

AD-A124 492

SOME TECHNIQUES IN UNIVERSAL SOURCE CODING AND DURING
FOR COMPOSITE SOURCES(U) ILLINOIS UNIV AT URBANA
COORDINATED SCIENCE LAB M S WALLACE DEC 81 R-929

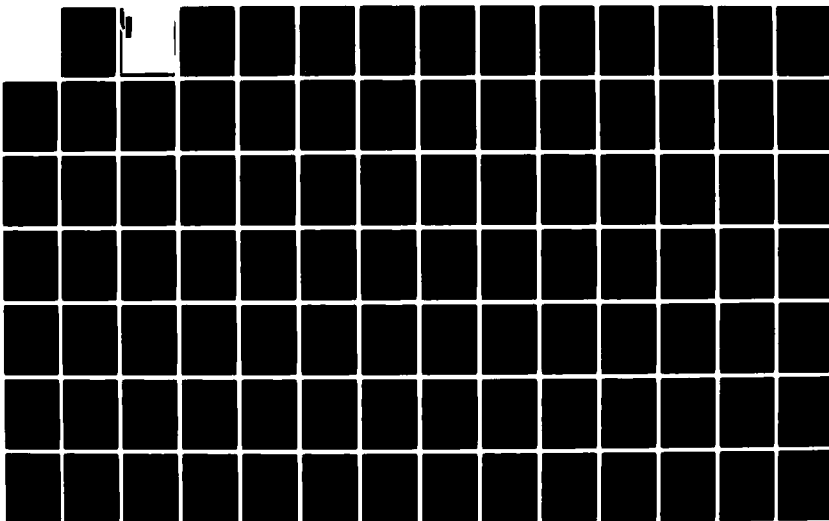
1/2

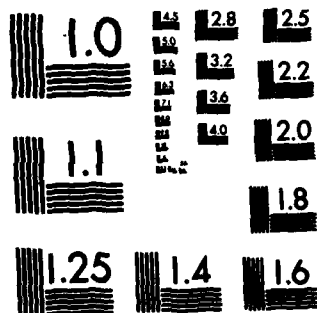
UNCLASSIFIED

N00014-79-C-0424

F/G 9/4 •

NL





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

DA 124492

Next we turn to source coding problems. The determination of the sources entropy is of interest as it provides a lower bound on the rate of any code. If the switching process is stationary then the output process is also stationary since it is a memoryless function of the switching process.

So the probability of a block $X_n = (X_1, X_2, \dots, X_n)$ of source outputs given as

UNCLASSIFIED

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

REPORT DOCUMENTATION PAGE		READ INSTRUCTIONS BEFORE COMPLETING FORM
1. REPORT NUMBER R-929	2. GOVT ACCESSION NO. AD-A124492	3. RECIPIENT'S CATALOG NUMBER
4. TITLE (and Subtitle) SOME TECHNIQUES IN UNIVERSAL SOURCE CODING AND CODING FOR COMPOSITE SOURCES		5. TYPE OF REPORT & PERIOD COVERED Technical Report
7. AUTHOR(s) Mark Stanley Wallace		6. PERFORMING ORG. REPORT NUMBER R-929; UILU-ENG 81-2260
9. PERFORMING ORGANIZATION NAME AND ADDRESS Coordinated Science Laboratory University of Illinois Urbana, IL 61801		8. CONTRACT OR GRANT NUMBER(s) N00014-79-C-0424
11. CONTROLLING OFFICE NAME AND ADDRESS Joint Services Electronics Program		10. PROGRAM ELEMENT, PROJECT, TASK AREA & WORK UNIT NUMBERS
14. MONITORING AGENCY NAME & ADDRESS (if different from Controlling Office)		12. REPORT DATE December 1981
		13. NUMBER OF PAGES 114
		15. SECURITY CLASS. (of this report) UNCLASSIFIED
		15a. DECLASSIFICATION/DOWNGRADING SCHEDULE
16. DISTRIBUTION STATEMENT (of this Report) Approved for public release; distribution unlimited		
17. DISTRIBUTION STATEMENT (of the abstract entered in Block 20, if different from Report)		
18. SUPPLEMENTARY NOTES		
19. KEY WORDS (Continue on reverse side if necessary and identify by block number) Universal Source Coding Composite Source Vector Quantization		
20. ABSTRACT (Continue on reverse side if necessary and identify by block number) We consider three problems in source coding. First, we consider the composite source model. A composite source has a switch driven by a random process which selects one of a possible set of subsources. We derive some convergence results for estimation of the switching process, and use these to prove that the entropy of some composite sources may be computed. Some coding techniques for composite sources are also presented and their performance is bounded.		

(over)

DD FORM 1 JAN 73 1473

UNCLASSIFIED

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

UNCLASSIFIED

SECURITY CLASSIFICATION OF THIS PAGE(When Data Entered)

Next, we construct a variable-length-to-fixed-length (VL-FL) universal code for a class of unifilar Markov sources. A VL-FL code maps strings of source outputs into fixed-length codewords. We show that the redundancy of the code converges to zero uniformly over the class of sources as the blocklength increases. The code is also universal with respect to the initial state of the source. We compare the performance of this code to FL-VL universal codes.

We then consider universal coding for real-valued sources. We show that given some coding technique for a known source, we may construct a code for a class of sources. We show that this technique works for some classes of memoryless sources, and also for a compact subset of the class of k-th order Gaussian autoregressive sources.



SECURITY CLASSIFICATION OF THIS PAGE(When Data Entered)

UILLU-ENG 81-2260

SOME TECHNIQUES IN UNIVERSAL SOURCE CODING
AND CODING FOR COMPOSITE SOURCES

by

Mark Stanley Wallace

This work was supported in part by the Joint Services Electronics
Program (U.S. Army, U.S. Navy and U.S. Air Force) under Contract
N00014-79-C-0424.

Reproduction in whole or in part is permitted for any purpose of
the United States Government.

Approved for public release. Distribution unlimited.



Accession For	
NTIS GRA&I	
DTIC TAB	
Unannounced	
Justification	
By	
Distribution/	
Availability Codes	
Dist	Avail and/or Special
A	

**SOME TECHNIQUES IN UNIVERSAL SOURCE CODING
AND CODING FOR COMPOSITE SOURCES**

BY

MARK STANLEY WALLACE

**B.Eng., McGill University, 1978
M.S., University of Illinois, 1980**

THESIS

**Submitted in partial fulfillment of the requirements
for the degree of Doctor of Philosophy in Electrical Engineering
in the Graduate College of the
University of Illinois at Urbana-Champaign, 1982**

Thesis Adviser: Professor Michael B. Pursley

Urbana, Illinois

ACKNOWLEDGEMENTS

I would like to thank my thesis adviser, Professor M. B. Pursley, for his advice and assistance throughout the research for this thesis. I am also grateful to Professor B. Hajek for many helpful suggestions and comments. I wish to thank Professors R. J. McEliece and R. B. Ash for serving on the doctoral examination committee.

Thanks are also due to Phyllis Young and Cathy Cassells for their expert typing.

TABLE OF CONTENTS

	Page
INTRODUCTION	1
CHAPTER 1 - STATE ESTIMATION AND CODING FOR COMPOSITE SOURCES ...	4
1.1 Introduction	4
1.2 Convergence of State Estimates for Two-State Composite Sources	8
1.3 Convergence of State Estimates for S-State Composite Sources	23
1.4 Generalization to Arbitrary Subsources	25
1.5 A Coding Technique for Composite Sources	27
1.6 Stability of Coding and Universal Coding for Composite Sources	33
1.7 Coding for an Infinite-State Composite Source	37
CHAPTER 2 - UNIVERSAL VL-FL CODING FOR MARKOV SOURCES	41
2.1 Introduction and Review of Previous Results	41
2.2 Universal VL-FL Code Construction	50
2.3 Performance Evaluation for Binary Memoryless Sources ..	55
CHAPTER 3 - UNIVERSAL CODING FOR REAL-VALUED SOURCES	63
3.1 Introduction	63
3.2 Code Construction	66
3.3 Generalization to Unbounded Distortion Measures	75
APPENDIX A - PROOFS OF THEOREMS FROM CHAPTER 1	78
APPENDIX B - PROOF OF REDUNDANCY BOUND FOR THE VL-FL CODE OF CHAPTER 2	96
APPENDIX C - PROOFS AND DERIVATIONS FOR CHAPTER 3	103
REFERENCES	112
VITA	114

SOME TECHNIQUES IN UNIVERSAL CODING AND

CODING FOR COMPOSITE SOURCES

Mark Stanley Wallace, Ph.D.
Coordinated Science Laboratory and
Department of Electrical Engineering
University of Illinois at Urbana-Champaign, 1982

ABSTRACT

We consider three problems in source coding. First, we consider the composite source model. A composite source has a switch driven by a random process which selects one of a possible set of subsources. We derive some convergence results for estimation of the switching process, and use these to prove that the entropy of some composite sources may be computed. Some coding techniques for composite sources are also presented and their performance is bounded.

Next, we construct a variable-length-to-fixed-length (VL-FL) universal code for a class of unifilar Markov sources. A VL-FL code maps strings of source outputs into fixed-length codewords. We show that the redundancy of the code converges to zero uniformly over the class of sources as the blocklength increases. The code is also universal with respect to the initial state of the source. We compare the performance of this code to FL-VL universal codes.

We then consider universal coding for real-valued sources. We show that given some coding technique for a known source, we may construct a code for any class of sources. We show that this technique works for some classes of memoryless sources, and also for a compact subset of the class of k -th order Gaussian autoregressive sources.

INTRODUCTION

The general problem in source coding is that of data compression. The data which is produced by some information source must be stored or transmitted. Since there is a cost associated with storage and transmission, it is of interest to encode the data into as small a number of bits as possible in order to minimize this cost. If the encoded data is to retain all of the original information then the problem is one in noiseless source coding. If there is some allowable distortion then the problem is one in rate-distortion theory or source coding with a fidelity criterion.

In these problems an information source is modeled as a discrete-time random process. The source output at each time i is a random variable X_i . The distribution of this random variable (which may depend on previous source outputs) determines the probability of a given source output. If the source outputs $(\dots, X_i, X_{i+1}, \dots)$ form a stationary random process, then the source is said to be stationary.

A code is defined as a function which maps blocks of source outputs into binary strings which are called codewords. The rate of a code is the expected number of bits which are used to encode a source output. If a source is stationary, then its entropy is defined. The entropy is a lower bound on the rate of any noiseless code, and noiseless codes exist with rates which are arbitrarily close to the entropy. The difference between the rate and the entropy is called the redundancy.

If the statistics of a source (i.e., the distribution of the source outputs) are known then a noiseless code for the source may be derived using Huffman's algorithm [1]. This algorithm gives fixed-length-to-variable-length (FL-VL) codes, a FL-VL code being one which maps a fixed

number of source outputs into a variable-length binary codeword. The redundancy of a blocklength n Huffman code is at most n^{-1} , so a Huffman code may be derived with rate as close to the entropy of a source as desired. A variable-length-to-fixed-length (VL-FL) algorithm (Tunstall's algorithm) is also known for a given source, and if the blocklength n is defined as the length of the codewords, then the redundancy of these codes also decreases as n^{-1} .

In practice the statistics of a source are seldom known exactly so these encoding algorithms do not apply. Universal source coding considers this problem. In universal source coding the source statistics are assumed to lie in some class. The goal is to design a code which performs well (i.e., one which has a small redundancy) for all of the sources in the class. A sequence of codes of increasing blocklength is called universal if the redundancy approaches zero as the blocklength increases for any source in the class.

There are a number of coding techniques which yield universal FL-VL codes for various classes of sources. Much less is known about universal VL-FL codes. In Chapter 2 a universal VL-FL coding technique for Markov sources is derived, and its redundancy is bounded.

A further generalization to the source model is to allow the source statistics to vary with time. So rather than having a source with fixed, but unknown statistics, a random process determines the statistics of the source. This random process, called the switching process, together with the set of possible source statistics is known as a composite source [24]. Composite sources of various types are considered in a number of papers,

e.g., [2], [3], and [8]. In Chapter 1 we consider composite sources in which the switching process is a Markov chain, and the possible sources are memoryless. (The outputs of a memoryless source at two different times are independent.) The state of the Markov chain determines the probabilities of the various source outputs, but the state cannot (in general) be determined by observing the source outputs. Some convergence properties for the estimate of the source statistics given the outputs are derived, and these are used to bound the accuracy of an algorithm to compute the entropy of some composite sources. Some coding techniques for composite sources are also presented.

In source coding with a fidelity criterion the rate of a code is to be minimized without exceeding some level of distortion. The fidelity criterion tells us the distortion incurred when one source output is reproduced as another output. There are a few possible approaches to coding in this case. The outputs may be quantized individually into some finite set of values and then encoded using a source model such as those used in Chapters 1 and 2. Another way is to design a code which maps blocks of source outputs directly to codewords. This is known as vector quantization. There are a number of techniques known for vector quantization under various constraints. In Chapter 3 we show how a technique of vector quantization for a known source may be used to generate a code for an entire class of sources.

CHAPTER 1

STATE ESTIMATION AND CODING FOR COMPOSITE SOURCES

1.1 Introduction

A composite source [24] consists of a set of subsources and a switching process which selects one of the subsources (see Fig. 1). We consider discrete-time composite sources with memoryless subsources and a switching process which is a Markov chain with state space $\mathcal{S} = \{1, 2, \dots, S\}$, $S < \infty$.

Define a state vector $\underline{Z}(i) = (Z_1(i), \dots, Z_S(i))$ by $Z_s(i) = 1$ if the switching process is in state s at time i and $Z_s(i) = 0$ otherwise. The i -th source output is a random variable X_i which takes on values in an alphabet A according to the distribution $\gamma_{\underline{Z}(i)}(\cdot)$. So the probability of a given source output is determined by the state of the switching process, and

$$P\{X_i = x | \underline{Z}(i), (X_{i-1}, \underline{Z}(i-1)), (X_{i-2}, \underline{Z}(i-2)), \dots\} = P\{X_i = x | \underline{Z}(i)\} = \gamma_{\underline{Z}(i)}(x). \quad (1)$$

We refer to $\underline{Z}(i)$ as the state of the source. The switching process is specified by an $S \times S$ matrix Q with elements

$$q(s' | s) = P\{Z_{s'}(i+1) = 1 | Z_s(i) = 1\}.$$

Note that the sequence of states $(\underline{Z}(0), \underline{Z}(1), \dots)$ is not determined by the outputs $(X_0, X_1, \dots)$ even if the state $\underline{Z}(0)$ is known. These sources are not unifilar Markov sources [1], pp. 187.

The composite source has been considered as a model for time-varying sources [2], [3], and for this application it is generally assumed that the switching process is slow. We do not assume this, in fact, all of our results are valid even if the source changes state with high probability after each source output.

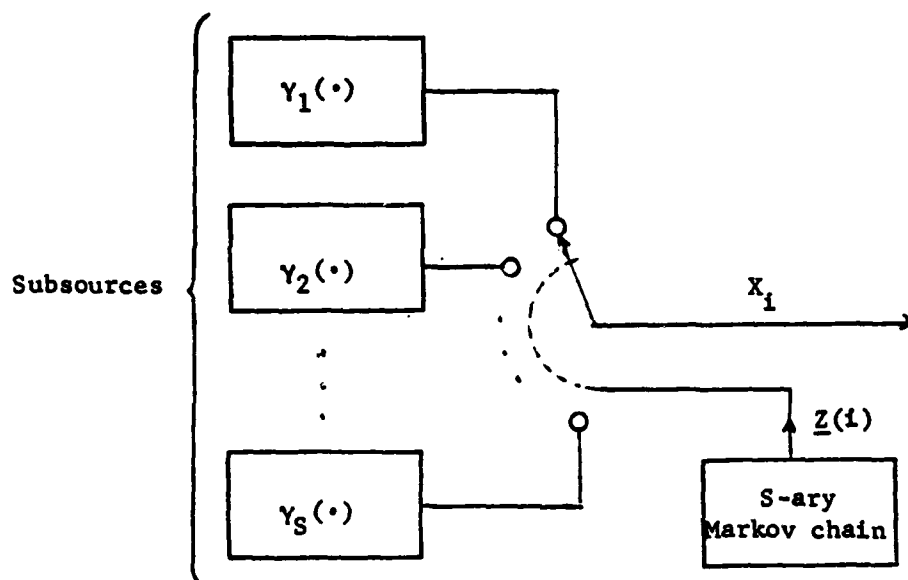


Figure 1. Diagram of composite source.

Since the state of such a composite source cannot in general be determined from the outputs, it is of interest to estimate it. Let

$$\hat{\underline{z}}(i) = E[\underline{z}(i) | X_1, X_{1-1}, \dots] \quad (2)$$

be the conditional mean estimate of the state given the past outputs. Since $\hat{\underline{z}}(i+1)$ is a sufficient statistic for X_{i+1} , $\hat{\underline{z}}(i+1)$ may be generated from X_{i+1} and $\hat{\underline{z}}(i)$ using Bayes rule [4]; however, some initial estimate is required.

The first part of the chapter is concerned with the properties of the estimation process $\hat{\underline{z}}(i)$. Although the method for generating the estimates recursively is well known, very little is known about the convergence properties of such processes. In Section 1.3 we consider the situation where no initial estimate is available, and prove that the estimates derived from any two initial estimates will converge. For composite source with only two states we show that the recursive computation of the estimates is stable. That is, small errors which are introduced in any actual implementation of the estimation procedure do not propagate. This result is not easily extended to include composite sources with a larger state space. The mean-square error of the estimate, or more generally the expected value of any function of $\hat{\underline{z}}(i)$ and $\underline{z}(i)$, is determined by the stationary distribution of the estimation process. However, in general this distribution is not known to be unique. We show that the estimation process has a unique stationary distribution, and give an algorithm which may be used to compute this distribution to any desired accuracy.

Next we turn to source coding problems. The determination of the sources entropy is of interest as it provides a lower bound on the rate of any code. If the switching process is stationary then the output process is also stationary since it is a memoryless function of the switching process. So the probability of a block $\underline{X} = (X_1, X_2, \dots, X_n)$ of source outputs given no previous outputs may be determined using the stationary distribution of the switching process as the initial estimate $\hat{Z}(0)$. Since the source is stationary we know that its entropy is

$$\lim_{n \rightarrow \infty} -n^{-1} \sum_{\underline{x} \in A^n} P\{\underline{X} = \underline{x}\} \log P\{\underline{X} = \underline{x}\} . \quad (3)$$

This does not imply that the estimation process has a unique stationary distribution. As previously mentioned, however, such a distribution exists if the source has two states, and in this case the entropy is

$$\int H(X_1 | \hat{Z}(0) = \underline{z}) \mu^*(d\underline{z})$$

where μ^* is the stationary distribution. For k -state composite sources, $k > 2$, we do not prove that a unique stationary distribution exists.

We construct fixed-length to variable-length (FL-VL) codes for composite sources and show that their redundancy is bounded by $n^{-1}(\lceil \log S \rceil + 1)$. (A FL-VL code maps fixed-length blocks of source outputs into variable-length codewords.) Again propagation of errors is a problem, and so it is not clear whether the technique is implementable for long blocklengths. For the two-state case we show that errors do not propagate. In addition the effect of inexact knowledge of the source parameters (i.e., switching probabilities and subsources statistics) is bounded. This result is used to construct a universal code for a class of two-state composite sources.

Finally we construct codes for a special class of composite sources with an infinite state space. This class of sources has the property that the probability of switching into any state is independent of the current state.

1.2 Convergence of State Estimates for Two-State Composite Sources

Let θ be a composite source consisting of two memoryless, finite-entropy subsources with alphabet A and a binary Markov switching process. The state at time t is Z_t , a random variable taking values in $\mathcal{S} = \{0,1\}$. The transition probability matrix $Q = \{q(z_t|z_{t-1})\}$ of the switching process $\{Z_t\}$ is specified by two values. For ease of notation let $\alpha = q(1|0)$ and $\beta = q(0|1)$, then $q(0|0) = 1 - \alpha$ and $q(1|1) = 1 - \beta$. The composite source θ is determined by α, β , and the two subsource distributions $\{\gamma_i(x); x \in A\}$ $i = 0,1$, so we write $\theta = (\alpha, \beta, \gamma_0, \gamma_1)$. Let Λ denote the class of such sources for a given alphabet A . Define the estimate

$$\hat{Z}_t = E[Z_t | X_t, X_{t-1}, \dots] \quad (4)$$

This estimate has the following property.

Lemma 1.1. Let $\underline{X} = (X_1, X_2, \dots)$. Then $\underline{X} - \hat{Z}_{-1} - Z_0$ forms a Markov chain.

Proof: If $z = E[Z_{-1} | \underline{X} = \underline{x}]$ then

$$\begin{aligned} P\{Z_0=s | \hat{Z}_{-1}=z, \underline{X}=\underline{x}\} &= \sum_{s'=0}^1 P\{Z_0=s | Z_{-1}=s', \hat{Z}_{-1}=z, \underline{X}=\underline{x}\} P\{Z_{-1}=s' | \hat{Z}_{-1}=z, \underline{X}=\underline{x}\} \\ &= \sum_{s'=0}^1 P\{Z_0=s | Z_{-1}=s'\} P\{Z_{-1}=s' | \underline{X}=\underline{x}\} \quad (5) \\ &= P\{Z_0=s | Z_{-1}=0\}(1-z) + P\{Z_0=s | Z_{-1}=1\}z \\ &= P\{Z_0=s | \hat{Z}_{-1}=z\} \end{aligned}$$

where (5) follows since $\hat{Z}_{-1}$ is a function of $\underline{X}$ and since the transition probabilities of the switching process do not depend on the outputs.

Given $\hat{Z}_t = z$ and $X_{t+1} = x$, then Lemma 1.1 implies that $\hat{Z}_{t+1}$ is given by Bayes rule [4], so $\hat{Z}_t$ is the conditional mean estimate of Z_t given observation of the source output up to time t .

$$\hat{Z}_{t+1} = f_x(z) \triangleq \frac{P\{X_{t+1} = x, Z_{t+1} = 1 | \hat{Z}_t = z\}}{P\{X_{t+1} = x | \hat{Z}_t = z\}} = \frac{\gamma_1(x)\eta_1(z)}{\sum_{i=0}^1 \gamma_i(x)\eta_i(z)} \quad (6)$$

where

$$\eta_i(z) \triangleq P\{Z_{t+1} = i | \hat{Z}_t = z\} = \begin{cases} \beta z + (1-\alpha)(1-z) & ; i=0 \\ (1-\beta)z + \alpha(1-z) & ; i=1 \end{cases} \quad (7)$$

If we define

$$p(x|z) = P\{X_{t+1} = x | \hat{Z}_t = z\} = \sum_{i=0}^1 \gamma_i(x)\eta_i(z) \quad (8)$$

then $p(x|z)$ is the probability that the new estimate will be $f_x(z)$ given that the old estimate was z . If μ is the distribution of $\hat{Z}_0$ then the distribution of $\hat{Z}_1$ is μT , where T is the measure transformation defined by

$$\mu T(B) = \sum_{x \in A} \int_{f_x^{-1}(B)} p(x|z) \mu(dz) \quad (9)$$

where $B \subset [0,1]$ and $f_x^{-1}(B) \triangleq \{z \in [0,1]: f_x(z) \in B\}$. The transformation T

has the following contraction property. The distance measure used here is the $\bar{\rho}$ -distance [5] which is defined by

$$\bar{\rho}(\mu, \nu) = \inf_{\pi \in P} \int |x-y| \pi(dx, dy) , \quad (10)$$

where P is the set of joint distributions with marginals μ and ν . We first prove the following theorem.

Theorem 1.1. (Contraction property of T in the $\bar{\rho}$ -metric)

If μ and ν are two distributions of the state estimate for a two-state composite source with memoryless subsources and if T is the transformation (9) for this source then

$$\bar{\rho}(\mu T, \nu T) \leq |\lambda| \bar{\rho}(\mu, \nu) \quad (11)$$

where $\lambda = 1 - \alpha - \beta$ and $\alpha = q(1|0)$ and $\beta = q(0|1)$ are transition probabilities for the switching process.

Proof: See Appendix A.

The following corollary is an immediate consequence of Theorem 1.1.

Corollary 1.1.

$$\bar{\rho}(\mu_1, \nu_1) \leq |\lambda|^1 \bar{\rho}(\mu_0, \nu_0) \quad (12)$$

where $\mu_1 = \mu_0 T^1$ and $\nu_1 = \nu_0 T^1$. We now show that a unique stationary distribution exists if $|\lambda| < 1$.

Theorem 1.2. The state estimate $\hat{Z}_t$ has a unique stationary distribution μ^* if $|\lambda| < 1$.

Proof: Since the space of possible distributions is compact in the $\bar{\rho}$ -metric we know that a subsequential limit exists. For any two distributions μ_0, ν_0

$$\bar{\rho}(\mu_0, \nu_0) \leq 1 \quad (13)$$

so

$$\bar{\rho}(\mu_i, \nu_i) \leq |\lambda|^i. \quad (14)$$

Let $\nu_0 = \mu_j$. Then (14) implies

$$\bar{\rho}(\mu_i, \mu_{i+j}) \leq |\lambda|^i \quad (15)$$

for any j . If there exist two subsequential limits μ' and μ'' with $\mu_{k_i} \rightarrow \mu'$ and $\mu_{j_i} \rightarrow \mu''$ for subsequences k_i and j_i then for arbitrary i

$$\bar{\rho}(\mu_{k_i}, \mu_{j_i}) \leq |\lambda|^{\min(j_i, k_i)}. \quad (16)$$

It follows that $\bar{\rho}(\mu', \mu'') = 0$, and thus $\{\mu_i\}$ has a unique limit. Since

$$\mu' = \lim_{i \rightarrow \infty} \mu_0 T^i = \lim_{i \rightarrow \infty} \mu_0 T^{i+1} = \mu' T \quad (17)$$

the limit is stationary. If the alphabet A is finite then we may compute this stationary distribution to any desired degree of

accuracy as follows. Let μ^* denote the stationary distribution. From a distribution $\tilde{\mu}_{j-1}$ concentrated on the set $\Omega \triangleq \{\frac{i-k}{n}; i = 1, \dots, n\}$ we generate a distribution $\hat{\mu}_j$ on $\Omega' \triangleq \{f_x(\frac{i-k}{n}); x \in A; i = 1, \dots, n\}$ using the recursive equation (9). Then a distribution $\tilde{\mu}_j$ concentrated on Ω is generated using

$$\tilde{\mu}_j(\{\frac{i-k}{n}\}) = \hat{\mu}_j((\frac{i-1}{n}, \frac{i}{n})) \quad ; \quad i = 1, \dots, n. \quad (18)$$

(We use $\{x\}$ to denote the set containing the point x .) This algorithm is clearly implementable since only a bounded number of points is considered,

and if we define

$$e_j \triangleq \bar{\rho}(\bar{\mu}_j, \mu^*) \quad (19)$$

then

$$e^* \triangleq \lim_{j \rightarrow \infty} e_j \leq \hat{e} \triangleq [2n[1 - |\lambda|]]^{-1} \quad (20)$$

and

$$e_j \leq |\lambda|^j + \hat{e}. \quad (21)$$

So we may compute $\bar{\mu}_j$ such that $\bar{\rho}(\bar{\mu}_j, \mu^*) \leq \epsilon$ for any $\epsilon > 0$ by choice of n and j sufficiently large. Equations (20) and (21) follow since (19) implies

$$\bar{\rho}(\bar{\mu}_j, \hat{\mu}_j) \leq (2n)^{-1} \quad (22)$$

and from Theorem 1.1

$$\bar{\rho}(\hat{\mu}_j, \mu^*) \leq |\lambda| \bar{\rho}(\bar{\mu}_{j-1}, \mu^*)$$

so

$$e_j \leq |\lambda| e_{j-1} + (2n)^{-1}. \quad (23)$$

In the limit as j goes to infinity (23) becomes

$$e^* \leq |\lambda| e^* + (2n)^{-1} \quad (24)$$

which gives (20) and subtracting $\hat{e}$ from both sides of (23) we have

$$e_j - \hat{e} \leq |\lambda| (e_{j-1} - \hat{e}) \quad (25)$$

which gives (21).

