{bmatrix} A^{-1}(D) = 0 . \end{aligned}$$

This important identity shows that $\hat{\underline{u}}(D)$, obtained in the minimization in Eq. (18), is the "best" correction factor for the received information

sequence, $\underline{z}(D) G^{-1}(D)$.

Again, let $\underline{e}_a(D)$ be the actual error sequence. Then, a substitution of $\underline{z}(D) = \underline{y}(D) + \underline{e}_a(D)$ into Eq. (19) yields

$$\begin{aligned} \hat{x}(D) &= [\underline{y}(D) + \underline{e}_a(D)] G^{-1}(D) - \hat{u}(D) \\ &= \underline{x}(D) + \underline{e}_a(D) G^{-1}(D) - \hat{u}(D), \end{aligned} \quad (20)$$

where $\underline{x}(D)$ is the original information or message sequence in Eq. (1).

Now, define E to be the set of all error sequences which can be corrected by Viterbi or error-trellis decoding, as in Eq. (18). Then, by Eq. (20), if $\underline{e}_a(D) \in E$, the estimate $\hat{u}(D)$ found by the minimization in Eq. (18) equals $\underline{e}_a(D)$, so that $\hat{u}(D) = \underline{e}_a(D) G^{-1}(D)$. Thus, the minimization in Eq. (18) needs only be taken over those sequences $\underline{u}(D)$ which belong to the set

$$E^{(-1)} = \left\{ \underline{u}(D) = \underline{e}(D) G^{-1}(D) \mid \underline{e}(D) \in E \right\}. \quad (21)$$

Hence, the most likely error sequence $\hat{e}(D)$ is found by

$$W_H(\hat{e}(D)) = \min_{\underline{u}(D) \in E^{(-1)}} W_H(\underline{u}(D) G(D) + \underline{z}(D) R(D)) \quad (22a)$$

with

$$\hat{e}(D) = \hat{u}(D) G(D) + \underline{z}(D) R(D), \quad (22b)$$

where $E^{(-1)}$ is the set of k -vectors, defined in Eq. (21).

Note that, if $\underline{e}_a(D) \in E$, then the most likely sequence found by Eq. (22) is equal to $\underline{e}_a(D)$, the actual error sequence. If $\underline{e}_a(D) \notin E$, a decoding error will be made in both Eq. (18) and Eq. (22).

In order to actually perform the minimization in Eq. (22) over the set $E^{(-1)}$, the sets E and $E^{(-1)}$ must be identified. This is generally impossible.

However, as will be shown in the following section for a certain restricted class of non-systematic CCs, a reasonable approximation to sets E and $E^{(-1)}$ can be found.

III. ERROR-TRELLIS DECODING OF CERTAIN NON-SYSTEMATIC CCs

To find a reasonable approximation to $E^{(-1)}$ in Eqs. (21) and (22), the set in which the most likely sequence $\hat{u}(D)$ can be found for message correction, consider first the error trellis in more detail. By Eq. (17), the error trellis is a coding trellis generated by the term $\underline{u}(D) G(D)$ plus the term $\underline{z}(D) R(D)^{-1}$, which depends on the received coded sequence.

Thus, the underlying coding trellis of the error trellis is

$$\underline{w}(D) = \underline{e}(D) - \underline{z}(D) R(D) = \underline{u}(D) G(D) ,$$

which can be conceived to be a linear sequential circuit with input $\underline{u}_j$ at frame j , where $\underline{u}_j = [u_{j1}, \dots, u_{jk}]$ is the j -th, k -vector, coefficient of the formal power series for $\underline{u}(D)$, i.e.

$$\underline{u}(D) = \sum_{j=0}^{\infty} \underline{u}_j D^j .$$

By expanding $\underline{u}(D)$, $G(D)$, and $\underline{w}(D)$ into their formal power series and equating coefficients,

$$\underline{w}_j = F(\underline{u}_j, \underline{u}_{j-1}, \dots, \underline{u}_{j-m}), \quad (23)$$

for $j \geq 0$ is the output of the coding trellis at frame j in terms of input $\underline{u}_j$ and the m previous values of $\underline{u}_j$, where initially $\underline{u}_j = 0$ for all $j < 0$. Evidently, function F is linear in each of its k -vector components so that $\underline{w}_0 = 0$, the zero n -vector.

where

$$G_0 = [1, 1, 1, \dots, 1] \text{ and}$$

$$G_1 = [g_{11}, g_{12}, \dots, g_{1n}] ,$$

with $g_{1i} \neq 0$ and $g_{1j} \in GF(2^K)$ for $(j = 1, 2, \dots, n)$.

From the above definition of a dual-K CC, the minimum distance of the code is $d = 2n$. Evidently also, the free distance is

$$d_{\text{free}} = d = 2n .$$

Hence, if no more than t symbol errors occur in the first two codeword frames, and

$$2t + 1 \leq d = 2n \text{ or}$$

$$t < n - 1/2 ,$$

then those errors which occur in the first frame can be corrected. In other words, the dual-K CC is a t-error-per-constraint length-correcting CC, where $t = [n - 1/2]$.

Consider now the example given by Odenwalder [4, Fig. 1].

Example: Although he does not say so explicitly, Odenwalder uses the Galois field $GF(2^3)$ generated by the polynomial $x^3 + x^2 + 1$, irreducible over $GF(2)$. If α is a root of this polynomial, then $\alpha, \alpha^2, \alpha^3, \alpha^4, \alpha^5, \alpha^6, \alpha^7 = 1$, and 0 are the eight elements of $GF(2^3)$. The representation of the elements as polynomials in α is given in Table 1.

The convolutional encoder of a rate 1/2, dual-3, CC is shown in Fig. 1. This is a more abstract version of the same encoder given by Odenwalder [4, Fig. 1].

$$B(D)^{-1} = \begin{bmatrix} I_k & -Q(D) \\ 0 & I_{n-k} \end{bmatrix},$$

and, as a consequence, $G(D)^{-1}$ is delay free and of form, Eq. (47), so that $\theta = 1$.

The number of states and transitions derived above in theorem 2 can be shown to be polynomial in mk . In fact, if $mk \gg t$, then it can be proved that

$$S(q, k, m, \theta t) \approx [(q - 1) e^{mk/\theta t}]^{\theta t}$$

and

$$T(q, k, m, \theta t) \approx [(q - 1) e^{(m+1)k/\theta t}]^{\theta t},$$

are the number of states and transitions per frame time required for error trellis decoding of a systematic CC. For standard Viterbi decoding, these numbers are exponential of form q^{mk} and $q^{(m+1)k}$, respectively. These results show that the efficiency of error-trellis decoding improves dramatically both with increased rate and mk over standard Viterbi decoding.

