

Abstract

DIRECT SEQUENCE SPREAD SPECTRUM

Adam Mirek Mankowski, M.S.E.

The University of Texas at Austin, 2000

Supervisor: Edward J. Powers

The concepts of spread spectrum with a focus on the widely used direct sequence spread spectrum modulation technique are presented. After examining the fundamental concepts of spread spectrum, direct sequence is discussed in detail, along with various pseudo-noise methodologies, concluding with several examples of commercial applications.

20010314 066

REPORT DOCUMENTATION PAGE			Form Approved OMB No. 0704-0188	
Public reporting burden for this collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington, VA 22202-4302, and to the Office of Management and Budget, Paperwork Reduction Project (0704-0188), Washington, DC 20503.				
1. AGENCY USE ONLY (Leave blank)		2. REPORT DATE 18.Jan.01		3. REPORT TYPE AND DATES COVERED MAJOR REPORT
4. TITLE AND SUBTITLE DIRECT SEQUENCE SPREAD SPECTRUM			5. FUNDING NUMBERS	
6. AUTHOR(S) 2D LT MANKOWSKI ADAM M				
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) UNIVERSITY OF TEXAS AUSTIN			8. PERFORMING ORGANIZATION REPORT NUMBER CI01-12	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES) THE DEPARTMENT OF THE AIR FORCE AFIT/CIA, BLDG 125 2950 P STREET WPAFB OH 45433			10. SPONSORING/MONITORING AGENCY REPORT NUMBER	
11. SUPPLEMENTARY NOTES				
12a. DISTRIBUTION AVAILABILITY STATEMENT Unlimited distribution In Accordance With AFI 35-205/AFIT Sup 1			12b. DISTRIBUTION CODE	
13. ABSTRACT (Maximum 200 words)				
14. SUBJECT TERMS			15. NUMBER OF PAGES 59	
			16. PRICE CODE	
17. SECURITY CLASSIFICATION OF REPORT	18. SECURITY CLASSIFICATION OF THIS PAGE	19. SECURITY CLASSIFICATION OF ABSTRACT	20. LIMITATION OF ABSTRACT	

Dedication

This Report is dedicated to my parents, Aleksander and Barbara, who have always encouraged me to do my best and have always been there for me when I needed them. Thank you for everything, especially for being my source of motivation when mine was lacking.

DIRECT SEQUENCE SPREAD SPECTRUM

by

Adam Mirek Mankowski, B.S.

Report

Presented to the Faculty of the Graduate School

of The University of Texas at Austin

in Partial Fulfillment

of the Requirements

for the Degree of

Master of Science in Engineering

The University of Texas at Austin

August 2000

DIRECT SEQUENCE SPREAD SPECTRUM

Approved by
Supervising Committee:



Edward J. Powers



Gary A. Hallock

Abstract

DIRECT SEQUENCE SPREAD SPECTRUM

Adam Mirek Mankowski, M.S.E.

The University of Texas at Austin, 2000

Supervisor: Edward J. Powers

The concepts of spread spectrum with a focus on the widely used direct sequence spread spectrum modulation technique are presented. After examining the fundamental concepts of spread spectrum, direct sequence is discussed in detail, along with various pseudo-noise methodologies, concluding with several examples of commercial applications.

Table of Contents

1. Introduction	1
2. Spread Spectrum	3
2.1 A Brief History of Spread Spectrum	3
2.2 What is Spread Spectrum?	7
2.3 Characteristics of Spread Spectrum Signals	10
2.3.1 Selective Addressing	11
2.3.2 Code Division Multiplexing	11
2.3.3 Low Power Density	12
2.3.4 Message Protection	13
2.3.5 High Resolution Ranging	13
2.3.6 Interference Rejection	14
2.3.7 You Can't Have Everything	14
2.4 Performance Specifications	15
2.4.1 Processing Gain	15
2.4.2 Jamming Margin	16
2.5 Two Popular Spread Spectrum Techniques	17
2.5.1 Direct Sequence	17
2.5.2 Frequency Hopped	18
2.5.3 Comparison of DS and FH Systems	19
2.6 Summary	19
3. Direct Sequence Spread Spectrum	21
3.1 Basic Concepts of DSSS	22
3.1.1 Block Diagram	22
3.1.2 Modulation	23
3.1.3 Demodulation	26
3.2 Characteristics of DSSS	28

3.2.1	Code Division Multiplexing	28
3.2.2	Low Power Density	30
3.2.3	Interference Rejection	32
3.3	Performance Specifications of DSSS	34
3.4	Pseudo-Noise Sequences	35
3.4.1	Average White Gaussian Noise vs. Pseudo-Noise	35
3.4.2	Properties of PN Sequences	36
3.4.3	Types of PN Sequences	39
3.4.3.1	Maximal Sequences	39
3.4.3.2	Gold Codes	43
3.5	Summary	46
4.	Commercial Applications	48
4.1	Global Positioning System	49
4.2	Personal Communications	51
4.3	Local Area Networks	52
4.4	Summary	53
5.	Conclusions	54
	Glossary	56
	References	58
	Vita	59

List of Tables

1. Tradeoffs in direct sequence and frequency hopped systems	20
2. Feedback connections for linear maximal sequence generators	41
3. Preferred pairs for Gold code generation	45

List of Figures

1. Power spectral density plots comparing equal power SS and CW signals for spread spectrum bandwidths of (a) 2 MHz and (b) 10 MHz	8
2. Frequency spectrum of a DSSS system	18
3. Frequency spectrum of a FHSS system	19
4. Block diagram of a DSSS system	22
5. Modulation of a DSSS signal showing both the time and frequency domain for (a) the data, (b) the PN sequence, and (c) the resultant spread spectrum output	25
6. Demodulation of a DSSS signal showing both the time and frequency domain for (a) spread spectrum input, (b) PN sequence, and (c) original data	27
7. Block diagram of a multiple access network via DSSS	29
8. Frequency spectra showing modulated and demodulated users on a multiple access DSSS network	30
9. Comparison between the frequency spectra of a DSSS modulated and unmodulated signal in the presence of Gaussian noise	31
10. Effect of narrowband interference on a DSSS signal	33
11. Effect of wideband interference on a DSSS signal	34
12. Autocorrelation function and frequency spectra of a PN sequence	38
13. Simple shift register generator	40
14. Autocorrelation of a maximal sequence	42
15. Cross-correlation of maximal sequences	42
16. Block diagram of Gold code generator	43
17. Autocorrelation of a Gold code	44
18. Cross-correlation of various Gold codes	45
19. Building blocks of GPS	49

1. Introduction

Although the foundations of spread spectrum can be traced as far back as 1924, its development can be attributed to the needs of the military during the Second World War. The desperate need for effective radar and secure communications drove the progression of spread spectrum technology at a furious pace. Since then, spread spectrum modulation methods have become a well-studied and mature technology. The evolution of VLSI design and signal processing combined with the vast demand for wireless communications services have led to the widespread application of spread spectrum. Besides its capability of providing secure communications, spread spectrum has many other advantages.

Spread spectrum modulation techniques, whose bandwidth is vastly greater than that required of the information being transmitted, have many distinct characteristics that distinguish this modulation technique from all other narrowband modulation methods. Some of the advantages of spread spectrum modulation methods include: selective addressing, code division multiplexing, signal hiding, message screening, high resolution ranging, and interference rejection. Unfortunately, all of these advantages cannot be employed at once. For example, it would be an exercise in futility to attempt to effectively hide a signal while maintaining good performance in the face of strong interference. However, applications and conditions arise that require the use of this technology.

The aim of this report is to provide the reader with a firm understanding of spread spectrum, stemming from Shannon's well-known work on information theory. After a brief discussion of the history of spread spectrum, we will delve into a study of the characteristics of spread spectrum listed above. Two useful performance specifications—process gain and jamming margin—will be derived and later applied to the examination of the direct sequence spread spectrum modulation technique.

Given a solid understanding of the concepts of spread spectrum, we will move on to the focus of this report—the study of direct sequence spread spectrum modulation. The most well-known and widely employed spread spectrum modulation technique, direct sequence makes use of noise-like sequences in order to spread the bandwidth of an information signal. It is this spreading that characterizes spread spectrum and yields the benefits of such modulation. The heart of direct sequence lies in the noise-like sequences.

Many different coding methods are employed in digital communications in order to provide robust and reliable data transfer. The difference between such codes and the sequences used in spread spectrum lies in the length of the codes. The extreme lengths of the sequences used by direct sequence provide privacy, protection against interference, and the reduction of the effects of noise. Two popular coding methods will be discussed: maximal sequences and gold codes.

Concluding the discussion of spread spectrum and direct sequence, several commercial applications of the technology will be examined. Both Global Positioning System receivers and Personal Communications Systems are widely employed throughout the world. Companies such as AT&T, Sprint, Motorola, and Samsung, to name but a few, have invested billions of dollars in the development of both infrastructure and devices that use spread spectrum modulation. Wireless local area networks are also coming of age, taking advantage of the multiple access capabilities of spread spectrum.

Spread spectrum modulation techniques can be seen all around us, and are a multi-billion dollar per year industry. Hopefully this report will provide the reader with a firm understanding of this fascinating technology; specifically the popular direct sequence spread spectrum modulation technique.

2. Spread Spectrum

Practical applications of spread spectrum modulation techniques emerged during the Second World War in response to the need for secure communications and effective radar systems. Over the past half-century, advances in VLSI, signal processing, and information theory combined with the demand for wireless communications services have led to the widespread use of spread spectrum systems.

This section presents the reader with a solid background on spread spectrum—beginning with a discussion of its roots in the history of radar, communication, and information theory. The various characteristics of spread spectrum signals are then presented, along with two performance specifications used to analyze the effectiveness of spread spectrum modulation techniques. Two popular spread spectrum modulation techniques—direct sequence and frequency hopped—are also presented, and comparisons made between their performance characteristics. The focus of this report, however, lies in the examination of the former modulation technique. Direct sequence will be discussed in detail in section three of this report.

2.1 A Brief History of Spread Spectrum

Like many other technologies, the history of spread spectrum has a well documented evolutionary trail that can be traced back as early as 1924. From its roots, spread spectrum technology progressed at a slow, steady pace as advances in radar, information theory, and communications were realized. As the struggle for freedom engulfed the world during the Second World War, a desperate need for secure communications and effective radar systems sparked renewed vigor in the development of this marvelous technology. Waged with jamming and antijamming tactics, World War II witnessed a battle for electronic supremacy the

likes of which had never before been encountered. Before we examine primitive spread spectrum systems, let us look briefly at a chronological list of the advances that led to their development [12].

Advances in Radar

- 1920's: E.V. Appleton and M.A.F. Barnett credited with birth of RADAR
- 1938: G. Guanella files a patent detailing the technical characteristics of a spread spectrum radar
- 1940: E. Huttman issued German patent for chirp pulse radar; later that decade Americans R.H. Dicke and S. Darlington are issued U.S. patents for a similar system
- 1940's: North, Van Vleck, and Middleton formulated the matched filter concept to maximize signal-to-noise ratio pulse detection
- 1950's: P.M. Woodward resolved issues of resolution, accuracy, and ambiguity properties of pulse waveforms

Information Theory

- 1930: N. Wiener develops the theory of spectral analysis for nonperiodic infinite-duration functions
- 1940's: C.E. Shannon begins to establish a fundamental theory of communication within a statistical framework and develops his well-known channel capacity theorem

Communications Systems

- 1924: A.N. Goldsmith issued patent for a communication method that counteracts the fading effects due to multipath—generally considered to be the first patent that is spread spectrum in nature

- 1935: P. Kotowski and K. Dannehl apply for German patent for a device that masks voice signals by combining them with an equally broad-band noise signal produced by a rotating generator
- Prior to WWII: H. Busignies, E. Deloraine, and L. deRosa applied for a patent on a facsimile communication system that is considered an early relative to time hopped spread spectrum systems.

