

Direct-Sequence Spread-Spectrum Communications for Multipath Channels

Michael B. Pursley, *Fellow, IEEE*

Invited Paper

Abstract—Direct-sequence spread spectrum has been adopted for many current and future cellular CDMA communication systems, and it is also used widely for military communication networks and systems. One of the motivations for employing direct-sequence spread spectrum is its ability to combat fading due to multipath propagation. The use of direct-sequence spread spectrum to resolve multipath signals is discussed and illustrated. The role of a rake receiver is described, and tradeoffs in the selection of the chip rate for the spread-spectrum system are discussed.

Index Terms—Fading, multipath channels, spread-spectrum communications.

I. INTRODUCTION

SPREAD-SPECTRUM techniques can provide effective means for communicating reliably over channels that exhibit multipath propagation. In order to develop a spread-spectrum system that accomplishes this objective, it is necessary for the system designer to know the key features of the propagation medium and choose the parameters of the spread-spectrum signal and demodulator with those features in mind. Some of the most important features of multipath channels are illustrated in this paper, and tradeoffs that arise in the design of spread-spectrum systems for such channels are described.

The emphasis is on intuitive explanations that are based on simple deterministic models of multipath propagation. The goal is to provide an aid to intuition rather than a precise evaluation of performance. Analyses of direct-sequence spread spectrum that account for the stochastic nature of multipath channels are available in previous publications (e.g., [14] and [16]).

A brief summary of some of the benefits of direct-sequence spread spectrum is provided in Section II, and mathematical models for direct-sequence spread spectrum signals are given in Section III. The multipath resolution capability of direct-sequence spread-spectrum systems is illustrated in Section IV, the effects of nonselective fading in a diffuse multipath channel are described in Section V, and a brief discussion of the rake re-

ceiver is provided in Section VI. Some of the primary issues involved in selecting the chip rate for the spread-spectrum system are discussed throughout this paper. It is not possible to tell the complete story in a paper of modest length, so several suggestions for additional reading are included among the references.

II. SOME BENEFITS OF DIRECT-SEQUENCE SPREAD SPECTRUM

Although there are many communication techniques that are classified as spread spectrum by all experts, there are others for which opinions may be divided. The classical definition of a spread-spectrum signal is a signal that occupies a wider band of frequencies than is required by the signal's data rate. This definition may be useful for some purposes, but it does have some flaws. For example, a standard binary phase-shift key (PSK) signal with one PSK pulse per data bit is not considered to be spread spectrum, but there are many modulation techniques that provide the same data rate with less bandwidth. In addition, the classical definition does not tell us how much the bandwidth must be expanded beyond the minimum in order for the resulting signal to be a spread-spectrum signal. Fortunately, a precise definition is not required for our discussion since we restrict attention to signals that are classified as spread spectrum by any reasonable definition and are among the most commonly used in spread-spectrum systems. These signals are obtained by direct modulation of the data by a relatively wide-band signal that is derived from a digital sequence.

The basic RF direct-sequence spread-spectrum signal in our illustrations employs binary amplitude-shift-key (ASK) modulation for data modulation and spectral spreading. The signal is defined by

$$s(t) = Aa(t)d(t)\cos(\omega_c t + \theta) \quad (1)$$

where A is the signal amplitude, $a(t)$ is a sequence of rectangular pulses of duration T_c , and $d(t)$ is a sequence of rectangular pulses of duration T that represent the digital data. In most applications, $T = NT_c$ for some integer N and the pulse transitions in $a(t)$ are aligned with those in $d(t)$. The signal $a(t)$ is referred to as the *spreading signal* or *signature signal*. The elemental pulses that make up the spreading signal are often referred to as *chips*, and the shape of the elemental pulse is referred to as the *chip waveform*. The chip rate is $R_c = 1/T_c$ and the data rate is $R_d = 1/T$. Throughout this paper, we assume

Manuscript received September 21, 2001. This work was supported by the U.S. Army Research Office under Grant DAAG55-98-1-0013 and by the Office of Naval Research under Grant N00014-99-1-0567 and MURI Grant N00014-00-1-0565.

The author is with the Department of Electrical and Computer Engineering, Clemson University, Clemson, SC 29634 USA (e-mail: pursley@eng.clemson.edu).

Publisher Item Identifier S 0018-9480(02)01986-5.

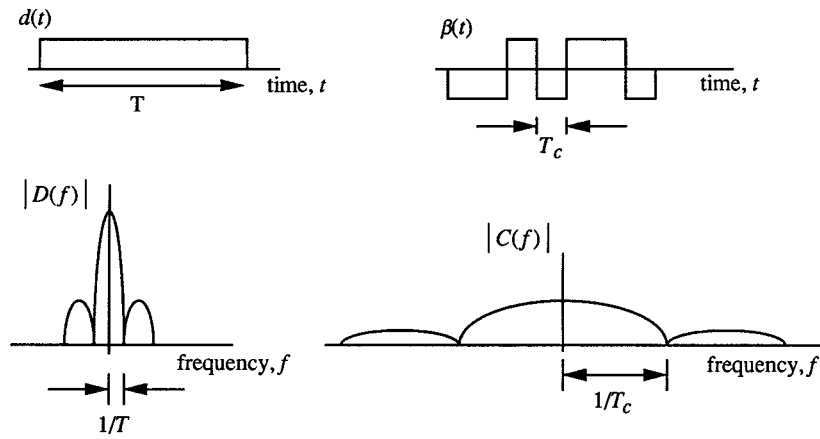


Fig. 1. Direct-sequence spread spectrum.

that the bandwidth of $a(t)$ is much smaller than the carrier frequency $f_c = \omega_c/2\pi$, which requires $R_c \ll f_c$.

The rectangular pulse of duration λ is defined by $p_\lambda(t) = 1$ for $0 \leq t < \lambda$ and $p_\lambda(t) = 0$ for all other values of t . The amplitudes of the pulses for $a(t)$ are obtained from a sequence $(x_n) = \dots, x_{-1}, x_0, x_1, x_2, \dots$, known as the *spreading sequence* or *signature sequence*. For a rectangular chip waveform, the spreading signal that is applied to a single data symbol can be expressed as

$$a(t) = \sum_{n=0}^{N-1} x_n p_{T_c}(t - nT_c), \quad 0 \leq t < T. \quad (2)$$

The data signal corresponding to (2) is $d(t) = d_0 p_T(t)$, where d_0 denotes the data symbol. For a message or data packet with several data symbols, the limits on the sum and the time interval for which $a(t)$ is defined are determined by the length of the message. The sequence (x_n) may be periodic, in which case, the sequence may repeat within the duration of the message. Alternatively, it may be that the sequence does not repeat or the period exceeds the message length.