The number of computations required increases linearly with both n , the size of the vector which approximates the joint distribution, and j , the number of iterations. The storage required increases linearly with n . If we fix j according to the limiting error $\hat{\epsilon}$ by

$$|\lambda|^j = \hat{\epsilon} \quad (26)$$

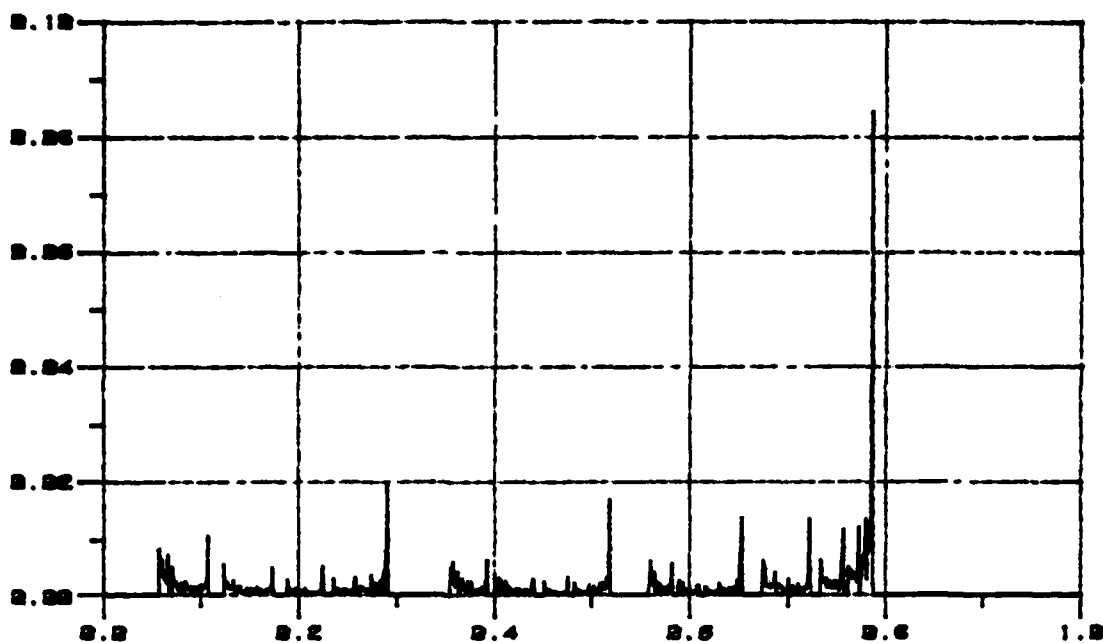
then j is of order $\log n$. So the number of computations required to derive a distribution within n^{-1} of the true stationary distribution increases as $n \log n$, and the storage required increases as n . This algorithm was implemented for $A = \{0,1\}$, i.e. binary memoryless subsources. Two computed distributions and their associated cumulative distribution functions are illustrated in Fig. 2. The distributions are concentrated on 1000 points, and the $\bar{\rho}$ -distance between these distributions and the stationary distributions is at most .006. The distributions are not smooth, and it does not appear likely that a closed form analytical description exists.

The computed distribution may be used to bound the performance of the estimator as follows. The mean-square estimation error is

$$\begin{aligned} E[(Z_t - \hat{Z}_t)^2] &= E[Z_t^2] - E[\hat{Z}_t^2] \\ &= E[Z_t] - E[\hat{Z}_t^2] \end{aligned} \quad (27)$$

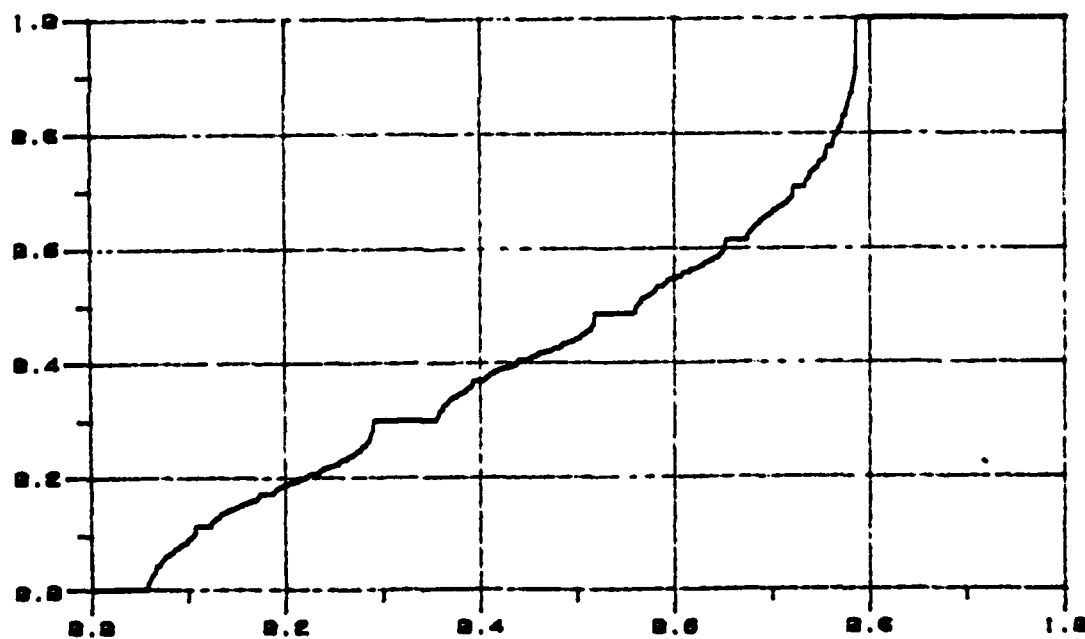
where (27) follows because $Z_t^2 = Z_t$. If the switching process is stationary and ergodic then

$$E[Z_t] = \frac{\alpha}{\alpha + \beta} \quad (28)$$



i) density (concentrated on 1000 points)

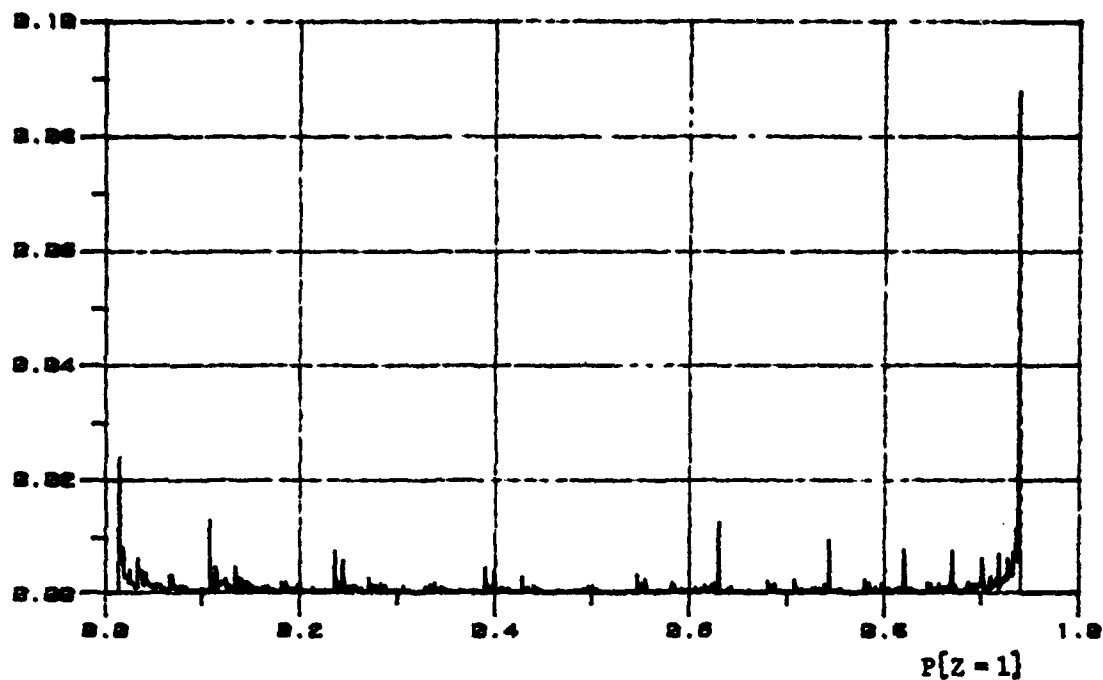
$P\{Z=1\}$



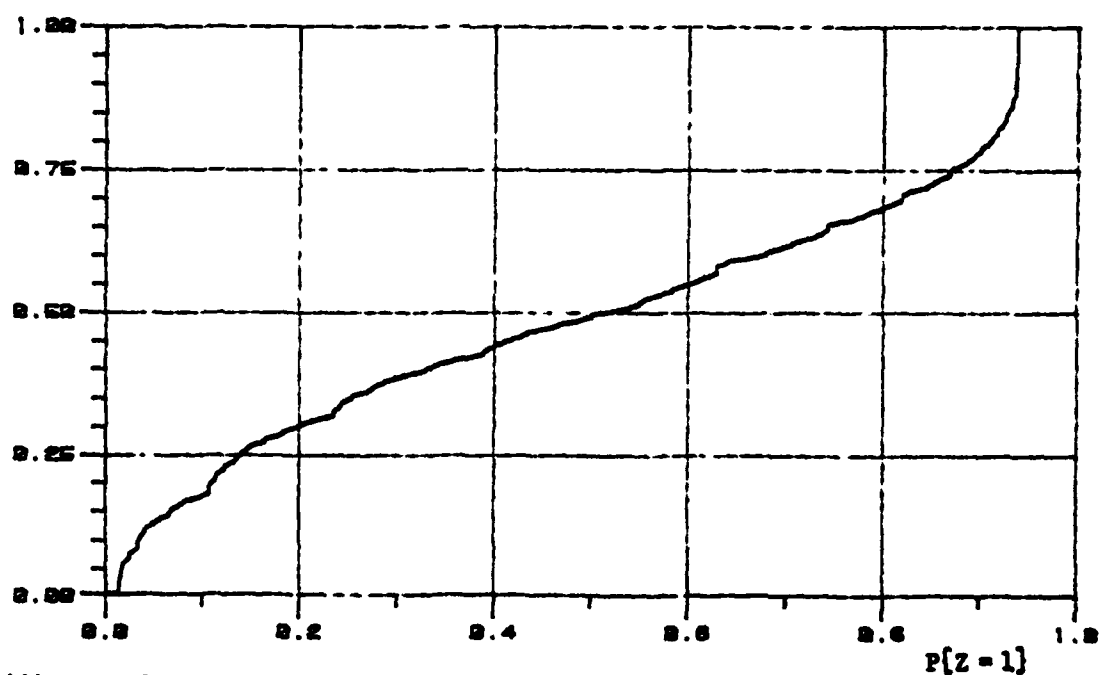
ii) cumulative distribution function

$P\{Z=1\}$

Figure 2a. Approximate stationary distribution of the state estimate for a composite source with $\alpha = \beta = 2$, $\gamma_0(0) = .5$, and $\gamma_1(1) = .9$.



i) density (concentrated on 1000 points)



ii) cumulative distribution function

Figure 2b. Approximate stationary distribution of the state estimate for a composite source with $\alpha = \beta = .05$, $\gamma_0(0) = .9$, and $\gamma_1(1) = .5$.

If μ^* is the stationary distribution of the estimate $\hat{Z}_t$ we have

$$E[\hat{Z}_t^2] = \int_0^1 z^2 \mu^*(dz) \quad (29)$$

Given a distribution $\tilde{\mu}$ such that $\bar{\rho}(\mu^*, \tilde{\mu})$ is small, we use the following theorem to bound

$$\Delta \triangleq \left| \int_0^1 z^2 d\tilde{\mu} - \int_0^1 z^2 d\mu^* \right| \quad (30)$$

in terms of $\bar{\rho}(\mu^*, \tilde{\mu})$.

Theorem 1.3. If μ and ν are probability measures on $[0,1]$ then

$$\left| \int_0^1 f d\mu - \int_0^1 f d\nu \right| \leq \sup_{x \in [0,1]} |f'(x)| \cdot \bar{\rho}(\mu, \nu) \quad (31)$$

Proof. The theorem follows directly from integration by parts. That is

$$\int_0^1 f d\mu = f(x) \Big|_0^1 - \int_0^1 f'(x) \mu[0, x] dx$$

so

$$\begin{aligned} \left| \int_0^1 f d\mu - \int_0^1 f d\nu \right| &= \left| \int_0^1 f'(x) (\nu[0, x] - \mu[0, x]) dx \right| \\ &\leq \int_0^1 |f'(x)| |\nu[0, x] - \mu[0, x]| dx \end{aligned} \quad (32)$$

$$\leq \sup_{x_0 \in [0,1]} |f'(x_0)| \bar{\rho}(\mu, \nu) \quad (33)$$

Equation (33) follows because for one-dimensional distributions [5]

$$\bar{\rho}(\mu, \nu) = \int |\mu[0, x] - \nu[0, x]| dx$$

So if $f(z) = z^2$ we have

We have

$$\begin{aligned}
 H'(X_0|Z_{-1}=z) &\triangleq \frac{d}{dz} H(X_0|Z_{-1}=z) \\
 &= -\left[\sum_{x \in A} p'(x|z) \log p(x|z) + p'(x|z) \log e \right] \\
 &= -\left[\sum_{x \in A} p'(x|z) \log p(x|z) \right] \quad (41)
 \end{aligned}$$

and

$$\begin{aligned}
 p(x|z) &= z[\gamma_1(x)(1-\beta) + \gamma_0(x)\beta] + (1-z)[\gamma_1(x)\alpha + \gamma_0(x)(1-\alpha)] \\
 &\geq \min\{[\gamma_1(x)(1-\beta) + \gamma_0(x)\beta], [\gamma_1(x)\alpha + \gamma_0(x)(1-\alpha)]\} \quad (42)
 \end{aligned}$$

If $p(x|z) = 0$ for some $x \in A$ and $z \in [0,1]$ and $p'(x|z) > 0$ then the theorem does not provide a bound on Δ' . However, if $\alpha, \beta \in (0,1)$

then

$$p(x|z) \geq \delta[\gamma_1(x) + \gamma_0(x)] \quad (43)$$

where

$$\delta \triangleq \min\{\alpha, \beta, 1-\alpha, 1-\beta\} \quad (44)$$

So

$$\begin{aligned}
 H'(X_0|Z_{-1}=z) &\leq - \sum_{x \in A} |p'(x|z)| \log[\delta[\gamma_1(x) + \gamma_0(x)]] \\
 &= |\lambda| \sum_{x \in A} |\gamma_1(x) - \gamma_0(x)| \log[\delta[\gamma_1(x) + \gamma_0(x)]] \\
 &\leq 2 |\lambda| \log \delta^{-1} + |\lambda| [\mathcal{K}(\gamma_0) + \mathcal{K}(\gamma_1)] \quad (45)
 \end{aligned}$$

where

$$\mathcal{K}(\gamma_i) \triangleq - \sum_{x \in A} \gamma_i(x) \log \gamma_i(x) \quad .$$

Recall $\mathcal{K}(\gamma_i)$ is assumed to be finite. If we do not have $\alpha, \beta \in (0,1)$ a bound may still be derived if $\gamma_i(x) \geq \epsilon > 0$, for all $x \in A$ and $i = 0,1$. Note

$$\Delta \leq 2\bar{\rho}(\bar{\mu}, \mu^*) , \quad (34)$$

hence the mean-square estimation error may be computed to any desired accuracy.

Under certain assumptions the entropy of the two state composite sources may be computed using the approximate stationary distribution. Lemma 1.1 and (1) imply that

$$\underline{X}^- \rightarrow \hat{Z}_{-1} \rightarrow X_0 \quad (35)$$

is a Markov chain, where $\underline{X}^- \triangleq (X_{-1}, X_{-2}, \dots)$. Since $\hat{Z}_{-1}$ is a function of $\underline{X}^-$ it follows from (35) that if $z = E[Z_{-1} | \underline{X}^- = \underline{x}^-]$ then

$$\begin{aligned} H(X_0 | \underline{X}^- = \underline{x}^-) &= H(X_0 | \hat{Z}_{-1} = z) \\ &= - \sum_{x \in A} p(x|z) \log p(x|z) . \end{aligned} \quad (36)$$

Then if μ^* is the stationary distribution of $\hat{Z}_{-1}$, the entropy of source θ is

$$\mathcal{H}_c(\theta) = \int_0^1 H(X_0 | \hat{Z}_{-1} = z) \mu^*(dz) . \quad (37)$$

If $\bar{\mu}$ is the computed distribution, we define

$$\begin{aligned} \hat{\mathcal{H}}_c(\theta) &= \int_0^1 H(X_0 | \hat{Z}_{-1} = z) \bar{\mu}(dz) \\ &= \sum_{i=1}^n H(X_0 | \hat{Z}_{-1} = z) \bar{\mu}(\{y_i\}) \end{aligned} \quad (38)$$

where

$$y_i \triangleq n^{-1}(i - \frac{1}{2}) . \quad (39)$$

We now use Theorem 1.3 to bound

$$\Delta' \triangleq |\mathcal{H}_c(\theta) - \hat{\mathcal{H}}_c(\theta)| . \quad (40)$$

that the alphabet must be finite in this case. Under this assumption $p(x|z) \geq \epsilon$ so

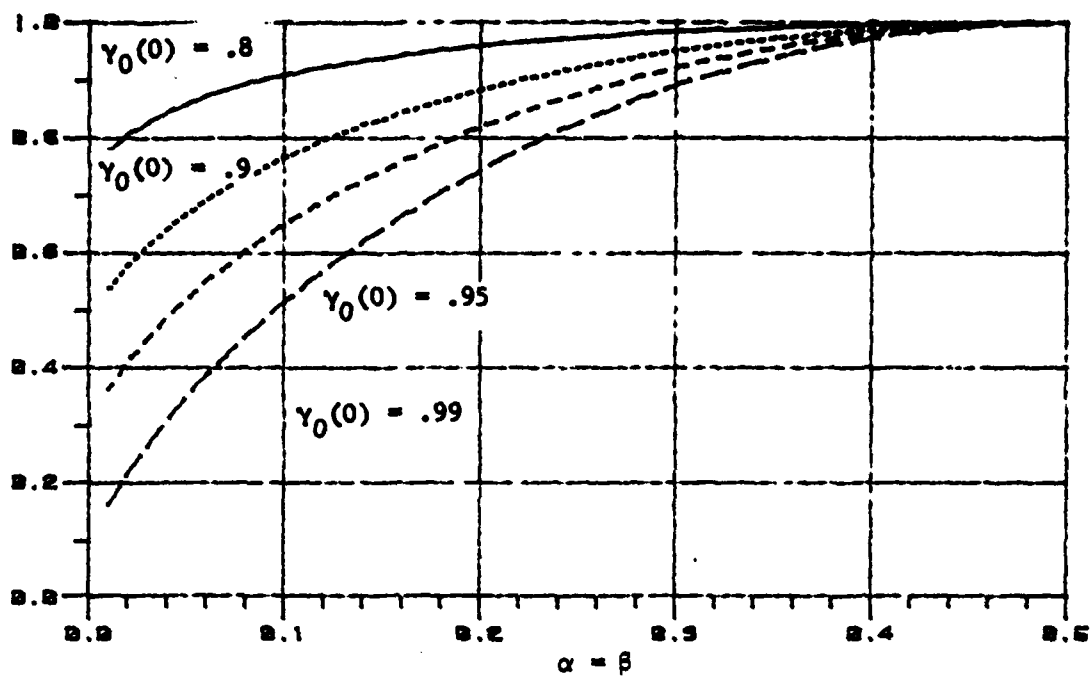
$$H'(X_0|Z_{-1}=z) \leq 2 |\lambda| \log \epsilon^{-1} . \quad (46)$$

In both of these cases Theorem 1.3 implies

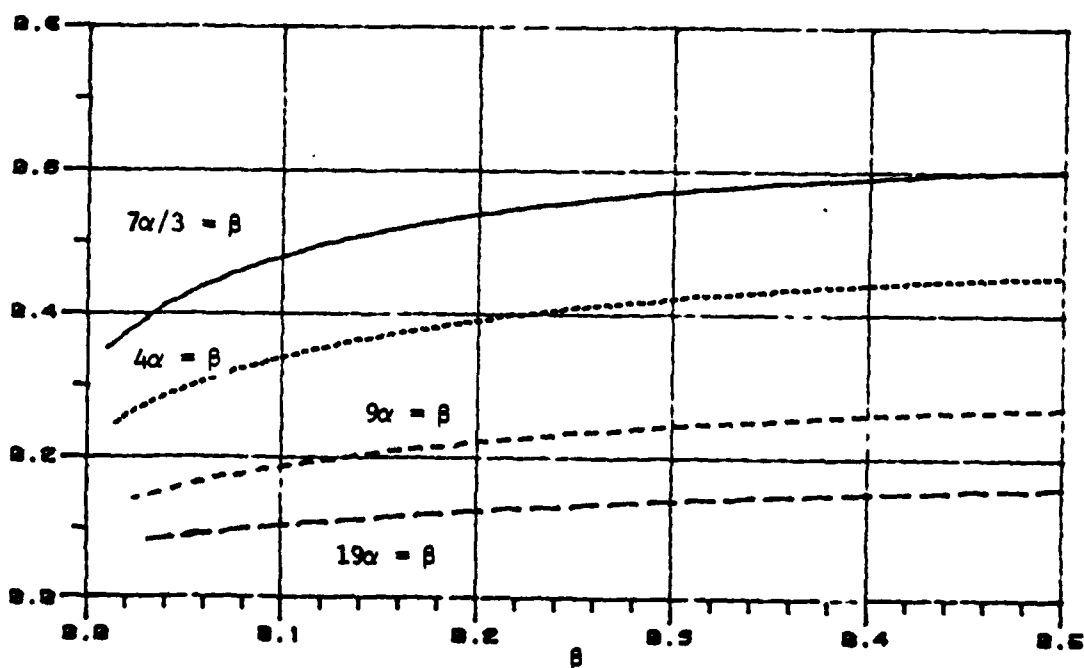
$$\Delta' \leq K \bar{\rho}(\bar{\mu}, \mu^*) \quad (47)$$

where $K < \infty$ depends only on the parameters of the source, so the entropy may be computed to any desired accuracy. Note that the complexity of the computation is the same as that of the computation of the stationary distribution. This algorithm was implemented and the entropy was computed for some two-state composite sources with binary alphabets. In Fig. 3a a family of curves is given. In each curve $\gamma_0(0) = \gamma_1(1)$ is fixed and $\alpha = \beta$ varies from 0 to .5. The entropy increases to one as the switching probabilities increase as would be expected. The same curves result if α and β are replaced by $1-\alpha$ and $1-\beta$. In Fig. 3b $\gamma_0(1) = .001$ and $\gamma_1(1) = .5$ for all curves. The ratio $\alpha/(\alpha+\beta)$ is fixed in each curve, and β varies from 0 to .5. So in each curve the proportion of time spent in state 1 is $\alpha/(\alpha+\beta)$. Again the entropy increases as the switching probabilities increase.

The $\bar{\rho}$ -convergence result may also be used to show that estimates which are derived using different initial estimates of the state converge. Consider two different initial state estimates, z_0 and $\hat{z}_0$. If $z_i(\underline{x}^i)$ and $\hat{z}_i(\underline{x}^i)$ are the estimates at time i derived using the recursion (6) when $\underline{x}^i = (x_1, \dots, x_i)$ is the output of the source, the following theorem shows that these estimates converge on the average. Define $\tilde{p}(\underline{x}^i|z_0) = \prod_{j=0}^{i-1} p(x_{j+1}|z_j(\underline{x}^j))$, where $p(x|z)$ is from (8). So $\tilde{p}(\underline{x}^i|z)$ is the probability that $\underline{x}^i$ is output given initial estimate z_0 .



a) $\alpha = \beta, \gamma_0(0) = \gamma_1(1)$



b) $\gamma_0(0) = .999, \gamma_1(0) = .5$

Figure 3. Entropies of some two-state binary composite sources.

Theorem 1.4. With $z_1(\underline{x}^1)$ and $\hat{z}_1(\underline{x}^1)$ as above (and $\lambda \triangleq 1 - \alpha - \beta$)

$$F_1(z_0, \hat{z}_0) \triangleq \sum_{\underline{x}^1 \in A^1} |z_1(\underline{x}^1) - \hat{z}_1(\underline{x}^1)| \tilde{p}(\underline{x}^1 | z_0) \leq |\lambda|^1 \cdot |z_0 - \hat{z}_0| \quad (48)$$

Proof. See Appendix A.

Corollary 1.2. If π is any joint distribution of z_0 and $\hat{z}_0$ then

$$E_{\pi}[F_1(z_0, \hat{z}_0)] = \int_0^1 \int_0^1 F_1(z_0, \hat{z}_0) d\pi \quad (49)$$

$$\leq |\lambda|^1 \quad (50)$$

since $|z_0 - \hat{z}_0| \leq 1$ for all z_0 and $\hat{z}_0$.

Theorem 1.4 implies that the recursive computation of the state estimate is stable. That is, suppose that some error e_i , where $|e_i| \leq \epsilon$, is introduced in the computation of the i -th estimate. Then if the initial estimate in the computation differs by some e_0 from the actual initial estimate (i.e., the estimate derived from observations of all past source outputs), the average error after i steps is bounded by

$$\sum_{\underline{x}^1 \in A^1} |z_1(\underline{x}^1) - \hat{z}_1(\underline{x}^1)| \tilde{p}(\underline{x}^1 | z_0) \leq |\lambda|^1 \cdot |e_0| + \epsilon [1 - |\lambda|]^{-1} \quad (51)$$

Here $\hat{z}_1(\underline{x}^1)$ includes the computational errors e_i .

Now suppose that the parameters of the composite source are not known precisely. That is, suppose that the source is $\theta = \{\alpha, \beta, \gamma_0, \gamma_1\}$ and we use the parameters for another source $\varphi = \{\alpha', \beta', \gamma'_0, \gamma'_1\}$ in the recursion (6). Under the assumption that the parameters for θ and φ are within ϵ the average error in the estimate derived using the parameters for φ is of order ϵ . Here we must assume that $\theta, \varphi \in \Lambda'(\delta)$ for some $\delta > 0$ where

$$\Lambda'(\delta) \triangleq \{\theta \in \Lambda: p(x|z) \geq \delta, \forall x \in \Lambda, z \in [0,1]\}. \quad (52)$$

This condition is satisfied if, for example, $\gamma_i(x) \geq \delta > 0$, for all $x \in \Lambda$ and $i = 0,1$. Again this implies that the alphabet Λ is finite. We also include computational errors e_i , $|e_i| \leq \epsilon$. Let z_0 and $\hat{z}_0$ be two initial estimates of the state Z_0 . Further let $z_1(\underline{x}^1)$ and $\hat{z}_1(\underline{x}^1)$ be the estimates derived from these initial estimates using the recursions for θ and φ respectively. Note that now these estimates are derived from different recursions and that $\hat{z}_1(\underline{x}^1)$ includes computational errors so

$$\hat{z}_1(\underline{x}^1) = \hat{f}_{x_1}(\hat{z}_{1-1}(\underline{x}^{1-1})) + e_1 \quad (53)$$

where $\hat{f}$ is defined as f , (6), (7), but with the parameters for φ . Then the following is true.

Theorem 1.5. If $\theta, \varphi \in \Lambda'(\delta)$ then

$$\begin{aligned} F_1(z_0, \hat{z}_0) &\triangleq \sum_{\underline{x}^1} |\hat{z}_1(\underline{x}^1) - z_1(\underline{x}^1)| \tilde{p}(\underline{x}^1 | z_0) \\ &\leq |\lambda_\theta|^1 + K_\epsilon \{1 - |\lambda_0|\}^{-1} \end{aligned} \quad (54)$$

where

$$K_\epsilon \triangleq \delta^{-2} [3\epsilon + 3\epsilon^2 + \epsilon^3] + \epsilon,$$

$p(\underline{x}^1 | z_0)$ is the probability that $\underline{x}^1$ is output from source θ if the initial estimate is z_0 , and $\lambda_\theta \triangleq 1 - \alpha - \beta$.

Proof. See Appendix A.

So the estimation procedure is robust; that is, small errors in source parameters do not cause unbounded errors in the estimates.