The discoveries and developments in these fields paved the way for the development of primitive spread spectrum systems. Although most of the spread spectrum systems at the time were developed under the shroud of secrecy, there remained a great deal of information exchange at the conceptual level [12]. The similarities between spread spectrum systems developed both during and after WWII can thus be explained by the freedom of information. We will now quickly examine a few of the early spread spectrum systems developed in response to the threats experienced during this war of electronic dominance [12].

WHYN

WHYN, an acronym standing for Wobulated HYperbolic Navigation, was developed by Sylvania in the midst of the 1940's. The goal of this project was to develop a guidance system for surface-to-surface missiles for the United States Army. Besides being one of the first frequency modulated spread spectrum systems, WHYN also employed correlation detection.

Hush-up

During the 1950's, Madison Nicholson of Sylvania headed the development of a communication system known as "Hush-Up," that used noise-like carriers to spread the bandwidth of an information signal. The method chosen was a pseudorandomly generated binary sequence that was PSK

modulated onto an RF sinusoid. Hush-up was one of the first direct sequence spread spectrum systems. During a series of tests at Wright-Patterson Air Force Base, the assigned carrier frequency was that of the base's tower frequency, and the system was set up within 100 yards of the tower's antennas. These tests were successful—neither the tower personnel nor the aircraft with which they were communicating were ever aware of the spread spectrum communications being conducted over their frequency bands.

Noise Wheels

The creation and exact reproduction of noise-like signals was one of the major difficulties in the early developments of spread spectrum systems. In the late 1940's, Mortimer Rogoff began experimenting with photographic techniques for storing a noise-like signal and for building an ideal cross-correlator. Spread spectrum systems employing such noise wheels were built and tested in early 1950. When properly synchronized at the transmitter and receiver, these noise wheels provided for the successful extraction of signals 35 dB below the interfering noise.

NOMAC

NOMAC, or NOise Modulation and Correlation, was developed at M.I.T. Lincoln Laboratory in the 1950's. NOMAC systems were used to compare the performances of transmitted- and stored-reference spread spectrum systems operating in the presence of broadband Gaussian noise. In stored-reference systems, the pseudo-noise sequence is locally generated at both the transmitter and receiver. In transmitted-reference systems, however, the sequence is transmitted on a second carrier. Such studies performed at M.I.T. led to the development of several transmitted-reference spread spectrum systems. Development of synchronization techniques has rendered transmitted-reference

systems obsolete.

The evolution of spread spectrum technology can only be attributed to the hard work and dedication of numerous gifted scientists and engineers. Since there have been so many contributions to this field of study, an adequate discussion of the history of spread spectrum lies beyond the scope of this paper. For those interested in a handsome presentation of the origins of spread spectrum, refer to Robert Scholtz's article found in [12].

2.2 What is Spread Spectrum?

Unlike conventional modulation techniques that confine their frequency spectra to a narrow band, a spread spectrum system is one whose transmitted signal is spread over a very wide frequency band [1]. The difference in the bandwidth of such signals can range from a few to several orders of magnitude. For example, conventional AM and FM signals are usually confined to bandwidths in the range of 10-100 kilohertz, whereas the spectra of spread spectrum systems often span tens to hundreds of megahertz [3]. The actual information signal buried within the spectrum remains only a few kilohertz. The key advantage of employing such a large bandwidth is that better signal-to-noise ratios are realized through this process. Figure 1 illustrates the bandwidth differences between a conventional and spread spectrum signal. It is important to note that the total power in both signals is equal.

Like all other modulation techniques, spread spectrum can be distinctly classified. As a matter of fact, for a signal to be classified as spread spectrum, it must satisfy two important criteria:

1. The bandwidth of the transmitted signal must be much greater than the bandwidth or rate of the information being sent.

2. The transmitted bandwidth must be determined by some function that is independent of the message being sent.

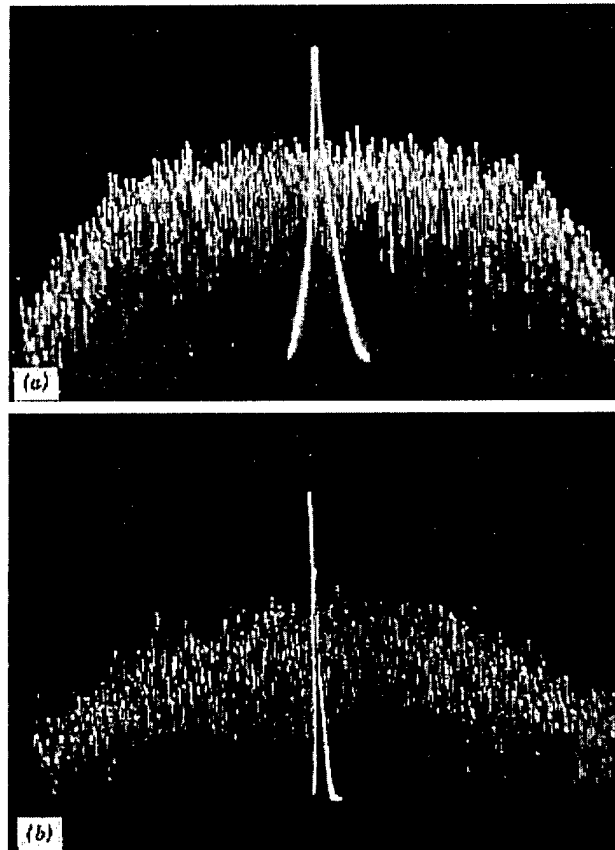


Figure 1. Power spectral density plots comparing equal power SS and CW signals for spread spectrum bandwidths of (a) 2 MHz and (b) 10 MHz [3].

The fundamental nature then, of spread spectrum, is to expand the bandwidth of a signal, transmit the expanded signal, and then recover the original information by despreading the received signal. The expansion of the bandwidth is achieved through the combination of noise-like coding sequences with the information signals. These codes are pseudorandom, and have noise-like

properties when compared with the information being transmitted [5].

In 1947, Claude E. Shannon—whose work in information theory forms the basis of much of modern digital communications—presented a very valuable expression for the capacity of a channel in the presence of additive white Gaussian noise [12]. The capacity of a channel, C , in bits per second can be calculated by equation (1):

$$C = W \log_2 \left(1 + \frac{S}{N} \right) \quad (1)$$

where W = bandwidth in hertz,

N = noise power, and

S = signal power.

This equation shows the ability of a channel to transfer error-free information as a function of the signal-to-noise ratio of the channel, and the bandwidth of the signal used to transmit the information.

In his book on spread spectrum systems, Robert Dixon employs this well-known theorem in showing the benefits yielded by spread spectrum [3]. This derivation provides a strong argument for spread spectrum, and is repeated here.

Fixing C as the desired system information rate, and changing to an exponential base, equation (1) can be rewritten as

$$\frac{C}{W} = 1.44 \log_e \left(1 + \frac{S}{N} \right).$$

Given a small signal-to-noise ratio (≤ 0.1), and taking the logarithmic expansion, we are left with

$$\frac{C}{W} = 1.44 \frac{S}{N}.$$

From this equation we can approximate

$$\frac{N}{S} = 1.44 \frac{W}{C} \approx \frac{W}{C}$$

which can be rearranged to yield

$$W = \frac{NC}{S} \quad (2).$$

Equation (2) allows us to see that, given any signal-to-noise ratio for a particular channel, we can have a low information-error rate by increasing the bandwidth used to transfer the desired information across the channel. The larger the ratio of the noise to the signal, the greater the bandwidth required to maintain low error communications. The results of equation (2) will aid in the understanding of the signal-hiding characteristic of spread spectrum signals. Signal hiding is but one of many characteristics of spread spectrum signals, all of which will be discussed in the following section.

2.3 Characteristics of Spread Spectrum Signals

In a world hungry for wireless communications services, it seems strange that the technology used to meet the demands is one that uses such an extraordinarily large bandwidth. It is important to understand that there are many valuable characteristics of spread spectrum signals. These useful characteristics can be attributed to the pseudorandom noise-like sequences used to code the

information signals. In the sections that follow, six characteristics of spread spectrum signals are presented, along with potential uses for such technology.

2.3.1 Selective Addressing

Spread spectrum signals are generated using noise-like sequences. These very sequences, however, can be used to select a single receiver or group of receivers by assigning to them unique codes that are different from those of other receivers. Selective addressing can then be implemented by simply transmitting the proper code sequence as modulation [3]. In order to take advantage of the capabilities of selective addressing, codes must be chosen which have good cross-correlation properties. Cross-correlation measures the degree of similarity between two different codes [6]. If two codes are too similar to one another, significant interference will occur, resulting in the degradation of system performance. Selective addressing is very similar to another characteristic of spread spectrum systems—that of the capability of multiple access through code division multiplexing.

2.3.2 Code Division Multiplexing

Although the application of spread spectrum for a single user seems preposterous, spread spectrum signals can also provide reliable, reduced error, and simultaneous communications for multiple users over the same channel. When multiple users employ such a channel, spread spectrum systems become much more bandwidth efficient [5]. Such capabilities are once again due to the noise-like sequences employed in the generation of spread spectrum signals. This type of multiplexed system is known as a code division multiple access (CDMA) system, and is widely used for satellite communications systems and terrestrial wireless personal communications systems.

Whereas the goal in selective addressing was to uniquely isolate a single

receiver or group of receivers, the goal of a CDMA system is to provide unique access to multiple users at the same time. Although similar in nature, there are some subtle differences in code selection for CDMA systems. In the selective addressing case we were not overly concerned with the interference produced by other users of the system. In CDMA systems, however, we are concerned with the additive interference created by the other users of the system. Therefore, the sequences employed for CDMA systems must have much better cross-correlation characteristics. Examples of such a sequence are Gold codes, which will be addressed later in this report. Another fact to note is that this capability not only allows for multiple users to share the same channel, it also allows for independent transmitting and receiving systems to operate over the same channel, simultaneously [3].

2.3.3 Low Power Density

This characteristic of spread spectrum signals has been one of the primary interests of military applications of the technology since the Second World War. The wideband spectra generated by code modulation produces a signal whose transmitted power over any narrow region is very low [3]. This power density can be driven so low that the signal spectral density cannot be readily detected in the presence of average white Gaussian noise. Signals that are virtually indiscernible from the noise are said to possess a low probability of intercept (LPI), and are called LPI signals [4].

Low spectral power density is also an important consideration for commercial applications of spread spectrum. With conventional signaling techniques, sufficient frequency spacing was required in order to minimize interference caused by one signal on another. This problem can be minimized with spread spectrum signals because the energy of the signal is spread over such a vast frequency range. The amount of interference caused to other users is thus

greatly reduced when compared to conventional signaling methods [3]. The interference cannot be ignored however. It is the limiting factor on the number of users who can simultaneously access a channel in CDMA applications.

2.3.4 Message Protection

Protection of transmitted information is of great importance in both military and civilian applications. Information transmitted via spread spectrum systems is protected against the casual eavesdropper due to the pseudonoise code sequences used to modulate the information signal. Not only do the codes protect the content of the message, they are also able to prevent unauthorized use of communications links [3]. If the codes selected for secure applications are sufficiently complex, it becomes very difficult for unauthorized users to break them and determine what the message actually is [1]. Unlike cryptographic applications, most codes are not specifically designed for the role of information security—message protection is merely a side effect of the use of codes for modulation. It must be stressed that the information transmitted via a spread spectrum system is not secure unless the spreading codes used are cryptographically secure.

2.3.5 High-Resolution Ranging

The ability for spread spectrum signals to be used for high-resolution ranging make it a particularly attractive technology for the enhancement of radar and navigation systems. Direct sequence spread spectrum techniques have been applied to ranging in space exploration programs since the early 1960s [3]. This capability stems from the fact that a broadband signal can be resolved in time much more precisely than a narrow-band signal [1]. Therefore, measurements of delay times and ranging information can be made much more accurately by transmitting a signal with a large bandwidth. One popular commercial application

that takes advantage of this characteristic of spread spectrum signaling is the Global Positioning System (GPS).

