

PART I

IEEE 802.11 WIRELESS LANs

COPYRIGHTED MATERIAL

IEEE 802.11 MEDIUM ACCESS CONTROL AND PHYSICAL LAYERS

KAVEH GHABOOSI, MATTI LATVA-AHO, and YANG XIAO

1.1 INTRODUCTION

A wireless local area network (WLAN) is an information system¹ intended to offer diverse location-independent network service access to portable wireless devices using radio waves instead of wired infrastructure. In corporate enterprises, WLANs are typically deployed as the ultimate connection between an existing cable infrastructure network and a cluster of mobile clients, giving them wireless access to the shared resources of the corporate network across a building or campus setting. Fundamentally, WLANs liberate customers from reliance on hard-wired access to the network backbone, giving them anywhere, anytime network services access. The pervasive approval of WLANs depends upon industry standardization to ensure product compatibility and reliability among various brands and manufacturers. Among existing system architectures, the IEEE 802.11 family is the most popular and accepted standard concerning medium access control (MAC) and physical (PHY) layers in WLANs; therefore, in this chapter, we briefly overview its basic features in both aforementioned layers. We start our investigation with the MAC layer and its fundamental components. Supported network types, different network services, and media access schemes are covered, accordingly. Subsequently, the physical layer and its basic characteristics are discussed. Different technologies,

¹ In telecommunications, an information system is any telecommunications and/or computer-related equipment or interconnected system or subsystems of equipment that are used in the acquisition, storage, manipulation, management, movement, control, display, switching, interchange, transmission, or reception of voice and/or data and includes software, firmware, and hardware (Federal Standard 1037C, MIL-STD-188, and National Information Systems Security Glossary).

including frequency hopping (FH), direct-sequence spread spectrum (DSSS) and its high rate (HR) counterpart (i.e., HR/DSSS), and orthogonal frequency division multiplexing (OFDM), recommended for the IEEE 802.11 physical layer are then explored. As a result, this chapter can be assumed as a comprehensive overview of the IEEE 802.11 standard.

1.2 IEEE 802.11 MAC PROTOCOL

In 1997, the IEEE 802.11 working group (WG) proposed the IEEE 802.11 WLAN standard and, subsequently, a revised version was released in 1999. The primary medium access scheme in IEEE 802.11 MAC is the distributed coordination function (DCF), a contention-based protocol which is based on the carrier sense multiple-access/collision avoidance (CSMA/CA) protocol. In the DCF, mobile terminals should contend for the shared wireless channel, and as a result, the medium access delay for each station (STA) cannot be bounded in heavy-traffic-load circumstances. Thus, the DCF is capable of offering only asynchronous data transmission on a best effort (BE) basis. In order to support real-time traffic such as voice and video, the point coordination function (PCF) scheme has been advised as a noncompulsory option. Basically, the PCF is based on a centralized polling scheme for which a point coordinator (PC) residing in an access point (AP) provides contention-free services to the associated stations in a polling list. In addition to the IEEE 802.11 standard [1], there is a well-known book by Gast [2] which is considered as a complete scientific review of 802.11 families. Due to the popularity of the aforementioned book, we will use it frequently throughout the section to refer the reader to more technical issues and discussions.

Recently, considerable interest in wireless networks supporting quality of service (QoS) has grown noticeably. The PCF is already available in IEEE 802.11 to offer QoS but has not yet been implemented in reality due to its numerous technical limitations and performance drawbacks. For that reason, the 802.11 WG initiated IEEE 802.11e activity to develop the existing 802.11 MAC to facilitate support of QoS. Regarding the 802.11e amendment, not only the IEEE 802.11e standard [3] but also recognized introductory and survey papers [4, 5] will be used repeatedly as key references. We cite many technical issues from these works and the references therein to more appropriately explain 802.11/802.11e-based system features.

1.2.1 Categories of 802.11 Networks

The key constructing component of an 802.11 network is the basic service set (BSS), a group of wireless terminals that communicate with each other over a common radio channel [1, 2]. Data transmission is accomplished within a *basic*

service area, defined by radio propagation characteristics of wireless channel. BSSs come in two categories, as illustrated in Fig. 1.1.

The *infrastructure BSS* networks are primary for mobile stations to access the Internet via an AP so that, in most of cases, communications between two stations within the same service set do not happen. In communications between mobile stations in the same service set, the AP deployment acts as an intermediate node for all information exchanges comprising communications. In other words, any data communication between two wireless clients should take two successive hops, i.e., source STA to AP and AP to destination STA. Obviously, the exploitation of APs in infrastructure networks brings two major advantages. On the one hand, no restriction is placed on the physical distance between mobile stations. On the other hand, allowing straight communication between wireless terminals would apparently preserve system capacity² but at the cost of increased physical and MAC layer complexity. The most important functions of an AP are to assist stations in accessing the Internet and help save battery power in associated wireless stations. If a mobile terminal is in the power-saving (PS) mode, the AP buffers those frames destined to reach the station during the period it will be in PS status. When the terminal exits the PS mode, the AP forwards the cached data frames to the station one by one. Hence, APs evidently play a key role in infrastructure networks to make implementation of PS mechanisms possible [1].

In the independent BSS (IBSS), mobile stations are allowed to communicate directly. Characteristically, IBSSs are composed of a few stations configured

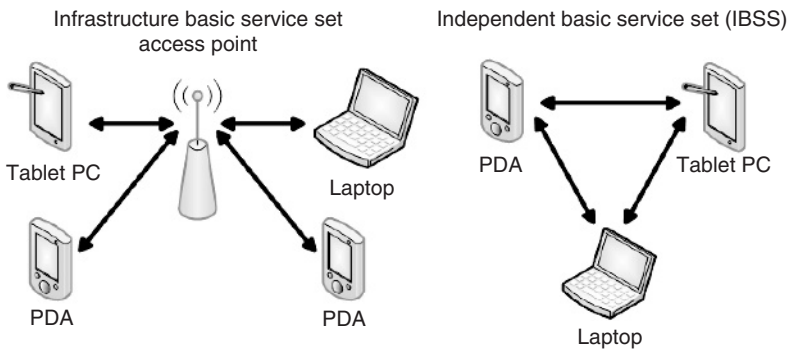


FIGURE 1.1 Infrastructure and independent basic service sets.

²In computer science, *channel capacity* is the amount of discrete information that can be reliably transmitted over a channel. By the noisy-channel coding theorem, the channel capacity of a given channel is the limiting information transport rate (in units of information per unit time) that can be achieved with vanishingly small error probability.

for a particular goal and for a short period of time. IBSSs are sometimes referred to as ad hoc BSSs or simply ad hoc networks.

IEEE 802.11 allows wireless networks of arbitrary size to be installed and utilized by introducing the extended service set (ESS) concept. Basically, the ESS is constructed by chaining neighboring BSSs and requires a backbone system that provides a particular set of services. Figure 1.2 illustrates an ESS as a combination of three neighboring BSSs. Switching between adjacent BSSs while being connected to the system is called a handoff.³ Stations with the same ESS are able to communicate with each other even if they are not in the same BSS or are moving from one point to another. For associated stations within an ESS, a wireless network should behave as if it were a single layer 2 local area network (LAN). In such an architecture, APs are similar to layer 2 bridges; consequently, the backbone network should be a layer 2 network as well (e.g., Ethernet). Several APs in a single area may be connected to a single switch or can even use virtual LANs (VLANs) if the link layer connection spans a larger area. 802.11 supplies link layer mobility within an ESS, but only if the backbone network is a single link layer domain, such as a shared Ethernet or a VLAN [2].

