

A. F. Kohn (2006). Autocorrelation and Cross-Correlation Methods. In: Wiley Encyclopedia of Biomedical Engineering, Ed. Metin Akay, Hoboken: John Wiley & Sons, pp. 260-283. ISBN: 0-471-24967-X.

AUTOCORRELATION AND CROSSCORRELATION METHODS

ANDRÉ FABIO KOHN
University of São Paulo
Sao Paulo, Brazil

1. INTRODUCTION

Any physical quantity that varies with *time* is a **signal**. Examples from physiology are being an electrocardiogram (ECG), an electroencephalogram (EEG), an arterial pressure waveform, and a variation of someone's blood glucose concentration along time. In Fig. 1, one can see examples of two signals: an electromyogram (EMG) in (a) and the corresponding force in (b). When a muscle contracts, it generates a force or torque while it generates an electrical signal that is the EMG (1). The experiment to obtain these two signals is simple. The subject is seated with one foot strapped to a pedal coupled to a force or torque meter. Electrodes are attached to a calf muscle of the right foot (*m. soleus*), and the signal is amplified. The subject is instructed to produce an alternating pressure on the pedal, starting after an initial rest period of about 2.5 s. During this initial period, the foot stays relaxed on the pedal, which corresponds to practically no EMG signal and a small force because of the foot resting on the pedal. When the subject controls voluntarily the alternating contractions, the random-looking EMG has waxing and waning modulations in its amplitude while the force also exhibits an oscillating pattern (Fig. 1).

Biological signals vary as time goes on, but when they are measured, for example, by a computerized system, the measures are usually only taken at pre-specified times, usually at equal time intervals. In a more formal jargon, it is said that although the original biological signal is defined in continuous time, the measured biological signal is defined in discrete time. For **continuous-time signals**, the time variable t takes values either from $-\infty$ to $+\infty$ (in theory) or in an interval between t_1 and t_2 (a subset of the real numbers, t_1 indicating the time when the signal

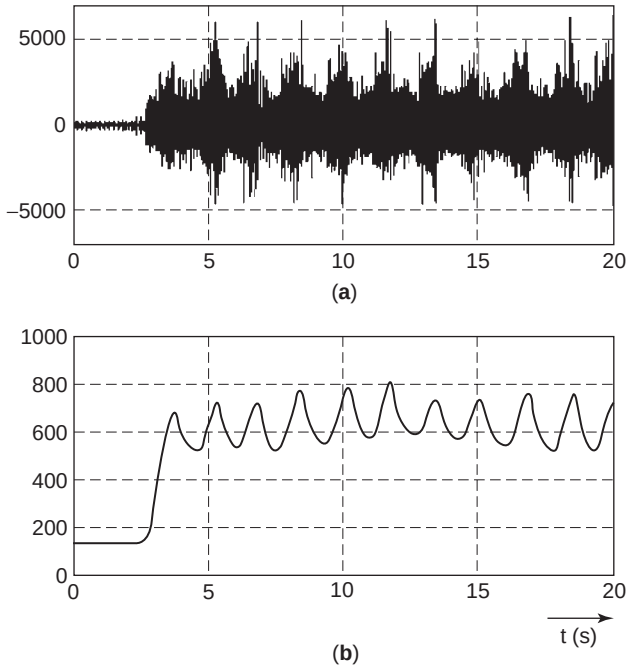


Figure 1. Two signals obtained from an experiment involving a human pressing a pedal with his right foot. (a) The EMG of the soleus muscle and (b) the force or torque applied to the pedal are represented. The abscissae are in seconds and the ordinates are in arbitrary units. The ordinate calibration is not important here because in the computation of the correlation a division by the standard deviation of each signal exists. Only the first 20 s of the data are shown here. A 30 s data record was used to compute the graphs of Figs. 2 and 3.

started being observed in the experiment and t_2 the final time of observation). Such signals are indicated as $y(t)$, $x(t)$, $w(t)$, and so on. On the other hand, a **discrete-time signal** is a set of measurements taken sequentially in time (e.g., at every millisecond). Each measurement point is usually called a sample, and a discrete-time signal is indicated by $y(n)$, $x(n)$, and $w(n)$, where the index n is an integer that points to the order of the measurements in the sequence. Note that the time interval T between two adjacent samples is not shown explicitly in the $y(n)$ representation, but this information is used whenever an interpretation is required based on continuous-time units (e.g., seconds). As a result of the low price of computers and microprocessors, almost any equipment used today in medicine or biomedical research uses digital signal processing, which means that the signals are functions of discrete time.

From basic probability and statistics theory, it is known that in the analysis of a *random variable* (e.g., the height of a population of human subjects), the mean and the variance are very useful quantifiers (2). When studying the linear relationship between two random variables (e.g., the height and the weight of individuals in a population), the correlation coefficient is an extremely useful quantifier (2). The correlation coefficient between N measurements of pairs of random variables, such as the weight w and height h of human subjects, may be

estimated by

$$\rho = \frac{\frac{1}{N} \sum_{i=0}^{N-1} (w(i) - \bar{w}) \cdot (h(i) - \bar{h})}{\sqrt{\left(\frac{1}{N} \sum_{i=0}^{N-1} (w(i) - \bar{w})^2\right) \cdot \left(\frac{1}{N} \sum_{i=0}^{N-1} (h(i) - \bar{h})^2\right)}}, \quad (1)$$

where $[w(i), h(i)]$, $i = 0, 1, 2, \dots, N-1$ are the N pairs of measurements (e.g., from subject number 0 up to subject number $N-1$); $\bar{w}$ and $\bar{h}$ are the mean values computed from the N values of $w(i)$ and $h(i)$, respectively. Sometimes ρ is called the *linear correlation coefficient* to emphasize that it quantifies the degree of linear relation between two variables. If the correlation coefficient between the two variables w and h is near the maximum value 1, it is said that the variables have a strong positive linear correlation, and the measurements will gather around a line with positive slope when one variable is plotted against the other. On the contrary, if the correlation coefficient is near the minimum attainable value of -1 , it is said that the two variables have a strong negative linear correlation. In this case, the measured points will gather around a negatively sloped line. If the correlation coefficient is near the value 0, the two variables are not linearly correlated and the plot of the measured points will show a spread that does not follow any specific straight line. Here, it may be important to note that two variables may have a strong nonlinear correlation and yet have almost zero value for the linear correlation coefficient ρ . For example, 100 normally distributed random samples were generated by computer for a variable h , whereas variable w was computed according to the quadratic relation $w = 300 * (h - \bar{h})^2 + 50$. A plot of the pairs of points (w, h) will show that the samples follow a parabola, which means that they are strongly correlated along such a parabola. On the other hand, the value of ρ was 0.0373. Statistical analysis suggests that such a low value of linear correlation is not significantly different to zero. Therefore, a near zero value of ρ does not necessarily mean the two variables are not associated with one another, it could mean that they are nonlinearly associated (see Extensions and Further Applications).

On the other hand, a *random signal* is a broadening of the concept of a random variable by the introduction of variations along time and is part of the theory of random processes. Many biological signals vary in a random way in time (e.g., the EMG in Fig. 1) and hence their mathematical characterization has to rely on probabilistic concepts (3–5). For a random signal, the mean and the autocorrelation are useful quantifiers, the first indicating the constant level about which the signal varies and the second indicating the statistical dependencies between the values of two samples taken at a certain time interval. The time relationship between two random signals may be analyzed by the cross-correlation, which is very often used in biomedical research.

Let us analyze briefly the problem of studying quantitatively the time relationship between the two signals shown in Fig. 1. Although the EMG in Fig. 1a looks erratic, its amplitude modulations seem to have some periodicity. Such slow amplitude modulations are sometimes

called the “envelope” of the signal, which may be estimated by smoothing the absolute value of the signal. The force in Fig. 1b is much less erratic and exhibits a clearer oscillation. Questions that may develop regarding such signals (the EMG envelope and the force) include: what periodicities are involved in the two signals? Are they the same in the two signals? If so, is there a delay between the two oscillations? What are the physiological interpretations? To answer the questions on the periodicities of each signal, one may analyze their respective autocorrelation functions, as shown in Fig. 2. The autocorrelation of the absolute value of the EMG (a simple estimate of the envelope) shown in Fig. 2a has low-amplitude oscillations, those of the force (Fig. 2b) are large, but both have the same periodicity. The much lower amplitude oscillations in the autocorrelation function of the absolute value of the EMG when compared with that of the force autocorrelation function reflects the fact that the periodicity in the EMG amplitude modulations is masked to a good degree by a random activity, which is not the case for the force signal. To analyze the time relationship between the EMG envelope and the force, their cross-correlation is shown in Fig. 3a. The cross-correlation function in this figure has the same period of oscillation as that of the random signals. In the more refined view of Fig. 3b, it can be seen that the peak occurring closer to zero time shift does so at a

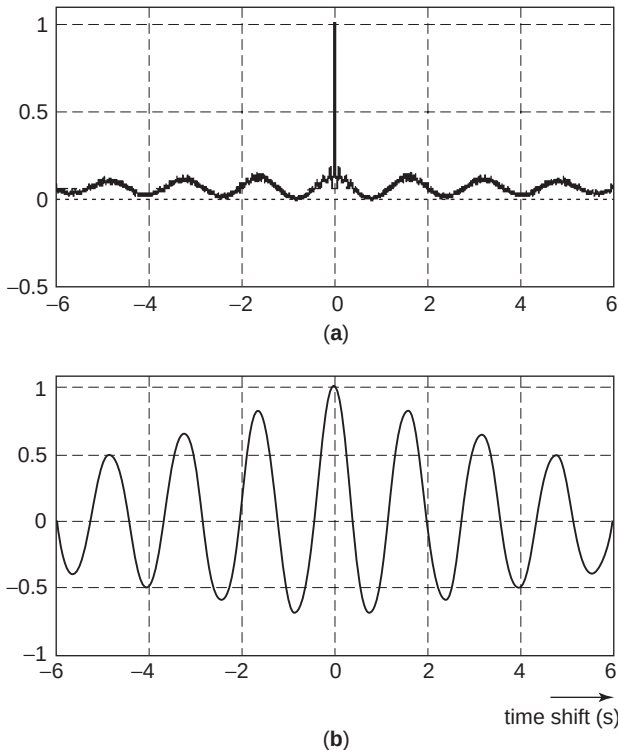


Figure 2. Autocorrelation functions of the signals shown in Fig. 1. (a) shows the autocorrelation function of the absolute value of the EMG and (b) shows the autocorrelation function of the force. These autocorrelation functions were computed based on the correlation coefficient, as explained in the text. The abscissae are in seconds and the ordinates are dimensionless, ranging from -1 to 1. For these computations, the initial transients from 0 to 5 s in both signals were discarded.

negative delay, meaning the EMG precedes the soleus muscle force. Many factors, experimental and physiologic, contribute to such a delay between the electrical activity of the muscle and the torque exerted by the foot.

Signal processing tools such as the autocorrelation and the cross-correlation have been used with much success in a number of biomedical research projects. A few examples will be cited for illustrative purposes. In a study of absence epileptic seizures in animals, the cross-correlation between waves obtained from the cortex and a brain region called the subthalamic nucleus was a key tool to show that the two regions have their activities synchronized by a specific corticothalamic network (6). The cross-correlation function was used in Ref. 7 to show that insulin secretion by the pancreas is an important determinant of insulin clearance by the liver. In a study of preterm neonates, it was shown in Ref. 8 that the correlation between the heart rate variability (HRV) and the respiratory rhythm was similar to that found in the fetus. The same authors also employed the correlation analysis to compare the effects of two types of artificial ventilation equipment on the HRV-respiration interrelation.

After an interpretation is drawn from a cross-correlation study, this signal processing tool may be potentially useful for diagnostic purposes. For example, in healthy subjects, the cross-correlation between arterial blood pressure and intracranial blood flow showed a negative peak at positive delays, differently from patients with a malfunctioning cerebrovascular system (9).

Next, the step-by-step computations of an autocorrelation function will be shown based on the known concept of correlation coefficient of statistics. Actually, different, but related, definitions of autocorrelation and cross-correlation exists in the literature. Some are normalized versions of others, for example. The definition to be given in this section is not the one usually studied in undergraduate engineering courses, but is being presented here first because it is probably easier to understand by readers from other backgrounds. Other definitions will be presented in later sections and the links between them will be readily apparent. In this section, the single term *autocorrelation* shall be used for simplicity and, later (see Basic Definitions), will present some of the more precise names that have been associated with the definition presented here (10,11).

The approach of defining an autocorrelation function based on the cross-correlation coefficient should help in the understanding of what the autocorrelation function tells us about a random signal. Assume that we are given a random signal $x(n)$, with n being the counting variable: $n = 0, 1, 2, \dots, N - 1$. For example, the samples of $x(n)$ may have been measured at every 1 ms, there being a total of N samples.

The mean or average of signal $x(n)$ is the value $\bar{x}$ given by

$$\bar{x} = \frac{1}{N} \sum_{n=0}^{N-1} x(n), \quad (2)$$

and gives an estimate of the value about which the signal varies. As an example, in Fig. 4a, the signal $x(n)$ has a mean that is approximately equal to 0. In addition to the

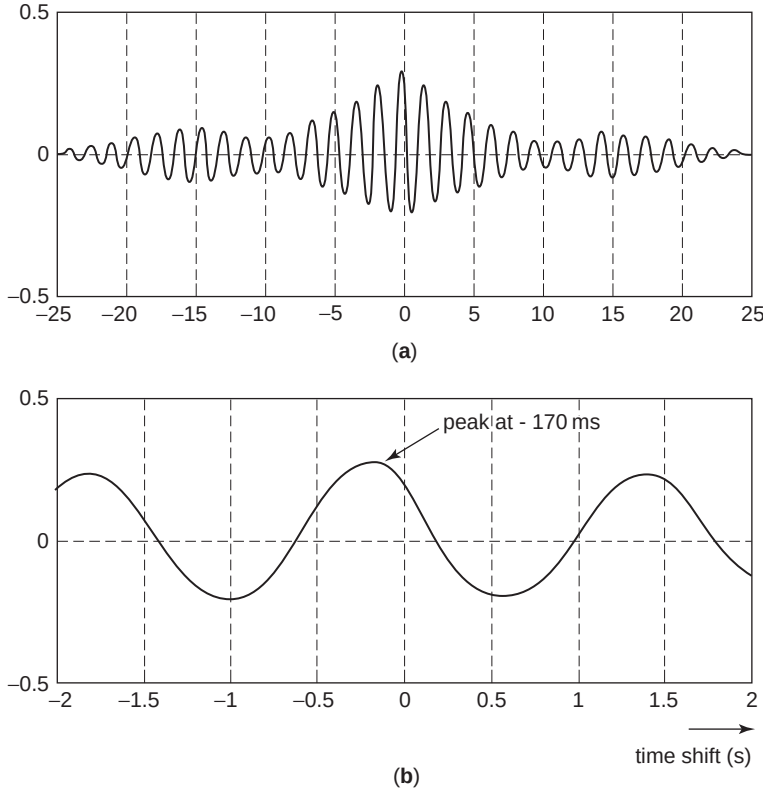


Figure 3. Cross-correlation between the absolute value of the EMG and the force signals shown in Fig. 1. (a) shows the full cross-correlation and (b) shows an enlarged view around abscissa 0. The abscissae are in seconds and the ordinates are dimensionless, ranging from -1 to 1 . For this computation, the initial transients from 0 to 5 s in both signals were discarded.

mean, another function is needed to characterize how $x(n)$ varies in time. In this example, one can see that $x(n)$ has some periodicity, oscillating with positive and negative peaks repeating approximately at every 10 samples (Fig. 4a). The new function to be defined is the autocorrelation $\rho_{xx}(k)$ of $x(n)$, which will quantify how much a given signal is similar to time-shifted versions of itself (5). One way to compute it is by using the following formula based on the definition (Equation 1)

$$\rho_{xx}(k) = \frac{\frac{1}{N} \sum_{n=0}^{N-1} (x(n-k) - \bar{x}) \cdot (x(n) - \bar{x})}{\frac{1}{N} \sum_{n=0}^{N-1} (x(n) - \bar{x})^2}, \quad (3)$$

where $x(n)$ is supposed to have N samples. Any sample outside the range $[0, N-1]$ is taken to be zero in the computation of $\rho_{xx}(k)$.