IV. PRUNED ERROR-TRELLIS DECODING OF DUAL-K CONVOLUTIONAL CODES

Dual-K convolutional codes are $(n, 1)$ CCs of rate $1/n$, of memory $m = 1$, and with symbols in the finite or Galois field $GF(2^K)$. See Odenwalder's paper [4]. The generating matrix $\underline{G}$ is

$$\underline{G} = \begin{bmatrix} G_0 & G_1 & & & \\ & G_0 & G_1 & & \\ & & G_0 & G_1 & \\ & & & \ddots & \ddots \\ & & & & \ddots & \ddots \end{bmatrix},$$

Theorem 2 shows that non-systematic (n, k) CCs with delay-free inverse generating matrices allow for the possibility of trellis pruning. In general, the smallest reduction factor is obtained for such codes with $\theta \equiv 1$ or inverse matrices of form

$$G^{-1} = \begin{bmatrix} I_k \\ p_{k+1} \\ \vdots \\ p_n \end{bmatrix}, \quad (47a)$$

where

$$0 \leq w_H(p_i) \leq 1, \text{ for } k+1 \leq i \leq n. \quad (47b)$$

Three classes of codes achieve $\theta = 1$ for greatest error trellis pruning.

They are:

- (i) All (n, k) systematic CCs.
- (ii) All $(n, 1)$ non-systematic CCs with delay-free inverse matrices.
- (iii) All (n, k) non-systematic CCs with delay free inverse matrices, where Eq. (47) is true.

All other CCs with a delay-free inverse generating matrix require $1 < \theta \leq k$.

To show (i), above, observe that all (n, k) systematic CCs have a generating matrix of form $G(D) = [I_k, Q(D)]$, and, as a consequence, a Smith normal form, $G(D) = [I_k, 0] B(D)$, where

$$B(D) = \begin{bmatrix} I_k, Q(D) \\ 0, I_{n-k} \end{bmatrix}.$$

Hence, $G(D)^{-1} = B(D)^{-1} [I_k, 0]^T = [I_k, 0]^T$, since

states S_j need to remain in the error trellis which have Hamming weight of at most θt . Since there are exactly $\binom{m+k}{j} (q-1)^j$ states S_j of weight j for $0 \leq j \leq \min(\theta t, m+k)$, where $\binom{n}{j}$ denotes the binomial coefficient, the first part of the theorem is proved.

Consider now the "change of state" equation, Eq. (25b), of the coding trellis, associated with the error trellis in Eq. (17). By Eqs. (25b), (28), and (30), a transition in the error trellis from state

$$S_j = (\sigma_1, \dots, \sigma_m)$$

at frame j yields, upon input $\underline{u}_j$, the next state,

$$S_{j+1} = (\sigma', \sigma_1, \dots, \sigma_{m-1}),$$

where $\sigma' = \underline{u}_j$ is the value of the input $\underline{u}_j$ at frame j . Thus, the transition is determined uniquely by input σ' and state $(\sigma_1, \dots, \sigma_m)$ or by the $(m+1)$ -tuple,

$$(\sigma', \sigma_1, \dots, \sigma_{m-1}, \sigma_m).$$

But, by Eq. (26), this $(m+1)$ -tuple or its equivalent, the transition, is equal to $(m+1) = L$ consecutive frames of the sequence $\underline{u}(D)$, namely,

$$(\sigma', \sigma_1, \dots, \sigma_m) = (\underline{u}_j, \underline{u}_{j-1}, \dots, \underline{u}_{j-m}),$$

where $\underline{u}_j = 0$ for $j < 0$. As in the first part of the theorem, only those sequences which are in $\Sigma_{m, \theta}^{-1}$ need to be considered. Therefore, a branching or transition from state $S_j = (\sigma_1, \dots, \sigma_m)$ to state $S_{j+1} = (\sigma', \sigma_1, \dots, \sigma_{m-1})$ needs to be in the error trellis if and only if $w_H(\sigma', \sigma_1, \dots, \sigma_m) \leq t$. Hence, the total number of transitions at frame j is given by Eq. (46) for $j > m$, and the theorem is proved.

If a probability measure were imposed on the set of all possible error sequences, it is probable that set E_1 would approximate E very closely in probability. Hence, for such cases, the performance of a Viterbi decoder to realize Eq. (44) would differ very little from the hypothesized generalized feedback decoder.

Realizing Eq. (44) with a Viterbi-like or sequential decoder has the advantage over standard Viterbi or sequential decoding in that it allows for a reduction of both the number of states and transitions in the error trellis. The following theorem quantifies the reductions achieved in the "pruned" trellis. The techniques for pruning the trellis are given in the proofs of the theorem.

Theorem 2. The number of states in the pruned error trellis of a q -ary (n, k) CC with memory m and a delay-free G^{-1} is

$$S(q, k, m, \theta t) = \sum_{j=0}^{\alpha} \binom{m+k}{j} (q-1)^j, \quad (45)$$

where $\alpha = \min(\theta t, m+k)$ and $t = \lceil (d_{\text{free}} - 1)/2 \rceil$. Also, for the same CC, the number of transitions in the pruned error trellis is

$$T(q, k, m, \theta t) = \sum_{j=0}^{\beta} \binom{(m+1)+k}{j} (q-1)^j, \quad (46)$$

where $\beta = \min[\theta t, (m+1)+k]$.

Proof: The minimization in Eq. (44) for error-trellis decoding is taken only over those sequences $\underline{u}(D)$ which belong to $\Sigma_{m, \theta}^{-1}$. That is, by Eq. (40), the minimization is taken only over those sequences $\underline{u}(D)$ which have Hamming weight of, at most, θt in every $L = m+1$ consecutive frames, i.e., $W_H(\underline{u}_j, \dots, \underline{u}_{j+m}) \leq \theta t$ for all $j \geq 0$. Hence, by Eq. (30), only those

For the other part, assume $\underline{u}(D) \in \Sigma_{m,1}^{(-1)}$. Then, $\underline{v}(D) = [\underline{u}(D), 0, \dots, 0] \in \Sigma_m$. But this implies that

$$\underline{v}(D) \cdot G^{-1} = [\underline{u}(D), 0, \dots, 0] \begin{bmatrix} I \\ p \end{bmatrix} = \underline{u}(D) \in \Sigma_m^{(-1)}.$$

Hence, $\Sigma_{m,1}^{(-1)} \subseteq \Sigma_m^{(-1)}$ and the theorem is proved.

