

CHAPTER 1

INTRODUCTION TO OPTICAL PACKET SWITCHED (OPS) NETWORKS

The demand for high bit-rate traffic is increasing rapidly each year, especially by Internet users [1]. Among different networking mechanisms, optical networks can support this overwhelming traffic demand. In optical networking, various amounts of traffic with different bit rates can be multiplexed into a single fiber and switched in the network. Optical fiber communication can offer much higher bandwidth than conventional copper wires used extensively as the transmission medium. In addition, optical fibers can provide low communication cost per kilometer. Moreover, signals transmitted on optical fibers encounter lower bit error rate than copper wires. In this chapter, an introduction is provided for optical networking, especially Optical Packet Switched (OPS) networks. The OPS networking can solve the mismatch between very high transmission capacity of WDM optical links and the processing power of routers/switches [2].

1.1 Optical Fiber Technology

Optical fibers are essentially very thin glass cylinders filaments that carry data and control signals in the form of light (i.e., optical signals). There are two common

fibers as single-mode and multi-mode as the basics of optical transmission. Single-mode fiber has a relatively small core diameter of about 8 to 10 μm . Using single-mode fiber effectively eliminates intermodal dispersion and enables a significant increase in the bit rates and distances between nodes. Using this fiber, a network needs a typical regenerator/amplifier spacing of about 40 km and it can operate at bit rates of a few tens gigabits per second. However, the distance between regenerators/amplifiers is primarily limited by fiber loss. The standard Single-Mode Fiber (SMF) is a common single-mode fiber used in optical technology. These fibers provide low dispersion and can be appropriate for long distances such as regional and long haul networks. On the other hand, multi-mode fibers have a relatively large core diameter of about 50 to 62.5 μm . These fibers have high dispersion and are only suitable for short distances such as metro networks.

Whereas in electronic networks different multiplexing techniques such as Time Division Multiplexing (TDM) and Frequency Division Multiplexing (FDM) are employed to use transmission resources efficiently, Wavelength Division Multiplexing (WDM) is used in optical networks as an effective technique to make use of the large amount of available bandwidth in optical fibers for applications with huge bandwidth demands. The technology of using multiple optical signals on the same fiber is called WDM. In WDM, several baseband-modulated channels are transmitted along a single fiber at the same time, but each channel is located at a different wavelength. The most important components of any WDM system are optical transmitters, optical receivers, optical amplifiers, optical switches, optical add drop multiplexers, WDM multiplexers, and WDM demultiplexers. All-optical networks using WDM appear to be the sole approach for transporting huge network traffic in future core networks.

In WDM optical networks, data are carried on wavelengths. Consider that there are W different wavelengths on a fiber, where each wavelength is able to carry data at different data rates such as $B = 2.5 \text{ Gbits/s}$, 10 Gbits/s , 40 Gbits/s , and even more. The aggregate system has a capacity of $W \times B$, and therefore the amount of total carried data on one optical fiber can increase up to several Tb/s. Therefore, WDM networks can be employed in Wide Area Networks (WANs) for the future backbone networks of Internet. Nowadays, WDM primarily uses the 1.55- μm wavelength region because of the inherent loss in optical fiber which is the lowest in this region; excellent optical amplifiers have been designed for this region. In Dense Wavelength-Division Multiplexing (DWDM), many wavelengths (say 160) are packed densely into a fiber with small channel spacing.

When selecting a fiber for an optical network design, two important issues must be considered: attenuation and dispersion. Attenuation is the optical signal loss when it travels through a fiber, whereas dispersion is the tendency of different wavelengths to travel with different speeds, which leads to pulse spreading [3].

Attenuation in an optical fiber depends on the wavelength an optical signal uses (see Fig. 1.1) [4]. As can be observed, not all wavelengths are suitable for transmission of an optical signal. Four principal windows of wavelength ranges with low attenuation are displayed in this figure, called O-band (displayed with first window), S-band (displayed with second window), C-band (displayed with third window), and L-band (displayed with fourth window). Optical lasers deployed in optical commu-

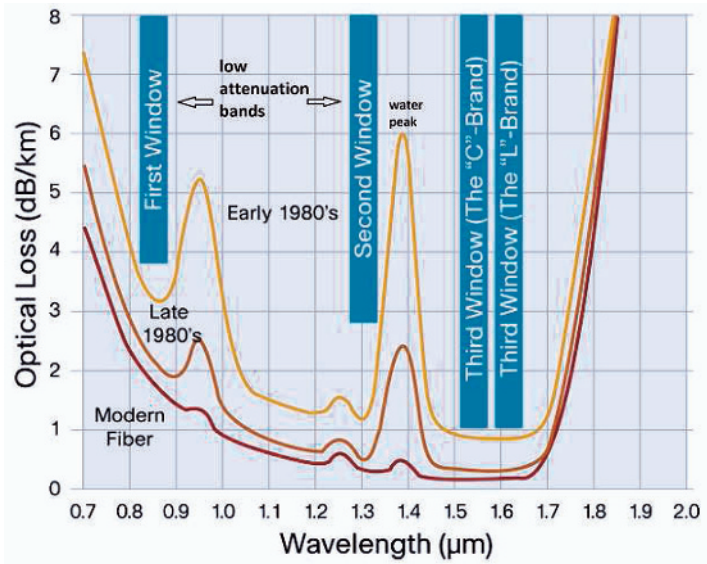


Figure 1.1: Attenuation versus wavelength curve [4]

communications typically operate at these bands. Note that in general, there are seven bands standardized for the attenuation-wavelength curve [3, 5]:

- 850 nm: wavelength range 770 nm to 910 nm. This range was used in the first generation of optical networks with opaque switches such as SONET/SDH and gigabit Ethernet networks.
- O-band (Original band) (known as 1310 nm): wavelength range 1260 nm to 1360 nm. This range is used for short range communications (say upstream wavelength in PON). The O-band provides high attenuation of 0.5 db/km, but no dispersion in standard single-mode fibers.
- E-band (Extended band): wavelength range 1360 nm to 1460 nm. This is the water-peak band that can only be used in modern fibers as they have reduced attenuation in this range. This range is only suitable for short-range communications.
- S-band (Short band): wavelength range 1460 nm to 1530 nm. This range is used for short-range communications (say downstream wavelength in PON).
- C-band (Conventional band) (known as 1550 nm): wavelength range 1530 nm to 1565 nm. This range has the lowest attenuation and was used for early WDM communications.
- L-band (Long band): wavelength range 1565 nm to 1625 nm. This range provides low attenuation and is used for WDM communications.

- U-band (Ultra-long band): wavelength range 1625 nm to 1675 nm. This range is used for WDM communications.

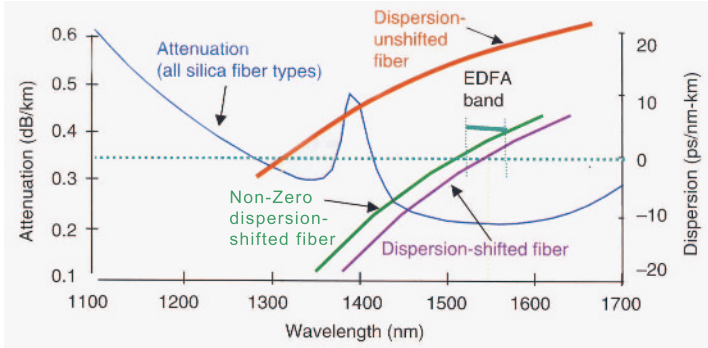


Figure 1.2: The SSF, DSF, and NZDSF fibers [6]

Chromatic dispersion in single-mode fibers causes pulse spreading because of the fact that various wavelengths travel at different speeds. When optical signals spread too far, they overlap and cannot be correctly detected at the receiving end of the network. Similar to attenuation, dispersion is a function of wavelength (see Fig. 1.2). Since Standard Single-Mode Fibers (SSMF), displayed with dispersion unshifted fiber in Fig. 1.2, provide a zero dispersion at wavelength 1310 nm, 1310 nm transmitters are not subject to chromatic dispersion. On the other hand, at 1550 nm, CWDM and DWDM transmissions over SSMF are affected by chromatic dispersion. In short, SSMF provides high attenuation at 1310 nm compared with other bands, but zero dispersion in this range, whereas it provides the lowest attenuation at 1550 nm, but high dispersion at this range [4, 6].

In order to reduce the dispersion encountered by transmissions within the C-band and L-band windows in SSMF, other fiber types have been developed. Dispersion-Shifted Fibers (DSF) provide a zero dispersion at 1550 nm (see Fig. 1.2). DSF with low loss in the L-band range is suitable for DWDM applications. Although DSF eliminates the dispersion problem for transmissions of single wavelengths at 1550 nm, it is still not appropriate for wavelength multiplexing applications since WDM transmissions can be affected by another non-linear effect called four-wave mixing. This has led to the development of Non-Zero Dispersion Shifted Fibers (NZDSF), where the zero dispersion is shifted just outside the C-band, around 1510 nm. This can limit the chromatic dispersion (as the zero dispersion remains close enough to the transmission band) and four-wave mixing. The NZDSF fiber is optimized for long-haul applications at 1550-nm range in terms of both attenuation and dispersion. It is suitable for distances over 70 km. There are two types of NZDSF known as $(-D)$ NZDSF (called negative NZDSF) and $(+D)$ NZDSF (called positive NZDSF) that can provide a negative and positive slope versus wavelength, respectively. Both positive and negative NZDSF are fine over distances of up to 200 km in the C-band

at data rate 10 Gbits/s. However, for 40 Gbits/s transmission rate, positive NZDSF is recommended [3, 4, 6].

To reduce the dispersion problem and improve optical signal transmission, there are various fiber types developed such as Dispersion-Compensated Fiber (DCF), Dispersion-Flattened Compensated Fiber (DFCF), Dispersion-Slope Compensated Fiber (DSCF), and Dispersion-Shift Compensated Fiber (DSCF) [6].

Optical fibers and components (such as amplifiers and filters) cause several kinds of impairments for optical signals. This results in signal quality degradation in optical receivers. To compensate signal loss, optical amplifiers must be used as essential components in transmission systems and networks. The most common optical amplifier could be the Erbium-Doped Fiber Amplifier (EDFA) operating in the C-band. In addition, L-band EDFAs and Raman amplifiers can be used. EDFAs can be used in almost all WDM networks, whereas Raman amplifiers can be used in addition to EDFA in many ultra-long-haul optical networks. The main advantage of an EDFA amplifier is that it can simultaneously amplify many WDM channels. To support high-capacity DWDM transmission, DCF fibers could be placed in between amplification stages [7, 8].

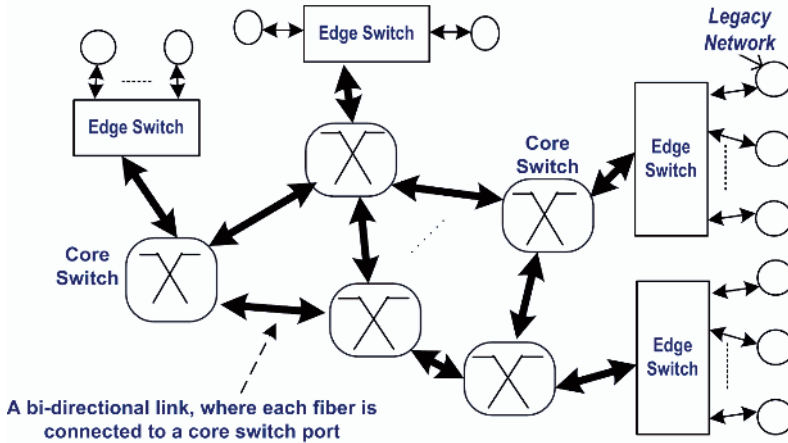


Figure 1.3: General optical network model

1.2 Why Optical Networks?

The Internet traffic demand is increasing 70 to 150 percent each year [9]. The need for high bandwidth and high data rate services are also increasing. Almost all communication applications tend to change their networks to IP; even voice applications use IP networks to carry their voice traffic. Traffic demands of various Internet applications (such as video conferencing, real-time medical imaging, emergency services, online gaming) continue to grow and this will consume more and more network bandwidth [10, 11]. In addition, current bandwidth-limited networks mostly sup-

port the best effort traffic, and therefore no differentiation can be easily made for real-time traffic. Moreover, network traffic patterns have a bursty nature, and not a uniform characteristic. Therefore, core networks in WANs are becoming bottleneck. It seems that the sole approach to overcome this huge demand and resolve the bottleneck is to deploy optical networks. Thus, the optical networking is the key technology for future communication networks.

The optical and electronic networks have essential differences in switching speed, buffering architecture, and bandwidth granularity issues. In the optical domain, switching speed is slower and building optical buffers is a more complex and a more expensive issue than electronic networks. On the other hand, an optical network can provide a higher bandwidth, a better signal quality, and a better security than an electronic network. Considering these major differences, different architectures and bandwidth management protocols must be used to utilize the huge bandwidth offered by all-optical networks [6, 12].

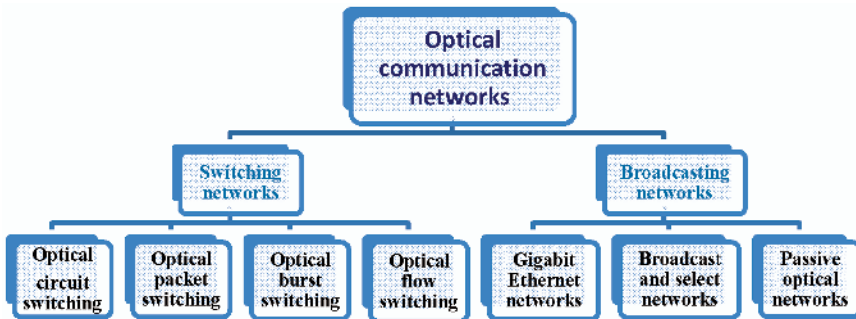


Figure 1.4: Hierarchy of optical communications networks

As Fig. 1.4 shows, optical communications networks can be divided into switching and broadcasting networking similar to electronic networks. Some common switching networks developed for data transmission over WANs are Optical Circuit Switching (OCS) networks [also called Wavelength Routed Networks (WRNs)], Optical Packet Switching (OPS) networks, Optical Burst Switching (OBS) networks, and Optical Flow Switching (OFS) networks. Similar to electronic domain in which packet switching is the most granular method of switching, the most promising technique for optical core networks could be OPS due to its high throughput and very good granularity and scalability. However, contention of optical packets in the core network is the major problem in OPS networks. On the other hand, gigabit Ethernet networks, broadcast and select networks, and Passive Optical Networks (PON) are common broadcasting networks developed for Local Area Networks (LAN) that broadcast data toward all destinations.

1.3 Optical Networking Mechanisms

There are two common mechanisms to share the bandwidth in an optical network: reservation-based and contention-based. Reservation-based mechanisms always reserve bandwidth in advance, say for a connection request, and are mostly suitable for (semi-) static traffic. The decision made in a reservation-based scheme can be either centralized or distributed. In centralized, only one central module is responsible for bandwidth provisioning, while a number of modules cooperate with each other to reserve bandwidth in the distributed manner. On the other hand, contention-based schemes perform no reservations and are appropriate for the networks with dynamic traffic. A contention-based scheme relies on a random access technique in order to access the network. Due to their lower complexity, they are easily scalable and can respond more quickly to bursty traffic than the reservation-based schemes. Contention-based schemes may suffer from collision, whereas reservation-based schemes may experience connection blocking.

An optical network is interconnected with a number of optical switches, each called a core switch (see Fig. 1.3). One edge switch operating in the electronic domain is connected to each optical switch. When an edge switch has the function of transmitting traffic to the optical network, it is called ingress switch. On the other hand, when an edge switch has the function of receiving traffic from the optical network, it is called egress switch.

The packets arriving at an ingress switch from legacy networks are called client packets. In this book, a legacy network denotes any network in the electronic domain including an old network, an Ethernet network, a SONET/SDH network, and a TCP/IP network. By buffering client packets in the electronic domain and forwarding them to the optical network, an edge switch provides interfacing between the electronic legacy networks and the optical network. The architecture of an optical network could be either single-fiber or Multi-Fiber (MF). In a single-fiber optical network, there is only one fiber used in each direction between two adjacent nodes (i.e., two fibers between two adjacent nodes). The adjacent nodes could be either (a) two optical switches or (b) an edge switch and an optical switch. On the other hand, in a multi-fiber optical network, there are say f fibers used in each direction between two adjacent nodes (i.e., $2 \times f$ fibers between two adjacent nodes), but with a small number of wavelength channels per fiber compared with the single-fiber architecture. For example, if a single-fiber network needs 48 wavelengths per fiber, then a multi-fiber architecture with $f = 3$ fibers on each connection link requires $W = 16$ wavelengths per fiber, i.e., $f \times W$ is constant.

A core switch size is displayed with say $N * N$, i.e., it has N input ports/links and N output ports/links. Note that term input/output port is used in single-fiber networks, and term input/output link is used in multi-fiber networks, where each link has f ports.

A multi-fiber architecture has a numbers of advantages. It can reduce connection blocking in circuit switching optical networks and traffic loss in optical packet/burst switching networking. In addition, the devices (such as dispersion compensator modules) required for a fiber with a large number of wavelengths are more expensive

than the devices used for a fiber with a small number of wavelengths [13–16]. For example, as Fig. 1.5 illustrates, for a fiber link between any two optical switches, Distributed Raman Amplifiers (DRA) with 82-km spacing can be employed. Each amplification span includes 70 km of SMF fiber whose dispersion and dispersion slope are compensated by 12 km of Dispersion Compensating Fibers (DCF). Finally, a multi-fiber architecture can provide a better protection against a single-fiber cut than can a single-fiber architecture [14].

Core switches used in an optical network may have different sizes according to the topology of the network. For example, at $f = 2$, one core switch may need 6 input ports and 6 output ports, the other core switch may need 10 input ports and 10 output ports, and so on. However, an optical switch must have a size of $N * N$, where we have $N = 2^x$ and x is an integer number. Therefore, some ports of the core switch may be unused. These unused ports can be connected to the relevant edge switch in order to reduce the receiver contention problem, i.e., to make an almost non-blocking receiver (detailed in Chapter 3).

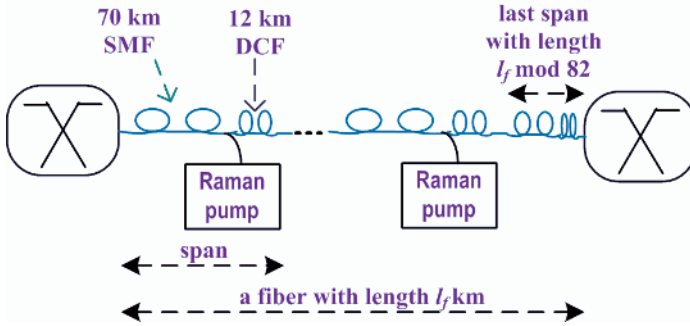


Figure 1.5: Link model in an optical network

Optical networks are categorized as all-optical (transparent) networks, opaque networks, and translucent networks [17–19]:

- **Opaque Networks:** Every optical switch, equipped with O/E/O (optical to electronic to optical) converters, converts an optical signal to the electronic domain, processes or saves it in the electronic domain, and then regenerates and converts it to the optical domain. In other words, opaque networks only apply optical fibers as their transmission medium, and processing and switching are carried out in the electronic domain. Electronic equipments used in O/E/O optical switches limit the data rate in opaque networks and high bandwidths cannot be achieved in such networks.
- **Transparent Networks:** An all-optical network uses a transparent optical signal transmission without any conversion of its data traffic into the electronic domain in its core switches, while it may process the header of data traffic in the electronic domain. Note that transparent optical switches are not only cheap but can also completely utilize the bandwidth offered by DWDM. Real-

izing all-optical networks needs very low-loss fibers, switches, and multiplexers/demultiplexers.

- **Translucent Networks:** In translucent optical networks, some optical switches are transparent and some other are opaque with O/E/O capability in such a way that a transmitted signal stays in optical domain in most of the core switches.

Different optical networking mechanisms have been proposed for managing network bandwidth, where each mechanism is based on a switching technique. In the following, a number of switching schemes used in synchronous (also called slotted) and asynchronous (pure) optical networks are studied. A survey of optical switching schemes can be easily found in [12, 20–24].

1.3.1 Asynchronous Optical Switching

In this section, different switching mechanisms introduced for asynchronous (pure) optical switching, where wavelength channels are not time-slotted.

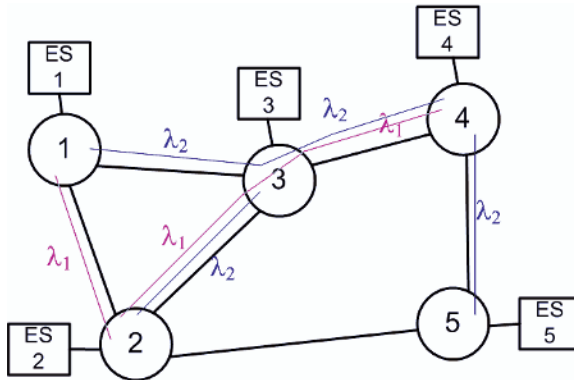


Figure 1.6: Light-path setup example in a WRN optical network

1.3.1.1 Optical Circuit Switching To transmit data in a Wavelength Routed Network (WRN), the OCS mechanism (as a two-way reservation-based mechanism) is used in which an optical connection, called light-path, must be set up between a source–destination pair using a Routing and Wavelength Assignment (RWA) technique prior to data transmission [25–30]. A light-path is setup on a wavelength channel in a path that may span a number of fiber links. Then, the entire bandwidth on this light-path is used for a connection until its termination. Since light-path setup is the most important operation to make a connection request between a pair of nodes, its effective establishment should be taken into account. If an RWA algorithm cannot establish a light-path due to lack of network resource, the relevant connection request will be blocked. For example, Fig. 1.6 shows a network in which four light-paths have been set up: (1) a light-path from Edge Switch (ES) 1 to ES 4 on

wavelength λ_2 crossing core switches 1, 3, 4; (2) a light-path from ES 1 to ES 4 on wavelength λ_1 passing through core switches 1, 2, 3, 4; (3) a light-path from ES 2 to ES 3 on wavelength λ_2 crossing core switches 2, 3; and (4) a light-path from ES 4 to ES 5 on wavelength λ_2 crossing core switches 4, 5.

