

Computer Networks
Prof. Sujoy Ghosh
Department of Computer Science and Engineering
Indian Institute of Technology – Kharagpur

Lecture 12
Protection and Restoration

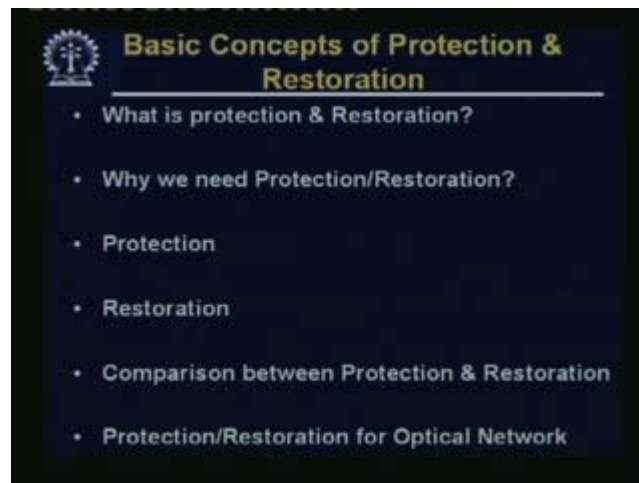
Good day. In this lecture we are going to discuss the various protection and restoration mechanisms which are usually employed in optical networks.

(Refer Slide Time: 00:57)



We have to discuss what is protection and what is restoration and why we need them.

(Refer Slide Time: 01:08)



What is protection, what is restoration, comparison between the two and the different schemes of protection would be our general outline of presentation.

(Refer Slide Time: 01:19)



Network is unreliable; and then so many failures can occur – a node may fail, a link may fail or a link may get cut, some fiber optic line may cut in-between because somebody cut it while digging a hole or something. A node might fail, there may be power failures at nodes and so there are failures in the network. But if you remember, one of the most important places where we deploy optical networks is in the core of the network. The core of the network connects so many people to so many other people and it is very vital that we cannot allow the services to be

severely disrupted because that would have very grave consequences and anyway the service provider attempts to give a high level of service. Although failures are unavoidable in real life, we have to find some way of combating this; that means if there is a failure we want to recover from it as soon or as fast as possible. That is what protection and restoration is all about. Another thing is that when we want to give some reliability, protection, restoration, etc., in some form or the other we always have to bring in some redundancy. Without any redundancy, a system cannot be protected, or it cannot be restored without any replacement, etc. There has to be some redundant capacity in some form in the network in order to achieve protection and restoration. So how this is done is what we will discuss.

(Refer Slide Time: 03:20)



Protection and restoration are the mechanism to recover from network failure: their difference will be discussed in the following parts.

(Refer Slide Time: 03:26)



Why we need protection and restoration is now clear: to recover from network failure, to prevent the lot of data loss. Now another point is what would we mean by prevent lot of data loss? The point is that we will be talking about protection at the optical level, at the lower, transport level, that is, the low physical level. Above this physical level, there are a whole lot of other layers like data link layer, network layer, transport layer, etc. They may have their own protection mechanism and they might be able to tolerate. In some cases, if such a protection is sort of implemented at higher level, such protection usually would tolerate a small amount of data loss, which they will retransmit or do something in the protocol to take care of that. We will be talking about this later on. Within that limit we do not have a 100% tight case, but we have something to play with. Within that limit, if the physical layer can come back up, then that is nice; then the end user who is sitting at the top of the application layer would not notice that a failure has actually occurred. So this is our general goal and, of course, to prevent lot of data loss.

(Refer Slide Time: 05:13)



To provide reliable communication service is a reason for having protection and restoration.

(Refer Slide Time: 05:19)



Protection is the primary mechanism; this is fast and routes are usually preplanned – we will come to this part later on – whereas restoration is the secondary mechanism, used to provide more efficient routes or additional resilience, etc., over and above protection. We will see both of these later on.

(Refer Slide Time: 05:40)



Techniques for protection: we could protect a path, which is called path protection; we could protect a link; that is called link protection. There are various schemes of protection: like 1+1 protection, 1:1 protection, 1:N protection, M:N protection. This depends on what kind of redundancy we have built into the network. So we will look at some of these techniques one by one.

(Refer Slide Time: 06:12)



So what are the considerations and tradeoffs that we have for protection? For support for fast protection, time is dictated by the client layer. This is what I was talking about earlier – in the client layer, whatever the higher layer we are talking about, all of them are clubbed together, and

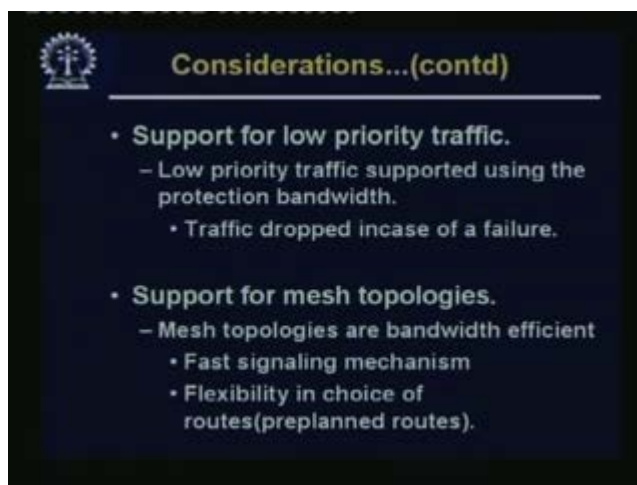
we are calling them client layer, let us say. In the client layer we may have some resilience; that means, we can tolerate some small amount of data loss. So depending on that, and beyond that, it will be taken that the service will fail. This is dictated by the client layer and that is the constraint within which we try to do our protection or restoration at the bottom layer, that is, the physical layer. We require some switching technologies as we will see, and how to implement protection.