Theoretically, extended service areas are the highestlevel abstraction supported by 802.11 wireless networks. In order to let non-802.11 network devices use the same MAC address to exchange data traffic with an associated station somewhere within the ESS, APs should mimic an absolute cooperative system. In Fig. 1.2, the illustrated gateway uses a single MAC address to deliver data frames to the targeted mobile stations in different BSSs. This is the MAC address of an AP with which the intended wireless station has been already associated. As a result, the gateway is unaware of the actual location of a tagged wireless terminal and relies only on the corresponding AP to forward data traffic [1, 2]. The backbone network to which APs are connected is called the distribution system (DS) since it makes delivery of information to and from the outside world possible.

It should be noted that technically different types of 802.11 networks may coexist at the same time. For instance, IBSSs might be constructed within the basic service area of an AP. Coexisting infrastructure BSSs and IBSSs should share the same radio channel capacity and, as a result, there may be adverse performance implications from colocated BSSs as well [2].

³ In cellular telecommunications, the term *handoff* refers to the process of transferring an ongoing call or data session from one channel connected to the core network to another. In satellite communications, it is the process of transferring satellite control responsibility from one earth station to another without loss or interruption of service. The British term for transferring a cellular call is *handover*, which is the terminology standardized within European-originated technologies such as Global System for Mobile (GSM) communications and Universal Mobile Telecommunications System (UMTS). In telecommunications, there are two reasons why a handoff (handover) might be conducted: if the mobile terminal has moved out of range from one cell site, i.e., base transceiver station (BTS) or AP, and can get a better radio link from a stronger transmitter or, if one BTS/AP is full, the connection can be transferred to another nearby BTS/AP.

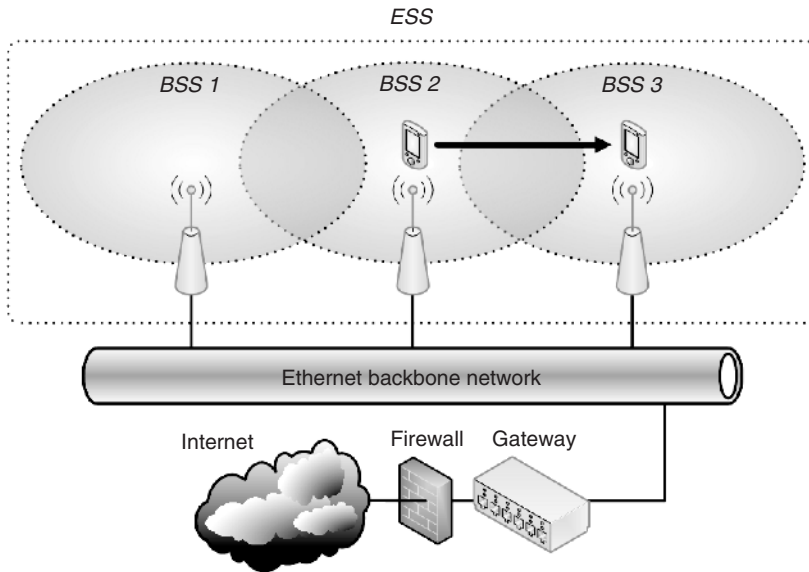


FIGURE 1.2 Extended service set.

1.2.2 IEEE 802.11 Networks Services

Generally the IEEE 802.11 standard provides nine dependent services: Three of these services are dedicated distinctively to data transfer purposes while the remaining ones are explicitly devoted to management operations enabling network systems to keep track of mobile stations and react in different circumstances accordingly.

The *distribution* service is exploited in infrastructure networks to exchange data frames. Principally, the AP, upon receiving a MAC protocol data unit (MPDU), uses this service to forward it to the intended destination station. Therefore, any communication with an AP should use a distribution service to be possible. *Integration* is a specific service provided by the distribution system that makes connection with a non-802.11 network possible. The integration function is not expressed technically by the standard, except in terms of the services it should offer.

MAC frame delivery to the associated terminals will not be possible unless the *association* service ensures that the AP and connected stations can work together and use the network services. Consequently, the distribution system is able to use the registration information to determine the AP with which a specific mobile station has been associated. In other words, unassociated wireless terminals are not permitted to obtain *any* service from the whole system. In an ESS, when a mobile station moves between different BSSs, there should be a set of handoffs to be accomplished in order to keep the station connected to the system. *Reassociation* is generally initiated by a wireless

terminal once the signal strength indicates that a different association is necessary. This means that handoff and reassociation requests are never commenced by APs. Upon completion of reassociation, the distribution system renews its location records to reflect the latest information about reachability of the mobile station. To terminate an existing association, wireless stations may possibly use the so-called *disassociation* service. Upon invocation of disassociation, any mobility information stored in the distribution system corresponding to the requesting station is removed at once.

Authentication is an obligatory prerequisite to association due to the fact that only authenticated users are authorized to use the network resources. If the APs of a distribution system have been configured in such a way as to authenticate any station, then the system is called an “open system” or an “open network.” These kinds of wireless networks can be found, for instance, at university campuses. *Deauthentication* terminates an authenticated relationship between an AP and a wireless station. Since authentication is required before system resources utilization, a side effect of deauthentication is termination of any existing association.

IEEE 802.11 offers a noncompulsory *privacy* service called wired equivalent privacy (WEP). WEP is not iron-clad security; in fact, it can be easily disabled. In response, the IEEE 802.11i task group (TG) is seeking an enhanced and stronger security scheme to be included in the next generation of 802.11 equipments. IEEE 802.11i, known as WiFi-protected access version 2 (WPA2), is an amendment to the 802.11 standard specifying security mechanisms for wireless networks. It makes use of the advanced encryption standard (AES) block cipher, while WEP and WPA (an earlier version) use the RC4 stream cipher. The 802.11i architecture contains the following components: 802.1X for authentication [entailing the use of extensible authentication protocol (EAP) and an authentication server], the robust security network (RSN) for keeping track of associations, and the AES-based counter mode with cipher block-chaining message authentication code protocol (CCMP) to provide confidentiality, integrity, and origin authentication. Another important element of the authentication process is an innovative four-way handshake.

The MPDU is a fancy name for 802.11 MAC frames. The MPDU does not, however, include PHY layer convergence procedure (PLCP) headers. On the other hand, the MAC service data units (MSDUs) are only composed of higher level data units [e.g., Internet protocol (IP) layer]. For instance, an 802.11 management frame does not contain an MSDU. Wireless stations provide the *MSDU delivery* service, which is responsible for getting the data to the actual recipient.

1.2.3 IEEE 802.11 Media Access Schemes

In what follows, we discuss the medium access rules defined in the 802.11 standard and its corresponding amendments. We begin the discussion with the contention-based *802.11 DCF* access scheme. Subsequently, a few paragraphs

are dedicated to the *802.11 PCF*, which is a contention-free channel acquisition technique. Finally, the supplementary QoS-aware amendment of the IEEE 802.11 standard, i.e., the *802.11e hybrid coordination function (HCF)*, is explored [1–3].