Assume for the moment that the minimization in Eq. (22) is accomplished by generalized feedback decoding, then, by Eqs. (22) and (35), the most likely error sequence is found by

$$W_H(\hat{e}(D)) = \min_{\underline{u}(D) \in \Sigma_j^{(-1)}} W_H(\underline{u}(D) G(D) + \underline{z}(D) R(D)) \quad (43a)$$

with

$$\hat{e}(D) = \underline{u}(D) G(D) + \underline{z}(D) R(D). \quad (43b)$$

But, in Eq. (43a), the same estimate of $\hat{e}(D)$ is obtained if the minimization is taken over any set which includes $\Sigma_j^{(-1)}$. Hence, since $\Sigma_j^{-1} \subseteq \Sigma_m^{-1}$, by Eq. (37), and $\Sigma_m^{-1} \subseteq \Sigma_{m,\theta}^{-1}$, it is sufficient to take the minimization in Eq. (43a) only over $\Sigma_{m,\theta}^{-1}$. Thus, the minimum error sequence $\hat{e}(D)$ in Eq. (43a) can be found, for all convolutional codes for which $G^{-1}(D)$ is delay free, by

$$W_H(\hat{e}(D)) = \min_{\underline{u}(D) \in \Sigma_{m,\theta}^{-1}} W_H(\underline{u}(D) G(D) + \underline{z}(D) R(D)) \quad (44a)$$

with

$$\hat{e}(D) = \hat{\underline{u}}(D) G(D) + \underline{z}(D) R(D), \quad (44b)$$

where $\hat{\underline{u}}(D)$ is a "best" message correction factor.

$$\begin{aligned}
 W_H(\underline{u}_j, \dots, \underline{u}_{j+m}) &\leq W_H\left(\left[\underline{v}_j^{(k)}, \dots, \underline{v}_{j+m}^{(k)}\right], \right. \\
 &\quad \left. \left[v_{j,k+1} \cdot \underline{p}_{k+1}, \dots, v_{j+m,k+1} \cdot \underline{p}_{k+1}\right], \right. \\
 &\quad \left. \left[v_{j,n} \cdot \underline{p}_n, \dots, v_{j+m,n} \cdot \underline{p}_n\right]\right) \\
 &= W_H\left(\underline{v}_j^{(k)}, \dots, \underline{v}_{j+m}^{(k)}\right) + W_H\left(v_{j,k+1} \cdot \underline{p}_{k+1}\right) + \dots + W_H\left(v_{j+m,k+1} \cdot \underline{p}_{k+1}\right) + \dots \\
 &\quad + W_H\left(v_{j,n} \cdot \underline{p}_n\right) + \dots + W_H\left(v_{j+m,n} \cdot \underline{p}_n\right) \\
 &= W_H\left(\underline{v}_j^{(k)}, \dots, \underline{v}_{j+m}^{(k)}\right) + W_H(\underline{p}_{k+1}) \left[W_H(v_{j,k+1}) + \dots + W_H(v_{j+m,k+1}) \right] + \dots \\
 &\quad + W_H(\underline{p}_n) \left[W_H(v_{j,n}) + \dots + W_H(v_{j+m,n}) \right] \\
 &\leq W_H\left(\underline{v}_j^{(k)}, \dots, \underline{v}_{j+m}^{(k)}\right) + \max_{k+1 \leq i \leq n} \left\{ W_H(\underline{p}_i) \right\} \left[W_H(v_{j,k+1}) + \dots \right. \\
 &\quad \left. + W_H(v_{j+m,k+1}) + \dots + W_H(v_{j,n}) + \dots + W_H(v_{j+m,n}) \right] \\
 &\leq \theta \left[W_H\left(\underline{v}_j^{(k)}, \dots, \underline{v}_{j+m}^{(k)}\right) + W_H(v_{j,k+1}) + \dots + W_H(v_{j+m,k+1}) + \dots \right. \\
 &\quad \left. + W_H(v_{j,n}) + \dots + W_H(v_{j+m,n}) \right], \tag{42}
 \end{aligned}$$

where θ is given in Eq. (41). After assembling these components of vectors $\underline{v}_j, \dots, \underline{v}_{j+m}$ within the Hamming weight function, this inequality becomes, finally,

$$\begin{aligned}
 W_H(\underline{u}_j, \dots, \underline{u}_{j+m}) &\leq \theta W_H(\underline{v}_j, \dots, \underline{v}_{j+m}) \\
 &\leq \theta t,
 \end{aligned}$$

where the last inequality follows, from Eq. (33). Hence, by Eq. (40),

$\Sigma_m^{(-1)} \subseteq \Sigma_{m,\theta}^{(-1)}$ and the first part of the theorem is proved.

$$\begin{aligned}\underline{u}(D) &= \sum_{j=0}^{\infty} \underline{u}_j D^j = \underline{v}(D) G^{-1} \\ &= \sum_{j=0}^{\infty} \left(\underline{v}_j G^{-1} \right) D^j.\end{aligned}$$

Thus, using Eq. (39) in lemma,

$$\begin{aligned}W_H(\underline{u}_j, \dots, \underline{u}_{j+m}) &= W_H(\underline{v}_j G^{-1}, \dots, \underline{v}_{j+m} G^{-1}) \\ &= W_H\left(\underline{v}_j^{(k)} + \sum_{i=k+1}^n v_{ji} \underline{p}_i, \dots, \underline{v}_{j+m}^{(k)} + \sum_{i=k+1}^n v_{j+m,i} \underline{p}_i\right),\end{aligned}$$

where $\underline{v}_j^{(k)} = (v_{j1}, \dots, v_{jk})$. This can re-expressed as

$$\begin{aligned}W_H(\underline{u}_j, \dots, \underline{u}_{j+m}) &= W_H\left(\left[\underline{v}_j^{(k)}, \dots, \underline{v}_{j+m}^{(k)}\right] \right. \\ &\quad \left. + \left[v_{j,k+1} \cdot \underline{p}_{k+1}, \dots, v_{j+m,k+1} \cdot \underline{p}_{k+1}\right] \right. \\ &\quad \left. + \right. \\ &\quad \left. \vdots \right. \\ &\quad \left. + \left[v_{j,n} \cdot \underline{p}_n, \dots, v_{j+m,n} \cdot \underline{p}_n\right]\right).\end{aligned}$$

But, since the Hamming weight of a sum of vectors is upper-bounded by the Hamming weight of the concatenation of the same set of vectors, one has

$\begin{bmatrix} I_k \\ P \end{bmatrix}$ where I_k is the $k \times k$ identity matrix and P is an $(n - k) \times k$ matrix of elements in $GF(q)$.