(Refer slide time: 07:13)



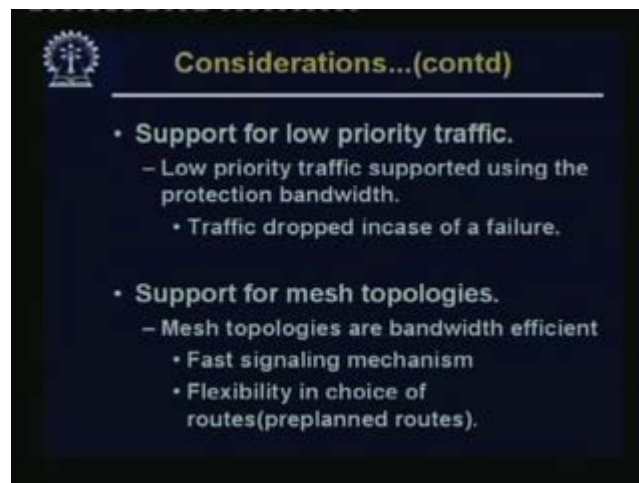
It could be through dedicated hardware or through software. Here we will only be talking about – since we are talking about the physical layer – hardware protection, where some hardware has been put in order to take care of this protection.

(Refer slide time: 07:34)



Support for low priority traffic – that is another consideration we might have. So low priority traffic supported using the protection bandwidth; traffic dropped in case of a failure. As I mentioned earlier that in order to give some protection capability into the network we have to build in some kind of redundancy over there. Now when there is no failure, naturally the redundant link would be idle. What you could do is that if you have some low priority traffic, which could be dropped whenever there is a problem, under normal circumstances you can use your provision in the network to carry the low priority traffic, and as soon as there is a failure, where naturally the provision is going to be used up for providing protection, restoration, etc., the low priority traffic would be discontinued. So that is one thing we could do.

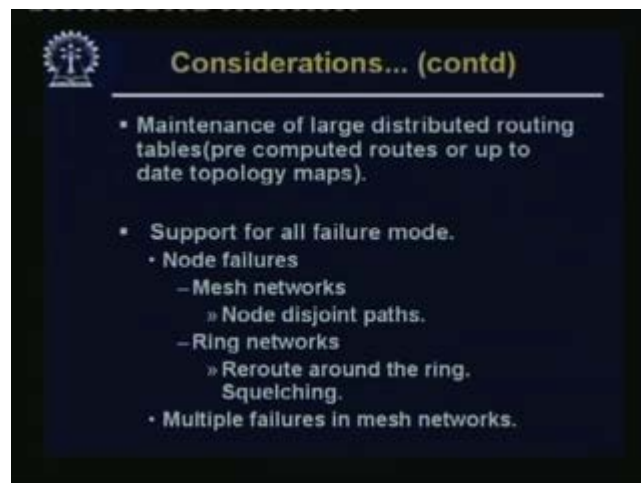
(Refer slide time: 08:46)



Support for mesh topologies – mesh topologies are bandwidth efficient, fast signaling mechanism, flexibility in choice of routes by preplanned routes, etc. If we have support for mesh topologies, it is also nice. As a matter of fact, if you remember our discussion about different types of topologies – we have star, tree, ring, mesh etc., so these are the different possible topologies. Out of all these, point-to-point connection, ring and mesh are the three topologies, which dominate most of the WAN. When you want to communicate between two points, you first set a point-to-point link. Now consider that link to be quite important. Then what you would like to do is that you would like to put in some kind of reliability over there, by making it into a ring. For example, we discussed SONET. A SONET gear quite often is put in the form of ring, rings that would touch, mutually touching rings. For example, the telecom people put up all these

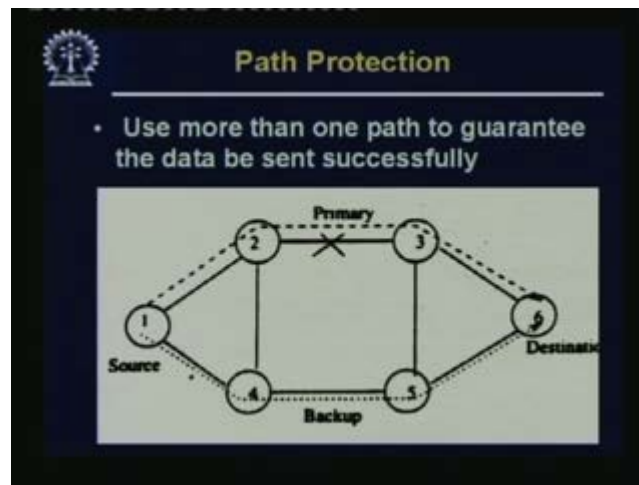
SONET rings through fiber optic networks. The rings are quite easy to handle and we will discuss the protection mechanism in ring quite a lot. But if you go to a wider geographical area ring, it may not be feasible due to various reasons. What you have is a mesh. A mesh, if you recall, is just a graph, where the nodes are connected in some fashion. It has to be a connected graph and actually later on we will see that for giving proper protection, it has to be not only a connected graph, it has to be biconnected graph as well. That means between the two nodes, which are communicating, there have to be two alternative paths, which are linked somehow. Otherwise, you will not be able to give protection. Giving protection in mesh networks is also an important consideration.

(Refer slide time: 11:17)



Other important considerations are maintenance of large distributed routing tables, that is, precomputed routes or up to date topology maps. So you have to maintain this dynamically, because a network is a very dynamic thing as links come up or go down or nodes come up or go down, this may have to be recomputed and stored in a distributed fashion. We would like to have support for all failure modes: node failure for mesh networks and for ring networks, so or it may be of course ring failures.