If we focus on the baseband signal $\beta(t) = Aa(t)d(t)$, we see that the direct-sequence spread-spectrum modulation process converts each data pulse into a spread-spectrum signal, as shown in Fig. 1. The resulting spread-spectrum signal conveys the same information as the original data signal, but the former occupies a wider band of frequencies than the latter if $N > 1$. The bandwidth expansion is proportional to N , as illustrated in Fig. 1, for $N = 7$. A graph of the main lobe and the two adjacent side lobes of the magnitude of $D(f)$, the Fourier transform of the data pulse, is shown in Fig. 1. Also shown is the corresponding graph of the magnitude of $C(f)$, the Fourier transform of the chip waveform. Note that the data rate $R_d = 1/T$ and the chip rate $R_c = 1/T_c$ correspond to the first nulls in $D(f)$ and $C(f)$, respectively. The increase in bandwidth demonstrated in Fig. 1 is the motivation for the names *spreading signal* and *spreading sequence* for $a(t)$ and (x_n) , respectively. If the bandwidth of the spread-spectrum signal is required to be much larger than the bandwidth of the data signal, it is necessary that the $R_c \gg R_d$, which requires $N \gg 1$.

In many applications, each element of the sequence (x_n) is either $+1$ or -1 . Such a binary sequence is often generated in

a logic device as a sequence of 0's and 1's and then converted according to $0 \rightarrow +1$ and $1 \rightarrow -1$. Some sequences that are employed as spreading sequences exhibit certain randomness properties, such as those described in [6] for maximal-length linear-feedback shift-register sequences (m -sequences). As a result of these randomness properties, m -sequences are sometimes referred to as pseudorandom sequences. Some authors apply the term *pseudorandom* more broadly to encompass other classes of sequences that are of interest for spread spectrum. Such sequences do not have all of the randomness properties of m -sequences, but many of them have other desirable properties that make them attractive for spread-spectrum systems [27].

Direct-sequence spread spectrum is one of the oldest forms of spread spectrum [28]. It has been employed in such well-known systems as the Global Positioning System (GPS) [30], the NASA Tracking and Data Relay Satellite System (TDRSS) [39], the Joint Tactical Information Distribution System (JTIDS) [39], the IS-95 mobile cellular communication system [33], and third-generation cellular systems [7]. One reason for the popularity of direct-sequence spread spectrum is that it can provide a means of reliable communication in the presence of various types of interference, including interference from transmitters that are part of the system, multipath interference, and RF interference from emitters that are not part of the system. Interference from transmitters that are part of the spread-spectrum system is referred to as multiple-access interference, and the ability of a system to accommodate simultaneous transmissions in the same frequency band is referred to as multiple-access capability.

To provide multiple-access capability in a direct-sequence spread-spectrum system, the sequence (x_n) is used to identify an individual signal, which is the reason that (x_n) is often given the name *signature sequence*. Multiple-access capability can be provided by direct-sequence spread spectrum [19] or frequency-hop spread spectrum [21]. By employing either of these forms of spread spectrum, multiple transmitter-receiver pairs can communicate reliably in the same frequency band at the same time [22]. With either type of spread-spectrum modulation, a well-designed spread-spectrum multiple-access (SSMA) system can maintain acceptably low levels of interference. In cellular communications and in some earlier applications, SSMA is referred to as CDMA.

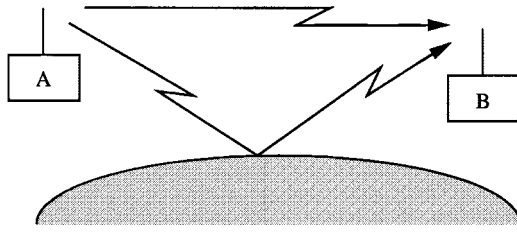


Fig. 2. Illustration of multipath.

Another form of interference that is encountered in mobile communications is multipath interference, for which a simple illustration is given in Fig. 2. Due to reflections off objects in the propagation path, multiple copies of the transmitted signal may be received. These different versions of the signal are offset in time and phase so that they may add constructively (in-phase) or destructively (out-of-phase). Direct-sequence spread-spectrum modulation combats multipath by permitting discrimination against unwanted multipath components that might otherwise cause destructive interference. When used in conjunction with a rake receiver, direct-sequence spread-spectrum modulation permits combining of at least some of the multipath components in such a way that they add constructively.

Among the advantages cited for direct-sequence CDMA cellular communications is the ability to communicate reliably in a multipath environment. In order to accomplish its task, the receiver must isolate the individual multipath components, and this places a lower bound on the chip rate of the direct-sequence spread-spectrum signal, as discussed in Section IV. Individual multipath signals that have been resolved may be combined in a rake receiver, as described briefly in Section VI. A related benefit is the ability to perform soft handoff for mobile terminals that are within range of two or more base stations. In effect, the signals transmitted from two or three base-stations can be treated as multipath components and combined in the mobile terminal's rake receiver.

A direct-sequence spread-spectrum system can provide some protection against jamming (e.g., see [29] and [35]) and it can operate with a low power-spectral density to facilitate coexistence with other systems [12]. Well-designed spread-spectrum signals are also difficult for unauthorized receivers to detect [13]. These and other benefits are described in several of the references, including [3], [18], and [23].

III. DIRECT-SEQUENCE SPREAD-SPECTRUM SIGNALS

Although we restrict attention to the signal defined by (1) and (2) for the illustrations in this paper, there are more general forms of direct-sequence spread-spectrum signals that are commonly used. One generalization is the signal defined by

$$s(t) = A\alpha(t) \cos(\omega_c t + \theta) \quad (3)$$

in which $\alpha(t)$ represents the data and also provides the spreading of the spectrum. If the sequence of data symbols is

denoted by (d_m) and the chip waveform is $\psi(t)$, then $\alpha(t)$ is defined by

$$\alpha(t) = \sum_n d_{\lfloor n/N \rfloor} x_n \psi(t - nT_c) \quad (4)$$

where $\lfloor u \rfloor$ denotes the integer part of the real number u and the range of the sum depends on the number of data symbols. As in (2), there are N chips per data symbol, so the sum in (4) is from 0 to $LN - 1$ if the data symbols $d_0, d_1, \dots, d_{L-1}$ are transmitted. The chip waveform $\psi(t)$ need not be rectangular, and its duration may exceed T_c .

