

AD-A247 968



2

NAVAL POSTGRADUATE SCHOOL

Monterey, California



DTIC
SELECTE
MAR 26 1992
S B D

THESIS

EFFECTS OF FADING AND DATA MODULATION
ON NONCOHERENT M-SEQUENCE ACQUISITION
SCHEMES

by

Patrick J. Vincent

March 1992

Thesis Advisor

Professor Alex W. Lam

Approved for public release; distribution is unlimited.

92-07629



92 07629

Unclassified

security classification of this page

REPORT DOCUMENTATION PAGE

1a Report Security Classification Unclassified			1b Restrictive Markings		
2a Security Classification Authority			3 Distribution Availability of Report		
2b Declassification Downgrading Schedule			Approved for public release; distribution is unlimited.		
4 Performing Organization Report Number(s)			5 Monitoring Organization Report Number(s)		
6a Name of Performing Organization Naval Postgraduate School		6b Office Symbol (if applicable) 32	7a Name of Monitoring Organization Naval Postgraduate School		
6c Address (city, state, and ZIP code) Monterey, CA 93943-5000			7b Address (city, state, and ZIP code) Monterey, CA 93943-5000		
8a Name of Funding Sponsoring Organization		8b Office Symbol (if applicable)	9 Procurement Instrument Identification Number		
8c Address (city, state, and ZIP code)			10 Source of Funding Numbers		
			Program Element No	Project No	Task No
			Work Unit Accession No		
11 Title (include security classification) EFFECTS OF FADING AND DATA MODULATION ON NONCOHERENT M-SEQUENCE ACQUISITION SCHEMES					
12 Personal Author(s) Patrick J. Vincent					
13a Type of Report Master's Thesis		13b Time Covered From To		14 Date of Report (year, month, day) March 1992	15 Page Count 110
16 Supplementary Notation The views expressed in this thesis are those of the author and do not reflect the official policy or position of the Department of Defense or the U.S. Government.					
17 Cosati Codes			18 Subject Terms (continue on reverse if necessary and identify by block number)		
Field	Group	Subgroup	sequential acquisition, spread spectrum, communications		
19 Abstract (continue on reverse if necessary and identify by block number)					
In direct-sequence spread-spectrum systems, successful communications require phase synchronization of the m-sequence in the incoming signal with a locally generated m-sequence at the receiver. Many acquisition schemes which extract the phase of an incoming m-sequence have been studied, but most of them assume coherent demodulation (which is usually not available during acquisition) and/or independent samples (which introduce a loss in the effective signal to noise ratio (SNR)). This thesis investigates the performance of two m-sequence acquisition schemes in the presence of fading and data modulation. A fixed sample size test and a truncated sequential test are studied without the usual assumptions of coherent demodulation or independent samples. The effects of fading and data modulation on our schemes' probability of false alarm, probability of detection and test length are thoroughly explored. We find that channel fading in effect induces a loss of signal SNR, but the desired power of the tests can be restored by suitable adjustments in the decision processor. We find that the effects of data modulation are less severe, but more problematic to correct.					
20 Distribution Availability of Abstract			21 Abstract Security Classification		
☒ unclassified unlimited ☐ same as report ☐ DTIC users			Unclassified		
22a Name of Responsible Individual Professor Alex W. Lam			22b Telephone (include Area code) (408) 646-3044		22c Office Symbol EC/LA

DD FORM 1473,84 MAR

83 APR edition may be used until exhausted
All other editions are obsolete

security classification of this page

Unclassified

Approved for public release; distribution is unlimited.

Effects of Fading and Data Modulation
on Noncoherent M-Sequence Acquisition Schemes

by

Patrick J. Vincent
Lieutenant, United States Navy
B.S., Polytechnic Institute of New York, 1984

Submitted in partial fulfillment of the
requirements for the degree of

MASTER OF SCIENCE IN ELECTRICAL ENGINEERING

from the

NAVAL POSTGRADUATE SCHOOL
March 1992

Author:

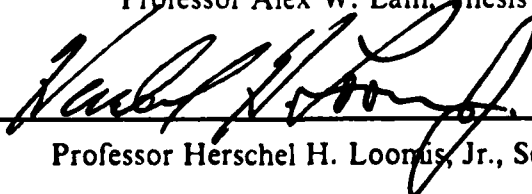


Patrick J. Vincent

Approved by:



Professor Alex W. Lam, Thesis Advisor



Professor Herschel H. Loomis, Jr., Second Reader



Professor Michael A. Morgan, Chairman,
Department of Electrical and Computer Engineering

ABSTRACT

In direct-sequence spread-spectrum systems, successful communications require phase synchronization of the m-sequence in the incoming signal with a locally generated m-sequence at the receiver. Many acquisition schemes which extract the phase of an incoming m-sequence have been studied, but most of them assume coherent demodulation (which is usually not available during acquisition) and/or independent samples (which introduce a loss in the effective signal to noise ratio (SNR)).

This thesis investigates the performance of two m-sequence acquisition schemes in the presence of fading and data modulation. A fixed sample size test and a truncated sequential test are studied without the usual assumptions of coherent demodulation or independent samples. The effects of fading and data modulation on our schemes' probability of false alarm, probability of detection and test length are thoroughly explored. We find that channel fading in effect induces a loss of signal SNR, but the desired power of the tests can be restored by suitable adjustments in the decision processor. We find that the effects of data modulation are less severe, but more problematic to correct.

Accession For	
NTIS GRA&I	☒
DTIC TAB	☐
Unannounced	☐
Justification	
By	
Distribution/	
Availability Codes	
Dist	Avail and/or Special
A-1	

TABLE OF CONTENTS

I. INTRODUCTION	1
A. TIME HOPPING (TH)	2
B. FREQUENCY HOPPING (FH)	2
C. DIRECT SPREADING (DS)	4
1. General	4
2. The Wideband Signal	5
3. Analog Direct Spreading	6
4. Digital Spread Spectrum	6
5. Matched Filters vs. Correlators	7
II. NONCOHERENT SEQUENTIAL ACQUISITION OF M-SEQUENCES	9
A. INTRODUCTION	9
B. NONCOHERENT ACQUISITION SCHEMES	10
1. Fixed-Sample-Size (FSS) Scheme	12
2. Sequential Probability Ratio Test (SPRT)	13
3. Truncated Sequential Probability Ratio Test (TSPRT)	13
C. DESIGN OF DECISION PARAMETERS	14
D. PERFORMANCE COMPARISONS	17
III. PERFORMANCE IN THE FADING CHANNEL	21
A. GENERAL DESCRIPTION OF FADING	21
B. MATHEMATICAL DESCRIPTION OF FADING	23
C. TEST STATISTIC DENSITY FUNCTION	28
D. SIMULATION TECHNIQUES	32
1. Monte Carlo Simulation	34
2. Numerical Integration	38
E. RECEIVER PERFORMANCE WITH FADING	40
1. Expected Results	40
2. Monte Carlo Simulation Results	42
3. Numerical Integration Simulation Results	46
4. Simulation For Various Fading Conditions	49

F. THRESHOLD ADJUSTMENT	52
IV. PERFORMANCE WITH MODULATION	60
A. GENERAL DESCRIPTION OF MODULATION	60
B. TEST STATISTIC DENSITY FUNCTION	62
C. RECEIVER PERFORMANCE WITH MODULATION	65
1. Expected Results	65
2. Numerical Integration Simulation Results: FSS Test	66
3. Numerical Integration Results: TSPRT	71
D. THRESHOLD ADJUSTMENT	81
V. CONCLUSIONS	91
APPENDIX GENERATION OF RICEAN RANDOM VARIATES	94
LIST OF REFERENCES	98
INITIAL DISTRIBUTION LIST	101

LIST OF FIGURES

Figure 1.	Autocorrelation Function of CHIRP	7
Figure 2.	Autocorrelation Function of an M-Sequence	8
Figure 3.	Noncoherent Serial Acquisition System	11
Figure 4.	ASN for Acquisition Schemes ($p_0 = p_1 = 0.5$)	18
Figure 5.	Power Functions ($p_0 = p_1 = 0.5$)	18
Figure 6.	ASN for Acquisition Schemes ($p_0 = 0.3, p_1 = 0.9$)	20
Figure 7.	Power Functions ($p_0 = 0.3, p_1 = 0.9$)	20
Figure 8.	Fast Fade Duration	23
Figure 9.	Envelope Distributions For n Equal-Amplitude Random-Phase Phasors	26
Figure 10.	Ricean distribution for various values of r	27
Figure 11.	TSPRT ASN vs ψ under H_1 (Monte Carlo)	43
Figure 12.	TSPRT PD vs ψ under H_1 (Monte Carlo)	43
Figure 13.	TSPRT $E(\text{ASN})$ vs number of runs: H_1	44
Figure 14.	TSPRT ASN vs ψ under H_0 (Monte Carlo)	45
Figure 15.	TSPRT PFA vs ψ under H_0 (Monte Carlo)	45
Figure 16.	MSE of ASN estimate for TSPRT under H_1	46
Figure 17.	TSPRT ASN vs ψ under H_1 (Numerical Integration)	47
Figure 18.	TSPRT PD vs ψ under H_1 (Numerical Integration)	47
Figure 19.	TSPRT ASN vs ψ under H_0 (Numerical Integration)	48
Figure 20.	TSPRT PFA vs ψ under H_0 (Numerical Integration)	48
Figure 21.	TSPRT PD vs ψ (Numerical Integration)	49
Figure 22.	TSPRT Power vs $ j + \gamma $	50
Figure 23.	TSPRT ASN vs $ j + \gamma $	50
Figure 24.	TSPRT ASN vs $ j + \gamma $ (various r)	51
Figure 25.	TSPRT PD vs $ j + \gamma $ (various r)	51
Figure 26.	FSS PD vs $ j + \gamma $ (various r)	52
Figure 27.	Fade Depth For Typical HF Link	55
Figure 28.	TSPRT Thresholds Adjusted For Fading	56
Figure 29.	Typical Trajectories For $r = 10$	57
Figure 30.	TSPRT PD Fading Performance	58
Figure 31.	TSPRT ASN Fading Performance	59

Figure 32. Expected Value of Z vs n	66
Figure 33. FSS PD vs Modulation Chip Boundary: H_1	68
Figure 34. FSS PFA vs Modulation Chip Boundary: H_0	68
Figure 35. FSS PD vs Modulation Chip Boundary	69
Figure 36. FSS Power vs $ j + \gamma $: Modulation	69
Figure 37. Calculated FSS Termination Point	71
Figure 38. Actual FSS Probability of Miss	72
Figure 39. Actual FSS Probability of False Alarm	73
Figure 40. Actual FSS Miss Probability With Modulation	74
Figure 41. FSS PM vs α for Fixed β With Modulation	75
Figure 42. FSS PM vs β for Fixed α With Modulation	75
Figure 43. Actual FSS False Alarm Probability With Modulation	76
Figure 44. TSPRT ASN vs Modulation Chip Boundary: H_1	77
Figure 45. TSPRT PD vs Modulation Chip Boundary	77
Figure 46. TSPRT ASN vs Modulation Chip Boundary: H_0	79
Figure 47. TSPRT PFA vs Modulation Chip Boundary	79
Figure 48. TSPRT Power vs $ j + \gamma $ With Modulation	80
Figure 49. TSPRT ASN vs $ j + \gamma $ With Modulation	80
Figure 50. TSPRT PD With Modulation	82
Figure 51. TSPRT ASN With Modulation	82
Figure 52. Density Function With Modulation: H_1	84
Figure 53. Area Under pdf From $z = 0$ to $z = 10.077$	85
Figure 54. Adjusted FSS Test Length	86
Figure 55. FSS Test Length vs β : New Thresholds	87
Figure 56. Actual FSS Miss Probability: New Thresholds	88
Figure 57. Actual FSS Miss Probability vs α : New Thresholds	89
Figure 58. Actual FSS Miss Probability vs β : New Thresholds	89
Figure 59. Actual FSS False Alarms Probability: New Thresholds	90

I. INTRODUCTION

This thesis explores the effects of fading and data modulation on sequential pseudonoise (PN) acquisition schemes. We begin by providing a brief introductory overview of spread spectrum communication systems. The purpose of this introductory chapter is twofold. The inclusion of this material enhances the degree to which this thesis is self-contained. Additionally, the notation introduced in this chapter will be used in subsequent sections.

A complete presentation of the theory and applications of spread spectrum communication systems would fill several volumes. The intent here is neither to be comprehensive in scope nor specific in detail. We present, in this chapter, only a thumbnail sketch of the subject, including only those details necessary to establish a foundation for the following chapters.

Spread spectrum signals allow us to [Ref. 1 : pp. 800-801]:

1. overcome or minimize the effects of interference, whether it be intentional (e.g., jamming), unintentional (e.g., interference from other channel users) or self-induced (e.g., multipath).
2. hide a signal by "burying" it below the background noise level, thus concealing the transmission from unintended listeners and eavesdroppers.
3. achieve message privacy.

In a conventional communication system employing amplitude modulation (AM) or phase modulation (PM), the communication signal bandwidth is approximately equal to twice the rate of data transmission. In spread spectrum (SS) systems the bandwidth of the communication signal is much greater than the data rate. In all SS systems, randomness is used to make interception and copying more formidable.

Prior to 1980, the field of spread spectrum communications fell almost entirely within the military domain. Techniques were developed to allow communications which were not detectable by adversaries and to combat the disruptive effects of jamming by an opponent. In 1983 the Federal Communication Commission (FCC) opened, for the first time, three frequency bands for commercial spread spectrum use [Ref. 2]. Cellular and indoor commercial networks using spread spectrum techniques are now becoming a reality.

In the following we will discuss several typical spread spectrum techniques.

A. TIME HOPPING (TH)

In a time hopping scheme, we transmit the communication signal at randomly chosen times for random burst durations. If η is the fraction of time the transmitter is active ($\eta < 1$) and r_c is the steady data rate from the user, then r_b , the transmission rate when the transmitter is activated, is given by

$$r_b = \frac{r_c}{\eta} \quad (1.1)$$

Note that the bandwidth of the signal that gets transmitted is proportional to r_b (which may be much larger than r_c) so we have distributed, or *spread*, the signal over a wider frequency band.

The primary advantage of TH is an improvement in the Signal to Jamming Ratio (SJR). Let S and J be the average signal power and average jamming power, respectively, at the receiver input. Then the SJR for a *non-SS* system is given by

$$SJR = \frac{S}{J} \quad (1.2)$$

For a TH system, let $\hat{S}$ be the input power at the receiver when a burst transmission occurs. Then $S = \hat{S}\eta$ and the SJR during burst reception is

$$SJR = \frac{\hat{S}}{J} = \frac{S}{\eta J} = \left(\frac{1}{\eta} \right) \frac{S}{J} \quad (1.3)$$

Thus we see that the SJR in a TH system is better than that for a non-SS system. The amplification factor ($= \frac{1}{\eta}$) is called the *processing gain* (PG).

The improvement in SJR in a time hopping system is countered by several disadvantages:

1. a data buffer must be employed in both the transmitter and receiver
2. the TH transmitter must radiate at higher power levels
3. the TH receiver must resynchronize on each new burst

B. FREQUENCY HOPPING (FH)

As previously noted, in all SS systems randomness is used to make the communications system more immune to signal detection and jamming. In a system employing frequency hopping, the transmitter carrier frequency changes value (hops) as a function of time. If we let t_h be the amount of time the carrier remains at each of the hopping

frequencies, then the bandwidth of the carrier is always constant and equal to $2/t_h = b$ Hz. If B is the total range of frequencies we are assigned to hop over, then the number of distinct carriers used is equal to B/b .

A FH receiver *despreads* the received signal by multiplying it by a locally generated hopping signal. Note that the hop sequence of the receiver's local oscillator must match the hop sequence of the transmitter in order to recover the signal.

As with time hopping, a major advantage of FH systems is an improved immunity to jamming. In a FH scheme, the jamming power is now spread over a bandwidth of B Hz. The signal bandwidth is approximately equal to b Hz. Although the entire signal power is present after dehoppping, the amount of jamming power present after dehoppping is $J(b/B)$. Thus, in a FH system,

$$SJR = \frac{S}{\left(\frac{b}{B}\right)^J} = \left(\frac{B}{b}\right) \left(\frac{S}{J}\right) \quad (1.4)$$

The processing gain is thus given by B/b .

In many military applications we would like to deny an adversary the knowledge that we are transmitting a signal. That is to say, besides making it difficult for an adversary to "break" our code once he intercepts our signal, we would like to make it difficult for him to detect that there exists a transmitted signal at all. (Obviously a signal that is not detected can not be exploited.) To achieve a low probability of detection in a communications scheme, the signal power level must be reduced below the noise power level. In other words, we need to "bury" the signal in the noise (or hide it behind another more powerful already-present signal) so that the adversary will "see" only noise. Such a scheme is termed a Low Probability of Detection (LPD) scheme.

A major disadvantage of the frequency hopping approach lies in the fact that the signal can not be reduced below the noise level; in fact, the signal must be maintained significantly above the noise level. As a result, frequency hopping schemes are used to provide anti-jamming, *not* low probability of signal detection.

In some communication environments, multiple transmission paths may arise between the transmitter and receiver. This phenomenon, termed *multipath*, often degrades communications since the signals arriving at the receiver via secondary paths destructively interfere with the signal on the primary link. In a FH receiver only the energy at a particular frequency is passed through the first stage filter at any interval of time. If the multipath signals arrive at the receiver with a sufficient time delay, then the

receiver may have "gone on to the next hop" by that time. In such instances the multipath signals are effectively ignored. This is called the *multipath immunity of frequency hopping*.

Finally, the utility of the frequency hopping method is limited by the speed with which the electrical components can alter frequency without generating excessive noise. As of 1990, frequency oscillators were able to vary frequencies at rates of 1 million hops per second over a 20 MHz bandwidth [Ref. 3].

C. DIRECT SPREADING (DS)

1. General

In a direct spreading scheme, the bandwidth spreading is accomplished on a bit-by-bit basis. We multiply the data signal $v_d(t)$ by a wideband voltage signal $v_s(t)$. Since multiplication in the time domain corresponds to convolution in the frequency domain, the resulting signal's bandwidth is *spread*. We can then use a matched filter or correlator in the receiver to detect the transmitted signal.

If a matched filter is used in the receiver, an expression can be obtained for the system's processing gain. Let E represent the energy per bit in the input signal to the receiver, let T represent the bit duration, let W be the bandwidth of the spread signal and let $N_o/2$ be the two-sided power spectral density of the noise. Then the input signal power is given by

$$S_i = \frac{E}{T} \quad (1.5)$$

The input noise power is given by

$$N_i = \left(\frac{N_o}{2} \right) (2W) = N_o W \quad (1.6)$$

and the input SNR is

$$SNR_{in} = \frac{S_i}{N_i} = \frac{E}{N_o} \left(\frac{1}{TW} \right) \quad (1.7)$$

A result from matched filter theory states that the output signal to noise ratio is given by

$$SNR_o = \frac{2E}{N_o} \quad (1.8)$$

Combining (1.7) and (1.8) yields

$$SNR_o = SNR_{in}(2TW) \quad (1.9)$$

Thus $2TW$ is the processing gain for direct spreading using a matched filter. To increase the processing gain we usually increase W . Note that there is no theoretical limit to the degree of improvement achievable.

It is a classic myth to believe that direct spreading improves the output SNR in receivers where thermal noise limits operation. This is not the case. If the receiver's performance is limited by thermal noise, then employing spread spectrum will not help. Looking at (1.9), we know that utilizing spread spectrum increases W and hence, we may be tempted to conclude, the output signal to noise ratio, SNR_o , must increase. However, increasing W increases the noise bandwidth of the receiver, and hence increases the input noise power since

$$N_i = kTW \quad (1.10)$$

where k is Boltzmann's constant and T is the effective noise temperature. The increase in N_i causes a decrease in SNR_{in} . The decrease in SNR_{in} counterbalances the increase in the $(2TW)$ term.

In jamming limited receivers $N_i = J = \text{constant}$. In this case we do get improvement by employing direct spreading.

Direct spreading affords us the ability to manage the power in the transmitted signal. We can distribute the power over any frequency interval without penalty. We can distribute the power over a wide frequency band so that the signal power density becomes less than the noise power density. In this way we can "bury" the signal in the noise and achieve covert communications.

To summarize, in the direct spreading technique, we represent each bit of duration T with a wideband signal $v_x(t)$ of duration T and bandwidth W .

2. The Wideband Signal

The value of the direct spreading scheme depends critically on the choice of $v_x(t)$, the wideband signal. The signal must have a good autocorrelation function (ACF) since we want a good matched filter output. A "good" ACF is one with a narrow, high magnitude main lobe and small side lobes. Such an ACF allows the use of a follow-on threshold device with minimum decision error. Additionally, $v_x(t)$ should have a

buildable matched filter and should have a near-white (flat) spectrum to allow for covert communications.

3. Analog Direct Spreading

In analog direct spreading, the wideband signal $v_x(t)$ is analog. The beauty of analog direct spreading is that the communication signal can be directly transmitted; there is no need for a separate modulating carrier. Perhaps the most commonly used analog spreading signal is the CHIRP signal. In this case $v_x(t)$ varies linearly in frequency from some initial frequency, f_1 , to some final frequency, f_2 , over the interval T . The formula for the spreading signal is

$$v_x(t) = A \cos \left[2\pi \left(f_1 t + \left\{ \frac{f_2 - f_1}{2T} \right\} t^2 \right) \right] \quad 0 < t < T \quad (1.11)$$

To send the bit "1" we may start at f_1 and terminate at f_2 (called UP-CHIRP if $f_2 > f_1$). To send the bit "0" we start at f_2 and terminate at f_1 (called DOWN-CHIRP if $f_2 > f_1$). Using Fresnel Integrals it can be shown that the bandwidth of the spread signal is approximately equal to $f_2 - f_1$. Additionally, the spectrum is approximately constant in the interval (f_1, f_2) and zero outside this interval.

The ACF of the analog CHIRP signal is quite good. Figure 1 on page 7 shows one period of the autocorrelation function of $v_x(t)$ given by (1.11) for the case of $2TW = 100$, $T = 1$ sec and $f_1 = 50$ Hz. The matched filter for CHIRPED signals is the surface acoustic wave device (SAW) [Ref. 4].

4. Digital Spread Spectrum

In digital direct spreading, the wideband signal $v_x(t)$ is digital. Specifically, for $v_x(t)$ we choose a binary sequence. Instead of transmitting a data bit (say a "1") we transmit a binary sequence (say, "1100011"). Each bit in the binary sequence that comprises $v_x(t)$ is called a *chip*. As of 1990, the maximum achievable chip rate was 50 million chips per second [Ref. 3].

The digital spreading signal is created by using the output of an n -stage shift register with feedback connections. An n -stage shift register can take on, at most, $2^n - 1$ different states (depending on the feedback connections), so the maximum length sequence attainable from the shift register is also equal to $2^n - 1$. These maximum length sequences, called *m-sequences*, have the best ACF attainable for sequences of their length. The theory of feedback shift register sequences is a well-developed subject, and the interested reader is referred to [Ref. 5 : Chapter 7].

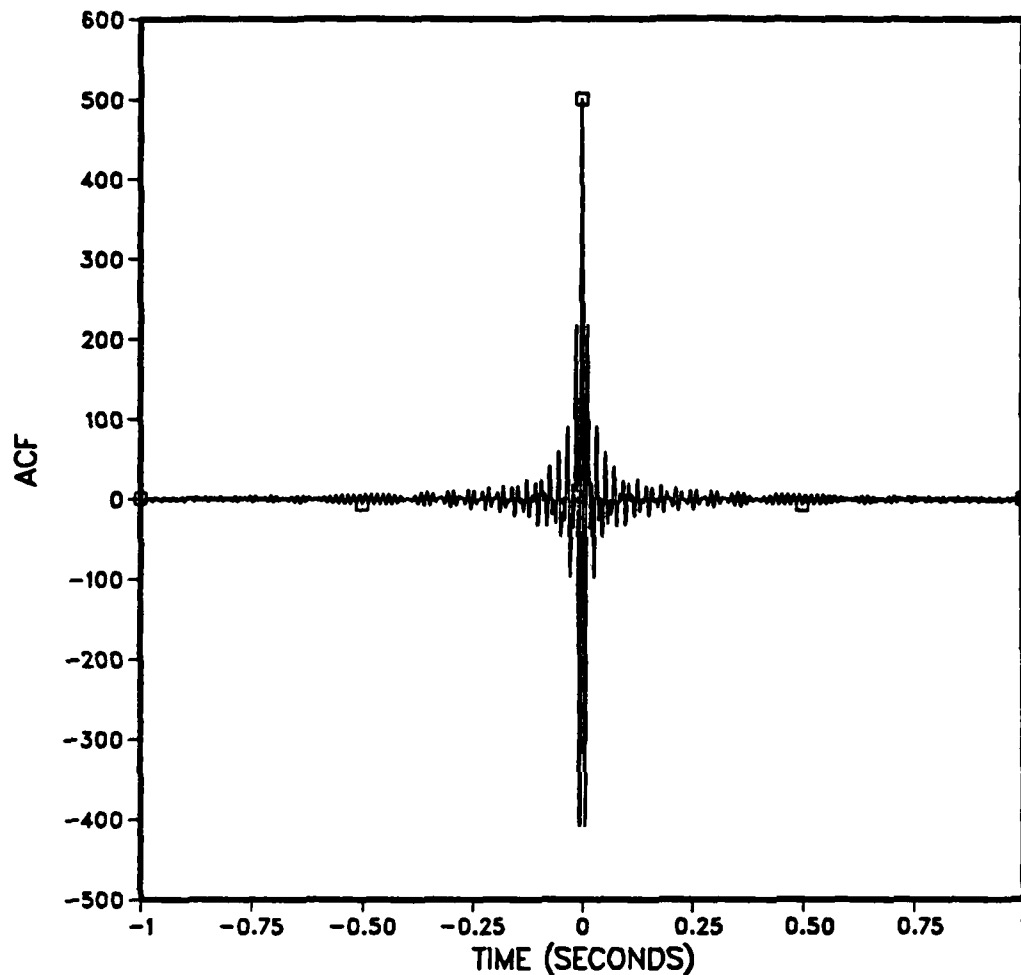


Figure 1. Autocorrelation Function of CHIRP

The ACFs of all m-sequences have the same form. The peak value is proportional to $2^n - 1$. The sidelobe value is proportional to -1. A normalized m-sequence ACF is shown in Figure 2 on page 8.

Sending data using a spreading sequence (e.g., an m-sequence) is called *Direct Sequence Spread Spectrum* (DSSS). To send a "1" we can send one period of the m-sequence. To send a "-1" we can send one period of the inverted m-sequence. For the remainder of this thesis we limit our consideration to DSSS systems.

5. Matched Filters vs. Correlators

Matched filters for binary sequences usually consist of some variant of a tapped delay line (realized as a tapped shift register). The received signal is clocked into a shift

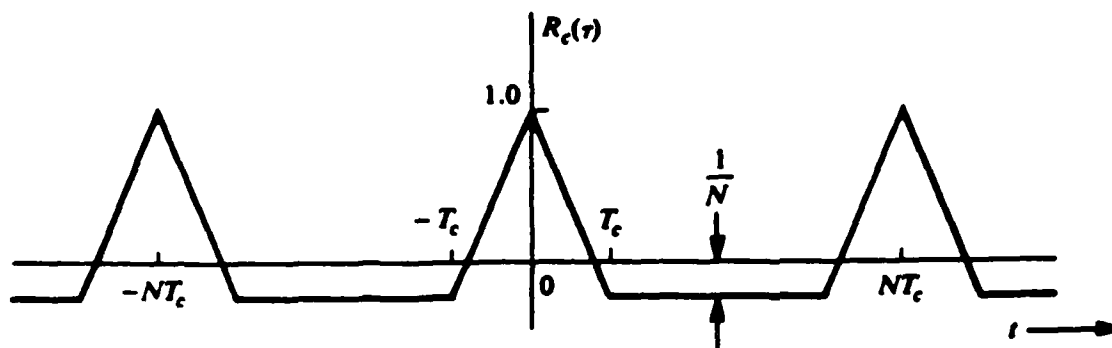


Figure 2. Autocorrelation Function of an M-Sequence

Source: Couch, L.W. II, *Digital and Analog Communication Systems*, Macmillan, 1990

register (of length equal to one m-sequence period) in the receiver. The contents of each individual register are tapped, inverted as appropriate and summed to form a discrete approximation to the ACF.

Unfortunately digital matched filters have many deficiencies. A large number of taps are required for desired values of processing gain, making the hardware implementation (on an integrated circuit (IC) chip) difficult if not impossible. Additionally, the chips are "analog" to the extent that they are buried in noise, and so digital hardware does not readily lend itself. However we can still use the ACF property of m-sequences and replace the matched filter with a three-port correlator. Unfortunately, but unavoidably, data recovery with a correlator requires a receiver-based locally generated matching m-sequence, phase synchronized with the received signal's m-sequence. Of course generating the local m-sequence is of no difficulty, but efficiently phase-locking it with the received signal has been, and continues to be, a major research concern.

II. NONCOHERENT SEQUENTIAL ACQUISITION OF M-SEQUENCES

In this chapter we consider schemes used to acquire m-sequences in direct sequence spread spectrum (DSSS) communication systems.¹

A. INTRODUCTION

In order to recover the data from a DSSS communications signal, the signal must be *despread*, i.e., the data must be extracted by removing the m-sequence. Because matched filters are impractical in this application, the autocorrelation property of m-sequences must be exploited. One way to remove the m-sequence (and preserve the data) is to multiply the incoming signal by a locally generated identical m-sequence which is phase synchronized with the incoming signal's m-sequence. Thus, a primary function of the receiver is determination of the phase of the m-sequence in the incoming signal. The phase is determined by correlating the incoming signal with the m-sequence generated in the receiver. The local sequence's phase is adjusted until synchronization is achieved. Aligning the two m-sequences is accomplished in two stages. The first stage, called *acquisition*, brings the m-sequences into coarse alignment. The following stage, called *tracking*, brings the sequences into precise alignment.

In this thesis we consider a *noncoherent* acquisition scheme. (*Noncoherent*, in this case, refers to the modulating carrier, not the m-sequence.) Although coherent schemes are simpler and more widely studied, they are not practical because m-sequence acquisition is usually performed prior to recovering the modulating carrier.

To *acquire* the incoming signal, we must examine, using some methodology, the various phases that the m-sequence can assume. One methodology would be to simultaneously examine all possible m-sequence positions by utilizing, in the receiver, a separate correlator for each potential phase. Such parallel schemes are characterized by short acquisition times but complex hardware requirements. On the other hand, we can examine each possible m-sequence phase one by one to determine if the local sequence is aligned with the incoming sequence. Such serial schemes are characterized by long

¹ This chapter, which presents the model and framework to be used in the subsequent chapters, is adapted from [Ref. 6], a paper that served as the original motivation for this research. This chapter, at various times, uses the words, assumptions, notations, equations, methods of organization and conclusions originally propounded in [Ref. 6]. This material is adapted and used with the permission of the authors.

acquisition times but simplified hardware requirements. In this thesis we consider a *serial* acquisition scheme.

B. NONCOHERENT ACQUISITION SCHEMES

A block diagram of the receiver's acquisition system is shown in Figure 3 on page 11. The receiver's input signal is comprised of the transmitted signal and corrupting additive white Gaussian noise (AWGN) with two-sided power spectral density $N_0/2$. The input signal at the receiver, assuming no modulation or doppler shifting, is

$$r(t) = A_o a(t + i\Delta T_c) \cos(\omega_o t + \theta) + n(t) \quad (2.1)$$

where A_o is the signal amplitude, $a(t)$ is the m-sequence signal waveform with phase $i\Delta T_c$ (i taken to be an integer without loss of generality), T_c is the chip interval, Δ is the amount that the phase of the local m-sequence is altered during the acquisition process, ω_o and θ are the frequency and phase of the carrier, and $n(t)$ is the AWGN. The locally generated m-sequence waveform is described as $a(t + (j + \gamma)\Delta T_c)$ where j is an integer and $|\gamma| < 0.5$. The incoming communication signal is multiplied by the locally generated signal and then noncoherently demodulated. The result Y_n is checked by the decision processor to determine if the local and incoming m-sequences are lined up to within $\Delta T_c/2$. If they are not aligned, the local m-sequence phase is altered by ΔT_c and the process repeats. If they are aligned, then j must equal i and the tracking circuit is then employed to reduce γ to zero.