(Refer slide time: 11:55)



First, let us talk about path protection. It uses more than one path to guarantee that data be sent successfully. If you look at this graph, it will have a 6-node graph where 1 and 6 are communicating and on the top, through the dashed line, we show the primary path, which is the primary connection between 1 and 6. Now what might happen is that the link between 2 and 3 might snap due to some reason. We have already got a backup route, which is calculated going through 1, 4, 5 and 6. So we will channel our communication through the backup link. Please note that this backup or secondary path from the source to the destination does not share any link with the primary path. That is the requirement: if any link in the primary path fails, assuming there is only one failure, which means that the backup link is all intact and you can switch to the backup path for this particular communication.

(Refer slide time: 13:06)



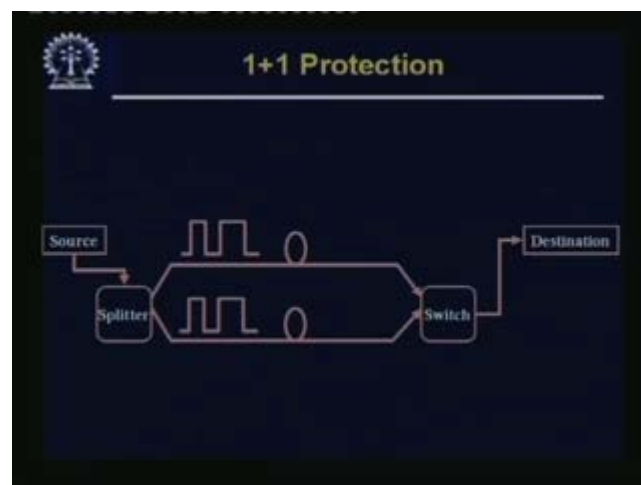
Now path protection: there are various ways to protect a path, that is, various ways of provisioning the extra bandwidth capacity in the network so that you can give protection. You need to build in some kind of redundancy in the network. So there are various ways you can build in the redundancy and the schemes may be divided as dedicated path protection, or shared path protection.

(Refer slide time: 13:40-13:50)



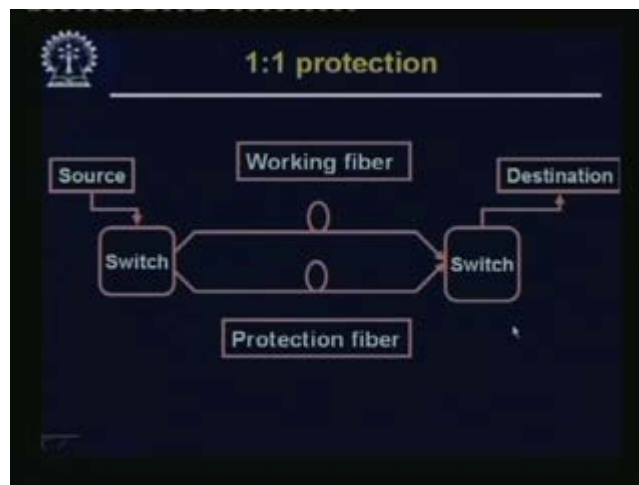
Dedicated path protection may be shown as 1+1 protection and shared path is 1:1 protection or 1:N protection. So we will look at this one by one.

(Refer slide time: 13:51)



This is an example of 1+1 protection. So you have the source on the left, and then, you have destination on the right. So it is communicating. What you can see is that from the source, the signal is coming to the splitter. If you remember, the splitter is going to split it into two halves, may be of equal power or something, and then the same signal is been carried through two different links and over there, we a switch. The switch determines which signal is better and may be switches to that and a communication is going on. If that particular link, say, this top one fails, it automatically switches to the other protection link. So this is some kind of hot redundancy we have. That means there is redundant source of information in the destination side; so if the primary one fails, the secondary one is already on. So protection would be very fast. The only trouble is that for each such path, you will have to give an alternate path, which is also being used at the same time. This is a dedicate path protection – for this particular path there is a dedicated alternative path; although this is very good, it is costly.

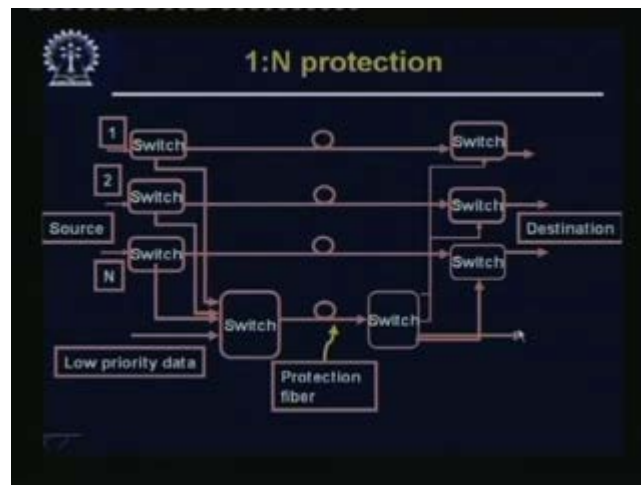
(Refer slide time: 15:32)



Now we come to shared mode protection; in the shared mode, the figure looks almost the same, except if you note, this splitter on the left has been replaced by a switch. So what we have is we have a working fiber and we have a protection fiber. Unlike the 1+1 protection, this is not a hot standby. This is cold in the sense that the protection fiber, to start with, is not carrying any signal, let us say. The source is just simply passing through the switch; the signal from the source is passing through the switch, then down the path through the destination switch to the final destination. If the working fiber goes down, then this switch will flip and the protection fiber

would be in place. Since this is not a hot standby, meaning that since it is not carrying any signal under normal circumstances, you could share this path with something else. You can use this to send some other channel or some other information, etc. So that is why this is a shared mode of protection.