The signal defined by (3) and (4) is an example of ASK modulation. If the values for d_m and x_n are limited to $+1$ and -1 and the chip and data pulse waveforms are rectangular, the resulting ASK signal is mathematically equivalent to the classical binary phase-shift key (BPSK) signal. In this special case, $\alpha(t) = a(t)d(t)$ so (3) and (4) give the same signal as (1) and (2). Two ASK signals can be combined to give a signal of the form

$$s(t) = A \{ \alpha_1(t) \cos(\omega_c t + \theta) - \alpha_2(t) \sin(\omega_c t + \theta) \} \quad (5)$$

which is a quadrature amplitude-shift key (QASK) signal. In general, $\alpha_1(t)$ and $\alpha_2(t)$ are defined by (4), but the two signals can have different sets of data symbols and different spreading sequences. With appropriate restrictions on the data symbols and the spreading signals, the QASK signal is mathematically equivalent to the classical quadriphase-shift key (QPSK) signal. The chips in the in-phase component of (5) may be offset in time from the chips in the quadrature component by any odd multiple of $T_c/2$ to produce offset QASK or offset QPSK. Other chip waveforms are available, such as the sine pulse [20] that gives minimum shift keying (MSK) if it is employed in (5) with such a time offset.

It is often convenient to express such signals in complex form. The complex signal corresponding to (5) is $\tilde{s}(t) = \beta(t)e^{j\omega_c t}$, where $\beta(t) = A\alpha(t)e^{j\theta}$ is the baseband equivalent signal and $\alpha(t) = \alpha_1(t) + j\alpha_2(t)$. The actual transmitted signal is the real part of the complex signal. It is easy to show that the signal in (5) satisfies $s(t) = \text{Re}\{\tilde{s}(t)\}$.

Many of the features of spread spectrum are derived from the properties of $\alpha(t)$, so it is convenient to work with the baseband equivalent $\beta(t) = A\alpha(t)e^{j\theta}$. This signal representation plays a role that is similar to that of the phasor representation of a sinusoidal signal. The baseband equivalent for (1) is

$$\beta(t) = Aa(t)d(t)e^{j\theta} \quad (6)$$

which is a real signal, except for the complex exponential that essentially keeps track of the phase of the RF signal. Notice that if a time delay Δ is introduced in (1), (3), or (5), the phase is changed by an amount $\phi_0 = -\omega_c \Delta$. Thus, for example, the baseband equivalent of the delayed signal is

$$\beta(t - \Delta) = Aa(t - \Delta)d(t - \Delta)e^{j(\theta + \phi_0)}. \quad (7)$$

IV. ILLUSTRATION OF MULTIPATH RESOLUTION

The illustration provided in this section is for a channel in which the multipath arises as a result of specular reflections off a number of objects and each individual multipath component is not dispersed in time. Such a multipath component is referred to as a *specular* component, and a *specular multipath channel* is a channel that produces only specular components. The output of a specular multipath channel consists of the sum of a number of attenuated time-delayed versions of the transmitted signal, each of which arrives at the receiver without distortion. There may or may not be a line-of-sight path.

Due to the possibility that the multipath signals produce destructive interference, the signal can undergo severe fading. As an illustration of the consequences of multipath, suppose that spread spectrum is not used, the transmitted signal is $A d(t)$, and the channel produces two specular components. The received signal is $r(t) = A_1 d(t) + A_2 d(t - \Delta) e^{j\phi}$, where Δ is the differential delay for the two paths and $\phi = -\omega_c \Delta$. Since the carrier frequency is much larger than $1/T$, Δ can satisfy $\omega_c \Delta \approx \pi$ and $\Delta \ll T$ simultaneously. If $d(t) = +1$ in the interval $0 \leq t < T$, then $r(t) \approx A_1 - A_2$ for most of the interval. This is an example of destructive interference because the second component of the received signal subtracts from the first. If $A_1 \approx A_2$, then $r(t) \approx 0$ for most of the interval so there is nearly a complete cancellation due to multipath. A typical receiver for the data pulse shown in Fig. 1 integrates the received signal over the interval from 0 to T . (This is the matched filter for the rectangular pulse of duration T .) It is clear that destructive interference can cause the output of such a receiver to be approximately zero. As illustrated by this example, multipath interference can produce fading of the received composite signal. This fading leads to an increase in the error rate for data transmission or a reduction in the voice quality for digital speech transmission.

A deterministic time-invariant linear-system model for a specular multipath channel provides a good illustration of one of the important benefits of direct-sequence spread spectrum. Consider the spread-spectrum signal given by (1) with $A = 1$, $\theta = 0$, $d(t) = p_T(t)$, and $T = NT_c$ for some integer N . The baseband equivalent is $\beta(t) = a(t)$ for $0 \leq t < T$; that is, the baseband equivalent of the transmitted signal is a spread-spectrum pulse consisting of N chips, which is illustrated in Fig. 1 for $N = 7$.

Suppose the baseband equivalent model for the RF channel is a time-invariant linear filter with impulse response $h(t)$. In practice, the channel may not be time invariant, but in some applications, the changes in the channel occur slowly compared to the message duration. In such an application, the channel can be modeled as a time-invariant linear filter, at least for the duration of several consecutive data bits and perhaps for the duration of an entire message. Even if the time-invariant model does not fit a particular channel precisely, it is instructive to examine this model as an aid to the understanding of the multipath resolution capabilities of direct-sequence spread spectrum.

Consider a receiver that uses a time-invariant linear filter with an impulse response $h_M(t)$ that is matched to the spread-spectrum signal $a(t)$. It follows that $h_M(t) = a(T_s - t)$ for $-\infty < t < \infty$, where T_s is the time that is selected to sample the output of the matched filter. The choice of T_s is

arbitrary, and it simplifies the presentation if we let $T_s = 0$. In this illustration, there is no attempt to match the filter to the output of the channel; instead, the filter is matched to the transmitted signal. The notion of matching the filter to the channel arises in the discussion of the rake receiver in Section VI.

As is well known, the output $w(t)$ of the channel is determined by convolving the functions a and h , i.e., $w = a * h$, which is shorthand for the operation $w(t) = \int a(\tau)h(t - \tau) d\tau$. (If limits are not indicated, the integral is from $-\infty$ to ∞ .) The output of the channel is the input to the matched filter, which is a time-invariant linear filter. The filter output $y(t)$ is given by $y = a * h_M * h$. This expression can be written as $y = R * h$, where $R = a * h_M$. Since $T_s = 0$, then $h_M(t) = a(-t)$, and the function R can be expressed as $R(t) = \int a(\tau)a(\tau - t) d\tau$. Thus, the function R is the autocorrelation function for the spreading signal $a(t)$. Note that $R(0)$ is the integral of $a^2(t)$, which is the energy E_a in the spreading signal.