By the above lemma, there is no loss in generality to assume that G^{-1} has the form

$$G^{-1} = \begin{bmatrix} I_k \\ \underline{p}_{k+1} \\ \vdots \\ \underline{p}_n \end{bmatrix}, \quad (39)$$

where $\underline{p}_j = [p_{j1}, \dots, p_{jk}]$ for $k + 1 \leq j \leq n$.

Now define a message-correction cylinder by

$$\begin{aligned} \Sigma_{m,\theta}^{(-1)} &= \{ \underline{u}(D) \\ &= [u_1(D), \dots, u_k(D)] \mid W_H(\underline{u}_j, \dots, \underline{u}_{j+m}) \leq \theta t \text{ for all } j \geq 0 \}, \end{aligned} \quad (40)$$

where

$$\theta = \max_{k+1 \leq i \leq n} \{ W_H(\underline{p}_i), 1 \}. \quad (41)$$

Evidently, $1 \leq \theta \leq k$. The following theorem now relates $\Sigma_m^{(-1)}$ with cylinders $\Sigma_{m,1}^{(-1)}$ and $\Sigma_{m,\theta}^{(-1)}$.

Theorem 1. $\Sigma_{m,1}^{(-1)} \subseteq \Sigma_m^{(-1)} \subseteq \Sigma_{m,\theta}^{(-1)}$, where cylinder $\Sigma_m^{(-1)}$ and $\Sigma_{m,\theta}^{(-1)}$

for $1 \leq \theta \leq k$ are defined in Eqs. (36) and (40), respectively.

Proof: First, to show $\Sigma_m^{(-1)} \subseteq \Sigma_{m,\theta}^{(-1)}$, suppose $\underline{u}(D) = \underline{v}(D) G^{-1} \in \Sigma_m^{(-1)}$, where $\underline{v}(D) \in \Sigma_m$. Then, since G^{-1} is delay free,

$\underline{u}_j, \underline{u}_{j-1}, \dots, \underline{u}_{j-m}$. Remember, by Eq. (36), that $\underline{u}(D) = \underline{v}(D) G^{-1}(D)$. Thus, in the special case that $G^{-1}(D)$ is delay-free of form

$$G^{-1} = \begin{bmatrix} \underline{g}_1 \\ \vdots \\ \underline{g}_n \end{bmatrix}, \quad \underline{g}_i = [g_{i1}, \dots, g_{ik}], \quad (38)$$

where $G_{ij} \in GF(q)$, one has

$$\underline{u}_j = \sum_{i=1}^n v_{ji} \underline{g}_i,$$

so that the k -vectors $\underline{u}_j$ are linear functions of the components of $\underline{v}_j$ for $(j = 0, 1, \dots, n)$. Hence, a CC for which G^{-1} is delay-free has the property that the states of the error trellis are dependent only on $m+1$ successive frames, $\underline{v}_j, \underline{v}_{j-1}, \dots, \underline{v}_{j-m}$, of error vectors. Therefore, in this case, by Eq. (33), one needs only consider the error cylinder Σ_m and its correspondent $\Sigma_m^{(-1)}$ when endeavoring to prune states from the error trellis.

Assume now that G^{-1} is delay-free, and that G^{-1} is an $n \times k$ matrix over $GF(q)$. Since a code sequence $\underline{y}(D) = [y_1(D), \dots, y_n(D)]$ has the same distance properties as a code $\underline{y}'(D) = [y_1'(D), \dots, y_n'(D)]$, where $y_j'(D) = y_{\pi(j)}(D)$, where $\pi(j)$ denotes a permutation of the integers $(j = 1, \dots, n)$, it is natural to call code $\underline{y}'(D)$ permutation equivalent to $\underline{y}(D)$.

Definition 1. Two codes are permutation equivalent if one can be obtained from the other by the same permutation of places or coordinates in all frames. The following lemma can be established for permutation equivalent CCs.

Lemma: Let C_1 be a CC with a delay-free inverse G^{-1} . Then C_1 is permutation equivalent to a code C_2 with a delay-free inverse of form

However, suppose the error trellis or its equivalent, a Viterbi decoder, is restricted to be a generalized feedback decoder wherein the correction of the ℓ -th frame is decision dependent on only the ℓ -th and the preceding J frames of data, namely $(\underline{z}_\ell, \dots, \underline{z}_{\ell-J})$ for $(\ell = 0, 1, \dots)$. Then, by construction, such a generalized feedback decoder has the property that it will correct any error sequence belonging to Σ_J . Hence, if set E_1 is defined to be the set of all error sequences which can be corrected by a generalized feedback decoder, using error trellis decoding or its equivalent, Viterbi decoding, with correction delayed by J frames, then

$$\Sigma_J = E_1. \quad (35)$$

Evidently, set E_1 is a reasonable approximation to set E .

The restriction of set E to its approximation E_1 makes it possible, by the definition of a state in Eq. (30), to often reduce the number of states as well as transitions needed for the error trellis. To show this, define first, by analogy with set $E^{(-1)}$ in Eq. (21), the sets

$$\Sigma_j^{(-1)} = \left\{ \underline{u}(D) = \underline{v}(D) G^{-1}(D) \mid \underline{v}(D) \in \Sigma_j \right\}, \quad (36)$$

corresponding to Σ_j in Eq. (33) for $0 \leq j \leq J$. These sets, associated with the cylinders Σ_j , also form a family of non-increasing sets, i.e.,

$$\Sigma_m^{(-1)} \supseteq \Sigma_{m+1}^{(-1)} \supseteq \dots \Sigma_J^{(-1)} = E_1^{(-1)}. \quad (37)$$

To prove this, suppose $\underline{u}(D) \in \Sigma_r^{(-1)}$, where $m \leq j \leq r \leq J$. Then, by Eqs. (36) and (34), $\underline{u}(D) = \underline{v}(D) G^{-1}(D)$, where $\underline{v}(D) \in \Sigma_r \subseteq \Sigma_j$ for $m \leq j \leq r$. Hence, by Eq. (36), $\underline{u}(D) \in \Sigma_j^{(-1)}$, so that $\Sigma_j \supseteq \Sigma_r$ and Eq. (37) is established.