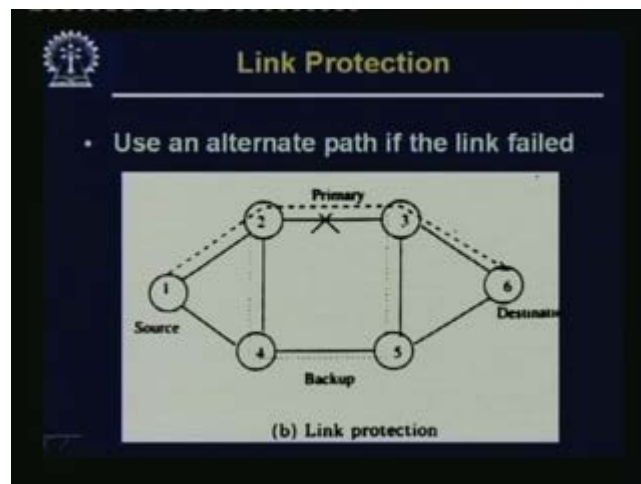
(Refer slide time: 16:58)



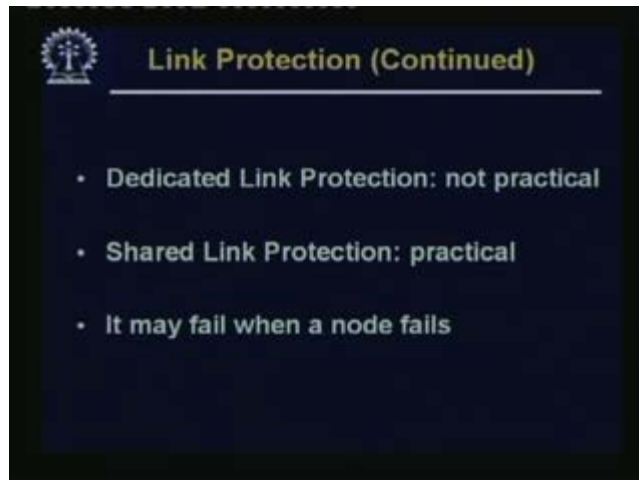
Generalizing this, we get this 1:N protection, where N line is sharing 1 protection line. We have the inputs from 1 to N lines, so these are the N sources, and let us say so many destinations. So they may be going – 1- to 2- and N to N- and so on; this is the normal mode of operation. This part of the network is for normal mode of operation and each of them is connected to a switch over here. Another link from the switch comes to one bigger switch, so this is an N input port over here and then there is a single link from this switch to this switch, on the destination side, which again feeds to all the switches. What would happen is that in case there is a failure in any one, let us say, ith link from 1,2 to N, out of that, the ith source to ith destination, which is going through the ith link and the ith link fails, what this ith switch could do is that the ith switch could switch the signal from the ith source to this particular switch at the bottom and now the protection fiber would be carrying the signal that was flowing down ith channel over there and then again it will feed it to the ith switch on that side. Of course these two switches have to communicate that the ith one has failed over here, so you switch it to your ith mode, or this switch may sense it that this line has gone down, so it will take signal from this line. This one protection fiber is been shared by these N working fibers, and as I mentioned earlier, when

everything is fine this protection fiber could be carrying some low priority data. When everything is fine these N working fibers are actually carrying the most of the important traffic so some low priority data could be flowing down the protection fiber. As soon as there is any failure anywhere, the low priority data would be stopped or it will be dropped and then this protection fiber would switch to give the services between the nodes that have experienced a link failure. Just as we have a path protection that means for a path you try to give an alternate path, similarly you could do it at the link level also. In general, it may be more efficient or more proper to do that because if you have N nodes and if it is slightly large, potentially you have very large number of paths through the networks and for each path having an alternate path may not always be a very good idea. So we concentrate at the link level and for each link level we may give a protection.

(Refer slide time: 20:21)



So use an alternate path if the link has failed. This is the primary where the link may have failed so I have an alternate path to that from 2 to 3, 2 to 4, 4 to 5, 5 to 3. This alternative, please note, is for the link from 2 to 3. We will have another diagram for that. (Refer slide time: 20:44-21:09)



Dedicated link protection is not always practical although sometimes we may have it; shared link protection is practical and it is quite often implemented. This link protection may fail because here you are only provisioning for the failure of a link, but if a node fails, then it may lead to some complication as we will see.

(Refer slide time: 21:10)



To compare between path switching and line switching, path switching is a coarser scheme and line switching is finer scheme; and line switching can again be a span protection. Span may be several links together; that may be span or a line protection.

(Refer slide time: 21:29)



In mesh networks of course restoration is possible only if the graph is two edge connected i.e., biconnected, which means that there are two edge disjoint paths between any pair of nodes so that no single edge failure can disconnect the network. This is a necessity and we usually try to keep it that way unless its very difficult or very cost it's not cost effective, etc.