1.2.3.1 IEEE 802.11 DCF. The fundamental IEEE 802.11 access scheme is referred to as the DCF and operates based upon a listen-before-talk (LBT) approach and CSMA/CA.

As indicated, MSDUs are transmitted using MPDUs. If the wireless station chooses to fragment a long MSDU into a number of MPDUs, then it should send the long MSDU through more than one MPDU over the radio system. 802.11 stations deliver MSDUs following a media detection procedure dealing with an idle wireless channel that can be acquired for data transmission. If more than one station senses the communication channel as being idle at the same time, they might commence their frame transmissions simultaneously, and inevitably a collision occurs subsequently. To minimize the collision risk, the DCF uses carrier sense functions and a binary exponential backoff (BEB) mechanism. In particular, two carrier sense schemes, namely physical and virtual carrier sense functions, are employed to simultaneously resolve the state of the radio channel. The former is offered by the physical layer and the latter by the MAC layer, called network allocation vector (NAV). The NAV records the duration that the medium will be busy based upon information announced before the control/data frames are captured over the air interface. If either function indicates a busy medium, the medium is considered busy (i.e., reserved or occupied); if not, it is considered idle. Subsequent to detection of wireless medium as *idle*, for a so-called DCF interframe space (DIFS) time duration, stations continue sensing the channel for an extra random time period called a *backoff period*. The wireless station begins traffic delivery whenever the shared medium remains idle over this further random time interval. The backoff time is determined by each station as a multiple of a pre defined slot time chosen in a stochastic fashion. This means that a fresh independent random value is selected for every new transmission. In the BEB algorithm, each station chooses a random backoff timer uniformly distributed in an interval $[0, CW - 1]$, where CW is the current *contention window* size. It decreases the backoff timer by 1 for every idle time slot. Transmission is started whenever the backoff timer reaches zero. When frame transmission fails due to any reason, the station doubles the CW until it reaches the maximum value $CW_{\max}$. Afterward, the tagged station restarts the backoff procedure and retransmits the MAC frame when the backoff counter reaches zero. If the maximum transmission retry limit is reached, the retransmission should be stopped, the CW should be reset to the initial value $CW_{\min}$, and the MAC frame is simply discarded. At the same time as a wireless station is counting down its backoff counter, if the radio channel becomes busy, it suspends its backoff counter decrement and defers from the media acquisition until the medium again becomes idle for a DIFS [1, 2, 4, 5].

Each MPDU requires the reception of an acknowledgment (ACK) frame to confirm its correct transmission over the wireless channel. If for any reason the intended ACK frame is not received right after the MPDU transmission, the source station concludes that the MPDU was not delivered successfully and may reiterate the transmission. Basically, the CW size of a contending station increases when the transmission fails. After an unsuccessful effort, the backoff procedure is restarted with a double-sized CW, up to a maximum value defined by $CW_{\max}$. Alternatively, subsequent to a successful transmission, the tagged station exploits another random backoff, even if there is no further queued MSDU to be delivered over the air interface. In the literature, this extra backoff is referred to as *post-backoff* given that it is executed subsequent to the data frame departure. There is an exception to the above-mentioned rule: If an MSDU arrives from layer 3 when (1) the transmission queue is vacant, (2) the latest post-backoff has finished, and (3) the medium has been idle for at least one DIFS, then it may be delivered at once with no further backoff procedure [4].

To overcome the so-called *hidden terminal* problem, the IEEE 802.11 DCF media access scheme utilizes a request-to-send/clear-to-send (RTS/CTS) mechanism which can be exploited optionally prior to MPDU transmission. As illustrated in Figure 1.3, the source station sends an RTS control frame to its intended destination. Upon reception of the RTS by the receiver, it sends a CTS frame back to the source station. The RTS and CTS frames include information on how long it will take to deliver the upcoming data frame (in the fragmentation case, it indicates the duration of the first fragment) and the corresponding ACK over the radio link. Upon reception of either the RTS or CTS, wireless stations located in the radio range of the transmitting node, in addition to those hidden to the source node and located in the transmission range of the destination station, set their local timer NAV, with the duration announced within the RTS/CTS frames. The RTS and CTS frames protect the MPDU from interferences due to other neighboring wireless nodes. Stations that receive these control frames will not initiate transmission until the above-mentioned NAV timer expires. Between two consecutive frames in the sequence

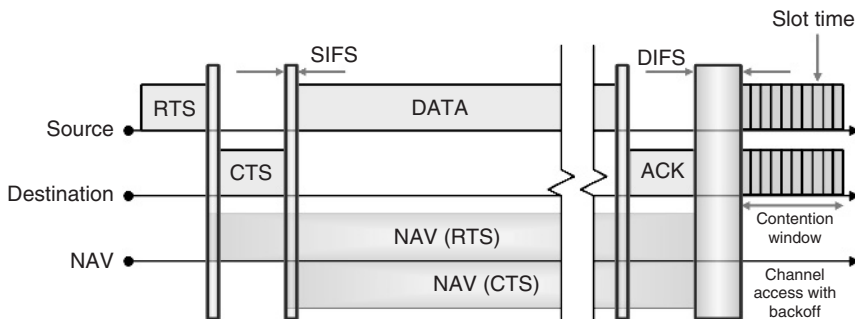


FIGURE 1.3 IEEE 802.11 RTS/CTS access scheme.

of RTS, CTS, MPDU, and ACK frames, a short interframe space (SIFS) gives transceivers time to switch (i.e., between transmitting and receiving modes). It is noteworthy that the SIFS is shorter than the DIFS, which gives the CTS and ACK the highest priority access to the wireless medium [1].

1.2.3.2 IEEE 802.11 PCF. IEEE 802.11 employs an optional PCF to support QoS for time-bound delay-sensitive services. The PCF offers techniques for prioritized access to the shared radio channel and is centrally coordinated by a PC station which is typically an AP. The PCF has higher priority than the DCF scheme. With the PCF, a contention-free period (CFP) and a contention period (CP) alternate periodically over time, where a CFP and the subsequent CP form an 802.11 superframe. The PCF is exploited throughout the CFP, while the DCF is used during the CP access phase. Each superframe is required to comprise a CP of a minimum length that allows at least one MSDU delivery (at least one frame exchange) of maximum size and at the slowest transmission rate under the DCF. A superframe is initiated by a beacon frame generated by the AP. The beacon frame is transmitted irrespective of whether or not the PCF is used. These frames are employed to preserve synchronization of the local timers in the associated stations and to deliver protocol-related parameters. The AP transmits these management frames at regular predefined intervals. Each station knows precisely when the subsequent beacon will arrive. These points in the time domain are referred to as target beacon transmission time (TBTT) and are announced in the previous beacon frame [1, 2].

During the CFP, there is no contention among wireless stations; instead, they are polled periodically by the AP. The PC polls a station requesting delivery of a pending data frame. Whenever the PC has a pending frame destined to an intended station, it utilizes a joint *data* and *poll* frame by piggybacking the CF-Poll frame onto the data frame. Upon reception of the so-called CF-Poll+Data, the polled station acknowledges the successful data reception and piggybacks an MPDU as well if it has any pending data frame targeted to the AP. If the PC does not receive a response from a polled station after waiting for a PCF interframe space (PIFS), it polls the next station or ends the CFP. Thus, no idle period longer than a PIFS occurs during a CFP. Bear in mind that a PIFS is longer than a SIFS but shorter than a DIFS. Since a PIFS is longer than an SIFS, a poll is never issued, e.g., between Data and ACK frames; hence a poll frame does not interrupt an ongoing frame exchange. The PC continues the aforementioned procedure until the CFP expires. A particular control frame, CF-End, is broadcast by the AP as the last frame within a CFP to indicate the end of the CFP (see Fig. 1.4) [1, 2, 4].