The computation steps are as follows:

- Compute the correlation coefficient between the N samples of $x(n)$ paired with the N samples of $x(n)$ and call it $\rho_{xx}(0)$. The value of $\rho_{xx}(0)$ is equal to 1 because any pair is formed by two equal values [e.g., $[x(0), x(0)]$, $[x(1), x(1)]$, ..., $[x(N-1), x(N-1)]$] as seen in the scatter plot of the samples of $x(n)$ with those of $x(n)$ in Fig. 5a. The points are all along the diagonal, which means the correlation coefficient is unity.
- Next, shift $x(n)$ by one sample to the right, obtaining $x(n-1)$, and then determine the correlation coefficient between the samples of $x(n)$ and $x(n-1)$ (i.e., for $n=1$, take the pair of samples $[x(1), x(0)]$, for $n=2$ take the pair $[x(2), x(1)]$ and so on, until the pair

$[x(N-1), x(N-2)]$). The correlation coefficient of these pairs of points is denoted $\rho_{xx}(1)$.

- Repeat for a two-sample shift and compute $\rho_{xx}(2)$, for a three-sample shift and compute $\rho_{xx}(3)$, and so on. When $x(n)$ is shifted by 3 samples to the right (Fig. 4b), the resulting signal $x(n-3)$ has its peaks and valleys still repeating at approximately 10 samples, but these are no longer aligned with those of $x(n)$. When their scatter plot is drawn (Fig. 5b), it seems that their correlation coefficient is near zero, so we should have $\rho_{xx}(3) \approx 0$. Note that as $x(n-3)$ is equal to $x(n)$ delayed by 3 samples, the need exists to define what the values of $x(n-3)$ are for $n=0, 1, 2$. As $x(n)$ is known only from $n=0$ onwards, we make the three initial samples of $x(n-3)$ equal to 0, which has sometimes been called in the engineering literature as zero padding.
- Shifting $x(n)$ by 5 samples to the right (Fig. 4c), generates a signal $x(n-5)$ still with the same periodicity as the original $x(n)$, but with peaks aligned with the valleys in $x(n)$. The corresponding scatter plot (Fig. 5c) indicates a negative correlation coefficient.
- Finally, shifting $x(n)$ by a number of samples equal to the approximate period gives $x(n-10)$, which has its peaks (valleys) approximately aligned with the peaks (valleys) of $x(n)$, as can be seen in Fig. 4d. The corresponding scatter plot (Fig. 5d) indicates a positive correlation coefficient. If $x(n)$ is shifted by multiples of 10, there will again be coincidences between its peaks and those of $x(n)$, and again the correlation coefficient of their samples will be positive.

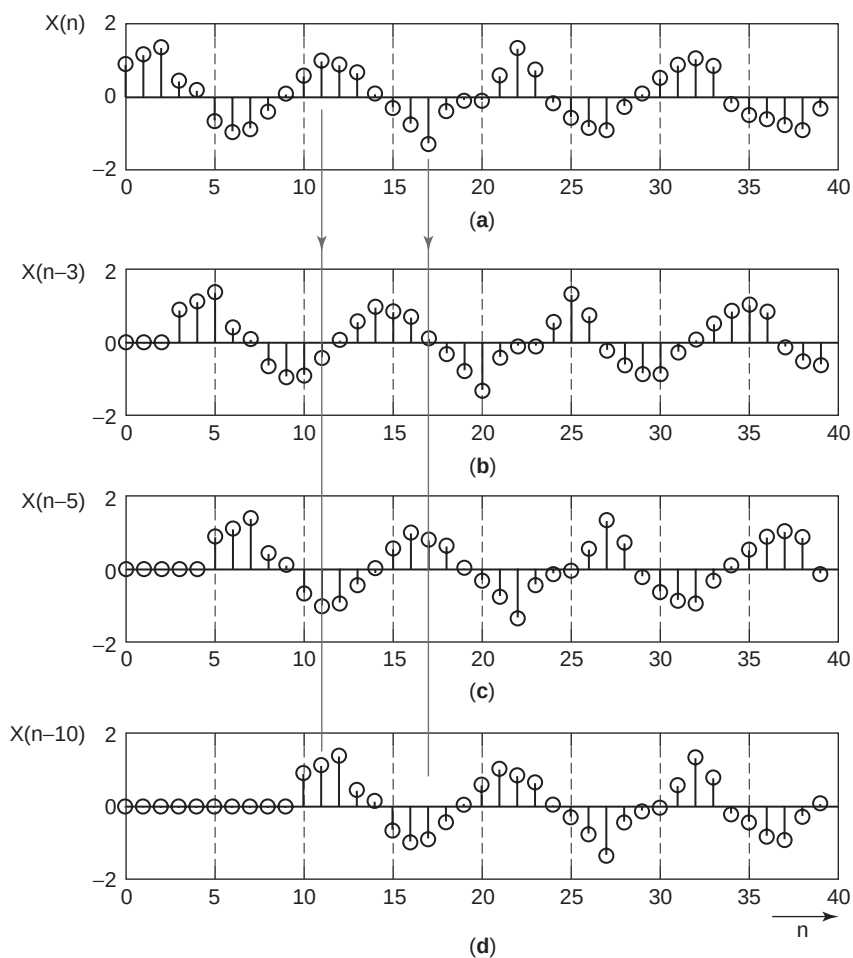


Figure 4. Random discrete-time signal $x(n)$ in (a) is used as a basis to explain the concept of autocorrelation. In (b)–(d), the samples of $x(n)$ were delayed by 3, 5, and 10 samples, respectively. The two vertical lines were drawn to help visualize the temporal relations between the samples of the reference signal at the top and the three time-shifted versions below. (This figure is available in full color at <http://www.mrw.interscience.wiley.com/ebe>.)

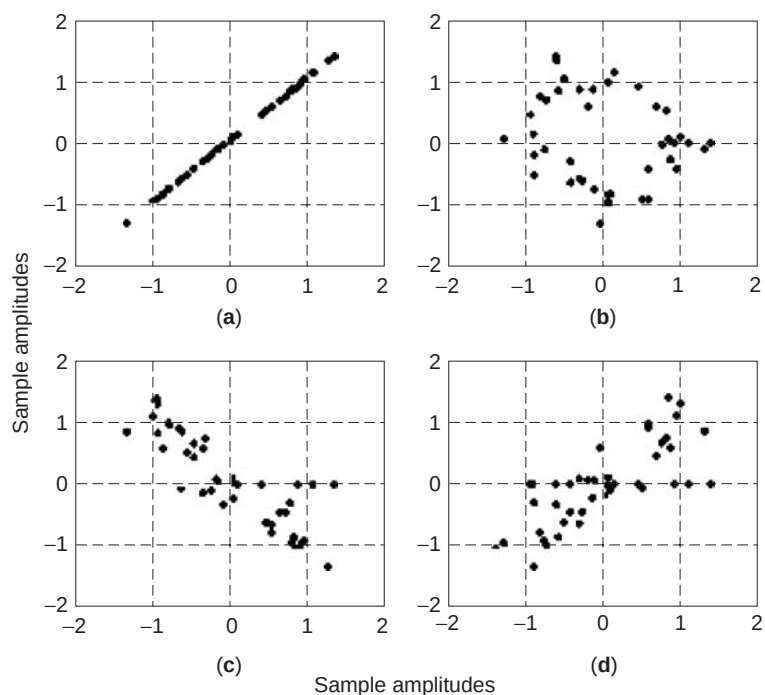


Figure 5. Scatter plots of samples of the signals shown in Fig. 4: (a) $x(n)$ and $x(n)$, (b) $x(n)$ in the abscissa and $x(n-3)$ in the ordinate, (c) $x(n)$ in the abscissa and $x(n-5)$ in the ordinate, (d) $x(n)$ in the abscissa and $x(n-10)$ in the ordinate. The computed values of the correlation coefficient from (a)–(d) were 1, -0.19 , -0.79 , and 0.66 , respectively.

Collecting the values of the correlation coefficients for the different pairs $x(n)$ and $x(n - k)$ and assigning them to $\hat{\rho}_{xx}(k)$, for positive and negative shift values k , the autocorrelation shown in Fig. 6 is obtained. In this and other figures, the hat “^” over a symbol is used to indicate estimations from data, to differentiate from the theoretical quantities, for example, as defined in Equations 14 and 15. Indeed, the values for $k = 0$, $k = 3$, $k = 5$ and $k = 10$ confirm the analyses based on the scatter plots of Fig. 5. The autocorrelation function is symmetric with respect to $k = 0$ because the correlation between samples of $x(n)$ and $x(n - k)$ is the same as the correlation between $x(n)$ and $x(n + k)$, where n and k are integers. This example suggests that the autocorrelation of a periodic signal will exhibit peaks (and valleys) repeating with the same period as the signal's. The decrease in subsequent peak values away from time shift $k = 0$ is because of the finite duration of the signal, which requires the use of zero-padding in the computation of the autocorrelation (this process will be dealt with in section 6.1 later in this chapter). In conclusion, the autocorrelation gives an idea of the similarity between a given signal and time-shifted replicas of itself. A different viewpoint is associated with the following question: Does the knowledge of the value of a sample of the signal $x(n)$ at an arbitrary time $n = L$, say $x(L) = 5.2$, give some “information-” as to the value of a future sample, say at $n = L + 3$? Intuitively, if the random signal varies “slowly,” then the answer is yes, but if it varies “fast,” then the answer is no. The autocorrelation is the right tool to quantify the signal variability or the degree of “information” between nearby samples. If the autocorrelation decays slowly from value 1 at $k = 0$ (e.g., its value is still near 1 for $k = 3$), then a sample value 5.2 of the given signal at $n = L$ tells us that at $n = L + 3$ the value of the signal will be “near” the sample value 5.2, with high probability. On

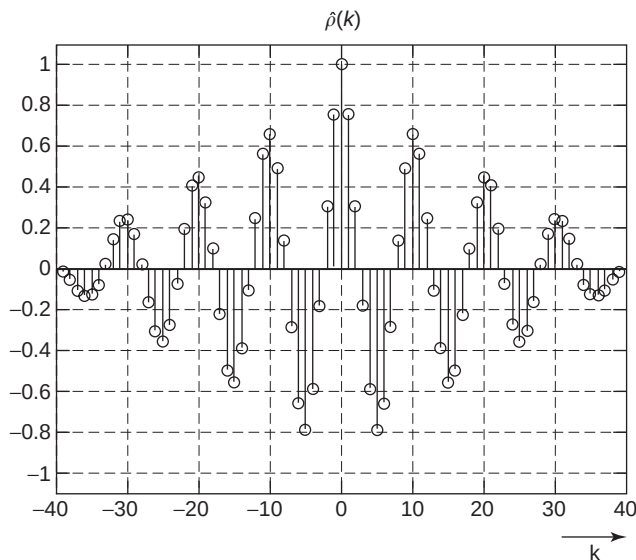


Figure 6. Autocorrelation (correlation coefficients as a function of time shift k) of signal shown in Fig. 4a. The apparently rhythmic behavior of the signal is more clearly exhibited by the autocorrelation, which indicates a repetition at every 10 samples.

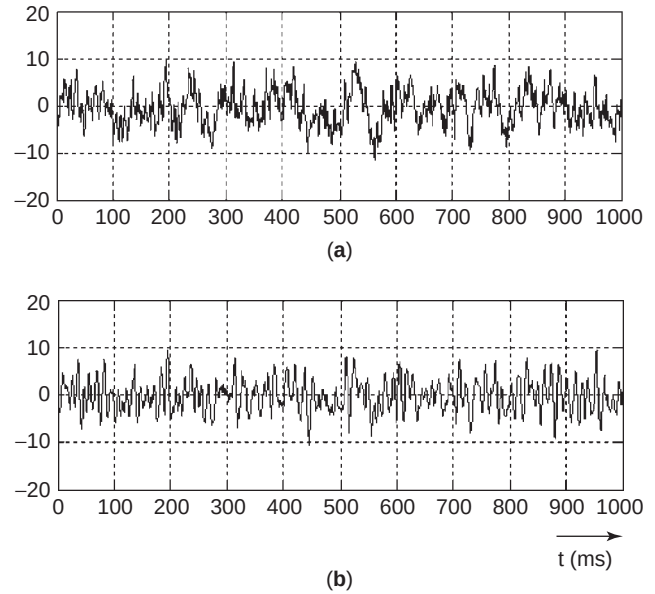


Figure 7. Two random signals measured from two different systems. They seem to behave differently, but it is difficult to characterize the differences based only on a visual analysis. The abscissae are in milliseconds.

the contrary, if the autocorrelation for $k = 3$ is already near 0, then the value of the signal at $n = L + 3$ has little or no relation to the value 5.2 attained three units of time earlier. In a loose sense, one could say that the signal has more memory in the former situation than in the latter. The autocorrelation depends on the independent variable k , which is called “lag,” “delay,” or “time shift” in the literature. The autocorrelation definition given above was based on a discrete-time case, but a similar procedure is followed for a continuous-time signal, where the autocorrelation $\rho_{xx}(\tau)$ will depend on a continuous-time variable τ .

Real-life signals are often less well-behaved than the signal shown in Fig 4a, or even those in Fig. 1. Two signals $x(t)$ and $y(t)$ shown in Fig. 7a and 7b, respectively, are more representative of the difficulties one usually encounters in extracting useful information from random signals. A visual analysis suggests that the two signals are indeed different in their “randomness,” but it is certainly not easy to pinpoint in what aspects they are different. The respective autocorrelations, shown in Fig. 8a and 8b, are monotonic for the first signal and oscillatory for the second. Such an oscillatory autocorrelation function would mean that two amplitude values in $y(t)$ taken 6 ms apart (or a time interval between about 5 ms and 10 ms) (see Fig. 8b) would have a negative correlation, meaning that if the first amplitude value is positive, the other will probably be negative and vice-versa (note that in Equation 3 the mean value of the signal is subtracted). The explanation of the monotonic autocorrelation will require some mathematical considerations, which will be presented later in this chapter in section 3. An understanding of the ways different autocorrelation functions may occur could be important in discriminating between the behaviors of a biological system subjected to two different experimental conditions, or between normal and pathological cases. In

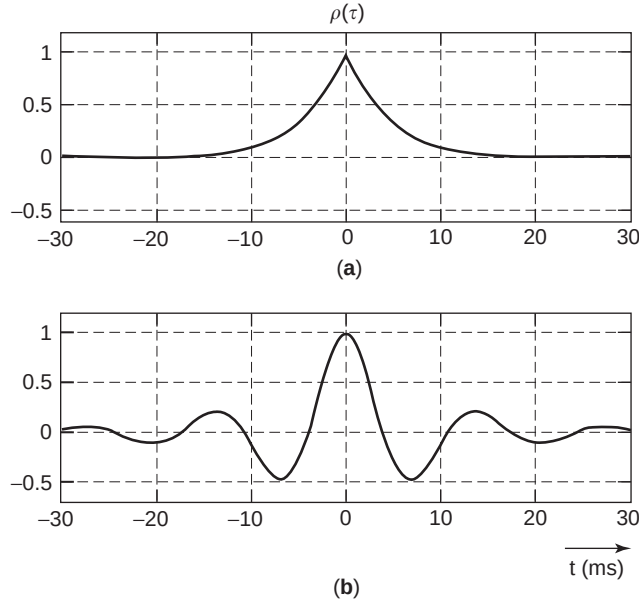


Figure 8. The respective autocorrelation functions of the two signals in Fig. 7. They are quite different from each other: in (a) the decay is monotonic to both sides from the peak value of 1 at $\tau=0$, and in (b) the decay is oscillatory. These major differences between the two random signals shown in Fig. 7 are not visible directly from their time courses. The abscissae are in milliseconds.

addition, the autocorrelation is able to uncover a periodic signal masked by noise (e.g., Fig. 1a and Fig. 2a), which is relevant in the biomedical setting because many times the biologically interesting signal is masked by other ongoing biological signals or by measurement noise. Finally, when validating stochastic models of biological systems such as neurons, autocorrelation functions obtained from the mathematical model of the physiological system may be compared with autocorrelations computed from experimental data obtained from the physiological system, see, for example, Kohn (12).