By Eq. (30), the state of the error trellis depends only on one constraint length or $L = m + 1$ successive values of $\underline{u}_j$, namely,

For most codes of practical interest, $J \geq m$.

Assume, as in Eq. (12), that $\underline{e}(D)$ is a possible error sequence and define

$$\Sigma_j = \left\{ \underline{e}(D) \mid W_H(\underline{e}_\ell, \dots, \underline{e}_{\ell+j}) \leq t \text{ for all } \ell \geq 0 \right\}, \quad (33)$$

where $t = [(d_{\text{free}} - 1)/2]$ and $[a]$ denote the largest integer less than or equal to a . Sets Σ_j for $(j = m, m+1, \dots, J)$ constitute a family of non-increasing sets, i.e.,

$$\Sigma_m \supseteq \Sigma_{m+1} \supseteq \dots \Sigma_J. \quad (34)$$

The smallest set Σ_J is the analogue of the classical error-correcting sphere in block codes. However, since the elements of Σ_j are infinite sequences, it is perhaps better to call Σ_j an error-correction cylinder rather than sphere. In fact, a set Σ_j , as defined in Eq. (33), is precisely the infinite intersection of what usually are called cylinder sets of form

$$C_\ell = \left\{ \underline{e}(D) \mid W_H(\underline{e}_\ell, \dots, \underline{e}_{\ell+j}) \leq t \right\}$$

for $(\ell = 0, 1, 2, \dots)$.

For a linear block code, the set of all error vectors which are correctable unambiguously by minimum distance decoding is equal to the error-correction sphere. As a consequence, one might suspect by analogy for convolutional codes that E , the set of all error error sequences which can be corrected by error-trellis decoding in Eq. (22) or its equivalent, Viterbi decoding, would equal Σ_j . However, the methods of Viterbi or sequential decoding require the processing of the entire sequence before a final decision is made. Hence, it is suspected that there are sequences correctable by standard Viterbi decoding which are not in Σ_j , and possibly vice versa.

$$\begin{aligned} \underline{e}_j &= \underline{w}_j + \underline{r}_j \\ &= F(\underline{u}_j, S_j) + \underline{r}_j \end{aligned} \quad (28)$$

in terms of inputs $\underline{u}_j$ and $\underline{r}_j$ and internal state S_j , where initially $\underline{u}_j = 0$ for $j < 0$ and $S_0 = 0$.

The states S_j of the error trellis, defined in Eqs. (24) and (25b), for an (n, k) CC of memory m are elements of the set of m -tuples of k -vectors as follows:

$$\Omega = \left\{ (\underline{\sigma}_1, \dots, \underline{\sigma}_m) \mid \underline{\sigma}_i \in V_k(F)_1, 1 \leq i \leq m \right\}. \quad (29)$$

By Eq. (26), a possible state $S_j = (\underline{\sigma}_1, \dots, \underline{\sigma}_m)$ in Ω is equal to the m -tuple of m consecutive past frames of the input $\underline{u}_j$. That is,

$$S_j = (\underline{\sigma}_1, \dots, \underline{\sigma}_m) = (\underline{u}_{j-1}, \dots, \underline{u}_{j-m}), \quad (30)$$

where $\underline{u}_j = 0$ for $j < 0$.

In convolutional codes, the j -th column distance, denoted by d_j , is defined [6] as the minimum Hamming weight of the first $(j + 1)$ codeword frames, where the 0-th information frame is non-zero. That is,

$$d_j = \min_{\underline{x}_0 \neq 0} W_H(\underline{x}_0, \dots, \underline{u}_j), \quad (31)$$

where $W_H(\cdot)$ is the Hamming weight. The column distance, d_j , forms a non-decreasing sequence for $(j = 0, 1, \dots)$. For a basic encoder, the limit of d_j is reached in a finite number, J , and equals the free distance, d_j . Thus,

$$\begin{aligned} d_j &< d_{\text{free}} \text{ for } j < J \\ &= d_{\text{free}} \text{ for } j \geq J. \end{aligned} \quad (32)$$

Equation (23) can be expressed in standard sequential circuit form by defining the state S_j of the circuit at frame j to be

$$S_j = (\underline{u}_{j-1}, \dots, \underline{u}_{j-m}), \quad (24)$$

the m -tuple of the previous m values of the k -vector $\underline{u}_j$ in sequence. Expressed as a sequential circuit, where initially $\underline{u}_j = 0$ for $j < 0$ and $S_0 = 0$, the coding trellis in Eq. (23) is given by

$$\underline{w}_j = F(\underline{u}_j, S_j) \quad (25a)$$

$$S_{j+1} = P(\underline{u}_j, S_j), \quad (25b)$$

where "P" denotes the projection along the first m components of $(\underline{u}_j, S_j) = (\underline{u}_j, \underline{u}_{j-1}, \dots, \underline{u}_{j-m})$ as follows:

$$P(\underline{u}_j, S_j) = P(\underline{u}_j, \underline{u}_{j-1}, \dots, \underline{u}_{j-m}) = (\underline{u}_j, \dots, \underline{u}_{j-m+1}). \quad (26)$$

The next state equation of the encoder sequential circuit, associated with the error trellis, is Eq. (25b) and its output equation is Eq. (25a).