The PCF has many problems that have been reported in the literature [4]. Among many others, erratic beacon frame delay and indefinite transmission duration of the polled stations are the most important drawbacks. At the TBTT, the PC schedules the beacon as the next frame to be transmitted, but the beacon can only be transmitted when the medium has been determined to be

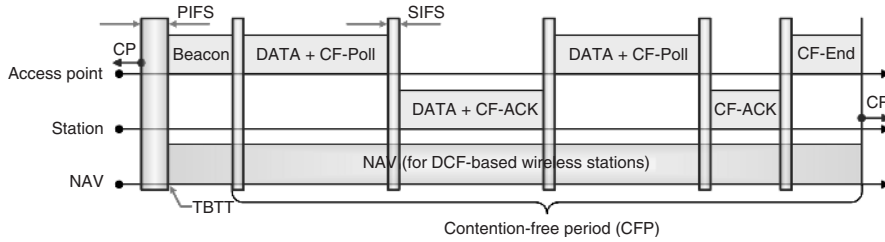


FIGURE 1.4 IEEE 802.11 PCF access scheme and TBTT.

idle for at least one PIFS. In IEEE 802.11, wireless stations are able to start their channel access even if the MSDU delivery is not finished before the upcoming TBTT. Depending upon whether the shared medium is idle or busy at the TBTT, a delay of the beacon frame might take place. The time the beacon frame is delayed from the TBTT determines the delay in a time-bounded MSDU transmission that has to be delivered in the CFP. This may rigorously influence the QoS as it introduces unpredictable time delays in each CFP. A further problem with the PCF is the unknown transmission duration of polled stations. A station that has been polled by the PC is allowed to deliver an MSDU that may be fragmented and of arbitrary length. In addition, different modulation and coding schemes are specified in the IEEE 802.11 family. Therefore, the duration of the MSDU is not under the control of the PC, which degrades the QoS offered to other stations polled during the rest of the CFP.

1.2.3.3 IEEE 802.11e: QoS Support in IEEE 802.11 MAC. The HCF, introduced in IEEE 802.11e, consists of two fundamental components: enhanced distributed-channel access (EDCA), an HCF contention-based channel access mechanism, and HCF controlled-channel access (HCCA). EDCA is the primary and mandatory access mechanism of IEEE 802.11e, while HCCA is optional and requires centralized polling and advanced scheduling schemes to distribute shared network resources among associated stations. According to the IEEE 802.11e, there can be two separate phases of operation within a superframe: CP and CFP. EDCA is used in the CP only, while HCCA is used in both phases. The HCF combines access methods of both the PCF and DCF, and this is why it is called *hybrid* [3].

The wireless station that operates as the central coordinator within a QoS-supporting basic service set (QBSS) is called a hybrid coordinator (HC). Similar to the PC, the HC resides within an 802.11e AP (i.e., QoS enabled access point (QAP)). There are *multiple* backoff entities operating in *parallel* within one QoS-aware 802.11e station (QSTA). A QSTA that is granted medium access opportunity should not occupy the radio resources for a time duration longer than a prespecified limit. This important characteristic of the 802.11e MAC protocol is referred to as transmission opportunity (TXOP). A TXOP is the

time interval during which a backoff entity has the right to deliver MSDUs and is defined by its starting time and duration. TXOPs obtained throughout the contention-based phase are referred to as EDCA-TXOPs. Alternatively, a TXOP obtained via a controlled medium access scheme is called an HCCA-TXOP or *polled* TXOP. The duration of an EDCA-TXOP is limited by a QBSS-wide parameter referred to as the *TXOPlimit*. This parameter is distributed regularly by the HC within an information field of the beacon frame. A further enhancement is that backoff entities of QSTAs are totally forbidden from transmitting across the TBTT. That is, a frame transmission is commenced only if it can be completed ahead of the upcoming TBTT. This reduces the expected beacon delay, which gives the HC superior control over the wireless media, especially if the noncompulsory CFP is exploited after the beacon frame. Moreover, an 802.11e backoff entity is allowed to exchange data frames directly with another backoff entity in a QBSS without involving communication with the QAP. Whereas within an 802.11-based infrastructure BSS all data frames are either sent or received by the AP, an 802.11e QSTA can establish a direct link with another 802.11e QSTA using the direct-link protocol (DLP) prior to initiating direct frame transmissions. It should be noted that here the backoff entity deals with the local backoff entity of a tagged QSTA; therefore, they are used interchangeably [3–5].

1.2.3.3.1 IEEE 802.11e: EDCA. In EDCA, QSTAs have up to four distinct and parallel queues for incoming traffic. Each queue is coupled with a specific access category (AC) and contends for the radio channel independent of the others. Collisions among a tagged station's queues are resolved internally, allowing the higher priority queue to commence its transmission while forcing the lower priority queue(s) to perform a collision response.⁴ Different levels of service are provided to each AC through a combination of three service differentiation mechanisms: arbitrary interframe space (AIFS), CW size, and TXOPlimit [3].

In contrast to the DCF access rules by which the backoff procedure is started after the DIFS from the end of the last indicated busy medium, EDCA backoff entities start at different intervals according to the corresponding AC of the traffic queue. As already pointed out, these time intervals are called AIFSs. The time duration of the interframe space AIFS[AC] is given by

$$\text{AIFS[AC]} = \text{SIFS} + \text{AIFSN[AC]} \times \text{aSlotTime}$$

where $\text{AIFSN[AC]} \geq 2$. Note that AIFSN[AC] should be chosen by the HC such that the earliest access time of 802.11e stations to be the DIFS, equivalent to IEEE 802.11. Note that the parameter aSlotTime defines the duration of a

⁴ In the literature, the internal collision between independent backoff entities is called *virtual collision*.

time slot. The smaller AIFSN[AC] corresponds to the higher medium access priority. The minimum size of CW, i.e., $CW_{\min}[AC]$, is another parameter which depends upon the AC. The initial value for the backoff counter is a random number taken from an interval defined by CW, which is exactly similar to the DCF case. Again, the smaller $CW_{\min}[AC]$ corresponds to a higher priority in acquiring a shared radio channel. An important difference between the DCF and 802.11e EDCA in terms of the backoff countdown rule is as follows: (1) The first backoff countdown takes place at the end of the AIFSN[AC] interval. (2) A frame transmission is initiated after a slot from the moment the backoff counter becomes zero. The CW increases upon unsuccessful frame exchanges but never exceeds the value of $CW_{\max}[AC]$. This parameter is defined per the AC as well as part of the EDCA parameter set. Note that the smaller $CW_{\max}[AC]$ corresponds to a higher medium access priority. The aforementioned system configurations, in addition to the EDCA parameter set, are illustrated in Fig. 1.5.

