

EEG BASIC PRINCIPLES—

PART II

Expected Graphoelements of the EEG: The Posterior Rhythm, the Anteroposterior Gradient, Vertex Waves of Sleep, and Sleep Spindles

Mark Libenson, MD

Senior Associate, Department of Neurology; EEG Laboratory Medical Director, END Technology Program Assistant Professor of Neurology, Harvard Medical School

The *posterior rhythm* (also called the posterior dominant rhythm) is one of the hallmark rhythms of the wakeful EEG. As expected from its name, it is dominant posteriorly, which is to say maximum in the occipital areas, but can spread as far forward as the F3 and F4 electrodes in some subjects. It is a rhythm seen when the eyes are closed, or at least when there is visual inattention. Conversely, if a patient happens to have eyes open throughout the whole waking recording, it may be difficult to find a good example of this rhythm unless the patient's visual fixation drifts off despite having eyes open. The posterior rhythm typically first appears between 3 and 4 months of age at 3–4 Hz, and increases in frequency by late childhood when it reaches adult ranges over 8 hertz. The posterior rhythm is often of higher voltage over the right hemisphere compared to the left. In occasional individuals a posterior rhythm is not identifiable, but when this is the only finding, its absence is not considered an abnormality. Therefore, if you want to win a bet with someone about whether a particular waveform represents the posterior rhythm, you should be able to demonstrate that it is 1) posterior-predominant, 2) only present during wakefulness, and 3) goes away when the eyes are open.

The *anteroposterior gradient* is a characteristic of the normal awake EEG. This really consists of two gradients. One is a voltage gradient, where higher voltages tend to be seen more posteriorly, and the second is a frequency gradient, where higher frequency activity tends to be seen more anteriorly. Therefore, in the normal waking EEG, frontal activity is low voltage and high-frequency, while posterior activity is higher voltage (e.g., the posterior rhythm) and lower frequency.

The first signs of drowsiness and sleep are a dropout of the posterior rhythm even when the eyes are closed, an increase in theta wave activity, and finally the appearance of *vertex waves of sleep* (these are maximum in the Cz electrode at the vertex and flanking C3 and C4 electrodes). This is followed soon after by the advent of stage II sleep which is marked by the appearance of *sleep spindles*: 14 Hz sinusoidal waves that also appear maximally in the Cz, C3, and C4 electrodes, but sometimes with a field that extends somewhat more frontally in some patients.

Describing the EEG

The first part of the description of an EEG event is generally its location on the scalp. If confined to one or just a few electrodes, the event can be described as *focal*. If an event occurs over a whole hemisphere at once, the term *lateralized* can be applied. The term *multifocal* generally implies that three or more independent locations are involved, and not all in the same hemisphere. The term *generalized* implies that all scalp electrodes are involved at the same time. Because even with generalized discharges, there is often a point of maximum voltage over the head, the term *generalized, maximum* is quite useful (e.g., “generalized, maximum frontal”). For instance, the spike-wave discharges of primary generalized epilepsy are often generalized, maximum in the F3 and F4 contacts.

Voltage

Voltage is ascertained from the height of the wave on the EEG page and the amplifier setting being used. Customarily, an amplifier’s strength is stated in

terms of sensitivity units in the world of EEG, but in most other fields, amplifiers are described in terms of gain. Consider the difference:

Amplifier Sensitivity versus Gain

Imagine you wanted to boast about how strong your amplifier is. Two general marketing strategies could be used to do this. One would be to say, “my amplifier is so strong that if you have a little 1-microvolt spike, it can move the pen a full 7 millimeters!” A yet stronger competitor might claim that with a little 1-microvolt spike, her amplifier could move the pen 10 millimeters! This is a way to state the *gain* of the amplifier, and the unit would be 7 or 10 millimeters per microvolts—increasing gain means bigger and bigger pen-sweeps for the same small electrical event.