When two signals x and y are measured simultaneously in a given experiment, as in the example of Fig. 1, one may be interested in knowing if the two signals are entirely independent from each other, or if some correlation exists between them. In the simplest case, there could be a delay between one signal and the other. The cross-correlation is a frequently used tool when studying the dependency between two random signals. A (normalized) cross-correlation $\rho_{xy}(k)$ may be defined and explained in the same way as we did for the autocorrelation in Figs. 4–6 (i.e., by computing the correlation coefficients between the samples of one of the signals, $x(n)$, and those of the other signal time-shifted by k , $y(n-k)$) (5). The formula is the following:

$$\rho_{xy}(k) = \frac{\frac{1}{N} \sum_{n=0}^{N-1} (x(n) - \bar{x}) \cdot (y(n-k) - \bar{y})}{\sqrt{\left(\frac{1}{N} \sum_{n=0}^{N-1} (x(n) - \bar{x})^2\right) \cdot \left(\frac{1}{N} \sum_{n=0}^{N-1} (y(n) - \bar{y})^2\right)}}, \quad (4)$$

where both signals are supposed to have N samples each. Any sample of either signal outside the range $[0, N-1]$

is taken to be zero in the computation of $\rho_{xx}(k)$ (zero padding).

It should be clear that if one signal is a delayed version of the other, the cross-correlation at a time shift value equal to the delay between the two signals will be equal to 1. In the example of Fig. 1, the force signal at the bottom is a delayed version of the EMG envelope at the top, as indicated by the cross-correlation in Fig. 3.

The emphasis in this chapter is to present the main concepts on the auto and cross-correlation functions, which are necessary to pursue research projects in biomedical engineering.

2. AUTOCORRELATION OF A STATIONARY RANDOM PROCESS

2.1. Introduction

Initially, some concepts from random process theory shall be reviewed briefly, as covered in undergraduate courses in electrical or biomedical engineering, see, for example, Peebles (3) or Papoulis and Pillai (10). A **random process** is an infinite collection or ensemble of functions of time (continuous or discrete time), called sample functions or realizations (e.g., segments of EEG recordings). In continuous time, one could indicate the random process as $X(t)$ and in discrete time as $X(n)$. Each sample function is associated with the outcome of an experiment that has a given probabilistic description and may be indicated by $x(t)$ or $x(n)$, for continuous or discrete time, respectively. When the ensemble of sample functions is viewed at any single time, say t_1 for a continuous time process, a random variable is obtained whose probability distribution function is $F_{X,t_1}(\alpha_1) = P[X(t_1) \leq \alpha_1]$, for $\alpha_1 \in \mathbb{R}$, and where $P[\cdot]$ stands for probability. If the process is viewed at two times t_1 and t_2 , a bivariate distribution function is needed to describe the pair of resulting random variables $X(t_1)$ and $X(t_2)$: $F_{X,t_1,t_2}(\alpha_1, \alpha_2) = P[X(t_1) \leq \alpha_1, X(t_2) \leq \alpha_2]$, for $\alpha_1, \alpha_2 \in \mathbb{R}$. The random process is fully described by the joint distribution functions of any N random variables defined at arbitrary times $t_1, t_2, \dots, t_N$, for arbitrary integer number N

$$P[X(t_1) \leq \alpha_1, X(t_2) \leq \alpha_2, \dots, X(t_N) \leq \alpha_N], \quad (5)$$

for $\alpha_1, \alpha_2, \dots, \alpha_N \in \mathbb{R}$.

In many applications, the properties of the random process may be assumed to be independent of the specific values $t_1, t_2, \dots, t_N$, in the sense that if a fixed time shift T is given to all instants t_i $i = 1, \dots, N$, the probability distribution function does not change:

$$P[X(t_1 + T) \leq \alpha_1, X(t_2 + T) \leq \alpha_2, \dots, X(t_N + T) \leq \alpha_N] \\ = P[X(t_1) \leq \alpha_1, X(t_2) \leq \alpha_2, \dots, X(t_N) \leq \alpha_N]. \quad (6)$$

If Equation 6 holds for all possible values of t_i , T and N , the process is called **strictsense stationary** (3). This class of random processes has interesting properties,

such as:

$$E[X(t)] = \mu = \text{const} \quad \forall t \quad (7)$$

$$E[X(t + \tau) \cdot X(t)] = \text{function}(\tau) \quad (8)$$

The first result in Equation 7 means that the mean value of the random process is a constant value for any time t . The second result in Equation 8 means that the second-order moment defined on the process at times $t_2 = t + \tau$ and $t_1 = t$, depends only on the time difference $\tau = t_2 - t_1$ and is independent of the time parameter t . These two (Equations 7 and 8) are so important in practical applications that whenever they are true, the random process is said to be **wide-sense stationary**. This definition of stationarity is feasible to be tested in practice, and many times a random process that satisfies Equations 7, 8 is simply called “stationary” and otherwise it is simply called “nonstationary.” The autocorrelation and cross-correlation analyses developed in this chapter are especially useful for such stationary random processes. All random processes considered in this chapter will be wide-sense stationary.

In real-life applications, the wide-sense stationarity assumption is usually valid only approximately and only for a limited time interval. This interval must usually be estimated experimentally (4) or be adopted from previous research reported in the literature. This assumption certainly simplifies both the theory as well as the signal processing methods. In the present text, it shall be assumed that all random processes are wide-sense stationary. Another fundamental property that shall be assumed is that the random process is **ergodic**, meaning that any appropriate time average computed from a given sample function converges to a corresponding expected value defined over the random process (10). Thus, for example, for an ergodic process (Equation 2) would give useful estimates of the expected value of the random process (Equation 7), for sufficiently large values of N . Ergodicity assumes that a finite set of physiological recordings obtained under a certain experimental condition should yield useful estimates of the general random behavior of the physiological system, under the same experimental conditions and in the same physiological state, which is of utmost importance because in practice all we have is a sample function and from it have to estimate and infer things related to the random process that generated that sample function.

A **random signal** may be defined as a sample function or realization of a random process. Many signals measured from humans or animals exhibit some degree of unpredictability, and may be considered as random. The sources of unpredictability in biological signals may be associated with (1) a large number of uncontrolled and unmeasured internal mechanisms, (2) intrinsically random internal physicochemical mechanisms, and (3) a fluctuating environment. When measuring randomly varying phenomena from humans or animals, only one or a few sample functions of a given random process are obtained. Under the ergodic property, appropriate processing of a random signal may permit the estimation of characteristics of the random process, which is why random signal processing techniques (such as auto and cross-correlation) are so important in practice.

The complete probabilistic description of a random process Equation 5 is impossible to obtain in practical terms. Instead, first and second moments are very often employed in real-life applications. Examples of these applications are the mean, the auto correlation, and cross-correlation functions (13), which are the main topics of this chapter, and the auto and cross-spectra. The auto-spectrum and the cross-spectrum are functions of frequency, being related to the auto and cross-correlation functions via the Fourier transform.

Knowledge of the basic theory of random processes is a pre-requisite for a correct interpretation of results obtained from the processing of random signals. Also, the algorithms used to compute estimates of parameters or functions associated with random processes are all based on the underlying random process theory.

All signals in this chapter will be assumed to be real and originating from a wide-sense stationary random process.

2.2. Basic Definitions

The mean or expected value of a *continuous-time* random processes $X(t)$ is defined as

$$\mu_x = E[X(t)], \quad (9)$$

where the time variable t is defined on a subset of the real numbers and $E[\cdot]$ is the expected value operation. The mean is a constant value because all random processes are assumed to be wide-sense stationary. The definition above is a mean calculated over all sample functions of the random process.

The autocorrelation of a *continuous-time* random processes $X(t)$ is defined as

$$R_{xx}(\tau) = E[X(t + \tau) \cdot X(t)], \quad (10)$$

where the time variables t and τ are defined on a subset of the real numbers. As was mentioned before, the nomenclature varies somewhat in the literature. The definition of autocorrelation given in Equation 10 is the one typically found in engineering books and papers. The value $R_{xx}(0) = E[X^2(t)]$ is sometimes called the average total power of the signal and its square root is the “root mean square” (RMS) value, employed frequently to characterize the “amplitude” of a biological random signal such as the EMG (1).

An equally important and related second moment is the autocovariance, defined for continuous time as

$$C_{xx}(\tau) = E[(X(t + \tau) - \mu_x) \cdot (X(t) - \mu_x)] = R_{xx}(\tau) - \mu_x^2 \quad (11)$$

The autocovariance at $\tau = 0$ is equal to the variance of the process and is sometimes called the average ac power (the average total power minus the square of the dc value):

$$C_{xx}(0) = \sigma_x^2 = E[(X(t) - \mu_x)^2] = E[X^2(t)] - \mu_x^2 \quad (12)$$

For a stationary random process, the mean, average total power and variance are constant values, independent of time.

The autocorrelation for a *discrete-time* random process is

$$R_{xx}(k) = E[X(n+k) \cdot X(n)] = E[X(n-k) \cdot X(n)], \quad (13)$$

where n and k are integer numbers. Any of the two expressions may be used, either with $X(n+k) \cdot X(n)$ or with $X(n-k) \cdot X(n)$. In what follows, preference shall be given to the first expression to keep consistency with the definition of cross-correlation to be given later.

For discrete-time processes, an analogous definition of autocovariance follows:

$$C_{xx}(k) = E[(X(n+k) - \mu_x) \cdot (X(n) - \mu_x)] = R_{xx}(k) - \mu_x^2, \quad (14)$$

where again $C_{xx}(0) = \sigma_x^2 = E[(X(k) - \mu_x)^2]$ is the constant variance of the stationary random process $X(k)$. Also $\mu_x = E[X(n)]$, a constant, is its expected value. In many applications, the interest is in studying the variations of a random processes about its mean, which is what the autocovariance represents. For example, to characterize the variability of the muscle force exerted by a human subject in a certain test, the interest is in quantifying how the force varies randomly around the mean value, and hence the autocovariance is more interesting than the autocorrelation.

The independent variable τ or k in Equations 10, 11, 13 or 14 will be called, interchangeably, *time shift*, *lag*, or *delay*.

It should be mentioned that some books and papers, mainly those on time series analysis (5,11), define the *autocorrelation function* of a random process as (for discrete time):

$$\rho_{xx}(k) = \frac{C_{xx}(k)}{\sigma_x^2} = \frac{C_{xx}(k)}{C_{xx}(0)} \quad (15)$$

(i.e., the autocovariance divided by the variance of the process). It should be noticed that $\rho_{xx}(0) = 1$. Actually, this definition was used in the Introduction of this chapter. The definition in Equation 15 differs from that in Equation 13 in two respects: The mean of the signal is subtracted and a normalization exists so that at $k=0$ the value is 1. To avoid confusion with the standard engineering nomenclature, the definition in Equation 15 may be called the *normalized autocovariance*, the *correlation coefficient function*, or still the *autocorrelation coefficient*. In the text that follows, preference is given to the term **normalized autocovariance**.

2.3. Basic Properties

From their definitions, the autocorrelation and autocovariance (normalized or not) are **even** functions of the time shift parameter, because, $X(t+\tau) \cdot X(t) = X(t-\tau) \cdot X(t)$ and $X(n+k) \cdot X(n) = X(n-k) \cdot X(n)$ for continuous- and discrete-time process, respectively. The property is indicated below only for the discrete-time case (for continuous-time, replace k by τ):

$$R_{xx}(k) = R_{xx}(-k), \quad (16)$$

and

$$C_{xx}(k) = C_{xx}(-k), \quad (17)$$

as well as for the normalized autocovariance:

$$\rho_{xx}(k) = \rho_{xx}(-k). \quad (18)$$

Three important inequalities may be derived (10) for both continuous-time and discrete-time random processes. Only the result for the discrete-time case is shown below (for continuous-time, replace k by τ):

$$|R_{xx}(k)| \leq R_{xx}(0) = \sigma_x^2 + \mu_x^2 \quad \forall k \in \mathbb{Z} \quad (19)$$

$$|C_{xx}(k)| \leq C_{xx}(0) = \sigma_x^2 \quad \forall k \in \mathbb{Z}, \quad (20)$$

and

$$|\rho_{xx}(k)| \leq 1 \quad \forall k \in \mathbb{Z}, \quad (21)$$

with $\rho_{xx}(0) = 1$.

These relations say that the maximum of either the autocorrelation or autocovariance occurs at lag 0.

Any discrete-time ergodic random process without a periodic component will satisfy the following limits:

$$\lim_{|k| \rightarrow \infty} R_{xx}(k) = \mu_x^2 \quad (22)$$

and

$$\lim_{|k| \rightarrow \infty} C_{xx}(k) = 0 \quad (23)$$

and similarly for continuous-time processes by changing k for τ . These relations mean that two random variables defined in $X(n)$ at two different times n_1 and n_1+k will tend to be uncorrelated as they are farther apart (i.e., the “memory” decays when the time interval k increases).

A final, more subtle, property of the autocorrelation is that it is positive semi-definite (10,14), expressed here only for the discrete-time case:

$$\sum_{i=1}^K \sum_{j=1}^K \alpha_i \alpha_j R_{xx}(k_i - k_j) \geq 0 \quad \text{for } \forall K \in \mathbb{Z}^+, \quad (24)$$

where $\alpha_1, \alpha_2, \dots, \alpha_K$ are arbitrary real numbers and $k_1, k_2, \dots, k_K \in \mathbb{Z}$ are any set of discrete-time points. This same result is valid for the autocovariance (normalized or not). This property means that not all functions that satisfy Equations 16 and 19 can be autocorrelation functions of some random process, they also have to be positive semi-definite.

2.4. Fourier Transform of the Autocorrelation

A very useful frequency-domain function related to the correlation/covariance functions is the power spectrum S_{xx} of the random process X (continuous or discrete-time),

defined as the Fourier transform of the autocorrelation function (10,15). For continuous time, we have

$$S_{xx}(j\omega) = \text{Fourier transform } [R_{xx}(\tau)], \quad (25)$$

where the angular frequency ω is in rad/s. The average power P_{xx} of the random process $X(t)$ is

$$P_{xx} = \frac{1}{2\pi} \int_{-\infty}^{\infty} S_{xx}(j\omega) d\omega = R_{xx}(0). \quad (26)$$

If the average power in a given frequency band $[\omega_1, \omega_2]$ is needed, it can be computed by

$$P_{xx} = \frac{1}{\pi} \int_{\omega_1}^{\omega_2} S_{xx}(j\omega) d\omega. \quad (27)$$

For discrete time

$$S_{xx}(e^{j\Omega}) = \text{Discrete Time Fourier transform } [R_{xx}(k)], \quad (28)$$

where Ω is the normalized angular frequency given in rad ($\Omega = \omega \cdot T$, where T is the sampling interval). In Equation 28 the power spectrum is periodic in Ω , with a period equal to 2π . The average power P_{xx} of the random process $X(n)$ is

$$P_{xx} = \frac{1}{2\pi} \int_{-\pi}^{\pi} S_{xx}(e^{j\Omega}) d\Omega = R_{xx}(0). \quad (29)$$

Other common names for the power spectrum are power spectral density and autospectrum. The power spectrum is a real **non-negative** and even function of frequency (10,15), which requires a positive semi-definite autocorrelation function (10), which means that not all functions that satisfy Equations 16 and 19 are valid autocorrelation functions because the corresponding power spectrum could have negative values for some frequency ranges, which is absurd.

The power spectrum should be used instead of the autocorrelation function in situations such as: (1) when the objective is to study the bandwidth occupied by a random signal, and (2) when one wants to discover if there are several periodic signals masked by noise (for a single periodic signal masked by noise the autocorrelation may be useful too).