1.3 Convergence of State Estimates for S-State Composite Sources

Now consider the more general case where the Markov chain has state space $\mathcal{S} = \{1, 2, \dots, S\}$ and selects one of S subsources which are discrete memoryless sources with alphabet A . Let $\gamma_s(x)$ be the probability that a letter $x \in A$ is output given that the Markov chain is in state $s \in \mathcal{S}$. Let Λ denote the class of such sources for a given A and $\mathcal{S}$. Define a state (row) vector $\underline{Z}(i) = (Z_1(i), Z_2(i), \dots, Z_S(i))$ by $Z_s(i) = 1$ if the chain is in state s at time i and $Z_s(i) = 0$ otherwise. Let $Q = \{q(i|j)\}$ be the state transition matrix. We define

$$\hat{\underline{Z}}(i+1) = E[\underline{Z}(i) | X_1, X_{i-1}, \dots] \quad (55)$$

where X_i is the output at time i . So $\hat{\underline{Z}}(i)$ is the conditional mean estimate of $\underline{Z}(i)$ given the outputs up to time i . A recursive equation for $\hat{\underline{Z}}$ is [6],

$$\hat{\underline{Z}}(i+1) = \hat{\underline{Z}}(i)T(x)[\hat{\underline{Z}}(i)T(x)\underline{1}]^{-1} \quad (56)$$

where $X_{i+1} = x$ is the source output,

$$T(x) \triangleq Q P(x), \quad (57)$$

$$P(x) = \begin{bmatrix} \gamma_1(x) & & & \\ & \gamma_2(x) & & \\ & & \ddots & \\ & & & \gamma_S(x) \end{bmatrix} \quad (58)$$

is a diagonal matrix and $\underline{1}$ is a column vector of 1's. The probability that source output $X_{i+1} = x$ given $\hat{\underline{Z}}(i) = z$ is

$$p(x|z) \triangleq zT(x)\underline{1} \quad (59)$$

Let $\underline{x} = (x_1, \dots, x_n)$, $x \in A$, and define a matrix $\tilde{T}$ by

$$\tilde{T}(\underline{x}) = \prod_{i=1}^n T(x_i) \quad (60)$$

Then if $\hat{\underline{z}}(0) = \underline{z}$ and $\underline{x}$ consists of the first n outputs, the n -th state estimate is

$$\hat{\underline{z}}(n) = \underline{z} \tilde{T}(\underline{x}) [\underline{z} \tilde{T}(\underline{x}) \underline{1}]^{-1} \quad (61)$$

and the probability of $\underline{x}$ given $\underline{z}$ is

$$\tilde{p}(\underline{x}|\underline{z}) = \underline{z} \tilde{T}(\underline{x}) \underline{1} \quad (62)$$

A source $\theta \in \Lambda$ is specified by Q and $\{P(x): x \in A\}$ so we write $\theta = \{Q, P(\cdot)\}$.

Let $IP(\mathcal{A})$ denote the set of probability distributions on $\mathcal{A}$. We now show that under certain conditions if $\underline{z}$ and $\hat{\underline{z}}$ are in $IP(\mathcal{A})$, then the estimates generated from (61) converge. Define

$$\hat{\Lambda}(\epsilon) = \{\theta \in \Lambda: q(i|j) \geq \epsilon > 0, i, j \in \mathcal{A}\}. \quad (63)$$

Then we have the following theorem.

Theorem 1.6. Let $\theta \in \hat{\Lambda}(\epsilon)$. If $\underline{z}$ and $\hat{\underline{z}}$ are probability vectors on $\mathcal{A}$ such that $\tilde{p}(\underline{x}|\underline{z}) > 0$ and $\tilde{p}(\underline{x}|\hat{\underline{z}}) > 0$, $\underline{x} = (x_1, \dots, x_n)$, then

$$\|\underline{z} \tilde{T}(\underline{x}) [\tilde{p}(\underline{x}|\underline{z})]^{-1} - \hat{\underline{z}} \tilde{T}(\underline{x}) [\tilde{p}(\underline{x}|\hat{\underline{z}})]^{-1}\| \leq \epsilon^{n-1} \quad (64)$$

where

$$\epsilon \triangleq \frac{(1 - (S-1)\epsilon)^2 - \epsilon^2}{(1 - (S-1)\epsilon)^2 + \epsilon^2} \quad (65)$$

and $\|\cdot\|$ is the norm defined by $\|\underline{u} - \underline{v}\| = \max\{|u_i - v_i|: 1 \leq i \leq n\}$.

Proof. See Appendix A.

The rate of convergence here is not as fast as that of the average convergence result for the two-state composite source, but the convergence bound holds for any sequence of outputs and not merely an average. The restriction that we must have $\tilde{p}(\underline{x}|\underline{z})$ and $\tilde{p}(\underline{x}|\hat{\underline{z}})$ positive is of no real importance, since if $\tilde{p}(\underline{x}|\underline{z})$ is zero, this means that the estimate $\underline{z}$ is incorrect so we may choose a new initial estimate $\underline{z}'$ such that $\tilde{p}(\underline{x}|\underline{z}') > 0$. A more important drawback here is that the theorem does not imply that the estimates converge at each step (in fact they do not in general), but only that after n steps they are within ϵ^{n-1} . The theorem does not imply that a computed estimate remains close to the true estimate despite small computational errors at each step.

Theorem 1.6 also applies in the more general case where the transition matrix Q depends on the current output x . So we have a family of matrices $\{Q_x: x \in A\}$. If we assume that the elements of Q_x are at least ϵ for all $x \in A$, then the theorem holds. The only change necessary in the proof is that Q is replaced by Q_x .

1.4 Generalization to Arbitrary Subsources

Some of the estimation results also hold for memoryless sources (not necessarily finite entropy) having an arbitrary alphabet A . Consider first the two state composite source. Where previously we assumed that the alphabet A was countable and that the subsources had finite entropy, here we assume that the sources are specified by two probability measures P_0 and P_1 on an alphabet A . If we define $\pi = \frac{1}{2}(P_0 + P_1)$ then the Radon-Nikodym derivative $\frac{dP_i}{d\pi}$ exists, $i = 0, 1$. Then given the i -th state estimate $\hat{z}_i = z$ and the $(i+1)$ -st source output X_{i+1} we have

$$\hat{z}_{i+1} = F_x(z) = \frac{\frac{dP_1}{d\pi}(x) \eta_1(z)}{\sum_{i=0} \frac{dP_i}{d\pi}(x) \eta_i(z)} \quad (66)$$

which is Bayes rule for this case. Define

$$p(x|z) = \sum_{i=0} \frac{dP_i}{d\pi} \eta_i(z) \quad (67)$$

so that

$$P\{X_{i+1} \in B | \hat{z}_i = z\} = \int_B p(x|z) \pi(dx) \quad (68)$$

Then if μ_0 is the distribution of $\hat{z}_0$, the distribution of $\hat{z}_1$ is given by

$$\mu_1(B) = \int_A \int_{f_x^{-1}(B)} p(x|z) \mu_0(dz) \pi(dx) \quad (69)$$

If we use the recursion (69) in place of (9) then Theorem 1.1 holds for these generalized subsources. The only modifications necessary to the proof of Theorem 1.1 are to replace $\gamma_i(x)$ by $\frac{dP_i}{d\pi}$, $i = 0, 1$, and to replace all summations over the alphabet A by integration with respect to the measure π . Corollary 1.1 and Theorem 1.2 follow directly from Theorem 1.1 so we know that the state estimate $\hat{z}_i$ has a unique stationary distribution. However, the computation of an approximation to this stationary distribution may not be performed as it was in Section 1.2 because the alphabet is not finite. The average convergence of Theorem 1.4 also holds, if we modify the proof in the same way as the proof of Theorem 1.1. Theorem 1.5 is not easily generalized though, as it was necessary to assume finite alphabet size.

The convergence result (Theorem 1.6) for S-state composite sources also generalizes. Let P_i , $i = 1, 2, \dots, S$ be probability measures on A for the subsources. Then if we define $\pi = S^{-1} \sum_{i=1}^S P_i$, the Radon-Nikodym derivative $\frac{dP_i}{d\pi}$ exists, $i = 1, 2, \dots, S$. If we replace $\gamma_i(x)$ by $\frac{dP_i}{d\pi}(x)$ in the definition of $P(x)$ (58) then Theorem 1.6 holds, and the same proof is valid.

1.5 A Coding Technique for Composite Sources

Let θ be a composite source as in the previous section. The switching process has state space $\mathcal{S} = \{1, 2, \dots, S\}$ and each subsource is a discrete memoryless source with alphabet A. If the state of the switching process is s then the probability of the source output x is $\gamma_s(x)$, independently of previous states and source outputs. Let $\mathcal{P}(\mathcal{S})$ be the set of probability distributions on $\mathcal{S}$ and define $\underline{e}^j \in \mathcal{P}(\mathcal{S})$ to be the probability (row) vector whose j -th element is one. If the switching process is in state j at time t we define the state $\underline{z}(t) = \underline{e}^j$. The transition probability matrix for the switching process will depend on the current state and the current source output. So we define

$$q_x(i|j) = P\{\underline{z}(t+1) = \underline{e}^i | \underline{z}(t) = \underline{e}^j, X(t) = x\} \quad (70)$$

and

$$Q_x = \{q_x(i|j) : i, j \in \mathcal{S}\}. \quad (71)$$

We do not require the elements of Q_x to be bounded by some $\epsilon > 0$ (as was the case in the previous section). Note that this class includes unifilar Markov sources, that is, sources where the next state is a deterministic function of the current state and source output. For these sources the elements of the matrices Q_x , $x \in A$, are either zero or one.

We now construct a variable rate code for a given composite source and bound its redundancy uniformly over all initial state estimates. The codes considered here are fixed-length to variable-length (FL-VL) codes, so they encode fixed-length blocks of source outputs into variable-length binary codewords. The blocklength of a FL-VL code is the number of source letters encoded in a block. The n -th order entropy of source θ given initial state estimate $\hat{z}(0) = \underline{z}$ is given by

$$H_n(\theta, \underline{z}) = -n^{-1} \sum_{\underline{x} \in A^n} \tilde{p}(\underline{x}|\underline{z}) \log \tilde{p}(\underline{x}|\underline{z}) \quad (72)$$

where

$$\tilde{p}(\underline{x}|\underline{z}) \triangleq \underline{z} \tilde{T}(\underline{x}) \mathbf{1} \quad (73)$$

and $\tilde{T}$ is as defined in (60). So $H_n(\theta, \underline{z})$ is a lower bound on the rate of any blocklength n code for source θ and initial state estimate $\underline{z}$. Let $l_n(\underline{x})$ be the length of the binary codeword for the output block $\underline{x} \in A^n$. Then the rate of the code is

$$R_n(\theta, \underline{z}) \triangleq n^{-1} \sum_{\underline{x} \in A^n} \tilde{p}(\underline{x}|\underline{z}) l_n(\underline{x}) \quad (74)$$

and the redundancy is

$$r_n(\theta, \underline{z}) \triangleq R_n(\theta, \underline{z}) - H_n(\theta, \underline{z}) \quad (75)$$

If we let $\underline{z} = (z_1, \dots, z_S)$ then

$$\tilde{p}(\underline{x}, \underline{z}) = \underline{z} \tilde{T}(\underline{x}) \mathbf{1} \quad (76)$$

$$= \sum_{i=1}^S z_i p(\underline{x} | \underline{e}^i) \quad (77)$$

$$\leq \max\{p(\underline{x} | \underline{e}^i) : i \in \mathcal{M}\} \quad (78)$$

So to design a code for θ and $\hat{z}(0) = \underline{z}$ we first design codes for initial state estimates $\underline{e}^i$, $i \in \mathcal{I}$, and combine these S codes into a single code by prefixing each codeword with $\lceil \log S \rceil$ bits. The code for initial state $\underline{e}^i$ is the Shannon code for probabilities $\tilde{p}(\underline{x}|\underline{e}^i)$, so the length of the codeword for $\underline{x}$ is

$$l_n^{(i)}(\underline{x}) = \lceil -\log \tilde{p}(\underline{x}|\underline{e}^i) \rceil \quad (79)$$

$$\leq 1 - \log \tilde{p}(\underline{x}|\underline{e}^i) . \quad (80)$$

The codeword for $\underline{x}$ in the combined code is then the shortest of the S possible codewords, so the length function of the combined code is

$$l_n(\underline{x}) = \min\{l_n^{(i)}(\underline{x}) : i \in \mathcal{I}\} + \lceil \log S \rceil \quad (81)$$

$$\leq -\log[\max\{\tilde{p}(\underline{x}|\underline{e}^i) : i \in \mathcal{I}\}] + 1 + \lceil \log S \rceil \quad (82)$$

$$\leq -\log[\tilde{p}(\underline{x}|\underline{z})] + 1 + \lceil \log S \rceil ; \forall \underline{z} \in \mathbb{P}(\mathcal{I}) . \quad (83)$$

The rate of the code when applied to θ with $\hat{z}(0) = \underline{z}$ is

$$R_n(\theta, \underline{z}) \leq n^{-1}\{1 + \lceil \log S \rceil - \sum_{\underline{x} \in A^n} \tilde{p}(\underline{x}|\underline{z}) \log \tilde{p}(\underline{x}|\underline{z})\} \quad (84)$$

$$= H_n(\theta, \underline{z}) + n^{-1}\{1 + \lceil \log S \rceil\} , \quad (85)$$

and so its redundancy is bounded by

$$r_n(\theta, \underline{z}) \leq n^{-1}\{1 + \lceil \log S \rceil\} , \quad (86)$$

for all $\underline{z} \in \mathbb{P}(\mathcal{I})$.

One problem with this coding technique involves the propagation of errors. To determine the codeword lengths l_n the probabilities of the codewords must be determined. This requires n matrix multiplications, and there is no guarantee that errors will not propagate. Theorem 1.6 implies that the effect of an error on one step will decrease exponentially, but does not imply that the effect of small errors made in each step will remain small. Propagation of errors is not a problem when coding for a unifilar Markov source with finite space and alphabet. For such a source the probability of a source vector $\underline{x}$ given initial state s_0 is

$$P(\underline{X} = \underline{x} | s_0) = \prod_{i=1}^n p(x_i | s_{i-1}) \quad (87)$$

$$= \prod_{x \in A} \prod_{s \in \mathcal{S}} p(x | s)^{N[(x,s), (\underline{x}, s_0)]} \quad (88)$$

where $N[(x,s), (\underline{x}, s_0)]$ is the number of times in the block $\underline{x}$ that the letter x occurs when the source is in state s given that the initial state is s_0 . The product (88) may be computed using at most $|A| \cdot |\mathcal{S}|$ multiplications for any n , so the effect of computational errors need not increase as n becomes large. Some further convergence result is required to show that the code for composite sources is implementable, although in view of the convergence result of Theorem 1.6 it is probable that the computation is stable.

Ott [7] considers the same coding problem but assumes that the encoder and decoder have the initial state estimate for the source. The code he constructs is simply the Huffman code for the source given a specific initial state estimate (and is optimal for that estimate), but it is not universal with respect to the initial state estimate. He does not prove any convergence results which would indicate that the computation is stable.

One modification which improves the code performance is as follows. Since only 2^n codewords of the $S \times 2^n$ possible codewords are used, the additional codewords may be removed and the remaining ones shortened. This technique is employed in Section 4 of [11]. Define

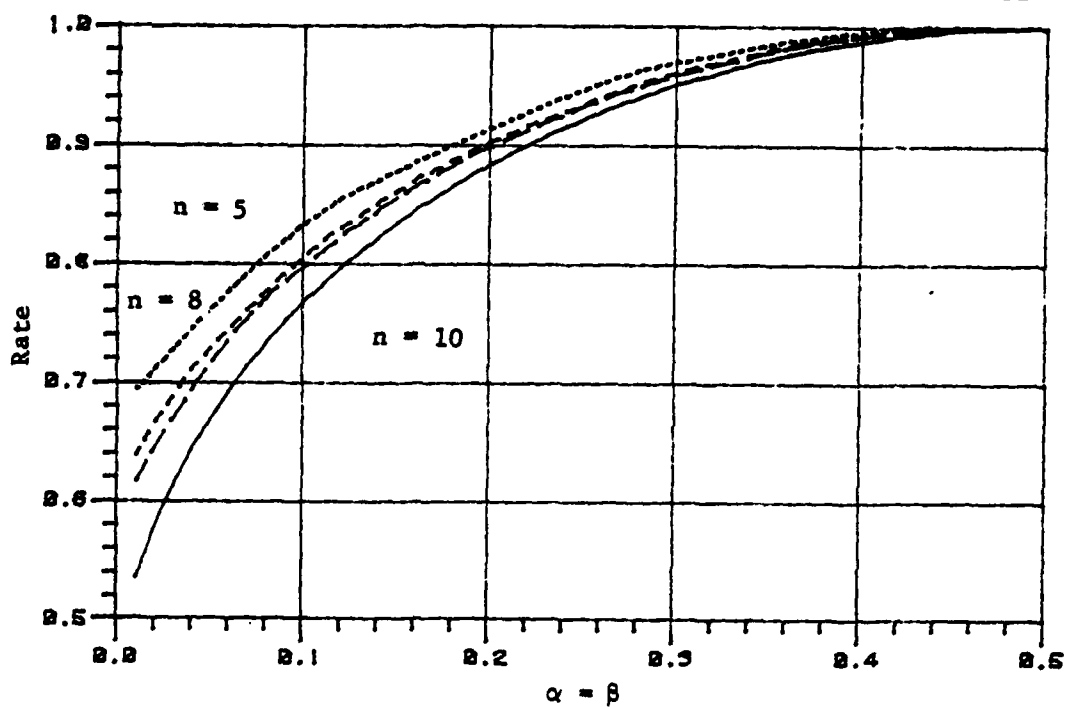
$$p_n^*(\underline{x}) = \frac{2^{-l_n(\underline{x})}}{\sum_{\underline{y} \in A^n} 2^{-l_n(\underline{y})}} \quad (89)$$

Then

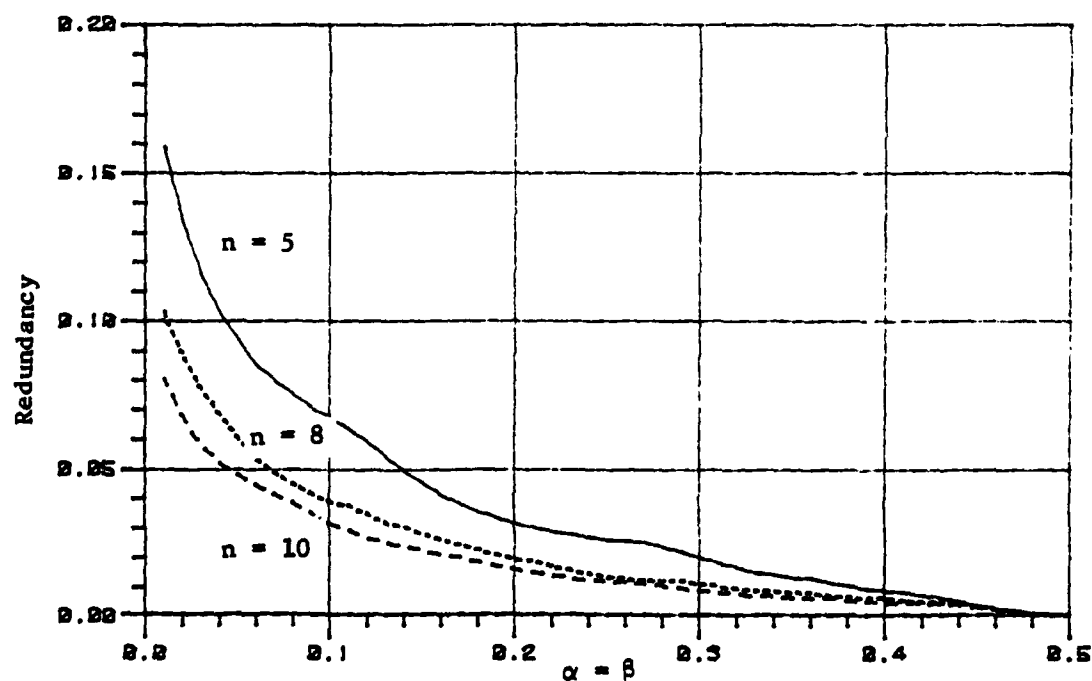
$$l_n(\underline{x}) \geq \lceil -\log p_n^*(\underline{x}) \rceil \quad (90)$$

so the Shannon code for p_n^* performs at least as well as l_n .

The performance of codes with blocklengths $n = 5, 8$, and 10 which incorporate this modification are presented in Fig. 4. The sources for which the codes are designed have $\gamma_0(0) = \gamma_1(1) = .9$ and $\alpha = \beta$ between 0 and $.5$. Each curve gives the performance of a set of codes of the same block length. The rates and source entropy are given in Fig. 4a, and the redundancies in Fig. 4b.



a) Rates of codes with blocklengths $n = 5, 8, \text{ and } 10$



b) Redundancy of codes with blocklengths $n = 5, 8, \text{ and } 10$

Figure 4. Performance of the coding technique for two-state binary composite sources with $\gamma_0(0) = \gamma_1(1) = 0.9$ and $\alpha = \beta$ in $[0, .5]$.

1.6 Stability of Coding and Universal Coding for Composite Sources

If the composite source consists of only two subsources as in Section 1.2 we may show that the effect of computational errors on the redundancy may be made small. This result is implied by the following theorem which bounds the mismatch redundancy; that is, the redundancy which results when a code designed for a source φ is applied to another source θ . The theorem includes the effect of computational errors. We return to the notation of Section 1.2. Let $\theta = \{\alpha, \beta, \gamma_0, \gamma_1\}$ and $\varphi = \{\alpha', \beta', \gamma'_0, \gamma'_1\}$ be two composite sources (recall that $\alpha = q(1|0)$ and $\beta = q(0|1)$ are the transition probabilities for the switching process). We assume that $\theta, \varphi \in \Lambda'(\delta)$ where

$$\Lambda'(\delta) \triangleq \{\theta \in \Lambda: p(x|z) \geq \delta, x \in A, z \in [0,1]\} \quad (91)$$

Let $\tilde{p}_\theta(\underline{x}^1|z_0)$ be the probability that $\underline{x}^1 = (x_1, \dots, x_1)$ is the output of source θ given initial estimate z_0 , and similarly for $\tilde{p}_\varphi(\underline{x}^1|z_0)$. Let $\hat{z}_1(\underline{x}^1)$ be the estimate of the state used in designing the code. So $\hat{z}_1(\underline{x}^1)$ includes computational errors e_1 as in Theorem 1.5 and we again assume that $|e_1| \leq \epsilon$. Then if the initial estimate is z_0 the mismatch redundancy is

$$r_n(\ell_n, \theta) = n^{-1} \left[1 + \sum_{\underline{x}^n \in A^n} \tilde{p}_\theta(\underline{x}^n|z_0) \left[\min_{k=0,1} \{-\log \tilde{p}_\varphi(\underline{x}^n|k)\} \right] + \log \tilde{p}_\theta(\underline{x}^n|z_0) \right] \quad (92)$$

Theorem 1.7. If $\theta, \varphi \in \Lambda'(\delta)$ and corresponding parameters (i.e., switching probabilities and subsources statistics) for sources θ and φ are within ϵ , then if ℓ_n is the code designed for source φ we have

$$r_n(\ell_n, \theta) \leq K n^{-1} + \hat{K} \epsilon \quad (93)$$

where

$$K \triangleq 2 + \delta^{-1} \log e [1 - |1 - \alpha - \beta|]^{-1} \quad (94)$$

and

$$K' = \delta^{-1} \log e [1 - |1 - \alpha - \beta|]^{-1} \{ \delta^{-2} [3\epsilon + 3\epsilon^2 + \epsilon^3] + 5\epsilon + 2\epsilon^2 \} \quad (95)$$

Proof. See Appendix A.

We may use this mismatch result to construct a sequence of minimax universal codes for any subset Φ of $\Lambda^1(\delta)$. A sequence of codes $\{l_n^*: n = 1, 2, \dots\}$ is said to be minimax universal for a class of sources Φ if the redundancy

$$r_n(l_n^*, \theta) \rightarrow 0 \quad (96)$$

uniformly on Φ as $n \rightarrow \infty$. We construct the code as follows. The alphabet A is assumed finite so let $A = \{1, 2, \dots, J\}$. Let i, j and $K(m, x)$; $m = 0, 1$, $x = 1, 2, \dots, J-1$ be nonnegative integers less than n . Define a set

$$B_n(i, j, K(\cdot, \cdot)) = \{ \theta \in \Phi: \alpha \in [in^{-1}, (i+1)n^{-1}], \beta \in [jn^{-1}, (j+1)n^{-1}], \\ \gamma_m(x) \in [K(m, x)n^{-1}, (K(m, x)+1)n^{-1}] \}. \quad (97)$$

Note that B_n has dimension $2 + 2(J-1) = 2J$ since each subspace is specified by $J-1$ parameters. From each non-empty set B_n choose an element φ called the design point source. The number of design point sources is bounded by n^{2J} , since there are at most this number of sets B_n . A Shannon code $l_{n, \varphi}$ is then constructed for each of the design point sources as in Section 1.5. A prefix of length $\lceil 2J \log n \rceil$ which identifies φ is attached to the codewords in the code $l_{n, \varphi}$. The universal code is then constructed by combining these codes. The universal code is uniquely decodable since the prefix

specifies φ and since the codes $l_{n,\varphi}$ are uniquely decodable. The encoding procedure is simply to choose the shortest codeword of the n^{2J} possible codewords for a given output block $\underline{x}$, so the length function for the universal code is

$$l_n^*(\underline{x}) = \lceil 2J \log n \rceil + \min_{\varphi} \{l_{n,\varphi}(\underline{x})\} \quad (98)$$

for any $\theta \in \mathcal{Q}$ there exists a design point source φ whose parameters are within n^{-1} of the parameters of θ . Let φ be this design point source for θ . Then we have

$$r_n(l_n^*, \theta) \leq n^{-1} \lceil 2J \log n \rceil + r_n(l_{n,\varphi}, \theta) \quad (99)$$

Since $r_n(l_{n,\varphi}, \theta)$ is the mismatch redundancy of Theorem 1.7 with $\epsilon = n^{-1}$ we have

$$r_n(l_n^*, \theta) \leq n^{-1} \{ \lceil 2J \log n \rceil + \tilde{K} + \tilde{K} \} \quad (100)$$

for all $\theta \in \mathcal{Q}$ and the sequence of codes l_n^* is minimax universal.

If some of the source parameters are fixed for all $\theta \in \Lambda$ so that Λ has dimension M , where $M < 2J$, then $2J \log n$ is replaced by $M \log n$ in (100).

To illustrate this procedure, Fig. 5 contains a graph of the redundancy of a blocklength 8 code for the class of binary two-state composite sources with $\gamma_0(0) = \gamma_1(1) = .9$ and $\alpha = \beta$ in $[0, 1]$. The code was constructed by combining codes designed for $\alpha = \beta = .05, .30, .70$, and $.95$ respectively. The redundancies of the codes for $\alpha = .05$ and $.30$ are also graphed over the class of sources. If these curves are reflected about $\alpha = .5$ then they become the curves for $\alpha = .95$ and $.70$. Note that the maximum redundancy of the combined code is much less than those of the other codes.

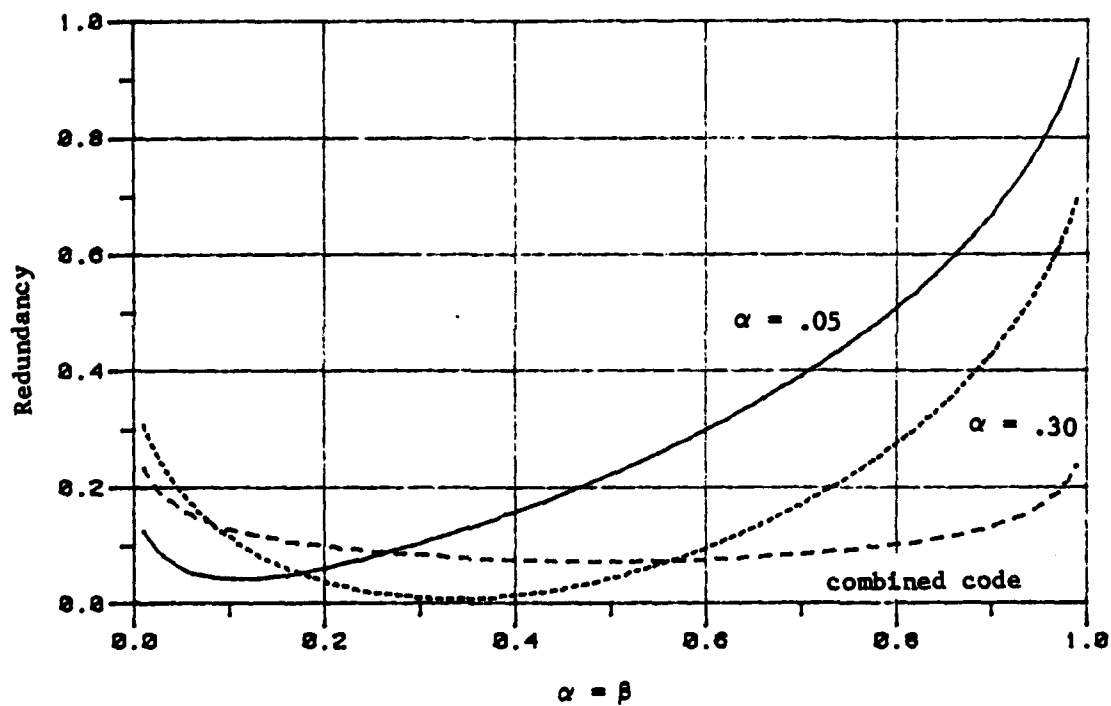


Figure 5. Redundancies of three codes over the class of two-state binary composite sources with $\gamma_0(0) = \gamma_1(1) = 0.9$ and $\alpha = \beta$.