Using the notation in Figure 3, and neglecting double frequency terms, we have

$$X_{i,n} = \int_0^{nT_c} r(t) a(t + (j + \gamma)\Delta T_c) \cos(\omega_o t) dt = \frac{A_o}{2} T_c S_n \cos \theta + N_{i,n} \quad (2.2)$$

$$X_{q,n} = \int_0^{nT_c} r(t) a(t + (j + \gamma)\Delta T_c) \sin(\omega_o t) dt = \frac{A_o}{2} T_c S_n \sin \theta + N_{q,n} \quad (2.3)$$

where

$$N_{i,n} = \int_0^{nT_c} n(t) a(t + (j + \gamma)\Delta T_c) \cos(\omega_o t) dt \quad (2.4)$$

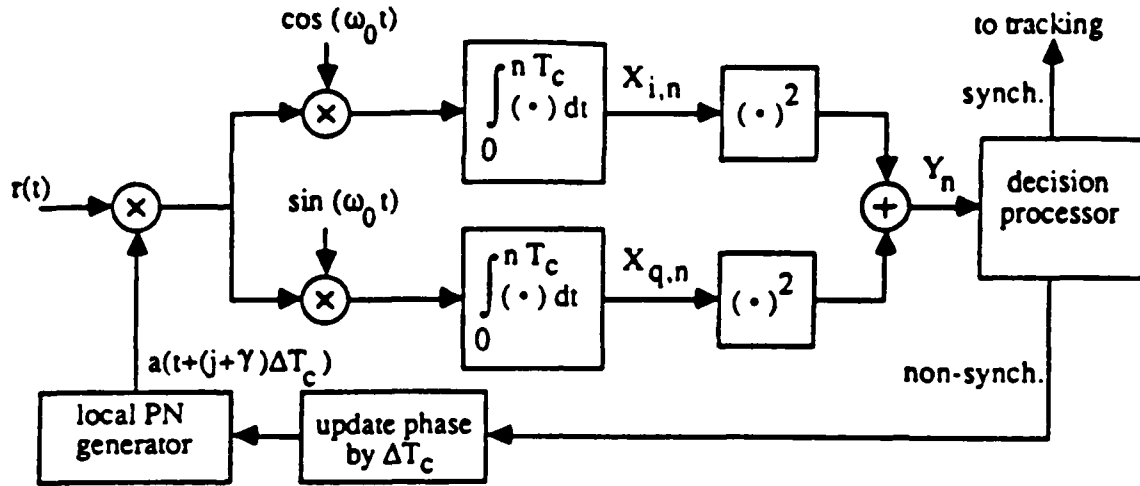


Figure 3. Noncoherent Serial Acquisition System

$$X_{q,n} = \int_0^{nT_c} r(t) a(t + (j + \gamma)\Delta T_c) \sin(\omega_0 t) dt \quad (2.5)$$

are independent, zero mean Gaussian random variables with variance $\sigma_n^2 = nT_c N_s / 4$ and where

$$S_n = \frac{1}{T_c} \int_0^{nT_c} a(t + i\Delta T_c) a(t + (j + \gamma)\Delta T_c) dt \quad (2.6)$$

For a fixed n , $X_{i,n}$ and $X_{q,n}$, the inphase and quadrature components, are independent Gaussian random variables, each with variance σ_n^2 and with means $(A_s/2)T_c S_n \cos \theta$ and $(A_s/2)T_c S_n \sin \theta$ respectively. Note that $X_{i,n}$ and $X_{i,m}$ (and similarly $X_{q,n}$ and $X_{q,m}$) are not independent for $n \neq m$. The test statistic for determining alignment is

$$Y_n = X_{i,n}^2 + X_{q,n}^2 \quad (2.7)$$

The random variable Y_n has a probability density function that is noncentral Chi-squared with two degrees of freedom [Ref. 1 : pp. 25-29], given by

$$f_{Y_n}(y_n) = \frac{1}{2\sigma_n^2} e^{-(y_n + \lambda_n) 2\sigma_n^2} I_0\left(\frac{\sqrt{\lambda_n y_n}}{\sigma_n^2}\right), \quad y_n \geq 0 \quad (2.8)$$

where $\lambda_n = [(A_o/2)T_c S_n \cos \theta]^2 + [(A_o/2)T_c S_n \sin \theta]^2 = (A_o^2/4)T_c^2 S_n^2$ and where $I_0(\cdot)$ is the modified Bessel function of order zero.

Using Y_n , the decision statistic, the acquisition scheme determines if the two m-sequence phases are aligned within $\Delta T_c/2$, i.e., $j = i = 0$ (we let $i = 0$ to simplify the notation), or if they differ by one or more chips, i.e., $|(j + \gamma) - i| \Delta T_c \geq T_c$. In other words, the decision processor decides between the following hypotheses:

$$\begin{aligned} H_0 \text{ (non-alignment)} : |j + \gamma| \geq \frac{1}{\Delta} \quad \text{and} \quad |\gamma| < \frac{1}{2} \\ H_1 \text{ (alignment)} : j = i = 0 \quad \text{and} \quad |\gamma| < \frac{1}{2} \end{aligned} \quad (2.9)$$

Note that the m-sequence phases may be such that our situation falls between H_0 and H_1 . In this case $\frac{1}{2} < |j + \gamma| < \frac{1}{\Delta}$.

The parameter λ_n takes on different values under H_0 and H_1 . We designate $\lambda_{n,0}$ and $\lambda_{n,1}$ as the worst-case values of λ_n under H_0 and H_1 respectively. Note that $\lambda_{n,1} > \lambda_{n,0} > 0$. Using these worst-case parameter values, the likelihood ratio for our test becomes

$$\Lambda_n(y_n) = \frac{f_{Y_n}(y_n | H_1)}{f_{Y_n}(y_n | H_0)} = \exp\left(\frac{\lambda_{n,0} - \lambda_{n,1}}{2\sigma_n^2}\right) \frac{I_0(\sqrt{(\lambda_{n,1}/\sigma_n^2)(y_n/\sigma_n^2)})}{I_0(\sqrt{(\lambda_{n,0}/\sigma_n^2)(y_n/\sigma_n^2)})} \quad (2.10)$$

We now describe the three decision schemes that we will use to resolve between the two hypotheses.

1. Fixed-Sample-Size (FSS) Scheme

In a fixed-sample-size (FSS) decision scheme, the length of integration is fixed a priori and a decision is made based on the resulting test statistic, Y_n . If the integration length is from $t = 0$ to $t = MT_c$, the FSS test can be described by

$$\text{FSS:} \quad \Lambda_M(y_M) \begin{cases} \geq \tau & \text{say } H_1 \\ < \tau & \text{say } H_0 \end{cases} \quad (2.11)$$

where τ is the threshold. Since the likelihood ratio is a monotonically increasing function of the variable y_n [Ref. 6], the FSS test is equivalent to the following test:

$$\text{ESS: } y_M \begin{cases} \geq \tau' = \Lambda_M^{-1}(\tau) & \text{say } H_1 \\ \leq \tau' = \Lambda_M^{-1}(\tau) & \text{say } H_0 \end{cases} \quad (2.12)$$

where $\Lambda_n^{-1}(\cdot)$ is the inverse function of $\Lambda_n(\cdot)$.

2. Sequential Probability Ratio Test (SPRT)

The sequential probability ratio test (SPRT) consists of testing the likelihood ratio against two thresholds for $n = 1, 2, 3, \dots$ until one of the thresholds is exceeded. The SPRT is optimal for independent Gaussian statistics. The length of integration increases by one chip each time n increases by one. The threshold that is reached first (upper or lower) determines the hypothesis that is accepted (H_1 or H_0). The test is described by

$$\text{SPRT: } \Lambda_n(y_n) \begin{cases} \geq A & \text{say } H_1 \\ \leq B & \text{say } H_0 \\ \text{otherwise, continue to next } n \end{cases} \quad (2.13)$$

Since Y_n and Y_{n+1} are not independent random variables we can not write $\Lambda_n(y_n)$ as a product of independent random variables. However, we can place bounds on the Type I error (false alarm probability) and the Type II error (miss probability) by using Wald's inequalities [Ref. 7]:

$$A \leq \frac{1 - \beta}{\alpha}, \quad \text{and} \quad B \geq \frac{\beta}{1 - \alpha} \quad (2.14)$$

where α and β are the resulting false-alarm and miss probabilities, respectively. If the average test length is large, the inequalities can be approximated by equalities. To make the real-time SPRT implementation more practical, we can rewrite the test as

$$\text{SPRT: } y_M \begin{cases} \geq A(n) \equiv \Lambda_n^{-1}(A) & \text{say } H_1 \\ \leq B(n) \equiv \Lambda_n^{-1}(B) & \text{say } H_0 \\ \text{otherwise, continue} \end{cases} \quad (2.15)$$

where the thresholds, which are functions of n , can be pre-computed.

3. Truncated Sequential Probability Ratio Test (TSPRT)

When using the SPRT it is possible that, for certain m-sequence phase disparities, the decision scheme may take an excessively long time to decide between H_0 and H_1 . Such an incident is especially likely when the difference between the phases of the two m-sequences is such that it corresponds to a state between H_1 and H_0 . To avoid a

very long test we impose an upper bound on the test length by truncating the SPRT (hence the name TSPRT) at a predetermined test duration, and conducting an FSS test at that point. The test is described by

$$\begin{aligned} \text{TSPRT:} \quad & \text{if } n \leq \hat{n}, \quad \Lambda_n(y_n) \begin{cases} \geq \hat{A} & \text{say } H_1 \\ \leq \hat{B} & \text{say } H_0 \\ \text{otherwise, continue} \end{cases} \\ & \text{if } n = \hat{n}, \quad \Lambda_n(y_n) = \Lambda_{\hat{n}}(y_{\hat{n}}) \begin{cases} \geq \hat{\tau} & \text{say } H_1 \\ < \hat{\tau} & \text{say } H_0 \end{cases} \end{aligned} \quad (2.16)$$

The test is truncated at $n = \hat{n}$ (provided it has not already terminated). Again we use the monotonicity of $\Lambda_n(\cdot)$ and rewrite the test as

$$\begin{aligned} \text{TSPRT:} \quad & \text{if } n \leq \hat{n}, \quad y_n \begin{cases} \geq \hat{A}(n) \equiv \Lambda_n^{-1}(\hat{A}) & \text{say } H_1 \\ \leq \hat{B}(n) \equiv \Lambda_n^{-1}(\hat{B}) & \text{say } H_0 \\ \text{otherwise, continue} \end{cases} \\ & \text{if } n = \hat{n}, \quad y_{\hat{n}} \begin{cases} \geq \Lambda_{\hat{n}}^{-1}(\hat{\tau}) & \text{say } H_1 \\ < \Lambda_{\hat{n}}^{-1}(\hat{\tau}) & \text{say } H_0 \end{cases} \end{aligned} \quad (2.17)$$

A key parameter in determining the utility of an acquisition scheme is the time that the test takes to *acquire* the incoming m-sequence. Generally, for a given false-alarm and miss probability, the test with the smallest sample size (n) will reach acquisition the soonest. The initial phase disparity between the incoming and local m-sequences can be regarded as a uniform random variable distributed over the sequence period. The hypothesis H_0 is, therefore, the most probable and so the expected test length under H_0 contributes most significantly to the value of the test acquisition time. Hence, it is best to minimize the value of the expected sample size under H_0 while keeping the expected sample size to within reasonable values when H_0 does not apply.

C. DESIGN OF DECISION PARAMETERS

In this section we determine the values of the thresholds to be used in the three acquisition schemes.

Denote the chips of the m-sequence by $c_k, k = \dots, 0, 1, 2, \dots$, where $c_k = \pm 1$, and let the period be $N = 2^m - 1$, so that $c_{k+N} = c_k$ for all k . The m-sequence spreading waveform can then be written as

$$a(t) = \sum_{k=-\infty}^{\infty} c_k p_{T_c}(t - kT_c) \quad (2.18)$$

where $p_{T_c}(t)$ is a rectangular pulse of width T_c and height 1, i.e., $p_{T_c}(t) = 1$ for $0 \leq t < T_c$ and it is zero elsewhere. Letting $i=0$, we obtain expressions for S_n by combining (6) and (18):

$$\begin{aligned} S_n &= \begin{cases} \sum_{k=0}^{n-1} c_k [(1 - |\gamma| \Delta) c_k + (|\gamma| \Delta) c_{k+\text{sgn}(\gamma)}] & \text{under } H_1 \\ \sum_{k=0}^{n-1} c_k [(1 - \delta) c_{k+l} + \delta c_{k+1+l}] & \text{under } H_0 \end{cases} \\ &= \begin{cases} = n(1 - |\gamma| \Delta) + (|\gamma| \Delta) \sum_{k=0}^{n-1} c_k c_{k+\text{sgn}(\gamma)} & \equiv S_{n,1} \text{ under } H_1 \\ = (1 - \delta) \sum_{k=0}^{n-1} c_k c_{k+l} + \delta \sum_{k=0}^{n-1} c_k c_{k+1+l} & \equiv S_{n,0} \text{ under } H_0 \end{cases} \end{aligned} \quad (2.19)$$

where $\text{sgn}(x)$ is one for $x \geq 0$ and -1 for $x < 0$, l is the largest integer no larger than $(j + \gamma)\Delta$ and $\delta = (j + \gamma)\Delta - l$. Note that $l \neq 0$ and $0 \leq \delta \leq 1$.

Defining the per-chip signal-to-noise ratio as

$$SNR = \frac{A_0^2 T_c}{2N_0} \quad (2.20)$$

we have

$$\frac{\lambda_n}{2\sigma_n^2} = \begin{cases} \frac{1}{n} (SNR) S_{n,1}^2 & \text{under } H_1 \\ \frac{1}{n} (SNR) S_{n,0}^2 & \text{under } H_0 \end{cases} \quad (2.21)$$

The acquisition system does not know the exact values of $S_{n,0}$ and $S_{n,1}$ in advance; however, we can use nominal worst-case values for these parameters when designing the thresholds. Simulation results on the schemes suggest that modeling $\sum_{k=0}^{n-1} c_k c_{k+\text{sgn}(\gamma)} \approx 0$ under H_1 and modeling $\sum_{k=0}^{n-1} c_k c_{k+l}$ and $\sum_{k=0}^{n-1} c_k c_{k+1+l} \lesssim \sqrt{n}$ under H_0 achieve the desired decision errors. Using these approximations, we have

$$\frac{\lambda_n}{2\sigma_n^2} = \begin{cases} \approx n(\text{SNR})(1 - |\gamma| \Delta)^2 & \equiv \frac{\lambda_{n,1}}{2\sigma_n^2} \text{ under } H_1 \\ \lesssim \text{SNR} & \equiv \frac{\lambda_{n,0}}{2\sigma_n^2} \text{ under } H_0 \end{cases} \quad (2.22)$$

The cumulative distribution function (cdf) of Y_n can be written in terms of the Q function as

$$F_{Y_n}(y_n) = \begin{cases} 0 & y_n < 0 \\ P(Y_n \leq y_n) = 1 - Q(\sqrt{\lambda_n' \sigma_n^2}, \sqrt{y_n' \sigma_n^2}) & y_n \geq 0 \end{cases} \quad (2.23)$$

where the Q function is defined as [Ref. 8]

$$Q(\zeta, \xi) = \int_{\zeta}^{\infty} x e^{-(x^2 + \zeta^2)/2} I_0(\zeta x) dx \quad (2.24)$$

An iterative algorithm for calculating the Q function is given in [Ref 9].

We now illustrate design techniques for determining a good choice of thresholds to use in the FSS test, the SPRT and the TSPRT such that our decision errors are less than our desired limits. We denote α and β as the desired false-alarm and miss probabilities respectively.

First consider the FSS test. Using the Q function, we can write the false alarm and miss probabilities as

$$P_{fa} = Q(\sqrt{\lambda_{M,0}' \sigma_M^2}, \sqrt{\tau' / \sigma_M^2}) \quad (2.25)$$

$$P_{miss} = 1 - Q(\sqrt{\lambda_{M,1}' \sigma_M^2}, \sqrt{\tau' / \sigma_M^2}) \quad (2.26)$$

From these equations, the values of τ' and M can be obtained by iteratively solving the equations simultaneously, such that $P_{fa} \leq \alpha$ and $P_{miss} \leq \beta$.

For the SPRT we can use (2.14) to calculate the thresholds by inserting equalities in place of the inequalities.

To devise thresholds for the TSPRT, we split this test into two parts: a SPRT with thresholds $\hat{A}$ and $\hat{B}$ and errors α_{spert} and β_{spert} , and a FSS test with sample size $\hat{n}$, threshold

$\hat{\tau}$ and errors α_{fss} and β_{fss} . It can be shown [Ref. 10] that the errors of the TSPRT are bounded by

$$\begin{aligned}\alpha_{ispri} &\leq \alpha_{spri} + \alpha_{fss} \\ \beta_{ispri} &\leq \beta_{spri} + \beta_{fss}\end{aligned}\tag{2.27}$$

So, to ensure $\alpha_{ispri} \leq \alpha$ and $\beta_{ispri} \leq \beta$, we split these errors into the sums of the errors due to the SPRT and the errors due to the FSS test. Specifically, we let

$$\begin{aligned}\alpha &= \alpha_{spri} + \alpha_{fss} = p_0\alpha + (1-p_0)\alpha \\ \beta &= \beta_{spri} + \beta_{fss} = p_1\beta + (1-p_1)\beta\end{aligned}\tag{2.28}$$

where p_0 and p_1 are constants in $[0,1]$. By setting $\alpha_{spri} = p_0\alpha$ and $\beta_{spri} = p_1\beta$, we can design the thresholds $\hat{A}$ and $\hat{B}$ according to (2.14). Similarly, by setting $\alpha_{fss} = (1-p_0)\alpha$ and $\beta_{fss} = (1-p_1)\beta$, we can design $\hat{\tau}$ and $\hat{n}$ using (2.25) and (2.26). Note that if we set $p_0 = p_1 = 0$, the TSPRT becomes the FSS test. If we set $p_0 = p_1 = 1$, the TSPRT becomes the SPRT. For values of p_0 and p_1 in $(0,1)$ the TSPRT can be considered a mixture of an SPRT and an FSS test.

D. PERFORMANCE COMPARISONS

We now present simulation results showing the expected sample size and power function for each of our three decision schemes. The expected sample size, called the *average sample number* (ASN), is the average number of chips that the test runs before deciding between H_0 and H_1 . (Note that for the FSS test the ASN is predetermined and equal to M). The power function is the probability of accepting H_1 . For all simulations in this thesis, the updating of the local m-sequence generator is half a chip (i.e., $\Delta = 1/2$), the value of γ used is 0.5, the m-sequence used has a length of 1023 with primitive polynomial $1 + x^4 + x^5 + x^8 + x^{10}$, and the values of p_0 and p_1 (the mixture weights for the TSPRT) are both 0.5.

Figure 4 on page 18 depicts the ASN for the FSS test, SPRT and TSPRT with $\alpha = \beta = 0.01$. Figure 5 on page 18 depicts the power functions for the same three tests. All functions are plotted in terms of $|j + \gamma|$. When $|j + \gamma| = 0.5$ we have hypothesis H_1 (since $j = 0$ and $\gamma = 0.5$). When $|j + \gamma| = 2.0$ we have hypothesis H_0 . The sample size M of the FSS test is 258. For all tests the resulting errors are equal to or slightly less than our design values chosen for α and β , so our threshold modeling (2.22) works

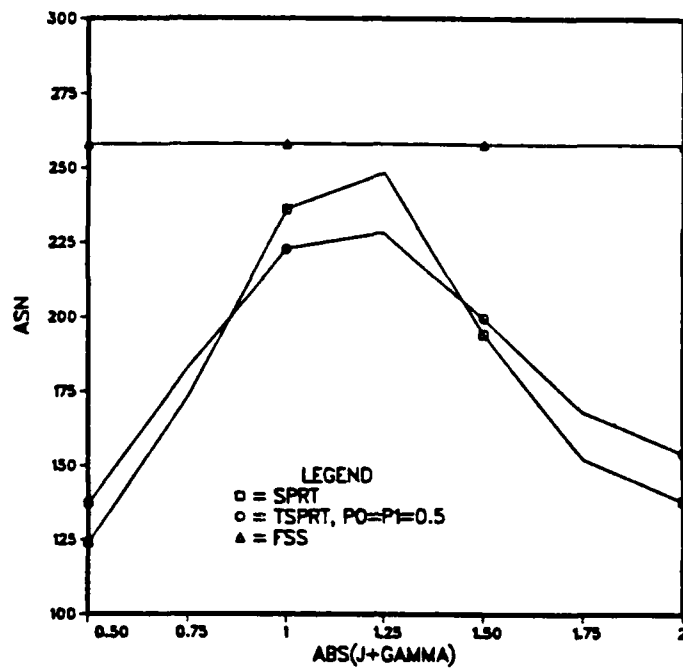


Figure 4. ASN for Acquisition Schemes ($p_0 = p_1 = 0.5$)

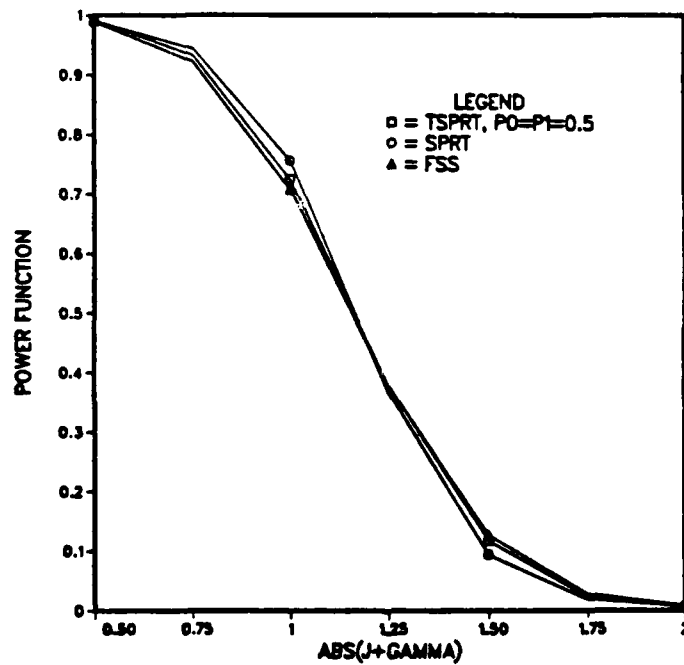


Figure 5. Power Functions ($p_0 = p_1 = 0.5$)

rather well. Under H_0 and H_1 the SPRT has the smallest ASN. Both the SPRT and the TSPRT perform much better than the FSS test under both H_0 and H_1 . It has been found that the ASN of the SPRT is larger than that for the TSPRT for some range of $|j + \gamma|$ between 0.5 and 2.0. (It was found that the ASN of the SPRT can even be larger than the sample size of the FSS test when the errors α and β are chosen to be very small).

Among all uncertainty phases to be tested in an m-sequence period, the synchronization condition occurs only once, the condition $0.5 < |j + \gamma| < 2.0$ occurs at one or a few uncertainty phases (depending on the value of Δ), and the non-synchronization condition (H_0) occurs at the remaining phases. So, to reduce the time taken to acquire the incoming m-sequence, we should minimize the ASN under H_0 . One way to minimize the ASN under H_0 is to utilize the SPRT acquisition scheme, however this results in high ASNs for the cases $0.5 < |j + \gamma| < 2.0$. It appears that the fastest acquisition will occur if we use the TSPRT with a large value of p_1 (2.28). This makes the ASN of the TSPRT approximately equal to that of the SPRT under H_0 and results in smaller values of ASN under the cases $0.5 < |j + \gamma| < 2.0$. Figure 6 on page 20 shows this comparison for the ASN function. Figure 7 on page 20 shows this comparison for the power function.

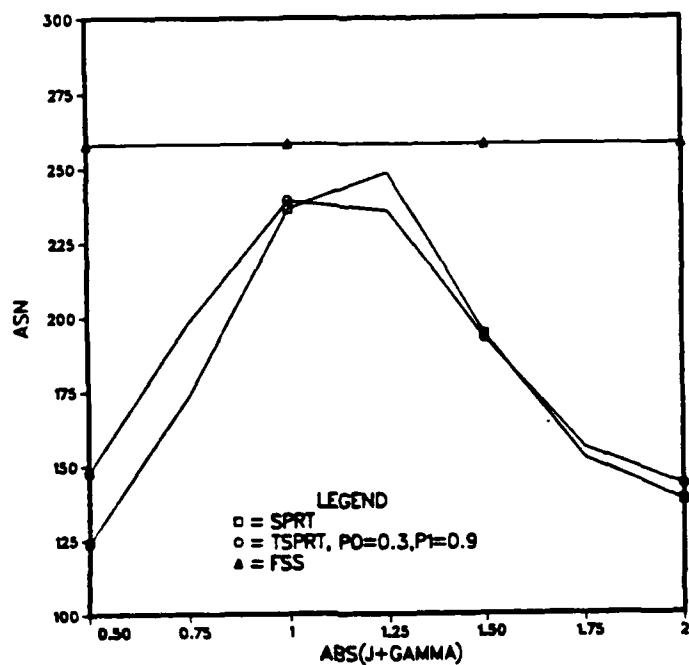


Figure 6. ASN for Acquisition Schemes ($p_0 = 0.3$, $p_1 = 0.9$)

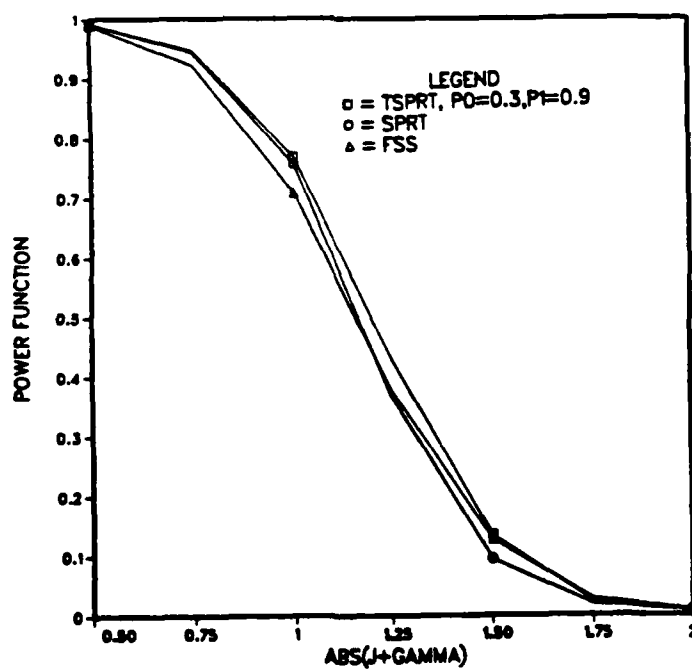


Figure 7. Power Functions ($p_0 = 0.3$, $p_1 = 0.9$)

III. PERFORMANCE IN THE FADING CHANNEL

A. GENERAL DESCRIPTION OF FADING

In a communication system, the received signal may suffer time varying power fluctuations. Since these power fluctuations most often represent an attenuation in the received signal (as opposed to an amplification), the phenomenon is termed *fading*. The power instability can be considered a result of instantaneous amplification or attenuation of the transmitted signal's amplitude.

In our previous discussion, the parameters and characteristics of the communications channel were considered to be fully known and time-invariant. In reality, especially in cases where part of the communications path consists of an unguided medium, the channel characteristics vary randomly with time. Utilizing an unguided medium (e.g., the atmosphere or ionosphere) as part of the channel unavoidably subjects the communications system to the irregular variations which often occur in nature. Channel parameters can be altered due to rain and humidity (which absorb microwave energy), atmospheric density variations (which refract and reflect electromagnetic waves), mountains and other physical obstacles (which reflect or dissipate microwaves), changes in the ionosphere's electron concentration distribution (which affects refraction of High Frequency waves), as well as other factors. Semiperiodic and totally random variations in the propagation attributes of the channel can be viewed as altering the channel's transfer function.

Although fading arises from a variety of causes, fading phenomena can often be modeled as causing multipath distortion. That is, many fading conditions can be modeled as causing several alternate transmission routes to arise between the transmitter and receiver. In addition to the transmission paths which the engineer incorporated into his design, additional unplanned transmission paths are excited randomly with time. A signal component arriving at the receiver via one path may be out of phase with the signal component arriving via a second (and additional) path(s). The arriving signal components interfere (constructively and destructively), resulting in fading.

Fading channels are classified as *flat fading* channels and *bandwidth-selective* channels. In a flat fading channel all of the frequencies present in the transmitted signal fade in exactly the same manner and the communication signal can be regarded as undistorted in relative shape. The received signal in such cases can be viewed as multiplied

by a random variable which accounts for the fading. One example of this type of fading occurs in space links passing through a turbulent atmosphere [Ref. 11 : pp.131-132]. Flat fading is further categorized as *slow fading* and *fast fading*. In slow fading the channel variations (which give rise to the signal power fluctuations) are slower than the longest period component in the waveform. In such cases we consider the received signal to be multiplied by a constant (although random) amplitude coefficient. In fast fading, the channel variations are comparable to the period components in the waveform. In such cases we consider the received signal to be multiplied by a fluctuating time function which effectively modulates the received signal.

In bandwidth-selective channels some of the transmitted signal's frequency components are affected to a greater degree than others. Such channel effects are equivalent to the insertion of a bandpass filter into the signal transmission path [Ref. 11 : p.132] which causes the signal's edge frequencies to have a larger or smaller fade than the center frequencies. In such cases a sudden attenuation of just the high frequency components may occur, or, alternately, a loss of only low frequency power may be observed.

It should be noted at this point that fading is a major communications obstacle which can cause signal attenuations of several tens of decibels. Figure 8 on page 23 shows the median duration of fast fades as a function of fade depth for a 4 GHz signal over a path length of 30-35 miles.

In the remainder of this paper we narrow our focus to consider only *flat, slow fading channels*.

Many methods have been devised to counter the effects of fading. The performance reduction that characterizes fading is caused by the received signal's amplitude being weakened in comparison to its design value. One technique for minimizing the received signal's amplitude fluctuations is to incorporate automatic gain control (AGC) into the receiving amplifier [Ref. 12 : p.112]. This is a negative feedback method from control theory in which the receiver amplifier output signal is employed to adjust the amplifier's gain. The circuit is designed so that the output level remains steady notwithstanding changes in the input amplitude. The major disadvantage of this technique is that the channel additive noise amplitude is also increased along with the signal amplitude.

A better technique to counter fading consists of dividing the communication's signal power between multiple subchannels which fade independently of each other. The probability that fading will be extreme in all subchannels simultaneously is very low, so if all subchannel outputs are used to reconstruct the signal, we expect better performance than over a single channel. The use of multiple transmission subchannels is called

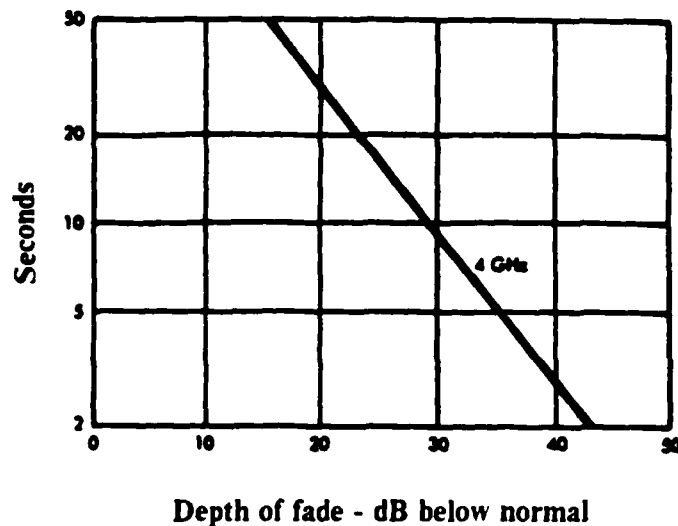


Figure 8. Fast Fade Duration

Source: Pierce, J.R., and Pasner, E.C., *Introduction To Communication Science And Systems*, Plenum Press, 1980

diversity transmission [Ref. 13 : pp.346-350]. Diversity transmission methods fall under three general areas [Ref. 14 : pp.632-634]:

1. *space diversity* utilizes several receivers in separate geographical areas. The receiver outputs are collected to reconstitute the desired signal. When the signal at one receiver fades it is hoped that the signal at a distant location (i.e., at the other receivers) will be unaffected.
2. *time diversity* employs retransmission of the same signal at spaced time intervals. It is hoped that if time-varying fading affects a signal at one time, it will not affect the signal at a later time.
3. *frequency diversity* employs multiple frequency channels to combat the effects of frequency selective fading. It is hoped that if the signal suffers fading in one frequency band, other frequency bands will still be usable.

B. MATHEMATICAL DESCRIPTION OF FADING

Thus far we have described fading phenomena in general qualitative terms. Although we have asserted that in slow, flat fading channels the received signal is effectively multiplied by a constant (albeit random) amplitude coefficient, we have yet to grant mathematical justification to this statement. In order to proceed any further with the analysis of our receiver in a fading channel, we must determine a suitable mathematical model to describe fading. Since fading phenomena largely depend on unpredictable variations in the environment, we can not model fading deterministically. (Indeed, if it could be modeled deterministically we would then simply account for it by

adding an amplification factor in our receiver design and no further consideration would be necessary.) The best that can be achieved is a statistical description of the amplitude coefficient.

Communication schemes are typically constructed to achieve a preset minimum standard of performance. For example, we may design a system so that the probability of false receiver synchronization is 0.01 and the probability of failing to detect synchronization is 0.01. Using a statistical description of fading we can simulate our system in a fading channel and determine if our minimum standards of performance are attainable. If the random channel attenuation factor degrades our system to the point where we fail to satisfy our performance standards we can employ three corrective options.

Our first option is to attempt to modify the receiver design by suitably changing the decision thresholds to meet the performance standards.

Alternately, a second course of action would be to increase our transmitter power to a level where it compensates for a worst-case power loss in the fading channel. In this option we meet our performance standards by allowing adequate "design margins."

Our third approach would be to relax our performance standards. This option may often be the most attractive if we consider other important design variables. As an illustration of how the third approach may be the most advantageous, consider again the system where we desire that the probability of false receiver synchronization and probability of failing to detect synchronization both be 0.01. Suppose further that a simulation shows that fading reduces both these specifications to 0.1. In our fading channel, with our degraded performance standards, it may be the case that the expected time to acquire synchronization is one minute. If we employ option one above and redesign the receiver, it may be the case that the performance standards are met, but the time to acquire synchronization is increased to one hour. In such a scenario we would probably sacrifice our performance standards for the sake of acquisition time.

Consider a multipath channel. In addition to the transmission path which the engineer incorporated into his design, additional unplanned transmission paths of various electrical lengths are excited. Portions of the power transmitted at a given time arrive "sequentially" at the receiver over each of the sundry transmission paths. Since we are restricting our consideration only to *flat* fading channels, it suffices to examine the channel effects on a single frequency component; all other frequencies will be affected in a similar way.

A particular frequency component is transmitted at a certain time. The received signals will arrive at the receiver over the various paths. Since the lengths of the paths

differ, the received sine waves will (presumably) be out of phase with each other, and when combined they will constructively or destructively interfere. The fading channel causes the electrical path lengths to vary continually and randomly and this in turn causes the phase differences between the component tones to vary. The superposition of all the tones (i.e., the received signal) will have an envelope and phase that will vary randomly as a result. So, in a fading channel, the received signal is the superposition of a number of *random phasors*.