Two types of traffic are considered for OCS networks. Under static traffic [31, 32], a set of fixed connection requests is given and the network should follow two main objectives in establishing these connections: (1) maximizing the number of established connections when the number of wavelengths is limited and (2) minimizing the number of wavelengths needed to set up the connections. Under dynamic traffic [33], a connection request arrives based on a random process with a random holding time, and therefore a light-path is established and then terminated after elapsing its connection holding time. Therefore, RWA decisions should be made rapidly when a connection request arrives at the network. Due to insufficient resources or unavailable wavelengths, the network may not be able to find a light-path for a given connection request, thus resulting in connection blocking. Therefore, the main objective of dynamic RWA is to find a route and choose a wavelength that maximizes the probability of establishing each connection request and at the same time attempt to minimize the blocking probability for future connections. Obviously, it is impossible to keep the resource utilization optimal in dynamic traffic.

There are three types of routing approaches used in RWA: fixed routing, fixed alternate routing, and adaptive routing. In fixed routing, there is only one fixed route (e.g., the shortest path or the shortest hop) between a pair of ingress and egress switches. The shortest path/hop for each source–destination pair is calculated in an offline manner using the Dijkstra or Belman–Ford algorithm. This is the simplest form of routing, but it may block a connection request if there is no feasible wavelength along the shortest route found for the request. In fixed-alternate routing, each ingress switch maintains a routing table that includes an ordered list of fixed routes to each egress switch (say the first-shortest-path route, the second-shortest-path route, the third-shortest-path route, and so on). The actual route for a connection can only be chosen from these fixed routes. In adaptive routing, any route between a source–destination pair can be used for a connection request provided that it has available wavelength resources. Adaptive routing techniques increase the probability of establishing a connection by using networks state information. Under adaptive routing, a routing decision is made dynamically based on the current wavelength usage information on each link (such as the shortest-cost path first or the least-congested path first). For example, under the least-congested path first adaptive routing, a path is selected that has the highest number of available wavelengths among other paths because the links that have a small number of available wavelengths are considered more congested [34].

After finding a route between a source–destination pair, a wavelength assignment algorithm should select an available wavelength that maximizes wavelength utilization. There are a number of heuristic wavelength assignment approaches used in RWA such as (1) the first-fit method that attempts to select the first available wavelength in numerical order, (2) the most used method that attempts to allocate the most utilized wavelength first, (3) the least used method that attempts to allocate the least

utilized wavelength first, (4) the best-fit method that picks a wavelength on which the light-path fits best (i.e., among the wavelengths available for a connection request, a wavelength is chosen that has the least free capacity remaining after accommodating the connection request), and (5) the random method that attempts to allocate a wavelength randomly. Under static traffic, the objective of wavelength assignment is to minimize the number of wavelengths, whereas in dynamic traffic it is assumed that the number of wavelengths is fixed and the network tries to minimize connection blocking [25, 35].

Wavelength continuity is a problem in WRN networks, where a connection request may have to be rejected (i.e., blocked) even though a route is available for the request because of non-availability of the same wavelength on all the links along the route. Connection blocking is increased when the traffic load goes up, thus reducing the performance of RWA. In summary, connection blocking is more likely to occur under the following cases:

- under dynamic traffic rather than static because there is no information about future connections.
- in the networks without Wavelength Converters (WCs) because of the wavelength continuity constraint. A wavelength converter is a device that can convert the wavelength of an incoming optical signal to another wavelength.
- lack of available wavelengths in the network.
- when the number of hops (path length) is high. In a large-diameter network, it will be harder to find the same wavelength free on long paths.
- when the network global information is not used in routing.
- when optical switches are not strictly non-blocking. Different architectures of strictly non-blocking optical switches have been studied in [36] such as Clos architecture, Cantor architecture, and so on. Note that a strictly non-blocking switch is a sort of switching module in which an unused input port can always be connected to an unused output port, without having to rearrange existing connections inside the switch. It should be noted that a Clos switch becomes strictly non-blocking under a specific condition. There are three switching stages in an $N \times N$ Clos switch. The first stage uses $\frac{N}{m}$ switches, each with size $m \times k$. There are k switches in the intermediate stage, each with size $\frac{N}{m} \times \frac{N}{m}$. Each of these k switches is connected to any switch in the first and the third stages. The third stage uses $\frac{N}{m}$ switches, each with size $k \times m$. A Clos switch becomes strictly non-blocking if $k \geq 2 \times m - 1$ [37].

The wavelength continuity problem can be resolved using wavelength converters, strictly non-blocking switches, and light-path rerouting. Using wavelength conversion, a light-path may use several wavelengths along its path when traversing through different fiber links. Light-path rerouting means the action of changing the physical

path and/or the wavelength(s) of an established light-path. Rerouting is a viable and cost-effective mechanism to improve the connection blocking performance [38, 39].

Bandwidth wastage and scalability issues are two drawbacks for OCS [20]. OCS suffers from low channel utilization since a single connection may not employ the whole bandwidth of a light-path. To resolve this drawback, grooming techniques are used to highly utilize the bandwidth of a light-path [40]. In grooming, a number of low bit rate connections are mixed and sent on a setup light-path. Traffic grooming performs a two-layer routing to effectively pack low-rate connections onto high-rate light-paths.

OCS networks perform two-way reservation for establishing a connection request before data transmission. The network can be thought of as a two-plane architecture consisting of a control plane and a data plane. One wavelength channel on each link is used for exchanging control information. The actual data are transmitted in the data plane.

There are two main signaling methods for light-path setup in WRN as centralized RWA and distributed RWA. In distributed RWA signaling [41, 42], each network element is in charge of light-path setup. Under distributed RWA, all connection requests are processed at different network core switches and each core switch makes its decision based on its available resources. In centralized RWA signaling [43], a central network element reachable by all network core switches is in charge of establishing all connection requests. The central element is aware of the complete network topology, physical parameters, and resource availability, thus leading to better performance results than the distributed RWA. However, the central element failure is a big concern. The centralized RWA could be suitable for static traffic and small networks, not for dynamic traffic because of much information that must be managed by the central element. Centralized RWA results in scalability and complexity problems for dynamic traffic. This is why distributed RWA is much more popular than centralized RWA. Distributed control can improve the scalability and reliability performances. A distributed RWA uses two strategies to find a suitable wavelength on a given route R between source node S and destination node D as follows [28, 44], where the purpose of both strategies is to find a continuous wavelength from source node S toward destination node D over a given physical path (i.e., no wavelength conversion is used in these algorithms):

- **Forward Reservation Method (FRM):** In FRM, node S reserves one of the unused wavelengths over its ongoing link toward the next node. For wavelength reservation, node S sends a message toward the next node over route R . The node that receives this message repeats the process of reserving the selected wavelength in other nodes until the message arrives at node D . If this reservation is blocked due to the unavailability of the selected wavelength along route R , node S will select another free wavelength and repeat the process.
- **Backward Reservation Method (BRM):** For a given connection request, node S initiates a wavelength set that contains all free wavelengths of its ongoing link towards the next node over route R . Every intermediate node modifies this set based on the intersection of the incoming wavelength set and the set of free

wavelengths of its ongoing link toward node D. The final set that reaches node D is the set that contains unused wavelengths over the whole path from S to D. Node D then processes this incoming set for finding an appropriate free wavelength over route *R*. If there is no wavelength in the final set, the call will be blocked. After these processes, node D initiates the reservation process in a backward manner towards node S.

Many studies on RWA have concentrated on establishing light-paths under the assumption of an ideal physical layer. However, this assumption could only be suitable for opaque networks, where a signal is regenerated at each intermediate optical switch. As an optical signal propagates along a light-path toward its destination in a transparent wavelength-routed optical network, the signal quality degrades since there is no conversion to the electronic domain and therefore no signal regeneration. This in turn increases Bit Error Rate (BER) of the signal. However, high BER is not acceptable by users. In addition, it is not acceptable if the establishment of a light-path results in increasing the BER of other existing light-paths. Establishing light-paths with low BER can reduce the number of retransmissions by high layers, thus increasing network throughput. Therefore, the RWA techniques that consider physical layer impairments for the establishment of a light-path, called Quality of Transmission-Aware (QoT-aware) RWA, could be much more practical.

The following RWA surveys can be found in literature. A survey of pure RWA techniques (without considering transmission impairments) can be found in [25]. The work in [45] has reviewed the techniques (including pre-emption, wavelength management, and routing) that can be utilized for the differentiation of light-paths in the RWA process. The RWA techniques suitable for translucent networks have been reviewed in [46], where a translucent network employs electrical regenerators at some intermediate nodes only when it needs to improve the signal quality. A review on the management and control planes required for QoT-aware RWA techniques has been provided in [32, 47]. The surveys in [32, 48, 49] have provided a review on static (offline) physical layer impairment-aware RWA algorithms in all-optical networks. The surveys presented in [27, 32, 48] have studied in general different topics related to optical networking and physical layer impairments, including DWDM technology, physical layer impairments, optical components, service level agreements, failure recovery, static impairment-aware RWA techniques, impairment-aware control plane techniques, and some dynamic impairment-aware RWA algorithms in both translucent and all-optical networks.

1.3.1.2 Optical Packet/Burst Switching Due to its finer granularity, a packet switching technique (as a contention-based mechanism) can obtain a higher channel utilization and can yield better bandwidth efficiency than OCS by two approaches as discussed below [21, 50]. Comparison of OPS and OBS switching techniques can be easily found in [23, 51, 52].

- **Asynchronous Optical Packet Switching (OPS):** An OPS network could be the promising network to employ the huge bandwidth provided by optical fiber

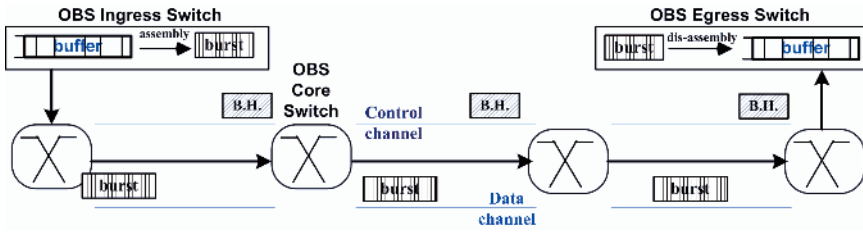


Figure 1.7: An example for OBS operation

technology. In asynchronous OPS, client packets are converted from the electronic domain to the optical domain in ingress switches (called optical packets), and then each optical packet is sent to the OPS network. An optical packet may include a single client packet or a number of client packets. In this switching, each optical packet is individually switched in a core network like its electronic counterpart while keeping payload in the optical domain and processing the optical packet header electronically [1, 2, 53]. Optical packets are switched to their desired ports as they arrive at an OPS switching node and there is no requirement for synchronization stages. In asynchronous OPS, client packets can have any size and thus are suitable for IP packets. Egress switches are the places for converting optical packets to electronic domain and sending them toward their final destination networks. The OPS network is necessary for both metro and backbone networks [21]. The OPS is also compatible with bursty nature of IP traffic and can increase throughput and efficiency of network bandwidth so that OPS needs less wavelengths and resources for handling the same traffic compared with an OCS network.

- **Asynchronous Optical Burst Switching (OBS):** This is a switching protocol [54, 55] with a finer granularity than OCS and a coarser granularity than OPS. An OBS network aggregates a number of client packets destined to a given egress switch in a burst, sends a Burst Header (BH) on a control channel for necessary resource reservation in the intermediate core switches along the egress switch, waits for an offset time, and then transmits the burst over an existing data wavelength without waiting for an acknowledgment from the egress switch. Each OBS core switch processes the BH and provides an output channel to forward the arriving burst while resolving the possible contention of the burst. A data burst in OBS can involve client packets with various data rates. A BH includes different information for the burst including its ingress switch address, its egress switch address, the number of client packets aggregated in the burst, the class of service for the client packets, the burst duration, the data wavelength on which the burst will be transmitted, and so on. Figure 1.7 depicts an example for OBS operation in which client packets are assembled in the ingress switch and dis-assembled in the egress switch. Since processing of the BH takes some time in each intermediate OBS core switch, the header and BH are getting closer to each other while they are at the egress switch. There-

fore, the offset time must take into account the processing time of the BH in all intermediate OBS core switches along the path to the egress switch.

Resource reservation in OBS can be implemented by two approaches. (1) In one-way reservation (called Tell and Go (TAG) signaling protocol), a control packet is transmitted by the ingress switch to allocate the necessary resources for a data burst. Then, the data burst follows its control packet without waiting for the reservation acknowledgment from the egress switch. This results in a high burst blocking rate, especially at high traffic load. However, the end-to-end delay is minimized. This type of reservation is ideal for delay-sensitive applications; and (2) In two-way reservation (called Tell and Wait (TAW) signaling protocol), a control packet is sent over the path of the burst in order to collect information on the availability of resources. At the egress switch, a resource assignment algorithm is executed. An acknowledgment packet is sent back to the intermediate OBS switches to reserve the necessary resources. If the reservation fails at an intermediate switch, a failure packet is sent to the egress switch. If the reply packet reaches the ingress switch, the burst data are transmitted. This eliminates the loss of data bursts, but it can also lead to high end-to-end delay.

In an ingress switch, an optical burst is created by the packet aggregation technique. The ingress switch is responsible for aggregating the client packets coming from legacy networks destined to a given egress switch into a burst. By receiving a burst, the egress switch unpacks all the client packets in the burst and then routes each packet individually toward its relevant destination. In each ingress switch, there is a dedicated queue to each egress switch. Packet aggregation can be either timer-based or threshold-based, or a combination of both timer-based and threshold-based. In the timer-based aggregation technique, a timer is started whenever a client packet arrives at the queue relevant to egress switch i and this queue is empty. The aggregation algorithm waits for an aggregation period and during this period, it collects all the client packets destined to egress switch i . A burst is created when timeout happens at the end of the aggregation period. On the other hand, in the threshold-based technique, when waiting traffic in a queue reaches a given threshold size, a burst can be formed. The aggregated traffic size in the burst must be equal to or less than the threshold size. Finally, in combined timer-based and threshold-based, a burst is created whenever there is enough traffic to fill a burst or whenever timeout occurs. Note that the timeout mechanism can limit the waiting time of packets in queues, but at the expense of creating small-size bursts [56–58].

There are two mechanisms for aggregation of class-based traffic. First, client packets of the same class of service can only be aggregated in a burst [58, 59]. Second, client packets from different classes of service can be aggregated in a burst (i.e., composite burst aggregation) [60]. Since the latter mechanism can aggregate bursts sooner than the former, it can lead to smaller end-to-end packet delay and burst loss rate than the former mechanism.

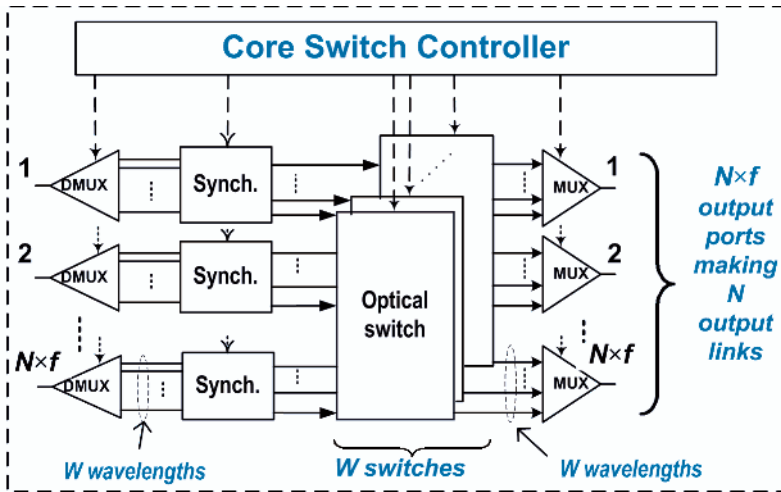


Figure 1.8: Synchronized multi-fiber OPS switch architecture

1.3.2 Synchronous (Slotted) Optical Switching

To provide the finest granularity and improve bandwidth usage, the Optical Time Division Multiplexing (OTDM) concept can be deployed in optical networks. Under OTDM, many source–destination pairs can share the network bandwidth. However, synchronization of traffic arrival at core switches is the limitation of this switching. In synchronized networks, fiber delay lines are required for synchronization issue at the input ports of core switches [53, 61].

1.3.2.1 Optical Circuit Switching In slotted OCS (as a reservation-based mechanism) [62–64], the bandwidth of a wavelength is divided into frames of fixed time slots and traffic for a given source–destination pair is periodically transmitted in pre-allocated time slot(s) in each frame. The Routing and Wavelength Assignment problem in the wavelength-routed networks is changed to the Routing, Wavelength, and Time-Slot Assignment problem in slotted OCS. Bandwidth is reserved for ingress switches by intermediate core switches in the optical network.

1.3.2.2 Optical Packet/Burst Switching There are different approaches to deploy the packet switching concept in slotted networks (as contention-based mechanisms):

- In conventional slotted OPS [53, 65], a fixed-length client packet together with a header makes an optical packet. Then, the optical packet is transmitted to the OPS network at a time slot boundary. A large-size client packet must be fragmented and transmitted in a number of time slots [2]. Optical packets are synchronized before entering an OPS switching module using switching delay lines. A slotted OPS network has a lower complexity in terms of switch control than an asynchronous OPS network [21]. In addition, it has a higher throughput

than asynchronous OPS due to a smaller contention rate in slotted OPS than in asynchronous OPS [21]. This is similar to slotted ALOHA in which the vulnerable period of packets is reduced, thus leading to a lower packet contention than with pure ALOHA. The simplest technique to manage wavelengths in this technique could be based on round robin. Under this scheme, a pointer in an egress switch points to the transmission wavelength to be used by the next incoming client packet. After sending the client packet in a time slot, the pointer is incremented by 1 modulo W , where W is the total number of wavelengths available in an ingress switch. Slotted OPS has lower complexity of switch control and lower optical Packet Loss Rate (PLR) than asynchronous OPS.

In slotted OPS, each time slot includes two time intervals. The optical packet interval for transmitting optical packet payloads is denoted by S_T ; and the time-gap interval plus the header interval for transmitting header information is denoted by S_O . Hence, time-slot size is equal to $S_{ts} = S_T + S_O$.

Variable length client packets cannot be efficiently utilized in slotted OPS networks because they may cause some bandwidth degradations in slotted OPS. For example, of the variable length IP packets, almost half of the packets are 40–44 bytes, almost 18% have the length of 1500 bytes, another 18% of packets are either 552 or 576 bytes, and a very small number of packets have a size larger than 1500 bytes [66]. Hence, the variance of the IP packet sizes is very high and choosing a time slot to carry 1500 bytes will result in a high bandwidth wastage. On the other hand, using a time-slot size to carry say 40 bytes requires not only a much faster optical switching speed, but also a faster processing at core switches to process a large number of packet headers in a very short time [21]. Clearly, this requirement will even be high when 40-Gbits/s or higher transmission rates are used in OPS networks. This issue also leads to the packet fragmentation problem that increases header complexity and induces cost due to the reconstruction of received packets at egress switches. The other problem with fragmentation is that when a part of an optical packet is dropped, the whole optical packet will be useless. To resolve the problems of the conventional slotted OPS network, one can use a relatively larger time slot in which a number of client packets of any type and class can be aggregated within a time slot before transmission to network.

In slotted OPS, a switch fabric can only be reconfigured at the start of each time slot for all input links. This needs alignment and synchronization of all optical packets coming from input ports according to a local clock. The reason for synchronization is due to the fact that different optical packets could arrive at an OPS switch at different point of times because of different distance of physical links, temperature fluctuations, and chromatic dispersion. Optical packets synchronization is the major disadvantage of slotted OPS.

- In slotted OBS [67], each burst is divided into multiple time slots and transmitted at fixed positions in a periodic frame structure. The control wavelengths carry burst header cells so that core switches along a path can set up a con-

nection. This architecture may need optical buffers in core switches in order to interchange time slots. The problem with the burst division is that the whole burst can be blocked due to the blocking of a small part of the burst, thus leading to an inefficient resource utilization.

- In Photonic Slot Routing (PSR) [68, 69], simultaneous slots (each slot may carry a number of client packets) transmitted on distinct wavelengths are aggregated in one photonic slot and routed through core switches. Since this technique does not require wavelength-selective optical switches, it is cost-effective. However, this network is not too flexible and the network bandwidth may not be fully utilized because the whole traffic of a fiber can only be switched to a specific output port of a core switch. Consequently, this switching usually finds applications in ring networks.