As mentioned earlier, in addition to the backoff parameters, the TXOPlimit[AC] is defined per the AC as part of the EDCA parameter set. Apparently, the larger TXOPlimit[AC] is, the larger is the share of capacity for this AC. Once a TXOP is obtained, the backoff entity may keep transmitting more than one MSDU consecutively during the same TXOP up to a duration of TXOPlimit[AC]. This important concept in 802.11e MAC is referred to as *continuation* of an EDCA–TXOP (4).

As already explained, four self-governing backoff entities with different EDCA parameter sets exist inside an 802.11e QSTA. In a tagged QSTA when counters of two or more backoff entities reach zero simultaneously, they

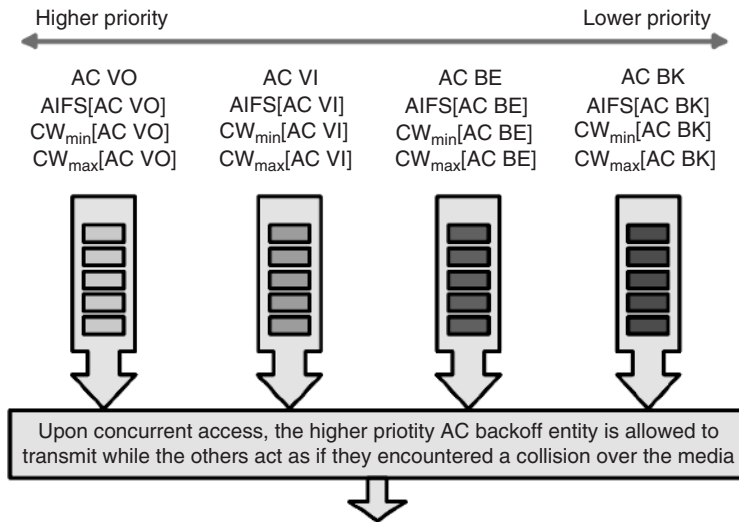


FIGURE 1.5 Four different access categories within a QSTA.

perform channel acquisition in the same time slot and consequently an internal *virtual collision* occurs. It should be noted that virtual collision is an abstract concept and there is no physical collision between contending backoff entities. When internal virtual collision occurs, the AC with the highest priority among collided entities is allowed to transmit, whereas all other backoff entities will act as if a collision has taken place on the shared radio channel.

1.2.3.3.2 IEEE 802.11e: HCCA. HCCA extends the EDCA medium access rules by assigning the uppermost precedence to the HC for the duration of both the CFP and CP. Basically, a TXOP can be attained by the HC through the controlled medium access stage. The HC may apportion TXOPs to itself in order to commence MSDU transactions whenever it requires, subsequent to detection of the shared wireless medium as being idle for PIFS, and without performing any further backoff procedure. To grant the HC a superior priority over legacy DCF and its QoS-aware counterpart, EDCA, AIFSN[AC] should be chosen such that the earliest channel acquisition for all EDCA stations can be the DIFS for any AC. During the CP, each TXOP of a QSTA begins either when the medium is determined to be available under the EDCA rules, i.e., after AIFS[AC] plus the random backoff time, or when a backoff entity receives a polling frame, the QoS CF-Poll, from the HC. The QoS CF-Poll is transmitted by the HC following a PIFS idle period and without any backoff procedure. On the other hand, for the duration of the CFP, the starting time and maximum duration of each TXOP is also specified by the HC, again by the use of QoS CF-Poll frames. In this phase, 802.11e backoff entities will not attempt to acquire the wireless media without being explicitly polled; hence, only the HC can allocate TXOPs by transmitting QoS CF-Poll frames or by immediately transmitting downlink data. Throughout a polled TXOP, the polled candidate mobile station can transmit multiple frames with a SIFS time space between two consecutive frames as long as the entire frame exchange duration does not exceed the dedicated maximum TXOPlimit. The HC controls the maximum duration of EDCA-TXOPs within its QBSS by the beacon frames. Thus, it is able to assign polled TXOPs at any time during the CP and the optional CFP [3].

Two supplementary schemes, namely block acknowledgement (BA) and DLP, which enhance the performance of the MAC protocol, have been taken into consideration in IEEE 802.11e [3, 4]. With the noncompulsory BA, the throughput efficiency of the protocol is improved. BA allows a backoff entity to send a number of MSDUs during one TXOP transmitted without individual ACK frames. The MPDUs delivered during the time of TXOP are referred to as a *block* of MPDUs in the literature and technical documents [4]. At the end of each block or in the next TXOP, all MPDUs are acknowledged at once by a bit pattern transmitted in the BA frame, and consequently the overhead of the control exchange sequences is reduced to a minimum of one ACK frame. On the other hand, each backoff entity is able to directly exchange information with any other backoff entity in the same QBSS without communicating

through the QAP. For IEEE 802.11 and within a BSS, all data frames are sent to the AP and received from the AP. However, it should be obvious that this procedure consumes at least twice the channel capacity in comparison to direct communication. For that reason, DLP is defined to enable pairs of 802.11e backoff entities to establish direct links between each other.

1.3 IEEE 802.11 PHYSICAL LAYER FAMILIES

In this section we first introduce the concepts utilized in radio-based 802.11 physical layers and then present detailed explanations of these physical layers.

The IEEE 802.11 physical layer is divided into two sublayers: the PLCP and the physical medium dependent (PMD). The PLCP receives incoming MSDUs from the MAC layer, adds its own designated header, and then gives them to the PMD. It is mandatory for delivered information to have a *preamble*, which has a pattern that depends on the modulation technique deployed in the physical layer. The PMD is responsible for transmitting every bit it receives from the PLCP over the wireless medium. The physical layer also incorporates a clear-channel assessment (CCA) function to inform the MAC layer when a carrier is detected [2]. Figure 1.6 illustrates the logical structure of the physical layer.

Three different physical layers were standardized in the initial revision of 802.11: frequency-hopping spread spectrum (FHSS), DSSS, and infrared (IR) light. Consequently, supplementary amendments 802.11a, 802.11b, and 802.11g were developed which are based on OFDM, high rate (HR)/DSSS, and the extended-rate PHY (ERP), respectively. Also, it is noteworthy to mention that 802.11n will be based on multi-input multi-output (MIMO) OFDM [6, 7].

In telecommunications, avoiding interference is a matter of law and the most imperative issue that should be taken into account. Thus, an official authority should impose rules on how the radio frequency (RF) spectrum is to be deployed. In the United States, the Federal Communications Commission (FCC) is responsible for regulating the use of the RF spectrum. European regulation is accomplished by the European Radio-communications Office (ERO) and the European Telecommunications Standards Institute (ETSI). The Ministry of Internal Communications (MIC) regulates radio exploitation in

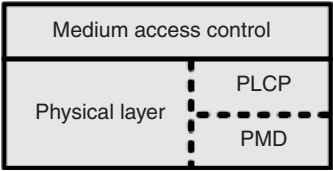


FIGURE 1.6 Physical layer logical structure.

Japan. Finally, worldwide regulation is done based upon recommendations of the International Telecommunications Union (ITU) [2].