Initially assume that each of the random phasors is of a comparable magnitude. Consider each of the individual phasors to be resolved into quadrature components. Since the fading channel randomly varies the path lengths, the quadrature components of one particular random phasor (at a particular time) will be independent of the quadrature components of all of the other random phasors. Additionally, the quadrature components of any particular phasor are uncorrelated (since the phasor is resolved into sine and cosine components).

Before finding the total received signal, first find the total signal in each of the quadrature directions. We do this by simply adding the particular quadrature components of all the random phasors. In doing this we are adding a number of independent random variables (from the same underlying distribution) together, and the sum, by the central limit theorem [Ref. 15 : p.210], approaches a Gaussian random variable as the number of contributing phasors grows large. Thus, at any given time, the two quadrature resultants are independent Gaussian random variables. The envelope of the received signal will therefore be Rayleigh distributed [Ref. 16 : pp.101-103]. (Note that the phase of the received signal is of no concern since we utilize a noncoherent detection scheme in the receiver.)

The received signal, considering all frequencies, will thus be effectively multiplied by a coefficient, ψ , having the probability density function (pdf)

$$f_{\Psi}(\psi) = \frac{\psi}{\sigma^2} e^{-\psi^2/2\sigma^2} \quad (3.1)$$

where σ^2 , determined by the nature of the channel, is the variance of each of the Gaussian quadrature random variates.

Since the above discussion utilized the central limit theorem, the question naturally arises: how many alternate transmission paths are required for the Rayleigh model to be valid? Surprisingly, it has been noted [Ref. 17 : pp.348-351] that as few as six sine

waves with independently varying random phases will combine to give a fluctuating resultant whose envelope closely approaches Rayleigh statistics. See Figure 9.

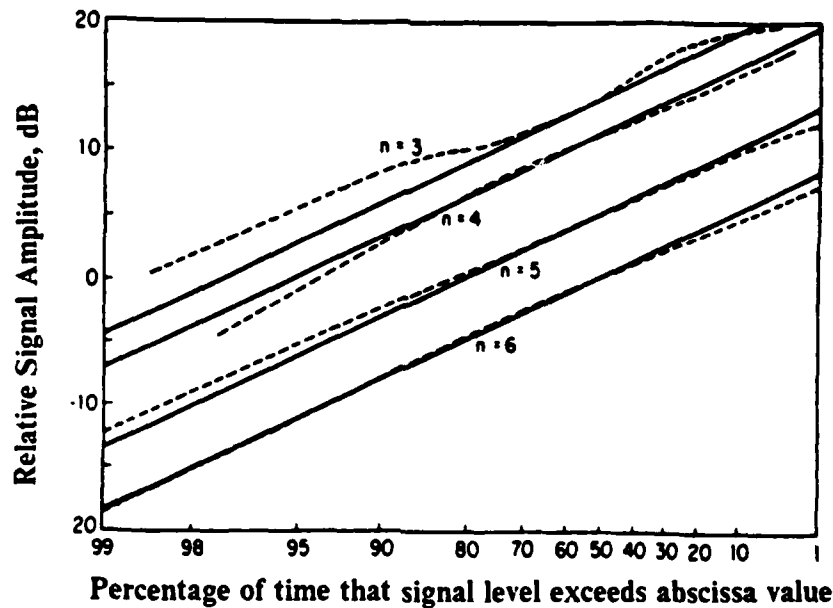


Figure 9. Envelope Distributions For n Equal-Amplitude Random-Phase Phasors: Solid Lines are Rayleigh Distributions.

Source: Schwartz, M., Bennett, W.R., Stein, S., *Communication Systems and Techniques*, McGraw Hill, 1966, p349.

In deriving the Rayleigh channel results, we assumed that each of the random phasors had a comparable magnitude. It is more realistic to assume that one of the paths, namely the path that we intend to design for (i.e., the no-fading path), is much stronger than the other paths. The power in the strong, intended path is termed the *direct component*. The total power in all the weak paths is termed the *diffuse component*. The net received signal will then be a combination of a steady tone (again, considering a single frequency since the channel is *flat*) and Gaussian quadrature random variables. The envelope statistics for "specular-plus-Rayleigh" fading, called *Ricean fading*, are given by [Ref. 17 : p.372]:

$$f_{\Psi}(\psi) = \frac{\psi}{\sigma^2} e^{-s^2/2\sigma^2} e^{-\psi^2/2\sigma^2} I_0\left(\frac{\psi s}{\sigma^2}\right) \quad (3.2)$$

where

ψ = instantaneous value of the fading coefficient
 s^2 = power in the direct component
 $2\sigma^2$ = power in the diffuse component

and with the total power having a normalized value of 1.

Letting $r = \frac{s^2}{2\sigma^2}$ we find that

$$f_{\Psi}(\psi) = 2\psi(1+r)e^{-r-\psi^2(1+r)}I_0(2\psi\sqrt{r(1+r)}) \quad (3.3)$$

Note from (3.2) that when the ratio $\frac{s^2}{2\sigma^2}$ approaches zero, the pdf approaches the Rayleigh pdf (3.1). When $r \rightarrow \infty$, we have no fading. Figure 10 shows the Ricean probability density function (3.3) for various values of r . Note that for large values of r the distribution of the fading coefficient is concentrated around unity.

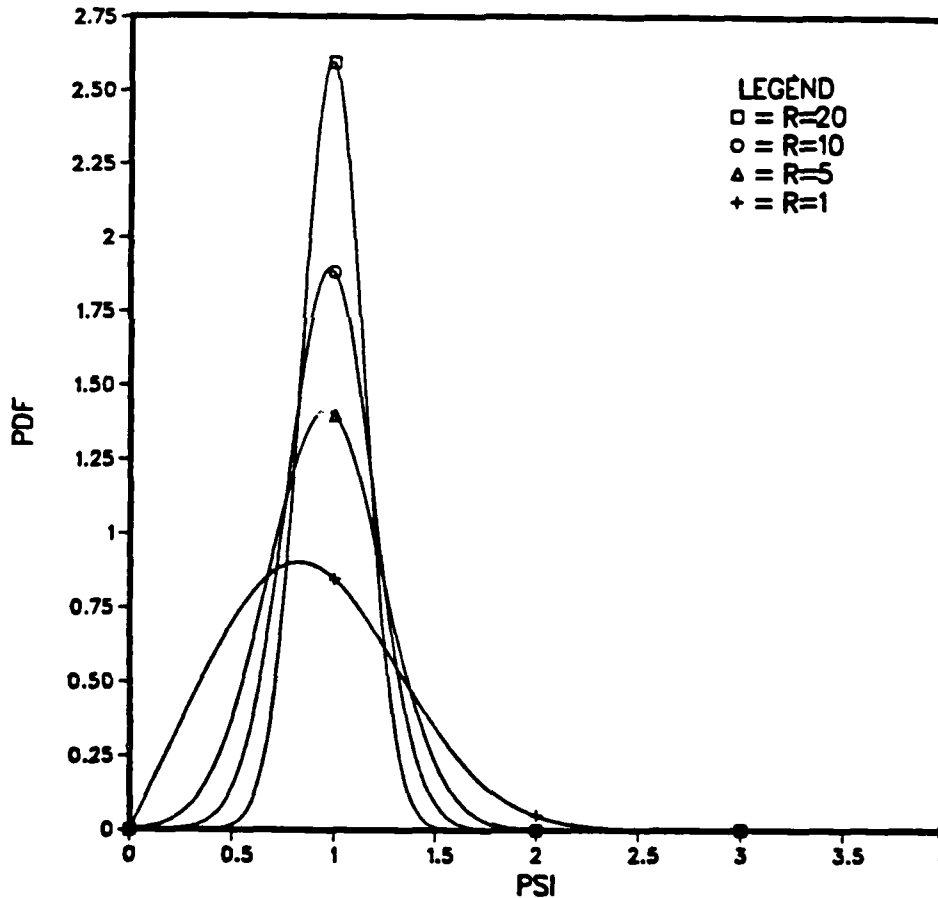


Figure 10. Ricean distribution for various values of r

C. TEST STATISTIC DENSITY FUNCTION

We now determine the probability density function of the receiver's test statistic in the presence of a slow, flat fading channel. Since we assume slow fading, ψ does not vary much within one data bit interval, and thus can be assumed constant during the acquisition process.

The received signal through the fading channel is given by

$$r(t) = A_o \psi a(t + i\Delta T_c) \cos(\omega_o t + \theta) + n(t) \quad (3.4)$$

where ψ is the fading random variable with a Ricean probability density function given by (3.2). Referring to our receiver structure (Figure 1), equations (2.2) and (2.3) become

$$X_{i,n} = \frac{A_o}{2} \psi T_c S_n \cos \theta + N_{i,n} \quad (3.5)$$

$$X_{q,n} = \frac{A_o}{2} \psi T_c S_n \sin \theta + N_{q,n} \quad (3.6)$$

where $N_{i,n}$ and $N_{q,n}$ are still defined by (2.4) and (2.5). As before, the test statistic for determining alignment is

$$Y_n = X_{i,n}^2 + X_{q,n}^2 \quad (3.7)$$

The conditional probability density function of y_n given ψ is noncentral Chi-squared:

$$f(y_n | \psi) = \frac{1}{2\sigma_n^2} e^{-(y_n + \lambda_n \psi^2) / 2\sigma_n^2} I_0\left(\sqrt{\lambda_n \psi^2 y_n} / \sigma_n^2\right), \quad y_n \geq 0 \quad (3.8)$$

where, as with no fading, we have $\lambda_n = (A_o^2/4)T_c^2 S_n^2$ and $\sigma_n^2 = nT_c N_o/4$.

We have readily determined the test statistic's conditional density function given ψ . We are interested in the density function without conditionality. To determine $f(y_n)$, we "integrate out" the ψ dependence:

$$\begin{aligned} f(y_n) &= \int_0^\infty f(y_n | \psi) f(\psi) d\psi \\ &= \int_0^\infty \frac{1}{2\sigma_n^2} e^{-(y_n + \lambda_n \psi^2) / 2\sigma_n^2} I_0\left(\frac{\sqrt{\lambda_n \psi^2 y_n}}{\sigma_n^2}\right) \frac{\psi}{\sigma^2} e^{-(s^2 + \psi^2) / 2\sigma^2} I_0\left(\frac{\psi s}{\sigma^2}\right) d\psi \end{aligned}$$

$$= \frac{1}{2\sigma_n^2\sigma} e^{-y_n 2\sigma_n^2} e^{-s^2 2\sigma^2} \int_0^\infty \psi e^{-\psi^2(\lambda_n 2\sigma_n^2 + 1 2\sigma^2)} I_0\left(\frac{\sqrt{\lambda_n \psi^2 y_n}}{\sigma_n^2}\right) I_0\left(\frac{\psi s}{\sigma^2}\right) d\psi \quad (3.9)$$

To simplify notation, we combine terms by letting

$$a = \frac{\sqrt{\lambda_n y_n}}{\sigma_n^2} ; \quad b = \frac{s}{\sigma^2} ; \quad c = \frac{\lambda_n}{2\sigma_n^2} + \frac{1}{2\sigma^2} \quad (3.10)$$

and

$$d = \frac{1}{2\sigma_n^2\sigma^2} e^{-y_n 2\sigma_n^2} e^{-s^2 2\sigma^2} \quad (3.11)$$

Substituting (3.10) and (3.11) into (3.9), we can rewrite the probability density function as

$$f(y_n) = d \int_0^\infty \psi e^{-\psi^2 c} I_0(a\psi) I_0(b\psi) d\psi \quad (3.12)$$

The Bessel function is defined as a power series expansion. Substituting the definition of the Bessel function into (3.12) yields

$$f(y_n) = d \int_0^\infty \psi e^{-\psi^2 c} \sum_{j=0}^\infty \frac{\psi^{2j} a^{2j}}{4^j (j!)^2} \sum_{k=0}^\infty \frac{\psi^{2k} b^{2k}}{4^k (k!)^2} d\psi$$

Since integration and summation are both linear operations, the above formula can be rewritten as

$$f(y_n) = d \sum_{j=0}^\infty \sum_{k=0}^\infty \frac{a^{2j}}{4^j (j!)^2} \frac{b^{2k}}{4^k (k!)^2} \int_0^\infty e^{-\psi^2 c} \psi^{2(j+k)+1} d\psi \quad (3.13)$$

It has been derived [Ref. 18 : p. 271] that

$$\int_0^{\infty} x^{2n+1} e^{-ax^2} dx = \frac{n!}{2a^{n+1}} \quad (3.14)$$

Applying (3.14) to (3.13) results in

$$f(y_n) = d \sum_{j=0}^{\infty} \sum_{k=0}^{\infty} \frac{a^{2j} b^{2k}}{4^j 4^k (j!)^2 (k!)^2} \frac{(j+k)!}{2c^{j+k+1}}$$

which is rewritten as

$$f(y_n) = \frac{d}{2c} \sum_{j=0}^{\infty} \frac{a^{2j}}{4^j (j!)^2 c^j} \sum_{k=0}^{\infty} \frac{b^{2k} (j+k)!}{4^k (k!)^2 c^k}$$

$$f(y_n) = \frac{d}{2c} \sum_{j=0}^{\infty} \left(\frac{a^2}{4c} \right)^j \frac{1}{j!} \sum_{k=0}^{\infty} \left(\frac{b^2}{4c} \right)^k \frac{1}{(k!)^2} \frac{(j+k)!}{j!} \quad (3.15)$$

The term $(j+k)! = (j+1)(j+2)(j+3)\dots(j+k)$ can be written using Appell's symbol [Ref. 19 : p. 7] as $(j+1, k)$. Similarly, using this notation, we can rewrite $k!$ as $(1, k)$. Adopting this notation, the density function in (3.15) is written as

$$f(y_n) = \frac{d}{2c} \sum_{j=0}^{\infty} \left(\frac{a^2}{4c} \right)^j \frac{1}{j!} \sum_{k=0}^{\infty} \frac{(j+1, k)}{(1, k)} \frac{(b^2/4c)^k}{k!} \quad (3.16)$$

The second summation is in the form of the confluent hypergeometric function [Ref. 20 : pp. 298-320]:

$$\sum_{k=0}^{\infty} \left(\frac{b^2}{4c} \right)^k \frac{1}{(k!)^2} (j+1, k) = {}_1F_1(j+1; 1; b^2/4c) \quad (3.17)$$

Substituting (3.17) into (3.16) yields

$$f(y_n) = \frac{d}{2c} \sum_{j=0}^{\infty} \left(\frac{a^2}{4c} \right)^j \frac{1}{j!} {}_1F_1(j+1; 1; b^2/4c) \quad (3.18)$$

To further simplify, we express ${}_1F_1$ in terms of a Laguerre Polynomial [Ref. 21 : p. 430] as

$${}_1F_1(c+n; c; x) = \frac{n!}{(c, n)} e^x L_n^{c-1}(-x) \quad (3.19)$$

Substituting (3.19) into (3.18), we find

$$f(y_n) = \frac{d}{2c} \sum_{j=0}^{\infty} \left(\frac{a^2}{4c} \right)^j \frac{1}{j!} e^{b^2/4c} L_j^0(-b^2/4c) \quad (3.20)$$

where $L_j^0(.) = L_j(.)$ is the Laguerre Polynomial of degree j . Using an identity involving Laguerre Polynomials [Ref. 21 : p. 316]:

$$\sum_{k=0}^{\infty} \frac{t^k}{k!} L_k(x) = e^t J_0(2\sqrt{xt})$$

we can rewrite (3.20) as

$$f(y_n) = \frac{d}{2c} e^{b^2/4c} e^{a^2/4c} J_0\left(2\sqrt{\frac{-a^2 b^2}{4^2 c^2}}\right) \quad (3.21)$$

where $J_0(.)$ is the Bessel function of the first kind of order zero. Using $i = \sqrt{-1}$ we have

$$f(y_n) = \frac{d}{2c} e^{b^2/4c} e^{a^2/4c} J_0\left(\frac{iab}{2c}\right) \quad (3.22)$$

But $J_0(ix) = I_0(x)$, and so we rewrite the probability density function $f(y_n)$ as

$$f(y_n) = \frac{d}{2c} e^{b^2 4c} e^{a^2 4c} I_0\left(\frac{ab}{2c}\right) \quad (3.23)$$

Substituting the values of a , b , c and d from (3.10) and (3.11) back into (3.23) we have our final formula for $f(y_n)$:

$$f(y_n) = \frac{1+r}{2\sigma_n^2[(1+r) + \lambda_n/2\sigma_n^2]} \exp\left[\frac{-[(1+r)(y_n/2\sigma_n^2) + r(\lambda_n/2\sigma_n^2)]}{(1+r) + \lambda_n/2\sigma_n^2}\right] \\ \cdot I_0\left(\frac{2\sqrt{r(1+r)(y_n/2\sigma_n^2)(\lambda_n/2\sigma_n^2)}}{1+r + \lambda_n/2\sigma_n^2}\right) \quad (3.24)$$

This interesting result states that the probability density function of y_n is still noncentral Chi-squared. The density function is the same form as (2.8) but the parameters λ_n and σ_n^2 are replaced by new values:

$$\lambda_n \rightarrow \lambda_n \frac{r}{1+r} \quad (3.25a)$$

$$\sigma_n^2 \rightarrow \sigma_n^2 \frac{1+r + \lambda_n/2\sigma_n^2}{1+r} \quad (3.25b)$$

The cumulative distribution function corresponding to (3.24) is

$$F_{y_n}(y_n) = 1 - Q\left(\sqrt{\frac{r(\lambda_n/\sigma_n^2)}{1+r + \lambda_n/2\sigma_n^2}}, \sqrt{\frac{(1+r)(y_n/\sigma_n^2)}{1+r + \lambda_n/2\sigma_n^2}}\right) \quad (3.26)$$

D. SIMULATION TECHNIQUES

Now that we have a precise description of the receiver test statistic's density function, we wish to evaluate the receiver's performance in a fading channel. Three receiver characteristics are of primary interest for each acquisition scheme: the actual probability of false alarm (PFA), the actual probability of detection (PD) and the number of chips needed to make a decision (called the average sample number (ASN)).² A fourth characteristic, acquisition time, which is the average amount of time needed to correctly

² Note that the actual probability of false alarm (PFA) and actual probability of detection (PD) are distinguished from the desired probability of false alarm (α) and desired probability of detection ($1 - \beta$).

phase synchronize the incoming and locally generated m-sequences, depends on the PFA, PD and ASN.

There are two general approaches that one may use to determine the PFA, PD and ASN for a given fading channel. The first approach is *mathematical analysis*. This method entails determining an equation or formula which allows one to directly compute the quantity of interest in terms of other known quantities. Needless to say, closed form analytical solutions for the receiver under consideration have eluded, and continue to elude, researchers. (The problem, in fact, has not even been solved for the mathematically simpler case where the correlators are reset to achieve independent identically distributed samples.)

The second approach used to determine the PFA, PD and ASN for given fading conditions is *simulation*. With this method, the threshold values for the various tests are stored in computer memory, and a computer program is used to simulate the receiver operation and the received signal. When the simulation shows that a threshold is exceeded, we stop the run and measure the ASN. By running the tests a large number of times under various conditions, we can confidently determine the PFA and PD. All results presented in this thesis for ASN are based on simulation.

To illustrate the simulation techniques, we will focus on attempting to determine the average ASN (for a particular scheme) with fading. The techniques for determining the average PFA and average PD are identical.

We are interested in determining the effect that fading has on the detection scheme's ASN. The average sample number depends on many factors (e.g., noise power spectral density, SNR, etc). To isolate the fading effect, we hold all factors stationary except for the fading magnitude. So, in other words, we are interested in determining the effects of fading on ASN with all other factors remaining constant.

With the above restriction, the ASN is a function of the fading random variable ψ ,

$$\text{ASN} = g(\psi) \quad (3.27)$$

The expected value of ASN in the presence of fading is thus

$$E[\text{ASN}] = E[g(\psi)] = \int_0^\infty g(\psi)f(\psi)d\psi \quad (3.28)$$

where $f(\psi)$ is the Rician probability density function describing the fading, given by (3.2) or (3.3).

This is precisely the point where the analysis stops. If we could solve the integral in (3.28) we would have the answer we seek. However, since the ASN function (3.27) is presently unknown, the integral is not solvable and we must resort to simulation. The only purpose in performing the following simulation is to solve (3.28) for $E[ASN]$ (and to solve similar equations for $E[PD]$ and $E[PFA]$).

1. Monte Carlo Simulation

We wish to determine a value for the integral in (3.28). The strong law of large numbers [Ref. 22 : pp. 88-90] states that for "large enough" n , we can determine $E[g(\psi)]$ (and hence (3.28)) from a sample mean:

$$\frac{1}{n} \sum_{i=1}^n g_i(\psi) \rightarrow E[g(\psi)] \quad (3.29)$$

So, our algorithm for computing $\int_0^\infty g(\psi)f(\psi)d\psi$ is very simple [Ref. 22 : pp. 181-194]:

1. Generate n independent samples of a Ricean random variable ψ whose density function describes the fading.
2. Determine $g(\psi_i)$, the ASN for a particular ψ_i , either directly or by simulation for $1 \leq i \leq n$.
3. The desired result is $\frac{1}{n} \sum_{i=1}^n g(\psi_i)$.

Step 1. above, generating independent samples of a Ricean random variable, is of critical importance. If the samples are not Ricean and/or not independent the algorithm will not be successful. The method used for generating such independent Ricean random variables is included in the Appendix at the end of this thesis.

We characterize the severity of fading by specifying a value of r , where r is the ratio of the power in the direct component to the power in the diffuse component and the Ricean density function is described by (3.3). For each Ricean random variate ψ_i , we run a simulation program to determine the ASN of a test given that value of ψ_i . We run the simulation program for a large number of independent variates ψ_i , and compute the average ASN by the summation in step 3. above.

The Monte Carlo method gives us only an approximation of the integral we are trying to evaluate. The sample mean that we compute in the third step of the algorithm is an estimate of the integral which becomes exact only as $n \rightarrow \infty$. We are interested in determining the mean square error of our estimator.

Recall that our goal is to evaluate $E[ASN]$ given by

$$\int_0^{\infty} g(\psi) f(\psi) d\psi \equiv I \quad (3.30)$$

where we designate the exact (albeit unknown) value of this integral by I . Designate the estimate of the integral by $\bar{g}_n$. That is,

$$\bar{g}_n = \frac{1}{n} \sum_{i=1}^n g(\psi_i) \quad (3.31)$$

Then the mean of the square of the error of our estimator is given by [Ref. 22 : pp. 194-196]

$$E[(\text{sample mean error})^2] = E[(\bar{g}_n - I)^2] \equiv \text{Var}(\bar{g}_n) \quad (3.32)$$

But by using the definition of $\bar{g}_n$, the variance of $\bar{g}_n$ can be rewritten as

$$\text{Var}(\bar{g}_n) = \text{Var}\left\{\frac{1}{n} \sum_{i=1}^n g(\psi_i)\right\} = \frac{1}{n^2} \text{Var}\left\{\sum_{i=1}^n g(\psi_i)\right\} \quad (3.33)$$

Substituting (3.33) into (3.32) and using the shorthand notation $V \equiv E[(\text{sample mean error})^2]$, we have

$$V = \frac{1}{n^2} \text{Var}\{g(\psi_1) + g(\psi_2) + \dots + g(\psi_n)\} \quad (3.34)$$

Since the variance of the sum of independent random variables is the sum of the variances, (3.34) is revised:

$$V = \frac{1}{n^2} \{\text{Var}[g(\psi_1)] + \text{Var}[g(\psi_2)] + \dots + \text{Var}[g(\psi_n)]\} \quad (3.35)$$

But each of the variance terms are equal to each other. That is to say, $\text{Var}[g(\psi_1)] = \text{Var}[g(\psi_2)] = \dots = \text{Var}[g(\psi_n)] = \text{Var}[g(\psi)]$. Using this fact

$$V = \frac{1}{n^2} \{n \text{Var}[g(\psi)]\} = \frac{1}{n} \text{Var}[g(\psi)] \quad (3.36)$$

But the variance of $g(\psi)$ is equal to $E[g^2(\psi)] - \{E[g(\psi)]\}^2 = E[g^2(\psi)] - I^2$. Thus

$$V = \frac{1}{n} \{ E[\{g(\psi)\}^2] - I^2 \} \quad (3.37)$$

But we know that

$$E[\{g(\psi)\}^2] = \int_0^\infty g^2(\psi) f(\psi) d\psi$$

so we substitute this quantity into (3.37) and arrive at

$$V = \frac{\int_0^\infty g^2(\psi) f(\psi) d\psi - I^2}{n} \quad (3.38)$$

Now, neither of the quantities in the numerator of (3.38) is known since we do not have an expression for $g(\psi)$. Thus, we again resort to the law of large numbers to approximate each of the numerator terms [Ref. 22 : pp. 196-197]:

$$\int_0^\infty g^2(\psi) f(\psi) d\psi \rightarrow \frac{1}{n} \sum_{i=1}^n g^2(\psi_i) \quad (3.39a)$$

$$I \rightarrow \frac{1}{n} \sum_{i=1}^n g(\psi_i) \quad (3.39b)$$

So, our final equation for the mean square error for our Monte Carlo estimator is given by

$$V = \frac{1}{n} \left\{ \frac{1}{n} \sum_{i=1}^n g^2(\psi_i) - \left[\frac{1}{n} \sum_{i=1}^n g(\psi_i) \right]^2 \right\} \quad (3.40)$$

What is so nice about the Monte Carlo method is that as we compute our estimator we can simultaneously compute just how good or bad our estimator is by using (3.40).

Again, the exact value of the integral that we are trying to determine is designated as I . We would like to be able to place an upper and lower bound on I with a certain degree of confidence. Mathematically, we would like to be able to specify a number d and then compute two constants c_1 and c_2 such that

$$P[c_1 < I < c_2] \geq 1 - d \quad (3.41)$$

Chebyshev's Inequality [Ref. 15 : p. 105] states that

$$P[\mu - t\sigma < X < \mu + t\sigma] \geq 1 - \frac{1}{t^2} \quad (3.42)$$

where μ is the mean of the random variable X and σ is its standard deviation. Applying the inequality to our scenario, we have

$$P[\bar{g}_n - t\sqrt{V} < I < \bar{g}_n + t\sqrt{V}] \geq 1 - \frac{1}{t^2} \quad (3.43)$$

Thus, to choose bounds for (3.41), simply let

$$t = \frac{1}{\sqrt{d}} \quad (3.44a)$$

$$c_1 = \bar{g}_n - t\sqrt{V} \quad , \quad c_2 = \bar{g}_n + t\sqrt{V} \quad (3.44b)$$

It is known that Chebyshev's Inequality gives a bound that is often pessimistic [Ref. 22 : p. 197]. A much tighter bound can be obtained if we use a "large enough" number of random variates ψ , so that the central limit theorem is applicable. In such an event it can be shown [Ref. 22 : p. 198] that

$$\frac{\bar{g}_n - I}{\sqrt{V}} \simeq N(0,1) = Y \quad (3.45)$$

where $N(0,1)$ is the standard normal random variable which we designate as Y . If we choose t such that

$$1 - F_Y = P[Y > t] = \frac{d}{2} \quad (3.46)$$

then rearranging (3.43) results in

$$\begin{aligned}
P(|I - \bar{g}_n| < t\sqrt{V}) &= P\left(\left|\frac{I - \bar{g}_n}{\sqrt{V}}\right| < t\right) \\
&= 1 - 2P\left(\frac{I - \bar{g}_n}{\sqrt{V}} > t\right) = 1 - \frac{2d}{2} = 1 - d \quad (3.47)
\end{aligned}$$

As an example of how to apply the above method, suppose we want to determine an upper and lower bound for I with 90% certainty. Using (3.41) we would choose $d = 0.1$. Using this value of d , we determine t (from tables) such that

$$P[Y > t] = \frac{d}{2} \quad (3.48)$$

where Y is the standard normal distribution. Now, with this value of t , we can be 90% certain that I lies between $\bar{g}_n - t\sqrt{V}$ and $\bar{g}_n + t\sqrt{V}$ where $\bar{g}_n$ and V are determined from (3.31) and (3.40) respectively.

2. Numerical Integration

The Monte Carlo method is but one way to evaluate $E[ASN]$ given by the integral in (3.28). A simpler and more direct technique involves estimating the integral by using a numerical method.

The integral we wish to evaluate is given by (3.30) and repeated below:

$$\int_0^{\infty} g(\psi)f(\psi)d\psi \quad (3.49)$$

We have a closed form analytical description of $f(\psi)$, given in (3.24), but we do not know the form of $g(\psi)$. Note that for all ψ , $g(\psi)f(\psi) \geq 0$.

Our approach is to view the integral as an area. We start by dividing the density function of $f(\psi)$ into a certain number of equal areas. For illustration, we will divide the density function into 25 sections, each with an area of $1/25 = 0.04$ square units. So the first area spans the abscissa from 0 to ψ_1 . The second area spans the abscissa from ψ_1 to ψ_2 , and so on. The last area spans the abscissa from ψ_{24} to ∞ . So we can write (3.49) as 25 separate integrals:

$$\int_0^{\infty} g(\psi)f(\psi)d\psi = \int_0^{\psi_1} g(\psi)f(\psi)d\psi + \int_{\psi_1}^{\psi_2} g(\psi)f(\psi)d\psi + \cdots + \int_{\psi_{24}}^{\infty} g(\psi)f(\psi)d\psi \quad (3.50)$$

Note that the values $\psi_1 \dots \psi_{24}$ are easy to determine since $f(\psi)$ is known. Also note that (3.50) is still exact.

Let the midpoints of each of the 25 intervals $(0, \psi_1), (\psi_1, \psi_2) \dots (\psi_{24}, \infty)$ be designated in general as ψ^* and specifically as $\psi_A, \psi_B, \psi_C, \dots \psi_Y$ respectively.³ Note that the majority of the mass of $f(\psi)$ is concentrated in the region between $\psi = 0$ and approximately $\psi = 3$ (depending on the value of r); see Figure 10. Because of this, we expect the 25 points $\psi_A, \psi_B, \dots \psi_Y$ to be clustered close together in the region $(0, 3)$. Putting this another way, we expect each of the first 24 intervals $(0, \psi_1), (\psi_1, \psi_2) \dots (\psi_{23}, \psi_{24})$ to be very "thin".

Now we must make an approximation. Suppose over each of the 25 areas, we assume $g(\psi) = g(\psi^*)$ where ψ^* is the midpoint of the interval (ψ_i, ψ_{i+1}) . Since ψ undergoes a very small variation between ψ_i and ψ_{i+1} , we assume $g(\psi)$ also undergoes only a very small variation over the interval, and is, essentially, constant with a value of $g(\psi^*)$. With this approximation, (3.50) can be rewritten as

$$\int_0^\infty g(\psi) f(\psi) d\psi \approx g(\psi_A) \int_0^{\psi_1} f(\psi) d\psi + g(\psi_B) \int_{\psi_1}^{\psi_2} f(\psi) d\psi + \dots + g(\psi_Y) \int_{\psi_{24}}^\infty f(\psi) d\psi \quad (3.51)$$

But, by the way that we chose $\psi_1, \psi_2 \dots \psi_{24}$, each integral has an area exactly equal to 1/25. Thus we have

$$\int_0^\infty g(\psi) f(\psi) d\psi \approx \frac{1}{25} [g(\psi_A) + g(\psi_B) + \dots + g(\psi_Y)] \quad (3.52)$$

Our numerical integration algorithm can be summarized as

1. Divide the probability density function $f(\psi)$ into X equal areas. Determine the midpoint of each area.
2. For each of the area midpoints, use a simulation program to determine $g(\psi^*)$.
3. The $E[ASN]$ is approximated by $\frac{1}{X} \sum_{i=1}^X g(\psi_i^*)$

A major disadvantage of the numerical integration technique is that we do not have bounds on the estimator (or its error) as we did with the Monte Carlo method.

³ Note that all of the midpoints are readily computed except for ψ_Y . Since the 25th integral is improper, technically we should have $\psi_Y = \infty$. To resolve this, we truncate the tail of the 25th integral when it goes below a certain very small value, perhaps 1×10^{-5} , and then use the mean value theorem for integrals to find an "effective" value for ψ_Y .

Most numerical integration schemes divide the abscissa into equal areas, and thus have well known error bounds [Ref. 23 : pp. 285-312]. In our technique the abscissa is specifically *not* divided up into equal areas, precisely so that we may solve each of the 25 integrals in (3.51).

Note that the numerical integration technique is not too unlike the Monte Carlo method. For our numerical integration scheme, the points are biased around the mean of the density function, as we would expect them to be if they were chosen randomly. More precisely, because we divided the density function into a certain number of *equal* areas, each of the ψ^* in the numerical integration scheme are *equally likely*, and thus we can well imagine that they were generated by our random number generator in the Monte Carlo method.

E. RECEIVER PERFORMANCE WITH FADING

We now use the Monte Carlo and numerical integration algorithms to examine the effects of fading on the receiver acquisition schemes. We begin by examining the performance of the TSPRT scheme in a fading channel characterized by $r=1$. Note that this is rather severe fading; there is just as much diffuse signal as direct signal. For our simulation we use $p_0 = p_1 = 0.5$ (2.28), $\text{SNR} = -10\text{dB}$, $\alpha = \beta = 0.01$, and a sequence of length 1023.