The second strategy to impress your friends as to how good your amplifier is would be to talk about how sensitive it is. “You know, most amplifiers on the market today need a spike of 10 microvolts to move the pen by one millimeter, but my amplifier is so good that a mere 7-microvolt spike will move the pen one millimeter!” This is stating the amplifier’s behavior in terms of its sensitivity, and this is way we describe an amplifier’s setting in the world of EEG. As you can see, the sensitivity units corresponding to this second descriptive strategy would be 10 microvolts per millimeter ($\mu\text{V}/\text{mm}$) and $7 \mu\text{V}/\text{mm}$. The awkwardness of using sensitivity is that as the amplifier setting gets stronger, the numerical unit for sensitivity goes down (e.g., in this example changing the sensitivity from $10 \mu\text{V}/\text{mm}$ to $7 \mu\text{V}/\text{mm}$ makes the waves bigger on the page—the lower number making the waves bigger can take a bit of getting used to). Since lower sensitivity numbers make bigger waves, $2 \mu\text{V}/\text{mm}$ and $1 \mu\text{V}/\text{mm}$ are just about as sensitive or “strong” as our EEG amplifiers are used in common practice.

Synchrony and Desynchronization

These words sound quite similar but actually refer to different concepts. The term synchrony can be used to talk about whether a wave or event that is occurring in one area of the brain is occurring in another area of the brain simultaneously, or *synchronously*. A spike can occur at the same time in F3 as in F4, and therefore it is occurring synchronously in those two locations.

(Alternatively it could occur independently in F3 and F4, which is to say first in one place, then then in the other.) Most often the term synchrony describes whether something occurring in one hemisphere of the brain is also occurring at the same time in the opposite hemisphere: interhemispheric synchrony, also called bisynchrony. A simple example of this is a 12-month-old patient who has a sleep spindle first in the left central area, followed by a spindle in the right central area. At this point the child can be said to have “asynchronous spindles.” By 2 years of age, spindles usually occur at the same time in both frontocentral areas, and the child can be said to have synchronous or bisynchronous spindles.

The concept of synchronization is most often referenced in terms of *desynchronization*, and has little to do with the discussion of synchrony in the previous paragraph. When we say the EEG has desynchronized, this is the same as saying that the voltage has attenuated, or gone down. Understanding desynchronization requires imagining what is going on under each individual EEG electrode. Imagine that all of the neurons being “read” by a single electrode are doing the exact same thing at the exact same time. All of their electrical efforts will be additive, and they will form a nice wave that can be recorded by the electrode. In the opposite circumstance, all the neurons below the electrode might be working just as hard, but every neuron is doing something different at any given time. In this situation, their individual activities tend to cancel each other out, and the electrode records a low voltage. An analogy for this would be everyone in a 30,000-seat baseball park singing a *different* song at the same time. A microphone above the park would hear the equivalent of white noise, and a waveform representing the volume of noise would remain fairly flat (desynchronization). If, on the other hand, everyone in the park sang “Sweet Caroline” in unison, a microphone above the park would record the distinctive waveforms for that song, including all the variations of volume (synchronization). In both examples, everyone in the park is singing just as strongly, but in the desynchronized baseball park, the volume waveform recorded by the microphone is pretty flat despite all the singers’ hard work.

Rhythmicity, Continuity, and Periodic Patterns

The term *rhythmicity* means pretty much what it sounds like. One way of looking at rhythmicity is to imagine that you can see three or four waves in a run but you have covered successive waves with your hand. If the run is rhythmic, you should be able to predict fairly well where the fifth wave will fall. The terms arrhythmic or irregular can be used for waves whose cadence is unpredictable. The term *semirhythmic* is used to describe a situation somewhere between rhythmic and irregular: there is some sense of regularity to the events, but they are not so regular as to be predictable most of the time.

The term *periodic* refers to an event that occurs repetitively, and the intervals between them are fairly similar. The term periodic also implies that the events are separated by time, as opposed to the peaks and troughs of a sine wave which are seen to occur in a continuous chain. Also, *periodic* implies slower frequencies of occurrence, typically from once per second to once every four seconds. Consider a door that is swinging open and closed rhythmically compared to a door that is opened and slammed shut every four seconds.

Spike-Wave Complexes, Sharp Waves, Spikes, and Transients

Spike-wave complexes are one of the most common epileptiform findings in the EEG. The full name is spike and slow-wave complex. In spike-wave complexes, the spike is followed by a slow wave in a fairly immediate and highly fixed interval. The term “spike” of course suggests a spike without a wave. In fact, some believe that there is always a wave to be found following a spike if you look closely enough. Still, when describing an EEG, it is sensible to use the term that best fits your visual impression—if you see what looks pretty much look like a pure spike, then that would be the best term. You may be surprised to learn that the difference between a spike and a sharp wave is defined not by how sharp vs. blunt the tip is, but how wide the base is. Sharp events with a base narrower than 70 milliseconds are referred to as spikes and those with bases broader than 70 milliseconds are called sharp waves. It is obvious that this definition is somewhat arbitrary, and it is also worth noting that few, if any epileptic syndromes would be distinguished by deciding whether a wave is a spike as opposed to a sharp wave.