Let the formal power series for $\underline{z}(D)$, $R(D)$ and $\underline{e}(D)$ be

$$\underline{z}(D) R(D) = \sum_{j=0}^{\infty} \underline{r}_j D^j \quad (27a)$$

and

$$\underline{e}(D) = \sum_{j=0}^{\infty} \underline{e}_j D^j. \quad (27b)$$

Then, by equating coefficients of D^j and using Eqs. (25) and (27), the output label at frame j of the error trellis is given by the n -vector

Table 1
Representations of $GF(2^3)$

α^k	$a_0 + a_1 \alpha + a_2 \alpha^2$
0	0
α	α
α^2	α^2
α^3	$1 + \alpha^2$
α^4	$1 + \alpha + \alpha^2$
α^5	$1 + \alpha$
α^6	$\alpha + \alpha^2$
α^7	1

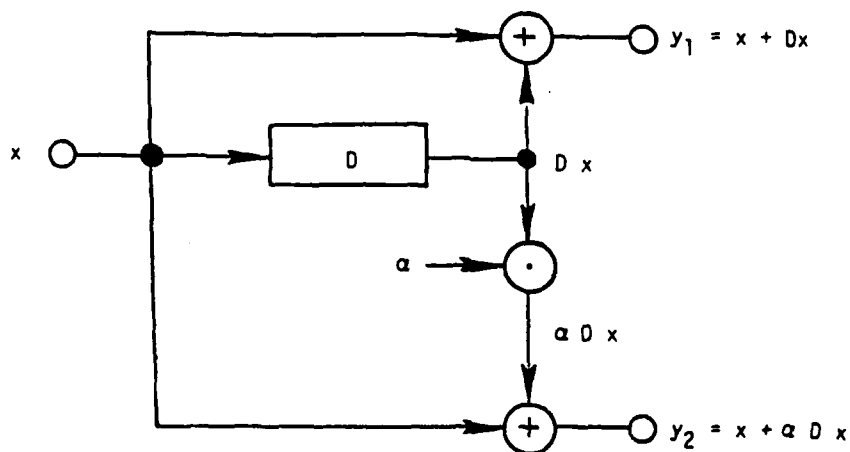


Fig. 1 - Rate 1/2, dual-3, convolutional encoder.

By Fig. 1, the output of the encoder, in terms of input, is

$$\underline{y}(D) = [y_1^{(D)}, y_2^{(D)}] = \underline{x}(D) [1 + D, 1 + \alpha D]$$

so that the generating matrix for the dual-3 CC is

$$G(D) = [1 + D, 1 + \alpha D].$$

If one applied elementary column operations to G , it is not difficult to show that

$$G(D) = [1, 0] \begin{bmatrix} 1 + D & 1 + \alpha D \\ 1 & \alpha \end{bmatrix}$$

is the Smith normal form, Eq. (5). Hence,

$$B(D) = \begin{bmatrix} 1 + D & 1 + \alpha D \\ 1 & \alpha \end{bmatrix} \text{ and}$$

$$B(D)^{-1} = \begin{bmatrix} \alpha^3 & \alpha^2 + \alpha^3 D \\ \alpha^2 & \alpha^2 + \alpha^2 D \end{bmatrix}$$

are the matrices needed in Eq. (6). By Eq. (17b),

$$\begin{aligned} R(D) &= \bar{B}_2(D) B_2(D) = \begin{bmatrix} \alpha^2 + \alpha^3 D \\ \alpha^2 + \alpha^3 D \end{bmatrix} [1, \alpha] \\ &= \begin{bmatrix} \alpha^2 + \alpha^3 D & \alpha^3 + \alpha^4 D \\ \alpha^2 + \alpha^2 D & \alpha^3 + \alpha^3 D \end{bmatrix} \end{aligned}$$

is the matrix $R(D)$ needed in the error trellis solution, Eq. (17a), of the syndrome equation.

It is easy to show from the above generating matrix $G(D)$ that $d_{\text{free}} = 4$ and $t = 1$. Thus, by Theorem 2, $S(q, k, m, t) = 8$ and $T(q, k, m, t_1) = 15$. Thus, the error trellis can be pruned, as shown in Fig. 2. In one constraint

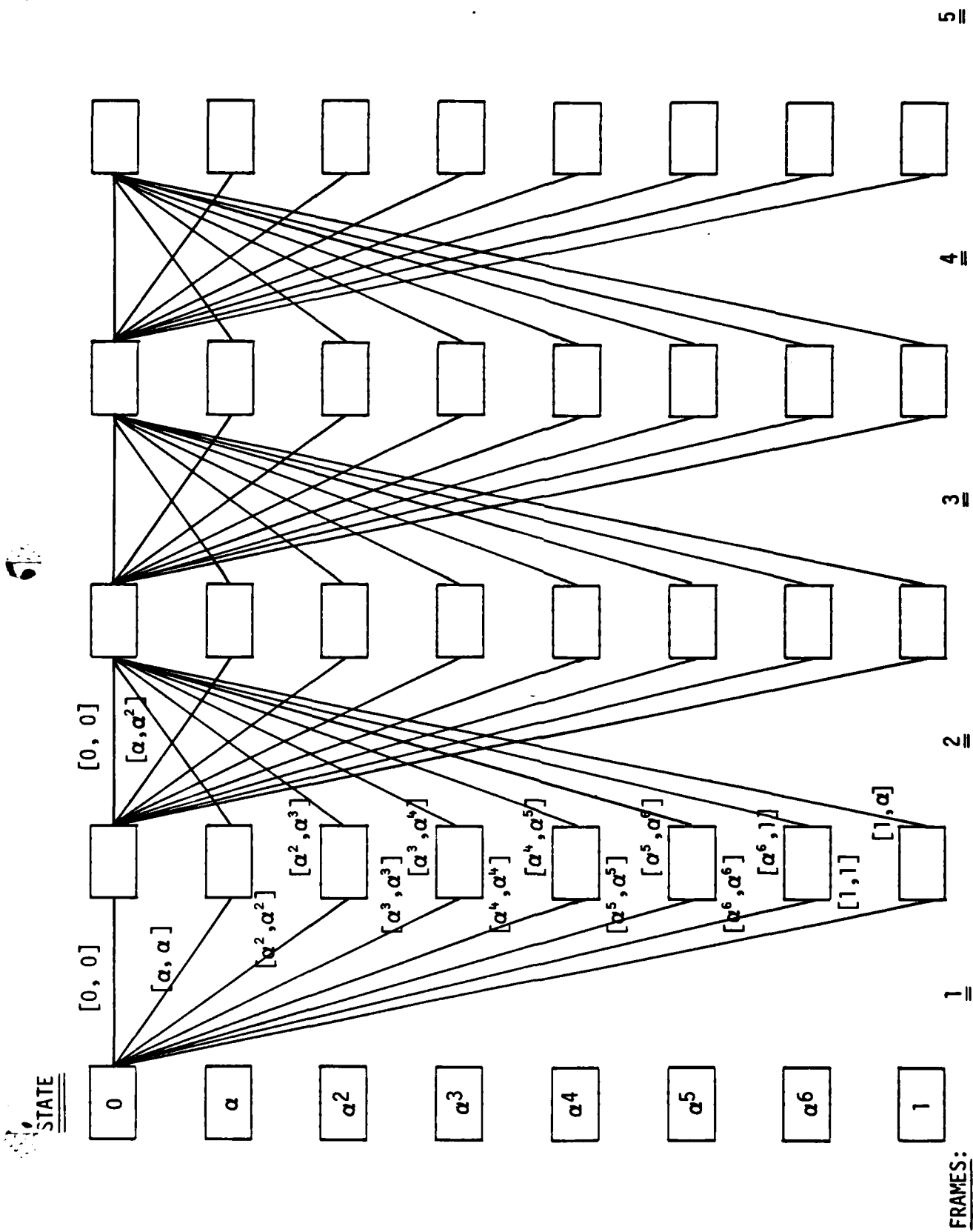


Fig. 2 — Pruned error trellis with no error outputs, $u(D) G(D)$.