1. Expected Results

Consider again the per-chip signal to noise ratio defined in (2.20) and repeated below:

$$\text{SNR} = \frac{A_0^2 T_c}{2N_0} \quad (3.53)$$

Rewriting (2.21) we have

$$\text{SNR} = \begin{cases} \frac{n}{S_{n,1}^2} \left(\frac{\lambda_n}{2\sigma_n^2} \right) & \text{under } H_1 \\ \frac{n}{S_{n,0}^2} \left(\frac{\lambda_n}{2\sigma_n^2} \right) & \text{under } H_0 \end{cases} \quad (3.55)$$

We know that fading changes the quantities λ_n and σ_n^2 in accordance with (3.25). Therefore, under both H_0 and H_1 it is the case that

$$SNR_{FADING} = SNR_{ORIGINAL} \left(\frac{r}{1 + r + \lambda_n' 2\sigma_n^2} \right) \quad (3.56)$$

In other words, fading has the effect of reducing the SNR by a factor of $r/(1 + r + \lambda_n' 2\sigma_n^2)$. Since $\lambda_n' 2\sigma_n^2$ is always greater than zero, this multiplicative factor will always be less than one, and thus the SNR is always effectively reduced below its original value, with the exact degree of reduction dependent upon the value of r .⁴

Consider again our model used in determining decision thresholds, presented in (2.22) and repeated below:

$$\frac{\lambda_n}{2\sigma_n^2} = \begin{cases} \approx n(SNR)(1 - |\gamma| \Delta)^2 & \equiv \frac{\lambda_{n,1}}{2\sigma_n^2} \text{ under } H_1 \\ \lesssim SNR & \equiv \frac{\lambda_{n,0}}{2\sigma_n^2} \text{ under } H_0 \end{cases} \quad (3.57)$$

We first examine H_0 . Without fading, our design anticipates $\lambda_n/2\sigma_n^2$ to be approximately equal to the per-chip SNR. With fading, the *actual* SNR is effectively reduced, so we expect our test to become *overdesigned* (e.g., *overly conservative*) under H_0 . In other words, whereas we accounted for a certain SNR in designing for a low probability of false alarm (see (3.57)), the actual SNR is even less than what we designed for, thus making our correlator output appear even more " H_0 -like". So, we expect fading to actually improve PFA (i.e., make it smaller) and lower the ASN under H_0 . We expect our tests, under H_0 , to perform better as the fading becomes worse.

Having said this, we quickly add that we expect the effects of fading to be only slight under H_0 . By (3.57), a reduction in effective SNR will mean that our test is overdesigned by some fixed amount. Since the two m-sequences are not synchronized under H_0 , we should, ideally, detect only noise after the correlation process in the receiver. Since we have only partial correlation in the decision process, we expect fading to influence the test (as described above), but only to a small extent.

The situation is much different under H_1 . As with H_0 , our design anticipates $\lambda_n/2\sigma_n^2$ to be proportional to the per-chip SNR. However, with fading, our actual effective SNR is less than what our test accounts for in attempting to satisfy the desired probability of detection. With fading, our signal appears less " H_1 -like". So, we expect

⁴ Note that for $r \rightarrow \infty$ the reduction factor becomes equal to one. This is expected, since $r = \infty$ represents no fading.

fading to cause a deterioration in the receiver's probability of detection, PD, for all acquisition schemes.

Although we expect the effects of fading to be slight under H_0 , we expect the effects to be severe under H_1 . Examining (3.57) we conclude that under H_0 the test is overdesigned by a fixed amount equal to the effective reduction in SNR. Under H_1 , however, the effective SNR reduction is made worse as n increases. Specifically, applying (3.56) we see that under H_0 the SNR is effectively reduced by a constant factor of $(r/[1+r+SNR])$. Under H_1 the SNR is effectively reduced by the factor

$$\left(\frac{r}{1+r+n(SNR)(1-|\gamma|\Delta)^2} \right)$$

which becomes smaller as n increases. Our scheme under H_1 becomes successively more underdesigned as n increases.

2. Monte Carlo Simulation Results

A 50 run Monte Carlo TSPRT simulation was performed under fading conditions with $r = 1$. After the data for all runs were collected, it was arranged in order of increasing ψ , the fading random variable. Figure 11 on page 43 displays ASN vs. ψ under H_1 . Figure 12 on page 43 displays PD vs. ψ under H_1 . Recall that $\psi = 1$ amounts to the absence of fading. As ψ increases beyond $\psi = 1$, PD remains high and ASN decreases dramatically. This is expected, since $\psi > 1$ corresponds to a stronger signal (see (3.4)) than expected. Our test is designed to meet PD = 0.99 with no fading; if we have "constructive fading" (i.e., $\psi > 1$) our signal is stronger than what we have designed for and therefore fewer chips are needed in the correlation process to meet the H_1 threshold.

As ψ goes below one, the ASN starts increasing rapidly while PD starts declining. The test runs longer in attempting to satisfy PD = 0.99, but increasingly decides H_0 by mistake. As ψ decreases, the H_1 condition increasingly appears to the receiver to be the H_0 condition. At about $\psi = 0.55$, PD drops below 0.5 and the decision scheme starts deciding H_0 more often than H_1 . As ψ drops below 0.55 the ASN drops as the TSPRT decides on H_0 earlier and earlier.

Figure 13 on page 44 shows $E[ASN]$, given by (3.31), vs. the number of runs. Note that $E[ASN]$ seems to settle down as the number of runs increases. Under the Monte Carlo scheme, $E[ASN]$ would give the exact answer as $n \rightarrow \infty$. Generally, the more runs used, the better the estimator is expected to be.

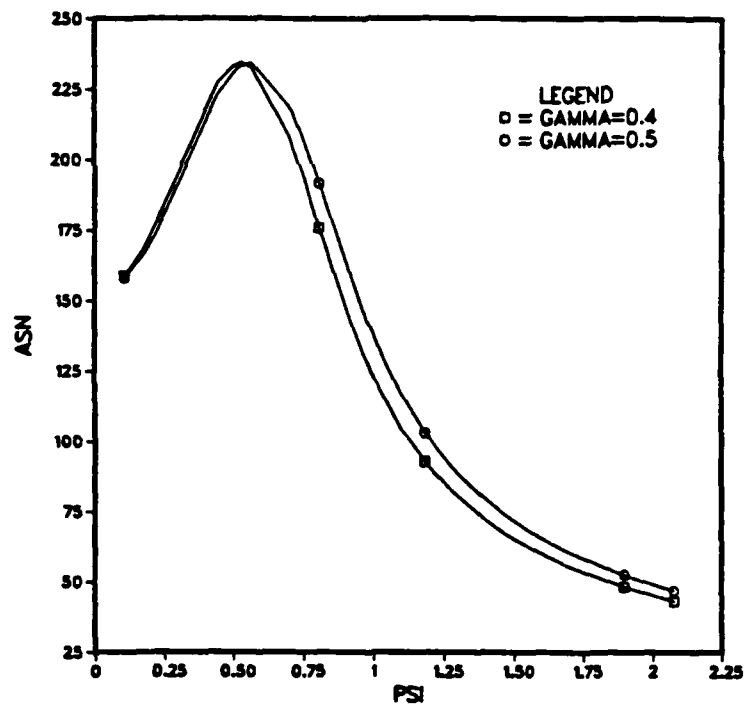


Figure 11. TSPRT ASN vs ψ under H_1 (Monte Carlo)

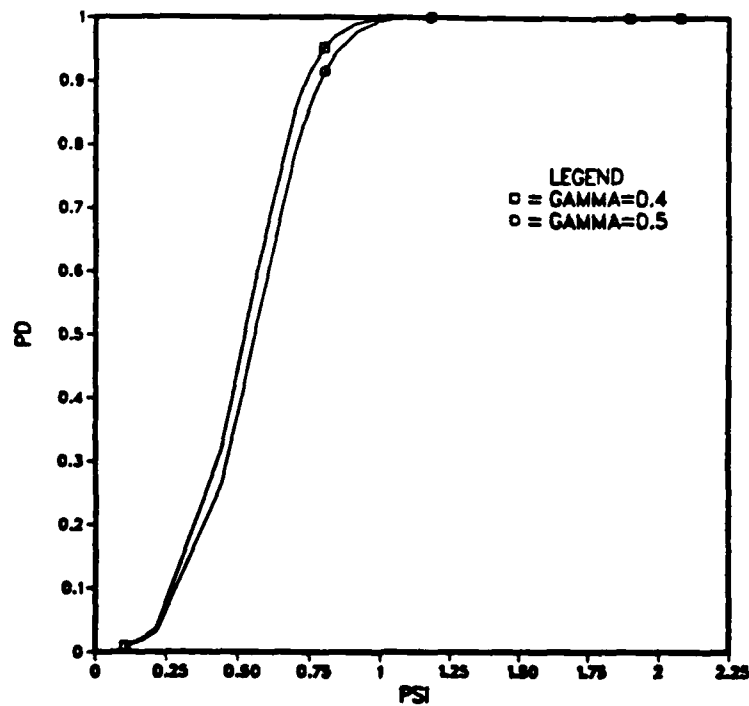


Figure 12. TSPRT PD vs ψ under H_1 (Monte Carlo)

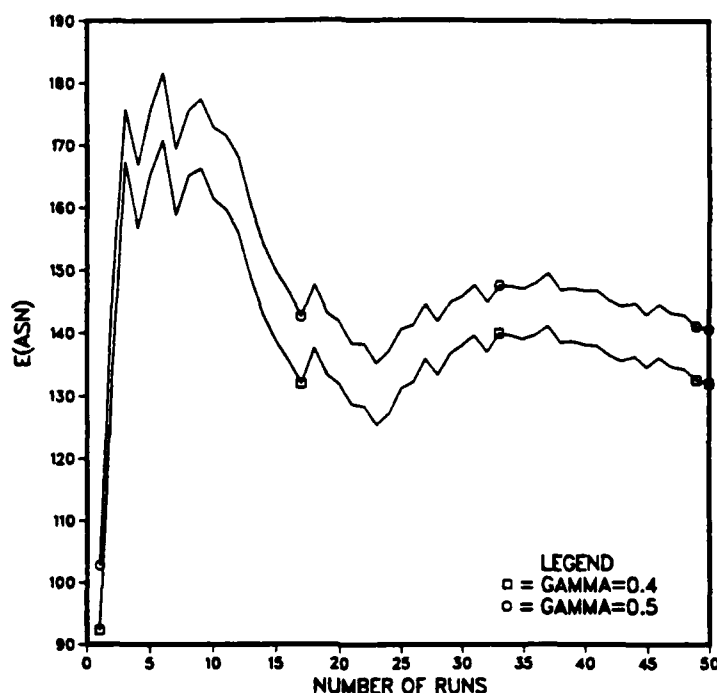


Figure 13. TSPRT $E(\text{ASN})$ vs number of runs: H_1

Figure 14 on page 45 displays ASN vs ψ under H_0 . Notice that the change in ASN over the range of ψ is only about ten chips, a small variation compared to the nearly 200 chip fluctuation seen under H_1 (Figure 11). Figure 15 on page 45 displays PFA vs. ψ under H_0 . Notice that PFA does not go above the design value of $\alpha = 0.01$ until $\psi \approx 1.5$. This confirms our expectation that fading would have only a small effect under H_0 .

How confident can we be that our Monte Carlo results are accurate? Figure 16 on page 46 shows the mean square error of our ASN estimate (under H_1) vs. the number of runs. We expect that as the number of runs increases, the mean square error, V , given by (3.40), will decrease, and we see that this is indeed the case. Ideally, $V \rightarrow 0$ as the number of runs $\rightarrow \infty$. After 50 runs, our Monte Carlo scheme yields $E[\text{ASN}] \approx 140$ with a mean square error ≈ 70 . Obviously, our results after 50 runs do not instill as much confidence as we would have desired.

A confidence interval calculation is instructive. Suppose we want to determine bounds for I with 90% confidence. Using (3.41), we see we require $d = 0.1$. Using (3.48) we find t such that $P[Y > t] = 0.05$ where Y is the standard normal distribution. Consulting a table yields $t = 1.645$. With $\bar{g}_n = 140$ and $V = 70$, (3.43) gives our desired confidence interval:

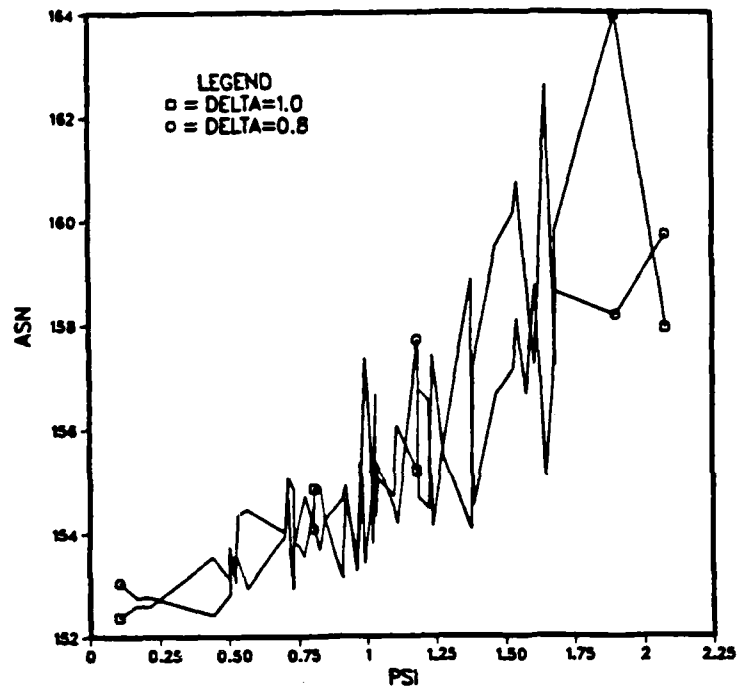


Figure 14. TSPRT ASN vs ψ under H_0 (Monte Carlo)

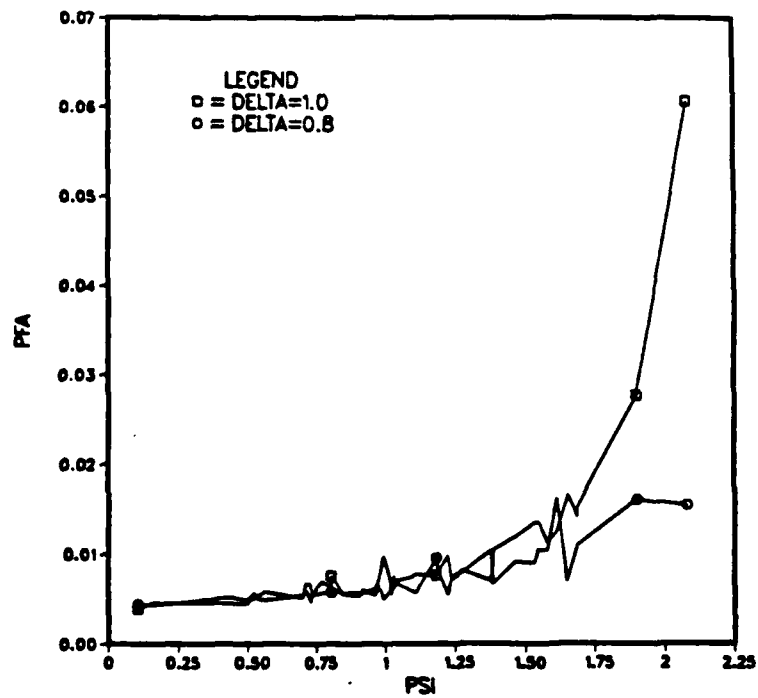


Figure 15. TSPRT PFA vs ψ under H_0 (Monte Carlo)

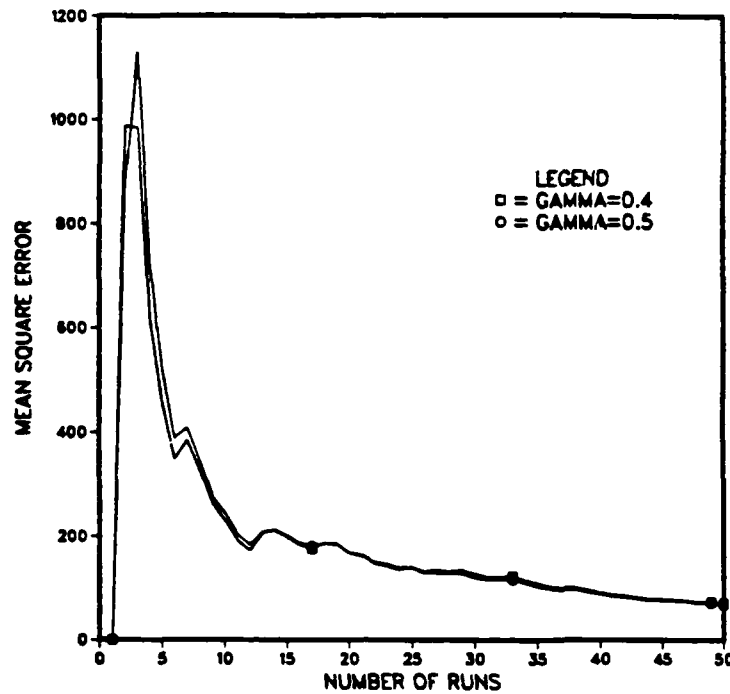


Figure 16. MSE of ASN estimate for TSPRT under H_1

$$P[126 < I < 154] \geq 90\%$$

In other words, we can be assured that $E[ASN]$ under H_1 is within the 28 chip band (126-154) with 90% confidence.

3. Numerical Integration Simulation Results

We reran the TSPRT simulation under $r = 1$ fading conditions by dividing the Ricean density into 25 equal areas. Figure 17 on page 47 shows the ASN vs ψ for H_1 . Note that these results, and the results to follow, closely resemble those obtained through Monte Carlo simulation. Figure 18 on page 47 displays PD vs. ψ under H_1 . Figure 19 on page 48 and Figure 20 on page 48 show ASN and PFA, respectively, vs. ψ under H_0 . Figure 21 on page 49 includes data for PD vs. ψ for points between H_0 and H_1 .

How confident can we be that our numerical integration results are accurate? Recall that there is no precise way to determine confidence intervals for our integration scheme. To gain insight into how stable our results are, we reran the integration scheme but divided the density function into 50 areas instead of 25. We found that, to the nearest chip, the value of $E[ASN]$ under H_1 for 50 areas was identical to the value obtained for 25 areas. Similarly, the value of $E[PD]$ using 50 areas was found to be equal

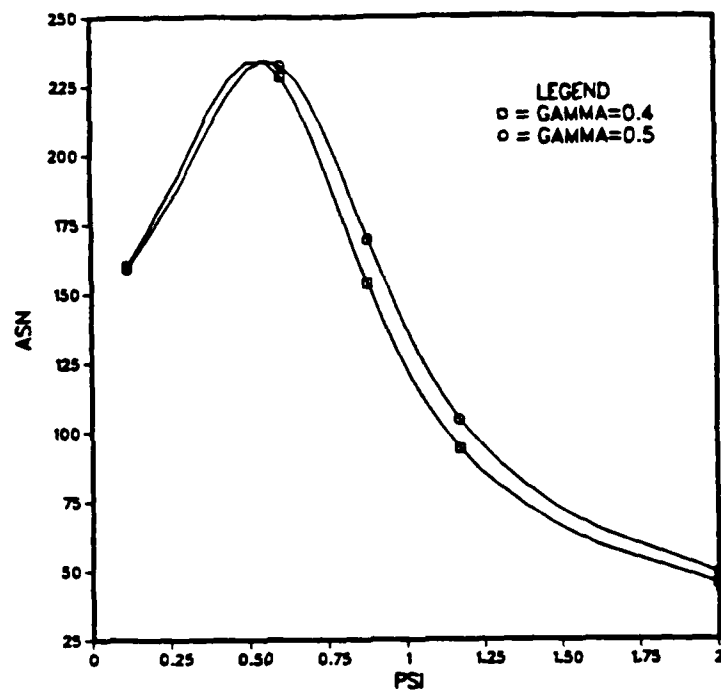


Figure 17. TSPRT ASN vs ψ under H_1 (Numerical Integration)

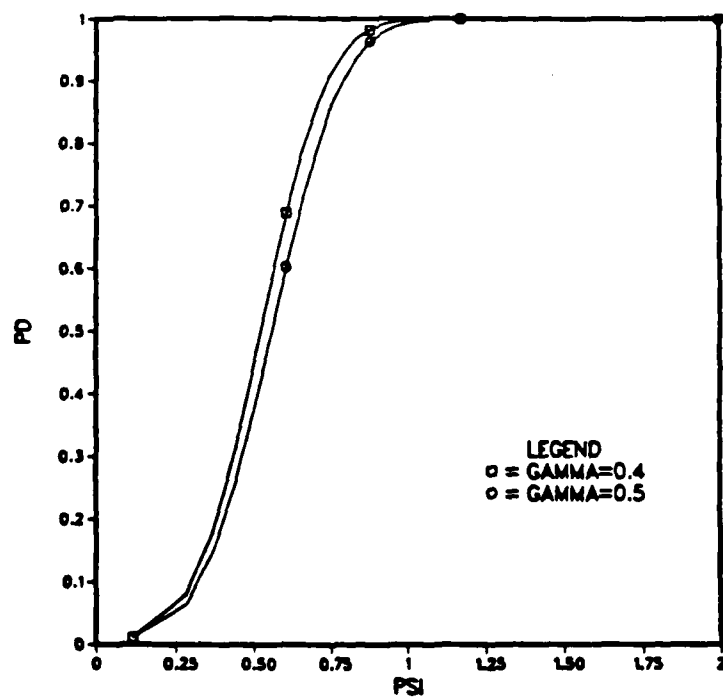


Figure 18. TSPRT PD vs ψ under H_1 (Numerical Integration)

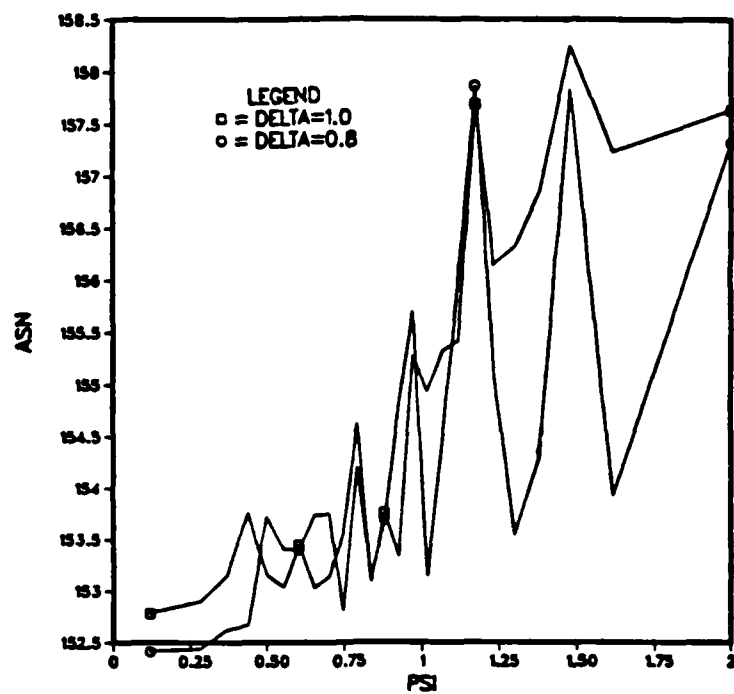


Figure 19. TSPRT ASN vs ψ under H_0 (Numerical Integration)

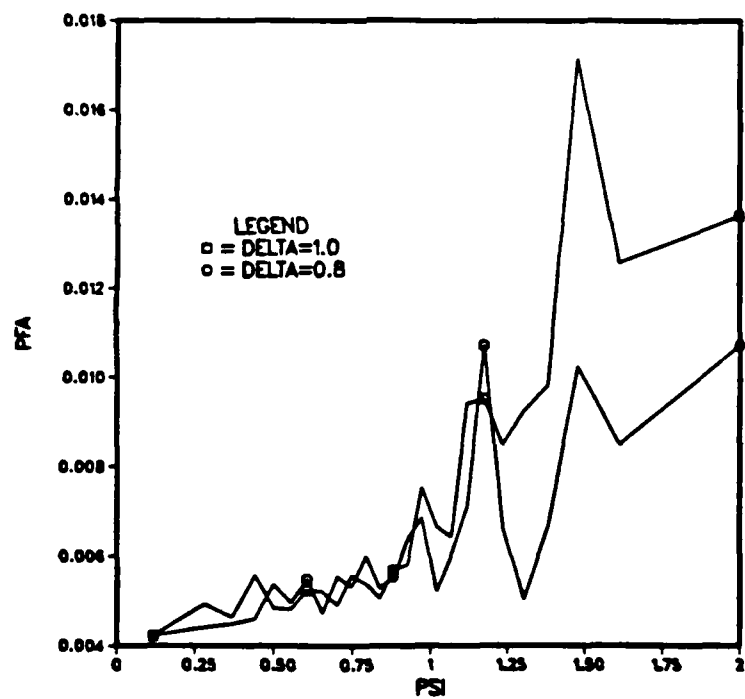


Figure 20. TSPRT PFA vs ψ under H_0 (Numerical Integration)

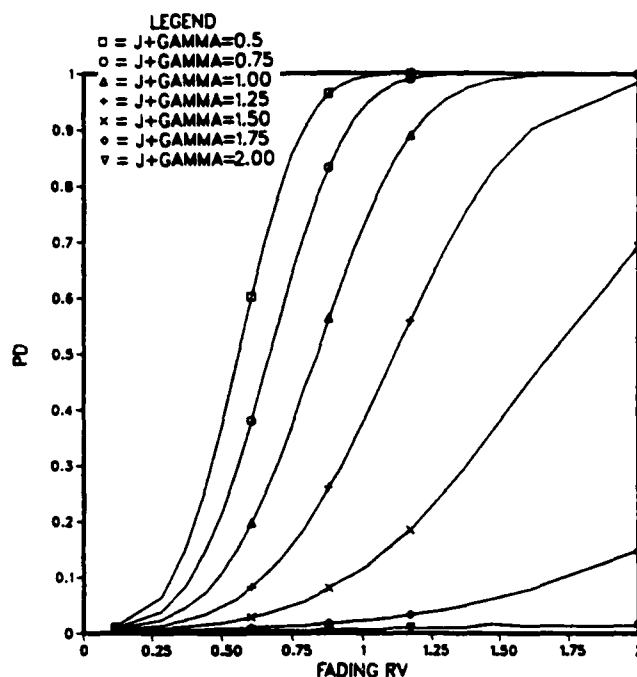


Figure 21. TSPRT PD vs ψ (Numerical Integration)

to the value obtained with 25 areas to within three decimal places. We conclude that our numerical integration scheme using 25 areas is accurate and, at this point, we abandon the Monte Carlo method for all future simulations in this thesis.

We summarize our results for the TSPRT with $r = 1$. Without fading, our expected value for the probability of detection under H_1 is 0.992. With fading, the expected value for the probability of detection dropped to 0.762. Our expected value of ASN increased from 138 without fading to 156 with fading.

Under H_0 our expected value of probability of false alarm decreased from the no-fading value of 0.01 to a value of 0.007 with fading. The ASN under H_0 decreased from 158 without fading to 155 with fading.

Figure 22 on page 50 shows the TSPRT power vs. $|j + \gamma|$ with and without fading. Recall that $|j + \gamma| = 0.5$ corresponds to H_1 while $|j + \gamma| = 2.0$ corresponds to H_0 . Figure 23 on page 50 summarizes the fading results for ASN vs. $|j + \gamma|$.

4. Simulation For Various Fading Conditions

We obtained simulation results for the TSPRT with fading for the additional cases of $r = 0, 10$ and 20 . In all cases the simulation technique consisted of numerical integration with the density function divided into 25 equal areas. Figure 24 on page 51 summarizes the results for ASN vs. $|j + \gamma|$ while Figure 25 on page 51 summarizes the

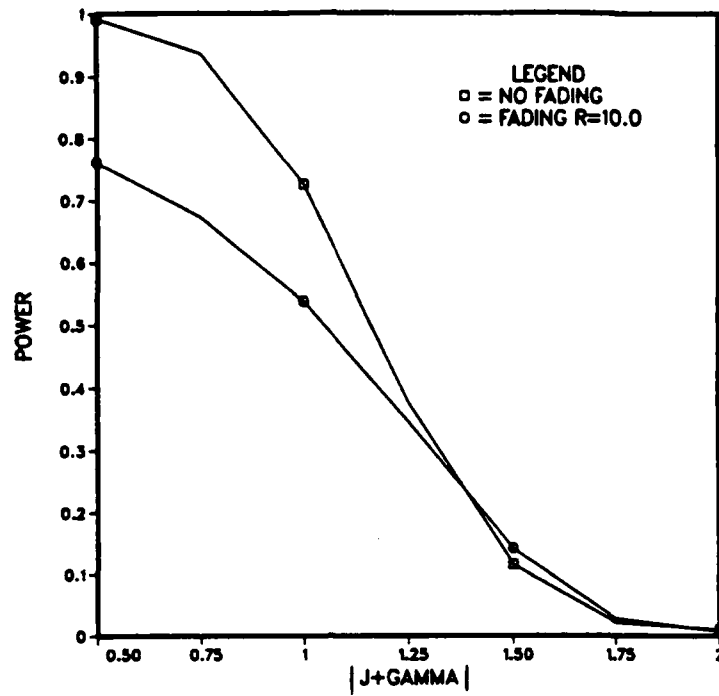


Figure 22. TSPRT Power vs $|j + \gamma|$

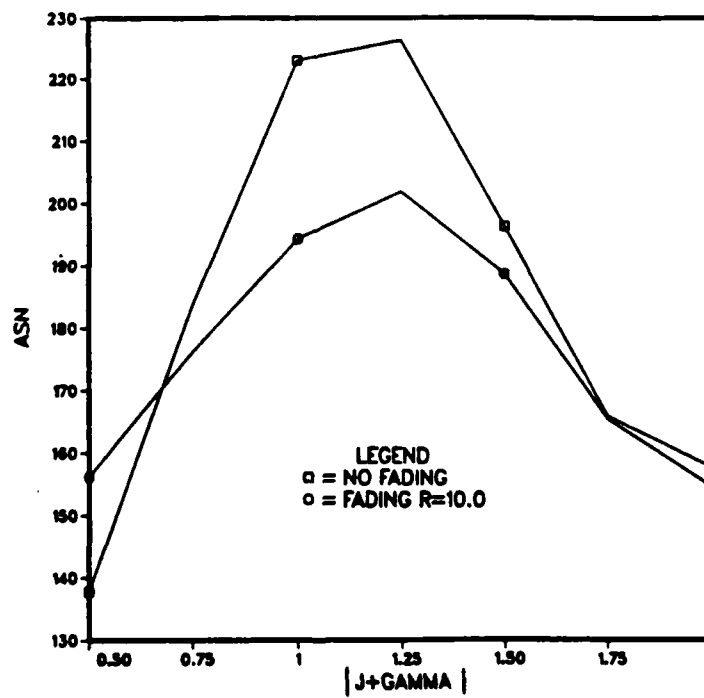


Figure 23. TSPRT ASN vs $|j + \gamma|$

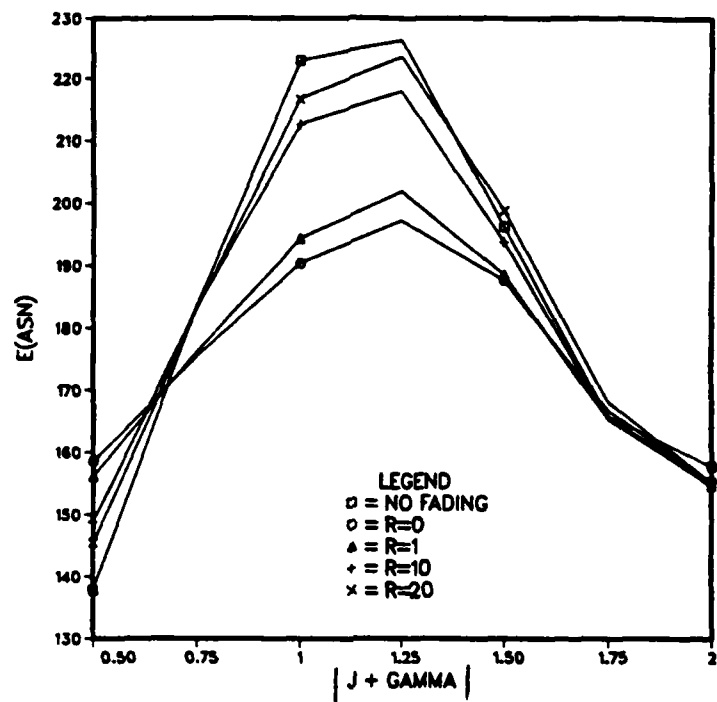


Figure 24. TSPRT ASN vs $|j + \gamma|$ (various r)

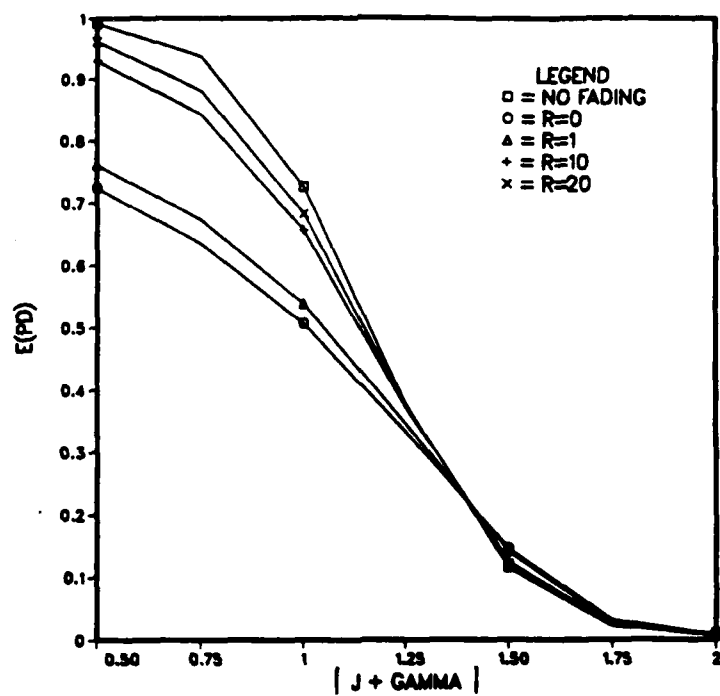


Figure 25. TSPRT PD vs $|j + \gamma|$ (various r)

results for PD vs. $|j + \gamma|$. Note that the test performs better as r increases, as expected. Additionally, it is seen that fading has very little effect on the H_0 hypothesis, as expected.