There is also a combined approach of asynchronous OPS and slotted OPS in which the network operation is asynchronous while switches operate synchronously. In this technique, the optical packets sent to the OPS network can have variable sizes and can arrive at a core switch at any time. However, the switching module inside an OPS switch begins to operate only in the start of time slots, where a time slot is the duration of transmitting the smallest packets size (say 40 bytes in IP networks). Here, large-size optical packets are transmitted in several consequent time slots. This approach provides the flexibility of having any packet size and makes easy contention resolution in OPS switches [70].

Note that OBS and OPS could be two dominant techniques for future optical core networks. Among the aforementioned switching schemes, OPS is not only scalable but also flexible and can dynamically allocate network resources with fine granularity [21]. OPS can efficiently use the network bandwidth that enables it to support different applications and services [71]. However, there are three implementation limitations for OPS [21]: (1) its inability to save optical data in memory with random access capability; (2) the lack of sophisticated processing in the optical domain, while high processing power is a requirement in the optical switches due to the larger amount of overhead; and (3) its requirement for fast switching speed (e.g., nanoseconds). For instance, in a slotted OPS with $B_c = 10$ Gbits/s, consider $S_T = 450$ ns and time gap $S_O = 50$ ns. In this example, one user packet of size 560 bytes can be carried within a time slot as an optical packet. Therefore, an optical switch must be able to perform header processing, contention resolution, and switching within 50 ns of the time gap.

To resolve the second and the third problems, a number of client packets should first be aggregated in a large-size optical packet and then transmitted to the network. This approach allows one to use relatively larger optical packet sizes and larger time gaps between optical packets (in slotted OPS), thus alleviating the need to use a very fast optical switch. This clearly reduces the complexity of OPS switches due to the lower number of entities per unit time to be processed [3]. For example, consider $S_T = 45,000$ ns and $S_O = 5000$ ns in slotted OPS. A time slot in this case can carry up to almost 56,000 bytes (usually more than one client packet). In addition, the

time for header processing, contention resolution, and switching increases to 5000 ns compared with 50 ns in the previous example.

1.4 Overview of OPS Networking

In the following, different aspects of OPS networks are detailed. Enabling technologies for OPS have been discussed in [65]. Generally speaking, an OPS core network is responsible for transferring clients' traffic inside optical packets between two edge switches. Each edge switch is in charge of collecting clients' traffic from legacy networks and delivering clients' traffic to the legacy networks.

OPS can be used in both connection-oriented and connectionless networks. In a connectionless OPS network, an optical packet can be routed toward its destination without path limitation through any path that an OPS switch decides. In connection-oriented OPS, each optical packet should follow the path determined in connection setup phase (i.e., virtual circuit). Each connection has a unique ID, and optical packets of that connection are switched in OPS switches based on this ID.

1.4.1 Network Topologies for OPS

An OPS networking was originally designed for regional/long-haul networks. However, nowadays, OPS is even desired for metro networks in order to efficiently use network resources for the following reasons [21]:

- A metro network should be able to provide more capacity to cope not only with the ever-increasing bandwidth demands of novel applications but also with the unexpected future demand growth [72, 73]. The simple reason for the huge amount of traffic in metro networks is the mass deployment of Fiber-To-The Home (FTTH) technologies in access networks [73].
- A metro network must provide different Quality of Service (QoS) levels for new applications, mostly based on the Internet.
- A metro network must provide agility in order to deploy bandwidth for traffic demands at a finer granularity since most of the traffic nowadays is Internet traffic.
- A metro network experiences high traffic dynamics due to Internet traffic.

Thus, all-optical packet-switched networks appear to be the sole approach to provide such capacity and agility. The general OPS network model has mesh topology similar to what is displayed in Fig. 1.3, suitable for regional/long-haul networks. However, for metro networks, star and ring topologies can also be used. The benefits and disadvantages of star topology over ring topology in metro networks have been stated in [74].

The metro architectures based on star topology could be based on three architectures:

- **Passive Broadcast-and-Select Star Couplers:** Although networks using passive star couplers have zero power consumption [26], they have three problems: (1) wavelength reuse problem [26], in which a wavelength cannot be used by several connections like in WRNs; (2) need for a large number of wavelengths to support even simple traffic patterns [75]; and (3) suffer from power loss problem due to splitting loss [26, 76, 77].
- **Central Passive Arrayed Waveguide Grating (AWG):** AWG can resolve the wavelength reuse and power loss problems. Thus, a highly efficient network architecture can be realized by using all wavelengths at all ports of an AWG simultaneously [72]. However, AWG acts like a static wavelength router. In AWG, an incoming to outgoing route is determined statically, where the routing pattern is a function of the incoming wavelength and input port [72]. In addition, the maximum number of wavelengths that can be used in the network is a function of the number of input ports of an AWG [74, 78].
- **Wavelength-Selective Cross-Connect Optical Switches (Known as Active Switches):** An active switch with the wavelength switching capability can dynamically switch each incoming wavelength from any input port to any output port. Obviously, an active switch can provide additional degree of freedom to route traffic by changing the setting of the switch module. Opposed to an AWG, the total number of wavelength channels to be used in the network based on an active switch is independent of the number of input ports of the switch and is easily scalable. This gives another freedom in designing a metro network based on an active switch. Therefore, by using an active switch, a more efficient and flexible star metro network can be realized. Using active switches, an OPS network can operate much more efficiently than an OPS network using passive switches due to the spatial wavelength reuse and splitting loss problems in passive components. Note that a power source must be present for controlling an active switch [74, 79–87].

The advantages of star topology in a metro network are: (1) simplicity in routing and configuration, (2) reduced complexity of control plane, (3) design simplicity, (4) good scalability as the network size grows physically since it is easy to add an edge switch to a star, (5) easy synchronization required in slotted OPS networks, and (6) reduced switch crosstalk because of having only one core switch along a transmission path compared to ring and mesh topologies [79, 88]. The main disadvantage of the single-hop network is its core switch failure. To provide robustness in a star network, the overlaid star topology (see Section 6.1.1) with a number of overlaid core switches can be used, where every edge switch is connected to each one of the core switches. Using this architecture, an edge switch can easily reroute its traffic to another star when a core switch fails. The bandwidth sharing mechanisms proposed for single-hop overlaid OPS networks will be detailed in Chapter 6 (see Section 6.1.3).

Different architectures have been proposed for metro networks based on the ring topology, e.g., [74, 89, 90], where optical packets are routed over a WDM ring network. Here, Media Access Control (MAC) protocols are proposed to enable the OPS

switches on the WDM ring to share the bandwidth of the network while preventing optical packet collisions. These protocols will be detailed in Chapter 6.

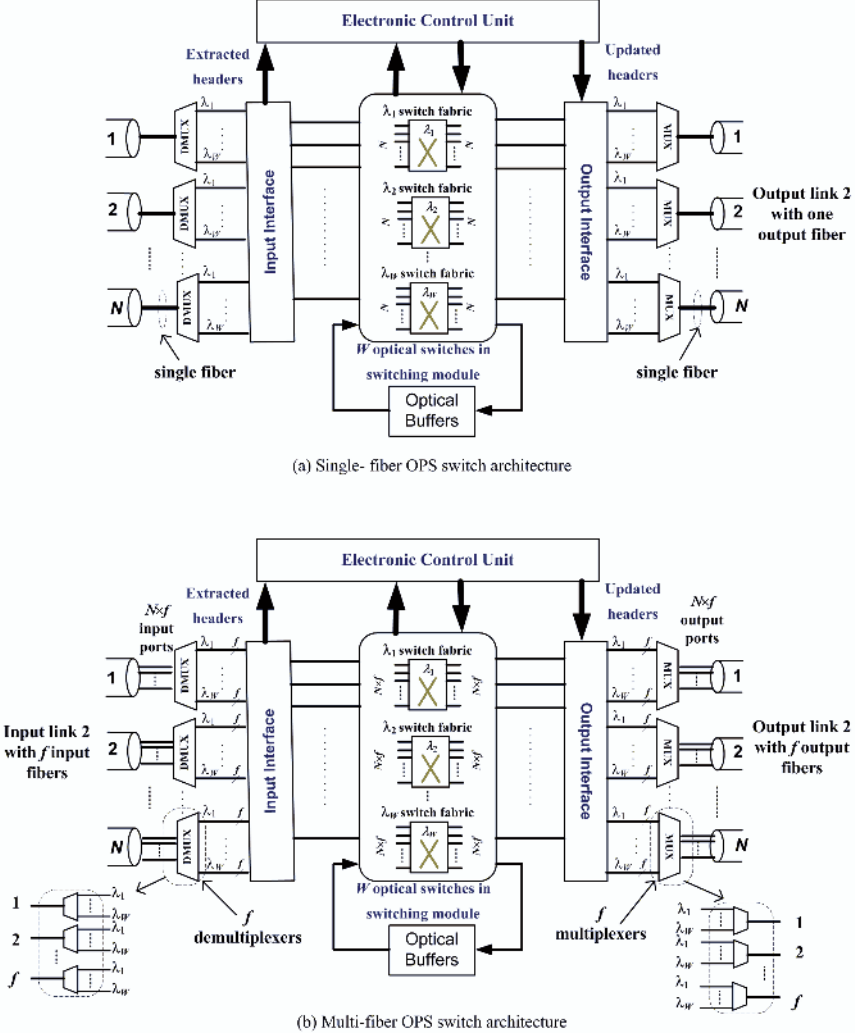


Figure 1.9: General OPS switch architecture

1.4.2 Core Switch Architecture in OPS Networks

An overview on the operation of OPS networks is discussed in this section. Figure 1.9a shows a generic single-fiber OPS switch architecture with N input ports and N output ports (i.e., an $N \times N$ OPS switch) and also shows W wavelength channels available per fiber [65, 70]. It includes different components as N demultiplexers,

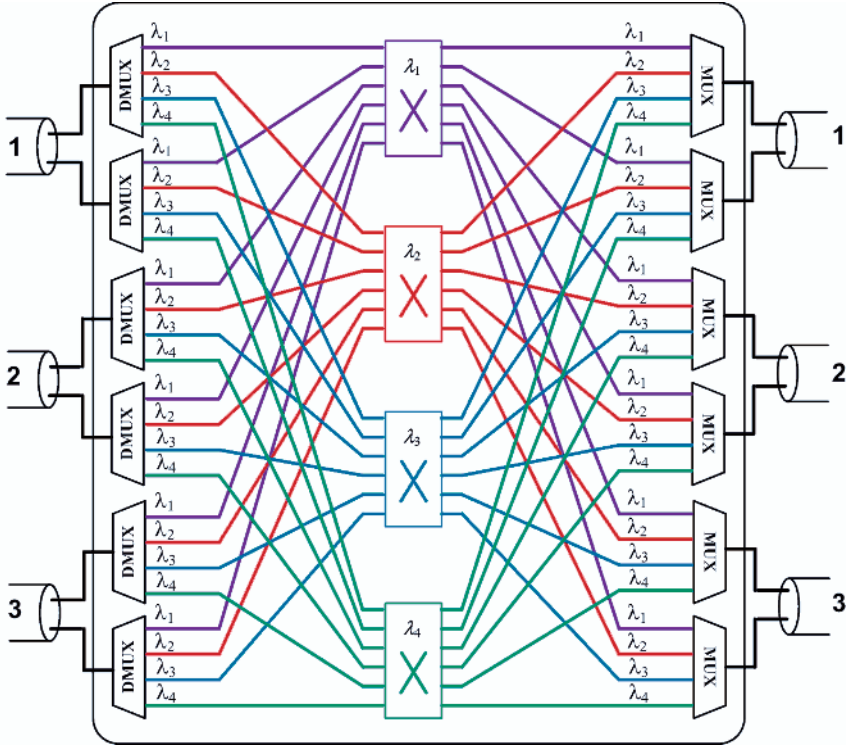


Figure 1.10: An example for switching module in an OPS core switch at $N = 3$, $f = 2$, and $W = 4$

an input interface, a switching module, an output interface, N multiplexers, optical buffers for delaying optical packets, and an electronic control unit. It should be mentioned that some efforts have been made to design the control unit in the optical domain [91]. The control unit keeps a routing table based on the shortest path mechanism in order to route optical packets destined to a given egress switch through the output links. In an OPS switch, the switching module includes a number of non-blocking wavelength-selective cross-connect optical switches. Here, it is considered that the control unit processes headers in the electronic domain. On the other hand, some efforts have been made to process an optical packet header in the optical domain [91–94]. An OPS switch will be referred to as a core switch from now on. On the other hand, Fig. 1.9b illustrates a generic multi-fiber OPS core switch architecture with $N \times f$ input ports and $N \times f$ output ports (i.e., an $(N \times f) \times (N \times f)$ OPS core switch), along with W wavelength channels available per fiber. It includes different components such as $N \times f$ demultiplexers, an input interface, a switching module, an output interface, $N \times f$ multiplexers, optical buffers for delaying optical packets, and an electronic control unit.

Each demultiplexer separates W wavelengths of each input fiber. The input interface detects the start and end of an optical packet header and payload (of at most

$N \times W$ optical packets in the single fiber architecture and of at most $N \times f \times W$ optical packets in the multi-fiber architecture), converts the header to the electronic domain, and sends the extracted header for further processing to the electronic control unit. Synchronization, re-amplification, and wavelength conversion of input wavelengths could also be carried out in the input interface, if necessary. For example, in slotted OPS, synchronization of optical packets to the beginning of time slots should be performed in the input interface. Even re-amplification of optical packets can be carried out in the input interface. For example, EDFA amplifiers can be used inside the core switch in order to compensate the internal loss of the core switch and the loss occurred on the fiber link from the last amplifier on the link and the core switch [7, 95].

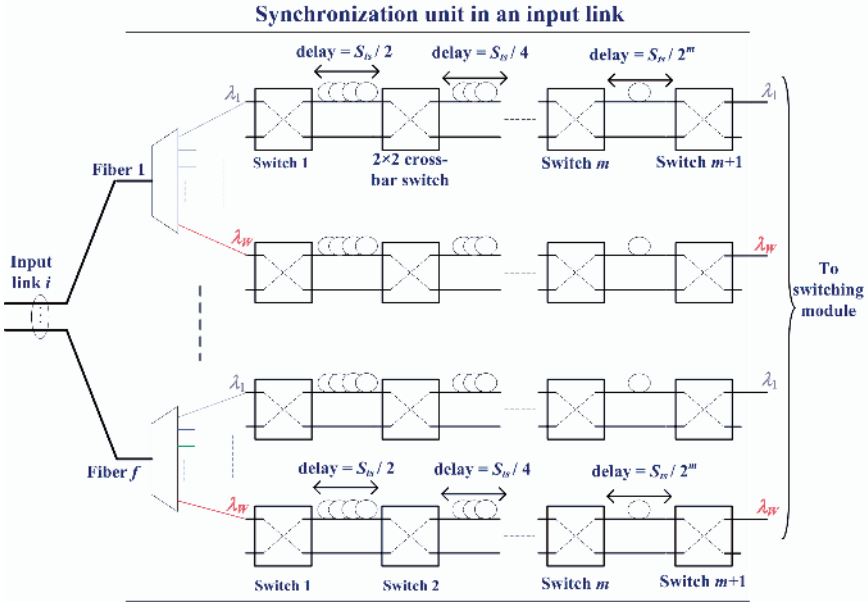


Figure 1.11: Synchronization unit for an input link of a multi-fiber OPS network

For proper operation of transmission and reception of optical packets in slotted OPS, we need to make an assumption where the fiber length between any pair of switches (edge-to-core or core-to-core) must be an integer number of time slot duration (i.e., $S_{ts} = S_T + S_O$). However, all fiber spans cannot be exactly designed as an integer number of time slot duration. In addition, the arrival of optical packets or control information at the core switches may be misaligned with each other due to the chromatic dispersion of different wavelengths, accumulative jitter of different paths inside optical switch fabrics, temperature variation, and other fiber transmission non-linearities that result in varying length of fibers. To provide synchronous switching operation, they must first be realigned by using a synchronizing unit at each input port of any core switch in order to synchronize the incoming optical packets boundaries to the local timing reference. Figure 1.8 shows a slotted multi-fiber

OPS switch architecture with f fibers in each input/output link, where synchronization modules are used to align optical packets to the beginning of time slots. This can be implemented using a finely calibrated set of optical delay lines with feedback control of the delays being provided through the system controller. Note that each core switch is operating with a reference to its own internal clock that can be derived from a network synchronization signal distributed throughout the network. Considering time-slot length of S_{ts} , Fig. 1.11 illustrates the synchronization block diagram in an OPS switch in which there are $f \times W$ synchronization modules. Each synchronization module for an input link include $m + 1$ switching elements, where each 2×2 switching element can be set either in bar state or in cross state. By controlling the switching elements, different fibers delay lines can provide different delays starting from at least 0 to at most $\sum_{j=1}^m \frac{S_{ts}}{2^j}$ for each wavelength channel [53, 65, 96–98].

The control unit processes the header of an optical packet to obtain the information about source and destination edge switches, and then it searches its switching table to find the suitable output port for forwarding the optical packet. Then, it sends control signals to the switching module for appropriately configuring it. In short, the control unit decides which optical packet should be switched to which output fiber and on which wavelength channel. Contention resolution is one of the most important responsibilities of the core switch. Contention happens when several optical packets must be switched to the same output fiber on the same wavelength channel at the same time.

The switching module is a wavelength-selective optical switch. This module in the single-fiber architecture includes W optical switching fabrics, where each switching fabric is used for switching a wavelength channel. Note that each switching fabric is a strictly non-blocking $N \times N$ optical switch. In the multi-fiber architecture, each switching fabric is a strictly non-blocking $(N \times f) \times (N \times f)$ optical switch. The switching module may use different contention resolution schemes (like wavelength conversion and optical buffering) for resolving the contention of optical packets that are going to be sent over the same fiber and wavelength channel at the same time. The optical buffering displayed in the figure is actually based on Fiber Delay Lines (FDLs) that are used for two purposes: delaying optical packets headers for contention resolution purposes and waiting for packet header processing. This switch module may also use Tunable Wavelength Converters (TWC) and other techniques for contention resolution. Figure 1.10 displays an example for switching module in a multi-fiber OPS switch with $N = 3$ input links and $N = 3$ output links, $f = 2$ fibers per link, and $W = 4$ wavelengths per fiber shown with different colors. As can be observed, four 6×6 optical switch fabrics are used in the switching module, one optical switch fabric for each wavelength channel. Each optical switch fabric is a strictly non-blocking switch. A strictly non-blocking switch can be built with different architectures such as Clos architecture, cross-bar architecture, Cantor architecture, and so on [36]. The important factor for an optical switching fabric is its reconfiguration speed which must be on the order of nanoseconds for OPS. The speed of switching fabric based on the Micro-Electro-Mechanical Systems (MEMS) technology is on

the order of milliseconds, which is not appropriate for OPS. The Semiconductor Optical Amplifier (SOA) and electro-optic Lithium Niobate (LiNbO₃) technologies are promising for OPS switching fabrics. SOAs have a switching speed on the order of a few nanoseconds and can be integrated to large scales, but they add some noises to optical signals. The LiNbO₃ switching fabrics have sub-nanosecond switching times, but with high insertion loss, thus limiting their integration scalability for only medium-scale optical switching fabrics. Other important factors for a switching fabric could be: scalability to large port numbers, ease of manufacturing, low cost, temperature independence, and being strictly non-blocking [21, 65].

After switching of an optical packet, the output interface attaches the updated header to the packet payload. The output interface can also perform signal amplification to overcome noise and degradation of optical signals and resynchronize the optical packet to the time slots in slotted OPS. Wavelength conversion is another responsibility of the output interface. Finally, multiplexers at the output side combine wavelengths on the output fibers.

It should be mentioned that different architectures have been proposed for OPS switches. The detailed architectures will be discussed in Chapters 3 and 4.

1.4.3 Edge Switch Architecture in OPS

An important issue in an OPS network is to provide network access for traffic coming from/delivering to legacy networks. As Fig. 1.12 depicts, an edge switch is used to isolate the optical domain of the OPS switch from the electronic domain of G legacy networks. An edge switch has two main functionalities. As an ingress switch, it collects clients' traffic coming from G legacy networks, redirects traffic destined to the same egress switch in the same buffer, schedules them, converts them to optical packets, and then transmits the optical packets to the OPS core network. As an egress switch, it receives optical packets from the OPS core network, converts them to the electronic packets, saves the traffic destined to the same legacy network in the same buffer, and then delivers the electronic packets to the G legacy networks. Borrowing from [99], a "**torrent**" is defined as the whole (class-based) traffic going to the same egress switch in an ingress switch. Therefore, Torrent- i traffic goes to egress switch i .

Figure 1.12 illustrates a multi-fiber OPS network in which there are f fibers on each connection link. The OPS switch has $N \times f$ input ports and $N \times f$ output ports, where each set of f input/output ports is called a link. Among these links, there are $N - 1$ transit links (i.e., with $(N - 1) \times f$ fibers) for switching transit optical packets and one link (i.e., with f fibers) for traffic add/drop from/to the edge switch. In this network, the edge switch is connected to the OPS switch with f fibers as add ports and with f fibers as drop ports. The edge switch can use all W wavelengths provided on each fiber for traffic transmission to the OPS switch and traffic receiving from the OPS switch. In this case, the edge switch should use $f \times W$ fixed optical transmitters and $f \times W$ fixed optical receivers. The bandwidth of each wavelength channel is B_C in bps.