A *transient* is an engineering term for any very short-lived (= fast) event in a signal. Therefore, sharp waves and spikes are electrical transients. In the world of EEG reporting, however, the term transient is generally reserved for events that look like a spike, but for one reason or another the reader does not want to imply that the wave is epileptiform, i.e., associated with epilepsy. Using the word “transient” avoids the sting of using the word “spike” even though the wave may be spike-shaped. Therefore, in the world of EEG usage/etiquette, the terms transient and spike can refer to similar waveforms, but transient is used if the reader chooses to underplay the wave’s association with epilepsy. Occasionally the term is used in specific expressions, such as “neonatal sharp transients” (a normal EEG graphoelement in newborns) probably again to impart the idea that this is something that looks like a spike, but is not actually epileptiform. The word transient may also be used to describe something that looks like a spike, but which you think may be an artifact, thus avoiding the more loaded word “spike.”

Frequency and Wavelength

Frequencies in EEG have historically been defined with the use of Greek letters, written out in English rather than using the letter from the Greek alphabet (therefore, “alpha,” not ‘α’). Delta: 0 to ≤ 4 Hz, theta: 4 to ≤ 8 Hz, alpha: 8-13 Hz, and beta: >13 Hz. For faster fast activity, generally above 30 hertz and perhaps up to 100 hertz, the term gamma activity is used, but no definition of the specific range has been universally accepted. Gamma and even faster activity (ripples or high frequency oscillations/HFOs) are also recognized, but are not yet a part of routine conventional EEG interpretation, though significant interesting information about epileptogenicity can be gleaned with special techniques looking at these very high-frequency ranges.

Frequency is directly related to wavelength in a reciprocal relationship: $f = 1/\lambda$, where f is the frequency in Hz (or cycles per second) and λ (lambda) is the wavelength stated in seconds. From this it follows that a wave that takes up (has a wavelength of) 1/5 of a second or 0.2 seconds has a frequency of $1/0.2 = 5$ Hz (cycles per second). Looking at it in the other direction, spindle waves at 14 hertz should have individual wavelengths of 1/14 of a second or approximately 0.07 seconds or 70 milliseconds. Therefore, you can determine the frequency of

a wave either by counting how many times it cycles in one second or by measuring the wavelength (peak to peak or trough to trough) and taking the reciprocal, even of a single wave.

EEG Reporting

Each EEG laboratory should have a standard reporting template which usually starts with identifying information about the patient and the date of the study. A brief statement of the chief complaint or the reason for the testing is included. This information is usually a combination of the clinical information submitted by the ordering provider on the requisition and history obtained by the EEG technologist. The medications the patient is taking at the time of EEG and the amount of sleep deprivation are also stated. The technical aspects of the study including the specific recording technique, electrode placement, and duration of the study are also described in this first paragraph. This part of the report is usually prepared from a template that reflects the standard recording technique that is used for the majority of patients seen in the lab, with appropriate edits when there are deviations from routine. Different templates can be used for particular patient types, such as newborns or ICU patients who may be in coma.

The goal of the Description portion of the report is ideally to allow the readers of the report to create a visual image in their mind of what the EEG looks like based on your prose description, without their having to view the tracing itself. Such descriptions are especially useful (and challenging) to create for recordings with nonstandard appearances, such as a tracing in coma. This could be something like: “a low-voltage 2 to 3 Hz irregular background pattern is interrupted every 3 to 4 seconds by a 150- μ V generalized polymorphic burst of activity lasting one second.” In theory, the reader would be able to draw the EEG more or less accurately based on the verbal description without ever having seen the tracing. Technical EEG terms such as electrode names (e.g., T7 and O2) may be used in the Description, since the writing style used in this portion of the report is generally aimed at persons familiar with reading EEGs and EEG terminology.