The autocorrelation function R for the spreading signal is determined by the chip waveform and the aperiodic autocorrelation function C_x for the sequence (x_n) . The function C_x is defined by $C_x(j) = x_0x_j + x_1x_{j+1} + \dots + x_{N-1-j}x_{N-1}$ for $0 \leq j \leq N-1$, $C_x(j) = C_x(-j)$ for $1-N \leq j \leq -1$, and $C_x(j) = 0$ for j outside these two ranges. The *peak* of the aperiodic autocorrelation function is $C_x(0)$, which is equal to the sum of the squares of the elements x_n for $0 \leq n \leq N-1$. The *sidelobes* are the values of $C_x(j)$ for $j \neq 0$. For illustrative purposes, suppose the sidelobes are negligible.

Next, suppose that the chip waveform is a rectangular pulse of duration T_c , as illustrated in Fig. 1. It follows that $R(jT_c) = C_x(j)T_c$ for each integer j . Furthermore, the graph of $R(t)$ for $jT_c \leq t \leq (j+1)T_c$ is the straight line connecting the two points $R(jT_c)$ and $R((j+1)T_c)$. Since the sidelobes are negligible, it follows that R is a triangular function that is centered at the origin and has base $2T_c$. That is, $R(t)$ is zero outside the interval from $-T_c$ to T_c , and $R(t) = R(0)[(T_c - |t|)/T_c]$ within that interval. This is the ideal autocorrelation function for a signal that consists of a sequence of rectangular pulses of duration T_c . Since $R(0) = E_a$ and $a^2(t) = p_T(t)$, it follows that $R(0) = T$. Thus, the height of the triangular pulse does not depend on the value of T_c .

The impulse response of a baseband channel that has only specular multipath is just the sum of Dirac delta functions, each of which may be multiplied by a complex exponential to account for the phase shift associated with the corresponding path. Thus, the impulse response of a channel with K multipath components is

$$h(t) = c_0 \delta(t - t_0) + c_1 \delta(t - t_0 - \tau_1) e^{j\phi_1} + \dots + c_{K-1} \delta(t - t_0 - \tau_{K-1}) e^{j\phi_{K-1}}. \quad (8)$$

If the only cause of a phase shift in the k th path is the propagation delay, then $\phi_k = -\omega_c \tau_k$.

Consider a three-path specular multipath channel for which the differential propagation delays are Δ and 2Δ , and assume $\omega_c \Delta$ is a multiple of 2π so that $\phi_1 = \phi_2 = 0$. The output of the matched filter is obtained from $y = R * h$, thus, (8) implies that

$$y(t) = c_0 R(t - t_0) + c_1 R(t - t_0 - \Delta) + c_2 R(t - t_0 - 2\Delta) \quad (9)$$

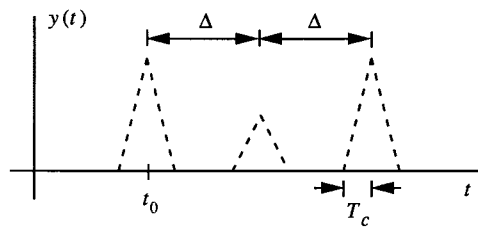


Fig. 3. Matched filter output for high chip rate.

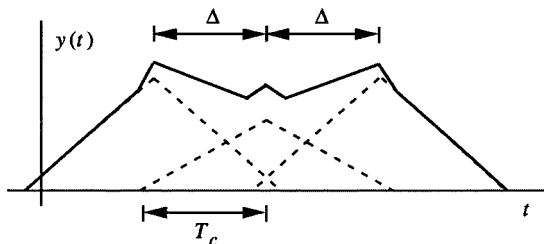


Fig. 4. Matched filter output for low chip rate.

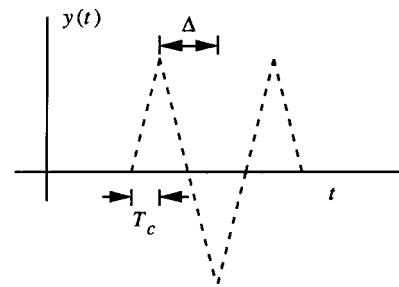


Fig. 5. Matched filter output for high chip rate.

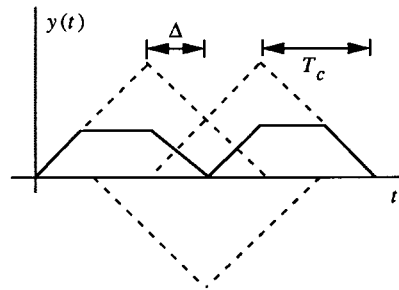


Fig. 6. Matched filter output for low chip rate.

as illustrated in Fig. 3 for $c_0 = c_2 = 1.0$ and $c_1 = 0.5$. If $T_c \leq \Delta/2$, as in Fig. 3, the three components of $y(t)$ do not overlap, thus, it is easy to identify three distinct multipath components in the matched filter output. By sampling the output at times t_0 , $t_0 + \Delta$, and $t_0 + 2\Delta$, we obtain three useful signal components that can be added to provide a larger signal than if we relied on only one component (e.g., the first component to arrive or the largest component). This is a simple illustration of multipath combining in a spread-spectrum receiver. In practice, the multipath components can be adjusted in amplitude and phase before they are combined. For example, if the receiver estimates the relative amplitudes of the three multipath components, it can apply gains to the components that are proportional to their estimated amplitudes. Regardless of the combining method, it is necessary to resolve the multipath components before they can be combined.

The consequences of having $T_c > \Delta/2$ are illustrated in Fig. 4, where the dashed lines represent the individual components of the matched filter output and the solid line represents the composite output. An observer at the output of the filter sees only the solid line, from which it is not possible to identify the three multipath components that are actually present. This is a result of the chip rate being too low to provide multipath resolution. (Note that $R_c < 2/\Delta$ in Fig. 4.) For a specular multipath channel it is desirable for the chip rate to be as high as possible, subject to the constraints on complexity and bandwidth. The bandwidth required by the direct-sequence system is proportional to the chip rate, thus, there is a tradeoff between the multipath resolution capability and the required clock rates and bandwidths for the devices in the spread-spectrum transmitter and receiver.