(Refer slide time: 21:58)



Protection in a mesh network is more complicated then a ring. Simple minded scheme would be two edge or node disjoint paths for each connection, 1+1. As is mentioned, it is not very efficient. There may be many paths and provisioning double the number of paths, which are pair wise mutually node or edge disjointed may be very difficult. That may be a lot of extra

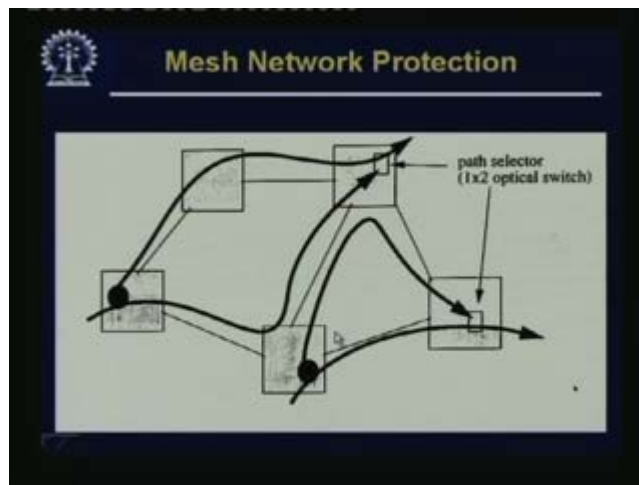
provisioning in the network. A better approach would be line protection, which of course has the problem of coordination. I will show that and protection cycles in mesh net later on.

(Refer slide time: 22:43)



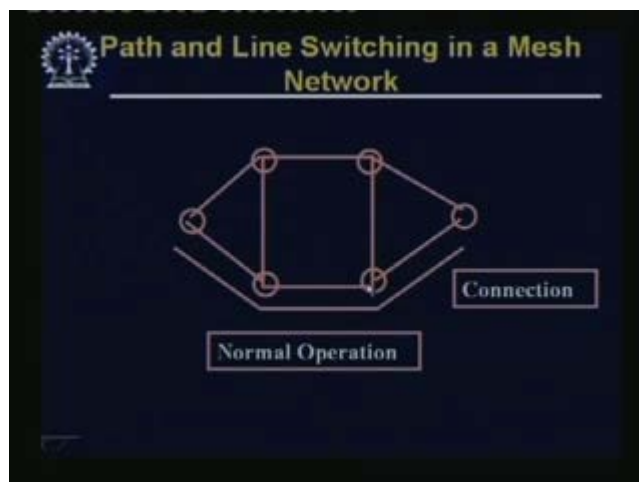
In the path layer and mesh protection, there is protection of mesh networks to protect the mesh at the single unit. Pre-computed routes means all possible routes and alternative routes are pre-computed. 1+1 protection is protection route per light path, protection route per failure. We will discuss this later but, as I said, this is a costly alternative. Or what we might do is that we can do on-the-fly route computation; that means it is not pre-computed. As there is a failure, centralized route computation and coordination route computation and coordination are done at end nodes or distributed route computations – all these are possibilities.

(Refer slide time: 23:25)



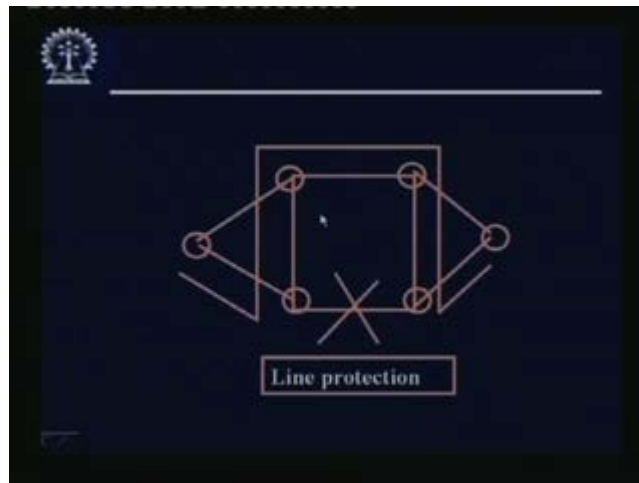
This is an example of mesh network, where let us say this is the primary path and this is an alternate path. This alternate path may be pre-computed or it may be computed on the fly when there is a failure. Similarly, from here to here this may be the primary path and this is an alternate path. Please note that this communication as well as this communication are going through the same fiber. Maybe they are going through different wavelengths or maybe they are combined together; so there are various way of handling this.

(Refer slide time: 24:04)



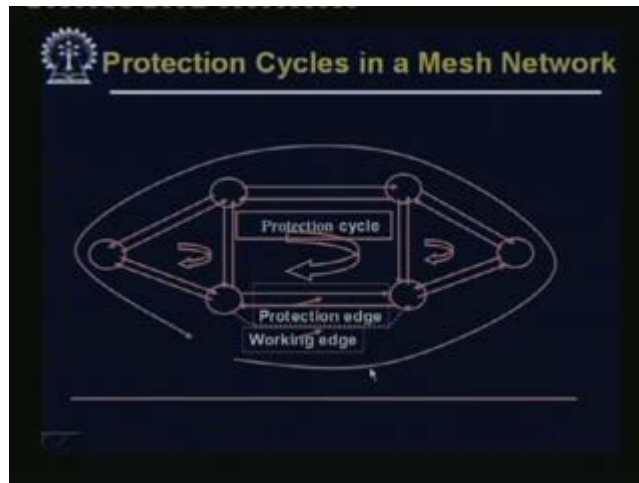
Let us look at some diagrams: this is a mesh network. Naturally once again, it has the same 6-node network and this is the normal operation, that is, communication from this node, node 1 to node 6. Now there is link failure over here; what you might do is you might switch the entire path like this. So this was the pre-computed path for this path. You switch to that – this is one possibility. Or this particular span, the span could be live from here to here or from here to here. That may have an alternative path, to which you switch.