length, only 15 transitions are needed in the error trellis, whereas, for the standard decoding trellis, 64 transitions are required. This yields a reduction of $15/84 \approx 1/4$ in the number of transitions needed for error-trellis decoding below the number required for standard Viterbi hard decoding.

The labels on the pruned error trellis shown in Fig. 2 correspond to the solution, Eq. (17a), of the syndrome equation for the actual error sequence equal to the all-zero sequence. That is,

$$\begin{aligned}\underline{e}(D) &= [e_1, e_2] = \underline{u}(D) G = \underline{u}(D) [1 + D, 1 + \alpha D] \\ &= [t + Dt, t + \alpha Dt]\end{aligned}$$

are the output of the trellis. For example, at frame time j and state 0, if $\underline{u}(D) = \alpha^4$, then

$$\underline{e}(D) = [\alpha^4 + 0, \alpha^4 + \alpha \cdot 0] = [\alpha^4, \alpha^4]$$

is the label on transition from state 0 to state 1. Such a transition represents an attempt to "cancel" a single error in the error-trellis equation, Eq. (17a). If such an error does, in fact, occur at frame j , then no further errors are allowed to occur at frame $j + 1$. Thus, a transition to other than state 0 must be followed by a transition back to state 0 in the next frame, as shown in Fig. 2.

Next, suppose a transition to state α^4 occurs, i.e., $Dt = \alpha^4$. Then, since $u(D) = 0$, the transition from state α^4 back to 0 is given by

$$\underline{e}(D) = [0 + \alpha^4, 0 + \alpha \cdot \alpha^4] = [\alpha^4, \alpha^5].$$

The remaining labels to the "pruned" error trellis are obtained in a similar manner.

To illustrate pruned error-trellis decoding of the dual-3 CC, let the generating function of message of information sequence be

$$\underline{x}(D) = 1 + \alpha D .$$

Then the codeword sequence is, by Eq. (3),

$$\underline{y}(D) = \underline{x}(D) G(D) = \left[1 + \alpha^5 D + \alpha D^2, 1 + \alpha^2 D^2 \right] .$$

Next, let the actual error sequence be

$$\underline{e}_a(D) = \left[D^2, \alpha \right]$$

so that

$$\underline{z}(D) = \underline{y}(D) + \underline{e}_a(D) = \left[1 + \alpha^5 D + \alpha^5 D^2, \alpha^5 + \alpha^2 D^2 \right] .$$

Hence, by Eq. (17b),

$$\begin{aligned} \underline{z}(D) R(D) &= \left[1 + \alpha^5 D + \alpha^5 D^2, \alpha^5 + \alpha^2 D^2 \right] \begin{bmatrix} \alpha^2 + \alpha^3 D, \alpha^3 + \alpha^4 D \\ \alpha^2 + \alpha^2 D, \alpha^3 + \alpha^3 D \end{bmatrix} \\ &= \left[\alpha^3 + \alpha^3 D + \alpha^2 D^2 + \alpha^3 D^3, \alpha^4 + \alpha^4 D + \alpha^3 D^2 + \alpha^4 D^3 \right] \\ &= \left[\alpha^3, \alpha^4 \right] + \left[\alpha^3, \alpha^4 \right] D + \left[\alpha^2, \alpha^3 \right] D^2 + \left[\alpha^3, \alpha^4 \right] D^3 . \end{aligned}$$

The finding of the minimum-weight error path $\hat{e}(D)$ in terms of $\hat{u}(D)$ is equivalent, by Eq. (17a), to finding the codeword $\underline{u}(D) G(D)$ which is closest to $\underline{z}(D) R(D)$, as given above. Hence, the minimum-weight error path can be found by applying the Viterbi decoding algorithm to the pruned error trellis in Fig. 2. To accomplish this, the frames of $\underline{z}(D) R(D)$ are added to the outputs $\underline{u}(D) G(D)$ in the pruned error trellis in Fig. 2, as shown in Fig. 3.