We then simulated the FSS test under identical conditions for $r=0,1,10$ and 20 . The results for PD vs. $|j + \gamma|$ are summarized in Figure 26. (Note that for the FSS test, the ASN is a constant independent of fading.)

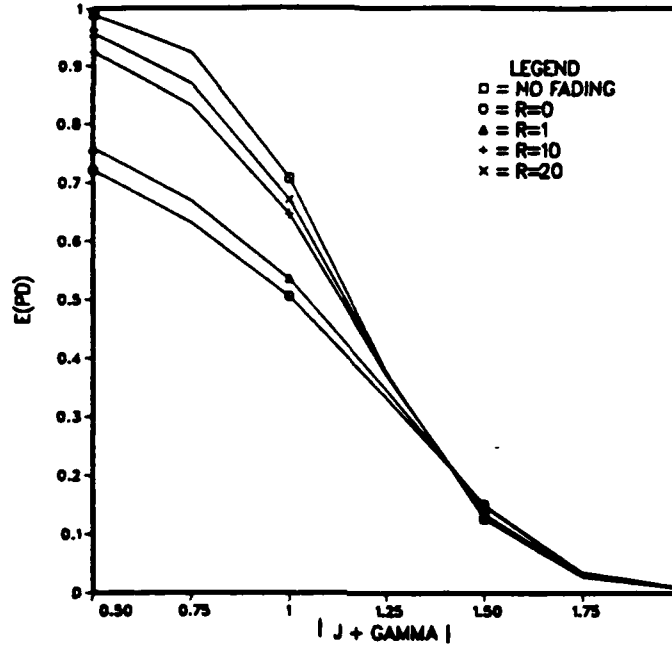


Figure 26. FSS PD vs $|j + \gamma|$ (various r)

F. THRESHOLD ADJUSTMENT

We have seen that fading causes no deterioration of the receiver performance under H_0 , but generates a considerable degradation under H_1 . Under worst possible fading, the Rayleigh case, the probability of detection is reduced from 0.992 to 0.724 for the TSPRT and from 0.990 to 0.721 for the FSS test.

Since we know the exact form of the test statistic's probability density function with fading (given by (3.24)), we can redesign our decision processor to account for the fading. Designating λ_n as $\lambda_{n,1}$ under H_1 , and similarly designating λ_n as $\lambda_{n,0}$ under H_0 , we write the likelihood ratio as

$$\Lambda_n(y_n) = \frac{1 + r + (\lambda_{n,0}/2\sigma_n^2)}{1 + r + (\lambda_{n,1}/2\sigma_n^2)} \exp \left[r \left(\frac{\lambda_{n,0}/2\sigma_n^2}{1 + r + \lambda_{n,0}/2\sigma_n^2} - \frac{\lambda_{n,1}/2\sigma_n^2}{1 + r + \lambda_{n,1}/2\sigma_n^2} \right) \right]$$

$$\begin{aligned}
& \cdot \exp \left[(1+r) \frac{y_n}{2\sigma_n^2} \left(\frac{1}{1+r+\lambda_{n,0}/2\sigma_n^2} - \frac{1}{1+r+\lambda_{n,1}/2\sigma_n^2} \right) \right] \\
& \cdot \frac{I_0 \left(\frac{2\sqrt{r(1+r)(\lambda_{n,1}/2\sigma_n^2)(y_n/2\sigma_n^2)}}{1+r+\lambda_{n,1}/2\sigma_n^2} \right)}{I_0 \left(\frac{2\sqrt{r(1+r)(\lambda_{n,0}/2\sigma_n^2)(y_n/2\sigma_n^2)}}{1+r+\lambda_{n,0}/2\sigma_n^2} \right)} \quad (3.58)
\end{aligned}$$

Regrouping terms, we can isolate all terms with the quantity $y_n/2\sigma_n^2$ into a test statistic and compare this test statistic to two thresholds:

$$\begin{aligned}
& (1+r) \frac{y_n}{2\sigma_n^2} \left[\frac{1}{1+r+\lambda_{n,0}/2\sigma_n^2} - \frac{1}{1+r+\lambda_{n,1}/2\sigma_n^2} \right] \\
& + \ln \left(\frac{I_0 \left(\frac{2\sqrt{r(1+r)(\lambda_{n,1}/2\sigma_n^2)(y_n/2\sigma_n^2)}}{1+r+\lambda_{n,1}/2\sigma_n^2} \right)}{I_0 \left(\frac{2\sqrt{r(1+r)(\lambda_{n,0}/2\sigma_n^2)(y_n/2\sigma_n^2)}}{1+r+\lambda_{n,0}/2\sigma_n^2} \right)} \right) \quad (3.59)
\end{aligned}$$

$$\begin{cases} \geq \ln A + \ln \left(\frac{1+r+\lambda_{n,1}/2\sigma_n^2}{1+r+\lambda_{n,0}/2\sigma_n^2} \right) + r \left(\frac{\lambda_{n,1}/2\sigma_n^2}{1+r+\lambda_{n,1}/2\sigma_n^2} - \frac{\lambda_{n,0}/2\sigma_n^2}{1+r+\lambda_{n,0}/2\sigma_n^2} \right) & \text{say } H_1 \\ \leq \ln B + \ln \left(\frac{1+r+\lambda_{n,1}/2\sigma_n^2}{1+r+\lambda_{n,0}/2\sigma_n^2} \right) + r \left(\frac{\lambda_{n,1}/2\sigma_n^2}{1+r+\lambda_{n,1}/2\sigma_n^2} - \frac{\lambda_{n,0}/2\sigma_n^2}{1+r+\lambda_{n,0}/2\sigma_n^2} \right) & \text{say } H_0 \end{cases}$$

where, as before, $\lambda_{n,1}/2\sigma_n^2$ and $\lambda_{n,0}/2\sigma_n^2$ are given by (3.57) and A and B are determined by the procedure detailed in Section C of Chapter II. For the FSS test $A = B$ and our two thresholds are equal.

Note that the new test statistic, as well as the thresholds, depends critically on the value of r . Adjusting the thresholds to counter the fading effect requires that we somehow measure quantitatively the value of r prior to profiting from the updated design. In other words, if we suspect that fading has caused our receiver performance to deteriorate, we must determine the degree of fading, r , prior to adjusting the thresholds. A detailed study of techniques that may be used to determine the fade depth is beyond the

scope of this thesis, but the following methods, individually or in combination, might be considered:

1. If the fading is due to multipath, we can employ a highly directive antenna to scan the received signal strength in azimuth and elevation. Using this information, we determine how much of the received signal power is arriving from directions outside the design path.
2. If the fading is due to ionospheric effects, we can use sounding techniques to measure the layer alterations, or we may avail ourselves of the historical information already gathered concerning the ionosphere. Figure 27 on page 55 [Ref. 24] shows the fade depth in dB for a typical HF channel at mid-latitudes, over a 200km path. We can use (3.56) to convert this information into an estimate for the value of r .
3. Prior to initiating communications, we can conduct an extensive channel survey to measure and assemble information on signal performance under varying conditions (e.g., rain, magnetic storms, sun spot activity, channel traffic, etc). If similar conditions arise again once the channel is in use, we will have a baseline estimate of the amount of fading to expect.

Howsoever determined, we assume for the remainder of this discussion that we are able to gauge the fade depth and calculate r .

Figure 28 on page 56 displays the degree to which the thresholds in (3.59) vary with fading for the TSPRT with $\text{SNR} = 0.1$, $\text{length} = 1023$ and $p_0 = p_1 = 0.5$. Note that the "original" thresholds correspond to $r \rightarrow \infty$, or no fading. As r decreases, fading becomes worse and the test runs for a longer time prior to truncation. This makes intuitive sense when one considers that fading effectively reduces the received signal's SNR, and so we must integrate more m-sequence chips to accumulate enough signal to make a decision. Note that if the fading becomes too severe, the TSPRT truncation point becomes greater than the sequence length.

Figure 29 on page 57 shows three typical test statistic trajectories and their interaction with the thresholds for the case of $r = 10$. Note that for the H_1 sample path with $\psi = 0.46$ (severe fading) the decision processor runs all the way to test truncation prior to resolving between the two hypotheses. Note that as the fading worsens from $r = \infty$, 100, 20 to 10 the test truncation increase from 306 to 332, 461 and 710 chips respectively. The truncation point for $r = 1$ (not shown) is an extraordinary length of 15,221 chips.

We reran the TSPRT with simulated fading, but with the thresholds adjusted in accordance with (3.59), for the cases of $r = 10$ and $r = 20$. As before, $\text{SNR} = 0.1$, $\alpha = \beta = 0.01$, p_0 and $p_1 = 0.5$ and the sequence length is 1023. The resultant power vs. $j + \gamma$, the phase offset, is displayed in Figure 30 on page 58. The power curves with no adjustments to the thresholds, presented in Figure 25, are repeated for purposes of

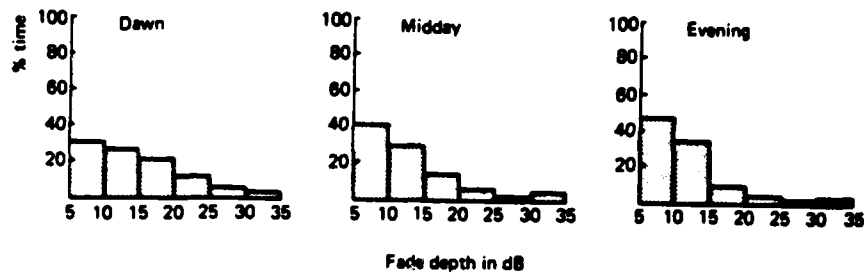


Figure 27. Fade Depth For Typical HF Link

Source: Maslin, N., *HF Communications: A Systems Approach*, Plenum Press, 1987

comparison. We see that the TSPRT is fully restored; adjusting the thresholds raises the probability of detection to greater than 0.99. Threshold adjustment also retains the probability of false alarm below 0.01, as desired.

While we may take satisfaction in restoring the TSPRT power, there is a severe price to be paid in employing our decision processor. Figure 31 on page 59 displays the TSPRT average sample number (the number of chips the decision circuit must correlate prior to reaching a decision) vs $|j + \gamma|$. Note that the ASN has increased drastically.

Recall that among all uncertainty phases to be tested in an m-sequence period, the synchronization condition occurs only once, the condition $0.5 < |j + \gamma| < 2.0$ occurs at one or a few uncertainty phases (depending on the value of Δ), and the non-synchronization condition (H_0) occurs at all the remaining phases. For our TSPRT simulation, the condition $0.5 < |j + \gamma| < 2.0$ occurs at three uncertainty phases while the H_0 condition occurs at 2042 uncertainty phases. To reduce the amount of time taken to acquire the incoming m-sequence it is critical that the ASN be minimized under H_0 . With no adjustments to the thresholds, $E[ASN]$ under H_0 is 158 chips. When the thresholds are adjusted, $E[ASN]$ increases to 215 chips for $r = 20$ and 308 chips for $r = 10$. In summary, the price to pay for meeting the probability of detection is a large increase in $E[ASN]$ under H_0 . Similar results were observed for the FSS test: the test length increased, favoring the probability of detection while sacrificing acquisition time. For $\alpha = \beta = 0.01$ and $SNR = -10\text{dB}$, the required FSS sample size increases from 258 (for no fading) to 361, 509, 6617 and 8944 when fading is characterized by $r = 20, 10, 1$ and 0 , respectively.

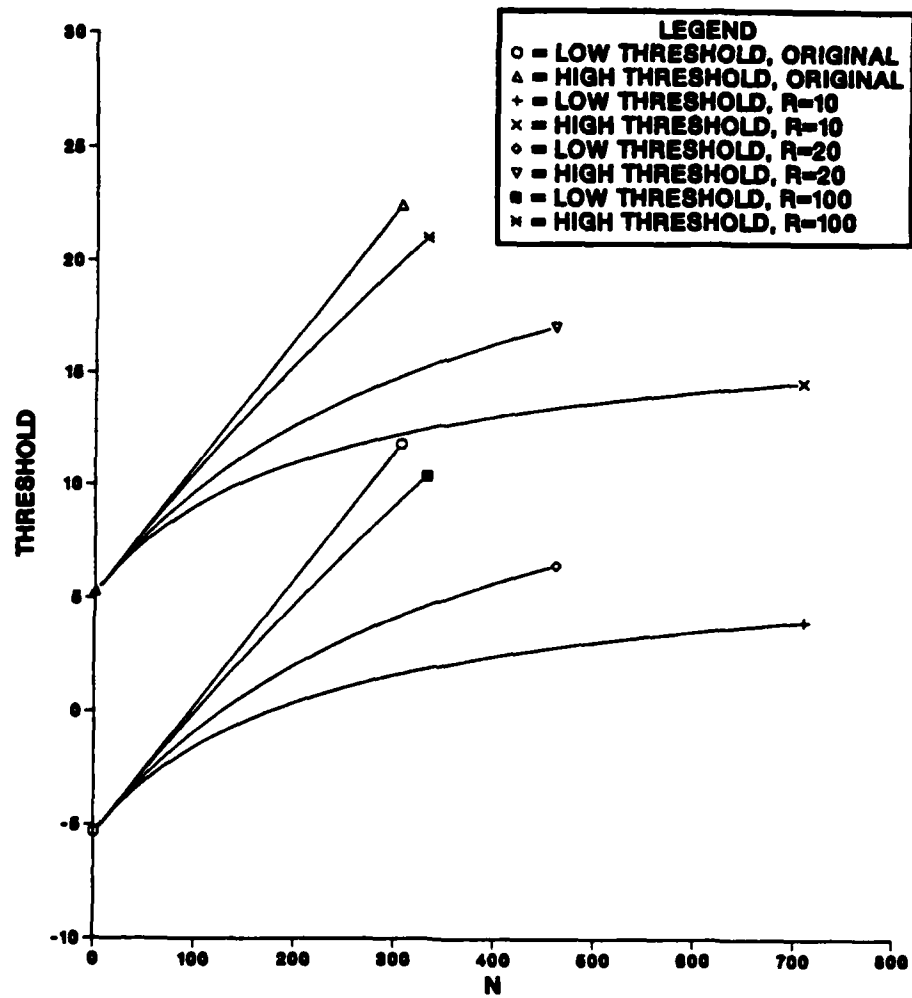


Figure 28. TSPRT Thresholds Adjusted For Fading

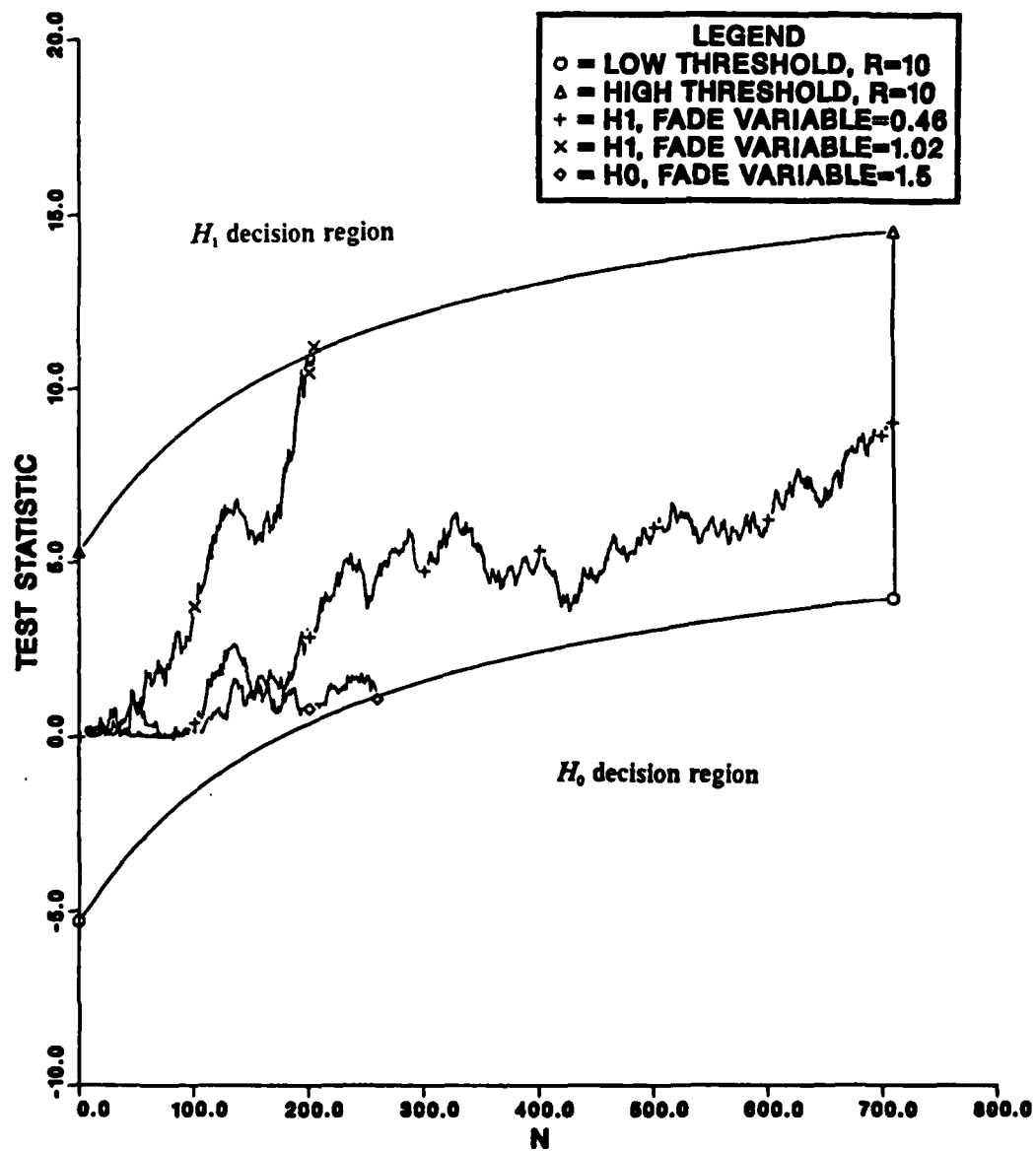


Figure 29. Typical Trajectories For $r = 10$

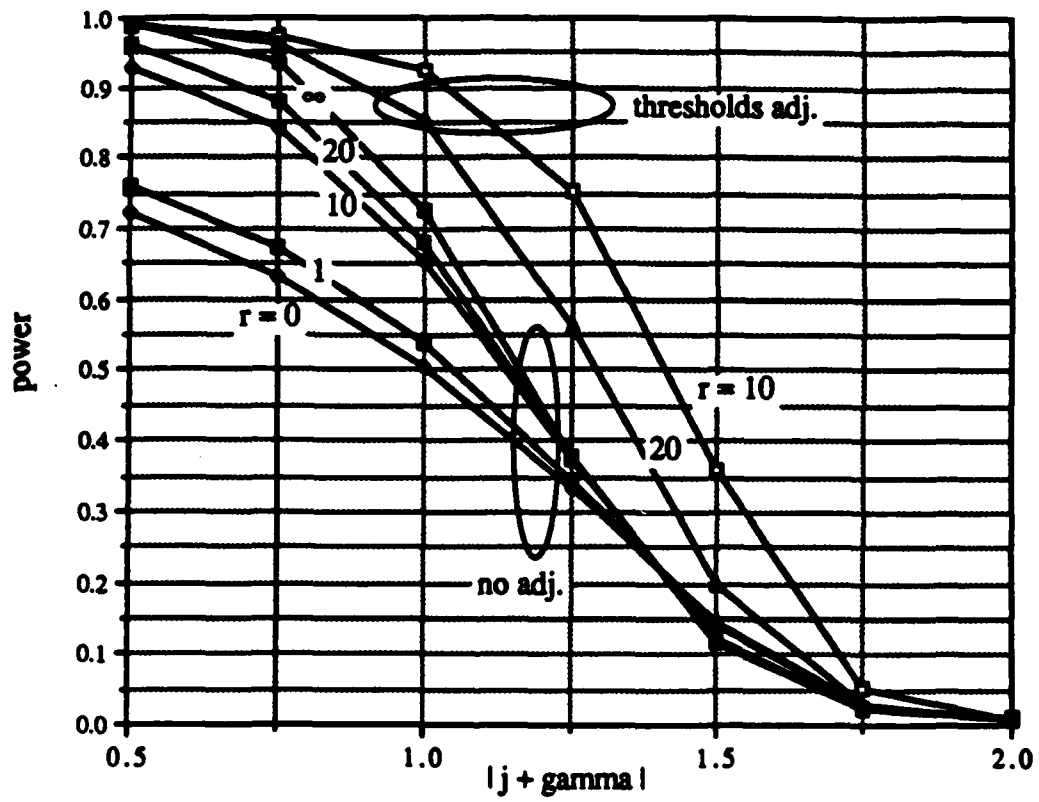


Figure 30. TSPRT PD Fading Performance

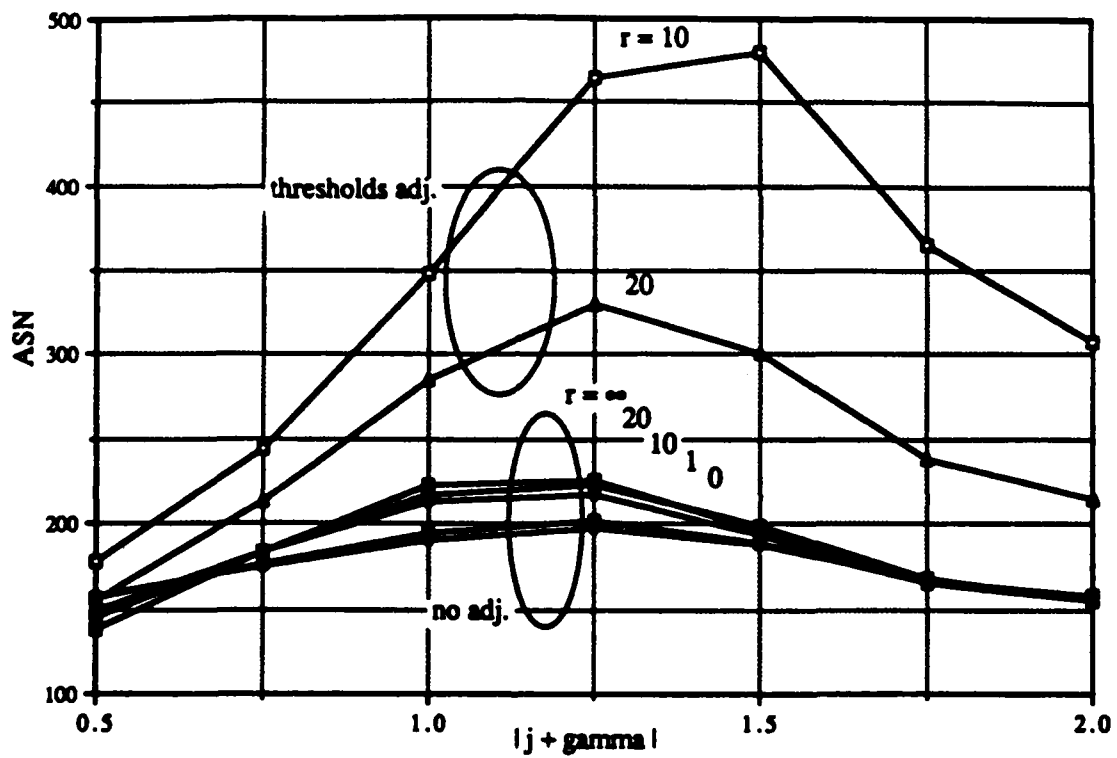


Figure 31. TSPRT ASN Fading Performance

IV. PERFORMANCE WITH MODULATION

A. GENERAL DESCRIPTION OF MODULATION

Up to this point, we have only discussed the acquisition of a communications signal that does not carry data. All the prior analysis assumed that our received signal was

$$r(t) = A_o a(t + i\Delta T_c) \cos(\omega_o t + \theta) + n(t) \quad (4.1)$$

where A_o is the signal amplitude, $a(t)$ is the m-sequence signal waveform with phase $i\Delta T_c$ (i taken to be an integer without loss of generality), T_c is the chip interval, Δ is the amount that the phase of the local m-sequence is altered during the acquisition process, ω_o and θ are the frequency and phase of the carrier, and $n(t)$ is the AWGN.

In this chapter we examine the performance of our detection schemes in the presence of data. By far the most commonly used modulation scheme is binary phase shift keying (BPSK). For our discussion we assume that the data rate is equal to the m-sequence chip rate divided by the sequence length. In other words, we assume we send one bit of information for each period of the m-sequence.

In the BPSK modulation scheme, to communicate a "1" we transmit one period of our chosen m-sequence. To communicate a "-1" we transmit the complement of the m-sequence. The m-sequence can thus be modulated by the data by using a single EXCLUSIVE-NOR gate, with the data bit stream and m-sequence chip stream as inputs. The modulated m-sequence is then multiplied by a high frequency carrier for transmission. Representing the data stream as $d(t)$, our received signal is now

$$r(t) = A_o d(t + i\Delta T_c) a(t + i\Delta T_c) \cos(\omega_o t + \theta) + n(t) \quad (4.2)$$

By examining Figure 3 and (2.6), we see that modulation may present a considerable problem for our receiver. Note that in the correlation process we multiply our locally generated m-sequence by the incoming m-sequence, and then integrate this result. If we transmit the data "1" during the duration of the correlation process, then (2.6) will apply just as before. If we transmit a "-1" during the duration of the correlation process, (2.6) will be preceded by a negative sign, which is subsequently removed by the squaring process in the decision statistic (2.7). No matter what data we send ("-1" or "1"), the receiver performs exactly as it did before, so long as the data remains unchanged during the correlation process.

The problem arises when a data boundary is encountered, and the data changes sign ("1" to "-1" or "-1" to "1") during the correlation process. Under the H_1 condition the correlator output builds up in magnitude until a decision is made. When a data sign change is encountered, the correlator output starts building in the opposite direction, wiping out the previous gains. As an example, suppose we start by sending a "-1". In the H_1 condition, the quantity S_n , given in (2.6), will start building "up" in the negative direction. (The fact that S_n becomes more and more negative is no concern since we are interested in the square of the quantity S_n .) Now, if the data bit changes to "1" during the correlation process, the per-chip S_n becomes positive. This positive quantity adds to (and thus diminishes) the previously accumulated negative running total.

It is true that we can utilize our communications scheme in such a way that we first acquire the m-sequence and then, after acquisition, start sending data. Such a method is, in effect, utilizing a data preamble of "1111..." to assist in m-sequence acquisition. We are interested in seeing how our acquisition schemes perform when data is transmitted from the outset, without a "1111..." preamble. Knowledge of performance with modulation is useful for the following reasons:

1. If we decide to transmit a signal without data for a certain fixed period of time in order to assist receiver acquisition, then it is likely that in many instances the receiver will be synchronized for a certain time before actual data transmission commences. Suppose, for example, that we transmit our signal without data for one minute in order to allow the receiver to acquire the m-sequence. Consider that when two m-sequences are synchronized, such a condition can be detected in about one second. Suppose further that it just so happens that, by chance, the receiver's m-sequence is immediately in phase with the incoming m-sequence at the start of the acquisition process, i.e., suppose that we are immediately in the H_1 state and we acquire in about one second. Then, in such a scenario, we "waste" 59 seconds; our receiver is "ready to receive" but no data is sent for a considerable time. If, on the other hand, we send data immediately, we expect that it may take longer to acquire but, once the receiver is synchronized, data transfer immediately commences. Now, it is true that if we send data immediately, some data sent before acquisition might be lost, and the time needed to acquire the signal might be greater. On the other hand, more data might be sent because we do not "waste" channel time. Knowledge of the performance with modulation may help us resolve this tradeoff.
2. The spreading waveform $a(t)$ is periodic and therefore is characterized by a line spectrum in the frequency domain. Multiplying $a(t)$ by the cosine carrier also results in a quantity that has a line spectrum in the frequency domain. If we are trying to conceal our communications from a potential adversary, we would prefer that our signal had a continuous spectrum instead of a more easily noticeable line spectrum. For this reason, in secure DSSS communications, it is not advisable to send preambles of all ones or all zeros. Once we start sending data, the random (i.e., non-periodic) nature of the data results in a continuous spectrum for our

communications signal. We conclude that sending data immediately affords us a more secure communication signal than one that begins with a "1111..." preamble.

3. No communications scheme is perfect, and we can expect that our scheme will, at some time, lose synchronization after data transfer has begun. Our system may be operating for a period of time, after which the locally generated and received m-sequences fall out of synchronization and make data recovery impossible. Causes for *losing lock* include drift in the oscillators in the transmitter and/or receiver, equipment failure and operator error. After such an incident has occurred, we must go through the acquisition process all over again. During such an event the transmitter will (most likely) not know that the receiver has lost lock, and so the transmitter will not cease sending data. The receiver will have to reacquire in the presence of data modulation.

B. TEST STATISTIC DENSITY FUNCTION

We now determine the probability density function of the receiver's test statistic in the presence of data modulation. We make one important presupposition: we assume that in every case the number of chips used by the decision processor is less than the m-sequence length. This is the same as saying that we assume that we will encounter *at most* one data boundary during the correlation process.

We assume that the test statistic under H_0 is unchanged when data modulation is introduced. When the sequences are out of phase we have, essentially, only noise at the correlator outputs and so we do not expect modulation to have much of an effect. Our confidence in this assumption is further reinforced by the fading simulation results which showed that fading had virtually no effect on the H_0 hypothesis. Under H_0 , with modulation, the pdf is still given by (2.8) with $\lambda_n/2\sigma_n^2 \approx \text{SNR}$. Similarly, the cumulative distribution function is still given by (2.23).

We now consider the H_1 hypothesis. When we transmit data there will be a data modulation boundary every N chips, where N is the number of chips in one period of the m-sequence. We assume that the location of the data boundary is uniformly distributed over the m-sequence length. Specifically, if we start the receiver correlation process at $t = 0$, then, if i is the number of chips integrated *prior* to encountering the data boundary, i is a discrete uniform random variable on $\{0, 1, 2, \dots, N-1\}$.

For each decision scheme, the correlation process will run for a certain (unknown) number of chips prior to exceeding a decision threshold. Let the number of chips in the integration interval for a particular test be n . In other words, our correlator runs from $t = 0$ and makes a decision at $t = nT_c$. The probability that the modulation boundary will *not* appear in the integration interval is thus $1 - n/N$. Now, suppose that a modulation boundary does appear in the integration interval, i.e., at some point in the first

$n - 1$ chips of the correlation process. Let the data prior to the boundary be designated d_0 and the data after the modulation boundary be designated d_1 where $d_0, d_1 \in \{1, -1\}$. Since the data is random, the probability that a sign change occurs at the data boundary is $1/2$: $P(d_0 \neq d_1) = 1/2$. Alternately, the probability that no sign change occurs at the modulation boundary is also $1/2$. Combining this with the fact that a modulation boundary will only appear in the integration interval with probability n/N , we conclude

$$\text{Prob(modulation has no effect)} = \left(1 - \frac{n}{N}\right) + \left(\frac{1}{2}\right) \frac{n}{N} = 1 - \frac{n}{2N} \quad (4.3)$$

When modulation has no effect, the pdf of the test statistic is, as before, given by (2.8) with $\lambda_n^2 2\sigma_n^2$ given by $n(\text{SNR})(1 - |\gamma| \Delta)^2$, in accordance with (2.22).

Now, we examine the one remaining case: a data boundary does occur and, at the boundary, the data changes sign. Reviewing Chapter II, we see that the only term in the description of the receiver decision processor which has changed is S_n , given formerly by (2.6), and previously modeled by (2.19). Using the same model as in Chapter II, we obtain a new expression for S_n in the presence of data modulation:

$$\begin{aligned} S_n = & d_0 \sum_{k=0}^{i-1} c_k [(1 - |\gamma| \Delta) c_k + |\delta| \Delta c_{k+sgn(\gamma)}] \\ & + d_1 \sum_{k=i}^{n-1} c_k [(1 - |\gamma| \Delta) c_k + |\gamma| \Delta c_{k+sgn(\gamma)}] \end{aligned} \quad (4.4)$$

Using the fact that $d_1 = -d_0$, we find

$$S_n = d_0 \left[(2i - n)(1 - |\gamma| \Delta) + |\gamma| \Delta \left\{ \sum_{k=0}^{i-1} c_k c_{k+sgn(\gamma)} - \sum_{k=i}^{n-1} c_k c_{k+sgn(\gamma)} \right\} \right]$$

As before, we assume each of the summations is approximately equal to zero, so it follows that

$$S_n \approx d_0 (2i - n)(1 - |\gamma| \Delta) \quad (4.5)$$

Since (2.21) still applies, we have

$$\frac{\lambda_n}{2\sigma_n^2} \Big|_{\text{data change}} = \frac{1}{n} (SNR) S_n^2 = \frac{1}{n} (SNR) (2i-n)^2 (1-|\gamma|\Delta) = \left(\frac{2i-n}{n} \right)^2 \frac{\lambda_{n,1}}{2\sigma_n^2} \quad (4.6)$$

where $\lambda_{n,1}/2\sigma_n^2$ is equal to $n(SNR)(1-|\gamma|\Delta)^2$.