(Refer slide time: 24:52)



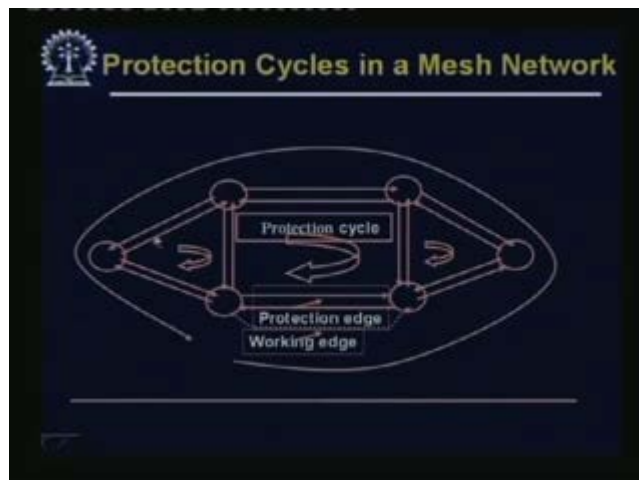
Or this particular line may have this alternative; that means from this node to this node. If this is the link, the alternative to this link is this section from here to here to here to here. It is trying to go through the same path, but then it takes a diversion when there is a link failure through the alternative, which is provided for this particular link. I have shown only one node coming into this but you can appreciate that this path may be very long so you do not have to re-compute the path over the entire length; rather locally, you can re-compute that for this particular link as an alternative. So the alternative may well be through some local nodes only. So this path just gets a slight diversion; that is line protection.

(Refer slide time: 25:51)



We talked about protection cycles in a mesh network, now for protection cycles, what we do is that for each of the paths we try to form a cycle – cycle of provisioning of light paths, let us say. The point in the cycle is that suppose we are going through some arc of the cycle and if some link in-between breaks, you can always go in the other direction. So you try to form these cycles in the mesh which you keep ready and pre-computed. So whenever there is a failure you can switch.

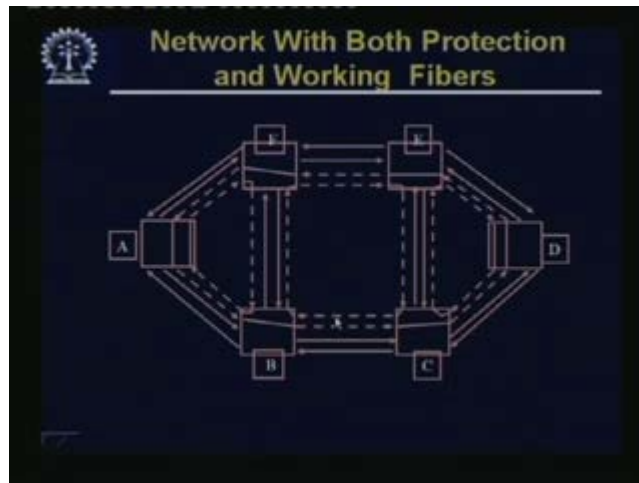
(Refer slide time: 26:38)



You can see here there are pairs of a fibers going in both the directions and there you can form cycles over there. One of them would be protection edge, maybe the inner one; and one of them

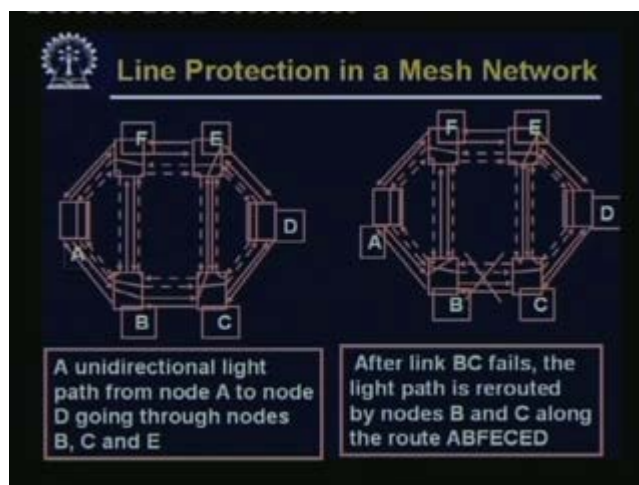
would be the working edge. If one of them fails, the other will automatically take over; the other direction will automatically take over.

(Refer slide time: 27:04)



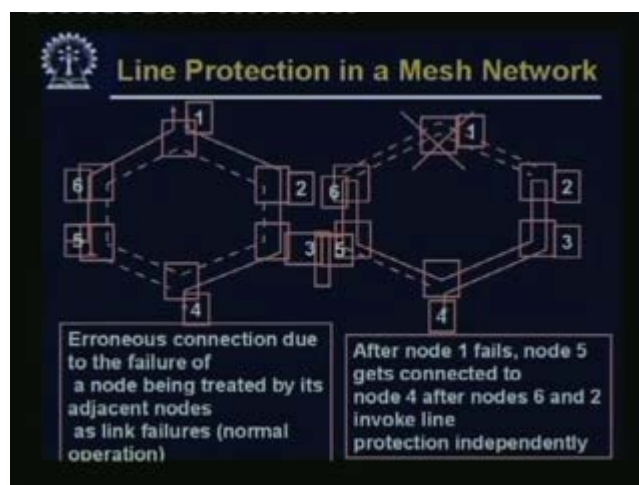
So this is another example of a network with both working and protection fibers. The working fibers have been shown in solid lines whereas protection fibers have been shown with dashed lines. Once again you must realize that you may not provide the same level of protection to all paths or to all parts of the network; depending on which part you consider more sensitive or more important, you may put your protection there. A mesh network may be partially protected. Some of the parts may be protected or some of the parts may not be protected.