2.5. White Noise

Continuous-time white noise is characterized by an autocovariance that is proportional to the Dirac impulse function:

$$C_{xx}(\tau) = \Psi \cdot \delta(\tau), \quad (30)$$

where Ψ is a positive constant and the Dirac impulse is defined as

$$\delta(\tau) = 0 \text{ for } \tau \neq 0 \quad (31)$$

$$\delta(\tau) = \infty \text{ for } \tau = 0, \quad (32)$$

and

$$\int_{-\infty}^{\infty} \delta(\tau) d\tau = 1. \quad (33)$$

The autocorrelation of continuous-time white noise is:

$$R_{xx}(\tau) = \Psi \cdot \delta(\tau) + \mu_x^2, \quad (34)$$

where μ_x is the mean of the process.

From Equations 12 and 30 it follows that the variance of the continuous-time white process is infinite (14), which indicates that it is not physically realizable. From Equation 30 we conclude that, for any time shift value τ , no matter how small, the correlation coefficient between any value in $X(t)$ and the value at $X(t + \tau)$ would be equal to zero, which is certainly impossible to satisfy in practice because of the finite risetimes of the outputs of any physical system. From Equation 25 it follows that the power spectrum of continuous-time white noise (with $\mu_x = 0$) has a constant value equal to Ψ at all frequencies. The name white noise comes from an extension of the concept of "white light," which similarly has constant power over the range of frequencies in the visible spectrum. White noise is non-realizable, because it would have to be generated by a system with infinite bandwidth.

Engineering texts circumvent the difficulties with the continuous-time white noise by defining a **band-limited white noise** (3,13). The corresponding power spectral density is constant up to very high frequencies (ω_c), and is zero elsewhere, which makes the variance finite. The maximum spectral frequency ω_c is taken to be much higher than the bandwidth of the system to which the noise is applied. Therefore, *in approximation*, the power spectrum is taken to be constant at all frequencies, the autocovariance is a Dirac delta function, and yet a finite variance is defined for the random process.

The utility of the concept of white noise develops when it is applied at the input of a finite bandwidth system, because the corresponding output is a well-defined random process with physical significance (see next section).

In discrete time, the white-noise process has an autocovariance proportional to the unit sample sequence:

$$C_{xx}(k) = \Psi \cdot \delta(k), \quad (35)$$

where Ψ is a finite positive real value, and

$$R_{xx}(k) = \Psi \cdot \delta(k) + \mu_x^2, \quad (36)$$

where $\delta(k) = 1$ for $k = 0$ and $\delta(k) = 0$ for $k \neq 0$. The discrete-time white noise has a finite variance, $\sigma^2 = \Psi$ in Equation 35, is realizable and it may be synthesized by taking a sequence of independent random numbers from an arbitrary probability distribution. Sequences that have independent samples with identical probability distributions are usually called **i.i.d.** (independent identically distributed). Computer-generated "random" (pseudo-random) sequences are usually very good approximations to a white-noise discrete-time random signal, being usually

of zero mean and unit variance for a normal or Gaussian distribution. To achieve desired values of Ψ in Equation 35 and μ_x^2 in Equation 36 one should multiply the (zero mean, unit variance) values of the computer-generated white sequence by $\sqrt{\Psi}$ and sum to each resulting value the constant value μ_x .

3. AUTOCORRELATION OF THE OUTPUT OF A LINEAR SYSTEM WITH RANDOM INPUT

In relation to the examples presented in the previous section, one may ask how may two random processes develop such differences in the autocorrelation as seen in Fig. 8. How may one autocorrelation be monotonically decreasing (for increasing positive τ) while the other exhibits oscillations? For this purpose, how the autocorrelation of a signal changes when it is passed through a time-invariant linear system will be examined.

If a continuous-time linear system has an impulse response $h(t)$ and a random process $X(t)$ with an autocorrelation $R_{xx}(\tau)$ is applied at its input the resultant output $y(t)$ will have an autocorrelation given by the following convolutions (10):

$$R_{yy}(\tau) = h(\tau) * h(-\tau) * R_{xx}(\tau). \quad (37)$$

Note that $h(\tau) * h(-\tau)$ may be viewed as an autocorrelation of $h(\tau)$ with itself and, hence, is an even function.

Taking the Fourier transform we conclude Equation 37, it is that the output power spectrum $S_{yy}(j\omega)$ is the absolute value squared of the frequency response function $H(j\omega)$ times the input power spectrum $S_{xx}(j\omega)$:

$$S_{yy}(j\omega) = |H(j\omega)|^2 \cdot S_{xx}(j\omega), \quad (38)$$

where $H(j\omega)$ is the Fourier transform of $h(t)$, $S_{xx}(j\omega)$ is the Fourier transform of $R_{xx}(\tau)$, and $S_{yy}(j\omega)$ is the Fourier transform of $R_{yy}(\tau)$.

The corresponding expressions for the autocorrelation and power spectrum for the output signal from a discrete-time system are (15)

$$R_{yy}(k) = h(k) * h(-k) * R_{xx}(k) \quad (39)$$

and

$$S_{yy}(e^{j\Omega}) = |H(e^{j\Omega})|^2 \cdot S_{xx}(e^{j\Omega}), \quad (40)$$

where $h(k)$ is the impulse (or unit sample) response of the system, $H(e^{j\Omega})$ is the frequency response function of the system, and $S_{xx}(e^{j\Omega})$ and $S_{yy}(e^{j\Omega})$ are the discrete-time Fourier transforms of $R_{xx}(k)$ and $R_{yy}(k)$, respectively. $S_{xx}(e^{j\Omega})$ and $S_{yy}(e^{j\Omega})$ are the input and output power spectra, respectively. As an example, suppose that $y(n) = [x(n) + x(n-1)]/2$ is the difference equation that defines a given system. This is an example of a finite impulse response (FIR) system (16) with impulse response equal to 0.5 for $n=0,1$ and 0 for other values of n . If the input is discrete-time white noise with unit variance, then from Equation 39 the output autocorrelation is a

triangular sequence centered at $k=0$, with amplitude 0.5 at $k=0$, amplitude 0.25 at $k=\pm 1$ and 0 elsewhere.

If two new random processes are defined as $U = X - \mu_x$ and $Q = Y - \mu_y$, it follows from the definitions of autocorrelation and autocovariance that $R_{uu} = C_{xx}$ and $R_{qq} = C_{yy}$. Therefore, applying Equation 37 or Equation 39 to a system with input U and output Y , similar expressions to Equations 37 and 39 are obtained relating the input and output autocovariances, shown below only for the discrete-time case:

$$C_{yy}(k) = h(k) * h(-k) * C_{xx}(k). \quad (41)$$

Furthermore, if $U = (X - \mu_x)/\sigma_x$ and $Q = (Y - \mu_y)/\sigma_y$ are new random processes, we have $R_{uu} = \rho_{xx}$ and $R_{qq} = \rho_{yy}$. From Equations 37 or 39 similar relations between the normalized autocovariance functions of the output and the input of the linear system are obtained shown below only for the discrete-time case (for the continuous-time, use τ instead of k):

$$\rho_{yy}(k) = h(k) * h(-k) * \rho_{xx}(k). \quad (42)$$

A monotonically decreasing autocorrelation or autocovariance may be obtained (e.g., Fig. 8a), when, for example, a white noise is applied at the input of a system that has a monotonically decreasing impulse response (e.g., of a first-order system or a second-order overdamped system). As an example, apply a zero-mean white noise to a system that has an impulse response equal to $e^{-\alpha t} \Pi(t)$, where $\Pi(t)$ is the Heaviside step function ($\Pi(t) = 1, t \geq 0$ and $\Pi(t) = 0, t < 0$). From Equation 37 this system's output random signal will have an autocorrelation that is

$$R_{yy}(\tau) = e^{-\alpha\tau} \Pi(\tau) * e^{\alpha\tau} \Pi(-\tau) * \delta(\tau) = \frac{e^{-\alpha|\tau|}}{8}. \quad (43)$$

This autocorrelation has its peak at $\tau=0$ and decays exponentially on both sides of the time shift axis, qualitatively following the shape seen in Fig. 8a. On the other hand, an oscillatory autocorrelation or autocovariance, as seen in Fig. 8b, may be obtained when the impulse response $h(t)$ is oscillatory, which may occur, for example, in a second-order underdamped system that would have an impulse response $h(t) = e^{-\alpha t} \cos(\omega_0 t) \Pi(t)$.

4. CROSS-CORRELATION BETWEEN TWO STATIONARY RANDOM PROCESSES

The presentation up to now was developed for both the continuous-time and discrete-time cases. From now on the expressions for the discrete-time case will only be presented.

4.1. Basic Definitions

Two signals are often recorded in an experiment from different parts of a system because an interest exists in analyzing if they are associated with each other. This association could develop from an anatomical coupling

(e.g., two interconnected sites in the brain (17)) or a physiological coupling between the two recorded signals (e.g., heart rate variability and respiration (18)). On the other hand, they could be independent because no anatomical and physiological link exists between the two signals. Start with a formal definition of independence of two random processes $X(n)$ and $Y(n)$.

Two random processes $X(n)$ and $Y(n)$ are **independent** when any set of random variables $\{X(n_1), X(n_2), \dots, X(n_N)\}$ taken from $X(n)$ is independent of another set of random variables $\{Y(n'_1), Y(n'_2), \dots, Y(n'_M)\}$ taken from $Y(n)$. Note that the time instants at which the random variables are defined from each random process are taken arbitrarily, as indicated by the set of integers n_i , $i = 1, 2, \dots, N$ and n'_i , $i = 1, 2, \dots, M$, with N and M being arbitrary positive integers. This definition may be useful when conceptually it is known beforehand that the two systems that generate $X(n)$ and $Y(n)$ are totally uncoupled. However, in practice, usually no *a priori* knowledge about the systems exists and the objective is to discover if they are coupled, which means that we want to study the possible association or coupling of the two systems based on their respective output signals $x(n)$ and $y(n)$. For this purpose, the definition of independence is unfeasible to test in practice and one has to rely on concepts of association based on second-order moments.

In the same way as the correlation coefficient quantifies the degree of linear association between two random variables, the cross-correlation and the cross-covariance quantify the degree of linear association between two random processes $X(n)$ and $Y(n)$. Their cross-correlation is defined as:

$$R_{xy}(k) = E[X(n+k) \cdot Y(n)], \quad (44)$$

and their cross-covariance as

$$\begin{aligned} C_{xy}(k) &= E[(X(n+k) - \mu_x) \cdot (Y(n) - \mu_y)] \\ &= R_{xy}(k) - \mu_x / \mu_y. \end{aligned} \quad (45)$$

It should be noted that some books or papers define the cross-correlation and cross-covariance as $R_{xy}(k) = E[X(n) \cdot Y(n+k)]$ and $C_{xy}(k) = E[(X(n) - \mu_x) \cdot (Y(n+k) - \mu_y)]$, which are time-reversed versions of the definitions above Equations 44 and 45. The distinction is clearly important when viewing a cross-correlation graph between two experimentally recorded signals $x(n)$ and $y(n)$, coming from random processes $X(n)$ and $Y(n)$, respectively. A peak at a positive time shift k according to one definition would appear at a negative time shift $-k$ in the alternative definition. Therefore, when using a signal processing software package or when reading a scientific text, the reader should always verify how the cross-correlation was defined. In Matlab (MathWorks, Inc.), a very popular software tool for signal processing, the commands `xcorr(x,y)` and `xcov(x,y)` use the same conventions as in Equations 44 and 45.

Similarly to what was said before for the autocorrelation definitions, texts on time series analysis define cross-

correlation as the normalized cross-covariance:

$$\rho_{xy}(k) = \frac{C_{xy}(k)}{\sigma_x \sigma_y} = E \left[\frac{(X(n+k) - \mu_x)}{\sigma_x} \cdot \frac{(Y(n) - \mu_y)}{\sigma_y} \right] \quad (46)$$

where σ_x and σ_y are the standard deviations of the two random processes $X(n)$ and $Y(n)$, respectively.

4.2. Basic Properties

As $X(n+k) \cdot Y(n)$ is in general different from $\pm[X(n-k) \cdot Y(n)]$, the cross-correlation and the cross-covariance have a more subtle symmetry property than the autocorrelation and autocovariance:

$$R_{xy}(k) = R_{yx}(-k) = E[Y(n-k) \cdot X(n)]. \quad (47)$$

It should be noted that the time argument in $R_{yx}(-k)$ is negative *and* the order xy is changed to yx . For the cross-covariance, a similar symmetry relation applies:

$$C_{xy}(k) = C_{yx}(-k) \quad (48)$$

and

$$\rho_{xy}(k) = \rho_{yx}(-k). \quad (49)$$

For the autocorrelation and autocovariance, the peak occurs at the origin, as given by the properties in Equations 19 and 20. However, for the crossed moments, the peak may occur at any time shift value, with the following inequalities being valid:

$$|R_{xy}(k)| \leq \sqrt{R_{xx}(0)R_{yy}(0)} \quad (50)$$

$$|C_{xy}(k)| \leq \sigma_x \sigma_y \quad (51)$$

$$|\rho_{xy}(k)| \leq 1. \quad (52)$$

Finally, two discrete-time ergodic random processes $X(n)$ and $Y(n)$ without a periodic component will satisfy the following:

$$\lim_{|k| \rightarrow \infty} R_{xy}(k) = \mu_x \mu_y \quad (53)$$

and

$$\lim_{|k| \rightarrow \infty} C_{xy}(k) = 0, \quad (54)$$

with similar results being valid for continuous time. These limit results mean that the “effect” of a random variable taken from process X on a random variable taken from process Y decreases as the two random variables are taken farther apart. In the limit they become uncorrelated.

4.3. Independent and Uncorrelated Random Processes

When two random processes $X(n)$ and $Y(n)$ are independent, their cross-covariance is always equal to 0 for any time shift, and hence the autocorrelation is equal to the

product of the two means:

$$C_{xy}(k) = 0 \quad \forall k \quad (55)$$

and

$$R_{xy}(k) = \mu_x \mu_y \quad \forall k. \quad (56)$$

Two random processes that satisfy the two expressions above are called **uncorrelated** (10,15) whether they are independent or not. In practical applications, one usually has no way to test for the independence of two random processes, but it is feasible to test if they are correlated or not. Note that the term uncorrelated may be misleading because the cross-correlation itself is not zero (unless one of the mean values is zero, when the processes are called orthogonal), but the cross-covariance is.

Two independent random processes are always uncorrelated, but the reverse is not necessarily true. Two random processes may be uncorrelated (i.e., have zero cross-covariance) but still be statistically dependent, because other probabilistic quantifiers (e.g., higher-order central moments (third, fourth, etc.)), will not be zero. However, for Gaussian (or normal) random processes, a one-to-one correspondence exists between independence and uncorrelatedness (10). This special property is valid for Gaussian random processes because they are specified entirely by the mean and second moments (autocovariance function for a single random process and cross-covariance for the joint distribution of two random processes).

4.4. Simple Model of Delayed and Amplitude-Scaled Random Signal

In some biomedical applications, the objective is to estimate the delay between two random signals. For example, the signals may be the arterial pressure and cerebral blood flow velocity for the study of cerebral autoregulation, or EMGs from different muscles for the study of tremor (19). The simplest model relating the two signals is $y(n) = \alpha \cdot x(n - L)$ where $\alpha \in \mathbb{R}$, and $L \in \mathbb{Z}$, which means that y is an amplitude-scaled and delayed (for $L > 0$) version of x . From Equation 46 $\rho_{xy}(k)$ will have either a maximum peak equal to 1 (if $\alpha > 0$) or a trough equal to -1 (if $\alpha < 0$) located at time shift $k = -L$. Hence, peak location in the time axis of the cross-covariance (or cross-correlation) indicates the delay value.