The radio spectrum is partitioned into distinct frequency bands dedicated to particular applications. Among all frequency bands, the FCC and its counterparts in other countries designated particular frequency bands for the use of *industrial, scientific, and medical* (ISM) equipment. For instance, the 2.4-GHz band is available worldwide for unlicensed use. The use of RF equipment in the ISM bands is usually license free and RF devices operating in such frequency bands typically do not emit significant amounts of radiation. For example, microwave ovens are high-powered home/office devices, but they have extensive shielding to restrict interfering radio emissions. WLAN equipment, i.e., spread-spectrum-based IEEE 802.11b and OFDM-based IEEE 802.11g, as well as Bluetooth, spread-spectrum cordless phones, and X10⁵ communication system was developed for the 2.4-GHz ISM band [2]. On the other hand, in addition to 2.4 GHz, the 5-GHz frequency band is another ISM spectrum band dedicated to OFDM-based IEEE 802.11a; The United States was the first country to allow unlicensed device use in the 5-GHz range, though both Japan and Europe followed.

1.3.1 IEEE 802.11 Spread-Spectrum-Based Physical Layers

In this section, we start the discussion of the 802.11 SS-based physical layer with a short introduction of the system concept of different spread-spectrum architectures; then, FH 802.11 and the popular 802.11b extension based on DSSS and HR/DSSS will be covered separately.

1.3.1.1 Overview of Spread Spectrum. Spread-spectrum techniques are methods by which energy generated at one or more discrete frequencies is deliberately spread or distributed in either the frequency or time domain. This is accomplished for a variety of goals, including establishing secure communications, increasing resistance to natural interference and jamming, and preventing detection. Spread-spectrum telecommunications is a signal-structuring technique that utilizes direct sequence (DS), FH, or a hybrid of these and can be used for multiple access and/or multiple functions. This technique reduces the potential interference to other receivers while achieving *privacy*. Spread spectrum generally makes use of a sequential noiselike signal structure (i.e., spreading code) to broaden the narrowband information signal over relatively wideband radio frequencies. The intended receiver correlates the

⁵ X10 is an international and open industry standard for communication among devices used for home automation and domestic. It primarily uses power line wiring for signaling and control, where the signals involve brief RF bursts representing digital information. A radio-based transport is also defined. X10 was developed in 1975 by Pico Electronics of Glenrothes, Scotland, in order to allow remote control of home devices and appliances. It was the first *domotic* technology and remains the most widely available. *Domotics* is the application of computer and/or robotic technology to household appliances and buildings.

received signals to retrieve the original information signal. Originally, there were two motivations regarding the spread-spectrum concept: to resist efforts to jam radio communications and to hide the fact that communication is taking place, sometimes called low probability of intercept (LPI). Ultra wideband (UWB) is another modulation technique that accomplishes the same purpose based on transmitting short duration pulses. IEEE 802.11 uses either FHSS or DSSS in its radio interface. FH systems jump from one frequency to another in a random pattern, transmitting a short burst at each subchannel. The 2-Mbps FH physical layer is specified in clause 14. On the other hand, direct-sequence systems spread the power out over a wider frequency band using mathematical coding functions. Two DS structures were specified. The initial specification in clause 15 standardized a 2-Mbps physical layer, and 802.11b added clause 18 for the HR/DSSS physical layer.

1.3.1.1.1 Frequency-Hopping Spread Spectrum. FHSS is a method of transmitting radio signals by rapidly switching a carrier⁶ among many frequency channels using a pseudorandom⁷ sequence known to both transmitter and receiver. FH is similar to the well-known frequency division multiple access (FDMA) but with a key difference. In FDMA systems, devices are allocated fixed orthogonal nonoverlapping frequencies (i.e., fixed while totally distinct center frequencies and bandwidths). On the other hand, in FH-based systems, the frequency is time dependent; each frequency is used for a short portion of time, i.e., the so-called *dwelt time*.

If two FH systems want to share the same frequency band, both of them can be configured such that they utilize different hopping sequences so that they do not interfere with each other. For the duration of each time slot, the aforementioned hopping sequences should be on dissimilar frequency slots. As long as the systems stay on different frequency slots, they do not encounter any interference due to the other party. In general, orthogonal hopping sequences maximize wireless network throughput while increasing system complexity.

Beacon frames on FH networks include a timestamp and the so-called FH Parameter Set element. The FH Parameter Set element includes the hop pattern number and a hop index. By receiving a Beacon frame, a station knows everything it needs to synchronize its hopping pattern.

⁶ A *carrier wave*, or simply *carrier*, is a waveform that is modulated to represent the information to be transmitted. This carrier wave is usually of much higher frequency than the baseband modulating signal, i.e., the signal which contains the information.

⁷ A *pseudorandom* process is a process that appears stochastic but is not. Pseudorandom sequences typically exhibit statistical randomness while being generated by an entirely deterministic causal process. Such a process is easier to produce than a genuine random one and has the benefit that it can be used again and again to produce exactly the same numbers, useful for testing and fixing software. To date there is no known method to produce true randomness. The random-number generation functions provided in many software packages are pseudorandom.

Finally, it is noteworthy to mention that adaptive frequency hopping (AFH) spread spectrum, as used in Bluetooth, enhances system resistance to RF interference by avoiding using crowded frequencies in the hopping sequence. This sort of adaptive modulation is much easier to implement with FHSS than with DSSS.

1.3.1.1.2 Direct-Sequence Spread Spectrum. DS transmission is an alternative spread-spectrum technique that might be utilized to transmit a narrow-band signal over a much wider frequency band. The fundamental approach of DS schemes is to cautiously spread the RF energy over a wide frequency band. The DS modulation scheme is accomplished by applying a chipping sequence (CS) to the information bit stream. A *chip* is a binary digit sequence employed by the DS system. These bits are higher level data while the chip signals are binary numbers used in encoding (transmitter side) and decoding (receiver side) procedures. Chipping streams, or the so-called pseudorandom noise (PN) codes, should have a much higher rate in comparison to the actual data stream.

For the PN code, IEEE 802.11 adopted an 11-bit Barker word meaning that each bit is encoded using the entire Barker word as a CS. Barker words have satisfactory autocorrelation, implying that the correlation function at the receiver operates as expected in a wide range of environments and is relatively tolerant to multipath delay spreads as incurred in multipath fading channels. The philosophy behind the deployment of exactly 11 bits is that most regulatory authorities usually necessitate a 10-dB processing gain in DS systems.

1.3.1.2 IEEE 802.11 FH Physical Layer. In 802.11 FH, the microwave ISM band is partitioned into a series of 1-MHz channels. Approximately 99% of the radio energy is confined to the channel. The modulation scheme employed by 802.11 encodes data bits as shifts in the transmission frequency from the channel center. Channels are defined by their center frequencies, which begin at 2.400 GHz for channel 0. Successive channels are derived by adding 1-MHz steps, meaning that channel 1 has a center frequency of 2.401 GHz, channel 2 has a center frequency of 2.402 GHz, etc., until channel 95, which has a center frequency at 2.495 GHz. Different regulatory authorities allow use of different parts of the ISM band. For example, the FCC in the United States and the ETSI in Europe (excluding France and Spain) allow channels 2–79 to be deployed while, in Japan, channels 73–95 might be utilized [1, 2].

The dwell time (see Section 1.3.1.1.1) in 802.11 FH systems is 390 time units, which is about 0.4 s. When an 802.11 FH physical layer hops between channels, the hopping process should take no longer than 224 μ s. The frequency hops are subject to extensive regulation, in terms of both the size of each hop and the rate at which hops must occur [2].