Next, the EEG should have an Abnormality List, assuming abnormalities are present. It is best to use a more “English” approach to this paragraph, for

instance saying left occipital spikes rather than O1 spikes, aiming this portion of the report at an audience less familiar with EEG terminology. It is useful to list abnormalities in order of importance, with the most important listed first.

The Clinical Correlation paragraph is the place where real clinical neurophysiologists “show their stripes.” Up until now, interpretation of the EEG is somewhat rote and mechanical. In the clinical correlation, the reader takes together the age of the patient, the indication for the EEG, and the clinical story as described, and within the limits of the information conveyed by the EEG, says what it all means. It is useful after generating an abnormality list to ask yourself, “what have I learned (and not learned) about this patient from this EEG result?” In some cases, there is little more to say about a spike that may have been found in the F7 electrode than that “the finding suggests the possibility of a lower seizure threshold in the left anterior temporal region.” Note that this interpretation does not go so far as to say that the patient has left-sided temporal lobe epilepsy, which would often be erroneous, as not all patients with spikes have seizures. Likewise, if right centrotemporal spikes are seen in an 8-year-old, it would be reasonable not just to say that this suggests a decreased seizure threshold in this area, but also to mention whether or not the discharge potentially fits into the category of a rolandic spike, since that is what the ordering physician would be most curious about with this finding. For instance, if there are slow-wave abnormalities or other abnormalities in the EEG inconsistent with rolandic epilepsy, this comment would be made.

Occasionally, an EEG is diagnostic, as in a case where absence seizures are recorded with clear clinical change. Note that the EEG can be considered diagnostic of absence seizures, but not necessarily of childhood absence epilepsy. That second question is for the clinician, who knows all of the patient’s background, imaging results, etc., to decide about. Even with a normal EEG result, what have we learned about the patient? Indeed, a normal EEG reduces the chances that the patient has epilepsy, but does not exclude the diagnosis.

Philosophy Regarding the Threshold for Calling a Finding Abnormal

Many of our discussions about EEG interpretation are cut and dry—either we saw a spike or we didn't, implying that the EEG is either abnormal or normal. From time to time, we come across situations where it's hard to decide whether to "call" a spike (or perhaps some other abnormality) or not. Indeed, it's interesting to consider the scenario where your confidence that an EEG is abnormal is exactly 50 percent—you just can't decide for sure about that spike. The question here is, what should your bias be? Should you waiver toward a tendency to "undercall" and bias toward calling unsure situations normal, or to "overcall," calling unsure situations abnormal. What this boils down to from the practical point of view is what the injury would be to the patient of making the wrong diagnosis in either direction.

Essentially, there are two types of error we can make in this 50:50 situation. In what we'll call a Type 1 error, the patient really has epilepsy, but we erroneously call the EEG normal. In a Type 2 error, the patient does not have epilepsy, the EEG is really normal, but we call the EEG abnormal. Which kind of error is worse in the field of epilepsy?

The specific disease process we are testing for is important when considering the potential impact of each type of error. For instance, in cancer diagnosis, if you feel there's a 50 percent chance that there is cancer on a test but you decide to call the test normal, the cost to 50% of the patients can be quite high. In EEG, however, the situation can be the opposite. The cost of being informed that the EEG is abnormal when it is truly normal can be quite high. Consider the scenario where a referring physician may have obtained an EEG for a low suspicion event (e.g, syncope vs. seizure). In such a case, the finding of an epileptiform discharge in the EEG may prompt a diagnosis of epilepsy and treatment with medication, often for two years or longer. This may also have impact on the patient's lifestyle in terms of driving, other activities, and even occupation. On the other hand, a patient who has a truly abnormal EEG that is read out as normal will likely experience a less deleterious impact. This is mostly due to the fact that practitioners know that a normal EEG does not rule out the diagnosis of epilepsy and the referring physician still has the option to consider pursuing the diagnosis (perhaps even with repeat EEG recordings which may show a clearer

picture in the future). While we always strive to produce as accurate an interpretation of the EEG as possible, in the case of a 50:50 situation this type of “Type 1” error of undercalling a finding has less of a chance of harming the patient compared to the “Type 2” overcalling error. Therefore, when unsure about the presence of an abnormal finding (i.e., “when in doubt...”), experienced electroencephalographers have a bias toward calling the study normal, but will describe the unclear findings within the text of the report.