1.7 Coding for an Infinite-State Composite Source

The coding technique derived in Section 1.5 applied to composite sources with a finite number of subsources. We now construct a code for a certain type of composite source with an infinite state space, and show that the rate of this code approaches the entropy of the source.

The state space $\mathcal{S}$ is the class of all memoryless sources with alphabet $A = (1, 2, \dots, J)$. We define $\mathcal{S}$ such that if $\underline{y} = (y_1, \dots, y_{J-1})$ is a source in $\mathcal{S}$ then

$$\gamma_{\underline{y}}(k) \triangleq P\{X_i = k | Z_i = \underline{y}\} = y_k \quad (101)$$

where

$$y_J \triangleq 1 - \sum_{k=1}^{J-1} y_k. \quad (102)$$

At each integer time i the switching process Z_i changes with probability α . If it does change then it takes on a new value according to a probability measure P^* on $\mathcal{S}$ which does not depend on the previous state. So each time the source changes state the effect of the past states is eliminated. We first assume that P^* has a density which we denote z^* . So if $Z_i = \hat{\underline{y}} \in \mathcal{S}$ then $Z_{i+1} = \hat{\underline{y}}$ with probability $1 - \alpha$, and

$$P\{Z_{i+1} \in B\} = P^*(B) \quad (103)$$

with probability α , where B is a subset of $\mathcal{S}$.

The estimate of the state Z_i given the past outputs $(x_1, x_{i-1}, \dots)$ is a probability measure P^i on $\mathcal{S}$ such that

$$P^i(B) = P\{Z_i \in B | X_1 = x_1, X_{i-1} = x_{i-1}, \dots\}. \quad (104)$$

If we assume that this measure also has a density $\hat{z}_1$ we may derive $\hat{z}_{i+1}$ from $\hat{z}_1$ and X_{i+1} using Bayes rule. Let p_z be the density of Z_1 given $\hat{z}_0 = z_0$. Then we have

$$p_z = \alpha z^* + (1-\alpha)z_0. \quad (105)$$

Further, if $p_{X,Z}$ is the joint density of Z_1 and X_1 given $\hat{z}_0 = z_0$ then

$$\begin{aligned} p_{X,Z}(\underline{y}, k) &= P[X_1 = k | Z_1 = \underline{y}] p_z(\underline{y}) \\ &= y_k p_z(\underline{y}). \end{aligned} \quad (106)$$

So if $X_1 = k$ and $\hat{z}_0 = z_0$ then $\hat{z}_1$ is given by

$$\begin{aligned} \hat{z}_1(\underline{y}) &= f_k(z_0, \underline{y}) = \frac{p_{X,Z}(\underline{y}, k)}{\int p_{X,Z}(\underline{y}, k) d\underline{y}} \\ &= \frac{\tilde{y}_k [\alpha z^*(\underline{y}) + (1-\alpha)z_0(\underline{y})]}{\alpha \int y_k z^*(\underline{y}) d\underline{y} + (1-\alpha) \int y_k z_0(\underline{y}) d\underline{y}} \end{aligned} \quad (107)$$

where $\tilde{\underline{y}} = (\tilde{y}_1, \dots, \tilde{y}_{J-1})$. Since $\hat{z}_1$ is of the form

$$\hat{z}_1(\underline{y}) = y_k [K z^*(\underline{y}) + K' z_0(\underline{y})], \quad (108)$$

where K and K' do not depend on $\underline{y}$, all subsequent densities $\hat{z}_i$ derived from $\hat{z}_1$ will be of the form

$$\hat{z}_i(\underline{y}) = \sum_{\substack{m_1, \dots, m_J \\ \sum m_j \leq i}} \prod_{j=1}^J [y_j]^{m_j} [K(m_1, \dots, m_J) z_0(\underline{y}) + K'(m_1, \dots, m_J) z^*(\underline{y})] \quad (109)$$

where $K(\cdot)$ and $K'(\cdot)$ do not depend on $\underline{y}$. So although the estimate $\hat{Z}_1$ is infinite dimensional, given z_0 and z^* only a finite number of constants $K(\cdot)$ and $K'(\cdot)$ are required to specify $\hat{Z}_1$ for any i . Further, knowledge of the moments of z_0 and z^* is sufficient to compute these constants.

The probability that $X_1 = k$ given $\hat{Z}_0 = z_0$, denoted $p(k|z_0)$, is the denominator of (107), and the probability of a block of source outputs $\underline{x}$, denoted $\tilde{p}(\underline{x}|z_0)$, may be computed by generating $\hat{Z}_1$ recursively from (107).

Given this estimation procedure we construct a code as follows. Compute the probabilities $\tilde{p}(\underline{x}|z^*)$ of output blocks $\underline{x} \in A^n$, where z^* is the density of P^* as previously defined. The code is then the Huffman code for these probabilities. So if the length function of the code is l_n then this code minimizes the redundancy

$$r_n(\theta, z^*) = \sum_{\underline{x} \in A^n} \tilde{p}(\underline{x}|z^*) [l_n(\underline{x}) + \log \tilde{p}(\underline{x}|z^*)] \quad (110)$$

Let $K(\theta) \triangleq H(X_0|X_{-1}, \dots)$ be the entropy of source θ . The probability of $\underline{x}$ given no previous source outputs is $\tilde{p}(\underline{x}|z^*)$ since z^* is the stationary distribution of the switching process. If we define $R_n(\theta)$ to be the average rate of the code l_n when applied to source θ then

$$R_n(\theta) = n^{-1} \sum_{\underline{x} \in A^n} \tilde{p}(\underline{x}|z^*) l_n(\underline{x}) \quad (111)$$

The following theorem gives an upper bound on the average redundancy of the code l_n .

Theorem 1.8. Let $r_n(\theta) = R_n(\theta) - K_c(\theta)$ be the redundancy of the code l_n . Then

$$r_n(\theta) \leq n^{-1} (1 - \frac{1}{\alpha} \log J) \quad (112)$$

Proof. See Appendix A.

We assume that both P^* and the estimates P^i have densities. If they do not the estimation procedure may be modified as follows. Let $\pi = \frac{1}{2}(P^* + P^0)$, where P^0 is the initial estimate. Then $\frac{dP^*}{d\pi}$ and $\frac{dP^i}{d\pi}$ exist for all $i \geq 0$. If we replace z^* and z_0 by $\frac{dP^*}{d\pi}$ and $\frac{dP^0}{d\pi}$ in (105)-(109) and integrate with respect to π , then (107) gives a recursion for $\frac{dP^i}{d\pi}$. The code l_n is then defined as before and the redundancy bound holds.

CHAPTER 2

UNIVERSAL VL-FL CODING FOR MARKOV SOURCES

2.1. Introduction and Review of Previous Results

An efficient universal noiseless source coding technique is presented in [11] for memoryless sources. It is extended to unifilar Markov sources in [12] and [13]. The codes constructed in these papers are fixed-length-to-variable-length (FL-VL) codes; that is, they encode fixed-length blocks of source outputs into variable-length binary codewords. We use the same basic technique to construct universal variable-length-to-fixed-length (VL-FL) codes for unifilar Markov sources. The performance of these VL-FL codes for binary memoryless sources is compared to that of the FL-VL codes constructed in [11]. We show that for medium blocklengths (~ 10) the VL-FL codes perform better and that for long blocklengths (~ 100) they perform about as well as the FL-VL codes.

Next a review of some terminology of universal noiseless coding [11] in a fixed-length-to-variable-length (FL-VL) framework may be helpful. Let Λ be a class of stationary sources. Each $\theta \in \Lambda$ has a probability function p_θ which gives the probability of the various possible strings of outputs.

A FL-VL code of blocklength n maps blocks of n source symbols into variable-length binary sequences. Let $\underline{x} = (x_1, \dots, x_n)$ be a block of source outputs.

A FL-VL code is specified for our purposes by the length function $l_n(\underline{x})$ which gives the length of the codeword for $\underline{x}$. The rate of a FL-VL code applied to a source θ is

$$R_n(l_n, \theta) = n^{-1} \sum_{\underline{x} \in A^n} l_n(\underline{x}) p_\theta(\underline{x}) \quad (113)$$

where A^n is the set of possible n -tuples from source θ . Defining the n -th order per-letter entropy of θ as

$$H_n(\theta) = -n^{-1} \sum_{\underline{x} \in A^n} p_\theta(\underline{x}) \log p_\theta(\underline{x}), \quad (114)$$

the n -th order redundancy of the code is

$$r_n(\ell_n, \theta) = R_n(\ell_n, \theta) - H_n(\theta). \quad (115)$$

Let

$$\hat{r}_n(\ell_n) \triangleq \sup\{r_n(\ell_n, \theta) : \theta \in \Lambda\}. \quad (116)$$

A sequence of codes $\ell_1, \ell_2, \dots$ is weakly universal if

$$R_n(\ell_n, \theta) \rightarrow H(\theta) \quad \forall \theta \in \Lambda \quad (117)$$

as $n \rightarrow \infty$ where $H(\theta) = \lim_{n \rightarrow \infty} H_n(\theta)$ is the entropy of the source θ . It is strongly universal if the convergence of (117) is uniform and minimax universal if $\hat{r}_n(\ell_n) \rightarrow 0$ as $n \rightarrow \infty$. Let $\mathcal{X}_n$ be the set of blocklength in FL-VL codes. We define the n -th order FL-VL minimax redundancy as

$$\mathcal{R}_F(n) = \inf\{\hat{r}_n(\ell_n) : \ell_n \in \mathcal{X}_n\}. \quad (118)$$

We now define similar quantities for VL-FL codes. A VL-FL code maps variable-length strings of source outputs into fixed-length binary codewords. The performance of a VL-FL code is determined by a set Γ which consists of the variable-length strings of source outputs which are encoded. The blocklength of a VL-FL code is the length of the codewords and is denoted by n . So $n = \lceil \log |\Gamma| \rceil$ where $|\Gamma|$ is the cardinality of the set Γ and $\lceil a \rceil$ represents

the smallest integer not less than a . Since Γ completely specifies the code we refer to Γ as the code. Let $l(\underline{x})$ be the number of letters in the string $\underline{x}$. The rate of a VL-FL code Γ applied to a source θ is

$$R_n(\Gamma, \theta) = n[\bar{l}_\theta(\Gamma)]^{-1} \quad (119)$$

where

$$\bar{l}_\theta(\Gamma) \triangleq \sum_{\underline{x} \in \Gamma} p_\theta(\underline{x}) l(\underline{x}) \quad (120)$$

is the expected length of the input strings. We may define a lower bound on the rate of this code as

$$\mathcal{K}(\Gamma, \theta) = - \sum_{\underline{x} \in \Gamma} p_\theta(\underline{x}) \log p_\theta(\underline{x}) [\bar{l}_\theta(\Gamma)]^{-1} . \quad (121)$$

So $\mathcal{K}(\Gamma, \theta)$ is the entropy of the set of strings $\underline{x} \in \Gamma$ divided by the expected length of these strings. The redundancy of the code Γ is defined as

$$r_n(\Gamma, \theta) = R_n(\Gamma, \theta) - \mathcal{K}(\Gamma, \theta) \quad (122)$$

and the maximum redundancy is

$$\hat{r}_n(\Gamma) = \sup\{r_n(\Gamma, \theta) : \theta \in \Lambda\} . \quad (123)$$

If $\mathcal{K}_n^*$ is the set of all VL-FL codes of blocklength n then define

$$R_v(n) = \inf\{\hat{r}_n(\Gamma) : \Gamma \in \mathcal{K}_n^*\} . \quad (124)$$

For the definitions (113)-(124) it is assumed that each source $\theta \in \Lambda$ is stationary. A unifilar Markov source is stationary only if it is in its steady-state distribution. We do not wish to assume that the sources are in their steady-state distributions since we are interested in applying these

codes to sources with slowly varying probabilities p_θ . For this reason the codes which we construct are universal with respect to the initial state of the source.

Let θ be a unifilar Markov source with alphabet $A = \{1, 2, \dots, J\}$ and a set of states $\mathcal{S} = \{1, 2, \dots, S\}$. The properties of the source θ are given by an initial state s_0 and a pair of $J \times S$ matrices $P_\theta = \{p_\theta(x|s)\}$ and $F_\theta = \{f_\theta(x, s)\}$ where $p_\theta(x|s)$ is the probability that letter x is output when the source is in state s , and $f_\theta(x, s)$ is the state into which the source moves following this event. The probability of a string $\underline{x} = (x_1, \dots, x_k)$ which starts with the first output letter is

$$p_\theta(\underline{x}) = \prod_{i=1}^k p_\theta(x_i | s_{i-1}) \quad (125)$$

where

$$s_i = f_\theta(x_{i-1}, s_{i-1}) . \quad (126)$$

If $\underline{x} = (x_{m+1}, \dots, x_{m+k})$ then

$$p_\theta(\underline{x}) = \sum_{j=1}^S P_\theta^*(j|m, s_0) \prod_{i=1}^k p_\theta(x_{m+i} | s'_{i-1})$$

where $P_\theta^*(j|m, s_0)$ is the probability of being in state j after m steps if the initial state is s_0 , $s'_0 = j$, and $s'_i = f_\theta(x_{m+i-1}, s'_{i-1})$, $i = 1, \dots, k-1$.

We assume that Λ is the class of all unifilar sources with a given alphabet A , state space $\mathcal{S}$, and transition matrix F_θ . (So F_θ is the same for all $\theta \in \Lambda$.) A source $\theta \in \Lambda$ is then specified by an initial state s_0 and a transition probability matrix P_θ . The sources in Λ are not stationary but the quantities defined in (113)-(116) are valid if we assume that $\underline{x} = (x_1, \dots, x_n)$ is the first block of n source outputs so that $p_\theta(\underline{x})$ is given by (125) and (126).

For VL-FL codes there are other difficulties. First $K(\Gamma, \theta)$ depends on the code Γ , so the code with the smallest redundancy $r_n(\Gamma, \theta)$ does not necessarily have the lowest rate $R_n(\Gamma, \theta)$. For memoryless sources $K(\Gamma, \theta)$ is the entropy of the source, so it is independent of Γ . This is not the case for unifilar Markov sources. In fact, even if a source is in its steady state distribution before the first string of source outputs is encoded, it need not be afterwards. The VL-FL code induces a distribution on the states. However, we may show that the lower bound of (121) is independent of Γ in the following sense.

Let $\{\theta_i, i=1,2,\dots,S\}$ be a set of sources in Λ with transition probabilities $p_{\theta_i} = p_\theta$ such that θ_i has initial state i . Suppose that some set of S codes with encoding sets Γ_i achieves $K(\Gamma_i, \theta_i)$ $i=1,\dots,S$. Then from the Kraft inequality and the fact that for any $\varphi \in \Lambda$

$$-\sum p_{\theta_i}(\underline{x}) \log p_\varphi(\underline{x}) \geq -\sum p_{\theta_i}(\underline{x}) \log p_{\theta_i}(\underline{x})$$

with equality if and only if $p_{\theta_i}(\underline{x}) = p_\varphi(\underline{x})$, the length of the codeword for $\underline{x} \in \Gamma_i$ must be $-\log p_{\theta_i}(\underline{x})$. (Note that this set of codes is VL-VL.) Now if we wish to determine the total length of the codewords used to encode a block $\underline{z}$ of m consecutive outputs with this set of codes the problem is that the end of the block $\underline{z}$ may be in the middle of an encoded string. However, due to the structure of the codes this problem may be resolved by dividing codewords. Suppose that one encoded string $\underline{x}$ has k letters within $\underline{z}$ and $l(\underline{x})-k$ outside $\underline{z}$. The length of the codeword for $\underline{x}$ is $-\log p_{\theta_i}(\underline{x})$ and since

$$p_{\theta_1}(\underline{x}) = \prod_{j=1}^{\ell(\underline{x})} p_{\theta}(x_j | s_{j-1}) \quad (127)$$

$$= \prod_{j=1}^k p_{\theta}(x_j | s_{j-1}) \prod_{j=k+1}^{\ell(\underline{x})} p_{\theta}(x_j | s_{j-1}), \quad (128)$$

the part of the codeword due to letters within $\underline{z}$ is $-\log \prod_{j=1}^k p_{\theta}(x_j | s_{j-1})$, independently of the following symbols. So the (not necessarily integer) number of bits used to encode $\underline{z}$ with initial state s_0 is

$$-\sum_{i=0}^j \log p_{\theta}(\underline{x}^{(i)} | s_i) = -\log p_{\theta}(\underline{z}) \quad (129)$$

where $\underline{z}$ is encoded as $\underline{x}^{(1)}, \underline{x}^{(2)}, \dots, \underline{x}^{(j)}$ ($\underline{x}^{(j)}$ is not necessarily an entire encoded string), and $p_{\theta}(\underline{x} | s)$ is the probability of $\underline{x}$ for source θ_s . The expected rate of this set of codes is

$$-m^{-1} \sum_{\underline{z} \in A^m} p_{\theta}(\underline{z} | s_0) \log p_{\theta}(\underline{z} | s_0) = H_m(\theta_j) \quad (130)$$

where $\theta_j = (p_{\theta}, s_0)$. So a set of codes which achieves $\mathcal{K}(\Gamma_i, \theta_i)$ for all initial states achieves the m -th order entropy given any initial state.

If we have a VL-FL code Γ^* such that

$$R_n(\Gamma^*, \theta_i) \leq \mathcal{K}(\Gamma^*, \theta_i) + \epsilon \quad (131)$$

where ϵ is independent of i , then the expected rate of this code over m outputs is

$$R_m(\theta_j) \leq H_m(\theta_j) + \epsilon \quad (132)$$

from (130) since ϵ is simply an extra per letter redundancy. So given the bound of (131) which depends on the set Γ^* we may derive a performance bound (132) which is independent of Γ^* .

If a sequence of codes Γ_n is minimax universal then from (132) and the fact that Λ contains sources with all possible P_θ and initial states s_0 , the rate of these codes approaches the m-th order entropy if we average as above. So if a VL-FL code Γ has

$$\hat{r}_n(\Gamma) = \epsilon \quad (133)$$

and a FL-VL code l_n has

$$\hat{r}_n(l_n) = \epsilon \quad (134)$$

then the two codes have approximately the same rate when averaged over a block of source outputs. A FL-VL code and a VL-FL code with the same number of codewords (due to the lack of structure in the codes the number of codewords is a good measure of complexity) have blocklengths n and $n \log J$ respectively. So if we wish to compare codes of the same complexity, then we should compare the performance of a blocklength n FL-VL code to that of a blocklength $n \log J$ VL-FL code.

In [14] a delay parameter d^* is defined by $d^* = n$ for FL-VL codes and by

$$d^* = \inf\{\bar{l}_\theta(\Gamma) : \theta \in \Lambda\} \quad (135)$$

for a VL-FL code Γ . The minimax redundancy, denoted $\tilde{R}_F(d)$ and $\tilde{R}_V(d)$ for FL-VL and VL-FL codes respectively, is defined as the minimum of $\hat{r}_n(l_n)$ or $\hat{r}_n(\Gamma)$ over all codes whose delay d^* does not exceed d . This may seem somewhat unnatural, but the number of codewords is approximately the same

for all codes with delay d^* , so this approach leads to the same comparison as that mentioned above. We show this as follows. First since $d^* = n$ for FL-VL codes we have

$$R_F(n) = \tilde{R}_F(n) . \quad (136)$$

However, $R_V(n)$ and $\tilde{R}_V(n(\log J)^{-1})$ are not quite the same. Any blocklength n VL-FL code Γ_n satisfies

$$d^* \leq n(\log J)^{-1} \quad (137)$$

if Λ includes the source θ^* which has all letters equiprobable in all states. This is because the entropy of θ^* is $\log J$. So we have

$$R_V(n) \geq \tilde{R}_V(n[\log J]^{-1}) . \quad (138)$$

Further, if Γ_n^* achieves a minimax redundancy $R_V(n)$ then

$$\bar{L}_\theta(\Gamma_n) \geq n[K(\Gamma_n, \theta) + R_V(n)]^{-1} \quad (139)$$

so

$$d^* \geq n[\log J + R_V(n)]^{-1} .$$

This implies

$$R_V(n) \leq \tilde{R}_V(n[\log J + R_V(n)]^{-1}) . \quad (140)$$

Since $R_V(n)$ is $O(n^{-1} \log n)$, we have

$$R_V(n) \approx \tilde{R}_V(n[\log J]^{-1}) . \quad (141)$$

So any bound on R_V may be used to derive a bound on $\tilde{R}_V$.

There are a number of papers with results on the minimax redundancy of FL-VL codes. In [15] and [16] asymptotic upper and lower bounds are derived for unifilar Markov sources which show

$$R_F(n) = \frac{1}{2} n^{-1} (J-1) S \log n + O(n^{-1}) . \quad (142)$$

These results are only asymptotic, however, as the $O(n^{-1})$ term is not evaluated. An upper bound

$$R_F(n) \leq \frac{1}{2} n^{-1} (J-1) S \log n + K n^{-1} \quad (143)$$

is derived in [12] and the constant K is given explicitly. For memoryless sources a lower bound is derived in [5] which shows that

$$R_F(n) \geq \frac{1}{2} n^{-1} (J-1) \log n - K' n^{-1} . \quad (144)$$

Again the constant K' is determined.

There are fewer results for VL-FL codes. Lawrence [17] derives a universal VL-FL coding technique for binary memoryless sources which has

$$\hat{r}_n(\Gamma_n) \leq n^{-1} \log n + K'' n^{-1} . \quad (145)$$

(This bound, however, does not appear in the paper.) In [14] results of Khodak are mentioned which state that

$$R_V(n) = O(n^{-1} \log n) \quad (146)$$

for memoryless sources. In the next section we show that

$$R_V(n) \leq (\log J) n^{-1} [\frac{1}{2} S (J-1) + 1] \log n + \hat{K} n^{-1} \quad (147)$$

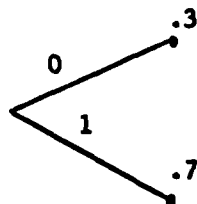
for unifilar Markov sources.

2.2. Universal VL-FL Code Construction

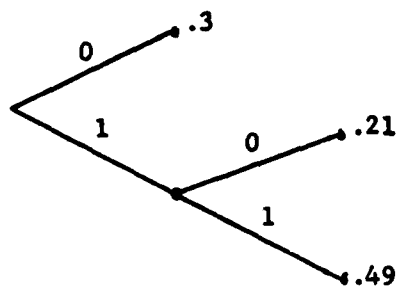
First we introduce the optimal VL-FL coding procedure (Tunstall's algorithm [18]) for memoryless sources. Let θ be a discrete memoryless source with alphabet $A = \{1, 2, \dots, J\}$ and let $p_\theta(x) = P\{X=x\}$, $x \in A$ be the probability that the letter x is output. A VL-FL code maps a variable number of source outputs into a fixed number of code symbols from an alphabet C . We will assume that $C = \{0, 1\}$, i.e., that the code is binary. Tunstall's algorithm generates a rooted tree whose terminal nodes (leaves) correspond to codewords. There are J branches leaving each non-terminal node, and these branches are labelled with the J source symbols. The encoding procedure consists of starting at the root node and traversing after each source output the branch with the corresponding label. When a leaf is reached, the codeword assigned to that leaf is sent and the procedure is repeated. So each leaf corresponds to a unique string $\underline{x} = (x_1, \dots, x_k)$ of source outputs and has probability $p_\theta(\underline{x}) = \prod_{i=1}^k p_\theta(x_i)$. The algorithm generates a larger optimal tree from a smaller one by adding J branches to the tree at the leaf with the highest probability. So the highest probability leaf is divided into a set of J leaves. It is easily seen that the ratio of the lowest probability leaf in the tree to be highest is not less than $\alpha \Delta \min\{p_\theta(x) : x \in \mathcal{X}\}$.

In Figure 6 this procedure is illustrated for a binary memoryless source with $p(1) = .7$ and $n = 2$ (4 leaves). The encoding tree is formed in three steps with the most probable leaf being extended at each step. Each of the final set of input strings Γ is assigned a codeword of length 2.

a. 2 leaves

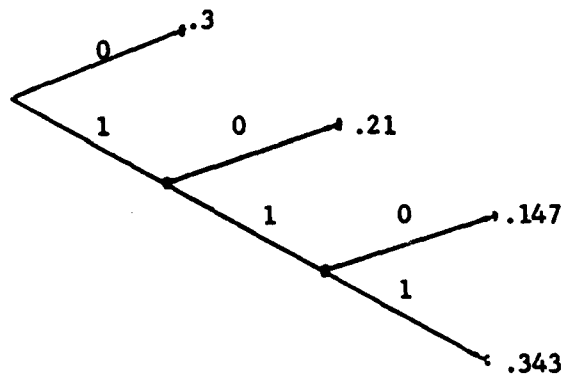


b. 3 leaves



c. 4 leaves

$\Gamma = \{0, 10, 110, 111\}$



$\underline{x}$	$l(\underline{x})$	Codeword for $\underline{x}$
0	1	00
10	2	01
110	3	10
111	3	11

Figure 6. Construction of a Tunstall code with blocklength 2 for a binary memoryless source with $p(1) = .7$.

We now extend this algorithm to coding for unifilar Markov sources. A VL-FL code for a unifilar Markov source θ is generated using an algorithm much like that for memoryless sources except that now each node of the tree has a state associated with it. The probabilities associated with the branches are given by $p_\theta(x|s)$ where x is the output letter which labels the branch and s is the state of the node which the branch is leaving. The algorithm again consists of extending the most probable node, where this probability is now given by the product of the transition probabilities of the branches traversed in reaching the node. It is not clear that this algorithm is optimal since in general to actually encode some block of source outputs S encoding trees are necessary, each designed for P_θ and F_θ but for different initial states. The structure of each of these S trees determines the probability of being in a particular state after encoding a string of source outputs, hence the probability that a particular tree is used is affected by the structure of all S trees. It is not necessarily true that generating these trees independently (as is done here) is the optimal encoding algorithm. However, the algorithm does yield code trees which have asymptotically good performance as will be seen later. Further, in each tree the ratio of the minimum probability leaf to the maximum is not less than [19]

$$\alpha = \min[p_\theta(x|s); x \in A, s \in \mathcal{S}] \quad (148)$$

We use the Tunstall algorithm for individual unifilar Markov sources to construct a universal code for a class Λ of sources as follows. Let $\Phi_m = \{\varphi_i; i = 1, \dots, \gamma_m\}$ be a finite subset of Λ such that if $\varphi \in \Phi_m$, then there is a source $\varphi_j \in \Phi_n$ with initial state j which has the same transition probability matrix as φ , for $j = 1, 2, \dots, S$. Let $\Gamma_m^{(1)}$ be the encoding set of a blocklength

m code designed for the i -th source $\varphi_i \in \mathcal{S}_m$. The codes Γ_m^* are constructed using Tunstall's algorithm as above. The universal code Γ_n^* is defined as follows (n is defined in terms of m in (150) below). A string $\underline{x}$ is an element of Γ_n^* if $\underline{x} \in \Gamma_m^{(i)}$ for some i and $(\underline{x} * \underline{y}) \notin \Gamma_m^{(j)}$ for any $j = 1, 2, \dots, \gamma_m$, where $\underline{y}$ is a non-empty string of source letters ($*$ represents concatenation). So the tree for Γ_n^* contains all nodes from the trees for $\Gamma_m^{(i)}$, $i = 1, 2, \dots, \gamma_m$. Now

$$|\Gamma_n^*| \leq \sum_{i=1}^{\gamma_m} |\Gamma_m^{(i)}| = \gamma_m \cdot 2^m \quad (149)$$

so the strings in Γ_n^* may be encoded into codewords of length

$$n \leq m + \lceil \log \gamma_m \rceil \quad (150)$$

The rate of this code Γ_n^* when applied to a source θ is

$$R_n(\Gamma_n^*, \theta) = n[\bar{\ell}_\theta(\Gamma_n^*)]^{-1} \quad (151)$$

$$\leq [m + \lceil \log \gamma_m \rceil][\bar{\ell}_\theta(\Gamma_m^{(k)})]^{-1} \quad (152)$$

for any $k = 1, 2, \dots, \gamma_m$. This follows because by its construction the expected length of the sequences in Γ_n^* must be at least that of any of the sets $\Gamma_m^{(i)}$. There are two sources of redundancy in (152) which we must bound in order to bound the redundancy of the code Γ_n^* . The first is the $\lceil \log \gamma_m \rceil$ term which is due to the fact that $|\Gamma_n^*|$ may be as large as $\gamma_m 2^m$. The second factor is the difference between

$$\max\{\bar{\ell}_\theta(\Gamma_m^{(i)}): i = 1, \dots, \gamma_m\} \quad (153)$$

and $\bar{\ell}_\theta(\tilde{\Gamma}_m)$, where $\tilde{\Gamma}_m$ is the Tunstall code designed for source θ . So the second factor is derived from the mismatch between the actual source θ and

the source for which the code $\Gamma_m^{(1)}$ was designed (that source being some $\varphi_1 \in \mathcal{S}_m$). As γ_m increases the effect of the first factor increases and that of the second decreases. Blumer [13] shows that for

$$\log \gamma_m = \frac{1}{2} S(J-1) \log m + K \quad (154)$$

where K does not depend on m , a set $\mathcal{S}_m$ may be constructed such that

$$\max\{\min\{J_C(\theta; \varphi_1) : 1 = 1, 2, \dots, \gamma_m\} : \theta \in \Lambda\} \leq m^{-1} \quad (155)$$

(Here $J_C(\theta; \varphi)$ is the entropy of source θ relative to source φ .) We use this result to bound the effect of the mismatch. The details of the derivation are in Appendix B. The final result is

$$\hat{r}_n(\Gamma_n^*) < n^{-1} \log J[\log n + \frac{1}{2} S(J-1) \log n] + K_1 n^{-1} \quad (156)$$

for $n > K_2 (\log n)^2$ where K_1 and K_2 are constants independent of n and θ given in Appendix B, equations (B.34) and (B.35). So the code is minimax universal.