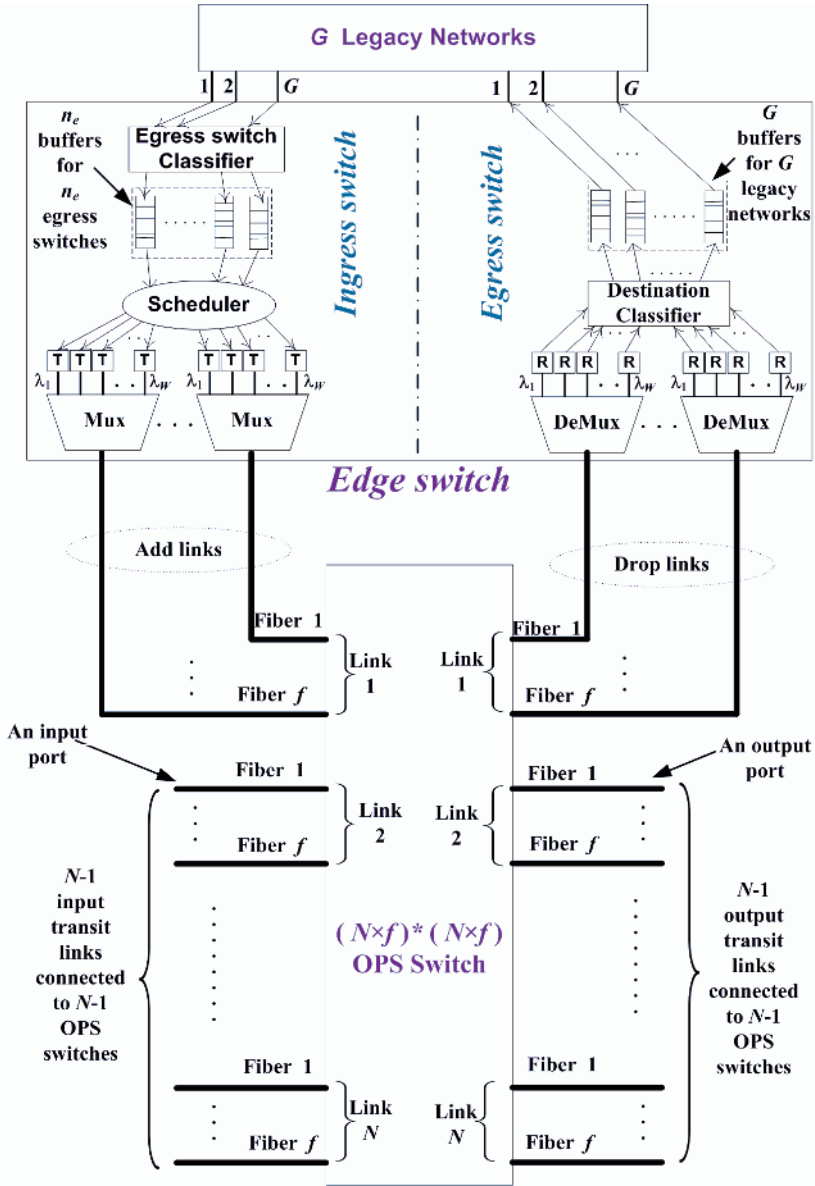


Figure 1.12: Edge switch architecture in a multi-fiber OPS network

In its simplest form, as Fig. 1.12 shows, an ingress switch contains one electronic buffer dedicated to each egress switch; $n - 1$ buffers for the OPS network with n OPS switches. When a client packet destined to a given egress switch arrives at an ingress switch, it is saved in the relevant egress switch FCFS (First-Come-First-

Served) buffer. There is a packet scheduler unit that schedules client packets from the buffers, converts each client packet to an optical packet, adds a header to the optical packet, and then sends the optical packet to the OPS network on an available wavelength channel immediately (in asynchronous OPS) or at time slot boundary (in slotted OPS).

For class-based traffic, packet differentiation can be implemented using differentiation buffers in an ingress switch. The network consists of M classes of traffic, where a class is defined according to the QoS requirements (e.g., $M = 3$ for EF, AF, and BE according to the Differentiated Services (DiffServ) model). Note that there are three general classes defined in DiffServ: Expedited Forwarding (EF) [100], Assured Forwarding (AF) [101], and Best Effort (BE). In this case, for a network supporting M classes of service, M buffers are used for each egress switch; $M \times (n - 1)$ buffers for the OPS network with n nodes. There is a QoS-based packet scheduler unit that schedules client packets from the class-based buffers, creates optical packets, and then transmits them to the network on available wavelengths.

As stated, an ingress switch can immediately send an arriving packet to the OPS network. However, this may increase traffic loss at the network due to the traffic burstiness, especially under IP traffic. There are two common mechanisms proposed for traffic shaping in ingress switches to cope with the bursty traffic and avoid contention in an OPS network:

- The first mechanism uses traffic shaping similar to electronic networks [102, 103].
- The second mechanism aggregates a number of client packets in an optical packet to reduce traffic burstiness in the OPS network [83, 104, 105], thus improving the network performance. This approach can be used in both single-class and multi-class traffic networks.

1.4.4 Signaling in OPS

In OPS, electronic control logic units are distributed in control units of OPS switches in order to process the header of each optical packet and route the optical packet toward its egress switch. Payload of an optical packet can contain one or more upper layer client packets, e.g., IP packets. It should be mentioned that the header of an optical packet may include some fields that must be updated at intermediate OPS core switches after switching the optical packet. A header of an optical packet may include different fields such as [65]

- Synchronization bits
- Ingress switch address generating the optical packet
- Egress switch address used for routing the optical packet in an OPS core network

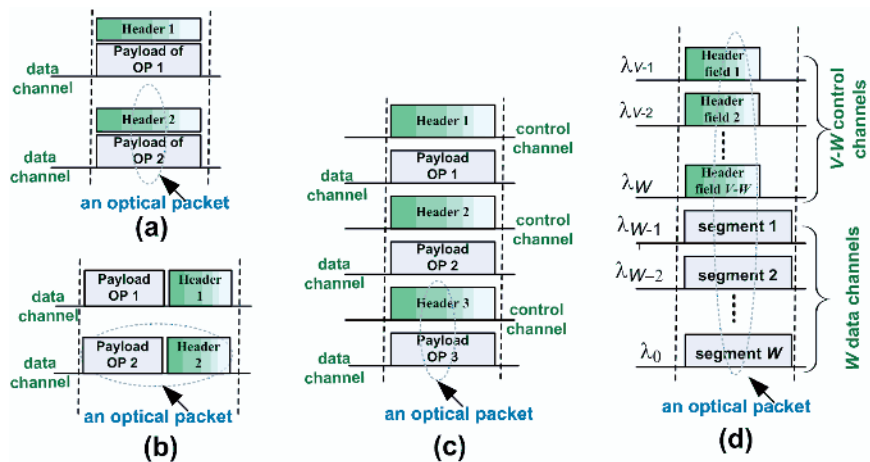


Figure 1.13: OPS signaling: (a) subcarrier multiplexing, (b) serial transmission, (c) separate wavelength, (d) header stripping

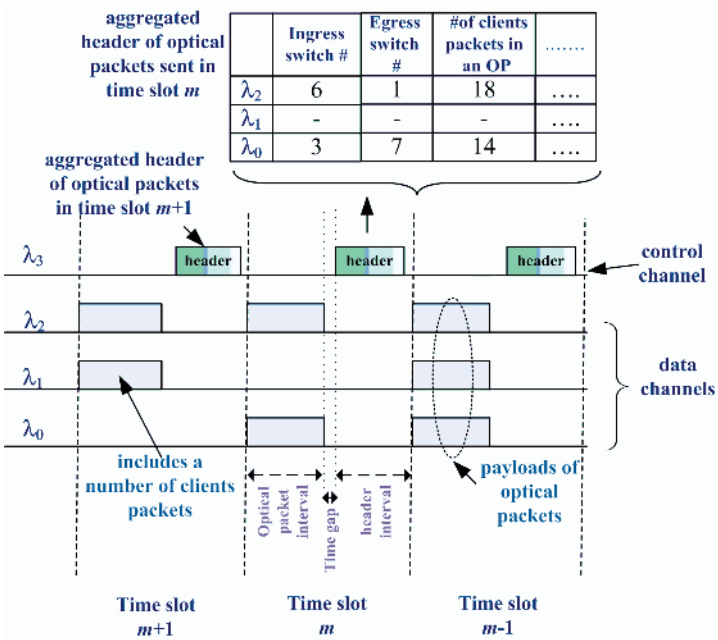


Figure 1.14: Aggregated header for three data wavelength channels

- Packet type that identifies the optical packet type and its priority for implementing QoS
- Packet sequence number to identify out-of-order or duplicated optical packets

- Operation, administration, and maintenance fields
- Header error correction field
- Number of client packets aggregated in the payload of the optical packet
- Fragmentation information for managing fragmented client packets in a number of optical packets
- Source routing information when the optical packet carries its route map

There are some time gaps (called guard bands) between the header and payload of an optical packet. Guard bands are used for waiting time of header processing in the electronic domain and for de-jittering of payloads. Different techniques have been proposed for transmitting data payload and header of an optical packet in an OPS network as follows:

- **Subcarrier Multiplexing:** In this approach (see Fig. 1.13a), an optical packet header (low-bandwidth) is placed on a subcarrier above the baseband frequencies occupied by the optical packet payload (high-bandwidth), and both are sent within the same time slot. In this technique, an optical packet includes only one client packet. Using this technique, header processing can take up to the entire data payload transmission time. However, by increasing the payload data rate, the baseband will expand and it might overlap with the header subcarrier frequency [21].
- **Serial Transmission:** In this simple technique (see Fig. 1.13b), before transmission of the data payload of an optical packet (that includes only one client packet), its header is sent serially on the same wavelength. There is a guard time between the header and the data payload to allow for the removal and the reinsertion of the new header at intermediate OPS switches [21].
- **Separate Wavelength:** Since in this approach (see Fig. 1.13c) the header and the payload of an optical packet (that carries only one client packet) are sent on two separate wavelength channels at the same time, extracting the header at the inputs of OPS switches is simpler than previous methods. There are two disadvantages for this method: crosstalk and dispersion. To compensate the dispersion effect on the delay, each intermediate OPS switch should realign the header and the payload [21].
- **Header Stripping:** In this approach (see Fig. 1.13d), optical packets are injected in an OPS network with a wavelength-striped packet format [106–108]. Different fields of an optical packet (such as packet length, quality of service, ingress switch address, and egress switch address) are separately encoded on dedicated control wavelengths, and the payload traffic is segmented and distributed over the rest of the available wavelengths. Each payload segment is modulated simultaneously at a high data rate (say at 40 Gbits/s per wavelength) to yield a high aggregate message bandwidth. The wavelength-striped optical

packets are routed at each OPS switch altogether. Note that in this technique, an optical packet carries only one client packet. In this case, each OPS switch must be a simple optical cross-connect switch without wavelength selection capability that transparently switches the whole optical packet from an input port to an output port. In the wavelength-striped OPS, wavelength conversion is no longer an appropriate contention resolution scheme in OPS core switches.

- **Aggregated Header:** This approach is suitable for synchronized OPS networks (see Fig. 1.14) [104], where there is one dedicated control wavelength channel and W data wavelength channels on each fiber. A time slot includes two parts: (1) optical packet interval for transmitting optical packet payloads (with duration S_T , called slot traffic) and (2) a time-gap interval plus a header interval for transmitting header information (with duration S_O , called slot offset). Note that a time slot interval is $S_{Ts} = S_T + S_O$.

For all W optical packets transmitted on W wavelength channels on a fiber within a time slot, an aggregated header is sent over the control channel during the header interval that includes traffic information (such as ingress/egress switch addresses) for each optical packet. It can be observed that there is a time gap between the transmission of a header and the transmission of relevant optical packets within a time slot. Note that in this technique, an optical packet may include an integer number of client packets. It should be mentioned that time slots are larger in this technique than the conventional synchronized OPS networks that use the serial transmission technique [104]. Figure 1.14 shows the information carried in the header for the optical packets sent within time slot m , where no optical packet is transmitted on wavelength λ_1 .

Slot offset S_O should consist of a guard time to allow for timing uncertainties (T_{guard}), a processing time (T_{proc}), and a switching time (T_{sw}) at a core switch. The processing time during the S_O interval consists of the time required to evaluate potential contention, to resolve the contention, to make a new aggregated header, and finally to make the core switch ready to switch the arriving optical packets toward their desired egress switches. Therefore, we should have $S_O > T_{guard} + T_{proc} + T_{sw}$. For instance, if we need $T_{guard} = 100$ ns, $T_{proc} = 700$ ns, and $T_{sw} = 400$ ns, then we should have $S_O > 1200$ ns.

Note that S_T must be chosen in such a way that the bandwidth wastage is minimized. In general, two bandwidth overheads can be found in slotted OPS networks: (1) slot-offset overhead of $O_1 = \frac{S_O}{S_{Ts}}$ and (2) client packet aggregation overhead of $O_2 = \frac{(B_C \times S_T) \bmod L_a}{B_C \times S_{Ts}}$, where B_C is wavelength channel in bps, L_a is average length of client packets, and $x \bmod y$ is the remainder of x divided by y . The second overhead occurs when a number of client packets can be carried in an optical packet within a time slot. For example, consider a slotted OPS network with wavelength channel bandwidth $B_C = 40$ Gbits/sec and average client packet size $L_a = 5200$ bits. For example, if one requires $S_O = 1$ μ s and wants at most 10% bandwidth overhead, then we should have $O_1 + O_2 \leq 0.1$. Hence, we obtain $36,000 + (40,000 \times S_T) \bmod 5200 - 4000 \times S_T \leq 0$. One

can find a list of answers for this inequality such as $S_T = 9.11 \mu\text{s}$, $S_T = 9.24 \mu\text{s}$, and so on. However, choosing a large value for S_T would increase queuing delay in ingress switches because of high optical packet inter-departure times. Therefore, to have both a desirable bandwidth overhead and small queuing delay in ingress switches, one should select a smaller value for S_T from the list of answers.

1.4.5 Contention Problem in OPS

Packet loss in an electronic switch is mainly due to buffer overflow and bit error rate that makes the packet corrupted. Although OPS networks are promising networks for future optical core networks, they still encounter some problems. In OPS networks, an OPS switch works in cut-through mode, and there is mainly no RAM in the switch, and therefore optical packets cannot be saved in buffers like the electronic networks. Hence, the most important problem could be contention of optical packets.

In OPS networks, there is no collaboration among ingress switches for packet transmission, and traffic is sent to the optical network without any coordination among transmitters. Assume that a number of optical packets arrive at an OPS switch, where they all must be switched to the same desired output port. Contention of optical packets occurs in a single-fiber OPS switch when more than one optical packet should be switched to the same output fiber on the same wavelength at the same time, thus damaging contending all optical packets. On the other hand, in a multi-fiber OPS switch, contention occurs when more than f optical packets should be switched to the same output fiber on the same wavelength at the same time. In any architecture, some of the contending optical packets will go through, while others should be dropped.

The contention problem is different in slotted and asynchronous OPS networks. As stated, optical packets only arrive at an OPS switch at the start of time slots in slotted OPS; therefore, contention could only happen at the starts of time slots. On the other hand, in asynchronous OPS, optical packets can arrive at any time; therefore, contention occurs when the desired wavelength of an output fiber is already occupied. This is why the probability of contention in asynchronous OPS is higher than slotted OPS. Note that, compared with asynchronous OPS, slotted operation can reduce the vulnerable period of information contention and can reduce traffic loss [21]. Another sort of contention happens at an egress switch in which the number of optical receivers at the egress switch is less than the number of received optical packets [109].

When contention happens in an OPS switch and there is no contention resolution mechanism, only one of the contending optical packets in a single-fiber architecture and f contending optical packets in a multi-fiber architecture can only be switched successfully to the desired output link, and the remaining contending optical packets must be dropped. This clearly increases PLR in OPS and reduces network throughput. Therefore, contention resolution mechanisms should be implemented in each OPS switch in order to decrease PLR.

Note that the same contention problem happens on an output port of an electronic switch in an electronic network (or even an opaque switch in an opaque optical network in which optical packets are converted to electronic packets in switches). The contention in this case is simply resolved by buffering all contending electronic packets in electronic queues located at the output port and then scheduling them on the output port. However, there is no optical queue in the optical domain, and optical packets cannot be stored like electronic packets.

In general, contention is the major problem for an OPS network. To deal with the contention problem, resolution and avoidance are two common schemes. Contention resolution schemes (as reactive schemes) resolve arisen collisions in OPS switches. On the other hand, the contention avoidance schemes (as proactive schemes) send optical packets in ingress switches in such a way to reduce the number of collision events in OPS switches. Due to the bursty nature of the Internet traffic, the contention resolution schemes may not be effective to maintain a reasonable network performance under higher traffic loads. Therefore, contention avoidance schemes should also be used to regulate traffic going to an OPS network, thus reducing the traffic loss. Contention avoidance includes a different range of operations such as load balancing, using multi-fiber OPS, traffic shaping, traffic aggregation in ingress switches, even packet transmission, and so on, which will be detailed in Chapter 2.

On the other hand, contention resolution schemes (detailed in Chapter 3) try to resolve an occurred contention in an OPS switch. There are two main approaches to resolve contention of optical packets. In the first approach, received optical packets are first converted to electronic packets in an OPS switch (i.e., opaque networks), and then the OPS switch switches them as in electronic networks. However, this can be a bottleneck because of limited speed and high cost of O/E/O converters [110]. The second approach keeps optical packets in the optical domain (i.e., transparent network). When contention happens in a transparent network, an OPS switch uses some approaches to resolve the contention by changing space (by using deflection routing), switching time (using optical buffers), retransmission, converting wavelength (using wavelength converters), or a combination of them.

It seems that the contention avoidance schemes are mostly cheaper than the contention resolution schemes since most of the avoidance schemes are software level operations. On the other hand, a multi-fiber OPS architecture without using wavelength converters inside any OPS switch could also be cheaper than a single-fiber OPS when using wavelength converters in OPS switches. In addition, devices such as dispersion compensators used for a fiber with a high number of wavelengths is more expensive than the devices used for a fiber with a small number of wavelengths [13, 14]. In other words, a multi-fiber OPS with a small number of wavelength channels per fiber is cheaper than a single-fiber OPS with a large number of wavelengths per fiber.

Finally, even the size of optical packets may influence PLR in asynchronous OPS [111]. Transmission of fixed-size optical packets leads to the lowest PLR when using FDL buffering in OPS switches. On the other hand, variable-size optical packets yield the smallest PLR when the arrival process is bursty and there is no FDL buffer in OPS switches.

1.4.6 Quality of Service (QoS) in OPS

Quality of Service (QoS) is a broad term with almost 40 definitions [112]. Based on these definitions, degree of satisfaction of users could be the main objective of QoS provisioning in networks. QoS can be categorized from two points of views [113]:

- **QoS Experienced by End Users:** End user perception of QoS denotes how the quality of a particular service is received from the network. This necessitates bandwidth guarantee, minimum delay, small PLR, and controlled jitter.
- **QoS from the Network Point of View:** Network's perception of QoS denotes how the network resources and capabilities are fully and efficiently utilized by end users.

Since implementing QoS is complicated, some may believe that QoS is not needed because OPS bandwidth will be high enough to provide good QoS for all applications. However, this needs an infinite bandwidth in network components in order to support all network traffic, which is an impossible issue. Although an OPS network can provide high bandwidth, it must also improve the parameters that affect the network QoS in order to fully utilize the network bandwidth. OPS networks should be QoS-capable in order to support a variety of services and traffic types (such as data, voice, video, and multimedia streaming) generated by the applications with very different requirements of network performance or QoS such as in the Internet. This is because a variety of real-time applications (such as video on demand, voice over IP, Internet Protocol TV (IPTV), video conferencing) with different QoS requirements are executed in the Internet domain nowadays. This results in providing QoS support to (a) the large variety of service requirements and (b) the requirement to carry multiple services efficiently across both IP and optical networks [113–116].

1.4.6.1 QoS Metrics Different QoS metrics can be stated for satisfaction of network users:

- **Network Reliability:** Reliability shows the probability that a network component will continue to perform its desired function under a given operating conditions over a predetermined period of time [117]. Different performance metrics include:
 - Network failure rate: This is the reverse of Mean Time To Failure (MTTF) [117].
 - Network repair rate: This is the reverse of Mean Time To Repair (MTTR) [117].
- **Network Security:** This metric is relevant to the strength of cryptographic algorithms for hiding information and integrity of information.
- **Network Performance:** Important performance metrics in an OPS network include:

- **Throughput:** The throughput metric is a measure of the number of packets that a network can successfully deliver. In practice, the throughput is less than the network bandwidth because of protocols overhead and network congestion. The throughput metric can be separately computed for client traffic and optical packet traffic because an optical packet may carry a number of client packets. The normalized throughput is the common performance metric in an OPS network which is computed in a network with n core switches (where one edge switch is connected to each core switch) by

$$Throughput = L \times \frac{\sum_{i=1}^n N_{dlv,i}}{\sum_{i=1}^n N_{sent,i}}, \quad (1.1)$$

where $N_{dlv,i}$ is the number of packets successfully delivered in egress switch i , $N_{sent,i}$ is the number of sent packets to the network by ingress switch i , and L is normalized traffic load of packets arrival at the network. Clearly, when the second term tends to 1.0 (i.e., almost all transmitted packets have been successfully delivered), throughput will tend to traffic load L .