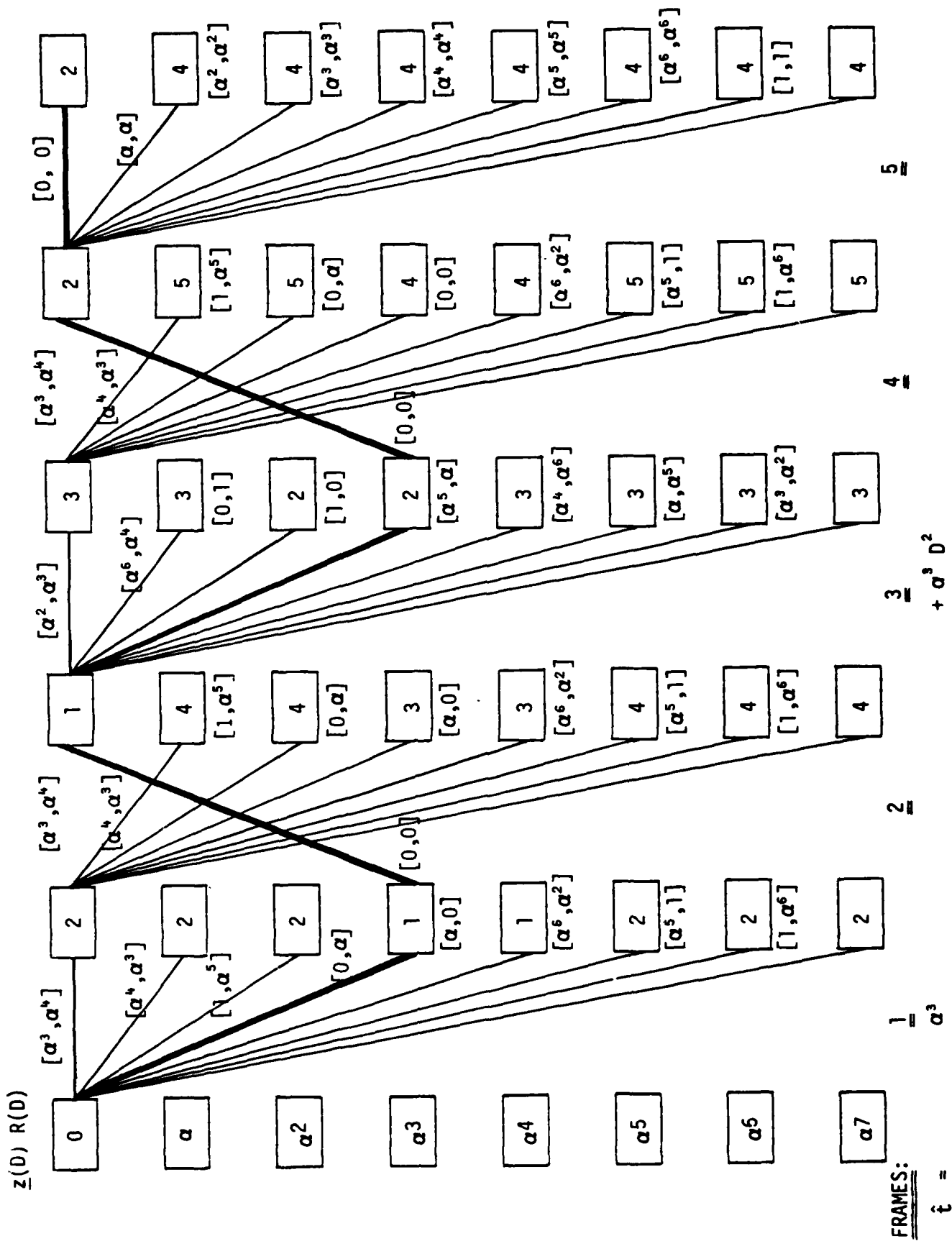


Fig. 3 - Minimum-error path, $\hat{u}(D)$, in pruned error trellis.

In order to illustrate the Viterbi algorithm as applied to the pruned error trellis, suppose the decoder has reached frame 4. The output of the transition from state α^3 to state 0 is

$$\begin{aligned} & \text{coef}_{D^3} \left[\underline{u}(D) G(D) + \underline{z}(D) R(D) \right] \\ &= \left[\alpha^3, \alpha^4 \right] + \left[\alpha^3, \alpha^4 \right] = \left[0, 0 \right] \end{aligned}$$

with Hamming weight 0. A similar calculation for the other seven possible transitions shows that the transition from α^3 to 0 is the only one with Hamming weight 0. The path segment from α^3 to 0 is chosen since it has minimum weight.

At frame 5 in Fig. 3, the minimum weight estimate of the D-transform of the error sequence is

$$\hat{\underline{e}}(D) = \left[0, \alpha \right] + \left[1, 0 \right] D^2 .$$

Hence, the estimate $\hat{\underline{u}}(D)$ of the message correction factor which achieves $\hat{\underline{e}}(D)$ is

$$\underline{u}(D) = \alpha^3 + \alpha^3 D^2 .$$

Finally, using the above results in Eq. (19) yields, by Table 1,

$$\begin{aligned} \hat{x} &= \underline{z}(D) G^{-1} - \hat{\underline{u}}(D) = \underline{z}(D) B^{-1} \begin{bmatrix} 1 \\ 0 \end{bmatrix} - \hat{\underline{u}}(D) \\ &= \underline{z}(D) \begin{bmatrix} \alpha^3, \alpha^2 + \alpha^3 D \\ \alpha^2, \alpha^2 + \alpha^2 D \end{bmatrix} \begin{bmatrix} 1 \\ 0 \end{bmatrix} - \hat{\underline{u}}(D) = \underline{z}(D) \begin{bmatrix} \alpha^3 \\ \alpha^2 \end{bmatrix} - \hat{\underline{u}}(D) \\ &= \left[1 + \alpha^5 D + D^2, \alpha^5 + \alpha^2 D^2 \right] \begin{bmatrix} \alpha^3 \\ \alpha^2 \end{bmatrix} + \left(\alpha^3 + \alpha^2 + D^2 \right) \\ &= \left(\alpha^3 + \alpha D + \alpha^3 D^2 + 1 + \alpha^4 D^2 \right) + \left(\alpha^3 + \alpha^3 D^2 \right) = 1 + \alpha D, \end{aligned}$$

the original encoded message.

V. CONCLUDING REMARKS

In this report, pruned error-trellis decoding of systematic and non-systematic convolutional codes with a delay-free inverse has been developed in detail, including quantitative formulas for the number of states and transitions which remain in the pruned error trellis. Currently, the problem of trellis pruning of other non-systematic CCs is being investigated. Finally, the reduced hardware requirements for pruned error-trellis decoding versus standard Viterbi decoding is being studied, and a preliminary architecture has already been found for the dual-K decoding algorithm developed in this report.

REFERENCES

1. Reed, I. S., *Sequential Syndrome Decoding of Convolutional Codes*, Third Quarterly Report submitted to the Naval Air Systems Command on Contract N00019-83-C-0075 by Adaptive Sensors, Inc., January 1984.
2. Reed, I. S., *Minimum-Error Trellis-Path Decoders for Convolutional Codes*, Final Report submitted to the Naval Air Systems Command on Contract N00019-83-C-0075 by Adaptive Sensors, Inc., March 1984.
3. Jensen, J. M., and I. S. Reed, "Error-Trellis Decoding of Convolutional Codes," paper submitted to *IEEE Transactions on Information Theory*.
4. Odenwalder, V. P., "Dual-K Convolutional Codes for Non-Coherently Demodulated Channels," *Proceedings of the International Telemetry Conference (ITC)*, Vol. 12, 1978, pp. 165-174.
5. Forney, G. D., "Convolutional Codes: Algebraic Structures," *IEEE Transactions on Information Theory*, IT-6, pp. 720-738.
6. Lin, S., and D. T. Costello, Jr., *Error Control Coding*, Prentice Hall, New Jersey, 1983.

END

FILMED

8-85

DTIC