As previously discussed, we wish to compare the performance of a blocklength n VL-FL code to that of a blocklength $n[\log J]^{-1}$ FL-VL code (these codes have the same number of codewords). For FL-VL codes (143) gives

$$R_F(n[\log J]^{-1}) \leq n^{-1} \log J[\frac{1}{2} S(J-1) \log n] + \tilde{K} n^{-1} \quad (157)$$

and (156) implies

$$R_V(n) \leq n^{-1} \log J[\log n + \frac{1}{2} S(J-1) \log n] + K_1 n^{-1} \quad (158)$$

so we see that the leading term in the redundancy bounds is the same except for a $\log n$ term which appears in the VL-FL bound. This additional term is present because there is no known bound on the redundancy of a Tunstall code which remains finite as the minimum letter probability of the source approaches zero. If the sources have all letter probabilities greater than some $\epsilon > 0$, then the $\log n$ term is replaced by $\log \epsilon^{-1}$.

Further if Λ is the class of binary memoryless sources (so $S = 1$ and $A = \{0,1\}$), then the $\log n$ term may be eliminated. The final result for this case is

$$R_V(n) \leq \frac{1}{2} n^{-1} \log n + K_3 n^{-1}. \quad (159)$$

The derivation of this result appears in Appendix B.

2.3. Performance Evaluation for Binary Memoryless Sources

In this section we construct VL-FL codes for the class of binary memoryless sources using the method presented in Section 2.2, and compare their performance to the performance of the FL-VL codes constructed in [11]. One modification to the basic code construction is given, and the performance of codes obtained from this modification is evaluated. Here $J = 2$ so $n \log J = n$ and we must compare VL-FL codes to FL-VL codes of the same blocklengths.

One difficulty which arises in designing a VL-FL code of blocklength n is that we do not know apriori the cardinality of Γ_n^* for a given m . We only have the upper bound of (149). So to actually construct a blocklength n VL-FL code we use the following iterative procedure. We choose an initial number N of codewords for the Tunstall codes $\Gamma_m^{(1)}$ which are designed for

sources in $\mathcal{S}_m$ (here $m \triangleq \log N$ is not necessarily an integer). We then combine these codes into a single code $\Gamma^*(N)$. We iterate this procedure to find

$$\bar{N} \triangleq \max \{N: |\Gamma^*(N)| \leq 2^n\} . \quad (160)$$

So $\bar{N}$ is the maximum number of codewords in the Tunstall codes $\Gamma_m^{(1)}$ such that the combined code has blocklength n . Then we set

$$\Gamma_n^* = \Gamma^*(\bar{N}) . \quad (161)$$

If we let the parameter θ for a binary memoryless source be the probability of a one, then the class of binary memoryless sources is the interval $[0,1]$. Because Λ is one-dimensional we may easily determine the optimum design point set $\mathcal{S}_m$ for any γ_m . These sets are given for some values of γ_m in Table 1 of [11] and may be determined for other γ_m using the technique described there. Codes of blocklengths 5, 8, and 10 were constructed using these sets $\mathcal{S}_m$. A graph of the redundancy of these codes is given in Figure 7. The curves are symmetric about $\theta = .5$. In Table 1 the maximum redundancies are compared to those of the FL-VL codes of [11]. We see that VL-FL codes have significantly better performance for blocklengths 8 and 10, and only slightly worse for blocklength 5. In Figure 8 we have graphed the redundancy of blocklength 8 and 10 VL-FL codes together with FL-VL codes of the same blocklengths. The VL-FL codes have lower redundancy for almost all values of θ . The largest difference occurs at $\theta = 0$ or 1. The reason for this is that in any universal FL-VL code the codewords for the all zeros and all ones output blocks must have length at least two, hence the redundancy at $\theta = 0$ or 1 must be at least $2n^{-1}$.

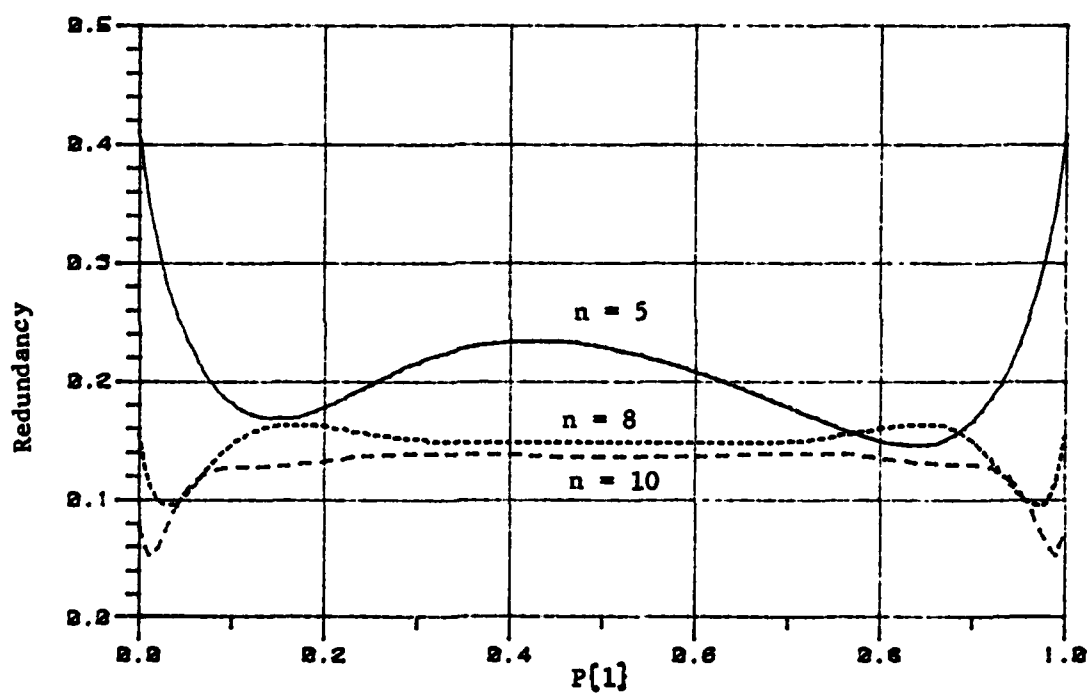


Figure 7. VL-FL universal code performance for blocklengths $n = 5, 8$, and 10 over the class of binary memoryless sources.

n	VL-FL	FL-VL
5	.409	.400
8	.164	.250
10	.139	.200

Table 1. Maximum redundancies for VL-FL and FL-VL codes of blocklengths $n=5, 8$, and 10 .

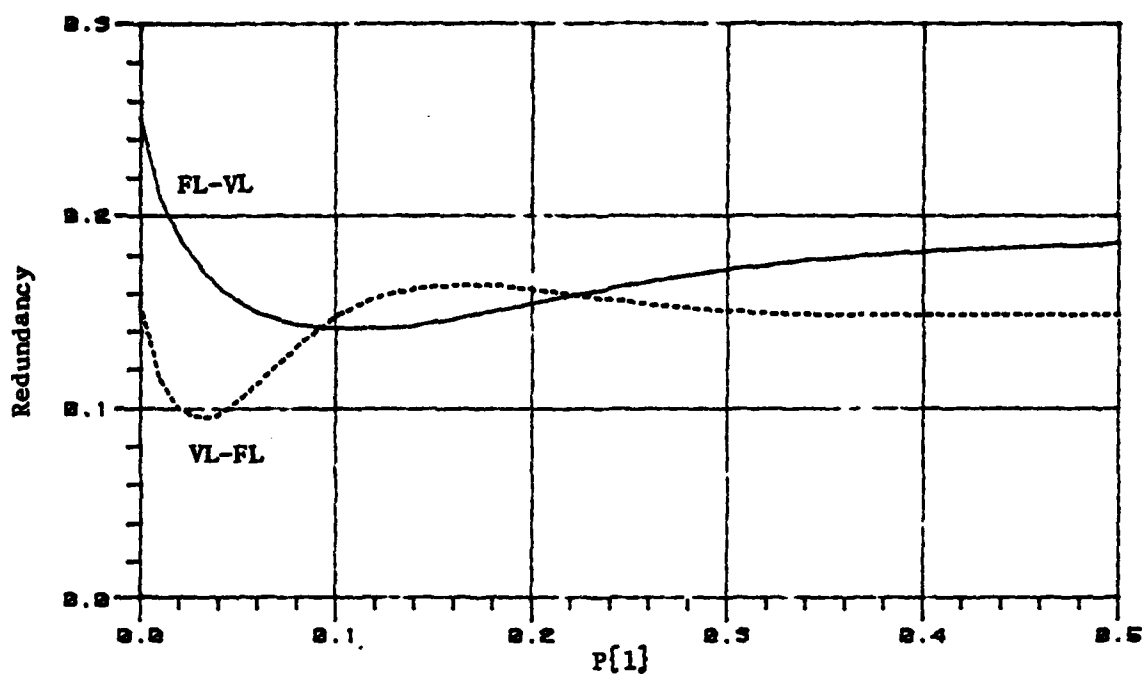
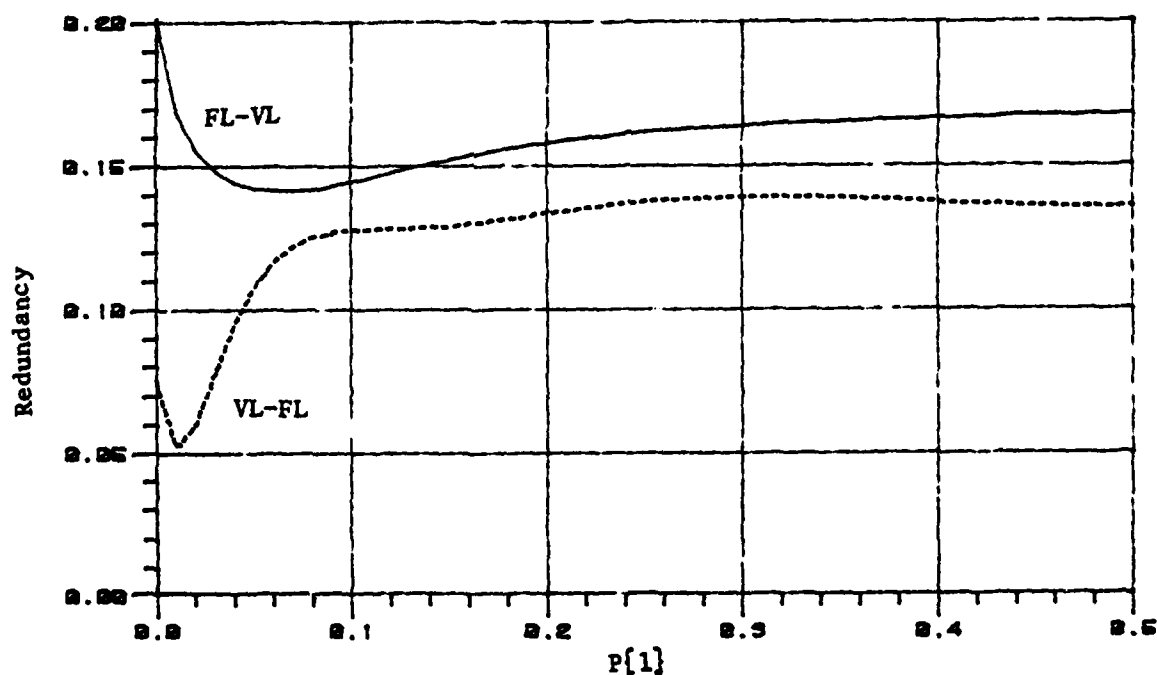
a) blocklength $n = 8$ b) blocklength $n = 10$

Figure 8. Comparison of the performance of VL-FL and FL-VL universal codes for binary memoryless sources. (The performance is symmetric about $P[1] = 0.5$.)

The lack of structure in these codes typically requires a table lookup scheme for decoding, so their complexity increases as 2^n , the total number of codewords. This limits n , and thus the achievable redundancy is also limited. To alleviate the problem of complexity we may adopt the following modified procedure. We design Tunstall codes $\Gamma_m^{(1)}$ of blocklength m for the sources $\varphi_j \in \Phi_m$. Instead of combining these codes we leave them as separate "subcodes". Then we encode the source outputs $(x_0, x_1, \dots)$ as follows. Each subcode encodes k strings from the source output. We use the subcode which has the lowest rate for this set of k strings, i.e., the one which encodes the largest number of source outputs. The codeword for this set of k strings is the concatenation of a prefix of length $\lceil \log \gamma_m \rceil$ which identifies the subcode we are using with the k codewords for the encoded strings from that subcode. The total number of codewords for this procedure is $\gamma_m 2^m$. The resultant blocklength is approximately km and would require about $\gamma_m 2^{km}$ codewords in the original coding procedure. So the complexity of this blocklength km code is approximately that of the blocklength m code previously considered. The reason that this new code performs better than a blocklength m code is that the redundancy due to combining the codes $\Gamma_m^{(1)}$ is of order $m^{-1} \log m$, whereas the other terms in the redundancy are of order m^{-1} . With the new procedure these terms are $(km)^{-1} \log m$ and m^{-1} respectively so that the dominant term is reduced with respect to the other terms.

A similar technique is used in [11] for longer blocklengths. A special code is used there for source with θ near 0 or 1, but the complexity remains about the same. In Table 2 we give the maximum redundancies of VL-FL codes of blocklengths 50, 80, and 100 which are constructed by encoding 10

n	VL-FL	FL-VL
50	.117	.080
80	.065	.050
100	.053	.050

Table 2. Maximum redundancies for VL-FL and FL-VL codes of blocklengths $n=50, 80$, and 100 .

strings with VL-FL subcodes at blocklengths 5, 8, and 10. Results from [11] for the same blocklengths are included for comparison. The FL-VL codes perform a little better, but there is no great difference.

CHAPTER 3

UNIVERSAL CODING FOR REAL-VALUED SOURCES

3.1 Introduction

Here we consider source coding for discrete-time real-valued sources. The source output for the i -th time interval is a real random variable X_i . In contrast with the previous chapter, the entropy of these sources is generally infinite, so noiseless source coding is not possible. The problem here is one in rate-distortion theory, so the goal is to find a code with low distortion for a given rate. We assume that we have a distortion measure $d_n(\underline{x}, \underline{y})$ for each positive integer n , where $\underline{x}$ and $\underline{y}$ are elements of $\mathbb{R}^n$, and that there is a maximum distortion $\bar{D} < \infty$ such that

$$n^{-1} d_n(\underline{x}, \underline{y}) \leq \bar{D}, \quad \underline{x}, \underline{y} \in \mathbb{R}^n \quad (162)$$

for all n . There are a number of papers on coding techniques for known sources of this type, e.g., [25], [28], and [29]. For some specific classes of sources we show how a code for an entire class may be constructed using a coding technique for single sources in Λ . We show that asymptotically this code performs as well on any source $\theta \in \Lambda$ as a code designed specifically for that source.

The codes which we consider are fixed rate; that is, all codewords have the same length. The codes consist of vector quantization followed by a mapping of the quantizer outputs into fixed length binary sequences. A block-length n M -point vector quantizer is a mapping $f_n: \mathbb{R}^n \rightarrow A$ where $A = \{a_i: i=1, \dots, M\}$ is a finite set with elements in $\mathbb{R}^n$. The elements of A are called output levels. The distortion which results when the outputs of a given source θ are quantized is defined as

$$D(f_n; \theta) = n^{-1} E_{\theta} [d_n(\underline{X}, f_n(\underline{X}))] \quad (163)$$

where $d_n(\cdot, \cdot)$ is the distortion measure and $\underline{X} = (X_1, \dots, X_n)$. The rate of quantizer f_n for source θ is defined as

$$R(f_n; \theta) = \frac{1}{n} \lceil \log M \rceil. \quad (164)$$

For our purposes a code is determined by its associated quantizer f_n , so we refer to the code as f_n . Then the rate and distortion of the code f_n when applied to source θ are defined by (164) and (163) respectively.

We assume that we have a coding technique for sources in a class Λ . So for any source $\theta \in \Lambda$ we may construct a blocklength n code f_n^{θ} with M output levels. For a given blocklength n and rate R let $\delta_{n,R}(\theta)$ be the distortion achieved by f_n^{θ} . We assume that

$$\delta_{n,R}(\theta) = D(f_n^{\theta}; \theta) \leq D(f_n^{\varphi}; \theta) \quad (165)$$

for $\theta, \varphi \in \Lambda$. So when applied to source θ , f_n^{θ} performs at least as well as a code designed for some other source in Λ . The coding technique here does not necessarily yield optimal codes; that is, $\delta_{n,R}(\theta)$ need not approach the distortion-rate function $D(R)$ for source θ as $n \rightarrow \infty$. For example, these codes may be derived from locally optimal quantizers (designed using the algorithm of [25]) or from optimal one-dimensional quantizers [29].

For some specific classes Λ we show that given a coding technique we may construct a sequence of codes of increasing blocklength $f_1^*, f_2^*, \dots$ such that

$$D(f_n^*; \theta) - \delta_{n,R}(\theta) \rightarrow 0$$

and

$$R(f_n^*; \theta) \rightarrow R \quad (166)$$

uniformly on Λ as $n \rightarrow \infty$. We call such a sequence of codes minimax universal with respect to the coding technique which yields $\delta_{n,R}(\theta)$. It is important to note that $\delta_{n,R}(\theta)$ is the distortion achieved by many different codes, each designed for a particular $\theta \in \Lambda$. In contrast to this, $D(f_n^*; \cdot)$ is the distortion of a single code over the class Λ .

First we consider classes of memoryless sources. A general result is derived for classes which are twice-differentiable with respect to their parameters θ . This result is in terms of an integral which is evaluated for some specific classes. For all of these classes the result is that

$$D(f_n^*; \theta) - \delta_{n,R}(\theta) \leq K_1 n^{-1}$$

and

$$R(f_n^*; \theta) - R \leq k n^{-1} \log n + K_2 n^{-1} \quad (167)$$

where k is the dimension of Λ and K_1 and K_2 are constants. We then show that a result of the same form holds for k -th order Gaussian autoregressive sources. These codes give upper bounds on the additional rate and distortion incurred when coding for a class Λ rather than a specific source $\theta \in \Lambda$.

An outline of the code construction and bounding of performance is as follows. For each integer n we have a finite subset Φ_n of sources in Λ . Codes for each source in Φ_n are constructed. These codes are then combined into a single code by adding a prefix to each codeword which identifies the source in Φ_n for which the code is designed. The rate of the resultant code is greater than the rate of the individual codes because of this prefix. The code has low distortion for the sources in Φ_n but may not for sources which are not in Φ_n . As the number of sources in Φ_n increases, the additional

rate increases and the distortion decreases. The first result bounds the mismatch distortion, i.e., the distortion which results when a code designed for one source is applied to another, in terms of the entropy of one source relative to the other. Next we show that if the relative entropy may be bounded then we may pick Φ_n to give a minimax universal code. The relative entropy is then bounded for some classes of memoryless sources and finally for Gaussian autoregressive sources.

3.2 Code Construction

We design codes $f_n^{(i)}$ of blocklength n and rate R for sources $\varphi_i \in \Phi_n$, where $\Phi_n = \{\varphi_i: i=1, \dots, \gamma_n\}$ is a subset of Λ . These codes take n -tuples of source outputs into codewords of length $\lceil \log M \rceil$, where M is the number of levels in the associated vector quantizer. These γ_n codes are then combined by adding a $\lceil \log \gamma_n \rceil$ -bit prefix to each codeword. We denote this combined code f_n^* . We know

$$D(f_n^{(i)}; \varphi_i) = \delta_{n,R}(\varphi_i) \quad (168)$$

where $R = n^{-1} \lceil \log M \rceil$. The encoding procedure for f_n^* is as follows. Set

$$f_n^*(\underline{x}) = f_n^{(i)}(\underline{x}) \quad (169)$$

for $i = \arg \min_i d_n(\underline{x}, f_n^{(i)}(\underline{x}))$. Then the codeword for $f_n^*(\underline{x})$ is the codeword for $f_n^{(i)}(\underline{x})$ with a prefix attached. So we have

$$D(f_n^*; \theta) \leq \min_i D(f_n^{(i)}; \theta); \quad (170)$$

that is, the distortion of f_n^* for a source θ is no greater than the distortion for any one of the codes $f_n^{(i)}$ from which it was constructed. The rate of f_n^* is

$$R(f_n^*; \theta) = n^{-1} \{ \lceil \log M \rceil + \lceil \log \gamma_n \rceil \}. \quad (171)$$

Now if $j = \arg \min_i D(f_n^{(i)}; \theta)$ we have

$$D(f_n^*; \theta) - \delta_{n,R}(\theta) \leq D(f_n^{(j)}; \theta) - D(f_n^{(j)}; \varphi_j) + D(f_n^{(j)}; \varphi_j) - \delta_{n,R}(\theta). \quad (172)$$

Let $\hat{f}_n$ be the code designed for θ . Now

$$D(f_n^{(j)}; \varphi_j) = \delta_{n,R}(\varphi_j) \leq D(\hat{f}_n; \varphi_j) \quad (173)$$

so we have

$$D(f_n^*; \theta) - \delta_{n,R}(\theta) \leq [D(f_n^{(j)}; \theta) - D(f_n^{(j)}; \varphi_j)] + [D(\hat{f}_n; \varphi_j) - D(\hat{f}_n; \theta)]. \quad (174)$$

The set Φ_n is designed such that the right hand side of (174) may be bounded uniformly for $\theta \in \Lambda$. Both of these terms are distortion mismatch terms, that is, they represent the distortion incurred when a quantizer designed for one source is used for another. The following theorem bounds the distortion mismatch in terms of the relative entropy.

Theorem 3.1. If f_n is a code and

$$\mathcal{K}_n(\theta; \varphi) + \mathcal{K}_n(\varphi; \theta) \leq \xi \quad (175)$$

then

$$|D(f_n; \theta) - D(f_n; \varphi)| \leq \xi^{\frac{1}{2}} \bar{D}(2 \log e)^{-\frac{1}{2}} \quad (176)$$

where $\mathcal{K}_n(\theta; \varphi)$ is the n -th order entropy of θ relative to φ [30],

$$\mathcal{K}_n(\theta; \varphi) \triangleq \sup_i \left[\sum_1 \mu_\theta(B_i) \log \frac{\mu_\theta(B_i)}{\mu_\varphi(B_i)} \right], \quad (177)$$

where the supremum is over all finite partitions $\{B_i\}$ of $\mathbb{R}^n$ and $\mu_\theta(B)$ is the probability that the source output $\underline{x} \in \mathbb{R}^n$ is in B for source θ .

Proof. See Appendix C.

Now suppose that Λ is a compact subset of $\mathbb{R}^k$. The following corollary bounds the rate and distortion of a universal code using Theorem 3.1. For $\theta = (\theta_1, \dots, \theta_k)$ and $\psi = (\psi_1, \dots, \psi_k)$ we define the norm

$$\|\theta - \psi\| = \max\{|\theta_i - \psi_i| : 1 \leq i \leq k\}. \quad (178)$$

Corollary 3.1. If $\|\theta - \psi\| \leq n^{-1}$ implies $\mathcal{K}_n(\theta; \psi) \leq \tilde{K}n^{-2}$ for $\theta, \psi \in \Lambda$ then a code f_n^* may be constructed such that

$$|D(f_n^*; \theta) - \delta_{n, \theta}(R)| \leq n^{-1} \tilde{K}^{\frac{1}{2}} (\log e)^{-\frac{1}{2}} \quad (179)$$

and

$$R(f_n^*; \theta) - R \leq n^{-1} [k \log n + 1 + \sum_{i=1}^k \log(\ell_i + n^{-1})], \quad (180)$$

for all $\theta \in \Lambda$ where

$$\ell_i = \max_{\theta, \psi \in \Lambda} |\theta_i - \psi_i|; \quad i = 1, \dots, k \quad (181)$$

is the maximum difference in the i -th components of θ and ψ for any $\theta, \psi \in \Lambda$.

Proof. We cover Λ with cubes of size n^{-1} , and then let $\mathcal{S}_n$ consist of one source from each of these cubes. There are at most

$$\prod_{i=1}^k \lceil \ell_i n \rceil \quad (182)$$

such cubes which gives (180) and clearly for any $\theta \in \Lambda$ there exists a source $\varphi \in \mathcal{S}_n$ such that $\|\theta - \varphi\| \leq n^{-1}$ so (179) follows from Theorem 3.1.

Let Λ be a class of memoryless sources. Since Λ is a subset of $\mathbb{R}^k$ a source $\theta \in \Lambda$ is specified by k parameters $\{\theta_i, i=1, \dots, k\}$, $\theta_i \in \mathbb{R}$, and we write $\theta = \{\theta_1, \dots, \theta_k\}$. If θ has a density p_θ , then we assume that

$$p_\theta^i \triangleq \frac{\partial}{\partial \theta_i} p_\theta$$

and

$$p_\theta^{ij} \triangleq \frac{\partial}{\partial \theta_i} \frac{\partial}{\partial \theta_j} p_\theta \quad (183)$$

exist for all $i, j = 1, \dots, k$. For memoryless sources

$$p_\theta(\underline{x}) = \prod_{i=1}^n p_\theta(x_i)$$

so

$$\mathcal{H}_n(\theta; \varphi) = \mathcal{H}_1(\theta; \varphi) \quad (184)$$

For such sources we have the following theorem.

Theorem 3.2. If $\|\theta - \psi\| < \epsilon$ then

$$\mathcal{H}_1(\theta; \psi) < k^2 K^* \epsilon^2 \quad (185)$$

where

$$K^* \triangleq (\log e) \sup_{\mathbf{R}} \left\{ \frac{p_\varphi^1(\mathbf{x}) p_\varphi^j(\mathbf{x})}{p_\varphi(\mathbf{x})} d\mathbf{x} : \varphi \in \Lambda, i, j = 1, \dots, k \right\} \quad (186)$$

for $\Lambda \subset \mathbb{R}^k$.