A slightly more realistic model in practical applications assumes that one signal is a delayed and scaled version of another but with an extraneous additive noise:

$$y(n) = \alpha x(n - L) + w(n). \quad (57)$$

In this model, $w(n)$ is a random signal caused by external interference noise or intrinsic biological noise that cannot be controlled or measured. Usually, one can assume that $x(n)$ and $w(n)$ are signals from uncorrelated or independent random processes $X(n)$ and $W(n)$. Let us determine $\rho_{xy}(k)$ assuming access to $X(n)$ and $Y(n)$. From

Equation 45,

$$C_{xy}(k) = E[(X(n+k) - \mu_x) \cdot (\alpha(X(n-L) - \mu_x) + E[(X(n+k) - \mu_x) \cdot (W(n) - \mu_w)], \quad (58)$$

but the last term is zero because $X(n)$ and $W(n)$ are uncorrelated. Hence,

$$C_{xy}(k) = \alpha C_{xx}(k+L). \quad (59)$$

As $C_{xx}(k+L)$ attains its peak value when $k = -L$ (i.e., when the argument $(k+L)$ is zero) this provides a very practical method to estimate the delay between two random signals: Find the time-shift value where the cross-covariance has a clear peak. If an additional objective is to estimate the amplitude-scaling parameter α , it may be achieved by dividing the peak value of the cross-covariance by σ_x^2 (see Equation 59).

Additionally, from Equations 46 and 59:

$$\rho_{xy}(k) = \frac{\alpha C_{xx}(k+L)}{\sigma_x \sigma_y}. \quad (60)$$

To find σ_y , we remember that $C_{yy}(0) = \sigma_y^2$ from Equation 12. From Equation 57 we have

$$C_{yy}(0) = E[\alpha(X(n-L) - \mu_x) \cdot \alpha(X(n-L) - \mu_x) + E[(W(n) - \mu_w) \cdot (W(n) - \mu_w)], \quad (61)$$

where again we used the fact that $X(n)$ and $W(n)$ are uncorrelated. Therefore,

$$C_{yy}(0) = \alpha^2 C_{xx}(0) + C_{ww}(0) = \alpha^2 \sigma_x^2 + \sigma_w^2, \quad (62)$$

and therefore

$$\rho_{xy}(k) = \frac{\alpha C_{xx}(k+L)}{\sigma_x \sqrt{\alpha^2 \sigma_x^2 + \sigma_w^2}}. \quad (63)$$

When $k = -L$, $C_{xx}(k+L)$ will reach its peak value equal to σ_x^2 , which means that $\rho_{xy}(k)$ will have a peak at $k = -L$ equal to

$$\rho_{xy}(-L) = \frac{\alpha \sigma_x}{\sqrt{\alpha^2 \sigma_x^2 + \sigma_w^2}} = \frac{\alpha}{|\alpha|} \frac{1}{\sqrt{1 + \frac{\sigma_w^2}{\alpha^2 \sigma_x^2}}}. \quad (64)$$

Equation 64 is consistent with the case in which no noise $W(n)$ ($\sigma_w^2 = 0$) exists because the peak in $\rho_{xy}(k)$ will equal $+1$ or -1 , if α is positive or negative, respectively. From Equation 64, when the noise variance σ_w^2 increases, the peak in $\rho_{xy}(k)$ will decrease in absolute value, but will still occur at time shift $k = -L$. Within the context of this example, another way of interpreting the peak value in the normalized cross-covariance $\rho_{xy}(k)$ is by asking what fraction $\Gamma_{x \rightarrow y}$ of Y is because of X in Equation 57, meaning the ratio of the standard deviation of the term $\alpha X(n-L)$ to

the total standard deviation in $Y(n)$:

$$\Gamma_{x \rightarrow y} = \frac{|\alpha| \sigma_x}{\sigma_y}. \quad (65)$$

From Equation 60

$$\rho_{xy}(-L) = \frac{\alpha \sigma_x}{\sigma_y}, \quad (66)$$

and from Equations 65 and 66 it is concluded that the peak size (absolute peak value) in the normalized cross-covariance indicates the fraction of Y because of the random process X :

$$|\rho_{xy}(-L)| = \Gamma_{x \rightarrow y}, \quad (67)$$

which means that when the deleterious effect of the noise W increases (i.e., when its variance increases), the contribution of X to Y decreases, which by Equations 64 and 67 is reflected in a decrease in the size of the peak in $\rho_{xy}(k)$.

The derivation given above showed that the peak in the cross-covariance gives an estimate of the delay between two random signals linked by model Equation 57 and that the amplitude-scaling parameter may also be estimated.

5. CROSS-CORRELATION BETWEEN THE RANDOM INPUT AND OUTPUT OF A LINEAR SYSTEM

If a discrete-time linear system has an impulse response $h(n)$ and a random process $X(n)$ with an autocorrelation $R_{xx}(k)$ is applied at its input, the cross-correlation between the input and the output processes will be (15):

$$R_{xy}(k) = h(-k) * R_{xx}(k), \quad (68)$$

and if the input is white the result becomes $R_{xy}(k) = h(-k)$.

The Fourier Transform of $R_{xy}(k)$ is the cross power spectrum $S_{xy}(e^{j\Omega})$ and from Equation 68 and the properties of the Fourier transform an important relation is found:

$$S_{xy}(e^{j\Omega}) = H^*(e^{j\Omega}) \cdot S_{xx}(e^{j\Omega}). \quad (69)$$

The results in Equations 68 and 69 are frequently used in biological system identification, when the system linearity may hold. Therefore, if the input signal is white, one may obtain an estimate of $R_{xy}(k)$ from the measurement of the input and output signals. Following Equation 68 the only thing to do is invert the time axis (what is negative becomes positive, and vice-versa) to get an estimate of the impulse response $h(k)$. In the case of nonwhite input, it is better to use Equation 69 to obtain an estimate of $H(e^{j\Omega})$ by dividing $S_{xy}^*(e^{j\Omega})$ by $S_{xx}(e^{j\Omega})$ (4,21) and then obtain the estimated impulse response by inverse Fourier transform.

5.1. Common Input

Let us assume that signals x and y , recorded from two points in a given biological system, are used to compute an estimate of $R_{xy}(k)$. Let us also assume that a

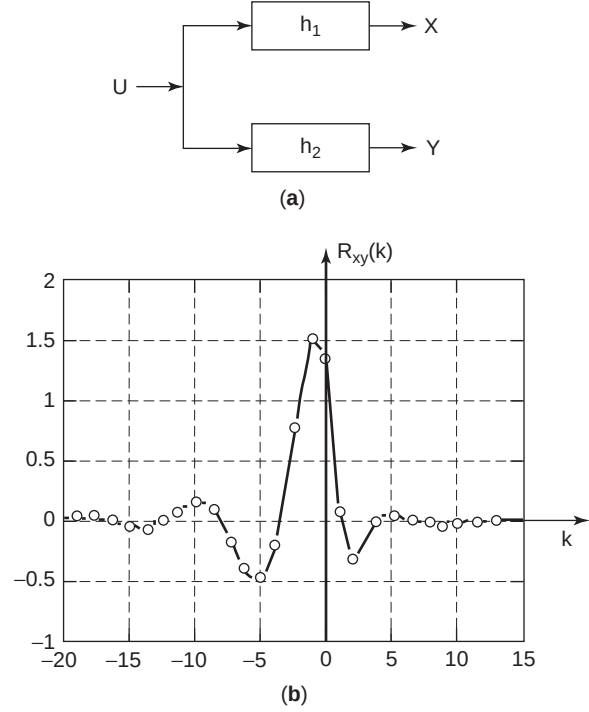


Figure 9. (a) Schematic of a common input U to two linear systems with impulse responses h_1 and h_2 , the first generating the output X and the second the output Y . (b) Cross-correlation between the two outputs X and Y of a computer simulation of the schematic in (a). The cross-correlation samples were joined by straight lines to improve the visualization. Without additional knowledge, this cross-correlation could have come from a system with input x and output y or from two systems with a common input, as was the case here.

clear peak appears around $k = -10$. One interpretation would be that signal x “caused” signal y , with an average delay of 10, because signal x passed through some equivalent (yet unknown) sub-system to generate signal y , and $R_{xy}(k)$ would follow from Equation 68. However, in biological systems, one notable example being the nervous system, one should never discard the possibility of a common source exerting effects on two subsystems whose outputs are the measured signals x and y . This situation is depicted in Fig. 9a, where the common input is a random process U that is applied at the inputs of two subsystems, with impulse responses h_1 and h_2 . In many cases, only the outputs of each of the subsystems X and Y (Fig. 9a) can be recorded, and all the analyses are based on their relationships. Working in discrete time, the cross-correlation between the two output random processes may be written as a function of the two impulse responses as

$$R_{xy}(k) = h_1(k) * h_2(-k) * R_{uu}(k), \quad (70)$$

which is simplified if the common input is white:

$$R_{xy}(k) = h_1(k) * h_2(-k). \quad (71)$$

Figure 9b shows an example of a cross-correlation of the outputs of two linear systems that had the same

random signal applied to their inputs. Without prior information (e.g., on the possible existence of a common input) or additional knowledge (e.g., of the impulse response of one of the systems and $R_{uu}(k)$) it would certainly be difficult to interpret such a cross-correlation.

An example from the biomedical field shall illustrate a case where the existence of a common random input was hypothesized based on empirically obtained cross-covariances. The experiments consisted of evoking spinal cord reflexes bilaterally and simultaneously in the legs of each subject, as depicted in Fig. 10. A single electrical stimulus, applied to the right or left leg, would fire an action potential in some of the sensory nerve fibers situated under the stimulating electrodes (st in Fig. 10). The action potentials would travel to the spinal cord (indicated by the upward arrows) and activate a set of motoneurons (represented by circles in a box). The axons of these motoneurons would conduct action potentials to a leg muscle, indicated by downward arrows in Fig. 10. The recorded waveform from the muscle (shown either to the left or right of Fig. 10) is the so-called H reflex. Its peak-to-peak amplitude is of interest, being indicated as x for the left leg reflex waveform and y for the right leg waveform in Fig. 10. The experimental setup included two stimulators (indicated as **st** in Fig. 10) that applied simultaneous trains of 500 rectangular electrical stimuli at 1 per second to the two legs. If the stimulus pulses in the trains are numbered as $n = 0, n = 1, \dots, n = N$ (the authors used $N = 500$), the respective sequences of H reflex amplitudes recorded on each side will be $x(0), x(1), \dots, x(N)$ and $y(0), y(1), \dots, y(N)$, as depicted in the inset at the lower right side of Fig. 10. The two sequences of reflex amplitudes were analyzed by the cross-covariance function (22).

Each reflex peak-to-peak amplitude depends on the upcoming sensory activity discharged by each stimulus and also on random inputs from the spinal cord that act on the motoneurons. In Fig. 10, a “U?” indicates a hypothesized common input random signal that would modulate synchronously the reflexes from the two sides of the spinal cord. The peak-to-peak amplitudes of the right- and left-side H reflexes to a train of 500 bilateral stimuli were measured in real time and stored as discrete-time signals x and y . Initial experiments had shown that a unilateral stimulus only affected the same side of the spinal cord. This finding meant that any statistically significant peak in the cross-covariance of x and y could be attributed to a common input. The first 10 reflex amplitudes in each signal were discarded to avoid the transient (nonstationarity) that occurs at the start of the stimulation.

Data from a subject are shown in Fig. 11, the first 51 H-reflex amplitude values shown in Fig. 11a for the right leg, and the corresponding simultaneous H reflex amplitudes in Fig. 11b for the left leg (R. A. Mezzarane and A. F. Kohn, unpublished data). In both, a horizontal line shows the respective mean value computed from all the 490 samples of each signal. A simple visual analysis of the data probably tells us close to nothing about how the reflex amplitudes vary and if the two sides fluctuate together to some degree. Such quantitative questions may be answered by the autocovariance and cross-covariance functions. The right- and left-leg H-reflex amplitude normalized autocovariances, computed according to Equation 15, are shown in Fig. 12a and 12b, respectively. The value at time shift 0 is 1, as expected from the normalization. Both autocovariances show that for small time shifts, some degree of correlation exists between the samples because

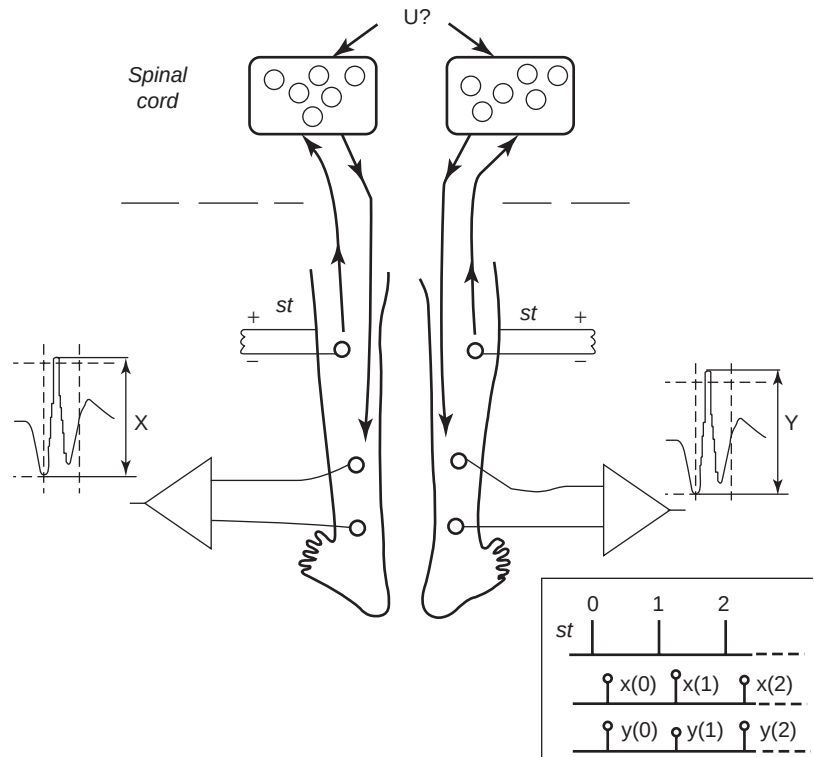


Figure 10. Experimental setup to elicit and record bilateral reflexes in human legs. Each stimulus applied at the points marked **st** causes upward-propagating action potentials that activate a certain number of motoneurons (circles in a box). A hypothesized common random input to both sides is indicated by “U?”. The reflex responses travel down the nerves located on each leg and cause each calf muscle to fire a compound action potential, which is the so-called H reflex. The amplitudes x and y of the reflexes on each side are measured for each stimulus. Actually, a bilateral train of 500 stimuli is applied and the corresponding reflex amplitudes $x(n)$ and $y(n)$ are measured, for $n = 0, 1, \dots, 499$, as sketched in the inset in the lower corner. Later, the two sets of reflex amplitudes are analyzed by auto and cross-covariance.