1.3.1.3 IEEE 802.11b Physical Layer. The IEEE 802.11b amendment to the original standard was ratified in September 1999. 802.11b has a maximum

raw data rate of 11 Mbps and uses the same CSMA/CA media access method defined in the original standard.

IEEE 802.11b products appeared on the market very quickly, since 802.11b is a direct extension of the DSSS modulation technique defined in the original standard. The 802.11b standard uses complementary code keying (CCK) as its modulation technique, which is a variation on code division multiple access (CDMA). Hence, chipsets and products were easily upgraded to support the 802.11b enhancements. The dramatic increase in throughput of 802.11b (compared to the original standard) along with substantial price reductions led to the rapid acceptance of 802.11b as the definitive WLAN technology.

Generally, 802.11b is used in point-to-multipoint configuration, wherein an AP communicates via an omnidirectional antenna with one or more clients that are located in a coverage area around the AP. Typical indoor range is 30 m (100 ft) at 11 Mbps and 90 m (300 ft) at 1 Mbps. With high gain external antennas, the protocol can also be used in fixed point-to-point arrangements, typically at ranges up to 8 km (5 miles), although some report success at ranges up to 80–120 km (50–75 miles) where line of sight (LOS) can be established. This is usually accomplished in place of costly leased lines or very cumbersome microwave communications equipment.

Channels for the 802.11b DS physical layer are much larger than the channels for the FH physical layer. The DS physical layer has 14 channels in the 2.4-GHz band, each 5 MHz wide. Channel 1 is placed at 2.412 GHz, channel 2 at 2.417 GHz, and so on, up to channel 13 at 2.472 GHz. Channel 14 was defined for operation in Japan and has a center frequency that is 12 MHz from the center frequency of channel 13 [2].

1.3.2 IEEE 802.11 OFDM-Based Physical Layers

In this section, we begin our discussion of 802.11 OFDM-based physical layers with a qualitative introduction to the basis of OFDM; subsequently, different 802.11 extensions, including 802.11a, g, h, and j, are covered separately.

1.3.2.1 Overview of OFDM. OFDM, essentially identical to coded OFDM (COFDM), is a digital multicarrier modulation scheme which deploys a large number of closely spaced orthogonal *subcarriers*. Each subcarrier is modulated with a conventional modulation scheme, e.g., quadrature amplitude modulation (QAM), at a low symbol rate, maintaining data rates similar to conventional *single-carrier* modulation schemes in the same bandwidth. In practice, OFDM signals are generated by the use of the fast Fourier transform (FFT) algorithm. The most important advantage of OFDM over single-carrier schemes is its ability to cope with severe channel conditions, e.g., multipath and narrowband interference, without complex equalization filters. Channel equalization is simplified due to the fact that OFDM may be viewed as using many slowly modulated narrowband signals rather than one rapidly modulated wideband signal. The orthogonality of the subcarriers results in zero cross-talk,

even though they are so close that their spectra overlap. Low symbol rate helps manage time domain spreading of the signal by allowing the use of a guard interval between successive symbols. The guard interval eliminates the need for a pulse-shaping filter.

OFDM necessitates a high level of accuracy in frequency synchronization between the receiver and transmitter. In other words, any deviation causes the subcarriers to no longer be orthogonal, resulting in intercarrier interference (ICI) and cross-talk between adjacent subcarriers. Frequency offsets are typically caused by mismatched transmitter and receiver oscillators or by *Doppler* shift due to mobile device movement. While Doppler shift alone may be compensated for by the receiver, the situation is worsened when combined with multipath, as reflections will appear at various frequency offsets, which is much harder to correct. This effect typically worsens as speed increases and is an important factor limiting the use of OFDM in high speed vehicles. Several techniques for ICI suppression have been suggested, but they may increase receiver complexity as well.

A key principle of OFDM is that, due to low symbol rate modulation, symbols are relatively longer than the channel time characteristics. As a result, it suffers less from intersymbol interference (ISI) caused by multipath. Since the duration of each symbol is long enough, it is feasible to insert a guard interval between the consecutive OFDM symbols, thus eliminating ISI. In addition, the guard interval also reduces the sensitivity to time synchronization problems. The *cyclic prefix*, which is transmitted during the guard interval, consists of the end of the OFDM symbol copied into the guard interval, and the guard interval is transmitted followed by the OFDM symbol. The guard interval consists of a copy of the end of the OFDM symbol so that the receiver will integrate over an integer number of sinusoid cycles for each multipath when it performs OFDM demodulation with FFT.

The stochastic effects due to channel frequency selectivity might be considered to be constant over an OFDM subchannel if the subchannel is sufficiently narrowbanded, i.e., if the number of subchannels is adequately large. This important feature makes equalization far simpler at the receiver side in OFDM systems than in conventional single-carrier modulation schemes. The equalizer simply has to multiply each subcarrier by a constant value, or even a rarely variable value. In addition, some subcarriers in some OFDM symbols might carry *pilot* signals for measurement of channel conditions, i.e., the equalizer gain for each subcarrier. In addition, pilot signals might be used for synchronization. If a differential modulation technique such as differential phase shift keying (DPSK) or differential quadrature phase shift keying (DQPSK) is applied to each subcarrier, equalization can be completely omitted, since these schemes are insensitive to slowly changing amplitude and phase distortion.

OFDM has been perpetually exploited in conjunction with channel coding schemes, i.e., forward error correction (FEC), and almost always uses frequency and/or time interleaving. On the one hand, frequency (subcarrier)

interleaving increases resistance to channel frequency selectivity. On the other hand, time interleaving ensures that bits that are initially close together in the bit stream are transmitted far apart in time, thus mitigating against severe fading, as would happen when traveling at high speed. However, time interleaving is of little benefit in slowly fading channels while frequency interleaving offers little to no benefit for narrowband channels that suffer from flat fading. The reason interleaving is used in OFDM is to attempt to spread the errors out in the bit stream presented to the error correction decoder. A common type of error correction coding scheme used with OFDM-based systems is *convolutional* coding, which is often concatenated with *Reed–Solomon* coding. Convolutional coding is used as the inner code and Reed–Solomon coding is used for the outer code, usually with additional interleaving on top of the time and frequency interleaving and between the two layers of coding. The motivation for this combination of error correction coding is that the Viterbi decoder used for convolutional decoding produces short error bursts when there is a high concentration of errors, and Reed–Solomon codes are inherently well suited to correcting bursts of errors.

Finally, it should be noted that *windowing* is another technique which helps OFDM-based transceivers cope with real-world effects. Transitions can be abrupt at symbol boundaries, causing a large number of undesired high frequency components. To make OFDM transmitters robust against the aforementioned problems, it is widespread to add *padding* bits at the beginning and end of transmissions to allow transmitters to ramp up and down from full power. Padding bits are frequently required when error correction coding is deployed. In the literature, padding streams are usually referred to as *training* sequences.

1.3.2.2 IEEE 802.11a/h/j 5-GHz Physical Layer. The IEEE 802.11a amendment to the original standard was ratified in September 1999 [8]. Basically, it utilizes the same core protocol as the original standard, operates in the 5-GHz band, and uses a 52-subcarrier OFDM with a maximum raw data rate of 54 Mbps, which yields realistic net achievable throughput in the mid-20 Mbps. The data rate is reduced to 48, 36, 24, 18, 12, 9, and then 6 Mbps if required. 802.11a has 12 nonoverlapping channels, 8 dedicated to indoor and 4 to point to point. It is not interoperable with 802.11b, except if using equipment that implements both standards. Due to the fact that the 2.4-GHz frequency band has been heavily deployed, operating in the 5-GHz band gives 802.11a the advantage of less interference. However, this high carrier frequency also brings disadvantages. It restricts the use of 802.11a to almost LOS, necessitating more AP deployment [2, 9].