Throughput can also be defined as the volume of traffic successfully delivered within a given period of time. For example, if a traffic of volume v bits successfully arrive at a destination within τ seconds, then throughput will be $\frac{v}{\tau}$ bits/s.

- **Fairness:** Fairness determines whether network users or applications are receiving a fair share of system resources. It is usually measured by the Jain's fairness index [118]:

$$J(x_1, x_2, \dots, x_m) = \frac{(\sum_{i=1}^m x_i)^2}{m \times \sum_{i=1}^m x_i^2}, \quad (1.2)$$

where m is the number of users/applications and x_i is the throughput or bandwidth received by user/application i .

- **Packet Loss Rate (PLR):** The packet loss rate can be separately computed for client traffic and optical packet traffic. This is because an optical packet may carry a number of client packets. When an optical packet carries only one client packet, these two packet loss rates will be the same. Note that it is common to consider the electronic buffers in ingress switches large enough so that no client packet is lost due to the buffer overflow in ingress switches. The general formula to compute PLR in a network with n core switches (where one edge switch is connected to each core switch) is given by

$$PLR = \frac{\sum_{i=1}^n N_{loss,i}}{\sum_{i=1}^n N_{sent,i}} , \quad (1.3)$$

where $N_{loss,i}$ is the number of lost packets in core switch i and $N_{sent,i}$ is the number of sent packets to the network by ingress switch i .

- **Delay:** An optical packet sent by an ingress switch experiences some delay until it arrives at its destination. This delay performance can be stated in two ways as queuing delay and end-to-end delay:
 - * The queuing delay is due to the waiting time of a packet in an electronic buffer of an edge switch and waiting time in the optical buffers of OPS core switches along the path to its destination. Clearly, when there is no optical buffers used in core switches, the queuing delay is limited to the waiting time in edge switches.
 - * The end-to-end delay t_{e2e} for a given optical packet is composed of four parts: (1) transmission delay time to push the optical packet bits onto a wavelength in the ingress switch (t_{tx}), (2) propagation delay of the optical packet due to the speed of light on the links along its path (t_{prop}), (3) processing delay to process the header of the optical packet and make switching module ready in all core switches along its path (t_{proc}), and (4) its queuing delay in the ingress switch and all the core switches along its path (t_{queu}) as stated above. In other words, we have

$$t_{e2e} = t_{tx} + t_{prop} + t_{proc} + t_{queu} . \quad (1.4)$$

Taking average among end-to-end delays of all successfully delivered optical packets results in average end-to-end delay in the network.

- **Jitter:** Jitter is computed from the end-to-end delay. This metric is the measure of variation of delay due to the fact that each packet in the network travels through different paths, and therefore the end-to-end delay varies. other reasons for creating jitter is network congestion and variable processing times at intermediate switches. In general, there are two methods to compute jitter. First, let d_i denote the end-to-end delay of packet i , and let m denote the number of packets from the same flow arrived at a destination. The first method to compute jitter is given by [119]

$$Jitter = E [|d_{i+1} - d_i|] = \frac{\sum_{i=1}^{m-1} |d_{i+1} - d_i|}{m-1} . \quad (1.5)$$

The second method to compute jitter is calculated by [120]

$$Jitter = \sqrt{\frac{\sum_{i=1}^m (d_i - \bar{d})^2}{m-1}}, \quad (1.6)$$

where $\bar{d} = \frac{1}{m} \sum_{i=1}^m d_i$ is average delay of m packets.

- **Bit Error Rate (BER) and Q Factor:** BER could be the most important metric that influences the end-end quality of an optical packet. BER is used to determine an optical packet signal quality. A transmitted optical packet may fail to arrive at its destination (i.e., lost due to high BER) because of signal degradation, packet corruption, faulty networking hardware, faulty network drivers, and distance between the source and destination. The Q factor is a good intermediate parameter for BER and Optical Signal-to-Noise Ratio (OSNR). For example, $BER = 10^{-11}$ is equivalent to $Q = 6.7$, and $BER = 2.86 \times 10^{-7}$ is equivalent to $Q = 5$. This factor is obtained from [7]

$$Q = \sqrt{\frac{B_o}{B_e}} \times \frac{2 \times OSNR}{1 + \sqrt{1 + 4 \times OSNR}}, \quad (1.7)$$

where B_o and B_e are optical and electrical bandwidths, respectively. We could have $\frac{B_o}{B_e} = 10$. The approximate relation between Q factor and BER is defined by the following equation when $Q \geq 3$ [121]

$$BER \approx \frac{e^{-Q^2/2}}{\sqrt{2 \times \pi \times (Q^2 + 2)}}, \quad (1.8)$$

where $e \approx 2.71828$ is the Euler's number.

1.4.6.2 QoS Provisioning QoS is related to traffic control and resource management. QoS provisioning in an OPS network can be divided into three categories.

1.4.6.2.1 QoS Improvement QoS improvement mechanism is defined as any mechanism that can improve the general performance of the network through traffic control and resource management mechanisms and provide predictable or guaranteed performance for delay, packet loss, jitter, fairness, reliability, and fault tolerance parameters. A QoS improvement mechanism controls the allocation of network resources to all applications traffic in a way to improve their performance parameters in order to satisfy their users. This is suitable for a network with single class traffic

of service in which all packets are treated equally. Applications like email and web browsing that are non-interactive can be counted in this category [114].

1.4.6.2.2 QoS Differentiation There is no guarantee in PLR and delay performance parameters in today's Internet services since Internet offers only the best effort and connectionless service model. For single class BE traffic service, no guarantees can be given to any packet regarding loss rate, delay, and jitter performance metrics since all traffic in the network is equally treated [114]. However, when network resources become inadequate, the network performance parameters will be degraded. This will in turn degrade the QoS requirements for the real-time traffic. Therefore, having an optical OPS network to differentiate the optical packets generated by different network applications/users is a requirement in the near future.

Under QoS differentiation, the network must be able to differentiate between high-important traffic and low-important traffic and provide better service for the high-important traffic. Traffic differentiation is required under two cases:

- **Differentiation Based on Application Traffic:** Here, high-important traffic includes the traffic generated by highly interactive applications (such as video conferencing and online gaming that need stringent operating requirements), and low-important traffic includes the traffic generated by non-interactive or semi-interactive applications (such as e-mail, web browsing). Since different network applications need different levels of QoS, service differentiation should also be considered for future optical networks.
- **Differentiation Based on User Traffic:** Here, high-important traffic includes the traffic generated by the companies that pay high premiums to ensure network reliability in order to transmit their critical transactions. On the other hand, low-important traffic includes the traffic generated by home users only that need cheap Internet access and can tolerate a lower service level [55].

In store-and-forward electronic routers and switches, buffers are easily used to provide packet differentiation. However, this is not suitable for all-optical OPS networks due to the lack of optical buffers in OPS switches (i.e., either having a little or no optical buffers). Therefore, there are interests in providing new approaches to provide service differentiation in all-optical OPS networks without using optical buffers [122–124]. The QoS in OPS implies high throughput, low latency, and low packet loss for high-priority class traffic. It is also desirable to keep the order of optical packets because optical packet reordering increases latency at egress switches.

A number of enhancements have been introduced to offer different levels of QoS for IP networks. There are mainly two common methods for service differentiation as IntServ [125] and Differentiated Services (DiffServ) [126]. IntServ can guarantee QoS by end-to-end bandwidth reservation for traffic flows and per-flow scheduling in all intermediate routers/switches in a network. On the other hand, DiffServ can guarantee QoS differentiation for different classes of traffic aggregates. Since DiffServ is more scalable than IntServ, optical networks use the DiffServ model for QoS provisioning [59, 122, 123, 127]. Under DiffServ, edge switches classify, mark, drop,

or shape client packets based on a Service Level Agreement (SLA) and prevent the DiffServ network from malicious attacks. On the other hand, core nodes perform high speed switching of differentiated packets and provide relative per-class QoS differentiation by allocating more bandwidth, lower delay, or lower loss to one class than another class. There are three general classes defined in DiffServ: Expedited Forwarding (EF) [100], Assured Forwarding (AF) [101], and Best Effort (BE). The EF service class is used for the applications that need low loss rate, low latency, low jitter, and bandwidth guarantees, while AF offers different levels of forwarding assurances to client packets. The remaining traffic is treated as BE without any QoS guarantee [128, 129].

In addition to EF, AF, and BE classes in the DiffServ model, some researches use general classes. In general, HP, MP, and LP denote high-priority, mid-priority, and low-priority traffic, respectively. Note that client packets are usually classified into two or three classes, and therefore we will have HP, MP, and LP client packets in the three-class scenario. If an optical packet carries one or more client packets from the same class c , the class of that optical packet is also considered as c . For instance, an optical packet carrying a number of LP client packets is called an LP optical packet.

The QoS provisioning can be grouped as either relative or absolute [130]:

- In absolute (hard) QoS provisioning, QoS is guaranteed to each traffic class based on the rules stated in an SLA. Note that the SLA is a contract between a user and the network provider, which defines the number of forwarding classes, level of service, and upper bounds on the values of network parameters (such as delay or loss). In other words, the absolute QoS provides the worst-case guarantee on the loss, delay, and bandwidth to applications, which is appropriate for delay and loss sensitive applications such as multimedia and mission-critical applications [131].
- The relative (soft) QoS differentiation model can provide quantitative differentiation between different service classes in such a way that a high-priority traffic is guaranteed to receive no worse service than a low-priority traffic. In other words, a high-priority optical packet should always obtain better service (such as lower loss, lower delay, and lower jitter) than a low-priority optical packet. It is clear that the amount of service received by a class depends on the current network load in each class. The main advantage of relative service differentiation over the absolute service differentiation is its simplicity and ease of deployment [116, 132].

The relative service differentiation model has been improved to the proportional differentiation QoS model in order to provide the network operators with quantitative QoS differentiation among service classes. With this QoS model, service differentiation between traffic classes can be adjusted according to pre-defined factors. The proportional differentiation model is more scalable than the relative service differentiation model. Define q_i to be the desired QoS metric for class i traffic, and define s_i to be the differentiation factor for class i traffic that is set by the network service provider. With M traffic classes in the network, we should have [132]

$$\frac{q_i}{q_j} = \frac{s_i}{s_j}, \quad (1.9)$$

for any class $i, j = 1, 2, \dots, M$. For instance, assume q_1 and q_2 to be packet loss rates of class 1 and class 2, respectively. In addition, consider that the network service provider has set $s_1 = 2$ and $s_2 = 3$. Then, we should have $\frac{q_1}{q_2} = \frac{2}{3}$; i.e., packet loss rate of class 1 must be almost two-third of class 2. The proportional differentiation model must be held for both long time and short time periods since the long-term average is not quite meaningful under bursty traffic. For this purpose, Eq. (1.9) is changed to the following equation [132]:

$$\frac{\bar{q}_i(t, t + \tau)}{\bar{q}_j(t, t + \tau)} = \frac{s_i}{s_j}, \quad (1.10)$$

where $\bar{q}_i(t, t + \tau)$ is the average QoS metric from time t to $t + \tau$. By this general equation, the quality of differentiation between different traffic classes can be defined as a function of various QoS metrics. However, in practice, we cannot expect the equality in Eq. (1.10), and some deviations must be allowed in this equation as [132]

$$\frac{\bar{q}_i(t, t + \tau)}{\bar{q}_j(t, t + \tau)} = \frac{s_i}{s_j} \pm \Delta, \quad (1.11)$$

where Δ is a deviation from the relation determined by $\frac{s_i}{s_j}$. By changing s_i and s_j accordingly, the differentiation between them can be controlled within a bounded deviation. When τ is small enough, high-priority traffic will always receive better service than low-priority traffic independent of traffic load fluctuation [132].

As stated, electronic networks can support class-based traffic based on the Diff-Serv model. This model has three major classes (EF, AF, and BE) and a number of sub-classes for AF (such as AF11, AF12, AF13, AF21, ..., AF33, AF41, AF42, and AF43) [133]. Hence, the number of traffic classes in the electronic domain could be high. On the other hand, the number of classes in the optical domain must be kept as small as possible in order to reduce operational complexity of OPS switches since complex scheduling algorithms for several classes of traffic may not be applicable in the optical domain because of the limitation of buffering in the this domain [134, 135]. Some works have considered two classes of traffic in the optical domain, e.g., [114, 135], and some have considered three traffic classes, e.g., [83, 95, 104]. Therefore, the class-based traffic in the electronic domain should be mapped to the available classes in the optical domain using the method stated in [114].

1.4.6.2.3 Quality of Transmission (QoT)-Based Scheduling Many OPS studies have assumed an ideal physical layer for transmission and switching of optical packets. However, the ideal assumption is suitable for opaque networks. As an optical packet propagates toward its destination in a transparent optical network, its signal quality degrades and its BER increases due to attenuation, noise, dispersion, crosstalk, jitter, and non-linear effects. This is because there is no signal regeneration on the path. Note that increasing BER of an optical signal decreases OSNR in optical receivers. This will finally result in optical packet drop in an egress switch or even in an intermediate OPS switch.

In the Quality of Transmission (QoT)-based scheduling mechanism, the network must switch network traffic in such a way that optical packets can be detectable at their egress switches. The QoT-based switching is different in OPS from OCS networks. Many studies on wavelength-routed networks have considered a non-ideal physical layer and provided complex models for evaluating different aspects of physical layer impairments for establishing a light-path since in these networks a light-path is first set up and then data are transmitted. However, using complex models for evaluating physical layer impairments for each optical packet cannot be practical in OPS networks because of its finest granularity. Therefore, evaluating complex models should be avoided in OPS. Instead, simple physical layer impairment models or indirect physical layer impairment models should be used in OPS. In the former, transmission impairments must be taken into account in switching of optical packets towards their destinations. To do this, OPS switches must be aware of the signal quality of an optical packet (i.e., QoT-aware OPS) and should not switch the optical packet toward its egress node if it does not have enough quality to traverse its path. In the latter, optical packets are routed through short-path or short-hop routes toward their destinations. This indirectly takes into account the impairments.

There are three kinds of Physical Layer Impairments (PLIs):

- Linear impairments (such as path loss, Chromatic Dispersion (CD) or Group Velocity Dispersion (GVD), Polarization Mode Dispersion (PMD), and insertion loss) are the impairments that do not depend on signal power.
- Nonlinear impairments (such as intra-channel Self-Phase Modulation (SPM), Stimulated Brillouin Scattering (SBS), Stimulated Raman Scattering (SRS), inter-channel crosstalk originated from the non-linear interaction within fiber spans of several signals co-propagating on different wavelengths such as inter-channel Cross-Phase Modulation (XPM) and inter-channel Four-Wave Mixing (FWM) crosstalk between channels) are signal power dependent and may have a time-varying property.
- Other impairments such as Amplified Spontaneous Emission (ASE) noise and intra-channel crosstalk (resulted from optical leaks in optical switches). Different physical impairments increase the BER of an optical packet.

For a detailed description of the impairments, an interested reader is referred to [6, 31, 32, 48, 136, 137].

1.4.6.3 QoS Support by GMPLS In conventional IP networks, each switch/router routes every packet toward its destination. However, there are some applications such as video on demand, IP Television, and voice over IP that generate a flow of packets destined to the same destination. In this case, the same routing must be repeated for every packet at each router. By providing a connection-oriented communication between the source and destination of a real-time communication and using the MPLS protocol, one can reduce the complexity of routing in routers/switches and improve packet performance in the core of the networks. MPLS forwards data from one network node to the next node based on short and fixed-length labels rather than long-length IP addresses. MPLS operates at layer 2.5 that lies between layer 2 (data link layer) and layer 3 (network layer). Generalized MPLS (GMPLS) has been developed to generalize the MPLS concept in optical networks. Both MPLS and GMPLS can provide QoS support in optical networks as discussed in the following.

1.4.6.3.1 MPLS Operation The MPLS maps IP routing into a type of link layer connection referred to as Label Switching Path (LSP), which is actually a unidirectional virtual circuit. In other words, an LSP is functionally an IP level route. However, the individual packet forwarding is optimized since most of the packets only follow the LSP at the link layer (i.e., switching) and do not need an IP level routing operation. MPLS-aware routers are in fact routers/switches, each called a Label Switching Router (LSR). A set of LSRs creates an MPLS domain [138].

MPLS can separately support various communication sessions between a given source–destination pair. For this purpose, MPLS introduces the concept of Forwarding Equivalent Class (FEC), where an FEC is the set of all packets sharing the same list of criteria (say belong to the same IP traffic class). By coupling FEC and LSP, a flexible way is offered to perform traffic engineering. Traffic engineering controls how traffic flows through a network, redirects the traffic flows in the core network to avoid congestion, provides QoS to end users, and optimizes the utilization of a network resources [138, 139].

An LSP is created by distributing labels to routers/switches on the path. Label distribution can be performed by extension of routing protocols (such as Open Shortest Path First (OSPF) and Border Gateway Protocol (BGP)) or by the Resource Reservation Protocol (RSVP). There is also a special protocol, called Label Distribution Protocol (LDP), designed for this purpose in order to distribute labels between Label Edge Routers (LERs) and LSRs [138].

An MPLS network includes a core with a number of LSRs and a number of Label Edge Routers (LERs) at the edge of the network. Each LER consists of two parts: ingress router (for managing traffic transmission) and egress router (for managing the received traffic). The functions of MPLS components are as follows [139]:

- **MPLS Ingress Router:** An ingress router can establish, modify, reroute, and terminate LSPs by using IP signaling and routing protocols such as BGP and RSVP. Note that an LSP is a kind of virtual circuit identified by its label that is formed by a sequence of LSRs.

There are two messages used by RSVP to establish an LSP. The RSVP PATH message is generated by the ingress router and is forwarded through the network along the path of a LSP. At each hop, the PATH message checks the availability of requested resources and stores this information. The RSVP RESERVATION message is created by the egress router in the MPLS domain and used to confirm the reservation request that was sent earlier with the PATH message.

After establishing an LSP, traffic transmission for that LSP can be performed. For an arriving IP packet, the ingress router determines its FEC, assigns it the label of its relevant LSP, encapsulates it in an MPLS header and creates a labeled IP packet (this process is called label pushing), and then sends the labeled IP packet to the MPLS core router.

- **MPLS Core Router:** An LSR does not examine IP header during forwarding. Instead, it forwards labeled IP packets according to the label swapping paradigm, where the LSR maps particular input label/port of an arriving labeled IP packet to a predetermined output label/port, which is provided during the LSP setup. In other words, packets forwarding decisions in LSRs are made solely based on their labels. An LSR may need to change the label of a labeled IP packet before forwarding it to the next LSR.
- **MPLS Egress Router:** An IP packet is restored from a labeled IP packet at the end of LSP by egress router. This process is called label popping.

The QoS tasks in an MPLS controller include the classification of client packets in ingress switches, the differentiated servicing of optical packets in OPS switches, and traffic engineering that performs operations such as LSP admission control and traffic load management. After creating an LSP, the forwarding hardware on the path is provisioned with the requested traffic engineering parameters to guarantee requested levels of QoS metrics such as traffic bandwidth, delay, jitter, and loss. When traffic of the LSP flows, network devices monitor and report the performance parameters and actual level of resource utilization at each core switch interface to the MPLS controller. The controller can use this information on setting up the future coming LSPs and improving the quality of the current LSP [6].

1.4.6.3.2 GMPLS Operation GMPLS is one of the key techniques to integrate IP over WDM. It can manage different types of switchings other than packet switching. A GMPLS-capable LSR is capable of handling five switching mechanisms such as [138]:

- **Fiber-Switch-Capable (FSC) Switching:** In space domain switching, data flow of an entire input fiber is forwarded to another output fiber. For this, a label represents a single-fiber in a bundle.
- **Waveband-Switch-Capable (WSC) Switching:** In waveband domain switching, data are carried by an incoming waveband and are forwarded to an outgoing waveband. For this, a label represents a single waveband within a fiber.

- **Lambda-Switch-Capable (LSC) Switching:** In wavelength domain switching, data are carried by an incoming wavelength channel and is forwarded to an outgoing wavelength channel. For this, a label represents a single wavelength within a waveband or fiber.
- **Time-Division Multiplex (TDM) Capable Switching:** In time domain switching, data are carried and switched in time slots. For this switching, a label represents a set of time slots within a wavelength.
- **Packet-Switch-Capable (PSC) Switching:** This is the original MPLS switching format where packets are switched based on their labels.

Since MPLS deals only with PSCs, there may be a number of LSPs on the same physical link, and a fine bandwidth allocation can be performed on them. For non-PSC LSPs in GMPLS, bandwidth allocation can be performed in discrete units with far fewer choices and small number of LSPs. In optical networks, capacity of a wavelength is too much and there are a number of wavelength on each fiber. Therefore, handling of huge number of LSPs at the PSC (or even TDM) level in GMPLS and communicating labels for them are complicated issues. One approach to alleviate label communication is to use the LSP hierarchy function of MPLS in which a number of hierarchical labels are assigned to a packet. In this hierarchy-label architecture [138],

- At the top level, there are FSC-LSPs so that each one can accommodate several LSC-LSPs or WSC-LSPs.
- At the next level, there are WSC-LSPs so that each one can group several LSC-LSPs.
- At the next level, an LSC-LSP can aggregate several TDM-LSPs.
- At the lowest level, several PSC-LSPs can be grouped into a TDM-LSP.