Combining (4.3) and (4.6) with (2.8), we determine the closed form probability density function for the test statistic under the H_1 hypothesis as

$$f_{Y_n}(y_n | H_1) = \left(1 - \frac{n}{2N}\right) \left(\frac{1}{2\sigma_n^2} e^{-(y_n + \lambda_n) / 2\sigma_n^2} I_0 \left[\frac{\sqrt{\lambda_n y_n}}{\sigma_n^2} \right] \right) + \left(\frac{n}{2N} \right) \cdot \left\{ \frac{1}{n} \sum_{i=0}^{n-1} \left(\frac{1}{2\sigma_n^2} \exp \left(\frac{-(y_n + \left[\frac{2i-n}{n} \right]^2 \lambda_{n,1})}{2\sigma_n^2} \right) I_0 \left[\frac{\sqrt{\left(\frac{2i-n}{n} \right)^2 \lambda_{n,1} y_n}}{\sigma_n^2} \right] \right) \right\} \quad (4.7)$$

The corresponding cumulative distribution function, using the Q function defined in (2.24), is given by

$$F_{Y_n}(y_n) = \left(1 - \frac{n}{2N}\right) \left(1 - Q \left[\sqrt{\frac{\lambda_{n,1}}{\sigma_n^2}}, \sqrt{\frac{y_n}{\sigma_n^2}} \right] \right) + \left(\frac{n}{2N} \right) \left(\frac{1}{n} \right) \sum_{i=0}^{n-1} \left(1 - Q \left[\sqrt{\left[\frac{2i-n}{n} \right]^2 \frac{\lambda_{n,1}}{\sigma_n^2}}, \sqrt{\frac{y_n}{\sigma_n^2}} \right] \right) \quad (4.8)$$

If we choose to work with a normalized random variable z , where $z = y_n / \sigma_n^2$, then (4.7) is rewritten as

$$f_z(z | H_1) = 0.5 \left(1 - \frac{n}{2N}\right) \exp \left[\frac{-(z + \lambda_{n,1} / \sigma_n^2)}{2} \right] I_0 \left[\sqrt{\frac{\lambda_{n,1}}{\sigma_n^2}} z \right] + \frac{1}{4N} \sum_{i=0}^{n-1} \exp \left(\frac{- \left[z + \left(\frac{2i-n}{n} \right)^2 \frac{\lambda_{n,1}}{\sigma_n^2} \right]}{2} \right) I_0 \left[\left(\frac{n-2i}{n} \right) \sqrt{\frac{\lambda_{n,1}}{\sigma_n^2}} z \right] \quad (4.9)$$

We see from (4.9) that the value of z will depend on the number of chips, n , in the integration interval. The expected value of z for a given value of n is

$$E[z] = \left(1 - \frac{n}{2N}\right)(2 + 2n(SNR)(1 - |\gamma| \Delta)^2) + \frac{1}{2N} \sum_{i=0}^{n-1} \left(2 + \left[\frac{2i-n}{n}\right]^2 (2n)(SNR)(1 - |\gamma| \Delta)^2\right) \quad (4.10)$$

The quantity $E[z]$ is plotted vs. n in Figure 32 on page 66

C. RECEIVER PERFORMANCE WITH MODULATION

We now use the numerical integration algorithm to examine the effects of fading on the receiver acquisition schemes. For all simulations we assume $SNR = -10\text{dB}$ and the m-sequence is of period 1023 chips. Additionally, for the TSPRT simulations, we assume $p_0 = p_1 = 0.5$.

1. Expected Results

We expect that modulation will not have a severe effect on the performance (characterized by PFA, PD and ASN) of our receiver acquisition schemes. When a modulation boundary is encountered, half the time it will not represent a data sign change and thus will be transparent to the receiver.

To reach a more quantitative estimate of the effects of modulation, we can examine the mean value of the quantity $\lambda_n/2\sigma_n^2$ considering its value when a data change is encountered (given by (4.6)) and when no data change is encountered (given by (2.22)):

$$\begin{aligned} \frac{\bar{\lambda}_{n,1}}{2\sigma_n^2} &= \left(1 - \frac{n}{2N}\right)n(SNR)(1 - |\gamma| \Delta)^2 \\ &+ \left(\frac{n}{2N}\right) \sum_{i=0}^{n-1} \left(\frac{2i-n}{n}\right)^2 (SNR)(1 - |\gamma| \Delta)^2 \end{aligned} \quad (4.11)$$

This equation can be further simplified as

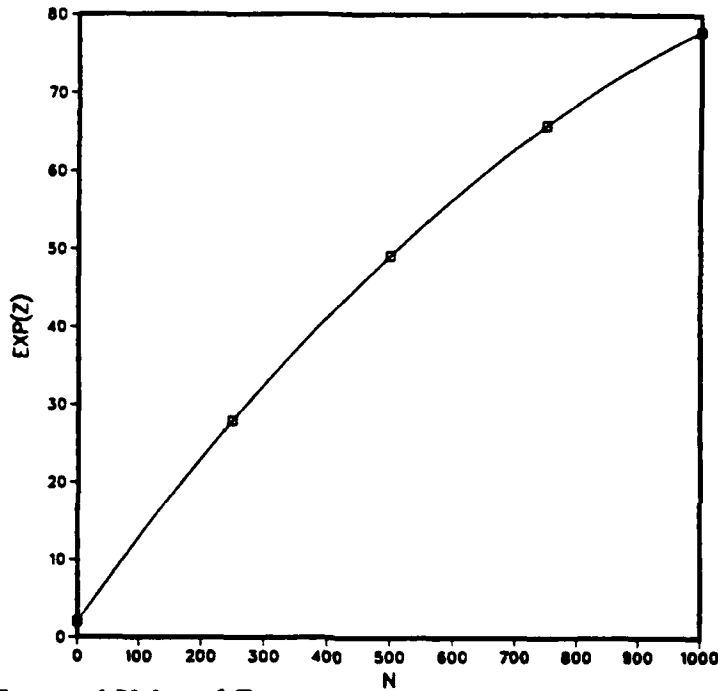


Figure 32. Expected Value of Z vs n

$$\frac{\bar{\lambda}_{n,1}}{2\sigma_n^2} = (SNR)(1 - |\gamma| \Delta)^2 \left[\left(1 - \frac{n}{2N}\right)n + \frac{1}{2Nn} \sum_{i=0}^{n-1} (2i - n)^2 \right]$$

$$\frac{\bar{\lambda}_{n,1}}{2\sigma_n^2} = n(SNR)(1 - |\gamma| \Delta)^2 \left(1 - \frac{1}{3} \frac{n}{N} + \frac{1}{3Nn}\right) \quad (4.12)$$

Comparing (4.12) to the original value of $\lambda_{n,1}/2\sigma_n^2$ given by (2.22), we see that modulation has the effect of reducing the SNR by approximately

$$\left(1 - \frac{1}{3} \frac{n}{N} + \frac{1}{3Nn}\right)$$

For large values of N , the reduction factor is small. In the worst case, $n = N$ and the SNR is reduced by $1/3$, or 1.8dB.

2. Numerical Integration Simulation Results: FSS Test

We examine the effects of modulation on the FSS test assuming the original receiver structure. In other words, without accounting for the fact that modulation alters the test statistic's density function, we will determine the receiver performance.

We first simulate the case where the data boundary is encountered and represents a change in the data sign (-1 to +1 or +1 to -1). If no data boundary is encountered, or if one is encountered and it does not represent a sign change, the modulation is transparent to the receiver and the resulting PD and PFA are as determined in Chapter II. (Recall that the ASN for the FSS test is predetermined by the desired probabilities of false alarm and miss: α and β .)

Let us examine a particular case in detail. Suppose we choose $\alpha = \beta = 0.01$. In this event, the test termination (found by iteratively solving (2.25) and (2.26)) turns out to be 258 chips. We then run the FSS test for each possible value of the data modulation boundary i (i.e., we run the test for $i = 0, 1, 2, \dots, 257$).

Figure 33 on page 68 displays the actual probability of detection vs. the modulation chip boundary under the H_1 hypothesis. We see that if the modulation boundary is encountered early in the integration process, the test can still reasonably recover and yield a relatively high PD (although significantly less than $1 - \beta$). Similarly, if the modulation boundary is encountered very late, the integration may have built up sufficiently to result in a relatively high PD. The worst case occurs when the modulation boundary is midway in the integration process. In such an event, we see from (4.5) that $S_n \approx 0$ and the decision scheme decides on H_0 .

Figure 34 on page 68 shows the actual probability of false alarm vs modulation chip boundary under the H_0 hypothesis. We see that modulation does not seem to have a major effect on PFA. Figure 35 on page 69 shows PD vs the modulation boundary for some points between H_1 and H_0 .

Figure 36 on page 69 summarizes the results for the FSS test power vs $|j + \gamma|$ in a scheme employing data modulation. With no data modulation, PD is 0.990. When data modulation is present PD drops to 0.921. The PFA remains at the same original value, to three decimal places, when modulation is present.

For our given test, with our given sequence length, the modulation effect does not seem very severe. As a matter of fact, simulation shows that for this particular set of conditions (SNR, α , β , N , etc), modulation causes the same performance degradation as fading with $r = 10$.

We now evaluate the performance of the FSS test with a variety of values of α and β . Specifically, we wish to evaluate the effects of modulation for $0.01 \leq \alpha, \beta \leq 0.20$.

To accurately evaluate the effects of modulation on the FSS test, we must first ascertain how the test performs without modulation for all cases of interest; in other

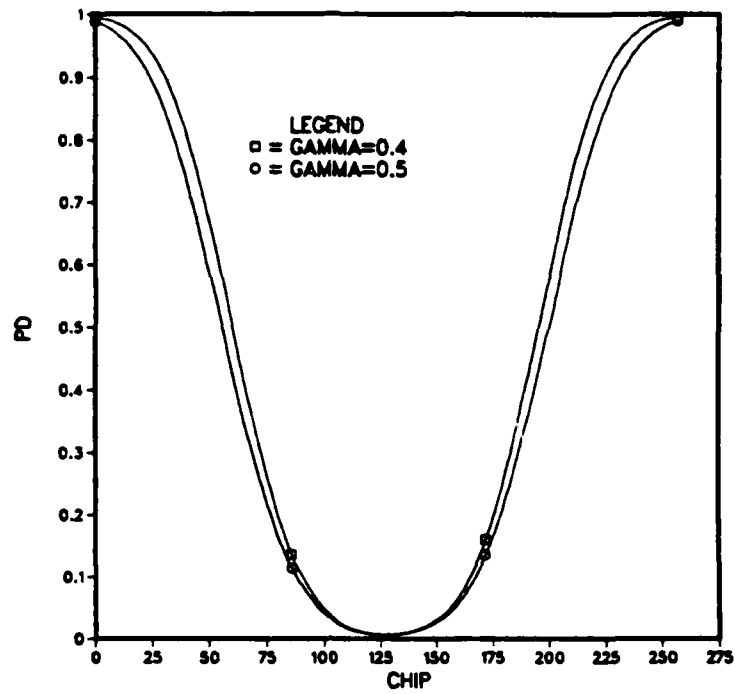


Figure 33. FSS PD vs Modulation Chip Boundary: H_1

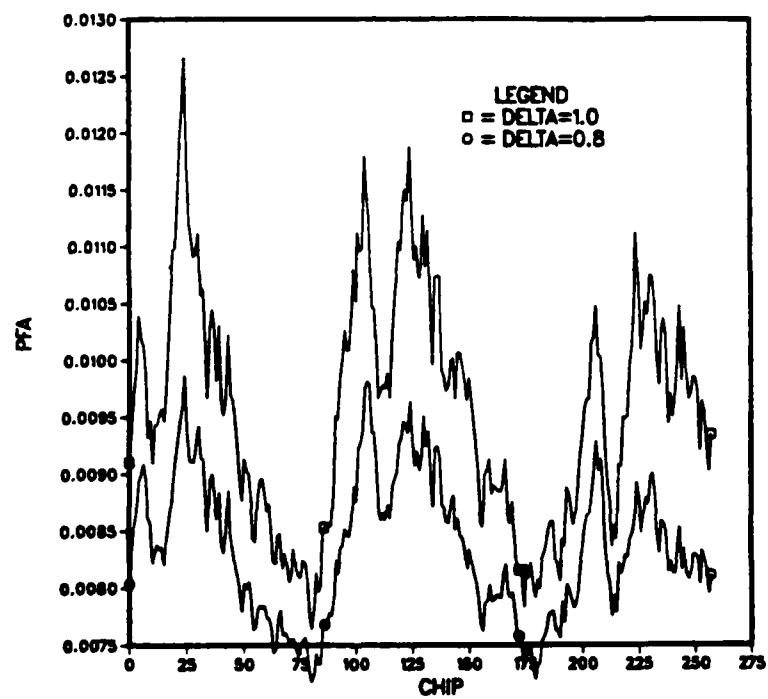


Figure 34. FSS PFA vs Modulation Chip Boundary: H_0

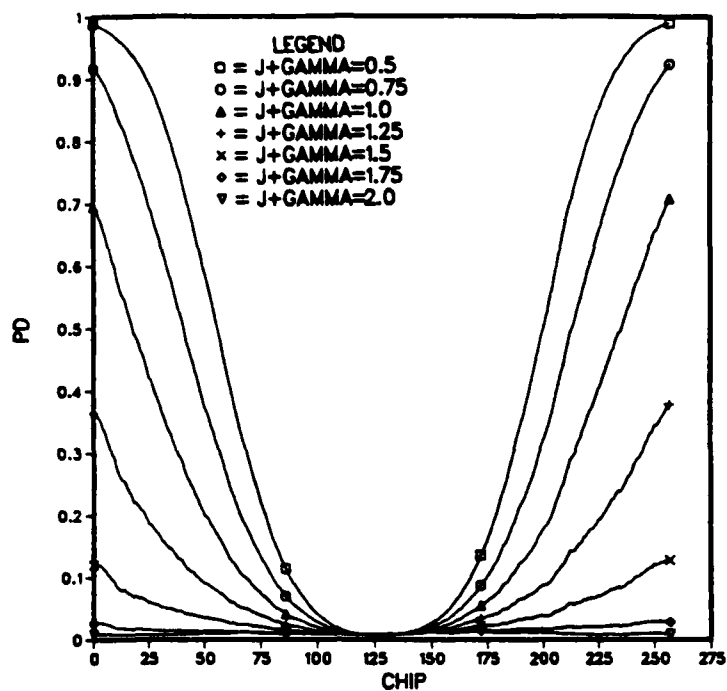


Figure 35. FSS PD vs Modulation Chip Boundary

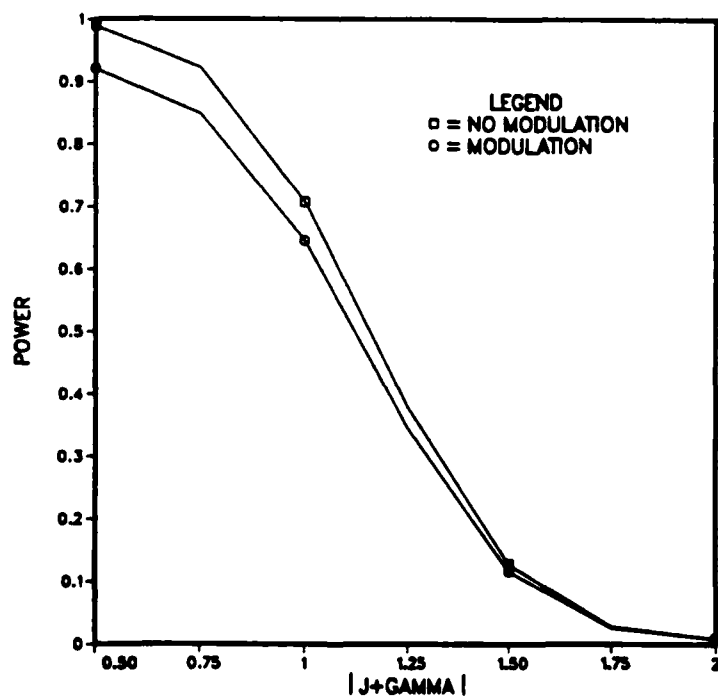


Figure 36. FSS Power vs $|j + \gamma|$: Modulation

words we must determine baseline values of PD, PFA and test termination M for all α and β which satisfy $0.01 \leq \alpha, \beta \leq 0.2$. Figure 37 on page 71 depicts the calculated test termination point (assuming no modulation) vs. the design probability of false alarm (α) and design probability of miss (β). The termination point is calculated using the methods of Chapter II. We note that the termination point seems to be slightly more dependent on β than on α .

Figure 38 on page 72 displays the actual probability of miss, 1-PD, hereafter designated as PM, vs. α and β in the absence of data modulation. The test performs as expected. For a given α , PM is approximately equal to β . For a given β , PM is approximately constant as α varies. Figure 39 on page 73 displays the results for PFA vs. α and β with no modulation.

We now run the FSS test for all cases of interest ($0.01 \leq \alpha, \beta \leq 0.2$) in the presence of data modulation. The test termination point vs. α and β is still as shown in Figure 37 since, at this point, we have not adjusted the thresholds. Figure 40 on page 74 shows the PM as a function of the design miss and false alarm probabilities. Compare this figure to Figure 38. Figure 41 on page 75 shows the variation of PM with α for a fixed value of β . The results can be explained by noting that as we relax our design probability of false alarm, the FSS terminates sooner (see Figure 37), and thus there is less chance that a modulation boundary will be encountered in a given integration interval. So, as we increase α , modulation tends to become more transparent to the receiver, and PM moves closer to β . Figure 42 on page 75 shows the variation of PM with β for a fixed α . The curves are linear (as desired), but they are offset higher than the ideal characteristic, which would be a line of slope 1 passing through the origin. Note that the degree to which the curves are offset from the ideal curve decreases as α increases. As we increase α , the test termination point decreases and modulation becomes more transparent to the receiver.

Figure 43 on page 76 displays the actual probability of false alarm vs α and β . Note that this curve for PFA is almost identical to the results obtained when no modulation is present, shown in Figure 39. Modulation seems to have almost a negligible effect on performance under H_0 .

To summarize our results for the FSS test, we note that modulation worsens the actual miss probability (Figure 38) and has a negligible effect on the actual false alarm probability. The degree to which the actual miss probability (with modulation) exceeds the desired miss probability depends on the design values chosen for α and β . As we increase α and/or β the test terminates earlier (Figure 37) and therefore modulation has

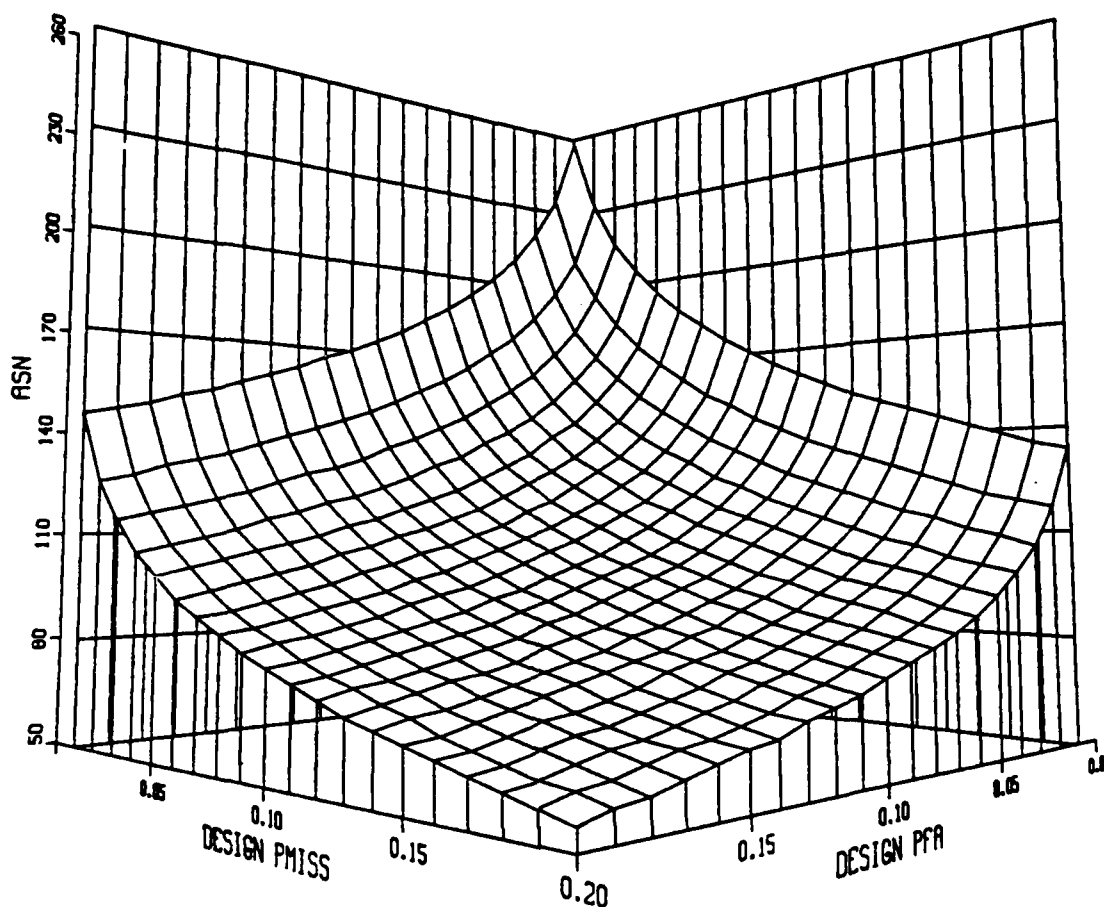


Figure 37. Calculated FSS Termination Point

less of an effect on receiver performance. As we increase α and/or β , the actual probability of miss more closely approaches its design value.

3. Numerical Integration Results: TSPRT

We examine the effects of modulation on the TSPRT, assuming the original receiver structure. Again, we designate i as the data modulation boundary, and we assume i is a uniform discrete random variable on $\{0, 1, 2 \dots N - 1\}$.

When we utilized the numerical integration algorithm to simulate modulation for the FSS test, we essentially divided the density function of i into 1023 equal areas.

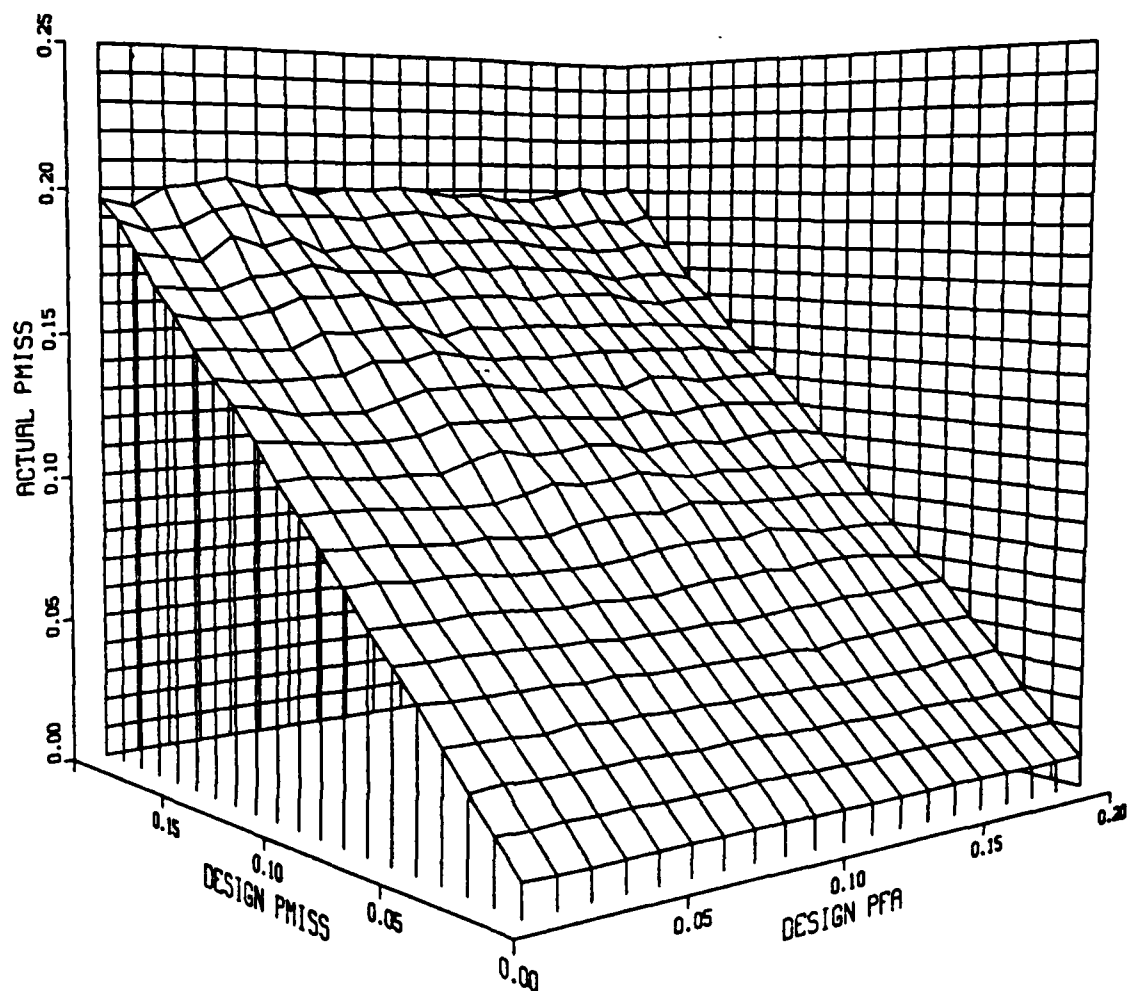


Figure 38. Actual FSS Probability of Miss

Since the TSPRT simulation program takes a considerably longer amount of time to execute, we can not divide the density function of i into so many equal segments. Instead, for all TSPRT simulations, we divide the uniform pdf into X equal areas, such that approximately twenty modulation boundaries may appear at equally spaced points in the integration interval. An example will illustrate our approach. For the TSPRT with $\alpha = \beta = 0.01$, the test is truncated after 306 chips have been integrated (if the test has not already ended). We divide the uniform density of i into 64 equal areas of 16 chips

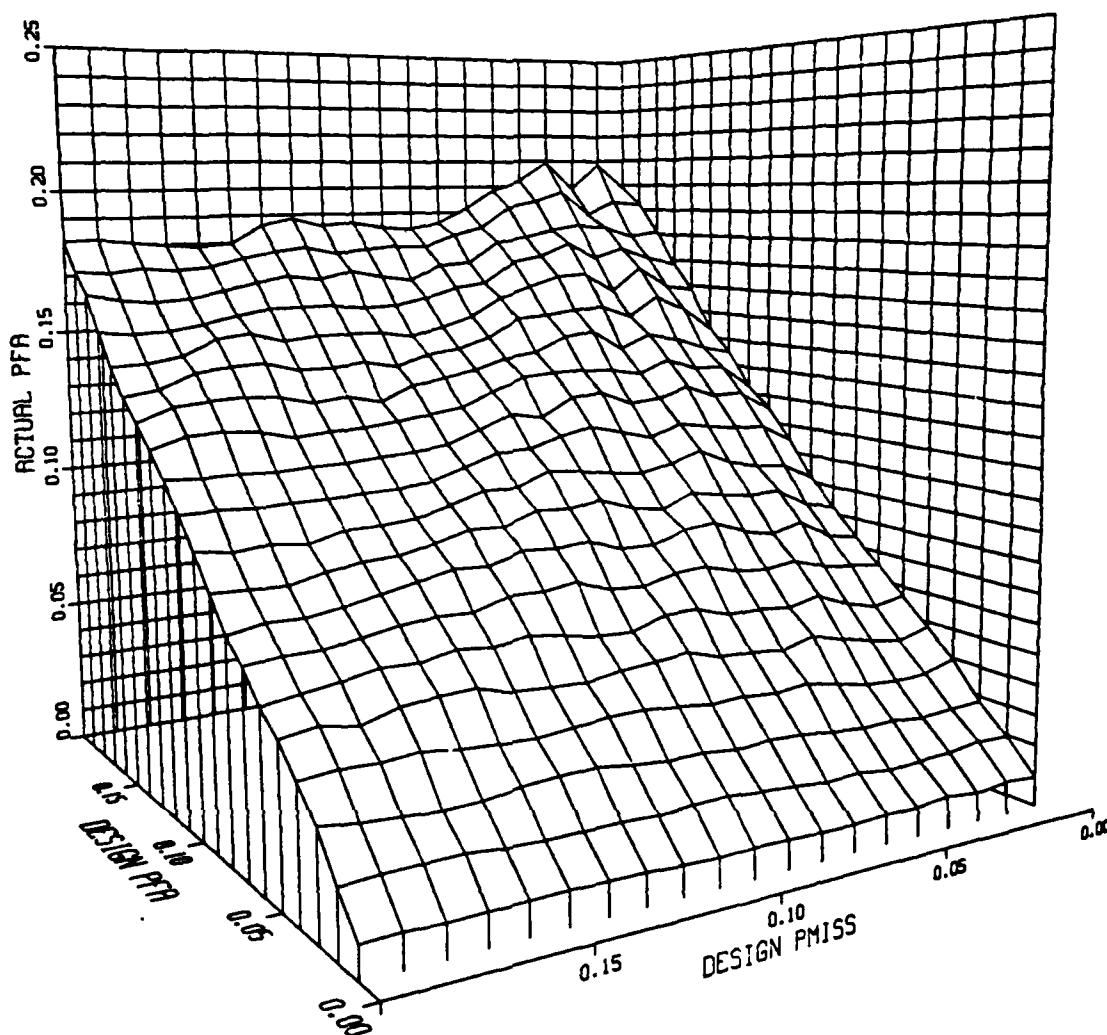


Figure 39. Actual FSS Probability of False Alarm

each.⁵ We then run the TSPRT with each modulation boundary: 8, 24, 40, 56...296; 19 runs in all for this case. Our final results for PD, PM ASN and PFA will represent values averaged over all the modulation boundaries.

⁵ Since the m-sequence length is always an odd number, we either pad the sequence with an extra chip, or subtract a chip, before dividing the pdf of i into a certain number of equal areas, each with an integer number of chips.

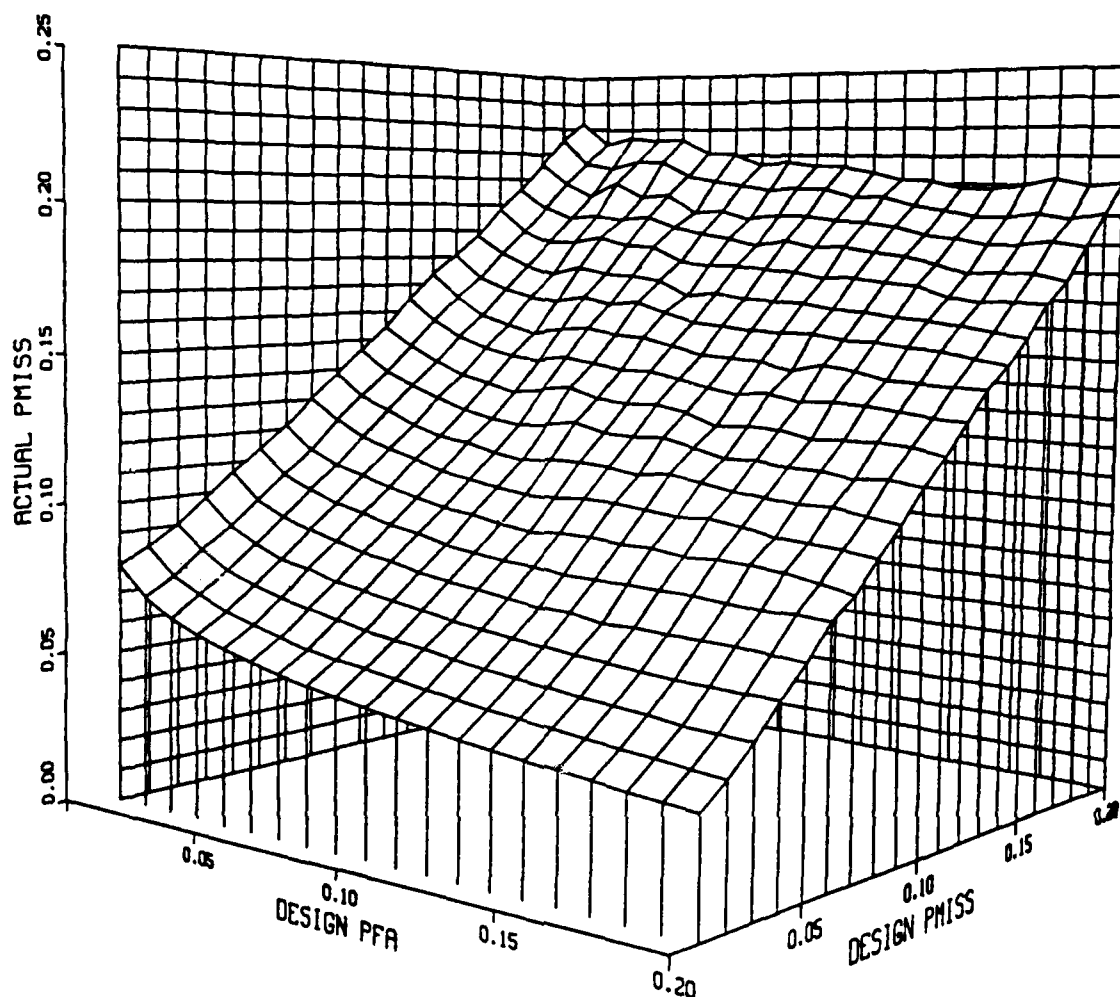


Figure 40. Actual FSS Miss Probability With Modulation

We start by examining in detail an arbitrary case of the TSPRT. We choose $\alpha = 0.04$ and $\beta = 0.14$. In such a communications scheme, we do not want to miss the H_1 condition more than once in seven times and we do not want to falsely declare the H_1 condition more than once in 25 opportunities. The truncation point occurs after 162 chips are integrated. We examine the effects of modulation with no threshold adjustment. The pdf of i is divided up such that there are 21 equally spaced potential modulation data boundaries within the integration interval.