2.3.6 Interference Rejection

There are two types of interference that pose problems to traditional communications methods. From the standpoint of the military, one significant interference source is the jamming of a signal by an unfriendly third party. The ability of a spread spectrum signal to mitigate the effects of such interference is vital to military command and control links [3]. This type of interference is called intentional interference, and can render communication systems useless. Another type of interference is known as unintentional interference, and resides on every communication channel. Unintentional interference can be caused by natural phenomena, multiple users on the same channel, or other manmade noise sources that are not produced maliciously.

Spread spectrum modulation techniques are capable of providing reliable communications in the presence of both intentional and unintentional interference. Since most interference sources are localized in frequency, their affect on spread spectrum signals is significantly reduced through the process of despreading. The amount that the interference is reduced is proportional to the amount of spreading achieved by the spread spectrum signal [1].

2.3.7 You Can't Have Everything

All of these characteristics of spread spectrum signals are very attractive, and have many benefits for military and civilian applications. Like everything else in engineering, obtaining the characteristics you want often leads to tradeoffs in other areas. The most significant tradeoff made by spread spectrum signals is bandwidth for performance. Some of the characteristics discussed above go hand in hand. For example, message protection and low power density are both

proportional to the amount of spreading. On the other hand, if you would like to have a very low power density, then you cannot expect to fare well in terms of interference rejection. In other words, you can't maximize the capabilities of all the characteristics of spread spectrum signals at the same time.

2.4 Performance Specifications

There are two parameters that are often used when specifying the performance of a spread spectrum signal in the presence of interference. These parameters stem from the derivations of Shannon's channel capacity that we presented earlier. The following sections detail these performance parameters.

2.4.1 Processing Gain

The processing gain is defined as the difference between the output and input signal-to-noise ratios of any processor. The processing gain for spread spectrum processors can be estimated by [1,3]:

$$G_p = \frac{BW_{RF}}{BW_{data}} = \frac{R_c}{R_b} = \frac{T_b}{T_c} = N_c \quad (3)$$

where BW_{RF} = bandwidth of the transmitted spread spectrum signal

BW_{data} = bandwidth of the information signal

R_c = chip rate of the pseudonoise sequence

R_b = bit rate of the information signal

T_b = bit time

T_c = chip time; and

N_c = number of chips per bit

The term chip is used to define a single segment of a code or sequence, whereas a single piece of information is often referred to as a bit or symbol. It is the output

of a code generator during one clock cycle. Therefore, the chip rate, R_c , is simply the rate at which chips are generated, and its period, T_c , is found by taking the inverse of the chip rate. The code rate determines the length of the code sequence since the information symbol (bit) period is fixed—there are always an integer number of chips per bit.

The processing gain can also be thought of as the power improvement factor that a receiver can achieve when it matches its locally generated pseudo-noise sequence with that of the transmitted message [10]. Although a spread spectrum system may possess a high processing gain, this does not imply that the system will function if the interference has a power sufficiently large to approach the processing gain. Another parameter, the jamming margin, is useful in examining such situations.

2.4.2 Jamming Margin

The jamming margin is the second parameter that is used to specify the performance of spread spectrum systems. This parameter takes into consideration both the internal losses of a system as well as the signal-to-noise ratios at the output. The jamming margin is therefore always less than the processing gain for any spread spectrum system. The jamming margin, in decibels, is estimated by [3]:

$$M_j = G_p - \left[L_{sys} + \left(\frac{S}{N} \right)_{out} \right] \quad (4)$$

where L_{sys} = system implementation losses; and

$(S / N)_{out}$ = signal-to-noise ratio at the output.

As shown, a consideration of the system implementation losses implies that the

jamming margin is the maximum allowable interference that can be present on the channel while still maintaining the specified error probability.

2.5 Two Popular Spread Spectrum Techniques

There are many types of spread spectrum systems that have been developed since the first applications during the Second World War. The two most prevalent are direct sequence and frequency hopped [3]. Both methods will be discussed briefly in order to provide the reader with a fundamental understanding of these two popular spread spectrum modulation techniques. Bear in mind that the focus of this report is on direct sequence spread spectrum.

2.5.1 Direct Sequence

Direct sequence is the most widely known and understood method of spread spectrum modulation [3]. It involves the modulation of a pseudo-noise (PN) sequence onto an information signal. The process of combining the PN sequence with the information signal produces a signal whose bandwidth is significantly greater than that of the information signal. The amount of spreading that is achieved is proportional to the length of the sequence used to modulate the information. This baseband spread spectrum signal is then modulated once more onto an RF carrier. Because the bandwidth of the spread signal is so large, its power density is very small. Therefore, the signal possesses the spectral equivalent of a noise signal. Figure 2 displays the frequency spectrum of a direct sequence spread spectrum (DSSS) modulated signal.

The simplicity of spreading the information bandwidth is one of the key attributes of DSSS modulation. Fortunately, the demodulation process is just as simple. Demodulation involves multiplying a locally generated PN sequence with the received signal. When the locally generated PN sequence identically matches and is synchronized with that of the modulated signal, the two are said to be

perfectly correlated. If correlation is achieved, the resulting output is the original information signal.

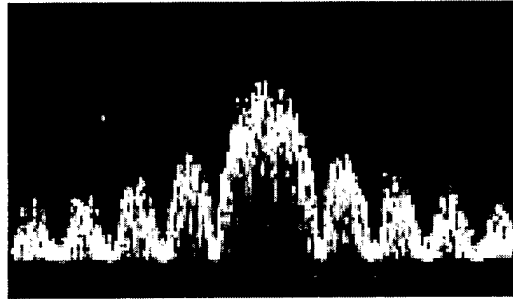


Figure 2. Frequency spectrum of a DSSS system [7].

2.5.2 Frequency Hopped

Whereas DSSS systems spread their signal energies over an extremely large bandwidth, frequency hopped spread spectrum (FHSS) systems rely on frequency diversity. If you were to instantaneously look at a DSSS modulated signal, you would find that it is always a broadband signal. On the other hand, FHSS signals look instantaneously narrowband—they transmit at full power on a particular frequency for a fixed amount of time. When the next hop time rolls around the frequency of transmission changes, or “hops.” Due to the number of frequencies employed by FHSS systems, their signals on average look broadband. Figure 3 shows the frequency spectrum of a frequency hopped spread spectrum signal.

The selection of the frequency for transmission is governed by a pseudo-noise sequence that feeds a frequency synthesizer. A sequence of n chips per bit dictates one of 2^n frequency positions. In order to calculate the bandwidth used by such a system, we need only to find the lowest and highest possible hop positions required of the synthesizer. Demodulation of FHSS modulated signals involves ensuring that both transmitter and receiver follow the same frequency

hops.



Figure 3. Frequency spectrum of a FHSS system [7].

2.5.3 Comparison of DS and FH Systems

In the preceding sections we have discussed two well-known spread spectrum modulation techniques. The two methods differ significantly, but each has characteristics that make them useful for different situations and conditions. For military systems, where it is often vital for a message to get through regardless of the environment through which it is being transmitted, frequency hopping is the modulation method of choice. Consumers, on the other hand, are often more concerned with quality of service—that is, minimum interruption and distortion. For this reason, direct sequence modulation is preferred for consumer applications. Besides these fundamental differences, there are many tradeoffs when selecting one modulation method over another. Table 1 lists many of the factors that must be taken into account when making such a decision.

2.6 Summary

The foundations of spread spectrum technology emerged from discoveries in radar, communications, and information theory during the early part of the twentieth century. The Second World War witnessed a flurry of development of spread spectrum modulation in order to secure communication channels as well as

develop more effective radar. In this section, we presented the fundamental concepts of spread spectrum signals, including the six noteworthy characteristics of: selective addressing, code division multiplexing, low power density, message protection, high resolution ranging, and interference rejection. The processing gain and jamming margin used to characterize the performance of spread spectrum systems were also introduced. Lastly, we discussed two popular modulation techniques—direct sequence and frequency hopped—and comparisons were made between the two.

In the following section, we will focus our attention on discussing direct sequence spread spectrum in detail.

Table 1. Tradeoffs in direct sequence and frequency hopped systems [3]

Direct Sequence	Frequency Hopped
Good range resolution	Poor range resolution
More noise like signal	Narrow instantaneous bandwidth
Can operate below ambient noise	Must have positive S / N ratio
Can operate without error correction codes	Requires error correction codes
Limited near-far performance	Best near-far performance
Requires linear signal path	Operates well with nonlinearity in signal path
Synchronization more difficult	Easier to synchronize

3. Direct Sequence Spread Spectrum

Of the various spread spectrum modulation techniques available today, the best-known and most widely used method remains direct sequence spread spectrum (DSSS). Dixon, in his third edition book on spread spectrum systems, attributes this popularity to the relative simplicity of the method—it does not need complex frequency synthesizers like the frequency hopping method requires [3]. In the previous section, we discussed the two requirements for a signal to be classified as spread spectrum. Direct sequence signals meet these requirements by employing a pseudo-noise sequence that is independent of the information data. These sequences are used to spread the spectrum over a bandwidth that is much greater than that of the information being transmitted.

In the following sections, we will examine direct sequence spread spectrum in great detail. We will first look at the basic concepts of DSSS signals, including baseband modulation and demodulation. The following section builds upon the knowledge of the characteristics of spread spectrum signals gained previously by going into greater detail on three of them: code division multiplexing, low power density, and interference rejection. After addressing the performance specifications of DSSS (processing gain and jamming margin), we will delve into pseudo-noise sequences. These sequences are responsible for all of the desirable characteristics of DSSS signals. Two types of PN sequences will be studied: maximal sequences and Gold codes.

Many figures will be presented in this section in order to facilitate an understanding of the concepts and ideas. These figures are found throughout the available literature on spread spectrum. For this report, I have chosen to use the figures presented in the white papers available from Sirius Communications in Rotselaar, Belgium for their clarity and ease of understanding. These papers were compiled by ir. J. Meel in December 1999 [8, 9].

3.1 Basic Concepts of DSSS

The purpose of this section is to familiarize the reader with the basic concepts of DSSS systems. The various components and functionality of a DSSS system are presented and discussed, leading to a simplified description of the modulation and demodulation of these types of signals. The figures presented in these sections are particularly useful for understanding the principles of frequency spreading because they show both the time and frequency domains.

3.1.1 Block Diagram

Throughout this report we have made references to DSSS systems and the components that create the spread spectrum signals. Figure 4 presents a block diagram of a DSSS system. This diagram also points out that all of the spreading and despreading is performed baseband—the bandpass region is only used to modulate the coded signal onto useable frequencies. As a matter of fact, we will be concerned solely with baseband when we examine the modulation and demodulation of DSSS signals.

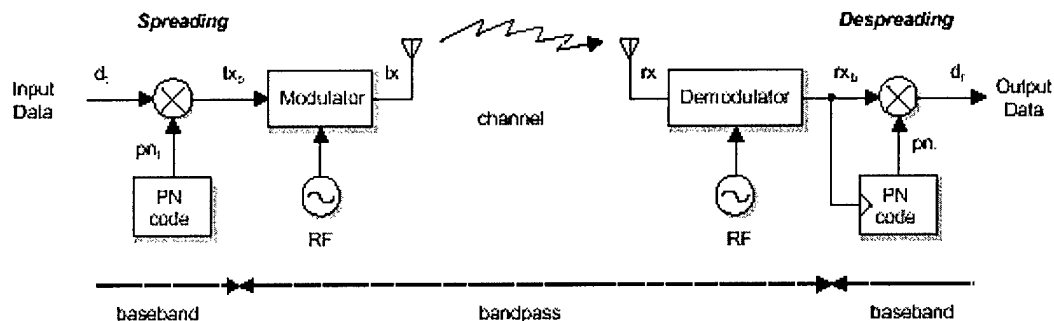


Figure 4. Block diagram of a DSSS system [9].

It is important to note that this is a digital system; therefore the input data

must be digitized before it is combined with the code sequence. The resultant output is also in digital format, which must also be taken into consideration.