Suppose now that the multipath channel with three paths has $c_0 = c_1 = c_2 = 1$ and the value of Δ is such that $\omega_c \Delta$ is an odd multiple of π , thus, $\phi_1 = \pi$ and $\phi_2 = 0$. The output of the matched filter is now $y(t) = R(t-t_0) - R(t-t_0-\Delta) + R(t-t_0-2\Delta)$, as illustrated in Fig. 5 for $T_c = \Delta/2$ and in Fig. 6 for $T_c = 2\Delta$. The situation depicted in Fig. 6 illustrates that destructive interference can occur even in a spread-spectrum system if the chip rate is too

low. If the chip rate is high enough to resolve the multipath components (i.e., if $T_c \leq \Delta/2$), the phases of the components can be estimated and the estimates can be employed to combine the multipath constructively. In the simple example of Fig. 6, the second component should be inverted before being added to the sum of the first and third components.

V. EFFECTS OF DIFFUSE MULTIPATH FADING

Some channels have one or more clusters of propagation paths for which each cluster has large number of paths with small differential delays. By this, we mean that, for each cluster, the differences among the path delays are small compared with the inverse of the chip rate of the spread-spectrum signal. The multipath that arises in a such a channel is referred to as *diffuse multipath*, and a multipath component that results from a cluster of this type is referred to as a *diffuse component* of the received signal. A diffuse component of the received signal is composed of a number of time-offset replicas of the transmitted signal. Thus, for a diffuse multipath channel, the width of a received pulse may exceed that of the transmitted pulse. If the amount of the increase in width is large enough to cause significant distortion to the pulse, the multipath channel is said to be *time dispersive*. When viewed in the frequency domain, the phenomenon that produces time dispersion is referred to as *frequency selectivity* [31].

In this section, we restrict attention to channels that have one path that produces a specular component and a cluster of paths that produce a diffuse component. It is also assumed that the differential delays among all paths are small enough to cause only negligible distortion of the chip waveform. If the differences among the path delays are large compared with the inverse of the carrier frequency, destructive interference may occur. We assume the composite signal is a linear combination of a number of signals whose time offsets are small enough to produce little

or no distortion in the chip waveform, but large enough to lead to destructive interference. Thus, the diffuse multipath signal may exhibit fading, which, in this situation, is referred to as nondispersive or nonselective fading. We focus on nonselective fading because of the simplicity of the error probability expressions for such fading. In addition, the issues that arise in examining the effect of the chip rate on error probability also arise in channels with selective fading (e.g., see [14]–[16]).

Nonselective fading leads to variations in the received energy, but it does not produce significant distortion in the received signal. If E_t is the energy in the transmitted signal, the energy in the received signal is $V^2 E_t$, where V is the random voltage gain applied by the fading channel. A model for nonselective fading that is often used for a multipath signal that is the sum of a specular component and a diffuse component is the nonselective Rician model. In this model, the random variable V has the Rician density function [26] which is given by

$$f_V(u) = \frac{2u}{\xi^2} \exp\left\{-(u^2 + c^2)/\xi^2\right\} I_0(2cu/\xi^2) \quad (10)$$

for $0 \leq u < \infty$, and $f_V(u) = 0$ for $u < 0$. In (10), the function I_0 is the zeroth-order modified Bessel function of the first kind and $E\{V^2\} = c^2 + \xi^2$. The nonnegative parameters c and ξ are related to the strengths of the specular and diffuse components, respectively. The energy in the specular component of the received signal is $E_{sp} = c^2 E_t$ and the average energy in the diffuse component is $E_d = \xi^2 E_t$. The total energy in the received signal is $E = E_{sp} + E_d$. A common model for purely diffuse multipath is the Rayleigh model, for which V is a Rayleigh random variable. The Rayleigh density function is given by (10) with $c = 0$, which corresponds to Rician fading with no specular component.

Many different digital modulation methods have been employed in wireless communications, and the choice as to which method should be used for a given application depends to a large extent on the characteristics of the channel. In order to simplify the presentation, we focus on a binary modulation scheme that is suitable for noncoherent demodulation. In particular, the binary signal set $\{v_0, v_1\}$ has signals of the form

$$v_i(t) = \sum_{k=0}^{M-1} y_{i,k} p_{T_0}(t - kT_0) \quad (11)$$

with $y_i = (y_{i,0}, y_{i,1}, \dots, y_{i,M-1})$ having the property that y_0 and y_1 are orthogonal. The data symbol duration is $T = MT_0$. The data signal of (1) and (6) is given by $d(t) = v_0(t - mT)$ for $mT \leq t < (m+1)T$ if 0 is sent in the m th time interval; otherwise, the data signal is given by $d(t) = v_1(t - mT)$ to represent the symbol 1 in this time interval. The baseband equivalent of the transmitted signal is $\beta_i(t) = a(t)v_i(t)$. If T_0 is a multiple of T_c , there are an integer number of chips of the spreading signal for each of the pulses $p_{T_0}(t - kT_0)$, but this integer could be as small as one. We assume for purposes of illustration that the autocorrelation function for each of the two signals $\beta_0(t)$ and $\beta_1(t)$ is the ideal autocorrelation function R defined in Section IV.

Among the possibilities for the vectors y_0 and y_1 are any two rows of a Hadamard matrix [5] of order 2^n for any integer n . For example, for $n = 3$, we can use $(+1, -1, +1, -1, -1, +1, +1, -1)$

and $(+1, -1, -1, +1, +1, -1, +1, -1)$. In general, the rows of a Hadamard matrix of order 2^n form a set of 2^n orthogonal signals, but only two of them are needed for our illustration. Some signals of this type are often referred to as Walsh functions or Hadamard–Walsh signals (e.g., [25], [34], and [38]).

For our illustration, we consider binary orthogonal modulation, a Rician nonselective fading channel, and a standard noncoherent receiver that employs matched filters followed by envelope detectors. The matched filters have low-pass equivalent impulse responses $\beta_0(-t)$ and $\beta_1(-t)$. Let $N_0/2$ denote the two-sided power spectral density of the thermal noise. The average probability of error for the noncoherent receiver is given by [36]

$$P_e = \frac{\exp\left\{-(E_{sp}/N_0)/\left[2 + (E_d/N_0)\right]\right\}}{2 + (E_d/N_0)}. \quad (12)$$

At the output of the matched filter, there are three signal-to-noise ratios (SNRs) of interest. The SNR for the specular component is E_{sp}/N_0 , and the SNR for the diffuse component is E_d/N_0 . The total SNR is $E/N_0 = (E_{sp}/N_0) + (E_d/N_0)$.