Proof. This follows directly from Taylor's formula. We know [30]

$$\mathcal{H}_n(\theta; \varphi) = \int_{\mathbf{R}} n p_\theta(\underline{x}) \log \frac{p_\theta(\underline{x})}{p_\varphi(\underline{x})} d\underline{x}$$

so

$$K^* = \frac{\alpha}{\beta_1^2} . \quad (192)$$

Next consider a class of mixture distributions. Let $\{q^i: i=1, \dots, k+1\}$ be a set of distributions on $\mathbb{R}$. Then define $\Lambda \subset \mathbb{R}^k$ by

$$\Lambda = \{\theta: p(x) = \sum_{i=1}^{k+1} \theta_i q^i(x), \theta_j \geq \varepsilon, j=1, \dots, k\} . \quad (193)$$

Since

$$\frac{\partial}{\partial \theta_1} p_\theta(x) = q^1(x)$$

and

$$\frac{q^1(x)}{p_\theta(x)} < \frac{1}{\theta_1} < \frac{1}{\varepsilon}$$

we have

$$K^* < \frac{1}{\varepsilon} . \quad (194)$$

Finally, consider the case where Λ is a compact set of k -th order Gaussian autoregressive sources [30]. We assume here that $d_n(\cdot, \cdot)$ is the minimum of $n\bar{D}$ and the r -th power of the Euclidean distance. We also assume that if $f_n^*(x) = \alpha_i$ then

$$d_n(x, \alpha_i) < d_n(x, \alpha_j) ; j = 1, 2, \dots, M. \quad (195)$$

This means that each source output is mapped to the closest output level by the quantizer, which is a necessary condition for a quantizer to be optimal. In a Gaussian autoregressive source the output is generated by adding a Gaussian r.v. to a weighted sum of the previous outputs. So the j -th output is given by

$$\begin{aligned} \mathcal{K}_1(\theta; \psi) &\leq \mathcal{K}_1(\theta; \theta) + \sum_{i=1}^h (\theta_i - \psi_i) \left. \frac{\delta}{\delta \psi_i} \mathcal{K}_1(\theta; \psi) \right|_{\psi=\theta} \\ &+ \sum_{i=1}^k \sum_{j=1}^k (\theta_i - \psi_i)(\theta_j - \psi_j) \sup_{\omega'=\omega} \left\{ \frac{\delta}{\delta \omega_i} \frac{\delta}{\delta \omega_j} \mathcal{K}_1(\omega; \omega') \right\} : \omega \in \Lambda. \end{aligned} \quad (187)$$

Now $\mathcal{K}_1(\theta; \theta) = 0$ and

$$\left. \frac{\delta}{\delta \theta_i} \mathcal{K}_1(\theta; \psi) \right|_{\psi=\theta} = \int_{\mathbf{R}} p_{\theta}^i(x) dx = 0 \quad (188)$$

So the theorem follows from

$$\left. \frac{\delta}{\delta \omega_i} \frac{\delta}{\delta \omega_j} \mathcal{K}_1(\omega; \omega') \right|_{\omega'=\omega} = (\log e) \int_{\mathbf{R}} \frac{p_{\omega}^i(x) p_{\omega}^j(x)}{p_{\omega}(x)} dx. \quad (189)$$

If K^* is finite for a class Λ then the hypothesis of Corollary 3.1 is true and the code is minimax universal. The performance of this code is given by (179) and (180) with $\tilde{K} = k^2 K^*$. If Λ is the class of Gaussian distributions with mean $\mu \in [\mu_1, \mu_2]$ and variance $\sigma^2 \in [\sigma_1^2, \sigma_2^2]$, $\sigma_1 > 0$, then k , the dimension of Λ , is 2, and computation of the integral gives

$$K^* = \max \{1, 2\sigma_1^{-2}\} \quad (190)$$

For the case where Λ is the class of exponential distributions with mean $\beta \in [\beta_1, \beta_2]$, $\beta_1 > 0$, we have $k=1$ and

$$K^* = \frac{1}{\beta_1^2}. \quad (191)$$

More generally, if Λ is the class of gamma distributions with $\alpha > 0$ fixed and $\beta \in [\beta_1, \beta_2]$, then

$$X_j = -\sum_{i=1}^k a_i X_{j-i} + Z_j \quad (196)$$

where $Z_j \sim \mathcal{N}(0, \sigma^2)$ are independent. We assume that the sources are asymptotically stationary and that $\sigma^2 \in [\sigma_1^2, \sigma_2^2]$. This is guaranteed if the zeros of

$$a(\lambda) = 1 + a_1 \lambda^{-1} + \dots + a_k \lambda^{-k}$$

have magnitudes less than one. In vector form (196) becomes

$$\underline{X}_j = \Psi \underline{X}_{j-1} + \underline{Z}_j \quad (197)$$

where $\underline{X}_j = (X_j, X_{j-1}, \dots, X_{j-k+1})^T$, $\underline{Z}_j = (Z_j, 0, \dots, 0)^T$, and

$$\Psi = \begin{bmatrix} -a_1 & -a_2 & \dots & -a_k \\ 1 & & & \\ & 1 & \dots & \\ & & \ddots & 1 & 0 \end{bmatrix}. \quad (198)$$

So the roots of $a(\lambda)$ are the eigenvalues of Ψ . A source $\theta \in \Lambda$ is determined by $(a_1, \dots, a_k)$ and σ^2 so we write $\theta = (a_1, \dots, a_k, \sigma^2)$. Each $\theta \in \Lambda$ must have $a_k \neq 0$; otherwise, it is not a k -th order autoregressive source.

The design procedure and the derivation of performance bounds are a little different here because the initial state of the source $(X_{-1}, \dots, X_{-k})$ is an issue. As in previous chapters we want the universal code to perform well for all initial states. Here however the initial state may lie anywhere in $\mathbb{R}^k$ so this is not possible. We must assume that the initial state lies in a compact subset of $\mathbb{R}^k$. Specifically we assume that $|X_{-j}| \leq \zeta < \infty$ for $j = 1, \dots, k$. To construct the universal code we design codes for various sources in Λ but only for a single fixed initial state $\underline{x}^0 = \underline{0} \triangleq (0, 0, \dots, 0)$. We first consider how a code designed for source φ with $\underline{x}^0 = \underline{0}$ performs when used for source θ . If $\|\theta - \varphi\| \leq \epsilon$ we have

$$K_n(\theta; \varphi) \leq \frac{1}{2} \epsilon^2 \log e[\sigma_1^{-4} + k^2 \sigma_1^{-2} \sigma_x^2] \quad (199)$$

where

$$\sigma_x^2 \triangleq \sup_{\theta \in \Lambda} \lim_{n \rightarrow \infty} E_{\theta}[X_n^2]. \quad (200)$$

Now $\sigma_x^2 < \infty$ since we assume that Λ is compact and that the sources in Λ are stationary. Details in the derivation of (199) are given in Appendix C.

Given (199) we can bound the rate and distortion of the code as before.

So we have now constructed a universal code for initial state $\underline{x}^0 = \underline{0}$. We bound the performance of this code for other initial states as follows.

Given a vector of source outputs $\underline{X} = (X_0, \dots, X_{n-1})$ from a source θ with initial state $\underline{x}^0 = (x_{-k}, \dots, x_{-1})$ we define a vector $\tilde{\underline{x}}$ by

$$\tilde{X}_i = X_i - \mu_i \quad (201)$$

where

$$\mu_i \triangleq [Y^{i+1} \underline{x}^0]_i; \quad i = 0, \dots, n-1. \quad (202)$$

The matrix Ψ is as in (198) for $\theta = (a_1, \dots, a_k, \sigma^2)$, and $[\underline{x}]_j$ is the j -th component of the vector $\underline{x}$. Then $\tilde{\underline{x}}$ has the same distribution as a vector generated by source θ with $\underline{x}^0 = \underline{0}$. So we know

$$E_{\theta} [d_n(\tilde{\underline{x}}, f_n^*(\tilde{\underline{x}}))] \leq n \delta_{n,R}(\theta) + K' \quad (203)$$

where

$$K' \triangleq \frac{1}{2} \bar{D} [\sigma_1^{-4} + k^2 \sigma_1^{-2} \sigma_x^2]^{\frac{1}{2}}. \quad (204)$$

From (195) we have

$$d_n(\underline{X}, f_n^*(\underline{X})) \leq d_n(\underline{X}, f_n^*(\tilde{\underline{x}})). \quad (205)$$

Let $\underline{X}^* = f_n^*(\underline{X})$. Then the expected distortion Δ (unnormalized) is bounded by

$$\Delta \leq E_{\theta} \left\{ \left[\sum_{i=0}^{n-1} (X_i - X_i^*)^2 \right]^{r/2} \right\}. \quad (206)$$

Now from (205)

$$X_i - \tilde{X}_i^* \leq (X_i - \tilde{X}_i) + (\tilde{X}_i - [f_n^*(\tilde{X})]_i)$$

so

$$\left[\sum_{i=0}^{n-1} (X_i - X_i^*)^2 \right]^{1/2} \leq \left[\sum_{i=0}^{n-1} (\tilde{X}_i - [f_n^*(\tilde{X})]_i)^2 \right]^{1/2} + \left[\sum_{i=0}^{n-1} \mu_i^2 \right]^{1/2}.$$

Since the eigenvalues of Ψ are strictly less than one $\Psi^i \rightarrow 0$ as $i \rightarrow \infty$.

So we may bound

$$\sum_{i=0}^{n-1} \mu_i^2 \leq h^2 \quad (207)$$

where h does not depend on n or θ . The details of this derivation are carried out in Appendix C. This gives

$$\Delta \leq E_{\theta} \left\{ \left[\left(\sum_{i=0}^{n-1} (\tilde{X}_i - [f_n^*(\tilde{X})]_i)^2 \right)^{1/2} + h \right]^r \right\}. \quad (208)$$

If $r \geq 1$ and $a, b \geq 0$, then

$$(a+b)^r \leq a^r + r b (a+b)^{r-1}$$

which implies

$$\Delta \leq n \delta_{n,R}(\theta) + K' + r h (\bar{D} + h)^{r-1}. \quad (209)$$

So, the distortion mismatch is bounded by

$$|D(f_n^*; \theta) - \delta_{n,R}(\theta)| \leq n^{-1} [K' + rh(\bar{D} + h)^{r-1}]$$

where K' is defined in (204), and from (180)

$$R(f_n^*; \theta) - R \leq n^{-1} [(k+1) \log n + 1 + \sum_{i=1}^k \log(b_i + n^{-1}) + \log \sigma_2^2 - \sigma_1^2 + n^{-1}]$$

where

$$b_i \triangleq \max_{\theta, \varphi \in \Lambda} |a_i - \hat{a}_i|$$

using $\theta = (a_1, \dots, a_k, \sigma^2)$ and $\varphi = (\hat{a}_1, \dots, \hat{a}_k, \hat{\sigma}^2)$. So the code is minimax universal. Notice that only $O(n^{-1})$ terms were added to the rate and distortion in going from a fixed initial state to an arbitrary initial state in some compact set. Again the additional distortion is $O(n^{-1})$ and the dominant term in the additional rate is the number of dimensions of the class Λ times $n^{-1} \log n$.

3.3 Generalization to Unbounded Distortion Measures

Under certain conditions we may remove the restriction that $d_n(\cdot, \cdot)$ is at most $n\bar{D}$, and still get minimax universal codes. In particular this may be done if d_n is a different distortion measure which does not increase exponentially and if the contribution of high distortion terms to the expected distortion is small for any $\theta \in \Lambda$.

Let $B(\omega)$ be a sphere in $\mathbb{R}^n$ with diameter ω and define

$$\tilde{d}_n(\|\underline{x} - \underline{y}\|) = n^{-1} d_n(\underline{x}, \underline{y})$$

where $\|\cdot\|$ is the Euclidean norm. If for all $\theta \in \Lambda$ we have

$$\int_{[B(\omega)]^c} \tilde{d}_n(\|\underline{x}\|) p_\theta(\underline{x}) d\underline{x} \leq f(\omega) \quad (210)$$

where $f(\omega) \rightarrow 0$ as $\omega \rightarrow \infty$ and in addition

$$\tilde{d}_n(\omega)e^{-\omega} \rightarrow 0 \quad (211)$$

as $\omega \rightarrow \infty$, then we may construct universal codes as before. We assume that the quantizer f_n^* has at least one output level in $B(\omega)$. If this is not the case we may add one output level to f_n^* . The effect which this has on the rate is small; M is simply replaced by $M+1$ in (171), so the dominant term is not affected. To bound the distortion we divide the outputs into two sets. For the set $[B(\omega)]^c$ we know that the total expected distortion is at most $f(\omega)$. For $\underline{x} \in B(\omega)$ we have

$$n^{-1} d_n(\underline{x}, f_n^*(\underline{x})) \leq \tilde{d}_n(\omega) \quad (212)$$

so we may bound the distortion as before using $d_n(\omega)$ in place of $\bar{D}$. So if we set $\omega = \log n$ then the distortion here is bounded by

$$D(f_n^*; \theta) \leq n^{-1} \tilde{d}_n(\log n) K^{\frac{1}{2}} (\log e)^{-\frac{1}{2}} + f(\log n) . \quad (213)$$

So from (211) it is clear that f_n^* is minimax universal.

If $\tilde{d}_n(\omega)$ does not increase exponentially with ω then all classes considered here (except for the mixture distributions) satisfy (210). For example, if d_n is the r -th power of the Euclidean distance and $p_\theta(x)$ decays as $e^{-\alpha x}$ then

$$f(\omega) = k \omega^r e^{-\alpha \omega}$$

so that

$$f(\log n) = k n^{-\alpha} (\log n)^r \quad (214)$$

and

$$\tilde{d}(\log n) = (\log n)^r .$$

Then we have

$$D(f_n^*; \theta) - \delta_{n,R}(\theta) \approx \hat{K} n^{-\alpha} (\log n)^r \quad (215)$$

where $\hat{K}$ is a constant. Note that the additional distortion is no longer $O(n^{-1})$ in this case.

APPENDIX A

PROOFS OF THEOREMS FROM CHAPTER 1

Theorem 1.1. Given any two initial distributions μ_0 and ν_0 on $[0,1]$, if μ_1 and ν_1 are generated from (38) then

$$\bar{\rho}(\mu_1, \nu_1) \leq |\lambda| \bar{\rho}(\mu_0, \nu_0) \quad (\text{A.1})$$

where $\lambda \triangleq 1 - \alpha - \beta$.

Proof. The proof is done in three stages. First we assume μ_0 and ν_0 are concentrated on individual points, then finite sets of points, and finally we generalize to arbitrary μ_0, ν_0 .

Assume that ν_0 is concentrated on a point z^* and μ_0 is concentrated on $z^* + \epsilon$, where $z^* \in [0,1]$ and $\epsilon \in (0, 1 - z^*]$. This gives $\nu_0(\{z^*\}) = 1$, $\mu_0(\{z^* + \epsilon\}) = 1$ and $\bar{\rho}(\mu_0, \nu_0) = \epsilon$. Now let μ_1 and ν_1 be the distributions generated from μ_0 and ν_0 using (9). So

$$\mu_1[\{f_k(z^* + \epsilon)\}] = p(k|z^* + \epsilon) \quad (\text{A.2})$$

and

$$\nu_1[\{f_k(z^*)\}] = p(k|z^*) \quad ; \quad k \in A. \quad (\text{A.3})$$

Since μ_1 and ν_1 are one-dimensional, $\bar{\rho}(\mu_1, \nu_1)$ is given by [5]

$$\bar{\rho}(\mu_1, \nu_1) = \int_0^1 |\mu_1[0, z] - \nu_1[0, z]| dz. \quad (\text{A.4})$$

So the $\bar{\rho}$ -distance is the area between the cumulative distributions for μ_1 and ν_1 . We first show that $\mu_1[0, z] - \nu_1[0, z]$ never changes sign. Define a set $B(z, z^*) \subset A$ by

$$B(z, z^*) \triangleq \{k \in A : z \geq f_k(z^*)\}. \quad (\text{A.5})$$

Then we have

$$\mu_1[0, z] = \sum_{k \in B(z, z^* + \epsilon)} p(k|z^* + \epsilon)$$

and

$$\nu_1[0, z] = \sum_{k \in B(z, z^*)} p(k|z^*) .$$

Now $f_k(z)$ is of the form

$$f_k(z) = \frac{a_k z + b_k}{c_k z + d_k} \quad (A.6)$$

where a_k, b_k, c_k , and d_k are constants, and its derivative is

$$f'_k(z) = \frac{a_k d_k - b_k c_k}{(c_k z + d_k)^2} \quad (A.7)$$

so $f_k(z)$ is monotonic. From the definition of f_k (6) we have

$$a_k d_k - b_k c_k = \gamma_0(k) \gamma_1(k) \lambda, \quad (A.8)$$

so if $\lambda \geq 0$ $f_k(z)$ is increasing, and if $\lambda \leq 0$ $f_k(z)$ is decreasing.

Assume that $\lambda \geq 0$. Then

$$f_k(z^* + \epsilon) \geq f_k(z^*) \quad (A.9)$$

for all $k \in A$ so

$$B(z, z^* + \epsilon) \subset B(z, z^*) .$$

We now show that

$$\sum_{k \in B(z, z^*)} [p(k|z^* + \epsilon) - p(k|z^*)] \leq 0 \quad (A.10)$$

for all $k \in A$.

Since $\lambda \geq 0$ (A.10) is equivalent to

$$\sum_{k \in B} [\gamma_0(k) - \gamma_1(k)] \geq 0 \quad . \quad (A.11)$$

(Here we use $B = B(z, z^*)$ for convenience.) If either B or its complement B^c is empty, then clearly (A.11) holds with equality. So we may assume that both B and B^c are non-empty. Suppose (A.11) is false. Then we have

$$\sum_{k \in B} \gamma_0(k) < \sum_{k \in B} \gamma_1(k)$$

and

$$\sum_{k \in B^c} \gamma_0(k) > \sum_{k \in B^c} \gamma_1(k) \quad , \quad (A.12)$$

which together imply

$$\sum_{k \in B^c} \gamma_0(k) - \sum_{k \in B} \gamma_1(k) > \sum_{k \in B} \gamma_0(k) - \sum_{k \in B^c} \gamma_1(k) \quad . \quad (A.13)$$

But if $k \in B$ and $j \in B^c$ then

$$f_j(z^*) > f_k(z^*) \quad (A.14)$$

(recall $B = B(z, z^*)$). So we have

$$\frac{\gamma_1(j)\eta_1(z^*)}{\sum_{i=0}^1 \gamma_i(j)\eta_i(z^*)} > \frac{\gamma_1(k)\eta_1(z^*)}{\sum_{i=0}^1 \gamma_i(k)\eta_i(z^*)} \quad (A.15)$$

which implies

$$\gamma_1(j)\gamma_0(k) > \gamma_1(k)\gamma_0(j) \quad . \quad (A.16)$$

If we sum both sides of (A.16) over all pairs (k, j) such that $k \in B$ and $j \in B^c$ we have

$$\sum_{j \in B^c} \gamma_1(j) \sum_{k \in B} \gamma_0(k) \geq \sum_{k \in B} \gamma_1(k) \sum_{j \in B^c} \gamma_0(j) . \quad (\text{A.17})$$

But this contradicts (A.13), hence (A.11) must be true.

So we have

$$\begin{aligned} \mu_1[0, z] &= \sum_{k \in B(z, z^* + \epsilon)} p(k|z^* + \epsilon) \\ &\leq \sum_{k \in B(z, z^*)} p(k|z^* + \epsilon) \\ &\leq \sum_{k \in B(z, z^*)} p(k|z^*) \\ &= \nu_1[0, z] \end{aligned} \quad (\text{A.18})$$

for all $z \in [0, 1]$ as desired.

For $\lambda \leq 0$ we have

$$B(z, z^*) \subset B(z, z^* + \epsilon) \quad (\text{A.19})$$

and since (A.11) still holds, it follows that

$$\sum_{k \in B(z, z^*)} [p(k|z^* + \epsilon) - p(k|z^*)] \geq 0 \quad (\text{A.20})$$

so in this case

$$\mu_1[0, z] \geq \nu_1[0, z] \quad ; \quad \forall z \in [0, 1] . \quad (\text{A.21})$$

In either case the absolute value in the definition of the $\bar{\rho}$ -distance (A.4) may be taken outside the integral, hence we have

$$\bar{\rho}(\mu_1, \nu_1) = \left| \int_0^1 [\mu_1[0, z] - \nu_1[0, z]] dz \right| . \quad (\text{A.22})$$

Now

$$\int_0^1 [1 - \mu_1[0, z]] dz$$

is the expected value of z under μ_1 so

$$\begin{aligned} \bar{\rho}(\mu_1, \nu_1) &= \left| \sum_{k \in A} [p(k|z^* + \epsilon) f_k(z^* + \epsilon) - p(k|z^*) f_k(z^*)] \right| \\ &= \left| \sum_{k \in A} \gamma_1(k) [\eta_1(z^* + \epsilon) - \eta_1(z^*)] \right| \\ &= |\lambda| \epsilon \\ &= |\lambda| \bar{\rho}(\mu_0, \nu_0) . \end{aligned} \quad (\text{A.23})$$

If μ_0 and ν_0 are concentrated on a finite set of points the result generalizes as follows. Let $\mathcal{P}$ be the class of distributions on $[0, 1]$ which are concentrated on a finite set of points. Let α^* be the joint distribution which achieves $\bar{\rho}(\mu_0, \nu_0)$. (Since μ_0 and ν_0 are one-dimensional α^* is easily determined.) Now α^* is also concentrated on a finite number of points, say N , so for some set $\{(x_i, y_i)\}_{i=1}^N$ we have

$$\alpha^*({(x_i, y_i)}) = \theta_i \geq 0 \quad (\text{A.24})$$

and $\sum_{i=1}^N \theta_i = 1$. Define probability measures $\mu_0^{(i)}$, $\nu_0^{(i)}$ and $\alpha_0^{(i)}$, $i = 1, 2, \dots, N$ by

$$\begin{aligned} \mu_0^{(i)}(\{x_i\}) &= 1 \\ \nu_0^{(i)}(\{y_i\}) &= 1 \\ \alpha_0^{(i)}(\{(x_i, y_i)\}) &= 1 . \end{aligned} \quad (\text{A.25})$$

Then let $\mu_1^{(i)}$ and $\nu_1^{(i)}$ be generated using (9). Let $\alpha_1^{(i)}$ be the joint distribution with marginals $\mu_1^{(i)}$ and $\nu_1^{(i)}$ which achieves $\bar{\rho}(\mu_1^{(i)}, \nu_1^{(i)})$. For each i (10) and (A.23) imply

$$E_{\alpha_1^{(i)}}[|x-y|] \leq |\lambda| \cdot |x_1 - y_1| \quad (\text{A.26})$$

since $\mu_0^{(i)}$ and $\nu_0^{(i)}$ are concentrated on single points. Further

$\alpha_1 \triangleq \sum_{i=1}^N \theta_i \alpha_1^{(i)}$ is a joint distribution with marginals μ_1 and ν_1 , hence

$$\begin{aligned} \bar{\rho}(\mu_1, \nu_1) &\leq E_{\alpha_1}[|x-y|] = \sum_{i=1}^N \theta_i E_{\alpha_1^{(i)}}[|x-y|] \\ &= \sum_{i=1}^N \theta_i |\lambda| \cdot |x_1 - y_1| \\ &= |\lambda| E_{\alpha^*}[|x-y|] \\ &= |\lambda| \bar{\rho}(\mu_0, \nu_0). \end{aligned} \quad (\text{A.27})$$

Now consider arbitrary distributions μ_0 and ν_0 . Define a sequence of distributions μ_0^N for $N = 1, 2, \dots$ by

$$\mu_0^N[0, x] = jN^{-1} \text{ if } \mu_0[0, x] \in ((j-1)N^{-1}, jN^{-1}] ; j = 1, \dots, N.$$

Then

$$\begin{aligned} \bar{\rho}(\mu_0^N, \mu_0) &= \int_0^1 |\mu_0^N[0, x] - \mu_0[0, x]| dx \\ &= \sum_{j=1}^N \int_{x_{j-1}}^{x_j} |\mu_0^N[0, x] - \mu_0[0, x]| dx \\ &\leq \sum_{j=1}^N \int_{x_{j-1}}^{x_j} |jN^{-1} - (j-1)N^{-1}| dx = N^{-1} \end{aligned}$$

where x_j is such that $\mu_0^N[0,x] \leq jN^{-1}$ for $x < x_j$ and $\mu_0^N[0,x] > jN^{-1}$ for $x > x_j$.

Now if μ_1 and μ_1^N are generated from μ_0 and μ_0^N using (9) we have

$$|\mu_1^N[0,x] - \mu_1[0,x]| = \left| \sum_{i \in A} \int_{f_i^{-1}([0,x])} p(i|z) (\mu_0^N(dz) - \mu_0(dz)) \right|. \quad (A.28)$$

Since f_i is monotonic $f_i^{-1}([0,x])$ is an interval, say $[w_i, y_i]$, and using integration by parts

$$\begin{aligned} \int_{w_i}^{y_i} p(i|z) (\mu_0^N(dz) - \mu_0(dz)) &= p(i|z) (\mu_0^N[0,z] - \mu_0[0,z]) \Big|_{w_i}^{y_i} \\ &\quad + \int_{w_i}^{y_i} (\mu_0^N[0,z] - \mu_0[0,z]) p'(i|z) dz. \end{aligned}$$

Now $\frac{d}{dz} p(i|z) = \lambda[\gamma_1(i) - \gamma_0(i)] \leq 1$ so

$$\begin{aligned} \int_{w_i}^{y_i} p(i|z) (\mu_0^N(dz)) &\leq \mu_0^N[0, w_i] - \mu_0[0, w_i] + \mu_0^N[0, y_i] - \mu_0[0, y_i] \\ &\quad + \int_{w_i}^{y_i} (\mu_0^N[0,z] - \mu_0[0,z]) dz \\ &\leq 2N^{-1} + \bar{\rho}(\mu_0^N, \mu_0) \\ &\leq 3N^{-1}. \end{aligned}$$

It follows from (A.28) that

$$|\mu_1^N[0,x] - \mu_1[0,x]| \leq 6N^{-1}$$

so

$$\bar{\rho}(\mu_1^N, \mu_1) \leq \int_0^1 6N^{-1} dx = 6N^{-1}. \quad (A.29)$$

If we define v_1^N similarly we have $\bar{\rho}(v_1^N, v_1) \leq 6N^{-1}$. We know

$$\bar{\rho}(\mu_1^N, v_1^N) \leq |\lambda| \bar{\rho}(\mu_0^N, v_0^N)$$

so by the triangle inequality,

$$\begin{aligned}
 \bar{\rho}(\mu_1, \nu_1) &\leq |\lambda| \bar{\rho}(\mu_0^N, \nu_0^N) + 12N^{-1} \\
 &\leq |\lambda| [\bar{\rho}(\mu_0, \nu_0) + 2N^{-1}] + 12N^{-1} \\
 &\leq |\lambda| \bar{\rho}(\mu_0, \nu_0) + 14N^{-1}.
 \end{aligned} \tag{A.30}$$

Since (A.30) holds for all N we have

$$\bar{\rho}(\mu_1, \nu_1) \leq |\lambda| \bar{\rho}(\mu_0, \nu_0)$$

for arbitrary distributions μ_0, ν_0 on $[0, 1]$.

Theorem 1.4.

$$F_1(z_0, \hat{z}_0) \stackrel{\Delta}{=} \sum_{\underline{x} \in A} |z_1(\underline{x}^1) - \hat{z}_1(\underline{x}^1)| \tilde{p}(\underline{x}^1 | z_0) \leq |\lambda|^1 \cdot |z_0 - \hat{z}_0| \tag{A.31}$$

where $z_1(\underline{x}^1)$ and $\hat{z}_1(\underline{x}^1)$ are derived from initial estimates z_0 and $\hat{z}_0$ using the recursion (6) 1 times.

Proof. We will use induction on i . For $i = 1$

$$F_1(z_0, \hat{z}_0) = \sum_{j \in A} |f_j(z_0) - f_j(\hat{z}_0)| p(j | z_0) \tag{A.32}$$

where f_j and $p(j | \cdot)$ are as defined in (6) and (8). Assume $\lambda(z_0 - \hat{z}_0) \geq 0$.

Then $f_j(z_0) \geq f_j(\hat{z}_0)$ so we may remove the absolute value brackets in (A.32).

Now $\sum_{j \in A} f_j(\hat{z}_0) p(j | z_0)$ is the expected value of a distribution which assigns

probability $p(j | z_0)$ to the point $f_j(\hat{z}_0)$; $j \in A$. Let $\tilde{\mu}$ be the probability measure for this distribution. Then

$$\begin{aligned}\tilde{\mu}[0, z] &= \sum_{j \in B(z, \hat{z}_0)} p(j|z_0) \\ &\geq v[0, z] \triangleq \sum_{j \in B(z, \hat{z}_0)} p(j|\hat{z}_0)\end{aligned}$$

where $B(z, \hat{z}_0)$ is as defined in (A.5) and the inequality follows from (A.20).