(Refer slide time: 27:51)



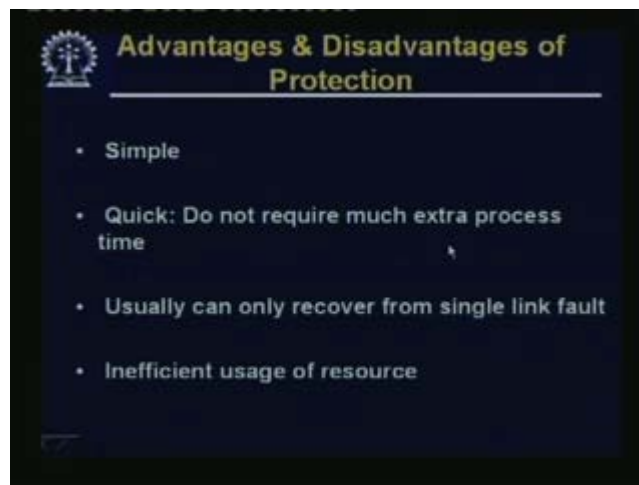
So some more cases: line protection in a mesh network. What we have is a unidirectional light path from node A to node D; so from node A to node D we have an unidirectional light path going may be like this through nodes B, C, and E. We are talking about this path A to B to C and then to E. Now after the link BC fails, the light path is rerouted by nodes B and C along the route A B F E C E D. The unidirectional light path was going from A B C E and then to D; so this was the path A B C E D. Now BC has failed; so A B F E C E D. You can see what has happened is that there was no point in coming from E to C and then back from C to E. What has happened was that we were doing line protection; that means, for BC my protection was from B to F to E to C because I wanted to go from B to C so I am going from B to F to E to C and so now BC has failed. So that is what I do and from that point I continue wherever I was going and actually I was going from C to E to D, so once again from C, I go back to E and then to D. Such a thing is possible because we are taking a local decision; that means, for this particular link what to do in the case of failure that has already been pre- computed? We are not taking a global picture; here I have shown you a very small graph so you can immediately see the entire global picture with your own eyes. But for local nodes we may not have the global pictures where this node is coming from and where this path is finally going to. So all the intermediate nodes may not have the global pictures because if all of them had the global pictures maybe they could have computed a better path but this is for line protection.

(Refer slide time: 30:41)



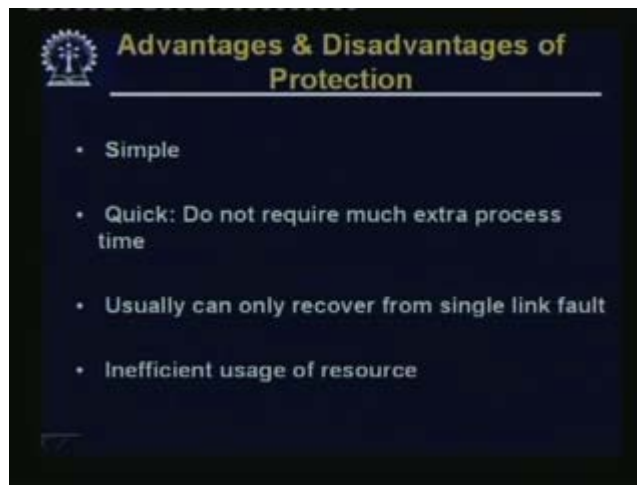
Now line protection in mesh network: here what might happen is that erroneous connection due to the failure of a node is being treated by its adjacent nodes as link failure. This is one case of the so-called race condition. What might happen is that node 1 has failed; so what would 6 do is that it may assume that the link from 6 to 1 has failed, whereas 2 may assume that the link from 2 to 1 has failed. So both of them perceive as two different failures, so they take some local decisions, and it might lead to a funny situation where you are going in a cycle. Maybe other kinds of things may happen after node 1 fails – node 5 gets connected to node 4 after node 6 and 2 invoke line protection independently. If they perceive the same failure, the actual failure was that of node 1; but if they perceived differently as failure of link 6 1 and failure of link 2 1 and they have independent actions, it may lead to race conditions.

(Refer slide time: 31:58)



The advantages and disadvantages of protection: we will also be talking a little bit of restoration. Protection is simple; it's quick; does not require much extra process time – and this is the important part, since this is quick. As I mentioned earlier, there will not be a lot of data loss. For example, a SONET ring would sort of come back up from a failure in a less than 15 milliseconds. So that is a benchmark; you are back in action in less than 15 milliseconds, so whatever little you might have lost during that time would be taken care of by the higher layer protocols.

(Refer slide time: 32:51)



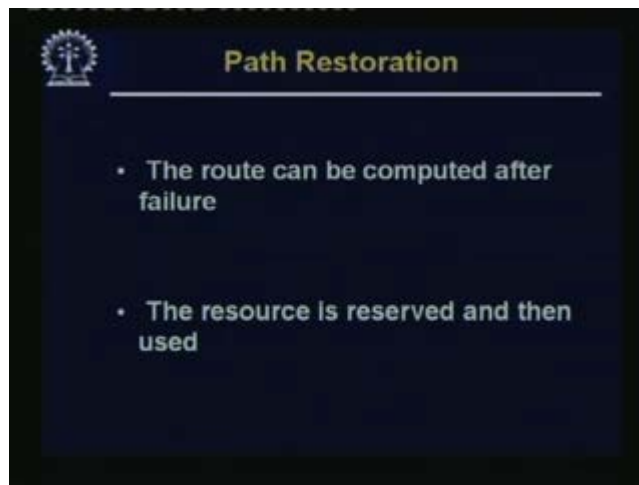
But usually they can only recover from single link faults. If there are multiple link faults, all kinds of funny things may happen. There is inefficient usage of resource, because protection needs a lot of resources, even if we are sharing them.