In the illustrations, the energy E_d must be interpreted as the effective energy in the diffuse component, and the effective energy depends on the filter that is used in the receiver. If the chip rate is changed, the filter must change to match the chip waveform, and that changes the effective energy in the diffuse component. The specular energy does not depend on the filter, however, for reasons given below. Thus, it is instructive to examine the performance of the receiver as a function of the diffuse SNR for a fixed specular SNR.

In Fig. 7, the average probability of error is shown as a function of the diffuse SNR for three different values of the specular SNR. The most significant conclusion to be drawn from the curves of Fig. 7 is that the probability of error is not a monotonic function of the diffuse SNR. If the specular SNR is, for example, 25 dB, we see that the probability of error increases as the diffuse SNR is increased to approximately 25 dB, and the probability of error decreases as the diffuse SNR is increased beyond 25 dB.

For a fixed value of the specular SNR, an increase in the diffuse SNR gives an increase in the total SNR, thus, the curves in Fig. 7 also show that the probability of error is not a monotonic function of the total SNR. As illustrated by the slopes of these curves in certain regions, an increase in the total SNR can degrade the performance of the receiver. The reason for this is that the diffuse signal is subject to fading from destructive interference, and this destructive interference can detract from the specular component. Thus, an increase in the effective energy of the diffuse component can degrade the performance that would otherwise be obtained if we were able to demodulate only the specular component. As the diffuse SNR is increased further, it eventually dominates and the error probability decreases in approximately an inverse linear fashion, as is characteristic of Rayleigh fading. That is, for very large values of diffuse SNR, (12) implies $P_e \approx [2 + (E_d/N_0)]^{-1}$.

This observation has implications regarding the choice of the chip rate that can be illustrated with a deterministic model.

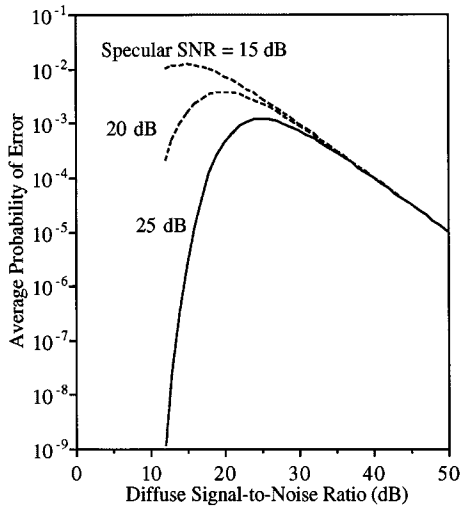


Fig. 7. Probability of error for a Rician fading channel.

Using the same analogy as in Section IV, we model the channel as a time-invariant linear filter with impulse response $h(t)$. Now, however, we let $h(t) = \delta(t) + h_c(t)$, representing a discrete component and a continuous component $h_c(t)$. The discrete component of the impulse response produces the specular component of the received signal, and the continuous component of the impulse response produces the diffuse component of the received signal. As an example, let $h_c(t) = 1$ for $-\varepsilon \leq t \leq \varepsilon$ and $h_c(t) = 0$ otherwise, where $0 < \varepsilon < T_c$.

Suppose $\beta_0(t)$ is the transmitted signal, and consider the diffuse component of the output of the filter that is matched to $\beta_0(t)$. Recall that the output of the matched filter is given by $y = R * h$. It follows that the sampled output of the matched filter is

$$y(0) = R(0) + \int_{-\varepsilon}^{\varepsilon} R(t) dt. \quad (13)$$

The value of $y(0)$ decreases as T_c decreases, as can be seen by observing the behavior of $R(t)$ as T_c decreases. Recall that the triangular autocorrelation function has a base of width $2T_c$ and its height is $R(0) = E_a$. As the chip rate is increased, the value of T_c is decreased. Thus, the triangular autocorrelation function becomes narrower, but its height is constant. The specular component does not change; it remains fixed at the value $R(0) = E_a$. However, the diffuse component decreases, as is clear from the integral in (13). As this illustration shows, a decrease in T_c , which corresponds to an increase in the chip rate, results in a decrease in the diffuse component of the output of the matched filter and a corresponding decrease in the diffuse SNR. This is consistent with results obtained for more general channels.

A more precise model of a diffuse multipath channel requires the impulse response to be random, thus, the analysis is more difficult than for the simple analogy. Also, just as stationary multipath can cause time dispersion, time-varying multipath can cause Doppler spread [31], which must be included in the analysis. Nevertheless, it is true that an increase in the chip rate results in a decrease in the effective energy for the diffuse component, even for more general multipath channels. As pointed

out in [11], several researchers have reported that the despread signal becomes more specular and less diffuse as the chip rate is increased. If the system is operating in a region in which a decrease in the diffuse SNR results in a decrease in the probability of error, then the chip rate should be increased to improve performance. This is the case in our illustration if the diffuse SNR is less than the specular SNR. More generally, we see from Fig. 7 that very low error probabilities require small values of the diffuse SNR or very large values of the diffuse SNR. Achieving a small probability of error by increasing the diffuse SNR requires a significant increase in transmitter power. For example, an error probability of 10^{-5} requires a diffuse SNR of approximately 50 dB if we increase the diffuse SNR rather than decrease it.

It should not be inferred that the diffuse SNR can be varied by a large amount by changing only the chip rate. For the illustration in which the diffuse multipath channel is modeled as a filter with a rectangular impulse response of width 2ε , the diffuse SNR changes by at most 6 dB over the range $0 < R_c < \varepsilon^{-1}$. Larger changes in the diffuse SNR require changes in the transmitter power, which also change the specular SNR. However, as seen from Fig. 7, there is a region of each curve for which even a modest decrease in the diffuse SNR results in a significant improvement in the error probability.

VI. RAKE RECEIVER

A significant step toward matching the filter to the received signal is accomplished by employing a rake receiver. The concept of the rake receiver originated in the mid 1950s and it has been applied in various forms throughout the four decades since its first description in the open literature [24]. Although each individual multipath component may undergo fading as the communication terminals move or the propagation conditions change, it is unlikely that all of the components fade simultaneously. Direct-sequence spread spectrum permits the resolution of a number of different components, thereby providing a certain amount of diversity. The receiver can take advantage of this diversity by proper combining of the multipath components that the spread-spectrum receiver is able to resolve. The amplitudes, time delays, and phases of the multipath components may change over time, and the combining method must adapt to the changes in these parameters.