Next

$$\int_0^1 v[0, z] dz = \sum_{j \in A} f_j(\hat{z}_0) p(j|\hat{z}_0)$$

so

$$\begin{aligned}F_1(z_0, \hat{z}_0) &\leq \sum_{j \in A} [f_j(z_0) p(j|\hat{z}_0)] \\ &= |\lambda(z_0 - \hat{z}_0)|.\end{aligned}\tag{A.33}$$

The same result follows if $\lambda(z_0 - \hat{z}_0) \leq 0$ using corresponding inequalities.

So the result holds for $i = 1$.

Now we assume $F_j(z_0, \hat{z}_0) \leq |\lambda|^j \cdot |z_0 - \hat{z}_0|$ for $j \leq i$. Then

$$\begin{aligned}F_{i+1}(z_0, \hat{z}_0) &= \sum_{\underline{x}^{i+1}} |z_{i+1}(\underline{x}^{i+1}) - \hat{z}_{i+1}(\underline{x}^{i+1})| \tilde{p}(\underline{x}^{i+1}|z_0) \\ &= \sum_{\underline{x}^i} \left\{ \sum_{x_{i+1}} |f_{x_{i+1}}(z_i(\underline{x}^i)) - f_{x_{i+1}}(\hat{z}_i(\underline{x}^i))| p(x_{i+1}|z_i(\underline{x}^i)) \right\} \tilde{p}(\underline{x}^i|z_0) \\ &\leq \sum_{\underline{x}^i} |\lambda| \cdot |z_i(\underline{x}^i) - \hat{z}_i(\underline{x}^i)| \tilde{p}(\underline{x}^i|z_0) \\ &= |\lambda| \cdot F_i(z_0, \hat{z}_0) \\ &\leq |\lambda|^{i+1} |z_0 - \hat{z}_0|\end{aligned}\tag{A.34}$$

where (A.34) follows from (A.33).

Theorem 1.5. If $\theta, \varphi \in \Lambda'(\delta)$ then

$$\begin{aligned} F_1(z_0, \hat{z}_0) &\triangleq \sum_{\underline{x}} |\hat{z}_1(\underline{x}^1) - z_1(\underline{x}^1)| \tilde{p}(\underline{x}^1 | z_0) \\ &\leq |\lambda_\theta|^1 + K_\epsilon \{1 - |\lambda_\theta|\}^{-1} \end{aligned} \quad (\text{A.35})$$

where

$$K_\epsilon \triangleq \delta^{-2} [3\epsilon + 3\epsilon^2 + \epsilon^3] + \epsilon, \quad (\text{A.36})$$

$\tilde{p}(\underline{x}^1 | z_0)$ is the probability that $\underline{x}^1$ is output from source θ if the initial estimate is z_0 , and $\lambda_\theta \triangleq 1 - \alpha - \beta$.

Proof. First by (6), (8), and (54)

$$\begin{aligned} |\hat{f}_j(z) - \tilde{f}_j(z)| &\leq \left| \frac{\eta_{j1}}{\eta_{j0} + \eta_{j1}} - \frac{\hat{\eta}_{j1}}{\hat{\eta}_{j1} + \hat{\eta}_{j0}} \right| + \epsilon \\ &\leq \delta^{-2} |\eta_{j1}\hat{\eta}_{j0} - \hat{\eta}_{j1}\eta_{j0}| + \epsilon \end{aligned} \quad (\text{A.37})$$

where $\eta_{ij} = \gamma_i(j)\eta_i(z)$ and $\hat{\eta}_{ij}$ is defined as η_{ij} using the parameters for source φ . Next, if we define

$$K = [(1-\beta)z + \alpha(1-z)] \cdot [\beta'z + (1-\alpha')(1-z)]$$

and

$$K' = [\beta z + (1-\alpha)(1-z)] \cdot [(1-\beta)'z + \alpha'(1-z)]$$

then $K - K' \leq \epsilon$. So since $\eta_{j1}\hat{\eta}_{j0} = \gamma_1(j)\gamma_0'(j)K$ and $\hat{\eta}_{j1}\eta_{j0} = \gamma_1'(j)\gamma_0(j)K'$ we have

$$\begin{aligned} |\eta_{j1}\hat{\eta}_{j0} - \hat{\eta}_{j1}\eta_{j0}| &\leq |(\gamma_1'(j) + \epsilon)(\gamma_0(j) + \epsilon)(K' + \epsilon) - \gamma_1'(j)\gamma_0(j)K'| \\ &\leq 3\epsilon + 3\epsilon^2 + \epsilon^3. \end{aligned} \quad (\text{A.38})$$

So we have $f_j(z) - \hat{f}_j(z) \leq \delta^{-2}[3\epsilon + 3\epsilon^2 + \epsilon^3]$ and

$$\begin{aligned} F_1(z_0, \hat{z}_0) &\leq \sum_j |f_j(z_0) - \hat{f}_j(\hat{z}_0)| p(j|z_0) + \epsilon \\ &\leq \sum_j |f_j(z_0) - f_j(\hat{z}_0)| p(j|z_0) + \sum_j |f_j(\hat{z}_0) - \hat{f}_j(\hat{z}_0)| p(j|z_0) + \epsilon \\ &\leq |\lambda_\theta| \cdot |z_0 - \hat{z}_0| + K_\epsilon \end{aligned} \quad (\text{A.39})$$

$$\leq |\lambda_\theta| + K_\epsilon. \quad (\text{A.40})$$

Equation (A.39) follows from Theorem 1.4. Next

$$F_{i+1}(z_0, \hat{z}_0) = \sum_{\underline{x}^i} \left\{ \sum_{x_{i+1}} |f_{x_{i+1}}(z_i(\underline{x}^i)) - \hat{f}_{x_{i+1}}(\hat{z}_i(\underline{x}^i))| p(x_{i+1}|z_i(\underline{x}^i)) \right\} \tilde{p}(\underline{x}^i|z_0)$$

so from (A.39)-(A.40)

$$\begin{aligned} F_{i+1}(z_0, \hat{z}_0) &\leq |\lambda_\theta| \sum_{\underline{x}^i} |z_i(\underline{x}^i) - \hat{z}_i(\underline{x}^i)| \tilde{p}(\underline{x}^i|z_0) + K_\epsilon \\ &= |\lambda_\theta| F_i(z_0, \hat{z}_0) + K_\epsilon \end{aligned} \quad (\text{A.41})$$

and we solve the recursion (A.41) with initial condition (A.40) to get

$$F_i(z_0, \hat{z}_0) \leq |\lambda_\theta|^i + K_\epsilon [1 - |\lambda_\theta|]^{-1}. \quad (\text{A.42})$$

AD-A124 492

SOME TECHNIQUES IN UNIVERSAL SOURCE CODING AND DURING
FOR COMPOSITE SOURCES(U) ILLINOIS UNIV AT URBANA
COORDINATED SCIENCE LAB M S WALLACE DEC 81 R-929

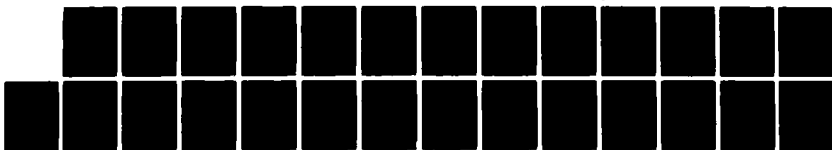
22

UNCLASSIFIED

NO0014-79-C-0424

F/G 9/4

NL



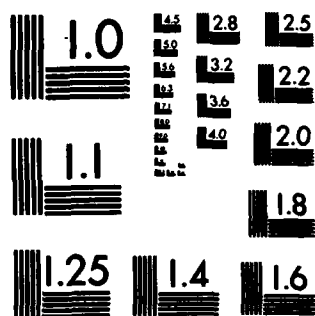
END

DATE

FILED

83

DTIC



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

Theorem 1.6. If $\theta \in \hat{\Lambda}(\epsilon)$, $\theta = (Q, \{P(x)\})$ then

$$D_n \triangleq \|\underline{z} \tilde{T}(\underline{x}^n) [\underline{z} \tilde{T}(\underline{x}^n)]^{-1} - \underline{\hat{z}} \tilde{T}(\underline{x}^n) [\underline{\hat{z}} \tilde{T}(\underline{x}^n)]^{-1}\| \leq \zeta^{n-1} \quad (\text{A.43})$$

for any probability row vectors $\underline{z}$ and $\underline{\hat{z}}$, where ζ is defined in (65) and $\underline{x}^n = (x_1, \dots, x_n)$.

Proof. Throughout this proof $\underline{x}^n$ will be a fixed vector of source outputs and $\underline{x}^i$ will be used to denote the first i components of $\underline{x}^n$ for $i \leq n$. Define $\underline{q}(j) = \tilde{T}(\underline{x}^j)_1$ and let C_j be the set of columns of $\tilde{T}(\underline{x}^j)$. Then

$$D_n = \max_{\underline{u} \in C_n} \left[\frac{\underline{z} \cdot \underline{u}}{\underline{z} \cdot \underline{q}(n)} - \frac{\underline{\hat{z}} \cdot \underline{u}}{\underline{\hat{z}} \cdot \underline{q}(n)} \right] \quad (\text{A.44})$$

$$\leq \Delta_n \triangleq \max_{\underline{u} \in C_n} \left[\max_{i \in V_n} \frac{u_i}{\sigma_i(n)} - \min_{i \in V_n} \frac{u_i}{\sigma_i(n)} \right] \quad (\text{A.45})$$

where

$$V_n \triangleq \{i: \sigma_i(n) > 0\}, \quad (\text{A.46})$$

since z_i , $\hat{z}_i$, and $\sigma_i(n)$ are non-negative. We will prove $\Delta_n \leq \zeta^{n-1}$ by induction. First $\Delta_1 \leq 1$ since $u_i \geq 0$ and $\sigma_i(1) \geq u_i$ for $u \in C_1$. Now assume $\Delta_{j-1} \leq \zeta^{j-2}$. Note that

$$\frac{[QP(x_j)\underline{u}]_1}{[QP(x_j)\underline{q}(j-1)]_1} = \frac{u'_1}{\sigma_1(j)} \quad (\text{A.47})$$

where $\underline{u} \in C_{j-1}$ and $\underline{u}' \in C_j$. Also if we define

$$\hat{V}_j = \{i: [P(x_j)\underline{q}(j-1)]_i > 0\} \quad (\text{A.48})$$

$\hat{v}_j \subset v_j$, and since $p(x_j|i) > 0$ implies

$$\frac{[P(x_j)\underline{u}]_i}{[P(x_j)\underline{q}(j-1)]_i} = \frac{u_i}{\sigma_i(j-1)} \quad (\text{A.49})$$

we have

$$\Delta_{j-1} \geq \max_{\underline{u} \in C_{j-1}} \left\{ \max_{i \in \hat{v}_j} \frac{[P(x_j)\underline{u}]_i}{[P(x_j)\underline{q}(j-1)]_i} - \min_{i \in \hat{v}_j} \frac{[P(x_j)\underline{u}]_i}{[P(x_j)\underline{q}(j-1)]_i} \right\}. \quad (\text{A.50})$$

Let $\underline{w} = P(x_j)\underline{q}(j-1)$ and $\underline{y} \triangleq P(x_j)\underline{u}$ for some fixed $\underline{x} \in C_{j-1}$. Let $a(\underline{u}) \triangleq \max_{i \in \hat{v}_j} [y_i/w_i]$ and $b(\underline{u}) \triangleq \min_{i \in \hat{v}_j} [y_i/w_i]$. Also define

$$\begin{aligned} \hat{a}(\underline{u}) &= \max_i \frac{[Q \underline{y}]_i}{[Q \underline{w}]_i} \\ \hat{b}(\underline{u}) &= \min_i \frac{[Q \underline{y}]_i}{[Q \underline{w}]_i}. \end{aligned} \quad (\text{A.51})$$

Then if $W \triangleq \{(k,m): y_k w_m > y_m w_k\}$

$$\begin{aligned}
\hat{a}(\underline{u}) - \hat{b}(\underline{u}) &= \max_{i,l} \frac{\sum_{k,m} q(i|k)q(l|m)(y_{k^w m} - y_{m^w k})}{\sum_{k,m} q(i|k)q(l|m)w_k w_m} \\
&= \max_{i,l} \frac{\sum_{(k,m) \in W} [q(i|k)q(l|m) - q(i|m)q(l|k)](y_{k^w m} - y_{m^w k})}{\sum_{(k,m) \in W} [q(i|k)q(l|m) + q(i|m)q(l|k)]w_k w_m} \\
&\leq \max_{i,l} \max_{(k,m) \in W} \frac{[q(i|k)q(l|m) - q(i|m)q(l|k)](y_{k^w m} - y_{m^w k})}{[q(i|k)q(l|m) + q(i|m)q(l|k)]w_k w_m} \\
&\leq \frac{(1 - (S-1)\epsilon)^2 - \epsilon^2}{(1 - (S-1)\epsilon)^2 + \epsilon^2} \max_{(k,m) \in W} \frac{(y_{k^w m} - y_{m^w k})}{w_k w_m} \\
&= \zeta(a(\underline{u}) - b(\underline{u})).
\end{aligned} \tag{A.52}$$

From (A.46) and (A.51) we have

$$\begin{aligned}
\Delta_j &= \max_{\underline{u} \in C_{j-1}} (\hat{a}(\underline{u}) - \hat{b}(\underline{u})) \\
&\leq \max_{\underline{u} \in C_{j-1}} [\zeta(a(\underline{u}) - b(\underline{u}))] \\
&\leq \zeta \Delta_{j-1} \\
&\leq \Delta^{j-1}
\end{aligned} \tag{A.53}$$

as desired.

Theorem 1.7. If $\theta, \varphi \in \Lambda'(\delta)$, and the parameters for θ and φ are within ϵ , then if ℓ_n is the code for φ we have

$$r_n(\ell_n, \theta) \leq Kn^{-1} + \hat{K}\epsilon.$$

Proof. For convenience let $\tilde{q}(\underline{x}^1|z_0) = \tilde{p}_\varphi(\underline{x}^1|z_0)$ and $\tilde{p}(\underline{x}^1|z_0) = \tilde{p}_\theta(\underline{x}^1|z_0)$ in this proof. Then

$$\begin{aligned} r_n(\ell_n, \theta) &= n^{-1} \{ 1 + \sum_{\underline{x}^n} \tilde{p}(\underline{x}^n|z_0) [\min_{k=0,1} \{ -\log \tilde{q}(\underline{x}^n|k) \}] + \log \tilde{p}(\underline{x}^n|z_0) \} \\ &\leq n^{-1} 2 + \sum_{\underline{x}^n} \tilde{p}(\underline{x}^n|z_0) \sum_{i=0}^{n-1} \log \frac{p(x_{i+1}|z_i(\underline{x}^1))}{q(x_{i+1}|\hat{z}_i(\underline{x}^1))} \end{aligned} \quad (A.54)$$

where $\hat{z}_i(\underline{x}^1)$ is the estimate derived from outputs $\underline{x}^i$ with $\hat{z}_0 = 0$. Now since $\log x \leq (x-1)\log e$ and $\tilde{p}(\underline{x}^1|z_0) = \sum_{x_{i+1}, \dots, x_n} \tilde{p}(\underline{x}^n|z_0)$

we have

$$\begin{aligned} r_n(\ell_n, \theta) &\leq n^{-1} \{ 2 + \sum_{i=0}^{n-1} \delta^{-1} \log e \sum_{\underline{x}^{i+1}} \tilde{p}(\underline{x}^{i+1}|z_0) \cdot \\ &\quad |p(x_{i+1}|z_i(\underline{x}^1)) - q(x_{i+1}|\hat{z}_i(\underline{x}^1))| \} \quad (A.55) \end{aligned}$$

Now

$$\begin{aligned} |p(x|z) - q(x|\hat{z})| &\leq |p(x|z) - p(x|\hat{z})| + |p(x|\hat{z}) - q(x|\hat{z})| \\ &= |\gamma_0(x) - \gamma_1(x)| \cdot |z - \hat{z}| + |\eta_{x1}(\hat{z}) + \eta_{x0}(\hat{z}) - \hat{\eta}_{x1}(\hat{z}) - \hat{\eta}_{x0}(\hat{z})| \end{aligned} \quad (A.56)$$

where η_{x1} is defined in (A.37).

The second term is at most $4\epsilon + 2\epsilon^2$ so

$$r_n(l_n, \theta) \leq 2n^{-1} + (n\delta)^{-1} \log e \sum_{i=0}^{n-1} \left\{ \sum_{\underline{x}} \tilde{p}(\underline{x}^{i+1} | z_0) \cdot |\gamma_0(x) - \gamma_1(x)| \cdot |z_i(\underline{x}^i) - \hat{z}_i(\underline{x}^i)| + 4\epsilon + 2\epsilon^2 \right\} \quad (\text{A.57})$$

and from Theorem 1.5

$$r_n(l_n, \theta) \leq Kn^{-1} + \hat{K}\epsilon \quad (\text{A.58})$$

where K and $\hat{K}$ are defined in (95) and (96).

Theorem 1.8. If $r_n(\theta) = R_n(\theta) - \mathcal{K}_c(\theta)$ is the redundancy of the code l_n then

$$r_n(\theta) \leq n^{-1} (1 - \frac{1}{\alpha} \log J) \quad (\text{A.59})$$

Proof. The entropy of blocks of source output $\underline{x} \in A^n$ given no previous outputs is

$$H_n(\underline{x}) = \sum_{\underline{x} \in A^n} \tilde{p}(\underline{x} | z^*) \log \tilde{p}(\underline{x} | z^*) \quad (\text{A.60})$$

since z^* is the stationary distribution of the switching process Z_i . Further we know

$$\begin{aligned} n^{-1} H_n(\underline{x}) &\geq H(X_1 | X_0, X_{-1}, \dots) \\ &= \mathcal{K}_c(\theta), \end{aligned} \quad (\text{A.61})$$

where $\mathcal{K}_c(\theta)$ is the entropy of the source. The average rate of the code l_n applied to source θ is

$$R_n(\theta) = n^{-1} \sum_{\underline{x} \in A^n} \tilde{p}(\underline{x} | z^*) l_n(\underline{x}), \quad (\text{A.62})$$

and since l_n is the Huffman code for $\{\tilde{p}(\underline{x} | z^*) : \underline{x} \in A^n\}$ we have

$$nR_n(\theta) \geq H_n(\underline{X}) \geq nR_n(\theta) - 1 \quad . \quad (A.63)$$

Let Z_0 be the initial state of the switching process and let T be the first switching time (we set $T = n$ if the first switch occurs after time n).

Then

$$\begin{aligned} nK_c(\theta) &= H_n(\underline{X}|X_0, \dots) \\ &\geq H_n(\underline{X}|Z_0, T, X_0, \dots) \\ &= H_n(\underline{X}|Z_0, T) \end{aligned} \quad (A.64)$$

since X_i , $i \geq 1$, is independent of $(X_0, X_{-1}, \dots)$ given Z_0 . Next

$$H_n(\underline{X}|Z_0, T) = E[-\sum_{\underline{x} \in A^T} \gamma_{Z_0}(\underline{x}) \log \gamma_{Z_0}(\underline{x}) - \sum_{\underline{x} \in A^{n-T}} \tilde{p}(\underline{x}|z^*) \log \tilde{p}(\underline{x}|z^*)] \quad (A.65)$$

where the expectation is over Z_0 and T , and

$$\begin{aligned} \gamma_{\underline{y}}(\underline{x}) &\triangleq \prod_{i=1}^m \gamma_{\underline{y}}(x_i) \\ &= \prod_{i=1}^m P[X_i = x_i | Z_0 = \underline{y}] \end{aligned} \quad (A.66)$$

for $\underline{x} \in A^m$. We know

$$H_{n-k}(\underline{X}) \geq H_n(\underline{X}) - k \log J$$

since the source alphabet has J letters. Since the second sum of (A.65) is the expectation over T of $H_{n-T}(\underline{X})$ and since the first sum is positive we have

$$\begin{aligned}
H_n(\underline{X}|Z_0, T) &\geq E[H_{n-T}(\underline{X})] \\
&\geq H_n(\underline{X}) - E[T] \log J \\
&\geq H_n(\underline{X}) - \frac{1}{\alpha} \log J .
\end{aligned}
\tag{A.67}$$

So from (A.61) and (A.64) we have

$$H_n(\underline{X}) \geq n\mathcal{K}_c(\theta) \geq H_n(\underline{X}) - \frac{1}{\alpha} \log J \tag{A.68}$$

hence from (A.63)

$$R_n(\theta) - \mathcal{K}_c(\theta) \leq n^{-1} (1 + \frac{1}{\alpha} \log J) \tag{A.69}$$

as desired.

APPENDIX B

PROOF OF REDUNDANCY BOUND FOR THE VL-FL CODE OF CHAPTER 2

In [13] a set $\Phi_n = \{\varphi_i; i = 1, \dots, \gamma_n\}$ is constructed with

$$\gamma_n = S \exp_2[S\{\lceil \frac{1}{2}(J-1) \log 9Jn \rceil + \lceil 2 \log J \rceil + J\}] \quad (\text{B.1})$$

such that if $\theta \in \Lambda$ has transition probabilities p_θ and initial state s_0 then there exists a $\varphi \in \Phi_n$ with transition probabilities p_φ and initial state s_0 such that

$$\mathcal{K}_r(p_\theta, p_\varphi; j) \stackrel{\Delta}{=} \sum_{x \in A} p_\theta(x|j) \log \frac{p_\theta(x|j)}{p_\varphi(x|j)} \leq n^{-1} \quad (\text{B.2})$$

for all $j = 1, 2, \dots, S$. Further, if φ is in Φ_n then

$$\min\{p_\varphi(x|j): j \in \mathcal{J}, x \in A\} = \frac{1}{9Jn}. \quad (\text{B.3})$$

Let Φ_n be this set. From (152), if $\hat{n} \leq n + \lceil \log \gamma_n \rceil$ then

$$R_n^*(\Gamma_n^*, \theta) \leq [n + \lceil \log \gamma_n \rceil] \cdot [\bar{l}_\theta(\Gamma_n)]^{-1} \quad (\text{B.4})$$

where Γ_n is the code for φ and $\mathcal{K}_r(p_\theta, p_\varphi; j) \leq n^{-1}$ for $j = 1, 2, \dots, S$. For convenience denote p_θ by ν and p_φ by η . Now 2^{-n} is the average probability of the strings in Γ_n , so by (148)

$$\eta(\underline{x})/2^{-n} \leq \frac{1}{\alpha} \leq 9Jn \quad (\text{B.5})$$

where α is the minimum transition probability of φ . So we have

$$\begin{aligned} n &= \sum_{\underline{x} \in \Gamma_n} \nu(\underline{x}) \{ \log[\eta(\underline{x})/2^{-n}] + \log[\nu(\underline{x})/\eta(\underline{x})] - \log \nu(\underline{x}) \} \\ &\leq \sum_{\underline{x} \in \Gamma_n} \nu(\underline{x}) \{ -\log \nu(\underline{x}) + \log[\nu(\underline{x})/\eta(\underline{x})] \} + \log 9Jn. \end{aligned} \quad (\text{B.6})$$

Next define

$$\begin{aligned} \Omega &\triangleq \sum_{\underline{x} \in \Gamma_n} v(\underline{x}) \log[v(\underline{x})/\eta(\underline{x})] = \sum_{\underline{x} \in \Gamma_n} v(\underline{x}) \sum_{i=1}^{l(\underline{x})} \log \frac{v(x_i | s_{i-1})}{\eta(x_i | s_{i-1})} \\ &= \sum_{i=1}^L \sum_{\alpha \in A} \sum_{\substack{\underline{x} \in B(\alpha) \\ l(\underline{x}) \geq i}} v(\underline{x}) \log \frac{v(x_i | s_{i-1})}{\eta(x_i | s_{i-1})} \end{aligned} \quad (B.7)$$

where $s_0 = j$, $s_i = f(x_{i-1}, s_{i-1})$, $L = \max\{l(\underline{x}) : \underline{x} \in \Gamma_n\}$ and

$$B(\alpha) \triangleq \{\underline{x} \in \Gamma_n : x_k = \alpha_k, 1 \leq k \leq i-1, l(\underline{x}) \geq i\}. \quad (B.8)$$

Note $B(\alpha) = \Gamma_n$ for $i = 1$. Next, since

$$v(\underline{x}) = v(\alpha) \prod_{k=1}^{l(\underline{x})} v(x_k | s_{k-1}), \quad (B.9)$$

if we split $\sum_{\underline{x} \in B(\alpha)}$ into $\sum_{\beta \in A} \sum_{\substack{\underline{x} \in B(\alpha) \\ x_i = \beta}}$ we have

$$\Omega = \sum_{i=1}^L \sum_{\alpha \in A} \sum_{\substack{\underline{x} \in B(\alpha) \\ l(\underline{x}) \geq i}} v(\underline{x}) \sum_{\beta \in A} v(\beta | s_{i-1}) \log \frac{v(\beta | s_{i-1})}{\eta(\beta | s_{i-1})} \sum_{\substack{\underline{x} \in B(\alpha) \\ x_i = \beta}} \sum_{k=i+1}^{l(\underline{x})} v(x_k | s_{k-1}). \quad (B.10)$$

Now in an encoding tree the total probability of the leaves which may be reached by passing through a given node is equal to the probability of that node. The innermost sum is the total probability of the paths from a node to all of its leaves, which is the total probability of the leaves divided by the probability of the node. Hence this sum is one unless it is empty. Further, since each non-terminal node has J successors (one for each $x \in A$) if no $\underline{x} \in B(\alpha)$ exists such that $x_i = \beta$ then all $\underline{x} \in B(\alpha)$ must be of length $i-1$. So if the innermost sum is zero, then

the sum over β is also zero. So we have

$$\Omega = \sum_{i=1}^L \sum_{\alpha \in A^{i-1}} v(\alpha) f(\alpha) \sum_{\beta \in A} v(\beta | s_{i-1})^{\log \frac{v(\beta | s_{i-1})}{\eta(\beta | s_{i-1})}} \quad (\text{B.11})$$

where

$$f(\alpha) = \begin{cases} 1 & \text{if } l(\underline{x}) \geq i \text{ for some } \underline{x} \in B(\alpha) \\ 0 & \text{otherwise} \end{cases}$$

Now the innermost sum is $K_r(p_\theta, p_\varphi; s_{i-1}) \leq n^{-1}$ for any s_{i-1} , so

$$\Omega < n^{-1} \sum_{i=1}^L \sum_{\alpha \in A^{i-1}} v(\alpha) f(\alpha) \quad (\text{B.12})$$

$$\begin{aligned} &= n^{-1} \sum_{i=1}^L \sum_{\substack{\underline{x} \in \Gamma_n \\ l(\underline{x}) \geq i}} v(\underline{x}) \\ &= n^{-1} \sum_{\underline{x} \in \Gamma_n} v(\underline{x}) l(\underline{x}) = n^{-1} \bar{l}_\theta(\Gamma_n) . \end{aligned} \quad (\text{B.13})$$

We substitute this into (B.6) to get

$$[n + \lceil \log \gamma_n \rceil] [\bar{l}_\theta(\Gamma_n)]^{-1} \leq \left\{ \sum_{\underline{x} \in \Gamma_n} v(\underline{x}) \log v(\underline{x}) + \log 9Jn + \lceil \log \gamma_n \rceil \right\} \cdot [\bar{l}_\theta(\Gamma_n)]^{-1} + n^{-1} \quad (\text{B.14})$$

$$\leq K(\Gamma_n, \theta) + \{\log 9Jn + \lceil \log \gamma_n \rceil\} \cdot [\bar{l}_\theta(\Gamma_n)]^{-1} + n^{-1}. \quad (\text{B.15})$$

If $n > \log 9Jn$, then (B.15) implies

$$\bar{l}_\theta(\Gamma_n) \geq \frac{n - \log 9Jn}{K(\Gamma_n, \theta) + n^{-1}} . \quad (\text{B.16})$$

So from (B.4), (B.15), and (B.16)

$$\hat{n}[\bar{\ell}_\theta(\Gamma_n)]^{-1} \leq \mathcal{K}(\Gamma_n, \theta) + n^{-1} + \frac{(\log 9Jn + \lceil \log \gamma_n \rceil) \mathcal{K}(\Gamma_n, \theta) + n^{-1}}{n - \log 9Jn} . \quad (\text{B.17})$$

Define the entropy of the set of strings Γ as

$$\hat{\mathcal{K}}(\Gamma, \theta) = - \sum_{\underline{x} \in \Gamma} v(\underline{x}) \log v(\underline{x}) . \quad (\text{B.18})$$

Since the encoding tree for Γ_n is a subset of the tree for Γ_n^* we have

$$\hat{\mathcal{K}}(\Gamma_n^*, \theta) \geq \hat{\mathcal{K}}(\Gamma_n, \theta) . \quad (\text{B.19})$$

Further

$$\bar{\ell}_\theta(\Gamma_n^*) \geq \bar{\ell}_\theta(\Gamma_n)$$

therefore

$$\begin{aligned} \hat{n}[\bar{\ell}_\theta(\Gamma_n)]^{-1} - \mathcal{K}(\Gamma_n, \theta) &= [\hat{n} - \hat{\mathcal{K}}(\Gamma_n, \theta)] [\bar{\ell}_\theta(\Gamma_n)]^{-1} \\ &\geq [\hat{n} - \hat{\mathcal{K}}(\Gamma_n^*, \theta)] [\bar{\ell}_\theta(\Gamma_n^*)]^{-1} \\ &= r_n^*(\Gamma_n^*, \theta) . \end{aligned} \quad (\text{B.20})$$

From (B.1) we have

$$\lceil \log \gamma_n \rceil \leq \frac{1}{2} S(J-1) \log n + C \quad (\text{B.21})$$

where

$$C \triangleq \log S + \frac{1}{2} S(J-1) \log 9J + 2S \log J + S(J+2) + 1 \quad (\text{B.22})$$

so we may rewrite (B.17) as

$$r_{\hat{n}}(\Gamma_{\hat{n}}, \theta) \leq \mathcal{K}(\Gamma_n, \theta) \left[\frac{\frac{1}{2} S(J-1) \log n + \log n}{n - \log 9Jn} \right] + \frac{C_n}{n - \log 9Jn} \quad (B.23)$$

$$C_n \triangleq (\log 9J + C) \mathcal{K}(\Gamma_n, \theta) + 2 + n^{-1} (\frac{1}{2} S(J-1) \log n + C) \quad (B.24)$$

Now if we follow the steps in (B.7)-(B.11) but with $\Omega' = -\sum v(\underline{x}) \log v(\underline{x})$ in place of Ω we get

$$\Omega' = -\sum_{i=1}^L \sum_{\alpha \in A} i-1 v(\alpha) f(\alpha) \sum_{\beta \in A} v(\beta | s_{i-1}) \log v(\beta | s_{i-1}) \quad (B.25)$$

and since the innermost sum is less than or equal to $\log J$

$$\mathcal{K}(\Gamma_n, \theta) = \Omega' [\bar{L}_\theta(\Gamma_n)]^{-1} \leq \log J \quad (B.26)$$

From (B.21) we have

$$C_n \leq K \triangleq (\log 9J + C)(\log J) + 2 + \frac{1}{2} S(J-1) + C \quad (B.27)$$

Let

$$\hat{K} \triangleq \frac{1}{2} S(J-1) + 1 \quad (B.28)$$

Then

$$\hat{n} \leq n + (\hat{K} - 1) \log n + C \quad (B.29)$$

and

$$\hat{r}_{\hat{n}}(\Gamma_{\hat{n}}^*) \leq \frac{\hat{K} \log J \cdot \log \hat{n} + K}{\hat{n} - \hat{K} \log \hat{n} - C'} \quad (B.30)$$

where

$$C' \triangleq C + \log 9J \quad (B.31)$$

Then if we define

$$K(\hat{n}) = \hat{n}^{-1} [\hat{K} \log \hat{n} + C'] \quad (B.32)$$

we have

$$\hat{f}_{\hat{n}}(\Gamma_{\hat{n}}^*) \leq \hat{n}^{-1} [\hat{K} \log J \cdot \log \hat{n} + K] \left[\sum_{i=0}^{\infty} [K(\hat{n})]^i \right] . \quad (\text{B.33})$$

So if

$$\hat{n} \geq 2(\hat{K}^2 + C')(\log \hat{n})^2 \quad (\text{B.34})$$

we have

$$\hat{f}_{\hat{n}}(\Gamma_{\hat{n}}^*) \leq \hat{n}^{-1} \log J \left[\frac{1}{2} S(J-1) \log \hat{n} + \log \hat{n} \right] + \hat{n}^{-1} (2K+1) . \quad (\text{B.35})$$

as desired.