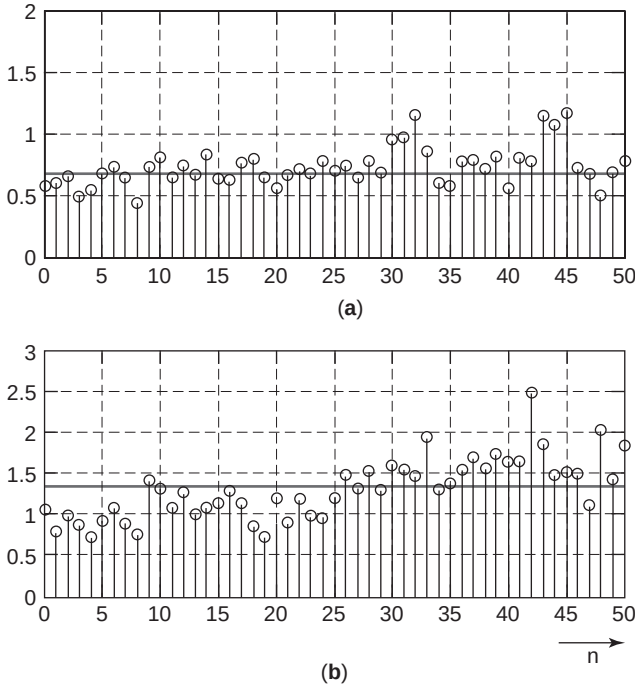


Figure 11. The first 51 samples of two time series $y(n)$ and $x(n)$ representing the simultaneously-measured H-reflex amplitudes from the right (a) and left (b) legs in a subject. The horizontal lines indicate the mean values of each complete series. The stimulation rate was 1 Hz. Ordinates are in mV. (This figure is available in full color at <http://www.mrw.interscience.wiley.com/ebe>.)

the autocovariance values are above the upper level of a 95% confidence interval shown in Fig. 12 (see subsection Hypothesis testing). This fact supports the hypothesis that randomly varying neural inputs exists in the spinal cord that modulate the excitability of the reflex circuits. As the autocovariance in Fig. 12b decays much slower than that in Fig. 12a, it suggests that the two sides receive different randomly varying inputs. The normalized cross-covariance, computed according to Equation 46, is seen in Fig. 13. Many cross-covariance samples around $k=0$ are above the upper level of a 95% confidence interval (see subsection Hypothesis testing), which suggests a considerable degree of correlation between the reflex amplitudes recorded from both legs of the subject. The decay on both sides of the cross-covariance in Fig. 13 could be because of the autocovariances of the two signals (see Fig. 12), but this issue will only be treated later. Results such as those in Fig. 13, also found in other subjects (22), suggested the existence of common inputs acting on sets of motoneurons at both sides of the spinal cord. However, as the peak of the cross-covariance was lower than 1 and as the autocovariances were different bilaterally, it can be stated that each side receives common sources to both sides plus random inputs that are uncorrelated with the other side's inputs. Experiments in cats are being pursued, with the help of the cross-covariance analysis, to unravel the neural sources of the random inputs found to act bilaterally in the spinal cord (23).

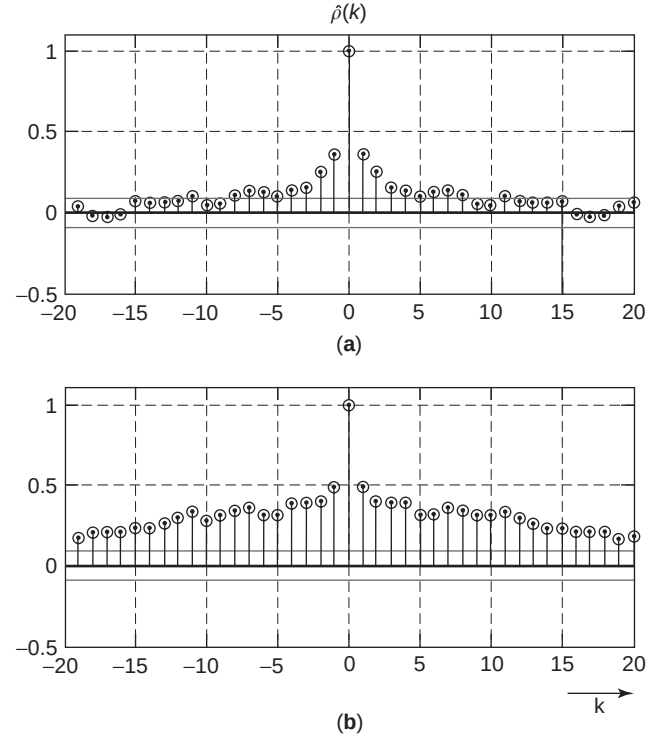


Figure 12. The respective autocovariances of the signals shown in Fig. 11. The autocovariance in (a) decays faster to the confidence interval than that in (b). The two horizontal lines represent a 95% confidence interval. (This figure is available in full color at <http://www.mrw.interscience.wiley.com/ebe>.)

6. ESTIMATION AND HYPOTHESIS TESTING FOR AUTOCORRELATION AND CROSS-CORRELATION

This section will focus solely on discrete-time random signals because, today, the signal processing techniques are almost all realized in discrete-time in very fast computers.

If *a priori* knowledge about the stationarity of the random process exists a test should be applied for stationarity (4). If a trend of no physiological interest is discovered in the data, it may be removed by linear or nonlinear regression. After a stationary signal is obtained, then the problem is to estimate first and second moments as presented in the previous sections.

Let us assume a stationary signal $x(n)$ is known for $n = 0, 1, \dots, N - 1$. Any estimate $\Theta(x(n))$ based on this signal would have to present some desirable properties, derived from its corresponding estimator $\Theta(X(n))$, such as unbiasedness and consistency (15). Note that an estimate is a particular value (or a set of values) of the estimator when a given sample function $x(n)$ is used instead of the whole process $X(n)$ (which is what happens in practice). For an **unbiased estimator** $\Theta(X(n))$, its expected value is equal to the actual value being estimated θ :

$$E[\Theta(X(n))] = \theta. \quad (72)$$

For example, if we want to estimate the mean μ_x of the process $X(n)$ using $\bar{x}$ Equation 2, the unbiasedness means that the expected value of the estimator has to equal μ_x . It

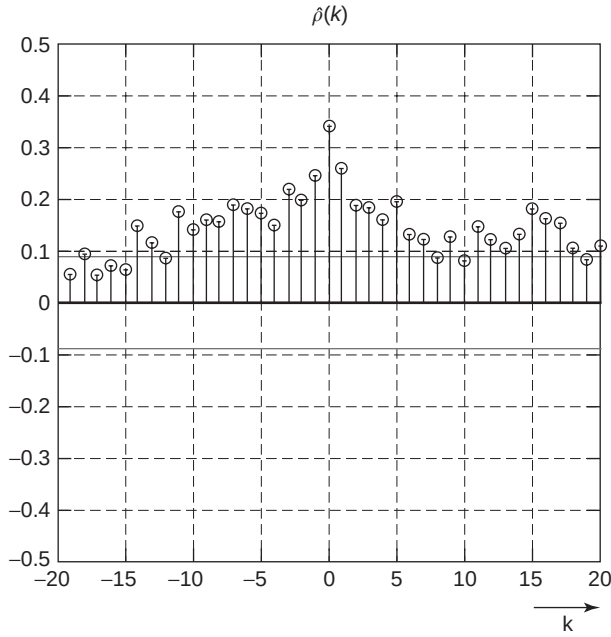


Figure 13. Cross-covariance between the two signals with 490 bilateral reflexes shown partially in Fig 11. Only the time shifts near the origin are shown here because they are the most reliable. The two horizontal lines represent a 95% confidence interval. (This figure is available in full color at <http://www.mrw.interscience.wiley.com/ebe>.)

is easy to show that the finite average $[X(0) + X(1) + \dots + X(N-1)]/N$ of a stationary process is an unbiased estimator of the expected value of the process:

$$\begin{aligned} E[\Theta(X(n))] &= E\left[\frac{1}{N} \sum_{n=0}^{N-1} X(n)\right] = \frac{1}{N} \sum_{n=0}^{N-1} E[X(n)] \\ &= \frac{1}{N} \sum_{n=0}^{N-1} \mu_x = \mu_x. \end{aligned} \quad (73)$$

If a given estimator is unbiased, it still does not assure its usefulness, because if its variance is high it means that for a given signal $x(n)$ - a single sample function of the process $X(n)$ - the value of $\Theta(x(n))$ may be quite far from the actual value θ . Therefore, a desirable property of an unbiased estimator is that its variance tends to 0 as the number of samples N tends to ∞ , a property called **consistency**:

$$\sigma_{\Theta}^2 = E[(\Theta(X(n)) - \theta)^2] \xrightarrow{N \rightarrow \infty} 0. \quad (74)$$

Returning to the unbiased mean estimator Equation 2, let us check its consistency:

$$\begin{aligned} \sigma_{\Theta}^2 &= E[(\Theta(X(n)) - \mu_x)^2] \\ &= E\left[\left(\frac{1}{N} \sum_{n=0}^{N-1} (X(n) - \mu_x)\right)^2\right], \end{aligned} \quad (75)$$

which, after some manipulations (15), gives:

$$\sigma_{\Theta}^2 = \frac{1}{N} \sum_{k=-N+1}^{N-1} \left(1 - \frac{|k|}{N}\right) \cdot C_{xx}(k). \quad (76)$$

From Equation 23 it is concluded that the mean estimator is consistent. It is interesting to observe that the variance of the mean estimator depends on the autocovariance of the signal and not only on the number of samples, which is because of the statistical dependence between the samples of the signal, quite contrary to what happens in conventional statistics, where the samples are i.i.d..

6.1. Estimation of Autocorrelation and Autocovariance

Which estimator should be used for the autocorrelation of a process $X(n)$? Two slightly different estimators will be presented and their properties verified. Both are much used in practice and are part of the options of the Matlab command `xcorr`.

The first estimator will be called the *unbiased autocorrelation estimator* and defined as

$$\begin{aligned} \hat{R}_{xx,u}(k) &= \frac{1}{N - |k|} \sum_{n=0}^{N-1-|k|} X(n + |k|)X(n); \\ |k| &\leq N - 1. \end{aligned} \quad (77)$$

The basic operations are sample-to-sample multiplications of two time-shifted versions of the process and arithmetic average of the resulting samples, which seems to be a reasonable approximation to Equation 13. To confirm that this estimator is indeed unbiased, its expected value shall be determined:

$$\begin{aligned} E[\hat{R}_{xx,u}(k)] &= \frac{1}{N - |k|} \sum_{n=0}^{N-1-|k|} E[X(n + |k|)X(n)]; \\ |k| &\leq N - 1, \end{aligned} \quad (78)$$

and from Equations 13 and 16 it can be concluded that the estimator is indeed unbiased (i.e., $E[\hat{R}_{xx,u}(k)] = R_{xx}(k)$). It is more difficult to verify consistency, and the reader is referred to Therrien (15). For a Gaussian process $X(n)$, it can be shown that the variance indeed tends to 0 for $N \rightarrow \infty$, which is a good approximation for more general random processes. Therefore, Equation 77 gives us a consistent estimator of the autocorrelation. When applying Equation 77 in practice, the available signal $x(n)$, $n = 0, 1, \dots, N-1$, is used instead of $X(n)$ giving a sequence of $2N - 1$ values $\hat{R}_{xx,u}(k)$.

The *unbiased autocovariance estimator* is similarly defined as

$$\begin{aligned} \hat{C}_{xx,u}(k) &= \frac{1}{N - |k|} \sum_{n=0}^{N-1-|k|} (X(n + |k|) - \bar{x}) \cdot (X(n) - \bar{x}); \\ |k| &\leq N - 1, \end{aligned} \quad (79)$$

where $\bar{x}$ is given by Equation 2. Similarly, $\hat{C}_{xx,u}(k)$ is a consistent estimator of the autocovariance.

Nevertheless, one of the problems with these two estimators defined in Equations 77 and 79 is that they may not obey Equations 19, 20 or the positive semi-definite property (14,15).

Therefore, other estimators are defined, such as the biased autocorrelation and autocovariance estimators $\hat{R}_{xx,b}(k)$ and $\hat{C}_{xx,b}(k)$

$$\hat{R}_{xx,b}(k) = \frac{1}{N} \sum_{n=0}^{N-1-|k|} (X(n+|k|)X(n)); \quad |k| \leq N-1 \quad (80)$$

and

$$\hat{C}_{xx,b}(k) = \frac{1}{N} \sum_{n=0}^{N-1-|k|} (X(n+|k|) - \bar{x}) \cdot (X(n) - \bar{x}); \quad |k| \leq N-1. \quad (81)$$

The only difference between Equations 77 and 80, or Equations 79 and 81 is the denominator, which means that $\hat{R}_{xx,b}(k)$ and $\hat{C}_{xx,b}(k)$ are biased estimators. The zero padding used in the time-shifted versions of signal $x(n)$ in Fig. 4b–d may be interpreted as the cause of the bias of the estimators in Equations 80 and 81 as the zero samples do not contribute to the computations in the sum but are counted in the denominator N . The bias is defined as the expected value of an estimator minus the actual value, and for these two biased estimators the expressions are:

$$\text{Bias}[\hat{R}_{xx,b}(k)] = -\frac{|k|}{N} R_{xx}(k); \quad |k| \leq N-1 \quad (82)$$

and

$$\text{Bias}[\hat{C}_{xx,b}(k)] = -\frac{|k|}{N} C_{xx}(k); \quad |k| \leq N-1. \quad (83)$$

Nevertheless, when $N \rightarrow \infty$, the expected values of $\hat{R}_{xx,b}(k)$ and $\hat{C}_{xx,b}(k)$ equal the theoretical autocorrelation and autocovariance (i.e., the estimators are asymptotically unbiased). Their consistency follows from the consistency of the unbiased estimators. Besides these desirable properties, it can be shown that these two estimators obey Equations 19 and 20 and are always positive semi-definite (14,15).

The question is then, which of the two estimators, the unbiased or the biased, should be chosen? Usually, that choice is less critical than the choice of the maximum value of k used in the computation of the autocorrelation or autocovariance. Indeed, for large k the errors when using Equations 80 and 81 are large because of the bias, as seen from Equations 82 and 83, whereas the errors associated with Equations 77 and 79 are large because of variance of the estimates, due to the small value of the denominator $N - |k|$ in the formulas. Therefore, the user should try to limit the values of $|k|$ to as low a value as appropriate for the application, for example, $|k| \leq N/10$ or $|k| \leq N/4$, because near $k=0$ the estimates have the

highest reliability. On the other hand, if the estimated autocorrelation or autocovariance will be Fourier-transformed to provide an estimate of a power spectrum, then the biased autocorrelation or autocovariance estimators should be chosen to assure a non-negative power spectrum.

For completeness purposes, the normalized autocovariance estimators are given below. The unbiased estimator is:

$$\hat{\rho}_{xx,u}(k) = \frac{\hat{C}_{xx,u}(k)}{\hat{C}_{xx,u}(0)}; \quad |k| \leq N-1, \quad (84)$$

where $\hat{\rho}_{xx,u}(0) = 1$. An alternative expression should be employed if the user computes first $\hat{C}_{xx,u}(k)$ and then divides it by the variance estimate $\hat{\sigma}_x^2$, which is usually the unbiased estimate in most computer packages:

$$\hat{\rho}_{xx,u}(k) = \left(\frac{N}{N-1} \right) \frac{\hat{C}_{xx,u}(k)}{\hat{\sigma}_x^2}; \quad |k| \leq N-1. \quad (85)$$

Finally, the biased normalized autocovariance estimator is:

$$\hat{\rho}_{xx,b}(k) = \frac{\hat{C}_{xx,b}(k)}{\hat{C}_{xx,b}(0)}; \quad |k| \leq N-1. \quad (86)$$

Assuming the available signal $x(n)$ has N samples, all the autocorrelation or autocovariance estimators presented above will have $2N-1$ samples.

An example of the use of Equation 84 was already presented in the subsection “Common input” above. The two signals, right- and left-leg H-reflex amplitudes, had $N=490$ samples, but the unbiased normalized autocovariances shown in Fig. 12 were only shown for $|k| \leq 20$. However, the values at the two extremes $k = \pm 489$ (not shown) surpassed 1, which is meaningless. As mentioned before, the values for time shifts far away from the origin should not be used for interpretation purposes. A more involved way of computing the normalized autocovariance that assures values within the range $[-1, 1]$ for all k can be found on page 331 of Priestley’s text (14).

From a computational point of view, it is usually not recommended to compute directly any of the estimators given in this subsection, except for small N . As the autocorrelation or autocovariance estimators are basically the discrete-time convolutions of $x(n)$ with $x(-n)$, the computations can be done very efficiently in the frequency domain using an FFT algorithm (16), with the usual care of zero padding to convert a circular convolution to a linear convolution. Of course, for the user of scientific packages such as Matlab, this problem does not exist because each command, such as `xcorr` and `xcov`, is already implemented in a computationally efficient way.