The pseudo-noise code is generated at the modulator, and multiplied with the input information signal. The result of combining the input information with the PN sequence is the spreading of the bandwidth of the data. The resultant bandwidth is equal to that of the PN sequence. After spreading the spectrum of the information, the signal is then ready to be modulated onto an RF carrier and transmitted over a channel.

The recovery of the original information is the goal of the demodulator and despreader. One of the characteristics of spread spectrum signals is that the signal cannot be detected by means of a conventional narrowband receiver. Signal recovery, then, is reserved for those receivers who know the proper sequence and are able to synchronize their local sequence generators with that being received. The process of despreading once again involves the multiplication of the spread signal with the PN sequence.

From Figure 4 and the preceding discussion, it should not be difficult to appreciate the simplicity of this implementation of DSSS. The following two sections go into more detail on the modulation and demodulation of DSSS signals.

3.1.2 Modulation

Direct sequence modulation is fundamentally the modulation of a carrier by a code sequence. The most common method of implementation is with a 180° biphasic phase-shift keying (PSK) carrier [3]. This carrier is pseudorandomly shifted by a pseudo-noise sequence generated at the modulator. These shifts characterize the signal and are responsible for the desired frequency spreading. Balanced forms of modulation such as PSK and binary PSK (BPSK) are useful for transmitting codes because one phase can be used to transmit a one, and the

opposite phase used to transmit a zero. There are three main reasons why balanced modulation methods akin to PSK are used as opposed to other forms of modulation [3]:

1. The suppressed carrier characteristic of these types of signals is difficult to detect.
2. More power is available for sending information because it is not wasted in a carrier.
3. The transmitted power efficiency is maximized for a given bandwidth because the signal has a constant level envelope.

Referring to Figure 4, let us examine the process of modulation for DSSS systems. In order to spread the bandwidth of a digital information signal, it must be combined with a pseudo-noise sequence. After multiplying the information with the pseudo-noise sequence, the bandwidth of the resultant signal is equal to the bandwidth of the PN sequence, but it contains the information signal within it. Figure 5 provides graphical examples of the process of bandwidth spreading in both the time and frequency domains.

The plots labeled by (a) in Figure 5 display a data signal that has a symbol period of T_s . Performing a Fourier transform of the signal, we see that it resembles a sinc function (keep in mind that we are looking at the power spectra, so it is actually a function of sinc^2). When examining the bandwidth of the signals, we are usually concerned with the 3-dB bandwidth, but in examining these ideas, we will look at the bandwidth of the first lobe. Therefore, the bandwidth of the symbol is equal to twice the symbol rate. The symbol rate is merely the inverse of the symbol period.

Plots labeled (b) in Figure 5 show the pseudo-noise sequence that will be used to spread the bandwidth of the information signal. Notice that the chip

period, T_c , is much shorter than the symbol period T_s . Examining the sequence closely, we can see that there are an integer number of chips per symbol—seven to be exact. The Fourier transform of this signal is much broader. This is to be expected since there are a greater number of variations in the signal. If we would like to reproduce this particular sequence using the Fourier series, we would require many more frequency components than the data signal in (a).

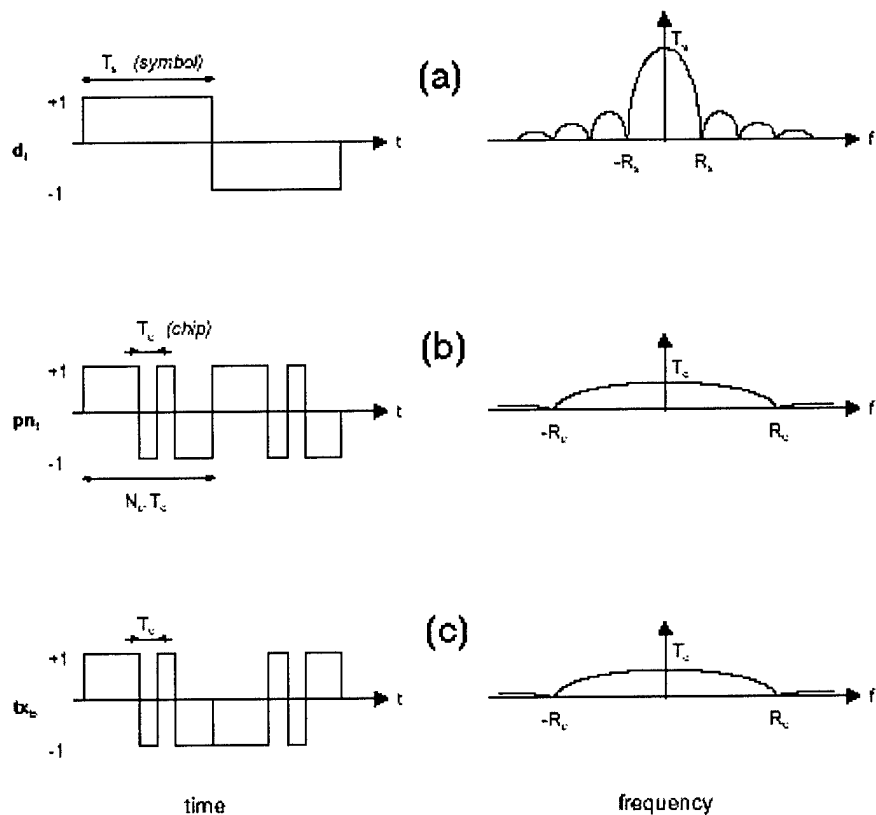


Figure 5. Modulation of a DSSS signal showing both the time and frequency domain for (a) the data, (b) the PN sequence, and (c) the resultant spread spectrum output [9].

The remaining plots displayed in Figure 5 show the result of combining

the data signal with the PN sequence. Notice that the resulting sequence for the data symbol 'one' is exactly opposite that for the data symbol 'zero.' Drawing our attention to the frequency spectrum, we see that the bandwidth of the combined data and PN sequence is equal to that of the PN sequence. This bandwidth is much larger than the original data, so spreading has been achieved. The spreading factor, or the ratio of the bandwidth of the PN sequence to that of the data, is (5):

$$S.F. = \frac{R_c}{R_s} = \frac{BW_{SS}}{BW_{data}} = N_c \quad (5)$$

Equation (5) should look familiar. It is simply the processing gain of a spread spectrum system and can be calculated easily by counting the number of chips per data symbol. In this section we were able to see how the modulation of a DSSS worked graphically.

3.1.3 Demodulation

Just as important as the process of modulating a DSSS signal is an understanding of how a DSSS modulated signal is demodulated. The key to successfully demodulating a DSSS modulated signal is to recreate the PN sequence that was used to code the data exactly, and to then synchronize the local PN sequence with the received signal. Once the local replica of the PN sequence is generated and synchronized, all that remains is to once again multiply the received data with the PN sequence. The result is a despread signal that identically matches the input data after narrow-band filtering [1]. If, however, the PN sequence is not an exact replica or synchronized properly, the output will not be indistinguishable from the original data.

Figure 6 graphically describes the demodulation of a DSSS modulated signal in both the time and frequency domains. Plots (a) display the received DSSS modulated signal. As we can see, the period is still the chip period of the encoding PN sequence. The Fourier transform reveals a wide bandwidth, equal to twice the chip rate of the pseudo-noise sequence generator.

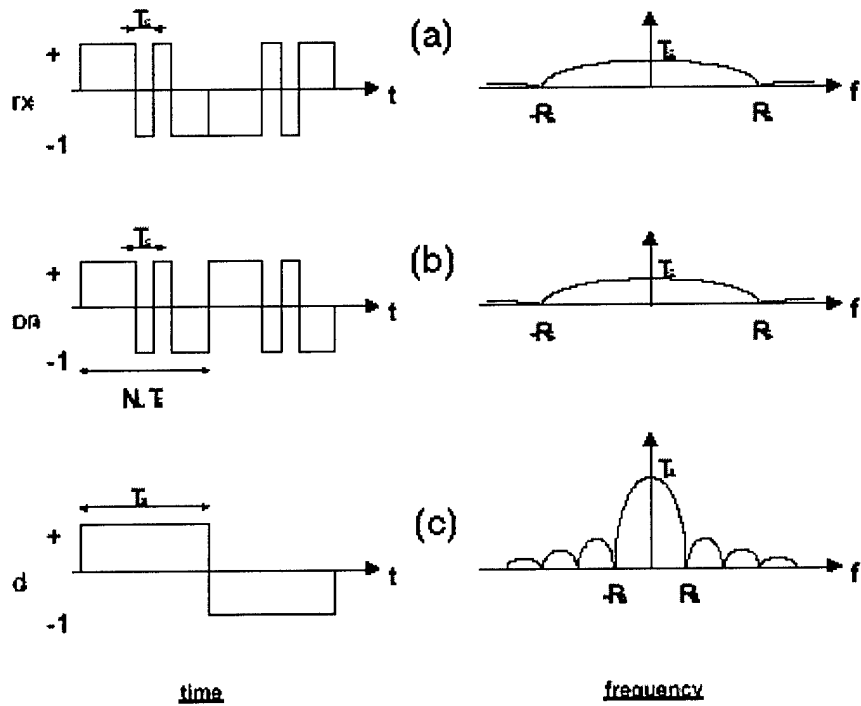


Figure 6. Demodulation of a DSSS signal showing both the time and frequency domain for (a) spread spectrum input, (b) PN sequence, and (c) original data [9].

The (b) plots shown in Figure 6 represent the locally generated pseudo-noise sequence that will be used to despread the received signal. In order to work properly, the sequence is an identical replica of that used to code the input data. It is also synchronized properly – that is, it begins when a symbol begins, and ends

when a symbol ends. As expected, the bandwidth of the PN sequence is equal to that of the received signal.

After multiplying the PN sequence with the received signal, we have recovered the original data, as shown in the remaining plots of Figure 6. The symbols have been recovered, and the bandwidth of the information signal is much smaller than that of the DSSS modulated signal.

3.2 Characteristics of DSSS

Six characteristics of spread spectrum modulation techniques were presented in section 2.3. In this section, we will expand on three of these characteristics and how they arise from DSSS modulation. The characteristics to be discussed are code division multiplexing, low power density, and interference rejection.

3.2.1 Code Division Multiplexing

As addressed in section 2.3.2, spread spectrum signals can be used to deliver reliable and simultaneous communications for multiple users over a single channel. The bandwidth efficiency of the system increases with the number of users who employ such a system. The method of using codes to uniquely identify users on such a multi-user network is known as code division multiple access (CDMA). The benefits of CDMA systems over frequency or time division multiple access systems are listed here [5]:

- CDMA does not require bandwidth allocation, as is the case with FDMA systems. Each user in a CDMA system is able to use the entire frequency spectrum available to the channel.
- CDMA does not require time synchronization, as is the case with TDMA systems. Users of CDMA systems are free to use the system at any time.

Although CDMA systems possess the capability to service multiple users over its entire frequency spectra at all times, it suffers from two noteworthy problems. Both of these problems stem from the interference caused by multiple users sharing the same bandwidth.

The first problem is that as the number of simultaneous users increase, the interference caused by them on a single receiver also increases. The noise floor is determined after decorrelation by the total power at a receiver due to multiple users [5]. The limit on the capacity of a CDMA system is determined by the amount of interference that can be withstood without disrupting communications. In comparison, the capacity of FDMA or TDMA systems is bandwidth limited. Figure 7 shows the multi-user environment of a CDMA system.