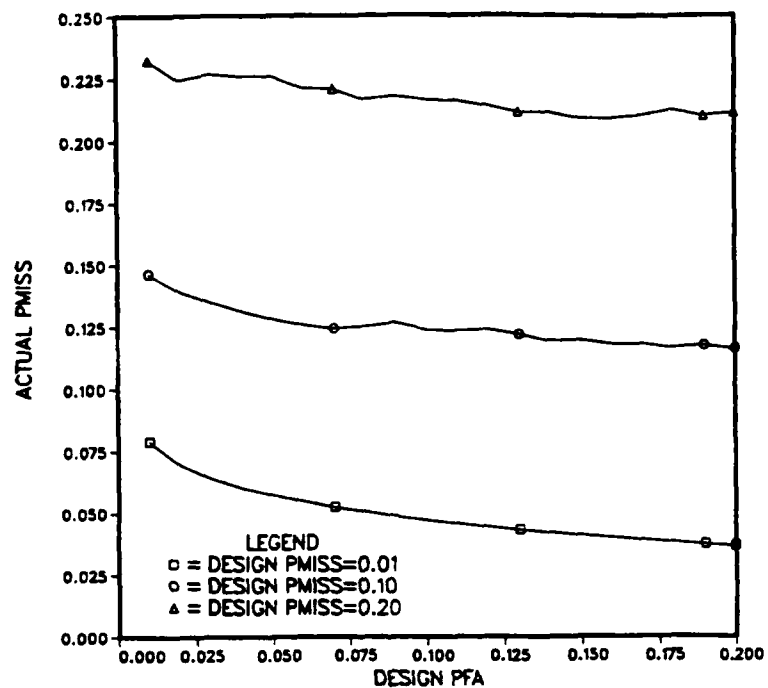


Figure 41. FSS PM vs α for Fixed β With Modulation

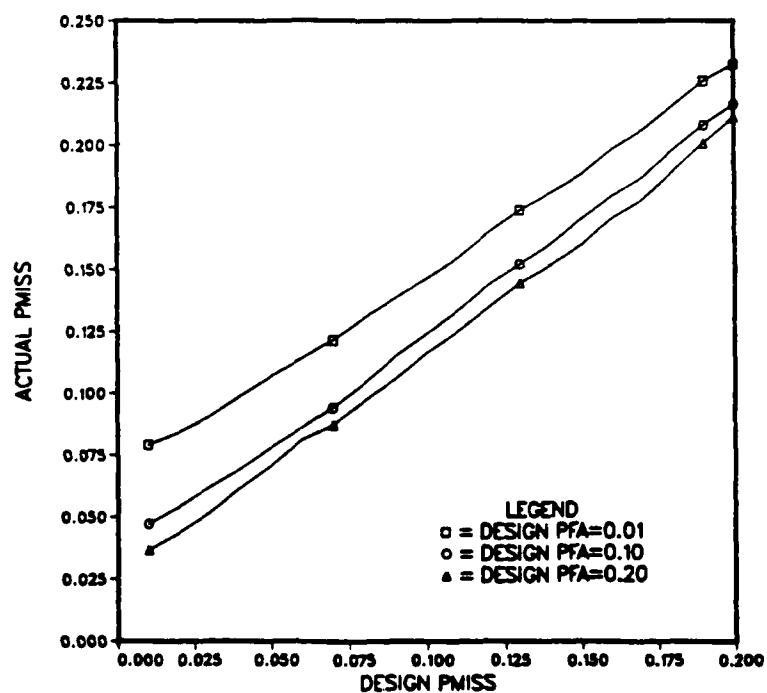


Figure 42. FSS PM vs β for Fixed α With Modulation

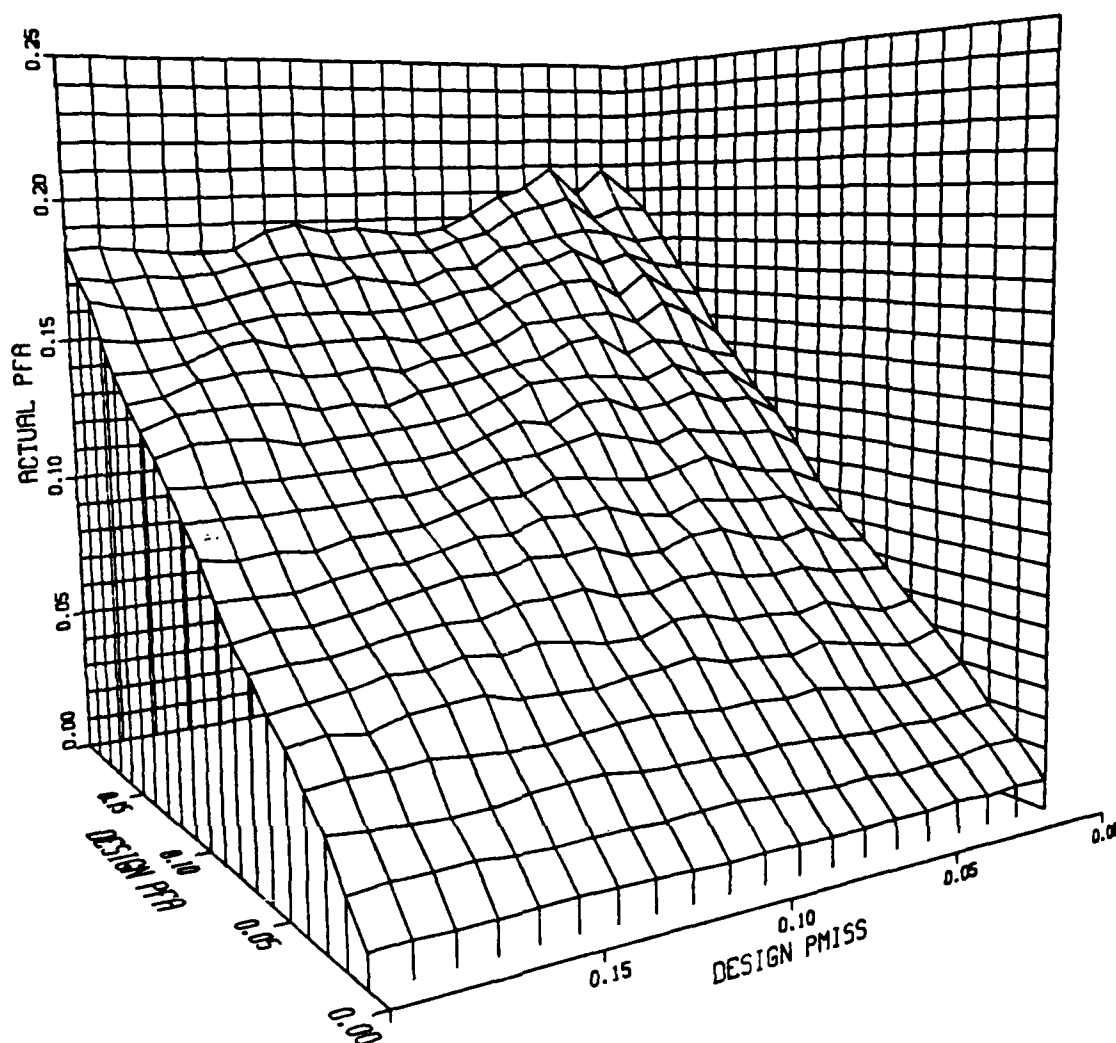


Figure 43. Actual FSS False Alarm Probability With Modulation

Figure 44 on page 77 displays the average sample number vs. modulation chip boundary under H_1 . Figure 45 on page 77 shows the actual probability of detection vs. the modulation chip boundary under H_1 . We see that if the modulation boundary is encountered very early, the test will run longer and attempt to decide the H_1 hypothesis. If the modulation boundary is encountered after the correlation process is well underway, PD drops dramatically. The ASN also drops since the test simply decides H_0 at an earlier and earlier point. The PD reaches a minimum if the data modulation boundary is encountered at approximately the 45th chip into the receiver correlation process. If

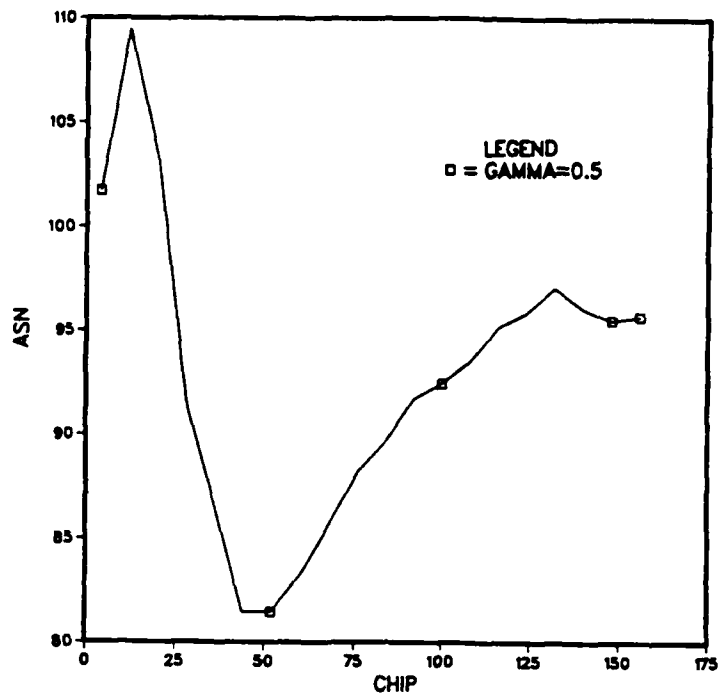


Figure 44. TSPRT ASN vs Modulation Chip Boundary: H_1

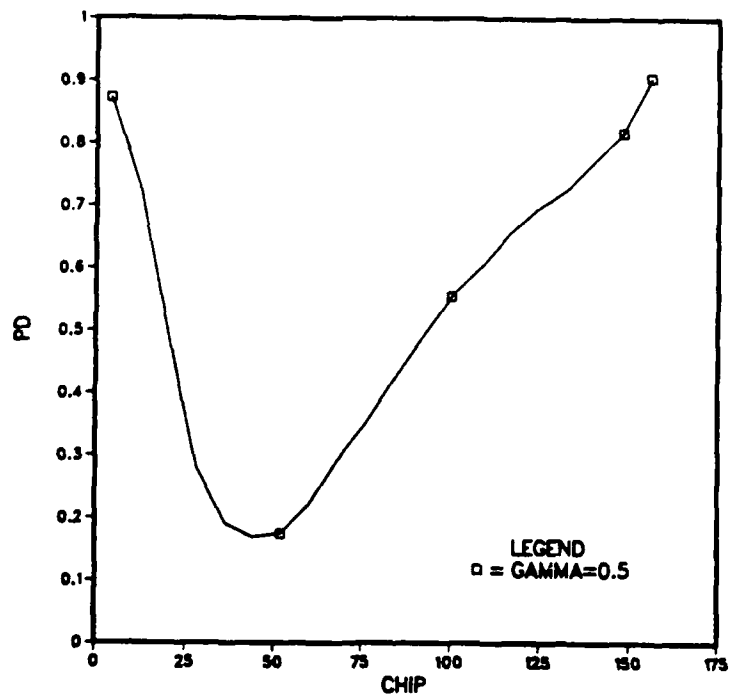


Figure 45. TSPRT PD vs Modulation Chip Boundary

the modulation boundary is encountered after this point, PD and ASN start increasing; the later the modulation boundary is encountered, the less damage done to the integration result that has built up prior to the boundary.

Figure 45 depicts the variation of PD vs. modulation chip for the TSPRT, while Figure 33 shows the same information for the FSS test. The minimum value of PD occurs in the FSS test at the midpoint between chip 0 and the termination chip. The minimum value of PD occurs in the TSPRT at a point far to the left of the midpoint between chip 0 and the test truncation value. This difference in the two curves arises from the fact that, whereas the FSS always runs the whole span from chip 0 to the termination chip, the TSPRT very rarely runs the whole span from chip 0 to the truncation chip. In actuality, for these TSPRT values of α and β , we found that *with no modulation* the expected value of ASN for the scheme is 95 chips. So, based on this, we would intuitively guess that modulation will cause the most damage if it occurs at the midpoint between chip 0 and chip 95. This is, in fact, what we observe.

Figure 46 on page 79 shows the ASN vs. modulation boundary under H_0 , and Figure 47 on page 79 shows the PFA vs. modulation boundary. Note that, as with the FSS test, modulation has little effect on receiver performance under the H_0 hypothesis.

We now summarize our TSPRT results for $\alpha = 0.04$, $\beta = 0.14$. Figure 48 on page 80 shows the test power vs. $|j + \gamma|$ with and without modulation. Without modulation, $E[PD]$ is 0.92. With modulation, $E[PD]$ drops to 0.89, but this is still greater than our desired value for the detection probability (0.86). This surprising result seems to indicate that our TSPRT is both *overdesigned and robust*. We only "request" that the test perform with $E[PD] = 0.86$. It actually performs with $E[PD] = 0.92$. If we then introduce modulation, the actual detection probability is decreased, but not enough such that its value is driven below our desired detection probability. This result can not be overemphasized: our original TSPRT scheme performs in accordance with our design values of α and β even in the presence of data modulation. Figure 49 on page 80 summarizes the expected value of ASN vs. $|j + \gamma|$ with and without modulation. The effect of modulation on ASN is extremely small.

The above results, while certainly encouraging, apply only to the particular case of $\alpha = 0.04$, $\beta = 0.14$. We wish to know just how robust and overdesigned the scheme is for other values of α and β .

We will pick a very stringent value of α and examine the TSPRT performance for various values of β . Specifically, we let $\alpha = 0.01$ and let β take on values from 0.01 to 0.12. (Equivalently, our desired probability of detection ranges from 0.88 to 0.99.)

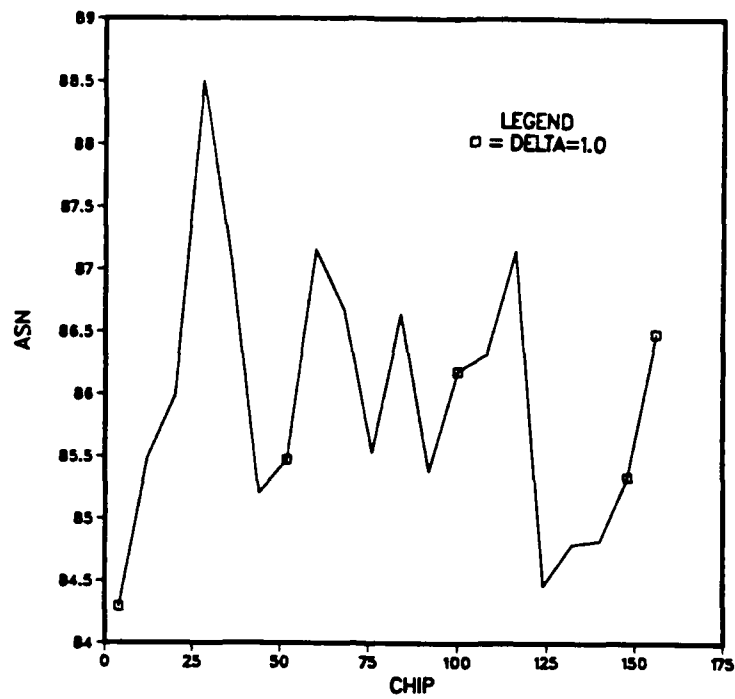


Figure 46. TSPRT ASN vs Modulation Chip Boundary: H_0

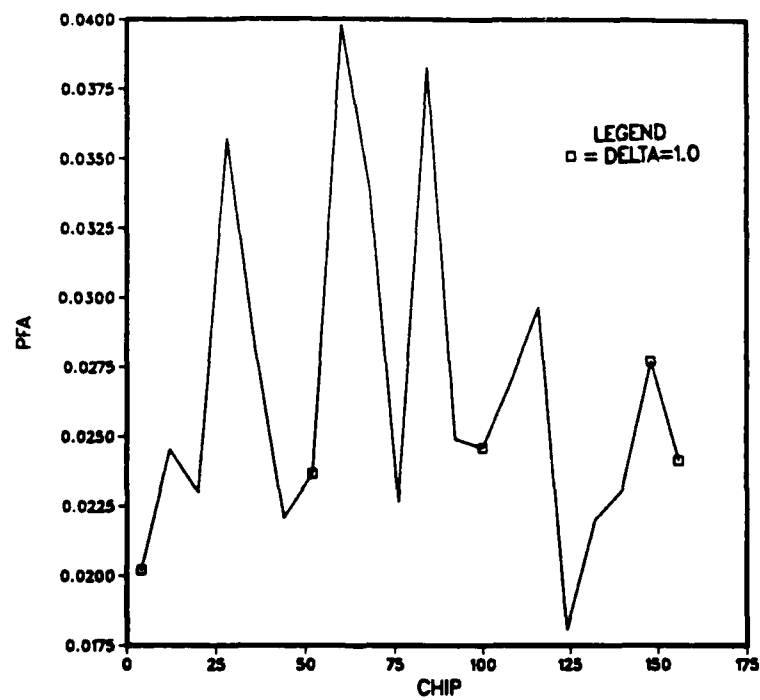


Figure 47. TSPRT PFA vs Modulation Chip Boundary

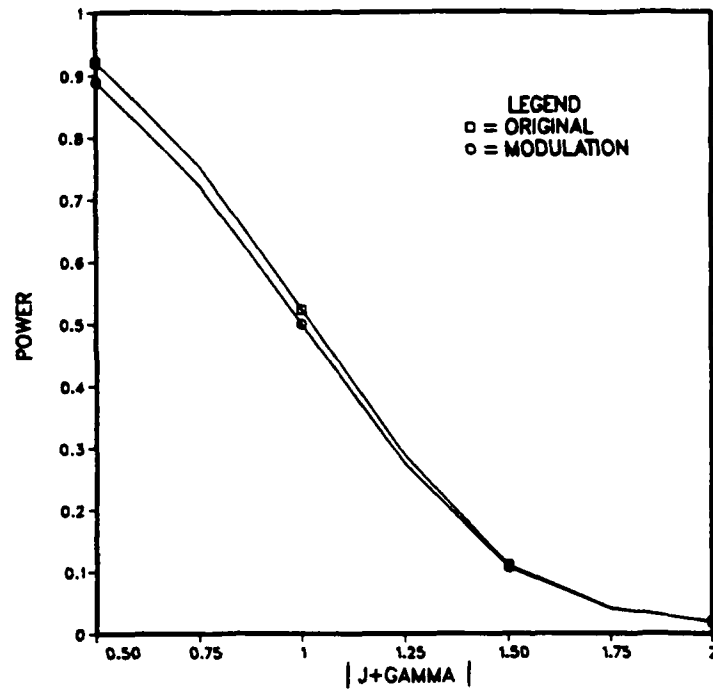


Figure 48. TSPRT Power vs $|j + \gamma|$ With Modulation

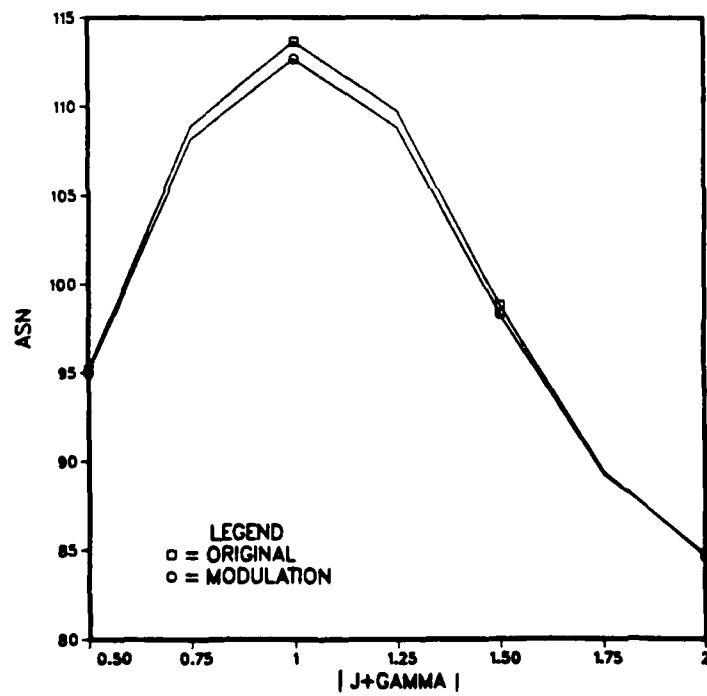


Figure 49. TSPRT ASN vs $|j + \gamma|$ With Modulation

Figure 50 on page 82 shows the actual detection probability (from simulation) vs. the desired detection probability with and without modulation. Also shown is the optimal curve for the actual vs. design probability of detection. Note that even with this very stringent value of α , the actual detection probability with modulation will satisfy our specification if our desired detection probability is less than 0.87. From Figure 50 we see that the degree to which the TSPRT is overdesigned increases as we relax the design value of detection probability. Figure 51 on page 82 displays the actual ASN vs. the design detection probability with and without modulation. We see that modulation has a negligible effect on the TSPRT ASN.

In summary, modulation does not greatly affect our TSPRT sequential detection scheme. With a very strict desired false alarm probability of 0.01 and a desired detection probability of 0.99, the actual detection probability in the presence of data modulation only drops to 0.95. If the desired detection probability is less than 0.88, the receiver will still perform in accordance with specifications. Note that this was not the case for the FSS test. For the FSS test, our results show that given any α and β , modulation will drive the receiver out of specification. (The degree to which the FSS receiver is out of specification *does*, however, depend on the values of α and β .)

D. THRESHOLD ADJUSTMENT

We have seen that modulation causes no deterioration of the receiver performance under H_0 , but does cause some degradation under H_1 , particularly for the FSS test. Since we know the exact form of the test statistic's pdf (given by (2.8) for H_0 and (4.7) or (4.9) for H_1) we can redesign our decision processor to account for the modulation. We consider the FSS test.

Using the Q function, we can write the FSS false alarm and miss probabilities as

$$P_{fa} = \alpha = Q\left(\sqrt{\frac{\lambda_{M,0}}{\sigma_M^2}}, \sqrt{\frac{\tau'}{\sigma_M^2}}\right) \quad (4.13)$$

$$P_{miss} = \beta = \left(1 - \frac{n}{2N}\right) \left(1 - Q\left[\sqrt{\frac{\lambda_{M,1}}{\sigma_M^2}}, \sqrt{\frac{\tau'}{\sigma_M^2}}\right]\right)$$

$$+ \left(\frac{n}{2N}\right) \left(\frac{1}{n}\right) \sum_{i=0}^{n-1} \left(1 - Q\left[\sqrt{\left[\frac{2i-n}{n}\right]^2 \frac{\lambda_{M,1}}{\sigma_M^2}}, \sqrt{\frac{\tau'}{\sigma_M^2}}\right]\right) \quad (4.14)$$

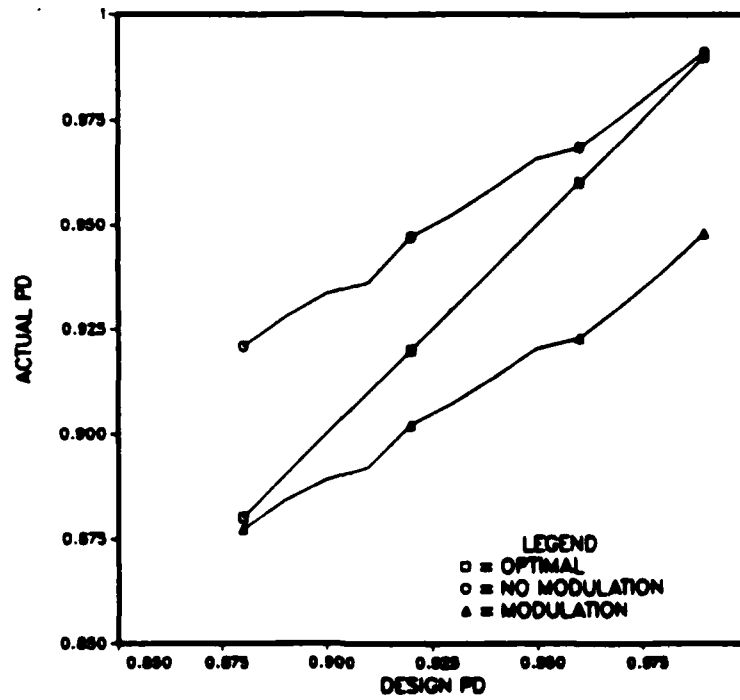


Figure 50. TSPRT PD With Modulation

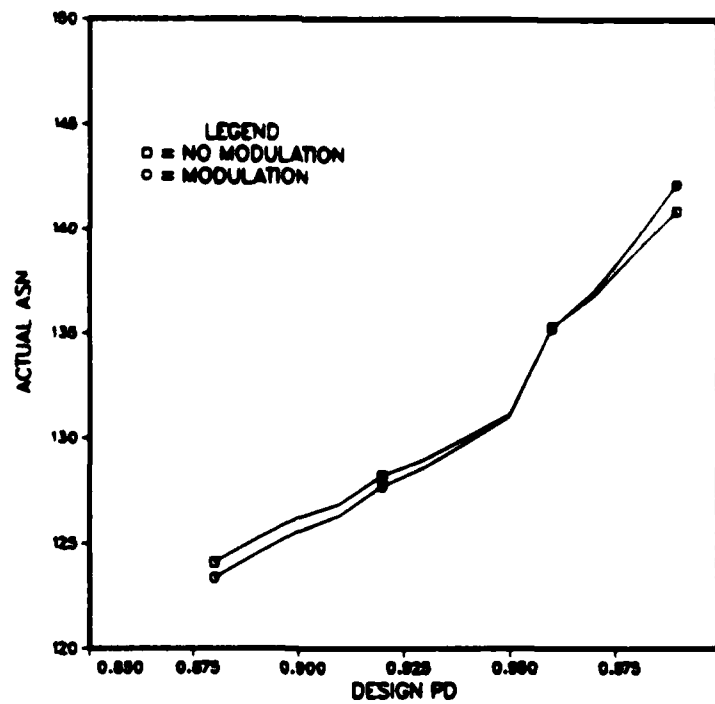


Figure 51. TSPRT ASN With Modulation

We attempt to iteratively solve these equations to determine the values of τ' and the termination point M . Unfortunately, we find that for many values of α and β the equations have no joint solution.

To investigate this dilemma, we first examine the test statistic pdf under H_1 (with modulation), given by (4.9), for various values of n , the number of chips in the integration process. We have two unknowns to determine: τ'/σ_M^2 (or z' in our normalized notation) and M . Equation (4.13) determines the value of z' , since z' is the only unknown in (4.13). This value of z' is used in (4.14) to determine M . We try consecutive values of M in (4.14) until we find the first value of M that satisfies the equation.

We find that sometimes no value of M exists which satisfies (4.14) once z' is determined from (4.13). Figure 52 on page 84 shows a plot of (4.9), the test statistic pdf under H_1 , for various values of n . Note that the value of the density function at $z = 0$ increases as n increases.

A particular example will illustrate the problem. Suppose we choose $\alpha = \beta = 0.01$ and we desire the solution of the pair of equations (4.13) and (4.14). Solving (4.13) we find that $z' = 10.0777$. Using this value of z' in (4.14), we attempt to adjust n until the equation is satisfied, but we find no solution. Instead of solving (4.14), let us try to determine a graphical solution.

Noting that (4.14) is the cumulative distribution function for the pdf under H_1 , we can rephrase (4.14) in words: determine a value of n such that the area under the probability density function, given by (4.9) and shown in Figure 52, from $z = 0$ to $z = 10.0777$ is less than or equal to 0.01. It turns out that the area from $z = 0$ to $z = 10.0777$ can not be made less than 0.01. We determined the exact value of the area under the pdf between $z = 0$ and $z = 10.0777$ for all values of n from $n = 0$ to $n = 1000$. The results are shown in Figure 53 on page 85. We find that the minimum area ever possible between $z = 0$ and $z = 10.0777$ is 0.0777 square units, which occurs when $n = 280$. Returning to (4.14), we conclude that if we attempt to adjust the thresholds, the best miss probability we can theoretically achieve (when $\alpha = 0.01$) is $\beta = 0.0777$, and this occurs if we run the FSS for 280 chips.

We draw two conclusions from our specific example.

1. If we choose $\alpha = 0.01$ then, no matter how we attempt to adjust the FSS scheme, we can not theoretically achieve a value of β less than 0.0777. If we choose a value of β greater than 0.0777, say 0.09 or 0.15, we should be able to adjust the FSS scheme such that it performs within specification when modulation is present.
2. The value of n obtained for the smallest β , $n = 280$ in our case, is the "optimal" n we can choose if we desire a β less than 0.0777. In other words, if we desire

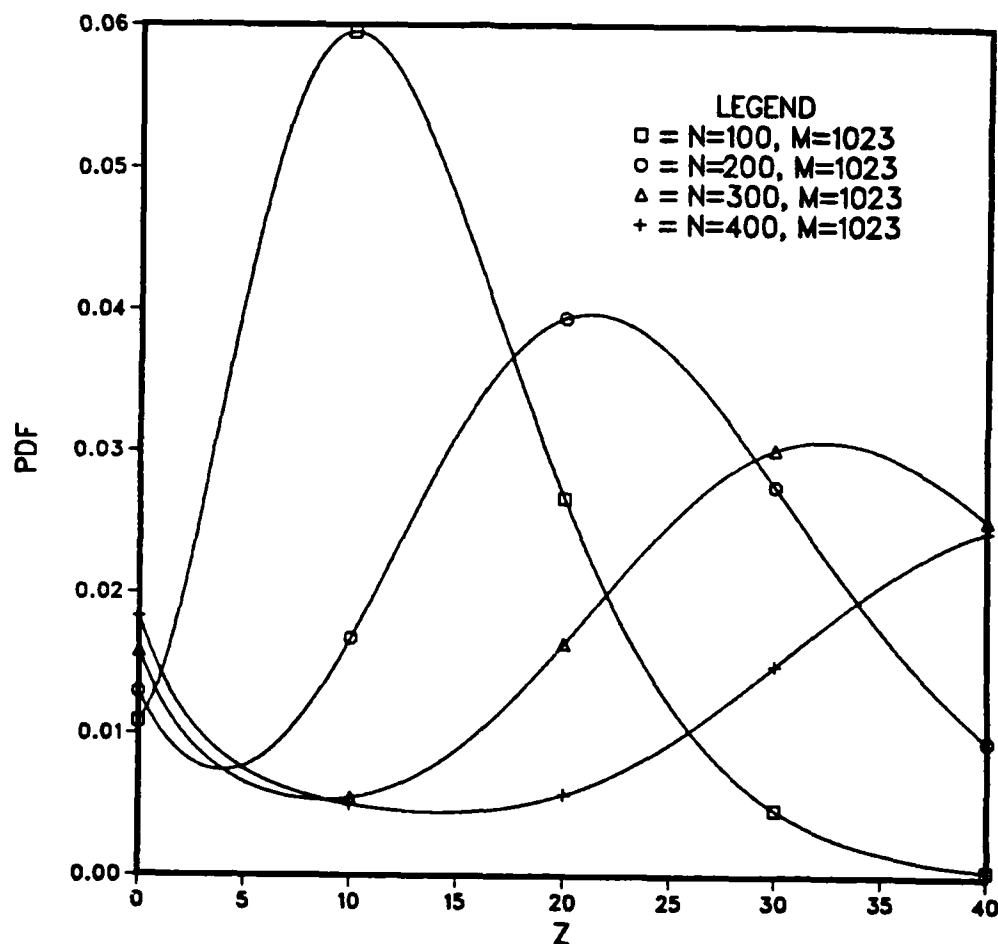


Figure 52. Density Function With Modulation: H_1

$\beta = 0.05$ then, if we pick a test length of 280 chips, although we will not achieve our desired β , any other value of M (less than or greater than 280) will perform even worse.

Note that the example we chose above is particularly strict. For many values of α and β (other than the ones we used above) the two equations (4.13) and (4.14) will have a joint solution which gives us a successful adjusted FSS test. Only for very small values of β do the equations not have a solution.

Figure 54 on page 86 shows the value of M , the adjusted FSS test length, for $0.01 \leq \alpha, \beta \leq 0.20$ with modulation. The optimal test length is found by solving (4.13) and (4.14) simultaneously for each pair (α, β) . The flat portion of the curve, for small values of design miss probability, represent "best" choices for M , i.e., choices for M which will not satisfy our desired α and β but are, nevertheless, the best we can do.

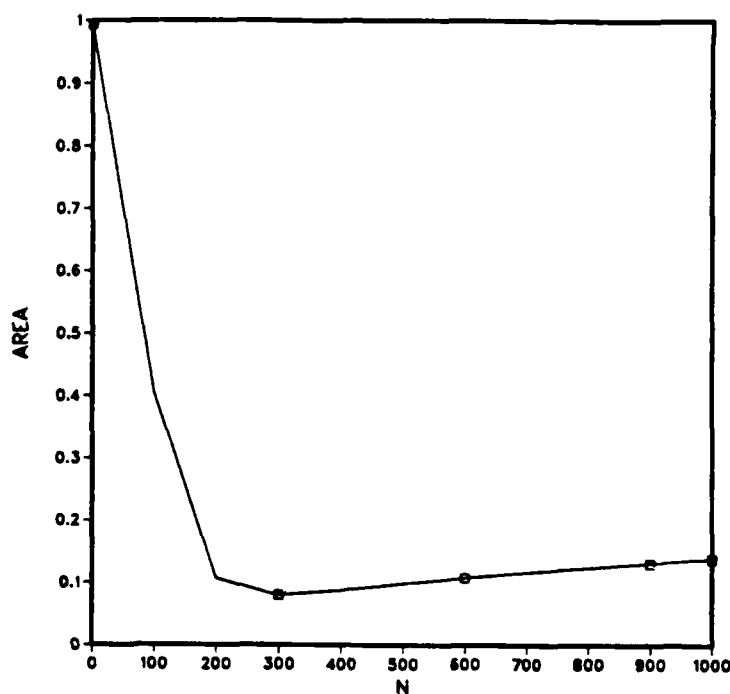


Figure 53. Area Under pdf From $z = 0$ to $z = 10.077$

Figure 55 on page 87 displays the required test length vs. desired miss probability for various fixed probabilities of false alarm. We expect the test to meet the specifications (α and β) for all values of the design miss probability except for points on the flat portions of the curves in Figure 55. For the flat portions, using the indicated value of M for the FSS test length will give us the best miss probability that can possibly be achieved, although we expect β will not be satisfied.