6.2. Estimation of Cross-Correlation and Cross-Covariance

The concepts involved in the estimation of the cross-correlation or the cross-covariance are analogous to those of the autocorrelation and autocovariance presented in the

previous section. An important difference, however, is that the cross-correlation and cross-covariance do not have even symmetry. The expressions for the unbiased and biased estimators of cross-correlation given below are in a form appropriate for signals $x(n)$ and $y(n)$, both defined for $n = 0, 1, \dots, N-1$. The unbiased estimator is

$$\hat{R}_{xy,u}(k) = \begin{cases} \frac{1}{N-|k|} \sum_{n=0}^{N-1-|k|} X(n+k)Y(n); & 0 \leq k \leq N-1 \\ \frac{1}{N-|k|} \sum_{n=0}^{N-1-|k|} X(n)Y(n+|k|); & -(N-1) \leq k < 0, \end{cases} \quad (87)$$

and the biased estimator is

$$\hat{R}_{xy,b}(k) = \begin{cases} \frac{1}{N} \sum_{n=0}^{N-1-|k|} X(n+k)Y(n); & 0 \leq k \leq N-1 \\ \frac{1}{N} \sum_{n=0}^{N-1-|k|} X(n)Y(n+|k|); & -(N-1) \leq k < 0. \end{cases} \quad (88)$$

In the definitions above, $\hat{R}_{xy,u}(k)$ and $\hat{R}_{xy,b}(k)$ will have $2N-1$ samples. Both expressions could be made more general by having $x(n)$ with N samples and $y(n)$ with M samples ($M \neq N$), resulting in a cross-correlation with $N+M-1$ samples, with k in the range $[-N, M]$. However, the reader should be cautioned that the outputs of algorithms that were initially designed for equal-sized signals may require some care. A suggestion is to check the outputs with simple artificially generated signals for which the cross-correlation is easily determined or known.

The corresponding definitions of the unbiased and biased estimators of the cross-covariance, for two signals with N samples each, are

$$\hat{C}_{xy,u}(k) = \begin{cases} \frac{1}{N-|k|} \sum_{n=0}^{N-1-|k|} (X(n+k) - \bar{x})(Y(n) - \bar{y}); & 0 \leq k \leq N-1 \\ \frac{1}{N-|k|} \sum_{n=0}^{N-1-|k|} (X(n) - \bar{x})(Y(n+|k|) - \bar{y}); & -(N-1) \leq k < 0 \end{cases} \quad (89)$$

and

$$\hat{C}_{xy,b}(k) = \begin{cases} \frac{1}{N} \sum_{n=0}^{N-1-|k|} (X(n+k) - \bar{x})(Y(n) - \bar{y}); & 0 \leq k \leq N-1 \\ \frac{1}{N} \sum_{n=0}^{N-1-|k|} (X(n) - \bar{x})(Y(n+|k|) - \bar{y}); & -(N-1) \leq k < 0 \end{cases} \quad (90)$$

where $\bar{x}$ and $\bar{y}$ are the means of the two processes $X(n)$ and $Y(n)$. The normalized cross-covariance estimators are given below with the $\odot$ meaning either b or u (i.e., biased or unbiased versions):

$$\hat{\rho}_{xy,\odot}(k) = \frac{\hat{C}_{xy,\odot}(k)}{\sqrt{\hat{C}_{xy,\odot}(0) \cdot \hat{C}_{yy,\odot}(0)}}. \quad (91)$$

If the user first computes $\hat{C}_{xy,u}(k)$ and then divides this result by the product of the standard deviation estimates, the factor $N/(N-1)$ should be used in case the computer

package calculates the unbiased variance (standard deviation) estimates

$$\hat{\rho}_{xy,u}(k) = \left(\frac{N}{N-1} \right) \frac{\hat{C}_{xy,u}(k)}{\hat{\sigma}_x \hat{\sigma}_y}, \quad (92)$$

which may be useful, for example, when using the Matlab command `xcov`, because none of its options include the direct computation of an unbiased and normalized cross-covariance between two signals x and y . The Equation 92 in Matlab would be written as:

$$\hat{\rho}_{xy,u}(k) = \text{xcov}(x, y, \text{'unbiased'}) * N / ((N-1) * \text{std}(x) * \text{std}(y)). \quad (93)$$

Note that Matlab will give $\hat{\rho}_{xy,u}(k)$ (or other alternatives of cross-covariance or cross-correlation) as a vector with $2N-1$ elements (e.g., with 979 points if $N=490$) when the two signals x and y have N samples. For plotting the cross-covariance or cross-correlation computed by a software package, the user may create the time shift axis k by $k = -(N-1) : (N-1)$ to assure the correct position of the zero time shift value in the graph. In Matlab, the values of k are computed automatically if the user makes `[C,k]=xcov(x,y,'unbiased')`. If the sampling frequency f_s is known and the time shift axis should reflect continuous time calibration, then vector k should be divided by f_s .

The Equation 93 was employed to generate Fig. 13, which is the normalized unbiased estimate of the cross-covariance between the reflex amplitudes of the right and left legs of a subject, as described before.

Similarly to what was said before for the autocorrelation and autocovariance, in practical applications only the cross-covariance (or cross-correlation) samples nearer to the origin $k=0$ should be used for interpretation purposes. This is because the increase in variance or bias in the unbiased or biased cross-covariance or cross-correlation estimates at large values of $|k|$.

6.3. Hypothesis Testing

One basic question when analyzing a given signal $x(n)$ is if it is white (i.e., if its autocovariance satisfies Equation 35). The problem in *practical applications* is that the autocovariance computed from a sample function of a white-noise process will be nonzero for $k \neq 0$, contrary to what would be expected for a white signal. This problem is typical in statistics, and it may be shown that a 95% confidence interval for testing the whiteness of a signal with N samples is given by the bounds $\pm 1.96/\sqrt{N}$ (11) when either of the normalized autocovariance estimators is used Equation 86, or Equation 84 for $k \ll N$, which comes from the fact that each autocovariance sample is asymptotically normal or Gaussian with zero mean and unit variance, and hence the range ± 1.96 corresponds to a probability of 0.95.

When two signals are measured from an experiment, the first question usually is if the two have some correlation or not. If they are independent or uncorrelated, theoretically their cross-covariance is zero for all values of k .

Again, in practice, this result is untenable as is exemplified in Fig. 14, which shows the normalized cross-covariance of two *independently generated* random signals $x(n)$ and $y(n)$. Each computer-simulated signal, $x(n)$ and $y(n)$, had $N = 1000$ samples and was generated independently from the other by a filtering procedure applied on normally distributed random samples. It can be shown (11) that under the hypothesis that two random signals are *white*, a 95% confidence interval for testing their uncorrelatedness is given by $\pm 1.96/\sqrt{N}$ when the normalized cross-covariance formula Equation 91 is used, either for the biased case or for the unbiased case if $k < N$. This confidence interval is drawn in Fig. 14, even without knowing beforehand if the signals are indeed white. Several peaks in the cross-covariance oscillations are clearly outside the confidence interval, which could initially be interpreted as a statistical dependence between the two signals. However, this cannot be true because they were generated independently. The problem with this apparent paradox is that the cross-covariance is also under the influence of each signal's autocovariance.

One approach to the problem of testing the hypothesis of independence of two arbitrary (nonwhite) signals is the **whitening** of both signals before computing the cross-covariance (11). After the whitening, the $\pm 1.96/\sqrt{N}$ confidence interval can be used. The whitening filters most used in practice are inverse filters based either on autoregressive (AR) or autoregressive moving average (ARMA) models. In the present example, two AR models shall be obtained for each of the signals, first determining the order by the Akaike information criterion (24) and then

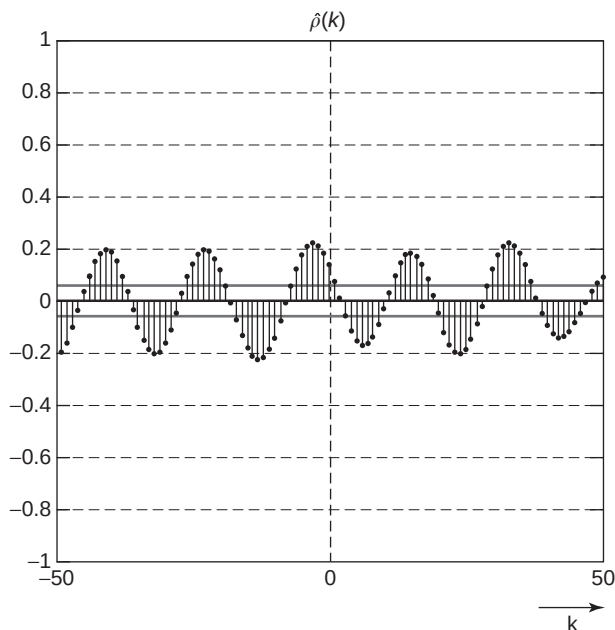


Figure 14. Cross-covariance of two signals $x(n)$ and $y(n)$, which were generated independently, with 1000 samples each. The cross-covariance shows oscillations above a 95% confidence interval indicated by the two horizontal lines, which would suggest a correlation between the two signals, but this conclusion is false because they were generated independently. (This figure is available in full color at <http://www.mrw.interscience.wiley.com/ebe>.)

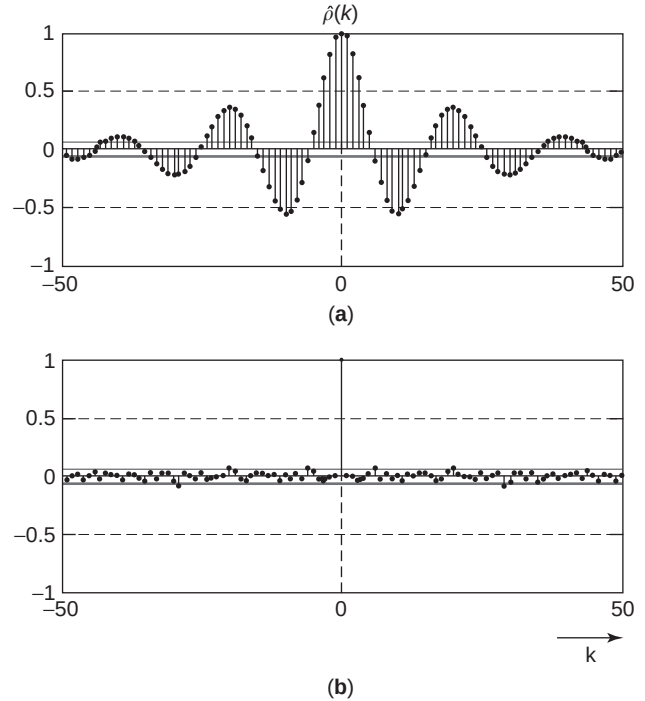


Figure 15. The autocovariance of signal $x(n)$ that generated the cross-covariance of Fig. 14 is shown in (a). It shows clearly that it is not a white signal because many of its samples fall outside a 95% confidence interval (two horizontal lines around the horizontal axis). (b) This signal was passed through an inverse filter computed from an AR model fitted to $x(n)$, producing a signal $x_o(n)$ that can be accepted as white as all the autocovariance samples at $k \neq 0$ are practically inside the 95% confidence interval. (This figure is available in full color at <http://www.mrw.interscience.wiley.com/ebe>.)

estimating the AR models by the Burg method (25,26). Each inverse filter is simply a moving average (MA) system with coefficients equal to those of the denominator of the transfer function of the respective all-pole estimated AR model. Figure 15a shows the autocovariance of signal $x(n)$, which clearly indicates that $x(n)$ is nonwhite, because most samples at $k \neq 0$ are outside the confidence interval. As $N = 1000$, the confidence interval is $\pm 6.2 \times 10^{-2}$, which was drawn in Fig. 15a and 15b as two horizontal lines. Figure 15b shows the autocovariance of the signal $x_o(n)$ at the output of the respective inverse filter that has $x(n)$ at its input. Similarly, Fig. 16a shows the autocovariance of signal $y(n)$ and Fig. 16b the autocovariance of the output $y_o(n)$ of the corresponding inverse filter that has $y(n)$ at its input. The two autocovariances in Fig. 15b and Fig. 16b show that, except for the sample at $k = 0$, which is equal to 1, the other samples stay within the confidence band, so that the hypothesis of white $x_o(n)$ and $y_o(n)$ can be accepted.

Next, the cross-covariance between $x_o(n)$ and $y_o(n)$ was computed and is shown in Fig. 17. Now, practically all the samples are contained in the confidence interval and the hypothesis that the two random signals $x(n)$ and $y(n)$ are uncorrelated (and independent, if they are Gaussian) can be accepted. For the few samples outside the confidence band, the user should use the fact that at each time shift k ,

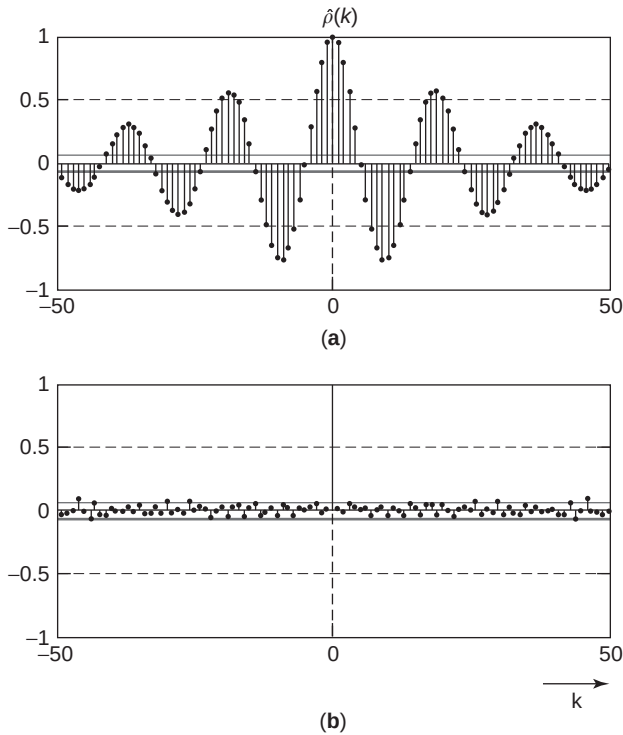


Figure 16. The autocovariance of signal $y(n)$ that generated the cross-covariance of Fig. 14 is shown in (a). It shows clearly that it is not a white signal because many of its samples fall outside a 95% confidence interval (two horizontal lines around the horizontal axis). (b) This signal was passed through an inverse filter computed from an AR model fitted to $y(n)$, producing a signal $y_0(n)$ that can be accepted as white as all the autocovariance samples at $k \neq 0$ are practically inside the 95% confidence interval. (This figure is available in full color at <http://www.mrw.interscience.wiley.com/ebe>.)

5% of the times a cross-covariance value could exist outside the 95% confidence interval, for independent signals. Besides this statistical argument, the biomedical user should also consider if the time shifts associated with samples outside the confidence band have some physiological meaning. Additional issues on testing the hypothesis of independency of nonwhite signals may be found in the text by Brockwell and Davis (11).

Let us apply this approach to the bilateral reflex data described in subsection “Common input.” Note that the autocovariance functions computed from experimental data, shown in Fig. 12, indicate that both signals are not white because samples at $k \neq 0$ are outside the confidence band. Therefore, to be able to properly test whether the two sequences are correlated, the two sequences of reflex amplitudes shall be whitened and the cross-covariance of the corresponding whitened sequences computed. The result is shown in Fig. 18, where one can see that a sample exists at zero time shift clearly out of the confidence interval. The other samples are within the confidence limits, except for two that are slightly above the upper limit line and are of no statistical or physiological meaning.

From the result of this test the possibility that the wide peak found in Fig. 13 is an artifact can be discarded and one can feel rather confident that a common input acting

synchronously (peak at $k=0$) on both sides of the spinal cord indeed exists.

An alternative to the whitening approach is to compute statistics based on surrogate data generated from each of the two given signals $x(n)$ and $y(n)$ (27). Each surrogate signal $x_i(n)$ derived from $x(n)$, for example, would have an amplitude spectrum $|X_i(e^{j\Omega})|$ equal to $|X(e^{j\Omega})|$, but the phase would be randomized. Here $X(e^{j\Omega}) = \text{DiscreteTimeFourierTransform}[x(n)]$. The phase would be randomized with samples taken from a uniform distribution between $-\pi$ and π . The corresponding surrogate signal would be found by inverse Fourier Transform. With this Monte Carlo approach (27), the user will have a wider applicability of statistical methods than with the whitening method. Also, in case high frequencies generated in the whitening approach cause some masking of a genuine correlation between the signals, the surrogate method is a good alternative.