In a GMPLS-enabled optical network, there are three types of service blocking [140, 141]:

- **Routing Blocking (RB):** blocking in the route computation phase when the path cannot be found between a source–destination pair.
- **Forward Blocking (FB):** blocking in the forward phase of RSVP (i.e., during sending RSVP PATH messages) due to resource unavailability (say wavelength unavailability) on the computed path links.
- **Backward Blocking (BB):** blocking in the backward phase of RSVP due to outdated information or reservation collision, e.g., more than two RSVP RESERVATION messages concurrently attempt to reserve the same wavelength or part of spectral segments in their common links.

GMPLS can be used in connection-oriented OPS networks, where each optical packet is labeled with a connection ID. In this case, a virtual connection must be setup. Each OPS core switch must be able to switch the optical packets belonging to the same connection through the path the connection passes. In GMPLS-based WDM optical packet-switched networks, optical packets are forwarded and switched in the optical domain, and the LSPs they follow are controlled and established by GMPLS protocols. Different techniques proposed for label coding and detection in optical packet headers under GMPLS have been analyzed and studied in [142].

A QoS routing algorithm in GMPLS can select a near-optimal path to meet the desired QoS of a traffic request, considering the combined topology and resource usage information of both IP and WDM layers [114]. GMPLS can set up connections that satisfy the QoS requirements of end users in terms of quality of transport level, reliability, and other QoS-related constraints. In the following, there is a list of LSP setup request examples made by an application generating high-priority traffic that will create a high-priority FEC:

- The application may request from GMPLS to set up an LSP on which packets will experience low crosstalk. This is possible when the LSP passes through the shortest hop path.
- The application may request from GMPLS to set up an LSP on a reliable path. This is possible when the LSP passes through the links with more protection fibers.
- The application may request from GMPLS to set up an LSP in a balanced manner in the network to reduce traffic loss and network congestion. This is possible when the LSP passes through less-congested links.
- The application may request from GMPLS to set up an LSP on which packets will experience low delay. This is possible when the LSP passes through the shortest distance path.
- The application may request from GMPLS to set up an LSP on a secure path. The GMPLS network must monitor the network and find the paths that are more secure and route a high-priority FEC LSP through those secure paths. Clearly, a low-priority FEC LSP can be routed through a less-secure path.

1.5 Optical OFDM-Based Elastic Optical Networking (EON)

In Section 1.3, a number of optical networking mechanisms (such as OPS, OBS, OCS) have been studied. These mechanisms are called fixed-grid optical networks since they provide rigid nature and inflexibly assign resources to users' traffic demands. For example, in a fixed-grid OCS network, a demand must occupy the whole capacity of a wavelength, even if the demand is less than the full wavelength capacity. Hence, the wavelength capacity may be wasted. The OBS is more flexible than OCS, but it still may waste some bandwidth since traffic is usually dynamic and

fluctuates over time. Note that it is common to over-provision network resources to protect network performance at high traffic loads, and this over-provisioning increases the network cost. In short, an optical network with fine-granularity can save network resources and use them better than a coarse-granularity optical network.

With the growth of users' traffic demands, the need for a network to flexibly assign its existing resources to the demands becomes important. Therefore, a new generation of optical networking called Elastic Optical Network (EON) has been introduced. An EON is a flex-grid optical network that can provide fine nature and flexibly assign the network resources to demands exactly according to the requested amount of arriving demands, thus avoiding wasting of the resources.

The Optical Orthogonal Frequency Division Multiplexing (OOFDM)-based EON is studied in the following. As a new and promising network architecture, there are a variety of issues that should be resolved in the OOFDM-based EON such as [143]:

- improving network planning, traffic engineering, control plane technologies
- enhancing current optical network standards
- developing novel Bandwidth-Variable (BV) transponders and BV Cross Connects (BV-WXCs)
- designing flexible spectrum allocation and routing algorithms
- designing traffic grooming approaches
- designing survivability mechanisms
- designing energy efficient mechanisms and components

1.5.1 OOFDM Modulation

The main difference between the fixed and flex grid optical networks lies in the modulation format. A flex-grid optical network can use the OOFDM modulation to provide superior advantages like high spectrum efficiency, superior tolerance to CD/PMD impairments, robustness against inter-carrier and inter-symbol interference, scalability to ever-increasing data rates based on its subcarrier multiplexing technology, and adaptability to channel conditions.

OOFDM is a method of encoding data on multiple carrier frequencies. OOFDM is a type of FDM used as a multi-carrier modulation method. A large number of closely spaced orthogonal subcarrier signals are used to carry traffic on several parallel channels. Each subcarrier can be modulated with a conventional modulation scheme at a low symbol rate while maintaining total data rate as the conventional single carrier modulation scheme in the same bandwidth.

Since different frequency components of an optical pulse travel with different speeds, the optical pulse is spread out after transmission (i.e., dispersion phenomenon). Therefore, an OOFDM symbol with a large delay spread may cross its symbol

boundary after traversing a long-distance, thus leading to interference with its neighboring symbols (called Inter-Symbol Interference (ISI)). Moreover, OOFDM symbols of different subcarriers are not aligned due to the delay spread, and therefore the critical orthogonality condition for the subcarriers may be lost, resulting in an Inter-Carrier Interference (ICI) penalty [143].

Using OOFDM as a bandwidth-variable and highly spectrum-efficient modulation format can provide flexible and scalable sub-wavelength and super-wavelength granularity in contrast to the fixed-grid optical networks [144, 145]. Unlike in WDM networks where all wavelength channels must be separated by a guard band in the frequency domain in order to prevent from interference, the subcarriers in EON can partially overlap in the frequency because of their orthogonalities. A comprehensive survey of OFDM-based optical high-speed transmission and networking technologies have been presented in [143].

Using OOFDM, EON can make the optical network more efficient and flexible. Besides, EON can generate elastic optical paths (the paths with variable bit rates), use bandwidth variable features, and divide available spectrum flexibly according to traffic demands of clients. In EON, if an arriving demand size is more than a wavelength capacity, multiple channels can be grouped together to construct a super-channel. A super-channel includes multiple very closely spaced channels which can transmit the demand as a single entity. As a multi-carrier system, the OOFDM modulation technique used in optical networks distributes data on several low bit rate subcarriers. The spectrum of orthogonally modulated neighbor subcarriers can overlap, thus increasing transmission spectral efficiency. The OOFDM can serve connections with fine-granularity by the elastic allocation of low bit rate subcarriers according to the connection demands [143, 146, 147].

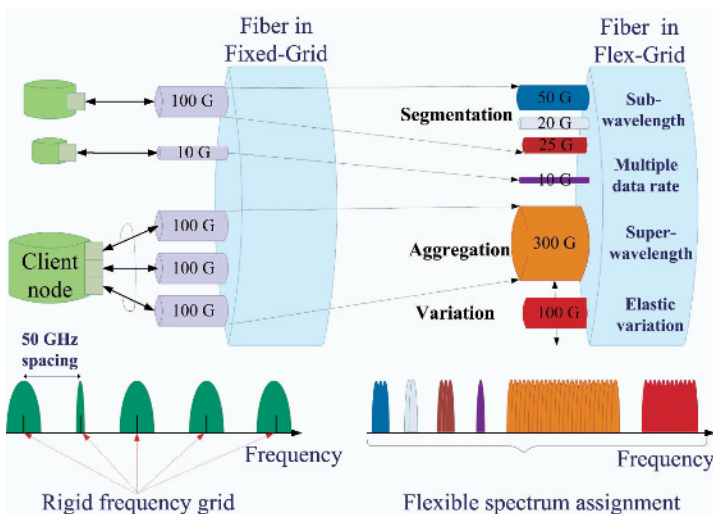


Figure 1.15: An example for spectrum assignment [148]

Figure 1.15 [148] illustrates an example for spectrum assignment in both fixed and flex-grid optical networks, where traffic from different clients with rates 10, 20, 25, 50, 100, and 300 Gbits/s arrive at an ingress switch. Assume data rate of wavelength channels to be 100 Gbits/s in a flex-grid optical network. In this case, the ingress switch must aggregate traffic 20, 25, 50 on one wavelength, use three wavelengths for transmitting 300 Gbits/s traffic, and use a wavelength for transmitting 10 Gbits/s traffic. Under flex-grid networks, these traffic can be separately assigned to different spectrum bands and transmitted (called sub-wavelength traffic accommodation). As it can be observed, the 300 Gbits/s traffic is sent by a number of subcarriers as a whole (called super-wavelength traffic accommodation), and the 10 Gbits/s traffic is directly accommodated in the spectrum using a smaller number of subcarriers, i.e., efficient accommodation of multiple data rates because of the flexible assignment of spectrum. Expansion and contraction of an elastic optical paths is one of the major unique features in EON and some subcarriers should be used for expansion purposes (i.e., variation bandwidth).

1.5.2 EON Components

In OOFDM networks, BV transponders and BV-WXCs are used to generate traffic in ingress switches and switch subcarriers in core nodes, respectively. To obtain high spectral flexibility and serve a client demand, a BV transponder generates an optical signal according to the traffic demand using just enough spectral resources in terms of subcarriers with an appropriate modulation level. Hence, it can provide efficient use of available spectrum resources. For instance, if a subcarrier can provide bandwidth of b bits/s (say 2 Gbits/s), and a client traffic demand needs bandwidth of r bits/s, then the BV transponder can only utilize $\lceil \frac{r}{b} \rceil$ contiguous subcarriers [144].

In traditional WDM networks, all light-paths are assigned the same spectrum width, regardless of the transmission distance of each light-path. This leads to an inefficient utilization of the spectrum. The distance-adaptive spectrum allocation concept can be used in OOFDM-based EON using adaptive modulation and bandwidth variable transponders to improve the spectrum efficiency and QoS for network traffic. There are different modulation formats to be used in OOFDM such as the M -PSK group [e.g., the BPSK (or 2-PSK) and the QPSK (or 4-PSK) groups] and the M -QAM group (e.g., 4-QAM, 8-QAM, 16-QAM, 64-QAM, and 256-QAM). A low-level modulation format with wide spectrum is adapted for long-distance paths, and a high-level modulation format with narrow spectrum can be used for short-distance paths. For this purpose, an adaptive modulation technology can be used to decide what modulation format to be used on which subcarrier under OOFDM according to channel conditions such as reach and OSNR. Here, a subcarrier with high OSNR should be loaded with a high-level modulation format (i.e., more bits loading per symbol), while a low-OSNR subcarrier should use a low-level modulation format (i.e., less bits loading per symbol). For example, let the capacity of a subcarrier using one-bit-per-symbol BPSK modulation be b bit/s. Then, QPSK with 2 bits per symbol provides a capacity of $2 \times b$ bit/s, and 64-QAM with 6 bits per symbol provides a capacity of $6 \times b$ bit/s [143, 147].

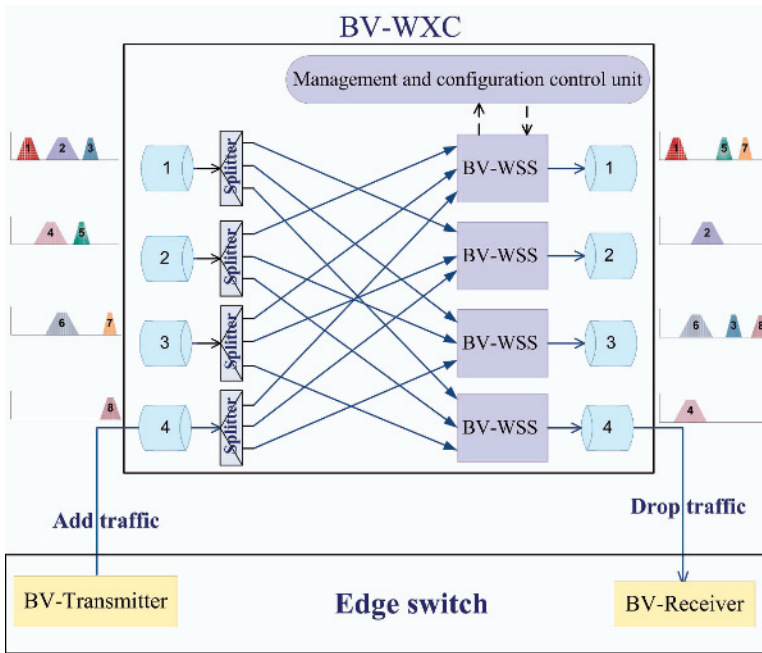


Figure 1.16: Bandwidth-variable WXC switch architecture [149]

A BV-WXC needs to flexibly configure its switching window according to the spectral width of the incoming optical signal. Figure 1.16 shows a 4×4 bandwidth-variable WXC switch architecture that consists of splitters at its input side and Bandwidth Variable Wavelength Selective Switches (BV-WSS) at its output side. The main functions of an intermediate BV-WXC in EON are add-and-drop for local optical signals, and switching/routing of transit optical signals. Splitters split incoming spectrums and send them to appropriate BV-WSSs. A BV-WSS is based on a Reconfigurable Optical Add/Drop Multiplexer (ROADM) to support the flexibility feature of EON networks. The BV-WSSs can receive the spectrums and perform multiplexing/demultiplexing and switching functions using integrated spatial optics. Liquid crystal or MEMS-based BV-WSSs can be employed as switching elements to realize an optical cross-connect with flexible bandwidth and center frequency. The management and configuration control unit assigns optical resources to the setup paths, and it routes the incoming signals by configuring the BV-WSS modules according to the information obtained from a forwarding table that includes the information of the traversing paths. The information of the table is updated when either a connection is established or terminated [143, 149].

Figure 1.16 also depicts switching of traffic coming on different subcarriers. For example, the traffic coming on spectrum segment 4 (where a segment is the contiguous subcarriers allocated to a given connection demand) at input port 2 is delivered to the local edge switch through output port 4, and the traffic on spectrum segment

7 from input port 3 is switched to output port 1. It should be noted that switching in spectrum is possible if there is no overlapping (i.e., no collision) between the switched segments in spectrum at the same output port. In Fig. 1.16, there is no collision among the switched segments at output ports 1 and 3.

1.5.3 Routing and Spectrum Assignment in EON Networks

When a connection demand arrives between a given source–destination pair, the OOFDM network should establish the connection using the Routing and Spectrum Assignment/Allocation (RSA) technique, which can be divided into two subproblems as routing and Spectrum Assignment (SA). There are different routing algorithms and various SA mechanisms developed for RSA. The conventional routing schemes used in OCS networks (such as fixed shortest path routing, fixed alternate routing and K -shortest path routing) can be used for the OOFDM network as well. The aim of routing is to find the shortest feasible path, where the feasible path is determined by the availability of common idle subcarriers along the shortest path.

In SA, there are two constraints that must be satisfied when assigning subcarriers to a connection as [149]:

- **Continuous Constraint:** Under this constraint, the same subcarriers must be used in all links of a path.
- **Contiguous Constraint:** Under this constraint, the subcarriers must be contiguous in the spectrum.

In SA, two issues must be taken into consideration. The first is to avoid spectrum collision, say by randomly assignment of different spectrum segments to different demands. The other issue is to avoid spectrum fragmentation and retain contiguous spectrums as much as possible. There are some heuristic SA mechanisms used for selecting the spectrum resources (i.e., subcarriers) for a given connection demand as [150–152]:

- **First-Fit (FF):** The simple and popular FF mechanism chooses the first idle segment in the spectrum and assigns the starting frequency of the connection to the starting point of the first idle segment. By this way, the FF selects frequency segments in the low-frequency domain and reduces the probability of spectrum fragmentation. However, collision is a significant problem in FF as different nodes may choose the same spectrum resource simultaneously, and this may lead to connection blocking.
- **Random Spectrum (RS):** The RS mechanism chooses the spectrum resource and the left or right side of spectrum randomly. Compared with FF, the RS tends to choose different resources, hence it can reduce the probability of spectrum collision. Obviously, RS splits the spectrum into more fragments, thus increasing the routing blocking.

- **Minimum Residual Spectrum (MR):** The MR mechanism selects the segment that has minimal, but sufficient bandwidth. Other segments with more contiguous subcarriers are kept for future coming connection demands.
- **Adaptive Frequency Assignment-Collision Avoidance (AFA-CA):** This mechanism minimizes the allocated frequency spectrum width in the network. The main idea behind AFA-CA is to assign the contiguous subcarriers to requests without happening any overlapping in contiguous spectrums.

Bandwidth in OOFDM networks can be provided for a connection demand by using either offline (static) RSA or online (dynamic) RSA [153]:

- **Offline Bandwidth Provisioning:** A traffic matrix is given with the requested transmission rates of all connections, and the network should allocate paths and spectrum resources in a way to both serve the connections and minimize the utilized spectrum. However, no spectrum overlapping on any given link is allowed among these connections. An offline RSA may be formulated as an Integer Linear Programming (ILP) problem that returns the optimum solution through a combined routing and spectrum assignment. The ILP formulations can find optimum or near-optimum solutions for small EON networks. However, they are not scalable to large networks [143, 151, 154–156].

For example, Balanced Load Spectrum Allocation (BLSA) provides routing with balancing load within a network in order to minimize the maximum number of subcarriers used on each fiber. The BLSA supports static traffic and performs routing and spectrum assignment according to a given traffic matrix. The BLSA includes three steps as path generation that uses the k -shortest path algorithm to generate k paths, path selection that selects the path based on load balancing within all the fibers in the network, and spectrum allocation [156].

- **Online Bandwidth Provisioning:** OOFDM provides bandwidth dynamically for the demands that their traffic rates fluctuate over time. In other words, each BV transmitter can adaptively change its modulation/coding format to deal with either physical layer impairments or traffic flow variances. In this case, OOFDM should assign a smaller or a larger number of subcarriers to the setup connection that its traffic rate is changing over time. However, the assigned subcarriers must be adjacent. Performance evaluation results show that OOFDM networks can provide lower connection blocking than OCS networks [141, 144, 149, 150, 157].

QoS issues should also be considered in RSA. QoS, as stated in Section 1.4.6.2, can be provided in three ways. There are few works that have considered QoS in RSA. In the following, RSA with QoS improvement, RSA considering QoS differentiation, and RSA considering QoT-aware are reviewed:

- As stated in Section 1.4.6.3.2, backward blocking (collision) is an important reason for blocking in GMPLS-enabled optical networks. The backward collision would happen in EON more frequently than that in OCS networking. Therefore, a collision avoiding routing mechanism with maximum

available spectrum allocation is proposed in [141], where the SA searches for the spectrum segment with minimum collision possibility in order to reduce backward blocking. This issue can improve the network performance parameters by distributing the traffic more evenly among the links (i.e., load-balancing), thus improving QoS.

- In allocating bandwidth, the network can pre-reserve r contiguous subcarriers for a connection demand to be established between a given source–destination pair. The number of reserved subcarriers depends on different parameters such as average traffic rate, maximum traffic rate, burstiness parameter, number of hops of the shortest-hop path, and quality of service type of the connection demand. For Best Effort (BE) connection demands, we usually have $r = 0$, while for guaranteed high-priority connection demands, we must have $r > 0$. Therefore, there are two types of subcarriers: pre-reserved (determined during connection setup) and shared (the subcarriers that can be allocated or deallocated to the connections according to their time-varying traffic rates). Note that if the traffic rate of a high-priority connection reduces and the number of required subcarriers m is less than r , the pre-reserved subcarriers are not released. However, when $m > r$, the network should assign $m - r$ additional contiguous subcarriers to the connection. However, considering both continuous and contiguous constraints, this extension may not be possible, and therefore the extension of bandwidth for the connection may be blocked. Now, assume that $m - r$ additional contiguous subcarriers can be allocated to the connection, and after a while traffic rate reduces. Then, the additional subcarriers can be released, but the network should always keep at least r subcarriers for the connection [144].
- The Quality of Transmission (QoT) can be taken into account in EON so that RSA tries to assign paths with the least amount of Optical Signal-to-Noise Ratio (OSNR), however without any traffic classification. In an EON, when a new connection demand arrives, the proposed dynamic QoT-aware RSA starts finding all the candidate paths and specifying the common unoccupied spectrum of each path with the appropriate modulation format. Then, it chooses a path that qualifies the required QoT constraint in terms of QoT metrics (such as OSNR) and spectrum efficiency. Finally, it assigns spectrum segment based on the allocated modulation level and requested bit rate using traffic balancing-spectrum assignment method, which uniformly distributes the traffic load over the optical spectrum. If there is no path available in the path computation phase, the demand is blocked. If no path can be chosen in the path selection phase, the blocking is because of the poor QoT conditions [149].

The impairment-aware routing and subcarrier allocation (IARSA) finds the shortest feasible path, where the availability of free subcarriers along a path determines the feasibility of the path by considering impairment levels and use of regenerators. The IARSA tries to balance traffic flows evenly across the network to reduce the blocking probability. It finds the path that does

not violate the QoT threshold, where maximum distance is considered as the impairment constraint [152].

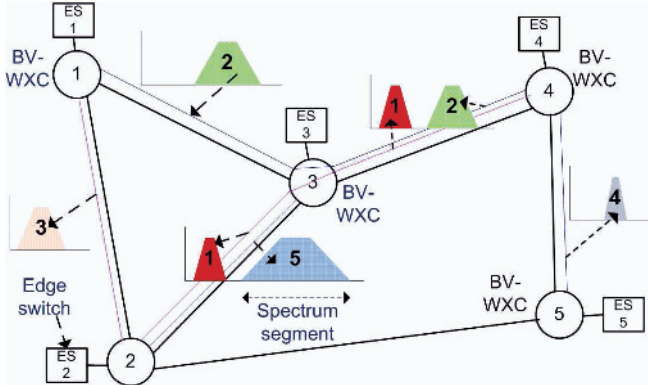


Figure 1.17: An example for connection establishment

Figure 1.17 displays an example for connection establishment under EON, where there are five setup connections: (1) one-hop connection from edge switch 1 (ES 1) to ES 2 on spectrum segment 3, (2) one-hop connection from ES 4 to ES 5 on spectrum segment 4, (3) one-hop connection from ES 2 to ES 3 on spectrum segment 5, (4) two-hop connection from ES 2 to ES 4 on spectrum segment 1, and (5) two-hop connection from ES 1 to ES 4 on spectrum segment 2. As can be observed, these connections are similar to the setup light-paths depicted in Fig. 1.6, where a whole wavelength channel is allocated to each light-path.