Figure 56 on page 88 shows the results of actual miss probability vs α and β with modulation present and thresholds suitably adjusted. Compare these results to those obtained without threshold adjustment, shown in Figure 40. We see that our threshold adjustment does indeed work where we expect it to, and does indeed fail where we predicted it would. Prior to threshold adjustment, we see from Figure 40 that the test always failed. Figure 57 on page 89 shows the actual miss probability vs. α for several fixed values of β . Note that if the desired miss probability is 0.1 or 0.2, the adjusted FSS test satisfies the performance criteria. If the desired miss probability is 0.01, we see that the test fails, but the curve in Figure 57 shows the best that is achievable.

Looking at the results another way, Figure 58 on page 89 displays the actual miss probability vs. the design miss probability for several fixed values of α . For each value

of α , the adjusted FSS test performs within specification until the design miss probability goes below a certain value.

Finally, Figure 59 on page 90 shows the actual false alarm probability vs. α and β for the FSS test with the thresholds adjusted to account for the data modulated received signal. Note that the results are similar to those obtained without threshold adjustment (Figure 43). Again, modulation does not significantly affect the test performance under the H_0 hypothesis.

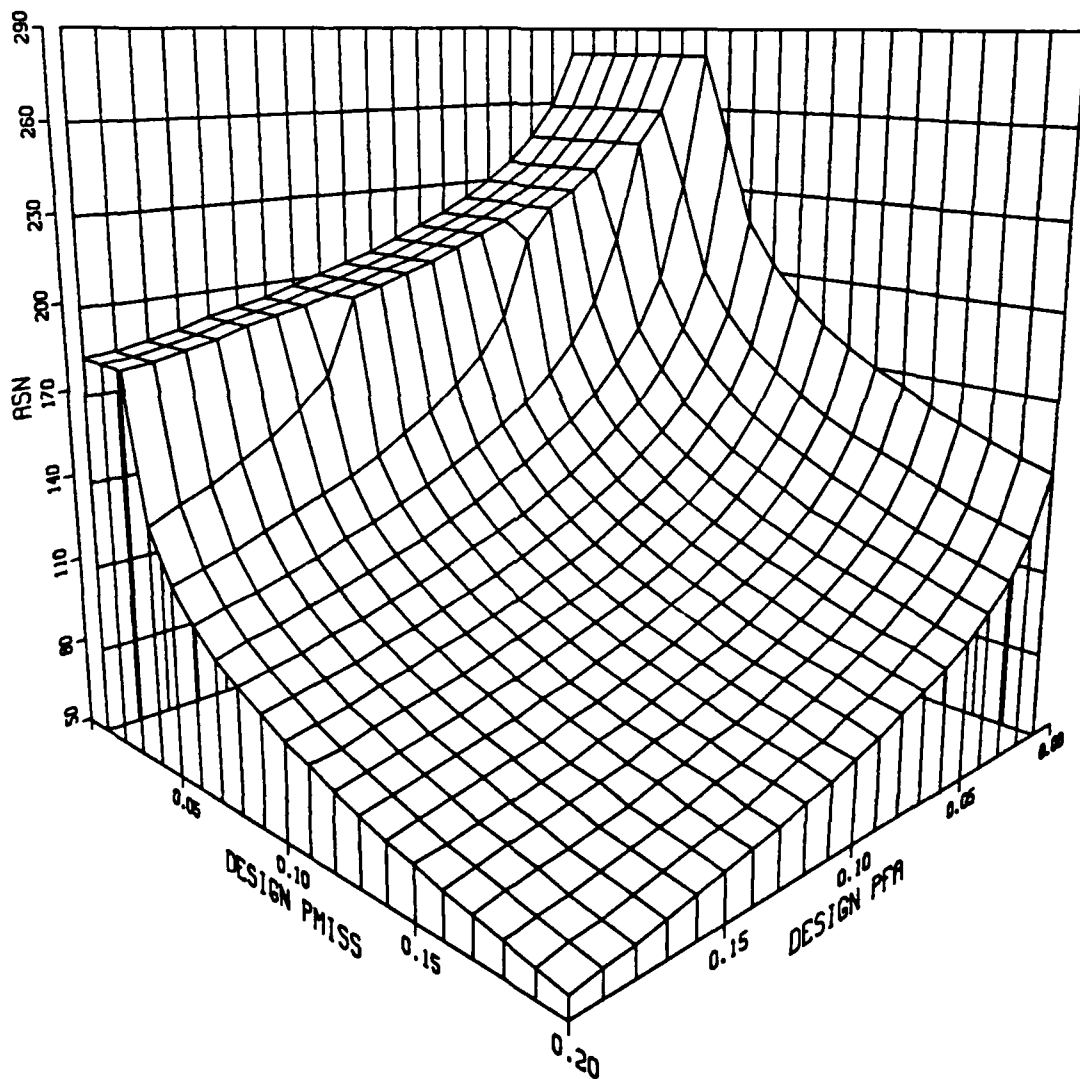


Figure 54. Adjusted FSS Test Length

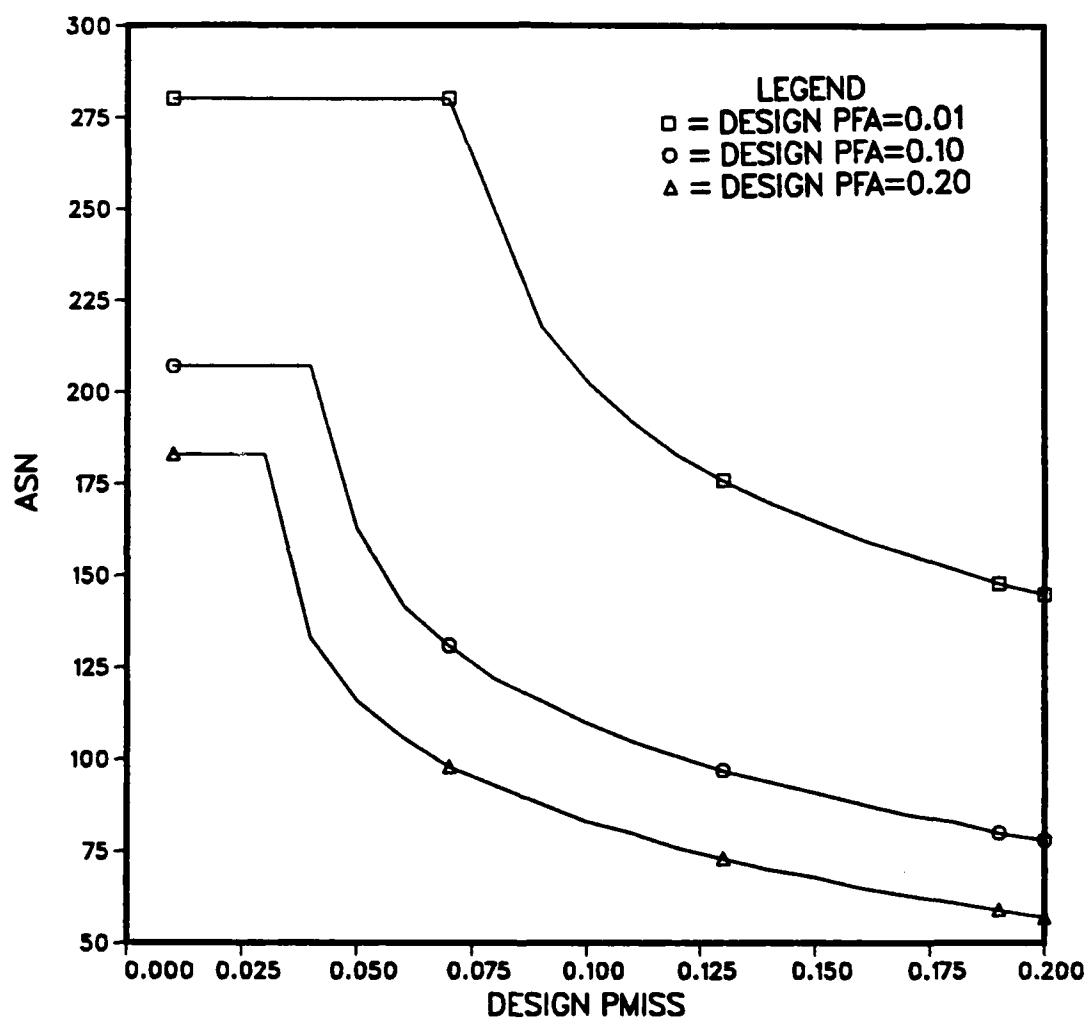


Figure 55. FSS Test Length vs β : New Thresholds

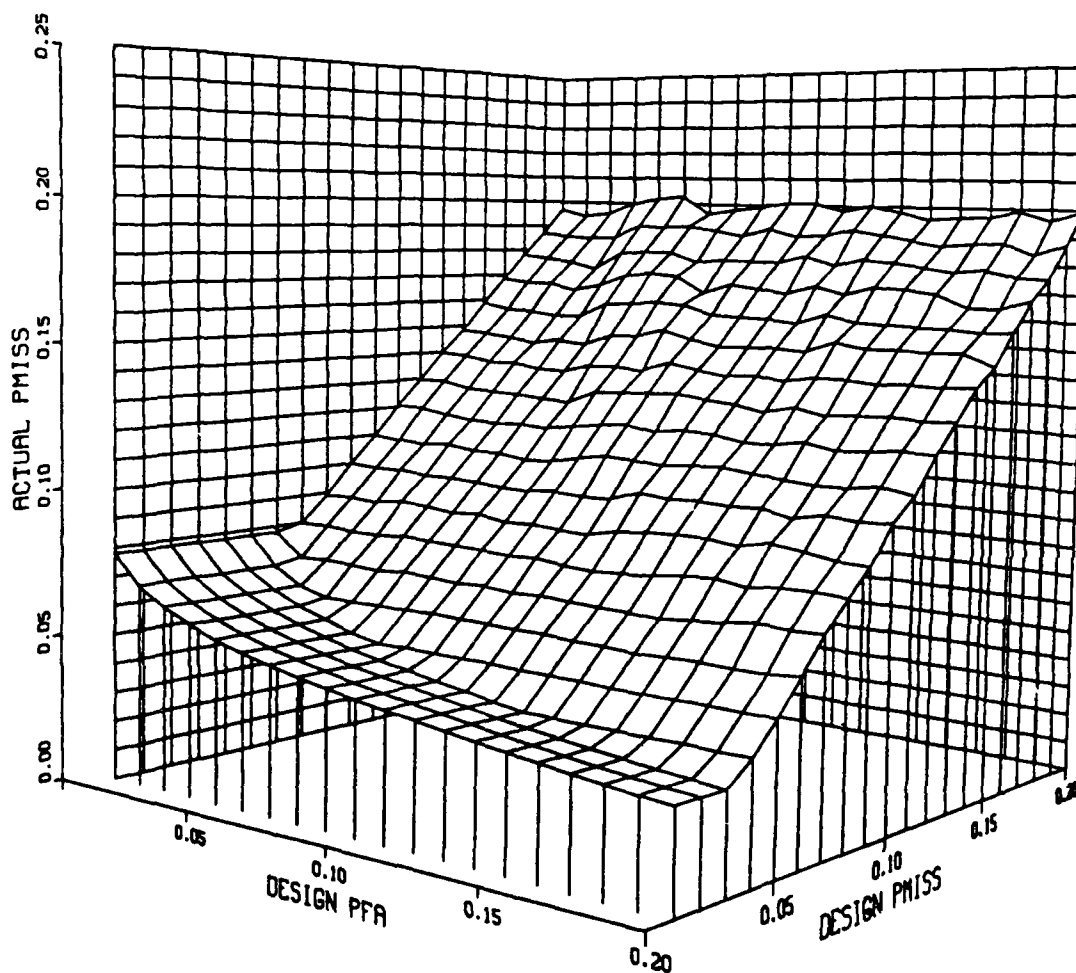


Figure 56. Actual FSS Miss Probability: New Thresholds

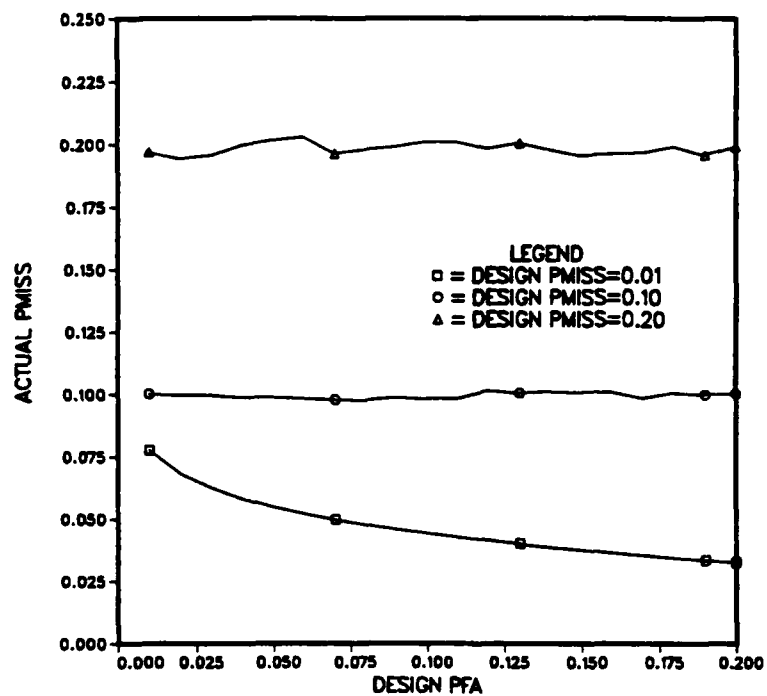


Figure 57. Actual FSS Miss Probability vs α : New Thresholds

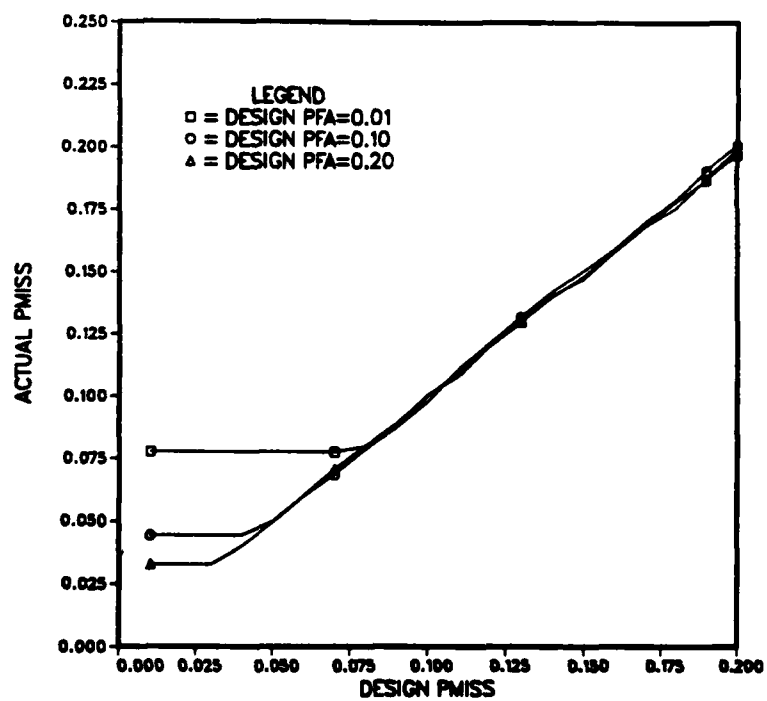


Figure 58. Actual FSS Miss Probability vs β : New Thresholds

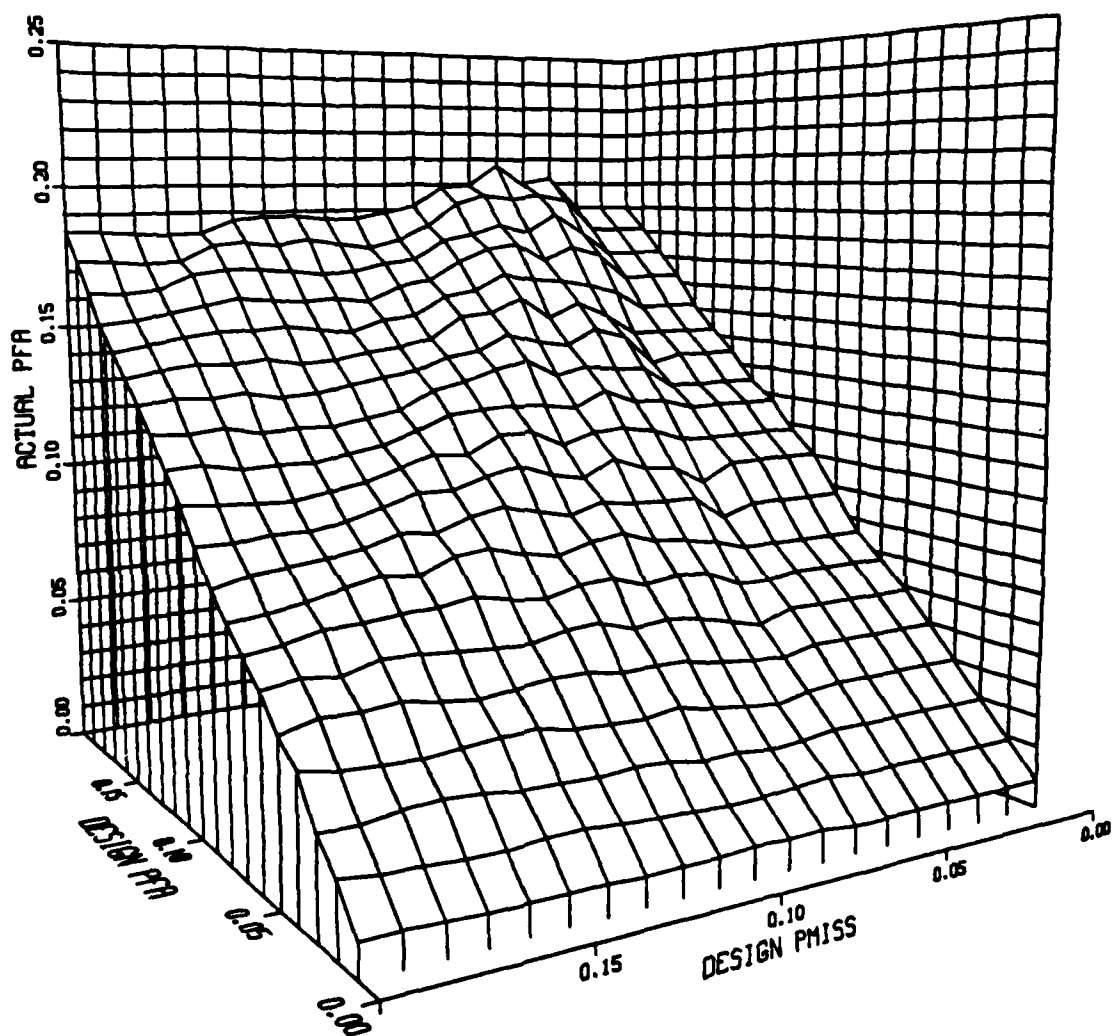


Figure 59. Actual FSS False Alarms Probability: New Thresholds

V. CONCLUSIONS

In this thesis, we have studied the effects of fading and modulation on noncoherent m-sequence acquisition schemes.

We have considered two schemes: the fixed sample size test (FSS) and the truncated sequential probability ratio test (TSPRT). Most prior studies have assumed that coherent demodulation is available during the acquisition process. We have dropped this unrealistic assumption since, in practice, m-sequence acquisition is generally performed by the receiver prior to carrier recovery. Additionally, most prior research has assumed that the samples used in the decision processor are independent and identically distributed. We do not reset the integrators during the correlation process, resulting in an enhanced effective signal to noise ratio.

The effects of channel fading on the performance of the synchronization tests were investigated. We considered a Ricean fading channel characterized by r , where r is the ratio of the received signal power in the direct component to the received signal power in the diffuse component. The probability density function for the receiver's test statistic was derived for the Ricean fading channel and found to be noncentral Chi-squared. It was determined that fading theoretically reduces the received signal's SNR by a factor of $r/[1 + r + \lambda_n/2\sigma_n^2]$.

To simulate fading we considered two simulation techniques. The first was a Monte Carlo algorithm derived from the strong law of large numbers. It was found that the Monte Carlo method did not provide a very tight confidence interval for a 50-run simulation. We then applied a numerical integration method that provided answers which converged within 25 runs.

We examined the effects of channel fading on the performance of the original decision processor. For the FSS test we observed that fading had virtually no effect on the false alarm probability, but drastically reduced the detection probability. With very severe fading, a detection probability of 0.99 can be decreased by more than 25%. As r increases, the detection probability approaches the original value without fading. For the TSPRT we also observed that the false alarm probability was essentially unaffected by the fading, while the detection probability was severely reduced. For the TSPRT, the ASN under H_0 also remains essentially the same, but the ASN under H_1 increases as the fading becomes more severe.

We then considered performance of the tests designed with the knowledge of the fading. We adjusted the sample size and thresholds of the decision schemes in accordance with the new test statistic density function. For the FSS test the performance criteria were then satisfied but the required sample size increased to a very large number when the fading was severe. As an example, for $\alpha = \beta = 0.01$ and an SNR of -10dB, the required FSS sample size increased from 258 (for no fading) to 361, 509, 6617 and 8944 when fading is characterized by $r = 20, 10, 1$ and 0 , respectively. Since the H_0 condition is unaffected by fading and occurs for all but a few of the uncertainty phases, adjusting the FSS thresholds greatly increases the acquisition time. Our results suggest that it may not be worthwhile to adjust the decision processor to account for fading.

When we adjusted the thresholds for the TSPRT test, a similar result was observed. The updated test performed in the fading channel in accordance with specifications, but at the expense of a much larger ASN under H_0 . As an example, for $\alpha = \beta = 0.01$ and an SNR of -10dB, the ASN under H_0 increased from 158 to 215 and 308 for $r = 20$ and 10 , respectively. To reduce receiver acquisition time it is critical that the ASN be minimized under H_0 . In adjusting the test thresholds, the price to pay for achieving the probability of detection is a large increase in ASN under H_0 . We conclude that it might be preferable to apply a test designed under no-fading conditions; we may simply accept any degradation in detection probability caused by the presence of fading in the channel.

We then investigated the effects of data modulation on the performance of the synchronization tests. Again, an expression for the density function of the receiver test statistic was derived. To see the effects of data modulation on the performance of the original decision processor, we first considered tests designed under no modulation and evaluated their performance in the presence of data modulation. For each scheme the modulation effect was found to be not severe. For example, if $\alpha = \beta = 0.01$ and SNR = -10dB, modulation causes almost the same effect on the FSS test as fading with $r = 10$. Modulation has virtually no effect on the FSS test false alarm probability, but does moderately reduce the detection probability. The degradation in the detection probability tends to lessen as we relax the design criteria (α and β).

We found that modulation had even less of an effect on the TSPRT scheme. For a stringent desired false alarm probability of 0.01, the scheme will meet any desired detection probability less than 0.87 even if data modulation is present. If the desired detection probability is greater than 0.87, the scheme will perform slightly worse than desired. For example, if the desired probability of detection for the original receiver is 0.99, it is only reduced to 0.95 if data modulation is present. The ASN of the TSPRT

is essentially unchanged with data modulation. We conclude that the TSPRT is rather robust in the presence of data modulation and threshold adjustment may be not warranted.

We finally considered performance of the FSS test designed with the knowledge of data modulation. We found that there are values of (α, β) for which it is impossible for the adjusted FSS receiver to perform within specification. For these points we determined the best sample size which resulted in the largest possible detection probability. For the other values of (α, β) , the test can be adjusted to perform satisfactorily at the expense of increasing the sample size.

APPENDIX GENERATION OF RICEAN RANDOM VARIATES

To generate random variates from an underlying probability distribution, one starts with uniformly distributed random variates (on the interval $[0,1]$) and uses a suitable transformation to arrive at the desired distribution. Uniform variates are always the logical starting point: the theory is very well developed and many algorithms have been written and thoroughly tested.

Obviously computer-generated "random" numbers are not random; they are, in fact, totally predictable. Furthermore, their predictability is often a desirable property [Ref. 25 : pp. 132-145]. The best we can achieve with a computer random number generator is a sequence of numbers which *appear to be random* when subjected to standard statistical tests. Generally, the best random number generators put out sequences of numbers which have long periods, i.e., which take a long time to start repeating.

Making use of abstract algebra, random number generators of the form $x_n = f(x_{n-1})$ have been developed which have the maximum possible period prior to repeating. The most commonly used maximum length generator is the *congruential generator* which generates uniform random variates u_i by using the algorithm [Ref. 26 : pp. 20-26]:

$$x_{i+1} = (ax_i + c) \pmod{m} \quad (A.1)$$

$$u_i = \frac{x_i}{m} \quad (A.2)$$

To get the period equal to m , we need to follow three rules [Ref. 22 : p. 13]

1. c must be relatively prime to m
2. $a \equiv 1 \pmod{p}$ for every prime p which divides m
3. $a \equiv 1 \pmod{4}$ if m is a multiple of 4

To optimize the statistical properties of the random variates, we adhere to a fourth rule [Ref. 27 : pp. 155-156]

4. the multiplier a should be larger than $\sqrt{m}$, preferably larger than $m/100$, but smaller than $m - \sqrt{m}$

The simulation programs were designed to run on a VAX computer using a base of 2 and a word size of 31. Thus theoretically, the best choice for m (to result in the longest sequence period) would be $m = 2^{31}$. However we can not choose $m = 2^{31}$ because the

largest number generated by the algorithm (A.1),(A.2) is $a(m-1) + c$. (This occurs when $ax_i + c = m-1$. In such a case $(ax_i + c)(\text{mod } m)$ will equal $m-1$ and then x_{i+1} will cause a number $a(m-1) + c$ to be generated.) If we choose $m = 2^{31}$ then, during the algorithm, at some point, the computer will want to work with a number larger than 2^{31} (potentially as large as $a(2^{31}-1) + c$), which will cause overflow.

Rule 4. above suggests a possible choice for the constants a , m and c . If we choose $a \approx \sqrt{m}$ then the largest number generated by the algorithm will be $a(m-1) + c \approx \sqrt{m} m = m^{3/2}$. We set this equal to the largest number obtainable in the computer:

$$m^{3/2} = 2^{31} \quad (A.3)$$

So, to the nearest power of two, a good choice for m is thus $m = 2^{20}$. The reason for choosing a power of two is to facilitate choosing the constant c which, by the third rule, must be relatively prime to m . Any odd number can now be chosen for c .

Summarizing the above results, we can satisfy rules 1.-4. and have a maximum period for our VAX computer by choosing:

1. $m = 2^{20} = 1048576$. This will be the number of uniform random variates generated prior to repeating.
2. Since we need $a > \sqrt{m}$ we must choose $a > 1024$. To satisfy $a \equiv 1 (\text{mod } 2)$ and $a \equiv 1 (\text{mod } 4)$ we choose $a = 1029$.
3. For c pick any odd number, subject to 4. below.
4. The maximum value generated by the algorithm is $a(2^{20}-1) + c \approx 2^{20} + c$. Since this number must be kept less than 2^{31} , we require that the number c be kept less than about one billion.

For the remainder of this discussion, suppose we have written two separate uniform random number generators (e.g., by applying the choices outlined above but with two different values of c) which output independent variates U_1 and U_2 . We now use these to generate other needed distributions.

It can be shown [Ref. 26 : pp. 67-68] that using the random variable U_1 , we can generate a random variate with an exponential distribution of parameter one by the transformation

$$E_1 = \ln\left(\frac{1}{1-U_1}\right) \quad (A.4)$$

Now, with this exponential variate E_1 and the independent uniform random variable U_2 , we can generate two independent normal variates with distribution $N(0,1)$ by using the procedure [Ref. 28]

$$N_1 = \sqrt{2E_1} \cos 2\pi U_2 \quad (A.5a)$$

$$N_2 = \sqrt{2E_1} \sin 2\pi U_2 \quad (A.5b)$$

Suppose that instead of standard normal variates N_1 and N_2 we desire normal variates described as $N(s, \sigma^2)$ and $N(0, \sigma^2)$. That is, we desire that both of our normal variates be independent but with variance σ^2 (instead of one) and with means s and 0 respectively. Such normal variates can be generated by using the basic properties of mean and variance [Ref 15 : pp. 96-103] as

$$N_A = \sigma N_1 + s \quad (A.6a)$$

$$N_B = \sigma N_2 \quad (A.6b)$$

We have now generated two normal variates with $N_A \sim N(s, \sigma^2)$ and $N_B \sim N(0, \sigma^2)$. To generate a Ricean distributed random variable, we need only to calculate [Ref 1 : pp. 30-31]

$$\psi = \sqrt{N_A^2 + N_B^2} \quad (A.7)$$

Such a random variable will have a Ricean density function given by (3.2) and repeated below:

$$f_{\Psi}(\psi) = \frac{\psi}{\sigma^2} e^{-s^2/2\sigma^2} e^{-\psi^2/2\sigma^2} I_0\left(\frac{\psi s}{\sigma^2}\right) \quad (A.8)$$

Recall from Chapter III that fading is usually characterized by specifying the value of r , where r is the ratio of the power in the direct component to the power in the diffuse component:

$$r = \frac{s^2}{2\sigma^2}$$

with the constraint that total power has a normalized value of 1 (i.e., $s^2 + 2\sigma^2 = 1$).

We consider an example. Suppose that we wanted to generate independent variates from a Ricean distribution characterized by $r = 1$. One way to choose s^2 and σ^2 would be

$$s^2 = 1/2 \quad , \quad \sigma^2 = 1/4$$

so that $s^2 + 2\sigma^2 = 1$. So starting with our standard normal variates N_1 and N_2 , we calculate, using (A.6):

$$N_A = \frac{N_1}{2} + \sqrt{\frac{1}{2}}$$

$$N_B = \frac{N_2}{2}$$

and then determine ψ from (A.7).

LIST OF REFERENCES

1. Proakis, J.G., *Digital Communications*, 2d ed, McGraw Hill, 1989
2. Schilling, D.L., and others, "Spread Spectrum For Commercial Communications," *IEEE Communications Magazine*, pp. 66-78, April 1991
3. Schilling, D.L., Pickholtz, R.L., and Milstein, L.B., "Spread Spectrum Goes Commercial," *IEEE Spectrum*, pp. 40-45, August 1990
4. Milstein, L.B., and Das, P.K., "Surface Acoustic Wave Devices," *IEEE Communications Society Magazine*, v. 17, n. 5, pp. 25-33, September 1979
5. Holmes, J.K., *Coherent Spread Spectrum Systems*, John Wiley and Sons, 1982
6. Tantaratana, S., and Lam, A.W., "Noncoherent Sequential Acquisition of PN Sequences for DS/SS Communications," presented at the 29th Allerton Conference on Communications, Control and Computing, Urbana-Champaign, October 1991
7. Wald, A., *Sequential Analysis*, John Wiley and Sons, 1947
8. Helstrom, C.W. *Statistical Theory of Signal Detection*, Pergamon Press, 1968
9. Brennan, L.E., and Reed, I.S., "A Recursive Method of Computing the Q Function," *IEEE Transactions on Information Theory*, Vol. IT-11, pp. 312-313, April 1965
10. Tantaratana, S., "Notes on the Design Of Truncated Sequential Tests With Non-i.i.d. Samples," submitted for publication
11. Gagliardi, R.M., *Introduction to Communications Engineering*, Wiley, 1988
12. Brown, J.B., and Glazier, E.V.D., *Telecommunications*, Chapman and Hall, 1964

13. Ziemer, R.E., and Tranter, W.H., *Principles of Communications*, Houghton Mifflin Company, 1976
14. Couch, L.W., *Digital and Analog Communication Systems*, 3d ed, Macmillan, 1990
15. Walpole, R.E. and Myers, R.H., *Probability and Statistics For Engineers and Scientists*, 4th ed, Macmillan, 1989
16. Whalen, A.D., *Detection of Signals in Noise*, Academic Press, 1971
17. Schwartz, M., Bennett, W.R., and Stein, S., *Communication Systems and Techniques*, McGraw Hill, 1966
18. Beyer, W.H., *CRC Standard Mathematical Tables and Formulae*, 29th ed., CRC Press, 1991
19. Carlson, B.C., *Special Functions of Applied Mathematics*, Academic Press, 1977
20. Andrews, L.C., *Special Functions for Engineers and Applied Mathematicians*, Macmillan, 1985
21. Hansen, E.R., *A Table of Series and Products*, Prentice Hall, 1975
22. Yakowitz, S.J., *Computational Probability and Simulation*, Addison Wesley, 1977
23. Gerald, C.F. and Wheatley, P.O., *Applied Numerical Analysis*, 4th ed., Addison Wesley, 1989
24. Sloggett, D.R., "Improving the Reliability of HF Data Transmissions," *IEE Colloquium Digest*, 48, 1979
25. Hamming, R.W., *Numerical Methods For Scientists and Engineers*, 2d ed., McGraw Hill, 1973
26. Rubinstein, R.Y., *Simulation and the Monte Carlo Method*, Wiley, 1981

27. Knuth, D.E., *The Art of Computer Programming: Seminumerical Algorithms*, Vol 12, Addison Wesley, 1969
28. Box, G.E. and Muller, M.E., "A Note on the Generation of Random Normal Deviates," *Annals of Mathematical Statistics*, v. 29, pp. 610-611, 1958

INITIAL DISTRIBUTION LIST

	No. Copies
1. Defense Technical Information Center Cameron Station Alexandria, VA 22304-6145	2
2. Library, Code 52 Naval Postgraduate School Monterey, CA 93943-5002	2
3. Chairman, Code EC Department of Electrical and Computer Engineering Naval Postgraduate School Monterey, CA 93943-5000	1
4. Professor Alex W. Lam, Code EC/La Department of Electrical and Computer Engineering Naval Postgraduate School Monterey, CA 93943-5000	5
5. Professor Herschel H. Loomis, Code EC/Lm Department of Electrical and Computer Engineering Naval Postgraduate School Monterey, CA 93943-5000	1
6. Patrick J. Vincent 87-08 80th Street Woodhaven, NY 11421	1