If Λ is the class of binary memoryless sources we may eliminate the $\log \hat{n}$ term from (B.35). We let Φ_n be defined as $\{\varphi_i: i = 1, \dots, \gamma_n\}$ where

$$\varphi_i \triangleq \begin{cases} 2i^2 \gamma_n^{-2} & , \text{ for } 1 \leq i \leq \frac{1}{2} \gamma_n \\ 1 - 2(\gamma_n - i + 1)^2 \gamma_n^{-2} & , \text{ for } \frac{1}{2} \gamma_n < i \leq \gamma_n \end{cases} \quad (\text{B.36})$$

and $\gamma_n = \lfloor 4\sqrt{n} \rfloor$. Then it is easily shown, ([12] equations (18)-(20)), that for $\theta \leq .5$ if $\theta \in [\varphi_i, \varphi_{i+1}]$ then $\mathcal{K}_r(\theta, \varphi_{i+1}) \leq 2n^{-1}$. We may replace $\log 9Jn$ in (B.6)-(B.23) with $\log \varphi_i^{-1}$, where $\theta \in [\varphi_{i+1}, \varphi_i]$ and (B.23) becomes

$$\hat{f}_{\hat{n}}(\Gamma_{\hat{n}}^*) \leq H(\theta) \left[\frac{\frac{1}{2} \log n + \log \varphi_i^{-1}}{n - \log \varphi_i^{-1}} \right] + \frac{K}{n - \log \varphi_i^{-1}} . \quad (\text{B.37})$$

But for $\theta \in [\varphi_{i-1}, \varphi_i]$ we have

$$\begin{aligned} (\log \varphi_i^{-1})H(\theta) &\leq \log \theta^{-1} H(\theta) \\ &\leq 1.69 \quad . \end{aligned} \tag{B.38}$$

By the same steps as (B.30)-(B.33) this implies

$$\hat{r}_{\hat{n}}(\Gamma_{\hat{n}}^*) \leq \frac{1}{2} \hat{n}^{-1} \log \hat{n} + K_3 \hat{n}^{-1} \tag{B.39}$$

where $K_3 \triangleq 2(K+1.69)$, for

$$\hat{n} \geq (\frac{1}{2} \log \hat{n} + 1.69 + K)(6 + 2 \log \hat{n}).$$

For $\theta > .5$ the same bound holds since $\hat{r}_{\hat{n}}$ is symmetric about $\theta = .5$.

APPENDIX C

PROOFS AND DERIVATIONS FOR CHAPTER 3

Theorem 3.1. If f_n is a code and

$$\mathcal{K}_n(\theta; \varphi) + \mathcal{K}_n(\varphi; \theta) < \xi$$

then

$$|D(f_n; \theta) - D(f_n; \varphi)| \leq \bar{D}(2 \log e)^{-\frac{1}{2}} \xi^{\frac{1}{2}}. \quad (C.1)$$

Proof of Theorem. Let J be a positive integer and define

$$A_m = \{ \underline{x} \in \mathbb{R}^n : n^{-1} d_n(\underline{x}, f_n(\underline{x})) \in (\frac{m-1}{J}, \frac{m}{J}] \}. \quad (C.2)$$

If we define $h(\underline{x}) \triangleq p_\theta(\underline{x}) - p_\varphi(\underline{x})$, $J' \triangleq \lceil n \bar{D} J \rceil$, and $H \triangleq \{ \underline{x} : h(\underline{x}) \leq 0 \}$ we have

$$\begin{aligned} n|D(f_n; \theta) - D(f_n; \varphi)| &= \left| \sum_{m=1}^{J'} \int_{A_m} d_n(\underline{x}, f_n(\underline{x})) h(\underline{x}) d\underline{x} \right| \\ &\leq \left| \sum_{m=1}^{J'} mJ^{-1} \int_{A_m} h(\underline{x}) d\underline{x} + J^{-1} \int_{A_m \cap H} h(\underline{x}) d\underline{x} \right| \end{aligned} \quad (C.3)$$

$$\leq \left| \sum_{m=1}^{J'} mJ^{-1} \int_{A_m} h(\underline{x}) d\underline{x} + J^{-1} \int_H h(\underline{x}) d\underline{x} \right| \quad (C.4)$$

$$\begin{aligned} &\leq \left| \sum_{m=1}^{J'} mJ^{-1} \int_{A_m} h(\underline{x}) d\underline{x} \right| + J^{-1} \\ &= \left| \sum_{m=1}^{J'} mJ^{-1} [\mu_\theta(A_m) - \mu_\varphi(A_m)] \right| + J^{-1}, \end{aligned} \quad (C.5)$$

where $\mu_\theta(B) \triangleq \int_B p_\theta(\underline{x}) d\underline{x}$ for $B \subset \mathbb{R}^n$.

For each of notation let $p_m \triangleq \mu_\theta(A_m)$ and $q_m \triangleq \mu_\varphi(A_m)$ for $m=1,2,\dots,J'$. By definition (177)

$$K_n(\theta;\varphi) \geq \sum_{i=1}^{J'} p_i \log \frac{p_i}{q_i}$$

so the problem reduces to finding

$$\max_{\{p_i, q_i\}} g(p, q) = \max_{\{p_i, q_i\}} \sum_{m=1}^{J'} m J' (p_m - q_m) \quad (C.6)$$

subject to the following constraints:

$$i) \quad K^*(p, q) \triangleq \sum_{i=1}^{J'} p_i \log \frac{p_i}{q_i} + \sum_{i=1}^{J'} q_i \log \frac{q_i}{p_i} \leq \xi$$

$$ii) \quad \sum_i p_i = \sum_i q_i = 1$$

$$iii) \quad p_i \geq 0 \text{ and } q_i \geq 0 \quad \forall i.$$

Since $\max_{\{p_i, q_i\}} g(p, q) = -\min_{\{p_i, q_i\}} g(p, q)$ by symmetry, we may bound (C.5) by

bounding (C.6). From i) we see that

$$p_i = 0 \Leftrightarrow q_i = 0.$$

Let

$$\begin{aligned} \hat{g}(p, q, \lambda, \lambda_p, \lambda_q, \nu_p, \nu_q, \dots) &= g(p, q) + \lambda K^*(p, q) \\ &\quad + \lambda_p \sum_i p_i + \lambda_q \sum_i q_i + \sum_i \nu_{pi} p_i + \sum_i \nu_{qi} q_i \end{aligned} \quad (C.7)$$

where $\lambda, \lambda_p, \lambda_q$ are Lagrange multipliers and ν_{pi} and ν_{qi} are zero if p_i and q_i are positive, and positive if $p_i = q_i = 0$. We must have

$\frac{\partial \hat{g}}{\partial p_i} = 0$ and $\frac{\partial \hat{g}}{\partial q_i} = 0$ at the maximizing $\{p_i, q_i\}$. Let W be the subset of

$\{1, 2, \dots, J'\}$ such that $i \in W$ implies p_i and q_i are positive. Then for $i \in W$ we have

$$\frac{\partial \hat{g}}{\partial p_i} = iJ^{-1} + \lambda \left[\log \frac{p_i}{q_i} + 1 - \frac{q_i}{p_i} \right] + \lambda_p = 0 \quad (C.8)$$

and

$$\frac{\partial \hat{g}}{\partial q_i} = -iJ^{-1} + \lambda \left[\log \frac{q_i}{p_i} + 1 - \frac{p_i}{q_i} \right] + \lambda_q = 0. \quad (C.9)$$

These imply

$$\lambda \left[2 - \frac{q_i}{p_i} - \frac{p_i}{q_i} \right] + \lambda_p + \lambda_q = 0$$

which is equivalent to

$$\frac{p_i}{q_i} + \frac{q_i}{p_i} = \hat{K} \quad (C.10)$$

where $\hat{K}$ is independent of i . Now (C.10) can have only two solutions for

$\frac{p_i}{q_i}$, some η and η^{-1} . But from (C.8)

$$\log \frac{p_i}{q_i} + 1 - \frac{q_i}{p_i} = \frac{-iJ^{-1} - \lambda_p}{\lambda}. \quad (C.11)$$

The left-hand side is a function of $\frac{p_i}{q_i}$, hence it has at most two distinct values, but the right-hand side is different for all i . So we must have only two elements in W ; that is, only two pairs, say (p_a, q_a) and (p_b, q_b) , may be non-zero. Now $p_a = 1 - p_b$ and $q_a = 1 - q_b$ from ii) and from (C.11) we have $\frac{p_a}{q_a} = \frac{q_b}{p_b}$ so

$$p_a - q_a = 2p_a - 1. \quad (C.12)$$

The values of a and b do not affect the relative entropy constraint i), but

$$\begin{aligned} \max_{\{p_i, q_i\}} g(p, q) &= \max_{\{p_i, q_i\}} J^{-1} [a(p_a - q_a) + b(p_b - q_b)] \\ &= \max_{\{p_i\}} J^{-1} (a-b)(2p_a - 1). \end{aligned} \quad (C.13)$$

So for a maximum we must have $a=1$, $b=J'$ or visa versa. The problem is now simplified to finding the maximum of $|2p_a - 1|$ subject to

$$\xi(p_a) \triangleq 2 \left[p_a \log \frac{p_a}{1-p_a} + (1-p_a) \log \frac{1-p_a}{p_a} \right] \leq \xi$$

and $1 > p_a > 0$. Now

$$\xi(p_a) \geq 2 \log e [2p_a - 1]^2 \quad (C.14)$$

so

$$\xi \geq 2 \log e [2p_a - 1]^2$$

or

$$|2p_a - 1| \leq \xi^{\frac{1}{2}} (2 \log e)^{-\frac{1}{2}}. \quad (C.15)$$

So from (C.13) we have

$$\begin{aligned} \max_{\{p_1, q_1\}} \sum_{m=1}^{J'} mJ^{-1} (p_m - q_m) &\leq J' J^{-1} \xi^{\frac{1}{2}} (2 \log e)^{-\frac{1}{2}} \\ &\leq n \bar{D} \xi^{\frac{1}{2}} (2 \log e)^{-\frac{1}{2}} \end{aligned}$$

which substituted into (C.5) gives

$$|D(f_n; \theta) - D(f_n; \varphi)| \leq \bar{D} \xi^{\frac{1}{2}} (2 \log e)^{-\frac{1}{2}} + (nJ)^{-1}. \quad (C.16)$$

Since (C.16) holds for all J

$$|D(f_n; \theta) - D(f_n; \varphi)| \leq \bar{D} \xi^{\frac{1}{2}} (2 \log e)^{-\frac{1}{2}} \quad (C.17)$$

as desired.

Derivation of Eq. (199). Let $\theta = (a_1, \dots, a_k, \sigma^2)$ and $\varphi = (a'_1, \dots, a'_k, \delta^2)$ and assume $\|\theta - \varphi\| \leq \epsilon$. Given initial state $\underline{x}^0 = (x_{-1}, \dots, x_{-k})$

$$\mathcal{K}_n(\theta; \varphi) = n^{-1} \int_{\mathbb{R}^n} p_\theta(\underline{x} | \underline{x}^0) \log \frac{p_\theta(\underline{x} | \underline{x}^0)}{p_\varphi(\underline{x} | \underline{x}^0)} d\underline{x} \quad (\text{C.18})$$

$$\begin{aligned} &= n^{-1} \sum_{j=0}^{n-1} \int_{\mathbb{R}^k} p_\theta(x_{j-1}, \dots, x_{j-k} | \underline{x}^0) \int_{\mathbb{R}} p_\theta(x_j | x_{j-1}, \dots, x_{j-k}) \\ &\quad \cdot \log \frac{p_\theta(x_j | x_{j-1}, \dots, x_{j-k})}{p_\varphi(x_j | x_{j-1}, \dots, x_{j-k})} dx_j, \dots, dx_{j-k}, \end{aligned} \quad (\text{C.19})$$

where

$$p_\theta(\underline{x} | \underline{x}^0) = \prod_{j=0}^{n-1} p_\theta(x_j | x_{j-1}, \dots, x_{j-k}).$$

The inner integral of (C.19) is the entropy of a Gaussian distribution $\eta(m_j, \sigma^2)$ relative to a distribution $\eta(m'_j, \delta^2)$ where

$$m_j = -\sum_{i=1}^k a_i x_{j-i}$$

and

$$m_j' = -\sum_{i=1}^k a_i' x_{j-i} \quad (C.20)$$

We may easily evaluate this integral to get

$$\begin{aligned} K_n(\theta; \varphi) &\leq \frac{1}{2} \log e [(\sigma^2 - \theta^2)^2 (\sigma\theta)^{-2} \\ &\quad + n^{-1} \theta^2 \sum_{j=0}^{n-1} \int_{\mathbb{R}^k} P_\theta(x_{j-1}, \dots, x_{j-k}) (m_j - m_j')^2 dx_{j-1} \dots dx_{j-k}] \\ &\leq \frac{1}{2} e^2 \log e [\sigma_1^{-4} + n^{-1} \sigma_1^{-2} \sum_{j=0}^{n-1} \sum_{i=1}^k \sum_{m=1}^k E_\theta[X_{j-i} X_{j-m}]]. \end{aligned} \quad (C.21)$$

Since

$$X_i = [\gamma^i \underline{x}^0 + \sum_{j=0}^{i-1} \gamma^j Z_{i-j}]_1$$

and $\underline{x}^0 = \underline{0}$ we have

$$\begin{aligned} E_\theta[X_i^2] &= E_\theta \left[\sum_{j=0}^{i-1} \gamma^j Z_{i-j} Z_{i-j}^T [\gamma^j]^T \right]_{1,1} \\ &= \sum_{j=0}^{i-1} [\gamma^j]_{1,1}^2 \sigma^2 \end{aligned} \quad (C.22)$$

Further

$$E[X_i X_j] \leq \max(E[X_i^2], E[X_j^2])$$

and from (C.33) $E_\theta[X_i^2] \geq E_\theta[X_j^2]$ for $i \geq j$, so

$$E_\theta[X_{j-i} X_{j-m}] \leq \sigma_x^2(\theta) \quad (C.23)$$

where

$$\sigma_x^2(\theta) \triangleq \lim_{j \rightarrow \infty} E_\theta[X_j^2] .$$

If we define

$$\sigma_x^2 = \max_{\theta \in \Lambda} \sigma_x^2(\theta)$$

we have from (C.20)

$$J_n(\theta; \varphi) \leq \frac{1}{2} e^2 \log e[\sigma_1^{-4} + \sigma_1^2 k^2 \sigma_x^2] \quad (C.24)$$

as desired.

Derivation of (207). Here we bound $\sum_{i=0}^{n-1} \mu_i^2$. To do this we first bound Ψ_{lm}^i where $\Psi^i = \{\Psi_{lm}^i\}_{l,m=1}^k$. Now Ψ^i has eigenvalue-eigenvector decomposition

$$\Psi^i = V^{-1} E^i V.$$

Let $V = \{v_{lm}\}$ and $V^{-1} = \{u_{lm}\}$ and define

$$\eta = \sup_{\theta \in \Lambda} \max_{l,m} \max \{|u_{lm}|, |v_{lm}|\}.$$

Since Λ is compact and V is invertable for all $\theta \in \Lambda$ (recall $a_k \neq 0$ for $\theta \in \Lambda$) we have $\eta < \infty$. Let λ be the maximum eigenvalue for any Ψ corresponding to a source $\theta \in \Lambda$. To bound Ψ_{lm}^i , the worst case is where Ψ has a single eigenvalue equal to λ . This gives

$$E = \begin{bmatrix} \lambda & 1 & 0 \\ & & 1 \\ 0 & & \lambda \end{bmatrix} .$$

Then we may compute $V^{-1} E^i$ to get [23, pp. 156]

$$\psi_{lm}^i = \begin{cases} \sum_{t=1}^k v_{tm} \sum_{j=1}^m u_{lj} \binom{i}{k-j} \lambda^{i-k+j} ; & i \geq k \\ \sum_{t=1}^k v_{tm} \sum_{j=1}^m u_{lj} \binom{i}{m-j} \lambda^{i-m+j} ; & i < k . \end{cases} \quad (C.25)$$

$$\sum_{t=1}^k v_{tm} \sum_{j=1}^m u_{lj} \binom{i}{m-j} \lambda^{i-m+j} ; \quad i < k . \quad (C.26)$$

So for $i < k$

$$\begin{aligned} \psi_{lm}^i &\leq k \eta^2 \sum_{j=0}^{m-1} \binom{i}{j} \lambda^{i-j} \\ &\leq k \eta^2 \sum_{j=0}^i \binom{i}{j} \lambda^j \\ &= k \eta^2 (1+\lambda)^i , \end{aligned} \quad (C.27)$$

and for $i \geq k$

$$\psi_{lm}^i \leq k \eta^2 \sum_{j=1}^m \binom{i}{k-j} \lambda^{i-k+j} .$$

Since

$$[\psi_{lm}^i]_1 \leq \max_{l,m} |\psi_{lm}^i|_k$$

we have

$$\sum_{i=0}^{n-1} \mu_i^2 \leq k^4 \eta^4 \left[\sum_{i=1}^{k-1} (1+\lambda)^{2i} + \sum_{i=k}^n \left(\sum_{j=1}^m \binom{i}{k-j} \lambda^{i-k+j} \right)^2 \right] . \quad (C.28)$$

Now $\binom{i}{k} \binom{k}{j} \geq \binom{i}{k-j}$ so

$$\begin{aligned}
\sum_{i=k}^n \left(\sum_{j=1}^m \binom{i}{k-j} \lambda^{i-k+j} \right)^2 &\leq \sum_{i=k}^n \binom{i}{k}^2 (\lambda^2)^{i-k} \left(\sum_{j=1}^k \binom{k}{j} \lambda^j \right)^2 \\
&\leq \sum_{i=k}^n \binom{i}{k}^2 (\lambda^2)^{i-k} [(1+\lambda)^k - 1]^2 \\
&\leq \left[\sum_{i=k}^n \binom{i}{k} \lambda^{i-k} \right]^2 [(1+\lambda)^k - 1]^2 \quad (C.29)
\end{aligned}$$

$$= (1-\lambda)^{-2(k+1)} [(1+\lambda)^k - 1]^2, \quad (C.30)$$

where (C.29) follows because $\sum x_i^2 \leq (\sum x_i)^2$ for $x_i \geq 0$. Finally we substitute (C.30) into (C.28) to get

$$\sum_{i=1}^{n-1} \mu_i^2 \leq h^2 \quad (C.31)$$

where

$$h \triangleq k^2 \eta^2 \{ (1-\lambda)^{-2(k+1)} [(1+\lambda)^k - 1]^2 + \sum_{i=1}^{k-1} (1+\lambda)^{2i} \}^{\frac{1}{2}} \quad (C.32)$$

and h is independent of n as desired.

REFERENCES

1. R. Ash, Information Theory, Wiley, New York, 1965.
2. L. Mulder and D. Cohn, "An adaptive variable-to-fixed encoding algorithm for discrete ergodic sources," Proceedings of the 16th Annual Allerton Conference on Communications, Control and Computing, pp. 972-981, 1978.
3. V. Ramamoorthy, "Speech coding based on a composite-Gaussian source model," Linkoping Studies in Science and Technology--Dissertation No. 60, Dept. of Electrical Engineering, Linkoping Univ., Sweden, 1981.
4. T. Ferguson, Mathematical Statistics: A Decision Theoretic Approach, Academic Press, New York, 1967.
5. R. M. Gray, D. L. Neuhoff, and P. C. Shields, "A generalization of Ornstein's d-distance with applications to information theory," Annals of Probability, vol. 3, no. 2, pp. 315-328, 1975.
6. R. D. Smallwood and E. J. Sondik, "The optimal control of partially observed Markov processes over a finite horizon," Operations Research, vol. 21, pp. 1071-1088, 1973.
7. G. Ott, "Compact encoding of stationary Markov processes," IEEE Transactions on Information Theory, vol. IT-13, no. 1, pp. 82-86, 1967.
8. R. J. Fontana, "Universal codes for a class of composite sources," IEEE Transactions on Information Theory, vol. IT-26, pp. 480-482, 1980.
9. S. A. Garde, "Communication of composite sources," Ph. D. Thesis, Univ. of California, Berkeley, CA, 1980.
10. E. N. Gilbert, "Capacity of a burst-noise channel," Bell System Technical Journal, vol. 39, pp. 1253-1266, 1960.
11. L. D. Davisson, R. J. McEliece, M. B. Pursley, and M. S. Wallace, "Efficient universal noiseless source codes," IEEE Transactions on Information Theory, vol. IT-27, pp. 269-279, 1981.
12. R. J. McEliece, M. B. Pursley, and M. S. Wallace, "Universal noiseless data compression for Markov sources," Proceedings of the 1980 Conference on Information Sciences and Systems, pp. 164-167, 1980.
13. A. Blumer, Ph. D. Thesis, in preparation.
14. R. E. Krichevsky and V. K. Trofimov, "The performance of universal encoding," IEEE Transactions on Information Theory, vol. IT-27, no. 2, 1981.
15. V. K. Trofimov, "Redundancy of universal coding of arbitrary Markov

- sources," Problems of Information Transmission, 12, pp. 289-295, 1976.
16. Yu. M. Starkov, "The coding of finite messages on output of a source with unknown statistic," Proceedings of the 5th Conference on Information Theory, Moscow-Gorkii, pp. 147-152, 1972.
 17. J. C. Lawrence, "A new universal coding scheme for the binary memoryless source," IEEE Transactions on Information Theory, vol. IT-23, no. 4, 1977.
 18. B. P. Tunstall, "Synthesis of noiseless compression codes," Ph. D. Thesis, School of Electrical Engineering, Georgia Inst. of Technology, 1967.
 19. F. Jelinek and K. S. Schneider, "On variable-length-to-block coding," IEEE Transactions on Information Theory, vol. IT-18, pp. 765-774, 1972.
 20. J. G. Dunham, "The principle of conservation of entropy," Proceedings of the 18th Annual Allerton Conference on Communications, Control, and Computing, pp. 440-445, 1980.
 21. D. Neuhoff, R. Gray, and L. Davisson, "Fixed rate universal source coding with a fidelity criterion," IEEE Transactions on Information Theory, vol. IT21, pp. 511-523, 1975.
 22. A. Gersho, "Asymptotically optimal block quantization," IEEE Transactions on Information Theory, vol. IT-25, pp. 373-380, 1979.
 23. W. L. Brogan, Modern Control Theory, Quantum Publishers, New York, 1974.
 24. T. Berger, Rate Distortion Theory: A Mathematical Basis for Data Compression, Prentice-Hall, Englewood Cliffs, N.J., 1971.
 25. Y. Linde, A. Buzo, and R. M. Gray, "An algorithm for vector quantizer design," IEEE Transactions on Communications, vol. COM-28, pp. 84-95, 1980.
 26. R. G. Gallager, Information Theory and Reliable Communication, Wiley, New York, 1968.
 27. R. M. Gray and L. D. Davisson, "Quantizer mismatch," IEEE Transactions on Communications, vol. COM-23, pp. 439-442, 1975.
 28. J. A. Bucklew and N. C. Gallagher, "Two-dimensional quantization of circularly symmetric densities," IEEE Transactions on Information Theory, vol. IT-25, pp. 667-671, 1979.
 29. J. Max, "Quantizing for minimum distortion," IRE Transactions on Information Theory, vol. IT-6, pp. 7-12, 1960.
 30. M. B. Pursley and R. M. Gray, "Source coding theorems for stationary, continuous-time stochastic processes," The Annals of Probability, vol. 5, pp. 966-986, 1977.

VITA

Mark Wallace was born in Cleveland, Ohio on October 7, 1956. He received the B. Eng. degree from McGill University, Montreal, Canada in 1978 and the M. S. degree in electrical engineering from the University of Illinois in 1980. He received a National Science Foundation Graduate Fellowship in 1979. He has co-authored the following papers:

"Efficient universal noiseless source codes," (with L. D. Davisson, R. J. McEliece, and M. B. Pursley). IEEE Trans. on Information Theory, vol. IT-27, pp. 269-279, May 1981.

"Universal noiseless data compression for Markov sources," (with R. J. McEliece and M. B. Pursley). Proceedings of the 1980 Conference on Information Sciences and Systems, Princeton Univ., Princeton, NJ, March 1980, pp. 164-167.

"Coding for correlated sources with unknown parameters," (with J. C. Kieffer, M. B. Pursley, and J. Wolfowitz). Proceedings of the 1980 Conference on Information Sciences and Systems, Princeton Univ., Princeton, NJ, March 1980, pp. 168-171.

"Universal noiseless source coding techniques for Markov sources," (with A. C. Blumer, R. J. McEliece, and M. B. Pursley). Presented at the 1981 IEEE International Symposium on Information Theory, Santa Monica, CA, February 1981.

LMEI

-8