Early implementations of rake receivers are based on the spacing of a large number of taps at equal intervals along a delay line. Processing is applied to the signal at each tap to produce a statistic for that tap, and a gain may be applied to each statistic. Phase shifts may be applied as well, so the gains may be considered to be complex numbers. The rake receiver processes a combination of the weighted statistics (weighted by the gains) in order to collect the energy that arrives at the receiver via a large number of propagation paths. A number of different combining methods are available (e.g., [32] and [37]), and the choice for a given application involves a tradeoff between performance and complexity. There are also a number of different algorithms for adjusting the gains to compensate for any changes that occur in the characteristics of the multipath channel as a result of motion of the communication terminals or fluctuations in the propagation medium.

Some of the more recent implementations for mobile radio systems employ a relatively small number of demodulators, each of which can track, despread, and demodulate one component of the received signal (e.g., [17]). One or more of the demodulators may be used to scan a range of time offsets in order to detect multipath components that are not currently being demodulated. Demodulators used for this purpose are often referred to as searchers. The remainder of the demodulators, referred to as data demodulators, are used to despread and demodulate multipath components that were previously detected. The outputs of the data demodulators are combined to provide decision statistics for subsequent processing, such as soft-decision decoding. In theory, at least, if there are no more paths than demodulators and if the system's chip rate is sufficiently high, the rake receiver can resolve each of the specular multipath components. If the chip rate is not high enough to isolate all of the multipath components, at least one of the data demodulators may have to process a signal that has two or more specular components. This can result in significant fading of the output of such a demodulator.

In the design of a rake receiver for a given application, it is important to know the features of the channels over which the communication system must operate. The chip rate determines the maximum number of resolvable components that can be obtained. Some channels may not generate this number of multipath components, however, in which case the system cannot achieve the maximum level of diversity that spread spectrum is capable of providing. In the event that there are more demodulators than multipath components, the extra demodulators can be used as searchers in order to accelerate the process of scanning for new multipath components.

The nature of the propagation medium for the system plays a major role in the design of a rake receiver. Statistics on the range of path delays and the number of paths that are likely to be encountered are very important. Such statistics should be used to determine the required chip rate, the number of demodulators in the rake receiver, and the time window that must be scanned in order to locate the multipath components. Including more demodulators than necessary increases system cost, and having too few demodulators results in unused energy in the received signal.

VII. CONCLUDING REMARKS

Deterministic models for multipath channels provide insight into the tradeoffs that face the spread-spectrum system designer. The presentations given in Sections IV and V are intuitive rather than rigorous, and they are not intended to be performance analyses. The hope is that the presentations are useful to those who make multipath propagation measurements or design and implement spread-spectrum systems that must cope with the effects of multipath propagation. It is also hoped that this paper serves as an introduction to the subject and a guide to the relevant literature. The list of references represents fewer than one-tenth of the extremely useful publications on the topics of spread-spectrum and multipath channels. However, many of the references in the list have extensive bibliographies that include many excellent choices for additional reading on the subjects that are mentioned only briefly in this paper.

Most of the original results on the mathematical modeling of fading channels are contained in [1] and [31]. An excellent overview of fading and its effects on the performance of communication systems is provided in [32]. Readers who wish to obtain more information on multipath channels and their characterizations are referred to [37]. Additional explanations of various aspects of direct-sequence spread spectrum are given in [2], [8], [18]–[20], [22], [23], and [29]. References on sequences for use in spread spectrum include [6] and [27]. For an excellent discussion of wide-band CDMA systems see [11], and further descriptions of the effects of the chip rate on the performance of direct-sequence spread-spectrum communications over multipath channels can be found in [14]–[16]. Applications to GPS systems are described in [4] and [30], and information on cellular CDMA systems is provided by [17], [25], [34], and [38]. Third-generation cellular CDMA systems are described in [7]. Two forthcoming issues of the *IEEE JOURNAL ON SELECTED AREAS IN COMMUNICATIONS* [9], [10] are highly recommended for those who seek additional information on channel modeling and propagation effects in wireless communications.

ACKNOWLEDGMENT

The author would like to thank several graduate students and faculty colleagues at Clemson University for comments that were very beneficial in the preparation of this paper. Dr. L. J. Greenstein and Prof. L. B. Milstein provided many excellent suggestions for improvements in the original manuscript.

REFERENCES

- [1] P. A. Bello, "Characterization of randomly time-variant linear channels," *IEEE Trans. Commun. Syst.*, vol. CS-11, pp. 8360–8393, Dec. 1963.
- [2] C. R. Cahn, "Spread spread applications and state-of-the-art equipments," in *Spread Spectrum Communications*, ser. AGARD Lecture 58: North Atlantic Treaty Org., July 1973.
- [3] R. C. Dixon, *Spread Spectrum Systems*, 2nd ed. New York: Wiley, 1984.
- [4] P. K. Enge, "The global positioning system: Signals, measurements, and performance," *Int. J. Wireless Inform. Networks*, vol. 1, no. 2, pp. 83–105, 1994.
- [5] S. W. Golomb *et al.*, *Digital Communications With Space Applications*. Englewood Cliffs, NJ: Prentice-Hall, 1964.
- [6] S. W. Golomb, *Shift Register Sequences*. San Francisco, CA: Holden-Day, 1967.
- [7] H. Holma and A. Toskala, *WCDMA for UMTS: Radio Access for Third Generation Mobile Communications*. New York: Wiley, 2000.
- [8] J. K. Holmes, *Coherent Spread Spectrum Systems*. New York: Wiley, 1982.
- [9] "Channel and propagation models for wireless system design: Part I," *IEEE J. Select. Areas Commun.*, May 2002, to be published.
- [10] "Channel and propagation models for wireless system design: Part II," *IEEE J. Select. Areas Commun.*, Aug. 2002, to be published.
- [11] L. B. Milstein, "Wideband code division multiple access," *IEEE J. Select. Areas Commun.*, vol. 18, pp. 1344–1354, Aug. 2000.
- [12] L. B. Milstein *et al.*, "On the feasibility of a CDMA overlay for personal communication networks," *IEEE J. Select. Areas Commun.*, vol. 10, pp. 655–668, May 1992.
- [13] D. L. Nicholson, *Spread Spectrum Signal Design: LPE & AJ Systems*. Rockville, MD: Computer Science, 1988.
- [14] D. L. Neneaker and M. B. Pursley, "On the chip rate of CDMA systems with doubly selective fading and rake reception," *IEEE J. Select. Areas Commun.*, vol. 12, pp. 853–861, June 1994.
- [15] —, "Selection of spreading sequences for direct-sequence spread-spectrum communications over a doubly selective fading channel," *IEEE Trans. Commun.*, vol. 42, pp. 3171–3177, Dec. 1994.