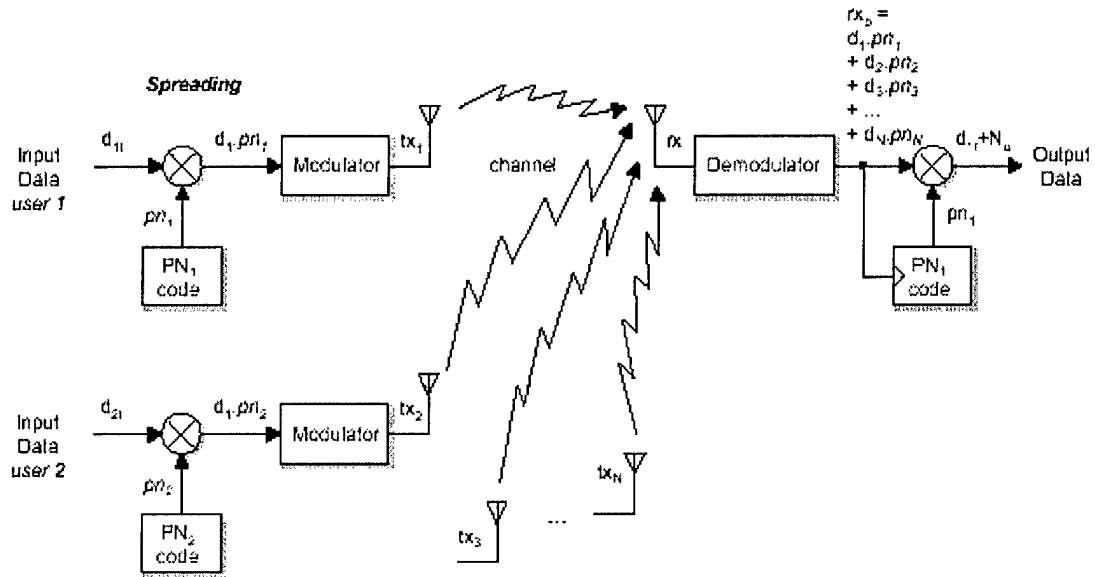


Figure 7. Block diagram of a multiple access network via DSSS [9].

The second problem of CDMA systems is known as the near-far problem

[5]. This situation occurs when many users share the same channel but are not uniformly spaced from the receiver. In this instance, a transmitter who is significantly closer to the receiver will generate more interference. In order to mitigate the problem, power control is used in most CDMA implementations [5]. By controlling the power, the signal level at the receiver can be made constant for all uses of the system. This produces a uniform interference from each user, preventing any single user from capturing the receiver.

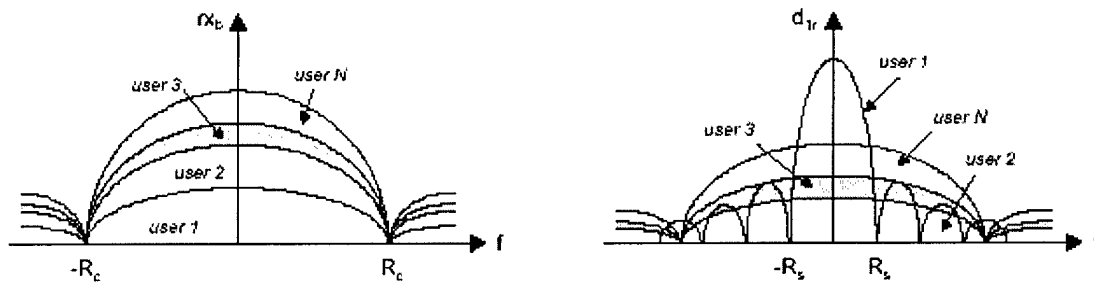


Figure 8. Frequency spectra showing modulated and demodulated users on a multiple access DSSS network [9].

In Figure 8, we see the interference effects of multiple users on a CDMA system. In the plot on the left, we have multiple users each spread using a unique pseudo-noise sequence. At the receiver, the PN sequence that matches the desired user reveals the original message from that user. The energy of other users, however, still contributes to the interference problem in the form of noise.

3.2.2 Low Power Density

As we discussed earlier, signal hiding has been one of the primary motivations of military research in the field of spread spectrum since the Second World War. We have shown that the amount of spreading of a signal is

proportional to the length of the PN sequence; we can use very long sequences to drive the average power extremely low. Direct sequence spread spectrum systems can operate below the ambient noise of the channel, as shown in Figure 9.

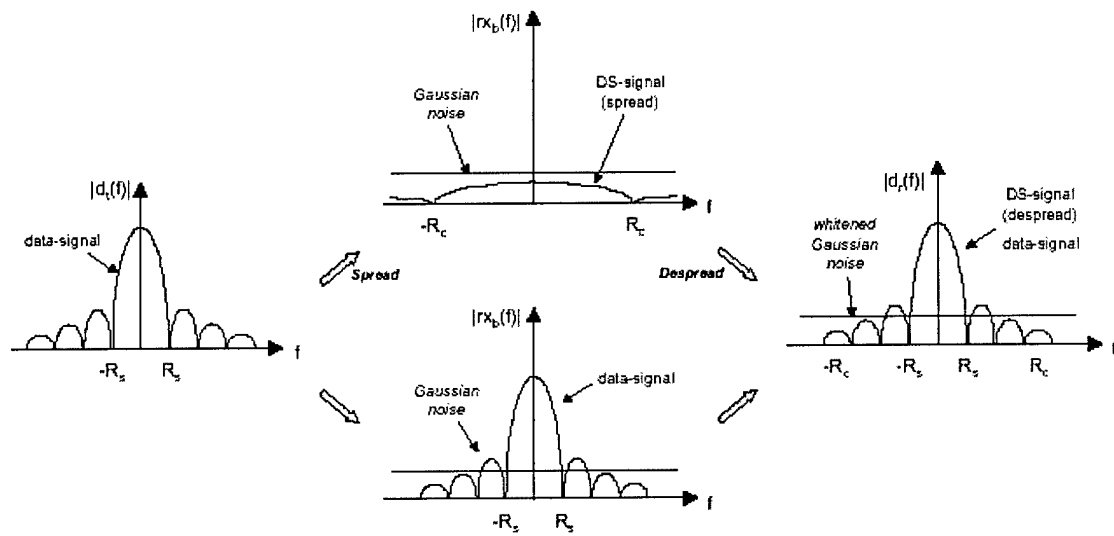


Figure 9. Comparison between the frequency spectra of a DSSS modulated and unmodulated signal in the presence of Gaussian noise [9].

The leftmost plot in Figure 9 shows the original data signal. This is the same spectrum we examined in previous sections. It is characterized by its main-lobe bandwidth. The two center plots are of greater interest to this section. The bottom plot shows what the data signal looks like on the channel in the presence of average white Gaussian noise. This signal is not hidden at all, and can be easily detected by casual eavesdroppers.

Examining the top plot shows us the benefits of DSSS modulation. In this plot, the transmitted signal has been spread so much that it is below the level of ambient noise on the channel. Although the signal has the same total power, the

average power is significantly reduced. As such, a casual eavesdropper would not notice the presence of any communications over this frequency range. The plot on the right shows the received and despread signal, displaying the noise remaining after despreading.

3.2.3 Interference Rejection

Previously, we introduced two types of interference that are capable of disabling conventional communication methods. The two types of interference are called either intentional or unintentional. The type of interference has no bearing on the nature of the bandwidth of the interfering signal. For example, an intentional jammer may cause interference over a very wide band, while a radio station may cause narrowband interference if the spread spectrum signal is operating directly under it. We have suggested that spread spectrum modulation techniques are capable of providing reliable communications in the face of such interference. In this section we will focus on how DSSS modulation techniques mitigate the effects of both narrowband and wideband interference.

Narrowband interference is localized in frequency. For comparison, the information signal before DSSS modulation is narrowband. When we modulated the spread spectrum signal by multiplying the data signal with the PN sequence, we effectively spread the bandwidth of the information signal. In the same manner, the narrowband interference is spread when the combined signal is demodulated in the receiver. Figure 10 displays the process by which the narrowband interferer is while the data remains unaffected.

Examining the plot on the left, we see that it is the same data signal we have been using throughout this report. After DSSS modulation, the data is spread and a narrowband interference source is introduced into the channel, as shown in the center frequency spectrum plot. After the combined signal is received, it is multiplied by the local pseudo-noise sequence that extracts the

original information signal. The narrowband interference, however, has been spread and therefore has a minimal impact upon the output data.

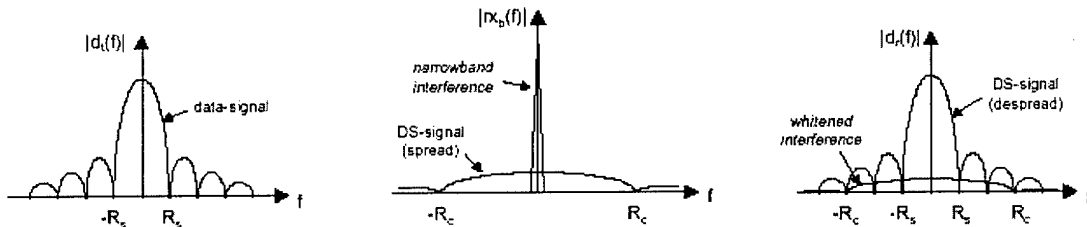


Figure 10. Effect of narrowband interference on a DSSS signal [9].

Although DSSS demodulators can eliminate virtually all of the effects of narrowband interference, this may not necessarily be the case with wideband interference. Lightning strikes are a good example of a source of wideband interference. They contain frequency components across a wide portion of the electromagnetic spectrum. Many users simultaneously accessing a CDMA network also cause wideband interference, as discussed earlier.

Fortunately, DSSS demodulators are also capable of reducing the effects of wideband interference. Recall that in each receiver is a pseudo-noise generator that recreates the specific sequence that is used to demodulate the received signal. If the local pseudo-noise sequence does not correlate well with the interfering signal, the result is a further spreading of the interference. In the case of narrowband interference, the reduction of the effects of the noise was significant because its total power was spread across a wide bandwidth. With intentional jammers or even in the near-far problem discussed before, the power density is much higher. If the interference source is strong enough, the original message may be unrecoverable.

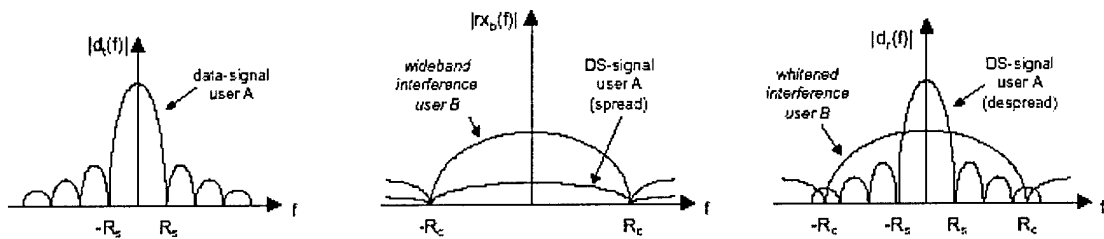


Figure 11. Effect of wideband interference on a DSSS signal [9].

In Figure 11, we graphically depict the process by which a signal is recovered in the presence of wideband interference. In the center plot, we see that the interfering source has a higher power spectral density than our spread data signal. Fortunately the data signal was recovered through despreading in the plot on the right, while the wideband interference is uncorrelated with the local PN sequence and is spread. Notice, however, that the power density of the interfering source remains significantly higher than that of a narrowband interferer. This is the reason behind limiting the number of users accessing CDMA network; there could be too much interference resulting in a loss of valuable information.

3.3 Performance Specifications of DSSS

Recall from section two that the processing gain provides us with information regarding the improvement in signal-to-noise ratios between the input data and output transmission of a system. For direct sequence systems, the process gain is a function of the RF bandwidth of the transmitted signal, as compared to the information bit rate [3]. When introducing DSSS modulation, we stated that the bandwidth of the direct sequence spread signal is equal to that of the main lobe of the sinc^2 power density spectrum. This value, however, is not equal to the bandwidth-spreading PN sequence clock rate. In fact, the bandwidth of the main lobe is approximately 88 percent of the clock rate of the pseudo-noise

generator. Therefore, when corrected, the processing gain equation for direct sequence spread spectrum systems becomes,

$$G_p = \frac{BW_{RF}}{BW_{data}} = \frac{0.88R_c}{R_b}.$$

The other parameter used to specify the performance of spread spectrum systems is the jamming margin. Since the jamming margin of a system takes into consideration the internal losses of that particular system, we cannot specify it in this generalized discussion of direct sequence spread spectrum.