In IEEE 802.11a, out of 52 OFDM subcarriers, 48 are for data and 4 are pilot subcarriers with a carrier separation of 0.3125 MHz (20 MHz/64). Each of these subcarriers can be binary phase shift keying (BPSK), quadrature phase shift keying (QPSK), 16-QAM, or 64-QAM. The total bandwidth is 20 MHz with an occupied bandwidth of 16.6 MHz and symbol duration is 4 μ s with a

guard interval of $0.8\text{ }\mu\text{s}$. The generation and decoding of orthogonal components are done in baseband using digital signal processing (DSP), which is then up converted to 5 GHz at the transmitter. Each subcarrier could be represented as a complex number. The time domain signal is generated by taking an inverse fast Fourier transform (IFFT). Correspondingly, the receiver down converts samples at 20 MHz and does an FFT to retrieve the original coefficients. The advantages of using OFDM include reduced multipath effects in reception and increased spectral efficiency [9].

In 802.11a, channels in the 5-GHz band are numbered starting every 5 MHz and each 20-MHz 802.11a channel occupies four channel numbers. Basically, 802.11a was originally designed for the United States. European channelization was added as part of 802.11h in late 2003, and subsequently Japanese operation was appended with 802.11j in late 2004 [2].

1.3.2.3 IEEE 802.11g 2.4-GHz Physical Layer. In June 2003, a third modulation standard was ratified: 802.11g. This extension exploits the 2.4-GHz frequency band (similar to 802.11b) but operates at a maximum raw data rate of 54 Mbps, or about 24.7 Mbps net throughput, similar to 802.11a. IEEE 802.11g hardware is compatible with its 802.11b counterpart. The modulation scheme used in 802.11g is OFDM for the data rates of 6, 9, 12, 18, 24, 36, 48, and 54 Mbps, and, on the one hand, reverts to CCK to achieve 5.5 and 11 Mbps while, on the other hand, switches to DBPSK/DQPSK + DSSS for 1 and 2 Mbps. The maximum range of 802.11g devices is slightly greater than that of 802.11b devices, but the range in which a client can achieve the full 54-Mbps data rate is much shorter than that which a 802.11b client can reach, 11 Mbps.

1.3.3 IEEE 802.11n Physical Layer

The emerging IEEE 802.11n specification differs from its predecessors in that it provides for a variety of optional modes and configurations that dictate different maximum raw data rates. This enables the standard to offer baseline performance parameters for all 802.11n devices while allowing manufacturers to enhance or tune capabilities to accommodate different applications and price points. With every possible option enabled, 802.11n could offer raw data rates up to 600 Mbps [6, 7].

In fact, the most widely celebrated component of 802.11n is the inclusion of MIMO technology. MIMO harnesses multipath transmission with a technique known as space division multiplexing. The wireless terminal basically splits a data stream into multiple concurrent divisions, called spatial streams, and delivers each one of them through separate antennas to corresponding antennas on the receiving end. The current 802.11n draft provides for up to four spatial streams, even though compliant hardware is not required to support that many. Doubling the number of spatial streams from one to two effectively doubles the raw data rate. There are trade-offs, however, such as increased power

consumption and, to a lesser extent, cost. The 802.11n specification will include a MIMO power-saving mode which mitigates power consumption by using multiple paths only when communication would benefit from the additional performance. The MIMO power-saving mode is expected to be a compulsory feature in the ratified IEEE 802.11n final specifications. A non compulsory mode in 802.11n that effectively doubles the offered data rate is deployment of the communication channel of 40 MHz bandwidth. The main trade-off here is less channel availability for other mobile stations. In the case of the 2.4-GHz frequency band, there is enough room for three nonoverlapping 20-MHz channels. Needless to say, a 40-MHz channel does not leave much room for other devices to join the network or transmit in the same airspace. This means that dynamic radio resource management is critical to ensure that the 40-MHz channel option improves the overall system performance by balancing the high bandwidth demands of some clients with the needs of other clients to remain connected to the network [2].

Note that in the MAC layer IEEE 802.11n offers BAs similar to the concepts recommended in the 802.11e amendment. By removing the need for one ACK for every data frame, the amount of overhead required for the ACK frames, as well as preamble and framing, is considerably reduced. BAs are helpful, but only if all the frames in a burst can be delivered without any problem. Missing one frame in the block or losing the ACK itself carries a steep penalty in protocol operations since the entire block must be retransmitted again. In addition, in 802.11n MAC, frame aggregation is also expected to be a mandatory MAC entity component. Combining several small layer 3 packets into a single relatively large frame improves the data-to-overhead ratio. Frame aggregation is often used with MAC header compression, since the MAC header on multiple frames to the same destination is quite similar.

1.4 SUMMARY AND CONCLUDING REMARKS

In this chapter, a brief overview of existing MAC and physical layers in WLANs was provided. Starting from the MAC layer, we reviewed the well-known IEEE 802.11 and its QoS-aware amendment 802.11e. We pursued our discussion of the IEEE 802.11 physical layer, and different types of physical layers with diverse system architectures and modulation schemes were explored.

ACKNOWLEDGMENTS

This work was partially supported by Nokia Foundation and Elisa Foundation.

REFERENCES

1. IEEE 802.11 WG, Part 11: "Wireless LAN medium access control (MAC) and physical layer (PHY) specification," IEEE, New York, Aug. 1999.
2. M. Gast, *802.11 Wireless Networks: The Definitive Guide*, 2nd ed., O'Reilly Media Inc., 2005.
3. IEEE 802.11e, Part 11: "Wireless LAN medium access control (MAC) and physical layer (PHY) specifications: amendment 8: Medium access control (MAC) quality of service enhancements," supplement to IEEE 802.11, IEEE, New York, Nov. 2005.
4. S. Mangold, S. Choi, G. R. Hiertz, O. Klein, and B. Walke, "Analysis of IEEE 802.11e for QoS support in wireless LANs," *IEEE Wireless Commun.* 10(6), 40–50 (2003).
5. Y. Xiao, "IEEE 802.11e: a QoS provisioning at the MAC layer," *IEEE Wireless Commun.* 11(3), 72–79 (2004).
6. Y. Xiao, "IEEE 802.11N: enhancements for higher throughput in wireless LAN," *IEEE Wireless Commun.* 12(6), 82–91 (2005).
7. Y. Xiao, "Efficient MAC strategies for the IEEE 802.11n wireless LANs," *Wireless Commun. Mobile Comput.* 6(4), 453–466 (2006).
8. IEEE 802.11b, Part 11: "Wireless LAN medium access control (MAC) and physical layer (PHY) specifications: High-speed physical layer extension in the 2.4GHz band," supplement to IEEE 802.11, IEEE, New York, Sept. 1999.
9. IEEE 802.11a WG, Part 11: "Wireless LAN medium access control (MAC) and physical layer (PHY) specification: high-speed physical layer in the 5GHz band," IEEE, New York, Sept. 1999.