One important technique to have a successful RSA is fragmentation. It is possible that there are no sufficient resources to accommodate a connection demand in available paths. One way to solve this problem is to fragment the demand into multiple low-bit-rate demands and route them from separate paths [153].

1.5.4 IP over OOFDM

As studied in previous sections, OOFDM networks focus on connection-oriented networks and can be a good replacement for OCS networks. The OFDM-based EON offers finer bandwidth granularity than WDM networks, and a coarser granularity than OPS networks. Hence, it can be considered as a middle-term alternative to the OPS technology. However, few works have considered packet-based networking under OOFDM [157–161].

A programmable and adaptive IP over OOFDM network architecture and subcarrier resource assignment for IP/OOFDM network have been proposed in [157, 159, 160], where virtual links are set up between every ingress–egress switch pair to build a virtual network topology among them. The virtual links are isolated from each other using different OOFDM subcarriers on each physical link. According to traffic rate variations, the bandwidth of each virtual link can be changed adaptively by

adjusting either the modulation format used on a subcarrier or the number of subcarriers. As stated in Section 1.5.2, an OOFDM network can use different modulation formats on each subcarrier; therefore, different bandwidths can be provided for the traffic on different subcarriers (e.g., using QPSK instead of BPSK doubles the subcarrier capacity). Backbone routers can periodically measure traffic for the virtual links. Then, when a significant change occurs in arriving traffic rate, the bandwidth adjustment decisions can be made. Performance evaluation results show that the adaptive IP/OOFDM network can save up to 30% of receivers and up to 10% total cost in comparison to an IP over TDM/WDM network [143, 159].

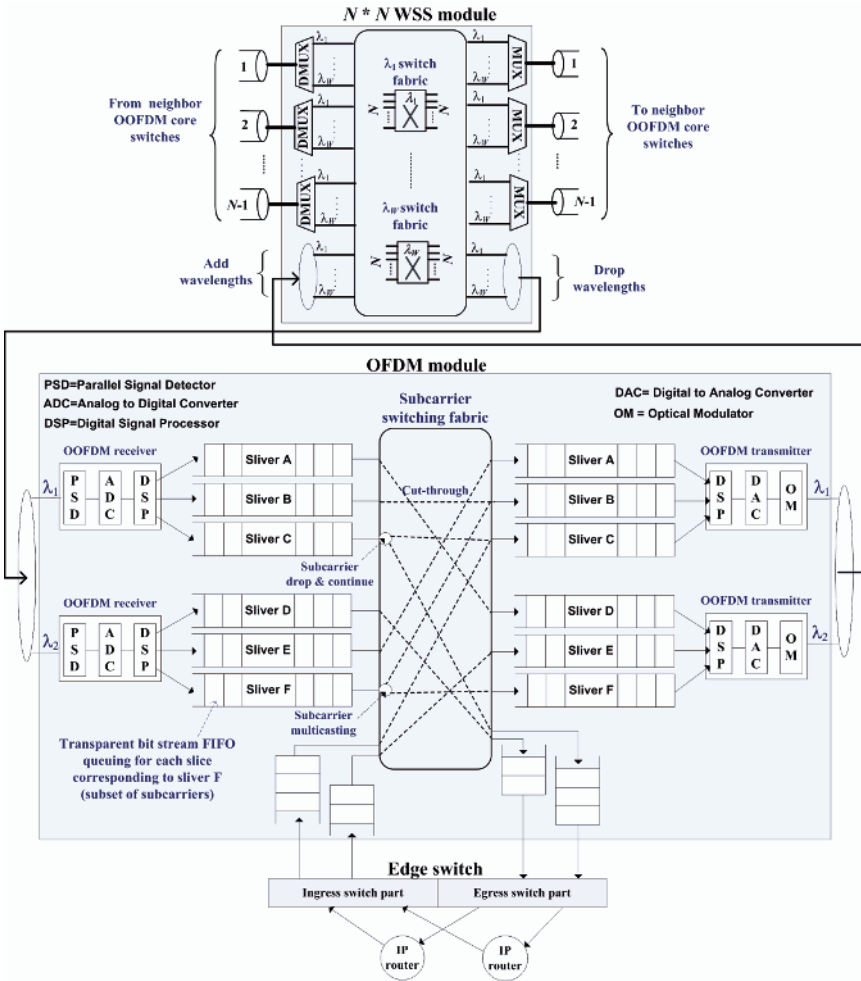


Figure 1.18: An OOFDM core switch model under IP over OOFDM [159]

To implement IP over OOFDM, a two-level hierarchical switching mechanism is used in OOFDM core switches (see Fig. 1.18). The top switching level (i.e.,

Wavelength-Selective Switch (WSS) module) switches wavelengths, whereas the bottom switching level (i.e., sub-wavelength OOFDM module) switches subcarriers. There are a number of dropped wavelengths from the WSS module to the OOFDM module and some add wavelengths from the OOFDM module to the WSS module. In addition, there are some add/drop traffic flow links from/to local ingress/egress switch. Each edge switch may be connected to a number of IP routers. Figure 1.18 shows the OOFDM module for the case that there are two wavelengths λ_1 and λ_2 in the network (i.e., $W = 2$).

Define sliver to be a set of subcarriers that can carry arbitrary format of user-defined data. Hence, each sliver may carry different types of packets, bursts, and even jumbo frames for various types of future Internet traffic. In Fig. 1.18, it is assumed that there are three slivers in each wavelength channel. For example, Slivers A, B, and C inside wavelength λ_1 may carry burst traffic, variable-length packet traffic, and fixed-length packet traffic, respectively.

Figure 1.18 presents the detailed architecture of the OOFDM module for two wavelengths and three slivers in each wavelength. The following components can be found in this architecture:

- OOFDM receivers including Parallel Signal Detector (PSD), high-speed Analog to Digital Converter (ADC), and Digital Signal Processing (DSP) for the functionalities such as QAM demodulation and Fast Fourier Transform (FFT) computation.
- Sliver input FIFOs for queuing of transparent bit stream.
- Subcarrier switching fabric for multi-rate sliver switch supporting functionalities such as subcarrier cut-through switching, subcarrier drop-and-continue switching, and subcarrier multicasting switching as depicted in the figure. The local traffic is also managed in the switching fabric as the source traffic flows from IP router and the sink traffic flows to IP router.
- Sliver output FIFOs for saving switched bit stream which is ready for optical OFDM multiplexing/modulation.
- OOFDM transmitters including DSP module for the functionalities such as QAM modulation and Inverse Fast Fourier Transform (IFFT), high-speed Digital to Analog Converter (DAC), and Optical Modulator (OM).

It should be noted that a time-division multiplexing/multiple access (TDM/TDMA)-based grooming switch (say SONET/WDM) can also be used as a subwavelength switching mechanism. However, the proposed OOFDM switch architecture provides unique flexibility with high resource efficiency when dealing with bandwidth resource sharing. It utilizes DSP for subcarrier multiplexing and grooming combined with adaptive transmission and parallel signal detection. In addition, each OOFDM transmitter can dynamically modify its modulation/coding methods in order to deal with packet traffic variance at the subwavelength level. Moreover, adaptive load balancing can be obtained by mapping virtual topology to variable packet flows in real

time. Besides, the OOFDM switch utilizes a receiver structure (i.e., PSD) that can receive multiple non-overlapping baseband signals on the same wavelength simultaneously. Note that no more than one wavelength can enter a PSD. Finally, no global synchronization and complicated multiplexing hierarchy are needed in the proposed OOFDM switch [162].

1.6 Summary

Future Internet requires both bandwidth and QoS because of an ever-increasing number of Internet users and a growing number of real-time applications (such as video on demand, voice over IP, IPTV, video conferencing, etc.) that require different levels of QoS. This could be a problem when the network bandwidth is limited, when the network supports only the best effort traffic, or when the Internet traffic does not have a uniform characteristic.

Different optical switching mechanisms in two versions of asynchronous and slotted have been detailed in this chapter. Among the proposed switching schemes, the OCS network is very popular in simplifying traffic transmission. However, it is neither efficient on the bandwidth usage nor scalable. Contention-based schemes seems to be appropriate in quickly responding to traffic dynamics such as Internet traffic. The OPS as a contention-based switching scheme appears to provide the finest granularity compared with OCS and OBS. Combination of huge bandwidth offered by optical technology with the flexibility of packet switching in OPS can be a very good candidate for future optical networks, for both metro and backbone networks.

Different aspects of OPS including OPS switch architecture, edge switch architecture, signaling in OPS, contention problem in OPS, and QoS in OPS have been studied in this chapter. As mentioned, contention is the major problem in OPS networks. Unfortunately, many proposed architectures resolve the contention problem in OPS by using immature optical buffering technology (using optical fiber delay lines), which is expensive, complex, and bulky. In addition, optical buffering degrades the signal power of an optical packet, increases its delay, and makes its arrival at its destination out of order. To reduce these problems, OPS core switches should reduce using optical buffers for contention resolution. In addition, slotted networks can provide a much finer granularity in bandwidth sharing and can improve the bandwidth usage. As stated, next generation optical networks must support the applications that require different levels of QoS. Based on all the above discussions, synchronized buffer-less all-optical packet-switched networking under the DiffServ domain could be a very good candidate for efficient transmission of traffic in optical networks.

In short, since QoS support is a requirement for any next-generation network and OPS networks are the promising next generation optical networks, this book will provide a comprehensive review of handling QoS issues in OPS networks.

Finally, OOFDM networks have been reviewed that can flexibly provide bandwidth for connection demands, while reducing the network bandwidth wastage. Using the OOFDM technology, the novel elastic optical network architecture with

great flexibility and scalability on allocating spectrum and accommodating data rate has been studied. Most of the current OOFDM network designs have focused on connection-oriented networks, and therefore they can be a good replacement for OCS networks. However, few works have considered packet-based networking under OOFDM.

REFERENCES

1. M. OMahony, D. Simeonidou, D. Hunter, and A. Tzanakaki. The application of optical packet switching in future communications networks. *IEEE Communications Magazine*, 39(3):128–135, Mar. 2001.
2. D. K. Hunter and I. Andonovic. Approaches to optical Internet packet switching. *IEEE Communications Magazine*, 38(9):116–122, Sept. 2000.
3. J. Ryan. Fiber considerations for metropolitan networks. *Alcatel Telecommunications Review*, pages 52–56, 2002.
4. Fiber types in gigabit optical communications. Cisco Public Information, 2008.
5. R. Ramaswami. Optical fiber communication: From transmission to networking. *IEEE Communications Magazine*, 40:138–147, 2002.
6. S. V. Kartalopoulos. *DWDM: Networks, Devices, and Technology*. Wiley-IEEE Press, 2003.
7. Y. Huang, J. P. Heritage, and B. Mukherjee. Connection provisioning with transmission impairment consideration in optical WDM networks with high-speed channels. *IEEE Journal of Lightwave Technology*, 23(3):982–993, 2005.
8. R. Ramaswami, K. Sivarajan, and G. Sasaki. *Optical Networks: A Practical Perspective*. 3rd edition, Morgan Kaufmann, 2009.
9. A. M. Odlyzko. Internet traffic growth: sources and implications. In *SPIE Optical Transmission Systems and Equipment for WDM Networking II*, volume 5247, pages 1–15, Orlando, FL, 2003.

10. X. Xiao and L. M. Ni. Internet QoS: a big picture. *IEEE Network*, 13(2):8–18, 1999.
11. F. Callegati, G. Corazza, and C. Raffaelli. Exploitation of DWDM for optical packet switching with quality of service guarantees. *IEEE Journal on Selected Areas in Communications*, 20:190–201, 2002.
12. G. I. Papadimitriou, M. S. Obaidat, and A. S. Pomportsis. Advances in optical networking. *International Journal of Communication Systems*, 15(2-3):101–113, 2002.
13. A. K. Somani, M. Mina, and L. Li. On trading wavelengths with fibers: a cost performance based study. *IEEE/ACM Transactions on Networking*, 12:944–951, 2004.
14. O. Gerstel and H. Raza. Merits of low-density WDM line systems for long-haul networks. *IEEE/OSA Journal of Lightwave Technology*, 21(11):2470–2475, Nov. 2003.
15. A. G. Rahbar and O. Yang. Fiber-channel tradeoff for reducing collisions in slotted single-hop optical packet-switched (OPS) networks. *Journal of Optical Networking*, 6:897–912, 2007.
16. G. Muretto and C. Raffaelli. Combining contention resolution schemes in WDM optical packet switches with multifiber interfaces. *Journal of Optical Networking*, 6:74–89, 2007.
17. M. Dhodhi, S. Tariq, and K. Saleh. Bottlenecks in next generation DWDM-based optical networks. *Computer Communications*, 24(17):1726–1733, 2001.
18. T. E. Stern, G. Ellinas, and K. Bala. *Multiwavelength Optical Networks: Architectures, Design, and Control*. Second edition, Cambridge University Press, 2008.
19. G. Shen and R. S. Tucke. Translucent optical networks: the way forward. *IEEE Communications Magazine*, 45(2):48–54, Feb. 2007.
20. M. J. Potasek. All-optical switching for high bandwidth optical networks. *Optical Networks Magazine*, 3(6):30–43, 2002.
21. C. Papazoglou, G. Papadimitriou, and A. Pomportsis. Design alternatives for optical-packet-interconnection network architectures. *OSA Journal of Optical Networking*, 3(11):810–825, 2004.
22. M. J. OMahony, C. Politi, D. Klonidis, R. Nejabati, and D. Simeonidou. Future optical networks. *IEEE Journal of Lightwave Technology*, 24(12):4684–4696, Dec. 2006.
23. M. Klinkowski, D. Careglio, and J. Sol-Pareta. Wavelength vs. burst vs. packet switching: comparison of optical network models. In *The e-Photon/ONe Winter School workshop*, Aveiro, Portugal, 2005.
24. G. N. Rouskas and H. G. Perros. A tutorial on optical networks. *Lecture Notes in Computer Science*, 2497:155–193, 2002.
25. H. Zang, J. P. Jue, and B. Mukherjee. A review of routing and wavelength assignment approaches for wavelength routed optical WDM networks. *Optical Networks Magazine*, 1(1):47–60, Jan. 2000.
26. T. E. Stern and K. Bala. *Multiwavelength Optical Networks: A Layered Approach*. Prentice Hall, 2000.
27. A. G. Rahbar. A review of dynamic impairment-aware routing and wavelength assignment techniques in all-optical wavelength-routed networks. *IEEE Communications Surveys and Tutorials*, 14(2):1065–1089, 2012.

28. S. Barkabi and A. G. Rahbar. MC-BRM: A distributed RWA algorithm with minimum wavelength conversion. *Elsevier Optical Fiber Technology*, 18(4):230–241, 2012.
29. S. Barkabi and A. G. Rahbar. LC-DLA: A fair congestion-aware distributed lightpath allocation mechanism. *Elsevier OPTIK—International Journal for Light and Electron Optics*, 123(15):1370–1377, Aug. 2012.
30. A. G. Rahbar. Dynamic impairment-aware RWA in multi-fiber wavelength-routed all-optical networks supporting class-based traffic. *IEEE Journal of Optical Communications and Networking*, 2(11):915–927, Nov. 2010.
31. I. Tomkos, D. Vogiatzis, C. Mas, I. Zacharopoulos, A. Tzanakaki, and E. Varvarigos. Performance engineering of metropolitan area optical networks through impairment constraint routing. *IEEE Communications Magazine*, 42(8):S40–S47, Aug. 2004.
32. C. Saradhi and S. Subramaniam. Physical layer impairment aware routing (PLIAR) in WDM optical networks: issues and challenges. *IEEE Communications Surveys & Tutorials*, 11(4):109–130, Dec. 2009.
33. A. Mokhtar and M. Azizoglu. Adaptive wavelength routing in all-optical networks. *IEEE/ACM Transactions on Networking*, 6(2):197–206, Apr. 1998.
34. H. Zang, J. P. Jue, L. Sahasrabudhe, R. Ramamurthy, and B. Mukherjee. Dynamic light path establishment in wavelength-routed WDM networks. *IEEE Communications Magazine*, pages 100–117, 2001.
35. A. Askarian, Y. Zhai, S. Subramaniam, Y. Pointurier, and M. Brandt-Pearce. QoT-aware RWA algorithms for fast failure recovery in all-optical networks. In *IEEE/OSA OFC'08*, San Diego, CA, 2008.
36. N. A. Shalmany and A. G. Rahbar. On the choice of all-optical switches for optical networking. In *IEEE International Symposium on High Capacity Optical Networks and Enabling Technologies (HONET07)*, UAE, Dubai, 2007.
37. C. Clos. A study of non-blocking switching networks. *Bell System Technical Journal*, 32:406–424, 1953.
38. X. Chu, H. Yin, and X. Y. Li. Lightpath rerouting in wavelength routed WDM networks. *Journal of Optical Networking*, 7(8):721–735, 2008.
39. G. Mohan and C. S. R. Murthy. A time optimal wavelength rerouting algorithm for dynamic traffic in WDM networks. *IEEE Journal of Lightwave Technology*, 17(3):406–417, Mar. 1999.
40. K. Zhu and B. Mukherjee. A review of traffic grooming in WDM optical networks: architectures and challenges. *Optical Networks Magazine*, 4(2):55–64, 2003.
41. J. Liu, G. Xiao, and W. Wang. On the performance of distributed light-path provisioning with dynamic routing and wavelength assignment. *Photonic Networks Communications*, 17:191–201, 2009.
42. G. S. Pavani, L. G. Zuliani, H. Waldman, and M. Magalhes. Distributed approaches for impairment-aware routing and wavelength assignment algorithms in GMPLS networks. *Computer Networks*, 52:1905–1915, 2008.
43. P. Singh, A. K. Sharma, and S. Rani. Minimum connection count wavelength assignment strategy for WDM optical networks. *Optical Fiber Technology*, 14:154–159, 2008.

44. C. S. R. Murty and M. Gurusamy. *WDM Optical Networks: Concepts, Design, and Algorithms*. Prentice-Hall India, 2002.
45. A. Szymanski, A. Lason, J. Rzasa, and A. Jaisczyk. Grade-of-service-based routing in optical networks. *IEEE Communications Magazine*, 45(2):82–87, Feb. 2007.
46. M. Gagnaire and S. Zahr. Impairment-aware routing and wavelength assignment in translucent networks: State of the art. *IEEE Communications Magazine*, 47(5):55–61, May 2009.
47. E. Salvadori, Y. Ye, C. Saradhi, A. Zanardi, H. Woesner, M. Carcagni, G. Galimberti, G. Martinelli, A. Tanzi, and D. La Fauci. Distributed optical control plane architectures for handling transmission impairments in transparent optical networks. *IEEE Journal of Lightwave Technology*, 27(13):2224–2239, Jul. 2009.
48. S. Azodolmolky, M. Klinkowski, E. Marin, D. Careglio, J. Pareta, and I. Tomkos. A survey on physical layer impairments aware routing and wavelength assignment algorithms in optical networks. *Computer Networks*, 53(7):926–944, May 2009.
49. S. Azodolmolky, Y. Pointurier, M. Klinkowski, E. Marin, D. Careglio, J. Sol-Pareta, M. Angelou, and I. Tomkos. On the offline physical layer impairment aware RWA algorithms in transparent optical networks: state-of-the-art and beyond. In *Optical Network Design and Modeling (ONDM2009)*, Braunschweig, Germany, 2009.
50. E. V. Breusegem, J. Cheyns, D. D. Winter, D. Colle, M. Pickavet, P. Demeester, and J. Moreau. A broad view on overspill routing in optical networks: a real synthesis of packet and circuit switching? *Elsevier Optical Switching and Networking*, 1(1):51–64, 2004.
51. L. Xu, H. Perros, and G. Rouskas. Techniques for optical packet switching and optical burst switching. *IEEE Communications Magazine*, 39(1):136–142, Jan. 2001.
52. M. Nord, S. Bjørnstad, and C. M. Gauger. OPS or OBS in the core network? In *IFIP Working Conference on Optical Network Design and Modeling*, Budapest, Feb. 2003.
53. S. Yao, B. Mukherjee, and S. Dixit. Advances in photonic packet switching: an overview. *IEEE Communications Magazine*, 38(2):84–94, Feb. 2000.
54. M. Yoo and C. Qiao. Optical burst switching (OBS): a new paradigm for an optical Internet. *Journal of High Speed Networks*, 8(1):69–84, 1999.
55. K. C. Chua, M. Gurusamy, Y. Liu, and M. H. Phung. *Quality of Service in Optical Burst Switched Networks*. Springer, 2007.
56. J. Liu and N. Ansari. The impact of the burst assembly interval on the OBS ingress traffic characteristics and system performance. In *IEEE International Conference on Communications (ICC)*, Paris, France, June 2004.
57. S. J. B. Yoo, F. Xue, Y. Bansal, J. Taylor, P. Zhong, C. Jing, J. Minyong, T. Nady, G. Goncher, K. Boyer, K. Okamoto, S. Kamei, and V. Akella. High-performance optical-label switching packet routers and smart edge routers for the next generation Internet. *IEEE Journal on Selected Areas in Communications*, 21(7):1041–1051, Sept. 2003.
58. C. Raffaelli and P. Zaffoni. TCP performance in optical packet-switched networks. *Photonic Network Communications*, 11(3):243–252, May 2006.
59. K. Long, R. S. Tucker, and C. Wang. A new framework and burst assembly for IP diffserv over optical burst switching networks. In *IEEE Globecom*, San Francisco, CA, 2003.