3.4 Pseudo-Noise Sequences

Pseudo-noise sequences provide a common thread between the various types of spread spectrum modulation. Their noise-like properties are responsible for the spreading of the information signals that give rise to the valuable characteristics of spread spectrum modulation discussed previously. In this section, we will address pseudo-noise sequences in detail, examining both their characteristics and properties. Two types of PN sequences will be investigated in this discussion: maximal sequences and Gold codes.

3.4.1 Average White Gaussian Noise vs. Pseudo-Noise

In a prior section, we introduced the concept of a white Gaussian noise source that is present on virtually every communication channel—especially wireless channels—during a discussion of the low power density characteristic of DSSS modulated signals. White noise is defined to be a stochastic process that has a uniform power spectral density over the entire frequency spectra. The autocorrelation, or degree of similarity between two identical signals, for a white noise signal is an impulse function. The impulse occurs only when the two

identical noise signals overlap exactly. Any cyclic shift of one of the signals causes the autocorrelation function to be reduced to zero. Restating this property mathematically, we have [4],

$$\phi_{nn}(k) = \frac{N_0}{2} \delta(k) = \begin{cases} N_0/2, & k = 0 \\ 0, & \text{otherwise} \end{cases}$$

where k denotes discrete time; and

N_0 represents the power spectral density of the noise.

In developing spread spectrum techniques, we attempt to reproduce this noise-like quality in our locally generated periodic sequences because such sequences offer maximum spreading of the information bandwidth. Because we must create reproducible periodic sequences, these sequences are not truly random. Instead, the PN sequences we use to modulate our information in a spread spectrum system are deterministic—the ability to recreate the sequence must be both known to both the transmitter and receiver. This is the reason why such sequences are known as pseudo-noise or pseudorandom sequences. As the length of the PN sequence increases, its characteristics become similar to a truly random signal.

3.4.2 Properties of PN Sequences

Pseudo-noise sequences possess several properties that distinguish them from other codes, such as those used for error-correction or cryptography for example. PN sequences are classified along with the set of sequences that are known as periodic sequences. A periodic sequence is one which consists of an infinite sequence of plus and minus ones divided into blocks of length N . The sequence contained in each of the blocks is identical—giving the sequence its

periodic nature. A sequence such as this is said to be pseudorandom if it satisfies the following three properties [1]:

Balance

A PN sequence is balanced when in each period, the number of ones and the number of minus ones differ by exactly one. The periodic sequence length, then, must be an odd number. The balance property is important because carrier suppression depends upon the symmetry of the modulating signal. The following equations explain the balance property mathematically.

$$\begin{aligned} N_{+1's} + N_{-1's} &= N \\ |N_{+1's} - N_{-1's}| &= 1 \end{aligned}$$

Run-length Distribution

When describing sequences, the term run is used to describe a sequence of a single type of binary digits. For example, if in a sequence there appear four consecutive ones, we can say that there is a run of four ones.

For a PN sequence to be considered pseudorandom, the distribution of runs in a period follows a specific distribution. One half of the runs in a period have length one, one forth have length two, one eighth have length three, and so on. In addition, the number of positive and negative runs is equal.

Autocorrelation

All correlation functions provide information on the degree of correspondence between a sequence and a phase-shifted replica of itself. The autocorrelation function is then defined as the number of agreements minus the number of disagreements in a term-by-term comparison of one full period of the

sequence with a cyclic shift of the sequence itself [3]. As we stated earlier, the autocorrelation of a white noise signal takes one the form of an impulse function. For our PN sequence to be considered pseudorandom, it too must have an autocorrelation that resembles an impulse. Therefore, the autocorrelation for a periodic sequence must be two-valued, as shown mathematically.

$$\phi_{aa}(k) = \sum_{n=1}^N a_n a_{n+k} = \begin{cases} N, & k = 0, N, 2N, \dots \\ -1, & \text{otherwise} \end{cases}$$

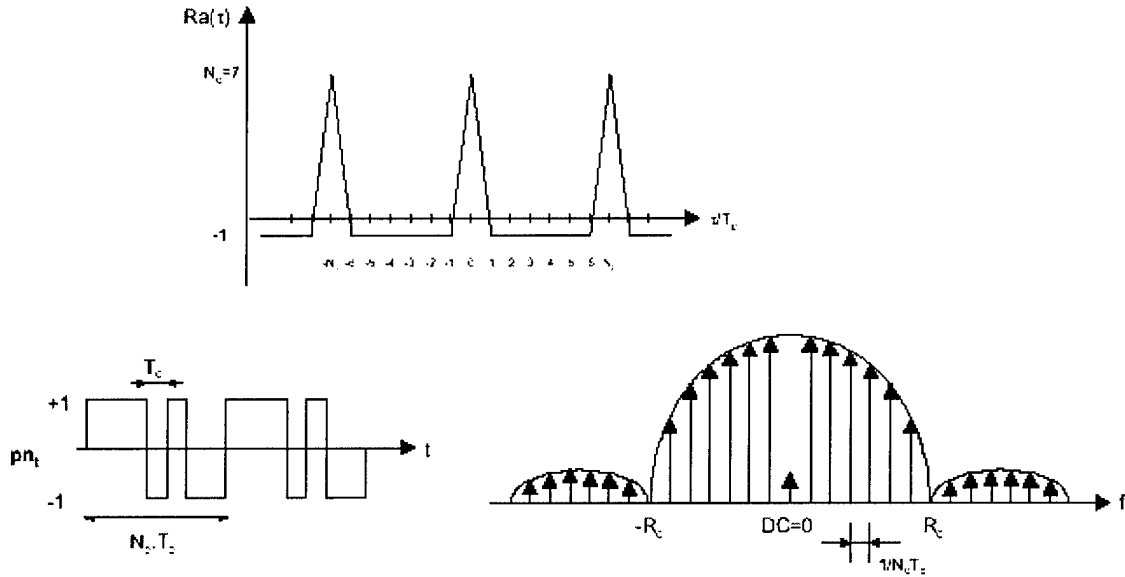


Figure 12. Autocorrelation function and frequency spectra of a PN sequence [9].

In Figure 12 we have a two periods of a seven-chip PN sequence. When two identical PN sequences are synchronized, the autocorrelation achieves its maximum value equal to the length of the sequence. For the given PN sequence, this autocorrelation is shown in the uppermost plot. The frequency spectrum plot shown in Figure 12 graphically explains the balance property of PN sequences.

The minimized DC component is due to the almost identical distribution of positive and negative ones that compose the sequence.

Cross-Correlation

Although not a necessary condition for a sequence to be classified as pseudorandom, another useful property is cross-correlation. While the autocorrelation measures the amount of agreement between two identical signals, the cross-correlation is used to measure the similarity between two different signals. For CDMA applications, where many users access the same bandwidth simultaneously, sequences with minimum cross-correlation values are vital to reducing mutual interference. If the cross-correlation between two sequences can be reduced to zero, the sequences are said to be orthogonal to one another; there is no interference from other users after despreading.

3.4.3 Types of PN Sequences

Based on our earlier discussions of the characteristics of spread spectrum signals, we will now present two types of pseudo-noise sequences that are widely used in a variety of communication systems. The first, called maximal sequences, are the most widely known binary PN sequences and have excellent autocorrelation properties that make them very useful for general spread spectrum communication systems. Another set of PN sequences that are particularly useful for CDMA applications due to their cross-correlation properties are known as Gold codes.

3.4.3.1 Maximal Sequences

By definition, maximal sequences are those sequences that possess the longest codes possible for a given shift register or delay element of a given length. Maximal length sequences are unexcelled for general use in communications and

ranging because of their fantastic autocorrelation properties and simplicity of design. They are widely used due to the ease with which they are generated using shift registers with a relatively small number of stages [1]. Other codes can simply do no better than equal their performance. It is important to note that maximal sequences are linear codes; they are easily decoded and as such should not be used for secure communications.

The length of a maximal sequence is governed by the number of shift register stages used in the construction of the generator. For an n stage shift register generator, the resulting maximal sequence length is $2^n - 1$ chips. If the sequence generator somehow enters the all-zero state, it will remain in that state indefinitely. This is why the maximal length is one less than 2^n —the generator must not enter the all-zero state. In Figure 13, we have a block diagram of a simple shift register generator (SSRGG). Such generators are composed of a series of shift registers with feedback from an even number of taps.

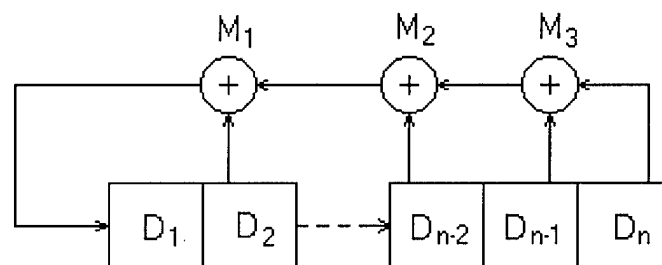


Figure 13. Simple shift register generator [3]

Table 2 lists several feedback connection configurations for generating maximal sequences of varying lengths using shift register generators.

As mentioned earlier, maximal length sequences have excellent autocorrelation properties. These sequences are thus extremely useful for general communications and ranging applications. Sequences that have good

autocorrelation properties have the least probability of a false synchronization. False synchronization occurs when the receiver incorrectly synchronizes its locally generated PN sequence with the sequence being received. The autocorrelation for a five stage, 31-chip maximal sequence is shown in Figure 14. Notice that the autocorrelation takes on only two distinct values over the period of interest, and that its maximum value is equal to the sequence length as discussed earlier in the properties of PN sequences.

Table 2. Feedback connections for linear maximal sequence generators [3]

Number of Stages	Code Length	Feedback taps for m-sequences	Number of m-sequences
2	3	[2,1]	2
3	7	[3,1]	2
4	15	[4,1]	2
5	31	[5,3] [5,4,3,2] [5,4,2,1]	6
6	63	[6,1] [6,5,2,1] [6,5,3,2]	6
7	127	[7,1] [7,3] [7,3,2,1] [7,4,3,2] [7,6,4,2] [7,6,3,1] [7,6,5,2] [7,6,5,4,2,1] [7,5,4,3,2,1]	18
8	255	[8,4,3,2] [8,6,5,3] [8,6,5,2] [8,5,3,1] [8,6,5,1] [8,7,6,1] [8,7,6,5,2,1] [8,6,4,3,2,1]	16

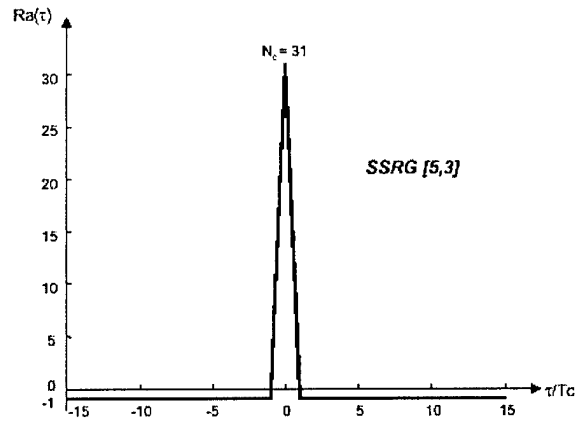


Figure 14. Autocorrelation of a maximal sequence [9].

The cross-correlation of the previous sequence with two other five-stage, 31-chip maximal sequences is shown in Figure 15. While the autocorrelation peak for maximal sequences is extremely high, these cross-correlation values are also quite large when compared to other code generation methods. This is of great concern for multiple access applications, especially if the transmitted power is high enough to raise the peak cross-correlation to a value that is near peak autocorrelation [3]. A way to alleviate concerns over the cross-correlation between sequences is to use a composite code, such as a Gold code.