- [16] —, "The effects of sequence selection on DS spread-spectrum with selective fading and rake reception," *IEEE Trans. Commun.*, vol. 44, pp. 229–237, Feb. 1996.
- [17] R. Padovani, "Reverse link performance of IS-95 based cellular systems," *IEEE Pers. Commun.*, vol. 1, no. 3, pp. 28–34, third quarter, 1994.
- [18] R. L. Peterson, R. E. Ziemer, and D. E. Borth, *Introduction to Spread Spectrum Communications*. Englewood Cliffs, NJ: Prentice-Hall, 1995.
- [19] M. B. Pursley, "Performance evaluation for phase-coded spread-spectrum multiple-access communication—Part I: System analysis," *IEEE Trans. Commun.*, vol. COM-25, pp. 795–799, Aug. 1977.
- [20] —, "Spread-spectrum multiple-access communications," in *Multi-User Communication Systems*, G. Longo, Ed. Berlin, Germany: Springer-Verlag, 1981, pp. 139–199.
- [21] —, "Frequency-hop transmission for satellite packet switching and terrestrial packet radio networks," *IEEE Trans. Inform. Theory*, vol. IT-32, pp. 652–667, Sept. 1986.
- [22] —, "The role of spread spectrum in packet radio networks," *Proc. IEEE*, vol. 75, pp. 116–134, Jan. 1987.
- [23] —, "Spread-spectrum communications," in *Encyclopedia of Telecommunications*. New York: Marcel Dekker, 1998, vol. 15, pp. 381–428.
- [24] R. Price and P. E. Green, Jr., "A communication technique for multipath channels," *Proc. IRE*, vol. 46, pp. 555–570, Mar. 1958.
- [25] M. Y. Rhee, *CDMA Cellular Mobile Communications and Network Security*. Englewood Cliffs, NJ: Prentice-Hall, 1998.
- [26] S. O. Rice, "Mathematical analysis of random noise," *Bell Syst. Tech. J.*, vol. 24, pp. 46–156, Jan. 1945.
- [27] D. V. Sarwate and M. B. Pursley, "Crosscorrelation properties of pseudorandom and related sequences," *Proc. IEEE*, vol. 68, pp. 593–619, May 1980.
- [28] R. A. Scholtz, "The origins of spread-spectrum communications," *IEEE Trans. Commun.*, vol. COM-30, pp. 822–852, May 1982.
- [29] M. K. Simon, J. K. Omura, R. A. Scholtz, and B. K. Levitt, *Spread Spectrum Communications*. Rockville, MD: Computer Science, 1985, 3 vols.
- [30] J. J. Spilker, Jr., "GPS signal structure and performance characteristics," *Navigat. J. Inst. Navigat.*, vol. 25, no. 2, pp. 121–146, Summer 1978.
- [31] S. Stein, *Communication Systems and Techniques*. New York: McGraw-Hill, 1966, pt. III, pp. 561–584.
- [32] —, "Communication over fading radio channels," in *Encyclopedia of Telecommunications*. New York: Marcel Dekker, 1992, vol. 3, pp. 261–303.
- [33] "TIA/EIA interim standard: Mobile station-base station compatibility standard for dual-mode wideband spread spectrum cellular system," Telecommun. Ind. Assoc., Washington, DC, TIA/EIA/IS-95, July 1993.
- [34] "TIA/EIA standard: Mobile station-base station compatibility standard for wideband spread spectrum cellular systems," Telecommun. Ind. Assoc., Washington, DC, TIA/EIA-95-B, Mar. 1999.
- [35] D. J. Torrieri, *Principles of Military Communication Systems*. Norwood, MA: Artech House, 1981.
- [36] G. L. Turin, "Error probabilities for binary symmetric ideal reception through nonselective slow fading and noise," *Proc. IRE*, vol. 46, pp. 1603–1619, Sept. 1958.
- [37] —, "Introduction to spread-spectrum antimultipath techniques and their application to urban digital radio," *Proc. IEEE*, vol. 68, pp. 328–353, Mar. 1980.
- [38] A. J. Viterbi, *CDMA: Principles of Spread Spectrum Communication*. Reading, MA: Addison-Wesley, 1995.
- [39] R. E. Ziemer and R. L. Peterson, *Digital Communications and Spread Spectrum Systems*. New York: Macmillan, 1985.



Michael B. Pursley (S'68–M'68–SM'77–F'82) received the B.S. degree (with highest distinction) and M.S. degree from Purdue University, West Lafayette, IN, respectively, both in electrical engineering, and the Ph.D. degree in electrical engineering from the University of Southern California, Los Angeles.

His industrial experience is primarily with the Space and Communications Group, Hughes Aircraft Company, from 1968 to 1974. In 1974, he served as an Acting Assistant Professor in the System Science Department, University of California at Los Angeles.

From June 1974 to July 1993, he was with the Department of Electrical and Computer Engineering and the Coordinated Science Laboratory, University of Illinois at Urbana-Champaign, where he became a Professor in 1980. He is currently the Holcombe Professor of Electrical and Computer Engineering at Clemson University, Clemson, SC. His research is in the general area of wireless communications with emphasis on spread-spectrum communications, applications of error-control coding, adaptive protocols for packet radio networks, and mobile wireless communication systems and networks. He is a member of the Editorial Advisory Board for the *International Journal of Wireless Information Networks*.

Dr. Pursley is a member of Phi Eta Sigma, Tau Beta Pi, and the Institute of Mathematical Statistics, and is an honorary member of the Golden Key National Honor Society. He held two three-year terms on the Board of Governors of the IEEE Information Theory Society, and he was elected President of that society in 1983. He was a member of the Editorial Board of the *PROCEEDINGS OF THE IEEE* (1984–1991). He is a senior editor for the *IEEE JOURNAL OF SELECTED AREAS IN COMMUNICATIONS*. He served as Technical Program chairman for the 1979 IEEE International Symposium on Information Theory, Grignano, Italy, and was co-chairman for the 1995 IEEE International Symposium on Information Theory, Whistler, Canada. He was the recipient of a 1984 IEEE Centennial Medal, the 1996 Ellersick Award of the IEEE Communications Society, and the 1999 IEEE Military Communications Conference Award for Technical Achievement. He was also the recipient of the IEEE Millennium Medal and the 2000 Clemson University Alumni Award for Outstanding Achievement in Research.