Finally, returning to the problem of delay estimation presented earlier, in section 4.4 it should be pointed out that several related issues were not covered in this chapter. For example, the standard deviation or confidence interval of the estimated delay may be computed to characterize the quality of a given estimate (4). A modified cross-covariance, based on a minimum least square method, may optimize the estimation of the time delay (20). Alternatively, the model may be different from Equation 57, (e.g., one of the signals is not exactly a delayed

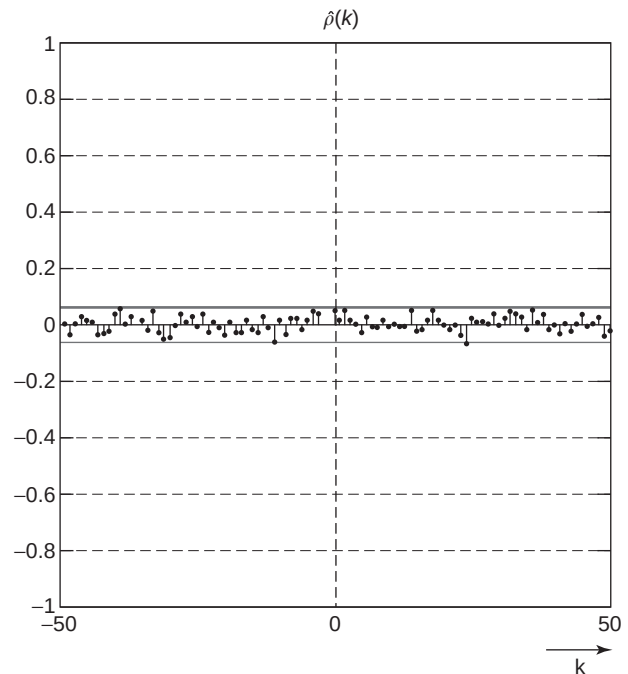


Figure 17. The cross-covariance estimate of the white signals $x_0(n)$ and $y_0(n)$ referred to in Figs. 15 and 16. The latter were obtained by inverse filtering the nonwhite signals $x(n)$ and $y(n)$ whose estimated cross-covariance is shown in Fig. 14. Here, it is noted that all the samples fall within the 95% confidence interval (two horizontal lines around the horizontal axis), which suggests that the two signals $x(n)$ and $y(n)$ are indeed uncorrelated. (This figure is available in full color at <http://www.mrw.interscience.wiley.com/ebe>.)

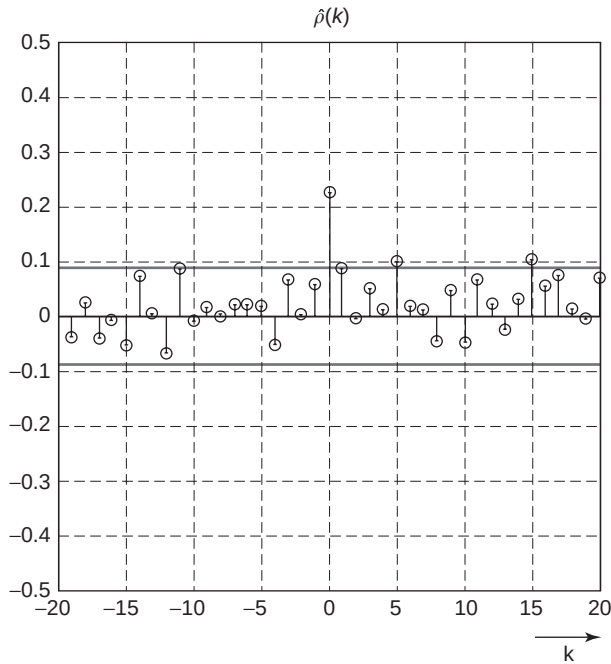


Figure 18. Cross-covariance computed from the whitened bilateral reflex signals shown in Fig. 11. Their original cross-covariance can be seen in Fig. 13. After whitening the two signals, the new cross-covariance shown here indicates that a statistically significant peak at time shift zero indeed exists. The two horizontal lines represent a 95% confidence interval. (This figure is available in full color at <http://www.mrw.interscience.wiley.com/ebe>.)

version of the other with added noise). One possible model would be $y(n) = \alpha q(n - L) + w(n)$, where $q(n) = x(n) * h(n)$. Here, $h(n)$ is the impulse response of a linear system that changes the available signal $x(n)$ into $q(n)$. The second available signal, $y(n)$, is a delayed and amplitude-scaled version of $q(n)$ with additive measurement noise. The concepts presented in section 5 earlier can be applied to this new model.

7. EXTENSIONS AND FURTHER APPLICATIONS

In spite of the widespread use of the autocorrelation and cross-correlation functions in the biomedical sciences, they have some intrinsic limitations. These limitations should be understood so that more appropriate signal processing tools may be employed or developed for the problem at hand. A few of the limitations of the auto and cross-correlation functions will be mentioned below together with some alternative signal processing tools that have been presented in the literature.

The autocorrelation and cross-correlation should not be applied to nonstationary random processes because the results will be difficult to interpret and may induce the user to errors. When the nonstationarity is because of a varying mean value, a subtraction of the time-varying mean will result in a wide-sense stationary random process (and hence the autocorrelation and cross-correlation

analyses may be used). If the mean variation is linear, a simple detrending operation will suffice. If the mean varies periodically, a spectral analysis may indicate the function to be subtracted. If the nonstationarity is because of time-varying moments of order higher than 1, it is more difficult to transform the random process into a stationary one. One of the approaches of dealing with these more complicated cases is the extension of the ideas of deterministic time-frequency analysis to random signals. An example of the application of such an approach to EEG signals may be found in Ref. 28.

When two stationary signals are, respectively, the input and the output of a linear system, the cross-correlation and the cross-spectrum presented above are very useful tools. For example, if the input is (approximately) white, the cross-correlation between the input and output provide us the impulse response of the system. A refinement is the use of the coherence function, which is the cross-spectrum divided by the square root of the product of the two signal autospectra (11,21). The squared magnitude of the coherence function is often used in biomedical signal processing applications (29,30). However, when significant *mutual interactions* exists between the two signals x and y , as happens when two EEG signals are recorded from the brain, it is very relevant to analyze the direction of the influences ($x \rightarrow y$; $y \rightarrow x$). One approach described in the literature is the partial directed coherence (31).

Several approaches have been proposed in the literature for studying nonlinear relations between signals. However, each proposed method works better in a specific experimental situation, contrary to the unique importance of the cross-correlation for quantifying the linear association between two random signals. The authors shall refer briefly to some of the approaches described in the literature in what follows.

In the work of Meeren et al. (32), the objective was to study nonlinear associations between different brain cortical sites during seizures in rats. Initially, a *nonlinear correlation coefficient* h^2 was proposed as a way to quantify nonlinear associations between two random variables, say x and y . The coefficient h^2 ($0 \leq h^2 \leq 1$) estimates the proportion of the total variance in y that can be predicted on the basis of x by means of a nonlinear regression of y on x (32). To extend this relationship to signals $y(n)$ and $x(n)$, the authors computed h^2 for every desired time shift k between $y(n)$ and $x(n)$, so that a function $h^2(k)$ was found. A peak in $h^2(k)$ indicates a time lag between the two signals under analysis.

The concept of *mutual information* (33) from information theory (34) has been used as a basis for alternative measures of the statistical association between two signals. For each chosen delay value the mutual information (MI) between the pairs of partitioned (binned) samples of the two signals is computed (35). The MI is non-negative, with $MI = 0$ meaning the two signals are independent, and attains a maximum value when a deterministic relation exists between the two signals. This maximum value is given by $-\log_2 \varepsilon$ bits, where ε is the precision of the binning procedure (e.g., for 20 bins, $\varepsilon = 0.05$, and therefore $MI \leq 4.3219$ bits). For example, Hoyer et al. (18) have employed such measures to better understand the statistical

structure of the heart rate variability of cardiac patients and to evaluate the level of respiratory heart rate modulation.

Higher-order cross-correlation functions have been used in the Wiener–Volterra kernel approach to characterize a nonlinear association between two biological signals (36,37).

Finally, the class of oscillatory biological systems has received special attention in the literature with respect to the quantification of the time relationships between different rhythmic signals. Many biological systems exist that exhibit rhythmic activity, which may be synchronized or entrained either to external signals or with another subsystem of the same organism (38). As these oscillatory (or perhaps chaotic) systems are usually highly nonlinear and stochastic, specific methods have been proposed to study their synchronization (39,40).

BIBLIOGRAPHY

1. R. Merletti and P. A. Parker, *Electromyography*. Hoboken, NJ: Wiley Interscience, 2004.
2. J. H. Zar, *Biostatistical Analysis*. Upper Saddle River, NJ: Prentice-Hall, 1999.
3. P. Z. Peebles, Jr., *Probability, Random Variables, and Random Signal Principles*, 4th ed., New York: McGraw-Hill, 2001.
4. J. S. Bendat and A. G. Piersol, *Random Data: Analysis and Measurement Procedures*. New York: Wiley, 2000.
5. C. Chatfield, *The Analysis of Time Series: An Introduction*, 6th ed., Boca Raton, FL: CRC Press, 2004.
6. P. J. Magill, A. Sharott, D. Harnack, A. Kupsch, W. Meissner, and P. Brown, Coherent spike-wave oscillations in the cortex and subthalamic nucleus of the freely moving rat. *Neuroscience* 2005; **132**:659–664.
7. J. J. Meier, J. D. Veldhuis, and P. C. Butler, Pulsatile insulin secretion dictates systemic insulin delivery by regulating hepatic insulin extraction in humans. *Diabetes* 2005; **54**:1649–1656.
8. D. Rassi, A. Mishin, Y. E. Zhuravlev, and J. Matthes, Time domain correlation analysis of heart rate variability in pre-term neonates. *Early Human Develop.* 2005; **81**:341–350.
9. R. Steinmeier, C. Bauhuf, U. Hubner, R. P. Hofmann, and R. Fahlbusch, Continuous cerebral autoregulation monitoring by cross-correlation analysis: evaluation in health volunteers. *Crit. Care Med.* 2002; **30**:1969–1975.
10. A. Papoulis and S. U. Pillai, *Probability, Random Variables and Stochastic Processes*. New York: McGraw-Hill, 2001.
11. P. J. Brockwell and R. A. Davis, *Time Series: Theory and Models*. New York: Springer Verlag, 1991.
12. A. F. Kohn, Dendritic transformations on random synaptic inputs as measured from a neuron's spike train — modeling and simulation. *IEEE Trans. Biomed. Eng.* 1989; **36**:44–54.
13. J. S. Bendat and A. G. Piersol, *Engineering Applications of Correlation and Spectral Analysis*, 2nd ed. New York: John Wiley, 1993.
14. M. B. Priestley, *Spectral Analysis and Time Series*. London: Academic Press, 1981.
15. C. W. Therrien, *Discrete Random Signals and Statistical Signal Processing*. Englewood Cliffs, NJ: Prentice-Hall, 1992.
16. A. V. Oppenheim, R. W. Schaffer, and J. R. Buck, *Discrete-Time Signal Processing*. Englewood Cliffs: Prentice-Hall, 1999.
17. R. M. Bruno and D. J. Simons, Feedforward mechanisms of excitatory and inhibitory cortical receptive fields. *J. Neurosci.* 2002; **22**:10996–10975.
18. D. Hoyer, U. Leder, H. Hoyer, B. Pompe, M. Sommer, and U. Zwiener, Mutual information and phase dependencies: measures of reduced nonlinear cardiorespiratory interactions after myocardial infarction. *Med. Eng. Phys.* 2002; **24**:33–43.
19. T. Muller, M. Lauk, M. Reinhard, A. Hetzel, C. H. Lucking, and J. Timmer, Estimation of delay times in biological systems. *Ann. Biomed. Eng.* 2003; **31**:1423–1439.
20. J. C. Hassab and R. E. Boucher, Optimum estimation of time delay by a generalized correlator. *IEEE Trans. Acoust. Speech Signal Proc.* 1979; **ASSP-27**:373–380.
21. D. G. Manolakis, V. K. Ingle, and S. M. Kogon, *Statistical and Adaptive Signal Processing*. New York: McGraw-Hill, 2000.
22. R. A. Mezzarane and A. F. Kohn, Bilateral soleus H-reflexes in humans elicited by simultaneous trains of stimuli: symmetry, variability, and covariance. *J. Neurophysiol.* 2002; **87**:2074–2083.
23. E. Manjarrez, Z. J. Hernández-Paxtián, and A. F. Kohn, A spinal source for the synchronous fluctuations of bilateral monosynaptic reflexes in cats. *J. Neurophysiol.* 2005; **94**:3199–3210.
24. H. Akaike, A new look at the statistical model identification. *IEEE Trans. Automat. Control* 1974; **AC-19**:716–722.
25. M. Akay, *Biomedical Signal Processing*. San Diego, CA: Academic Press, 1994.
26. R. M. Rangayyan, *Biomedical Signal Analysis*. New York: IEEE Press - Wiley, 2002.
27. D. M. Simpson, A. F. Infantosi, and D. A. Botero Rosas, Estimation and significance testing of cross-correlation between cerebral blood flow velocity and background electro-encephalograph activity in signals with missing samples. *Med. Biolog. Eng. Comput.* 2001; **39**:428–433.
28. Y. Xu, S. Haykin, and R. J. Racine, Multiple window time-frequency distribution and coherence of EEG using Slepian sequences and Hermite functions. *IEEE Trans. Biomed. Eng.* 1999; **46**:861–866.
29. D. M. Simpson, C. J. Tierra-Criollo, R. T. Leite, E. J. Zayen, and A. F. Infantosi, Objective response detection in an electroencephalogram during somatosensory stimulation. *Ann. Biomed. Eng.* 2000; **28**:691–698.
30. C. A. Champlin, Methods for detecting auditory steady-state potentials recorded from humans. *Hear. Res.* 1992; **58**:63–69.
31. L. A. Baccala and K. Sameshima, Partial directed coherence: a new concept in neural structure determination. *Biolog. Cybernet.* 2001; **84**:463–474.
32. H. K. Meeren, J. P. M. Pijn, E. L. J. M. van Luijckelaar, A. M. L. Coenen, and F. H. Lopes da Silva, Cortical focus drives widespread corticothalamic networks during spontaneous absence seizures in rats. *J. Neurosci.* 2002; **22**:1480–1495.
33. N. Abramson, *Information Theory and Coding*. New York: McGraw-Hill, 1963.
34. C. E. Shannon, A mathematical theory of communication. *Bell Syst. Tech. J.* 1948; **27**:379–423.
35. N. J. I. Mars and G. W. van Arragon, Time delay estimation in non-linear systems using average amount of mutual information analysis. *Signal Proc.* 1982; **4**:139–153.
36. P. Z. Marmarelis and V. Z. Marmarelis, *Analysis of Physiological Systems. The White-Noise Approach*. New York: Plenum Press, 1978.

37. P. van Dijk, H. P. Wit, and J. M. Segenhout, Dissecting the frog inner ear with Gaussian noise. I. Application of high-order Wiener-kernel analysis. *Hear. Res.* 1997; **114**:229–242.
38. L. Glass and M. C. Mackey, *From Clocks to Chaos. The Rhythms of Life*. Princeton, NJ: Princeton University Press, 1988.
39. M. Rosenblum, A. Pikovsky, J. Kurths, C. Schafer, and P. A. Tass, Phase synchronization: from theory to data analysis. In: F. Moss and S. Gielen, eds., *Handbook of Biological Physics*. Amsterdam: Elsevier, 2001, pp. 279–321.
40. R. Quian Quiroga, A. Kraskov, T. Kreuz, and P. Grassberger, Performance of different synchronization measures in real data: a case study on electroencephalographic signals. *Phys. Rev. E* 2002; **65**:041903/1–14.