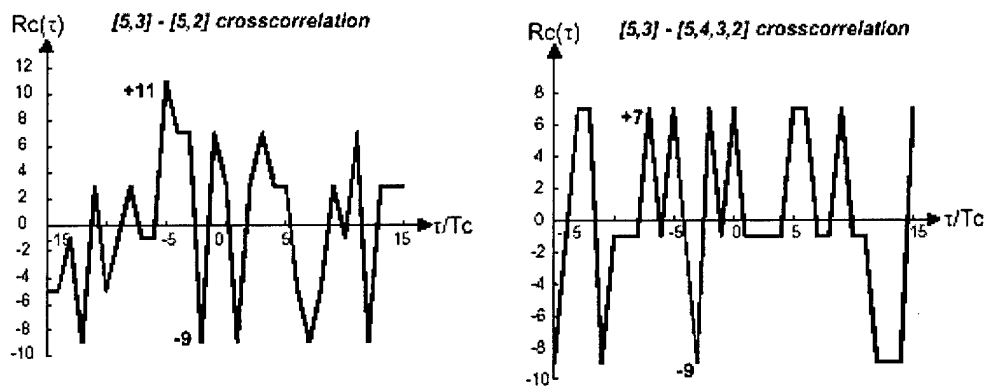


Figure 15. Cross-correlation of maximal sequences [9].

3.4.3.2 Gold Codes

Characterized by their well-defined cross-correlation properties, Gold codes are much more effective in multiple access environments such as CDMA [4]. Although Gold codes are constructed from maximal sequences, they are not themselves maximal sequences. Another advantage of Gold codes is that they allow construction of families of $2^n - 1$ codes that are of equal length and have specified-level cross-correlation values [3].

Gold codes are generated via modulo-2 addition of a pair of maximal sequences. Figure 16 provides a block diagram depicting such a Gold code generator. The maximal sequences used to generate Gold codes are the same length, which allows the two generators to maintain an identical phase relationship. The resultant codes are the same length as the codes that are added together, but are not maximal.

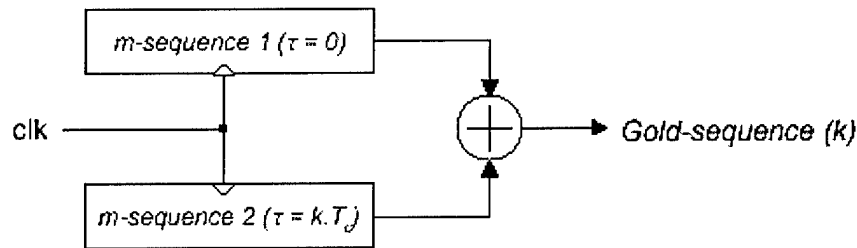


Figure 16. Block diagram of a Gold code generator [9].

Gold code generators take advantage of the fact that for every change in phase position between the two generators, a new sequence is generated. Going back to Table 2, we see that for a five-stage maximal sequence generator, we can create only six unique sequences. Since there are $2^n - 1$ possible phase shifts—recalling that the all-zero state must always be avoided—we are able to create $2^n - 1$ unique Gold codes. In addition to the Gold codes created, we can also use the

base maximal sequences used by the generator. Therefore, for a two-register, five-stage Gold code generator, we can create 33 unique Gold codes. This is a vast improvement over the number of codes possible for maximal sequences, which means that in a multi-user environment more people can access the system.

Due to the fact that Gold codes are not maximal, their autocorrelation values are significantly different than those of maximal sequences. A plot of the autocorrelation of a two-register, five-stage Gold code is shown in Figure 17. Notice that the resulting autocorrelation is not two-valued, as was the case for maximal sequences.

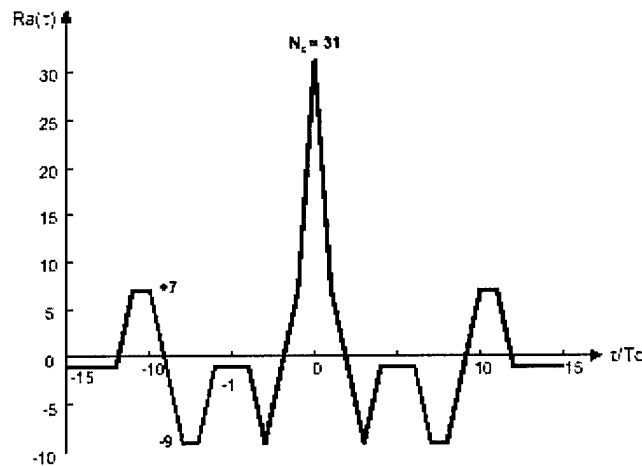


Figure 17. Autocorrelation of a Gold code [9].

Besides the creation of a vast number of unique sequences, Gold codes also possess well-defined cross-correlation properties that make them very useful for multiple access applications. These properties arise from the proper selection of maximal sequences, known as preferred maximal sequences. When such sequences are used, the resulting Gold codes have a three-valued cross-correlation. This important subset of Gold codes is called the Preferred Pair Gold

codes [9]. Table 3 lists several preferred pairs of maximal sequences that can be used to generate Preferred Pair Gold codes, while Figure 18 plots the cross-correlation for various Gold codes. These gold codes were created by the two five-stage preferred pair maximal sequences shown in Table 3.

Table 3. Preferred pairs for Gold code generation [3]

Number of Stages	Preferred pairs of m-sequences	cross-correlations		
5	[5,2] [5,4,3,2]	7	-1	-9
6	[6,1] [6,5,2,1]	15	-1	-17
7	[7,3] [7,3,2,1] [7,3,2,1] [7,5,4,3,2,1]	15	-1	-17
8	[8,7,6,5,2,1] [8,7,6,1]	31	-1	-17
9	[9,4] [9,6,4,3] [9,6,4,3] [9,8,4,1]	31	-1	-33

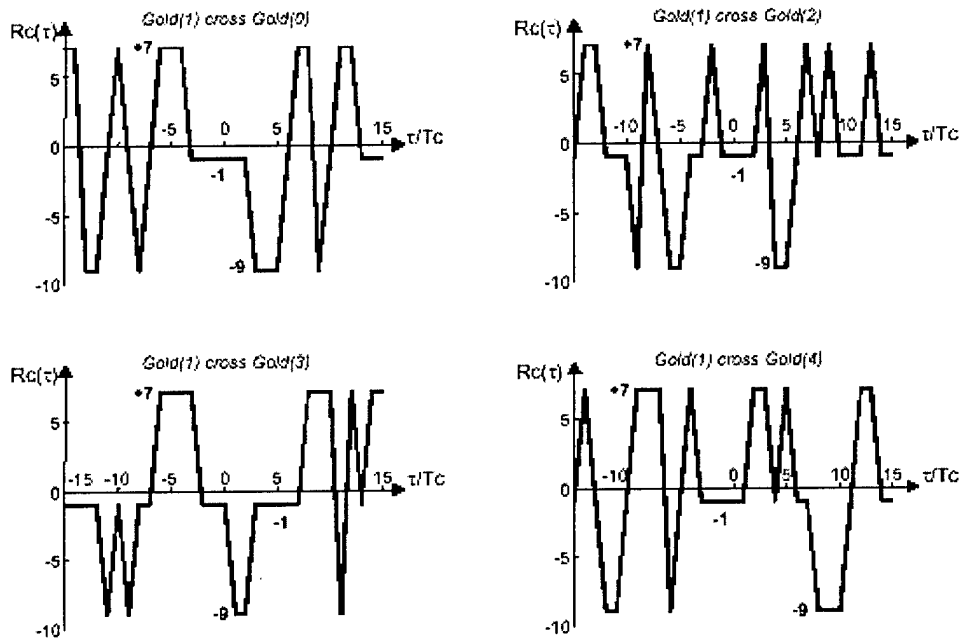


Figure 18. Cross-correlation of various Gold codes [9].

Notice that the cross-correlation values shown in Figure 18 all fall into one of three discrete levels: +7, -1, and -9, as shown in Table 3. By knowing the cross-correlation values beforehand, receivers can be designed with a better understanding of potential interference implications of using Gold codes.

3.5 Summary

In this section, we examined the direct sequence spread spectrum modulation technique in detail, examining modulation and demodulation, as well as several key characteristics. The use of unique pseudo-noise sequences provides this modulation method with the ability to function superbly in a multiple access environment. Implementing longer sequences allows for an even greater spreading of the information bandwidth—so much so that the spectral power density of the signal can drop below the level of the ambient noise on the channel. The signal has thus, by all intents and purposes, disappeared—it has been effectively hidden from the casual, uninformed observer. The last characteristic that was discussed is perhaps the most important for both military and commercial applications. DSSS modulation techniques are able to withstand the effects of both narrowband and wideband interference due to the spreading of the signal energy.

After briefly correcting the general definition of processing gain introduced in section two, pseudo-noise sequences were then investigated. All of the benefits and characteristics of direct sequence are the direct result of the use of these pseudorandom noise-like sequences. Having addressed several properties of PN sequences, we then looked at two popular sequences—maximal sequences and Gold codes. Maximal sequences are unexcelled in general communication and ranging applications due to their autocorrelation properties. All other codes can do no better as far as autocorrelation is concerned. Gold codes, although

generated from maximal sequences, are not themselves maximal. However, Gold codes have very well defined cross-correlation properties, which makes them perfect for CDMA applications.

Next, we will present various examples of commercial applications for spread spectrum technology.

4. Commercial Applications

For the greater part of a half-century, spread spectrum systems were used primarily by the military. Many commercial applications have arisen as spread spectrum technology became declassified and available to the public. The widespread demand for wireless communication networks, along with advances in VLSI and signal processing techniques have pushed spread spectrum to the forefront of commercial development. In May 1995, the Federal Communications Commission released the ISM (Industrial, Scientific, and Medical) bands for direct sequence and frequency hopped spread spectrum use [3]. Since that historic occasion, many commercial products have entered the marketplace.

In previous sections, we have presented several characteristics of spread spectrum systems that make it a very attractive technology for commercial development. Its ability to withstand interference allows for communications in high interference environments, such as a manufacturing floor. In view of the fact that SS signals have a very low power density, they can be overlaid on bands where systems are already operating without causing interference to those systems. We have also shown that CDMA spread spectrum systems offer great flexibility and capacity for wireless multi-user networking applications.

In this section of the report, we briefly describe three valuable commercial applications of spread spectrum technology. Although the Global Positioning System (GPS) was developed for military navigation and location, it stands as the most famous implementation of spread spectrum technology, and many commercial receivers can be purchased today. The greatest commercial interests in spread spectrum, however, lie in the areas of Personal Communications Systems (PCS) and Wireless Local Area Networks (WLANs).

4.1 Global Positioning System

The Global Positioning System is a satellite navigation system that is funded and controlled by the U.S. Department of Defense. Civilian applications of GPS are numerous and include surveying, monitoring of earthquake fault lines, positioning and navigation, and time synchronization. Three building blocks compose the GPS system: the space segment, user segment, and control segment [8]. Figure 19 illustrates these building blocks, as well as the flow of information in the system.

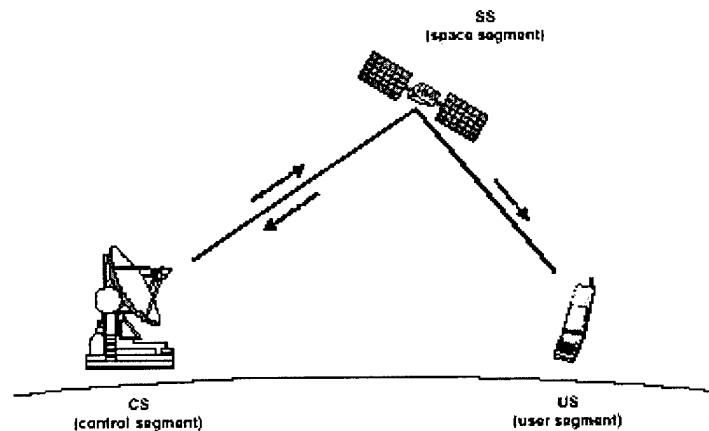


Figure 19. Building blocks of GPS [8].