60. V. Vokkarane and J. P. Jue. Prioritized burst segmentation and composite burst assembly techniques for QoS support in optical burst switched networks. *IEEE Journal on Selected Areas in Commun.*, 21(7):1198–1209, Sept. 2003.
61. S. Seo, K. Bergman, and P. R. Prucnal. Transparent optical networks with time-division multiplexing. *IEEE Journal on Selected Areas in Communications*, 14(5):1039–1051, June 1996.
62. N. F. Huang, G. H. Liaw, and C. P. Wang. A novel all optical transport network with time-shared wavelength channels. *IEEE Journal on Selected Areas in Communications*, 18(10):1863–1875, Oct. 2000.
63. A. Chen, A. K. Wong, and C. T. Lea. Routing and time-slot assignment in optical TDM networks. *IEEE Journal on Selected Areas in Communications*, 22(9):1648–1657, Nov. 2004.
64. B. Wen and K. M. Sivalingam. Routing, wavelength and time-slot assignment in time division multiplexed wavelength-routed optical WDM networks. In *IEEE Infocom*, New York, 2002.
65. T. S. El-Bawab and J. Shin. Optical packet switching in core networks: between vision and reality. *IEEE Communications Magazine*, 40(9):60–65, Sept. 2002.
66. Caida. Packet size reports. Technical report, <http://www.caida.org/research/traffic-analysis/fix-west-1998/packetsizes/>, 1998.
67. J. Ramamirtham and J. Turner. Time sliced optical burst switching. In *IEEE Infocom*, pages 2030–2038, San Francisco, Mar. 2003.
68. I. Chlamtac, V. Elek, A. Fumagalli, and C. Szabo. Scalable WDM network architecture based on photonic slot routing and switched delay lines. In *IEEE Infocom*, Kobe, Japan, Apr. 1997.
69. I. Chlamtac, A. Fumagalli, and G. Wedzinga. Slot routing as a solution for optically transparent scalable WDM wide area networks. *Photonic Network Communications*, 1(9):9–21, 1999.
70. M. Maier. *Optical switching networks*. Cambridge University Press, 2008.
71. J. Fang, W. He, and A. K. Somani. Optimal light trail design in WDM optical networks. In *IEEE ICC*, pages 1699–1703, Paris, June 2004.
72. M. Maier and M. Reisslein. AWG-based metro WDM networking. *IEEE Communications Magazine*, 42(11):S19–S26, Nov. 2004.
73. S. Bjornstad, M. Nord, D. R. Hjelle, and N. Stol. Optical burst and packet switching: node and network design, contention resolution and quality of service. In *IEEE International Conference on Telecommunications (ConTEL2003)*, Zagreb, Croatia, June 2003.
74. H. S. Yang, M. Herzog, M. Maier, and M. Reisslein. Metro WDM networks: performance comparison of slotted ring and AWG star networks. *IEEE Journal on Selected Areas in Communications*, 22(8):1460–1473, Oct. 2004.
75. A. Aggarwal, A. Bar-Noy, D. Coppersmith, R. Ramaswami, B. Schieber, and M. Sudan. Efficient routing in optical networks. *Journal of the ACM*, 43(6):973–1001, 1996.
76. H. Y. Tyan, J. C. Hou, B. Wang, and C. C. Han. On supporting temporal quality of service in WDMA-based star-coupled optical networks. *IEEE Transactions on Computers*, 50(3):197–214, 2001.

77. P. Sarigiannidis, G. Papadimitriou, and A. Pomportsis. CS-POSA: a high performance scheduling algorithm for WDM star networks. *Photonic Network Communications*, 11(2):211–227, 2006.
78. C. Fan, S. Adams, and M. Reisslein. The FT-FR AWG network: a practical single-hop metro WDM network for efficient uni- and multicasting. *IEEE Journal of Lightwave Technology*, 23(3):937–954, 2005.
79. M. Jin and O. Yang. APOSN: operation, modeling and performance evaluation. *Computer Networks*, 51(6):1643–1659, 2007.
80. C. Peng, G. v. Bochmann, and T. J. Hall. Quick Birkhoff–von Neumann decomposition algorithm for agile all-photonic network cores. In *IEEE ICC*, Istanbul, Turkey, 2006.
81. A. G. Rahbar and O. Yang. An integrated TDM architecture for AAPN networks. In *SPIE Photonics North*, volume 5970, Toronto, Canada, Sept. 2005.
82. A. G. Rahbar. EBvN: efficient BvN in multi-fiber/multi-wavelength overlaid-star optical networks. *Springer Journal of Annals of Telecommunications*, 67(11-12):575–588, Nov. 2012.
83. A. G. Rahbar and O. Yang. Agile bandwidth management techniques in slotted all-optical packet interconnection networks. *Elsevier Computer Networks*, 54(3):387–403, Feb. 2010.
84. N. Saberi and M. Coates. Scheduling in overlaid star all-photonic networks with large propagation delays. *Photonic Network Communications*, 17(2):157–169, 2009.
85. M. Jin and O. Yang. A TDM solution for all-photonic overlaid-star networks. In *CISS2006*, pages 1691–1695, Princeton, N. J., 2006.
86. L. Mason, A. Vinokurov, N. Zhao, and D. Plant. Topological design and dimensioning of agile all-photonic networks. *Computer Networks*, 50(2):268–287, Feb. 2006.
87. S. A. Paredes, G. Bochmann, and T. J. Hall. Deploying agile photonic networks over reconfigurable optical networks. In *IEEE Symposium on Computers and Communications (ISCC)*, pages 182–187, Sousse, Tunisia, 2009.
88. F. J. Blouin, A. W. Lee, A. J. M. Lee, and M. Beshai. Comparison of two optical-core networks. *Journal of Optical Networking*, 1(1):56–65, Jan. 2002.
89. A. Carena, V. D. Feo, J. M. Finochietto, R. Gaudino, F. Neri, C. Piglioneri, and P. Poggiolini. Ringo: an experimental WDM optical packet network for metro applications. *IEEE Journal on Selected Areas in Communications*, 22(8):1561 – 1571, Oct. 2004.
90. I. M. White, M. S. Rogge, K. Shrikhande, and L. Kazovsky. A summary of the HOR-NET project: A next-generation metropolitan area network. *IEEE Journal on Selected Areas in Communications*, 21(9):1478–1494, Nov. 2003.
91. L. F. K. Lui, L. Xu, P. K. A. Wai, L. Y. Chan, C. C. Lee, H. Y. Tam, and M. S. Demokan. 1*4 all-optical packet switch with all-optical header processing. In *IEEE Optical Fiber Communication Conference*, California, Mar. 2005.
92. N. Calabretta, M. Presi, G. Contestabile, and E. Ciaramelli. Compact all-optical header processing for DPSK packets. In *European Conference on Optical Communications*, 2006.
93. P. K. AWai, L. Y. Chan, L. F. K. Lui, L. Xu, H. Y. Tam, and M. S. Demokan. All-optical header processing using control signals generated by direct modulation of a DFB laser. *Optics Communications*, 242(1-3):155–161, Nov. 2004.

94. J. P. Wang, B. S. Robinson, S. A. Hamilton, and E. P. Ippen. 40-gbit/s all-optical header processing for packet routing. In *Optical Fiber Communication Conference*, 2006.
95. A. G. Rahbar. Impairment-aware merit-based scheduling in QoS-capable multi-fiber OPS networks. *Journal of Optical Communications*, 33(2):103–121, June 2012.
96. A. Pattavina. Architectures and performance of optical packet switching nodes for IP networks. *IEEE/OSA Journal of Lightwave Technology*, 23(3):1023–1032, Mar. 2005.
97. B. Bostica, M. Burzio, P. Gambini, and L. Zucchelli. *Synchronization issues in optical packet switched networks*, pages 362–376. Springer-Verlag, 1997.
98. J. R. Feehrer and L. H. Ramfelt. Packet synchronization for synchronous optical deflection-routed interconnection networks. *IEEE Transactions on Parallel and Distributed Systems*, 7(6):605–611, 1996.
99. N. Brownlee and K. C. Claffy. Understanding Internet traffic streams: dragonflies and tortoises. *IEEE Communications Magazine*, 40(10):110–117, 2002.
100. V. Jacobson, K. Nichols, and K. Poduri. An expedited forwarding PHB. Technical report, RFC 2598, 1999.
101. J. Heinanen, F. Baker, W. Weiss, and J. Wroclawski. Assured forwarding PHB group. Technical report, RFC 2597, 1999.
102. H. Elbiaze, T. Chahed, T. Atmaca, and G. Hbuterne. Shaping self-similar traffic at access of optical network. *Performance Evaluation*, 53(3-4):187–208, 2003.
103. V. Sivaraman, D. Moreland, and D. Ostry. A novel delay-bounded traffic conditioner for optical edge switches. In *IEEE Workshop on High Performance Switching and Routing (HPSR)*, Hong Kong, May 2005.
104. A. G. Rahbar and O. Yang. Distribution-based bandwidth access scheme in slotted all-optical packet-switched networks. *Elsevier Computer Networks*, 53(5):744–758, Apr. 2009.
105. A. G. Rahbar. Improving throughput of long-hop TCP connections in IP over OPS networks. *Springer Photonic Network Communications*, 17(3):226–237, June 2009.
106. A. Shacham and K. Bergman. An experimental validation of a wavelength-striped, packet switched, optical interconnection network. *Journal of Lightwave Technology*, 27(7):841–850, Apr. 2009.
107. C. P. Lai and K. Bergman. Network architecture and test-bed demonstration of wavelength-striped packet multicasting. In *OFC*, Mar. 2010.
108. Z. Lu, D. K. Hunter, and I. D. Henning. Contention resolution scheme for slotted optical packet switched networks. In *International conference on Optical Network Design and Modelling (ONDM)*, Milan, Italy, Feb. 2005.
109. M. R. Salvador. *MAC protocols for optical packet-switched WDM rings*. PhD thesis, Centre of Telematics and Information Technology, University of Twente, 2003.
110. S. Yao, B. Mukherjee, S. J. B. Yoo, and S. Dixit. A unified study of contention-resolution schemes in optical packet-switched networks. *IEEE Journal of Lightwave Technology*, 21(3):672, 2003.
111. H. Øverby. How the packet length distribution influences the packet loss rate in an optical packet switch. In *Advanced International Conference on Telecommunications and International Conference on Internet and Web Applications and Services (AICT/ICIW 2006)*, 2006.

112. What is quality of service (QoS). <http://www.igi-global.com/dictionary/quality-of-service-qos/24296>. Accessed: 2014-04-30.
113. A. Kavitha. Performance of optical networks: a short survey. *International Journal of Engineering Science and Technology (IJEST)*, 4(2):600–605, Feb. 2012.
114. W. Wei, Q. Zeng, Y. Ouyang, and D. Lomone. Differentiated integrated QoS control in the optical Internet. *IEEE Communications Magazine*, 42(11):S27 – S34, Nov. 2004.
115. N. Golmie, T. D. Ndousse, and D. H. Sue. A differentiated optical services model for WDM networks. *IEEE Communications Magazine*, 38(2):68–73, Feb. 2000.
116. H. Øverby, N. Stol, and M. Nord. Evaluation of QoS differentiation mechanisms in asynchronous bufferless optical packet-switched networks. *IEEE Communications Magazine*, 44(8):52–57, 2006.
117. P. Ellerman. Calculating reliability using FIT & MTTF: Arrhenius HTOL model. Technical report, Microsemi, 2012.
118. R. Jain, D. M. Chiu, and W. Hawe. "a quantitative measure of fairness and discrimination for resource allocation in shared computer systems". DEC research report tr-301. Technical report, Digital Equipment Corporation, 1984.
119. H. Dahmouni, A. Girard, and B. Sans. An analytical model for jitter in IP networks. *Annals of Telecommunications*, 67:81–90, 2012.
120. N. Olifer and V. Olifer. *Computer Networks: Principles, Technologies and Protocols for Network Design*. John Wiley & Sons, 2006.
121. L. F. Mollenauer and J. P. Gordon. *Solitons in Optical Fibers: Fundamentals and Applications*. Academic Press, 2006.
122. A. Kaheel, T. Khatatb, A. Mohamed, and H. Alnuweiri. Quality-of-service mechanisms in IP-over WDM networks. *IEEE Communications Magazine*, 40(12):38–43, Dec. 2002.
123. H. Øverby and N. Stol. Quality of service in asynchronous buffer-less optical packet switched networks. *Telecommunication Systems*, 27(2-4):151–179, Oct.–Dec. 2004.
124. D. Careglio, J. Sol-Pareta, and S. Spadaro. Novel contention resolution technique for QoS support in connection-oriented optical packet switching. In *IEEE ICC*, Seoul, Korea, May 2005.
125. R. Braden, D. Clark, and S. Shenker. Integrated services in the Internet architecture: an overview. Technical report, Internet informational RFC 1633, 1994.
126. S. Blake, D. Black, M. Carlson, E. Davies, Z. Wang, and W. Weiss. An architecture for differentiated services. Technical report, RFC 2475, 1998.
127. A. Striegel and G. Manimaran. Packet scheduling with delay and loss differentiation. *Computer Communications*, 25(1):21–31, Jan. 2002.
128. P. Siripongwutikorn, S. Banerjee, and D. Tipper. A survey of adaptive bandwidth control algorithms. *IEEE Communications Surveys and*, 5(1):2–14., 2003.
129. A. Habib, S. Fahmy, and B. Bhargava. Monitoring and controlling QoS network domains. *ACM/Wiley International Journal of Network Management*, 15(1):11–29, Jan. 2005.
130. W. Zhao, D. Olshefski, and H. Schulzrinne. Internet quality of service: An overview. Technical report, Columbia University, 2001.

131. Q. Zhang, V. M. Vokkarane, B. Chen, and J. P. Jue. Early drop scheme for providing absolute QoS differentiation in optical burst-switched networks. In *High Performance Switching and Routing (HPSR)*, pages 153–157, Torino, Italy, 2003.
132. Y. Chen, C. Qiao, M. Hamdi, and D. H. K. Tsang. Proportional differentiation: a scalable QoS approach. *IEEE Communications Magazine*, 41:52–58, 2003.
133. J. Heinanen, F. Baker, W. Weiss, and J. Wroclawski. AF-forwarding. <http://tools.ietf.org/html/rfc2597>, Jun 1999. RFC2597.
134. D. K. Hunter, M. C. Chia, and I. Andonovic. Buffering in optical packet switching. *IEEE/OSA Journal of Lightwave Technology*, 16(10):2081–2094, Dec. 1998.
135. F. Callegati, W. Cerroni, G. Corazza, and C. Raffaelli. Optical packet switching: a network perspective. In *IEEE Conference on Optical Fiber communication/National Fiber Optic Engineers Conference*, San Diego, CA, Feb. 2008.
136. R. Martinez, C. Pinart, F. Cugini, N. Andriolli, L. Vakarengi, P. Castoldi, L. Wosinska, J. Comellas, and G. Junyent. Challenges and requirements for introducing impairment-awareness into the management and control planes of ASON/GMPLS WDM networks. *IEEE Communications Magazine*, 44(12):76–85, Dec. 2006.
137. G. Agrawal. *Fiber-Optic Communication Systems*. 3rd edition, John Wiley & Sons, 2002.
138. K. Chen and W. Fawaz. *Optical Networks: New Challenges and Paradigms for Quality of Service*, chapter 5, pages 115–134. John Wiley & Sons, 2009.
139. X. Xiao, A. Hannan, B. Bailey, and L. M. Ni. Traffic engineering with MPLS in the Internet. *IEEE Network*, 14(2):28–33, Aug. 2000.
140. A. Giorgetti, N. Sambo, I. Cerutti, N. Andriolli, and P. Castoldi. Suggested vector scheme with crankback mechanism in GMPLS-controlled optical networks. In *Optical Network Design Modeling (ONDM)*, Kyoto, Japan, 2010.
141. B. Guo, J. Li, Y. Wang, S. Huang, Z. Chen, and Y. He. Collision-aware routing and spectrum assignment in GMPLS-enabled flexible-bandwidth optical network. *IEEE Journal of Optical Communications and Networking*, 5(6):658–666, 2013.
142. J. V. Gontn, P. P. Mario, J. V. Alonso, and J. G. Haro. A feasibility study of GMPLS extensions for synchronous slotted optical packet switching networks. In *The V Workshop in G/MPLS Networks*, pages 89–100, Girona, Spain, 2006.
143. G. Zhang, M. D. Leenheer, A. Morea, and B. Mukherjee. A survey on OFDM-based elastic core optical networking. *IEEE Communications Surveys and Tutorials Journal*, 15(1):65–87, FIRST QUARTER 2013.
144. K. Christodoulopoulos, I. Tomkos, and E. Varvarigos. Dynamic bandwidth allocation in flexible OFDM-based networks. In *IEEE Optical Fiber Communication Conference*, Los Angeles, California, Mar. 2011.
145. J. Armstrong. OFDM for optical communications. *IEEE Journal of Lightwave Technology*, 27(3):189–204, Feb. 2009.
146. O. Gerstel, M. Jinno, A. Lord, and S. J. Ben Yoo. Elastic optical networking: a new dawn for the optical layer. *IEEE Communications Magazine*, 50(2):s12–s20, Feb. 2012.
147. M. Jinno, B. Kozicki, H. Takara, A. Watanabe, W. Imajuku, Y. Sone, T. Tanaka, and A. Hirano. Distance-adaptive spectrum resource allocation in spectrum-sliced elastic optical path network. *IEEE Communications Magazine*, 48(8):138–145, Aug. 2010.

148. M. Jinno, H. Takara, B. Kozicki, Y. Tsukishima, Y. Sone, and S. Matsuoka. Spectrum-efficient and scalable elastic optical path network: architecture, benefits, and enabling technologies. *IEEE Communications Magazine*, 47(11):66–73, Nov. 2009.
149. H. Beyranvand and J. A. Salehi. A quality-of-transmission aware dynamic routing and spectrum assignment scheme for future elastic optical networks. *Journal of Lightwave Technology*, 31(18):3043–3054, Sept. 2013.
150. X. Wan, N. Hua, H. Zhang, and X. Zheng. Study on dynamic routing and spectrum assignment in bitrate flexible optical networks. *Photonic Network Communications*, 24:219–227, 2012.
151. M. Klinkowski and K. Walkowiak. Routing and spectrum assignment in spectrum sliced elastic optical path network. *IEEE Communications Letters*, 15:884–886, 2011.
152. S. Yang and F. Kuipers. Impairment-aware routing in translucent spectrum-sliced elastic optical path networks. In *17th European Conference on Networks and Optical Communication (NOC)*, June 2012.
153. S. Talebi, F. Alam, I. Katib, M. Khamis, R. Salama, and G. N. Rouskas. Spectrum management techniques for elastic optical networks: a survey. *Optical Switching and Networking*, 13:34–48, Feb. 2014.
154. K. Christodoulopoulos, I. Tomkos, and E. Varvarigos. Spectrally/bitrate flexible optical network planning. In *36th European Conference and Exhibition on Optical Communication (ECOC)*, 2010.
155. K. Christodoulopoulos, I. Tomkos, and E. Varvarigos. Elastic bandwidth allocation in flexible OFDM-based optical networks. *IEEE Journal of Lightwave Technology*, 29:1354–1366, 2011.
156. Y. Wang, X. Cao, Q. Hu, and Y. Pan. Towards elastic and fine-granular bandwidth allocation in spectrum-sliced optical networks. *IEEE Optical Communications and Networking*, 4(11):906–917, Nov. 2012.
157. W. Wei, C. Wang, and X. Liu. Adaptive IP/Optical OFDM networking design. In *IEEE Optical Fiber Communication (OFC)*, San Diego, CA, Mar. 2010.
158. Y. Yoshida, T. Kodama, S. Shinada, N. Wada, and K. Kitayama. Fixed-length elastic-capacity OFDM payload packet: Concept and demonstration. In *IEEE OFC/NFOEC*, 2013.
159. W. Wei, J. Hu, D. Qian, P. N. Ji, T. Wang, X. Liu, and C. Qiao. PONIARD: A programmable optical networking infrastructure for advanced research and development of future Internet. *IEEE Journal of Lightwave Technology*, 27(3):233–242, Feb. 2009.
160. W. Wei, C. Wang, J. Yu, N. Cvijetic, and T. Wang. Optical orthogonal frequency division multiple access networking for the future Internet. *IEEE Journal of Optical Communications and Networking*, 1:236–246, 2009.
161. W. Wei, J. Hu, C. Wang, T. Wang, and C. Qiao. A programmable router interface supporting link virtualization with adaptive Optical OFDMA transmission. In *IEEE OFC/NFOEC*, 2009.
162. W. Wei, C. Wang, and J. Yu. Cognitive optical networks: Key drivers, enabling techniques, and adaptive bandwidth services. *IEEE Communications Magazine*, 50:106–113, 2012.