Of interest to us are the signals generated in the space segment of the GPS system. Each GPS satellite transmits two signals. The first, known as the L1 frequency (1575.42 MHz) carries the navigation messages and the Standard Positioning System (SPS) code signals. The SPS is provided to civilian users worldwide, with the accuracy of the positioning intentionally degraded by the DOD [8]. The second signal transmitted by each satellite is called the L2 frequency (1227.60 MHz), and is used to measure delays caused by the ionosphere. The carrier phase of the L1 and L2 signals are shifted by three binary

codes listed below [8].

Coarse Acquisition Code

The Coarse Acquisition (C/A) code modulates the phase of the L1 carrier. It is a periodic 1023 chip pseudorandom noise code that is generated at a rate of 1.023 Mchip/second. The bandwidth of the information is spread by this code to over one megahertz. Using the processing gain equations presented earlier, we find that such a length code has a processing gain of 30 dB.

The GPS system also employs code division multiplexing techniques. Each GPS satellite has a unique C/A code that is used to identify them, even though they all share the same L1 and L2 frequencies. To avoid any interference problems, a Gold code is used to minimize the cross-correlation between individual C/A codes.

Precise Code

The precise code, or P-code, is used to modulate both L1 and L2 carrier phases. The precise code is used for the Precise Positioning System (PPS), whose authorized users require special cryptographic equipment and keys in order to use them. The accuracy of the information provided by the PPS is significantly better than that of the SPS.

This code is the same for all satellites, and is extremely long. The P-code is a pseudorandom noise code whose length is 6.19×10^{12} chips. At a transmission rate of 10.23 million chips per second, it takes seven days to transmit the entire sequence.

Navigation Message

The navigation message is simply a 1500 bit message that provides information describing the GPS satellite orbits, clock corrections, and other

system parameters. It is transmitted at 50 bits per second, and modulates the L1-C/A code signal.

Theoretically, signals from three satellites can locate a specific position in three-dimensional space. In practice though, differences in the clock accuracy between the satellites and the hand-held receivers require the use of a fourth signal. This fourth signal is used as a crosscheck, providing a correction factor used to reestablish the accuracy capable of three perfectly timed signals.

4.2 Personal Communications

Communication device companies are very interested in increasing the capacity of their voice-oriented digital cellular and personal communication services (PCS). For this reason, companies are turning to CDMA technologies as opposed to the more traditional TDMA and FDMA methods. CDMA based on spread spectrum would increase the wireless digital network's capacity, allow for smoother hand-off of cellular connections, and provide more reliable service [8]. As discussed earlier, CDMA can be implemented easy using DSSS by choosing the sequences which minimize the mutual interference caused by all users of a multiple access system.

The North American CDMA Digital Cellular (IS-95) standard is one example of the use of spread spectrum technology for personal communications. Interest in this standard sprung from the desire to improve the capacity of cellular networks over then traditional analog systems (AMPS). The IS-95 air-interface specification was release in July of 1993, and employs a spread spectrum signal with 1.2288 MHz spreading bandwidth [8]. Due to the increased capacity of these networks, the interference caused by multiple users had to be resolved. The near-far problem, discussed earlier, requires complex transmission power control algorithms built into the control stations. The pseudo-noise sequences used in the

IS-95 standard are called Walsh codes, and provide 64 orthogonal sequences

4.3 Local Area Networks

The availability of the unlicensed ISM bands provide spread spectrum technology the opportunity to be employed for wireless local area networks (WLANs). IEEE standard 802.11 defines three versions of the physical layer used in local area networks: infrared, and two RF transmissions in the 2.4 GHz ISM band. Only the two RF transmission methods have a significant presence in the market for WLANs. This standard requires the implementation of either DSSS or FHSS modulation. Let us briefly examine the characteristics of each of these modulation methods as prescribed by IEEE 802.11 [8].

DSSS

The DSSS modulation method employs an 11-bit Barker sequence in spreading the data before transmission. Both Barker sequences and the Walsh sequences mentioned earlier are beyond the scope of this report, but information regarding these sequences can be found in [9]. The 11-bit sequence provides a 10.4 dB processing gain. The sequence is generated at a rate of 11 million chips per second, and is modulated onto the carrier frequency within the 2.4 GHz ISM band. Data rates in the 1-2 Mbps range are achieved using this method.

FHSS

In FHSS WLANs, the information is first modulated using Gaussian Frequency Shift Keying. The carrier frequency in the 2.4 GHz ISM band then hops from channel to channel, following a pattern selected from among 78 different prearranged pseudorandom hop patterns. According to FCC regulation, the minimum hop rate is 2.5 hops per second, corresponding to a dwell time of no more than 400 ms. A 1-2 Mbps data rate is routinely achieved using this method.

4.4 Summary

The purpose of this section was to provide for the reader several examples of commercial applications of spread spectrum. For the greater part of a half-century, development of spread spectrum technology and devices was restricted to the military. Advances in VLSI and signal processing techniques over the past several decades have—combined with the exhausting demand for wireless communication services—brought spread spectrum to the forefront of commercial development. Since the release of the ISM bands by the FCC in May of 1995, commercial development of spread spectrum services has entered a frenzied pace, as companies try to edge each other out in the telecommunications and wireless networking industry.

One of the oldest and most well known implementations of spread spectrum lies in the Global Positioning System. Developed by the DOD for positioning and navigation, it employs the direct sequence modulation technique to spread the information signals used by terrestrial receivers. The greatest commercial interests, though, lie in the fields of Personal Communications Services and Wireless Local Area Networks. Standardization within these areas provides the industry with a focus and direction for the development of these spread spectrum devices.

5. Conclusions

Spread spectrum, developed by the military in secret during the Second World War, has become one of the most useful and widely implemented modulation techniques for wireless multiple access networks. Characterized by a bandwidth that is significantly larger than that required of its information, spread spectrum modulation has several properties that make them very valuable for a variety of communication and ranging purposes. Low power density, code division multiplexing, and interference rejection are but three of these properties.

The goal of this paper was to provide a solid foundation of the concepts and ideas behind spread spectrum technology. The most well known method of spread spectrum modulation—direct sequence—was investigated through the use of many visual examples. Direct sequence employs pseudorandom noise-like sequences to spread the bandwidth of an information signal. The characteristics of spread spectrum modulation are a direct result of these pseudo-noise sequences.

Coding methods play an important part in all types of digital communication systems. The codes used for spread spectrum applications are different in that they are often extremely long and noise-like. Using these sequences involves consciously trading bandwidth for processing gain. Two commonly used types of sequences were discussed in this presentation. Maximal sequences have outstanding autocorrelation properties and are without peer for use in general communication and ranging applications. For multiple access environments, however, the cross-correlation between two signals plays a more significant role. For this reason, Gold codes are used due to their well-defined cross-correlation properties.

A mature technology, spread spectrum has become the modulation method of choice for multi-user communication systems throughout the world. Both

Wireless Local Area Networks and Personal Communications Systems employ spread spectrum for their increased capacity and robustness against interference. Another well-known application of spread spectrum is in the worldwide position and navigation system known as the Global Positioning System. As civilian applications and standards of spread spectrum continue to emerge, we can expect that this modulation method will be used to an even greater extent than it is today.

Glossary

Autocorrelation. The degree of agreement between a pair of identical signals. Calculated through the multiplication of a signal with a time-delayed replica of itself.

CDMA. Code division multiple access. Any signal multiplexing technique employing codes to separate signals.

Chip. The output of a code generator during one clock interval.

Chip Rate. The rate at which chips are output from a code generator.

Correlation. The degree of agreement between a pair of signals.

Cross-correlation. The degree of agreement between a pair of dissimilar signals. Calculated by multiplying one signal with a time-delayed version of the second signal.

Direct Sequence (DS). A form of spread spectrum modulation where a code sequence is used to directly modulate a carrier, usually by phase-shift keying.

FDMA. Frequency division multiple access. A signal multiplexing technique using frequency to separate signals.

Frequency Hopped (FH). A type of spread spectrum modulation in which the wideband signal is generated by hopping from one frequency to another over a large number of frequency choices. The frequencies are pseudorandomly chosen by a code similar to those used in direct sequence systems.

Gold Code. A particular type of composite code generated from maximal sequences. Oriented to multiple access due to its well-defined cross-correlation properties.

Interference. Any signal that tends to hamper the normal reception of a desired signal.

Jamming Margin (M_j). The amount of interference a system is able to withstand

while producing the required output signal-to-noise ratio or bit-error rate.

Maximal Sequence. A type of code sequence that is the longest code that can be generated by a feedback code generator. Characterized by its unparalleled autocorrelation properties.

Process Gain (G_p). The gain or signal-to-noise improvement achieved by a spread spectrum system due to coherent band spreading and remapping of the desired signal.

Pseudonoise. The term used to signify any of a group of code sequences that exhibit noise-like properties.

Sequence. A code or train of chips.

Shift Register. A sequentially connected group of delay elements used to generate code sequences in spread spectrum systems.

Shift Register Generator. The combination of shift registers and modulo-2 adders used for code sequence generation.

Spread Spectrum. Any of a group of modulation formats in which an RF bandwidth much wider than necessary is used to transmit an information signal so that a signal-to-noise improvement may be gained in the process.

Synchronization. Timing agreement or the act of gaining timing agreement between a spread spectrum transmitter and its receiver.

TDMA. Time division multiple access. A signal multiplexing technique using time division to separate signals.

References

- [1] Cooper, George R., McGillem, Clare D. Modern Communications and Spread Spectrum. McGraw-Hill, New York 1986.
- [2] Dillard, Robin A., Dillard, George M. Detectability of Spread-Spectrum Signals. Artech House, Inc., Norwood 1989.
- [3] Dixon, Robert C. Spread Spectrum Systems with Commercial Applications, 3rd edition. John Wiley & Sons, Inc., New York 1995.
- [4] Proakis, John G. Digital Communications, 3rd edition. McGraw-Hill, Boston 1995.
- [5] Rappaport, Theodore, S. Wireless Communications: Principles & Practice. Prentice Hall PTR, Upper Saddle River 1996.
- [6] Stremmler, Ferrel G. Introduction to Communication Systems, 3rd edition. Addison-Wesley, Reading 1990.
- [7] "Spread Spectrum Primer." Spread Spectrum Scene Online Magazine. Available online at www.sss-mag.com/primer.html. August 1998.
- [8] Meel, J. "Spread Spectrum: Applications." Sirius Communications White Paper. Rotselaar, Belgium 1999.
- [9] Meel, J. "Spread Spectrum: Introduction." Sirius Communications White Paper. Rotselaar, Belgium 1999.
- [10] Pickholtz, R.L., et al. "Theory of Spread-Spectrum Communications." IEEE Transactions on Communications, Vol. COM-30, May 1982, pp. 855-884.
- [11] Schilling, Donald L., et al. "Spread Spectrum goes Commercial." IEEE Spectrum, August 1990.
- [12] Scholz, Robert A. "The Origins of Spread-Spectrum Communications." IEEE Transactions on Communications, Vol. COM-30, May 1982, pp. 822-854.

Vita

Adam Mirek Mankowski was born in Naperville, Illinois on February 6, 1977, the son of Aleksander Mankowski and Barbara Mankowski. After completing his work at A.E. Stevenson High School, Lincolnshire, Illinois in 1995, he was awarded an appointment to the United States Air Force Academy in Colorado Springs, Colorado. He received the degree there of Bachelor of Science in both Electrical Engineering and Physics in June of 1999. Upon graduation, he was commissioned as a second lieutenant in the United States Air Force and entered active duty service. In August of 1999, he entered graduate school at The University of Texas at Austin to pursue the Master of Science degree in Electrical Engineering.

Permanent address: 29529 N. Waukegan Road, Apt. 108
Lake Bluff, Illinois 60044-5442

This report was typed by the author.