(Refer slide time: 33:08)



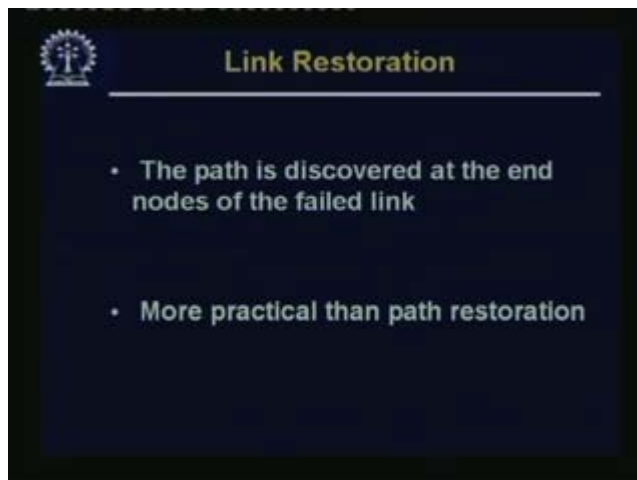
Dedicated protection needs even more resources; we talk about path restoration and link restoration.

(Refer slide time: 33:20)



What we do is that we compute the path after the failure and the resource is reserved and then used. So in restoration what you try to do that is that you try to look at the current actual situation. You try to have some protocol for keeping track of the present situation and then you compute some root and then you reserve the path and then restore the original service. This has got some software parts, so some data has to be processed, etc. This usually takes more time, so this has a hardware as well as software aspect to do it.

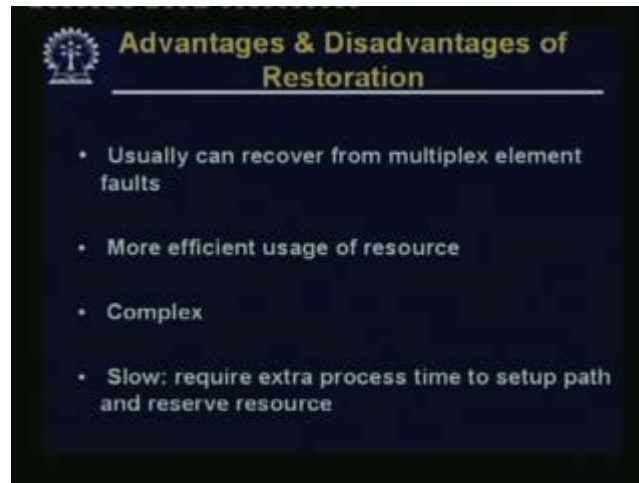
(Refer slide time: 34:00)



The path is discovered at the end nodes of the failed link; but this is more practical than path restoration. We have both path restoration and link restoration. Path restoration means the path

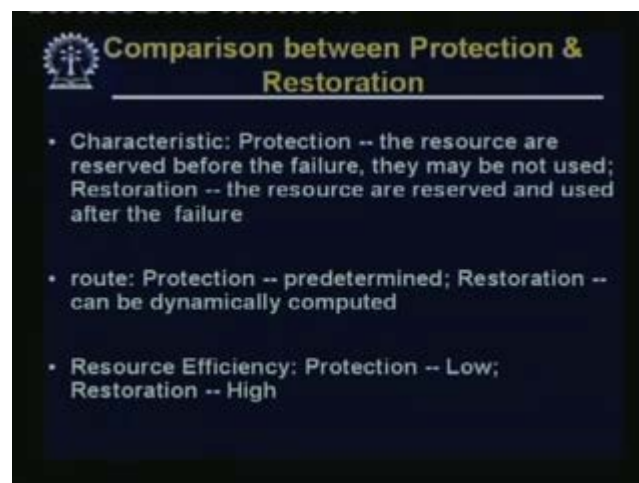
may be long and to find out an alternative, it may be more difficult, whereas links are between the adjacent nodes so they may quickly find an alternative.

(Refer slide time: 34:26)



Advantages and disadvantages of restoration are the following: usually it can recover from multiplex element faults because you are sort of having some protocol to exchange information and then find the current situation and form the alternatives, etc. There is more efficient usage of resources; it is more complex; it is slower; it requires extra process time to set up path or reserve resources.

(Refer slide time: 34:54)



So for comparison between protection and restoration, in protection the resources are reserved before the failure; they may not be used. In restoration the resources are reserved and used after the failure. So this is the main difference between the two. Route: in protection it is predetermined; in restoration it can be dynamically computed. Resource efficiency: in protection it is naturally low and the restoration is comparatively high.

(Refer slide time: 35:20)



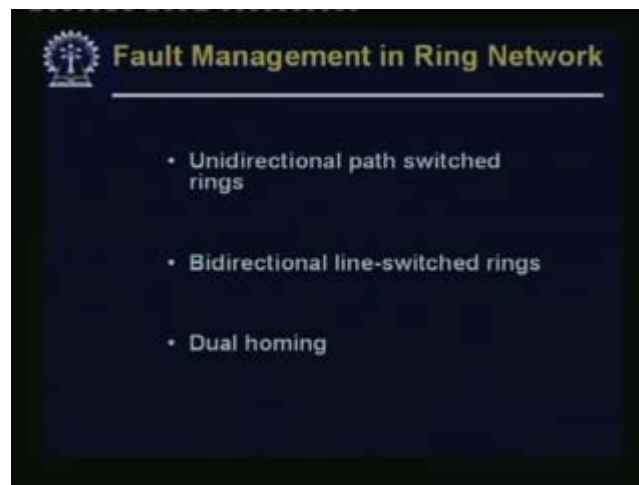
Time used for protection is short; for restoration it is longer. Reliability: protection is mainly for single fault, whereas restoration can survive under multiple faults. Well, it is not that it can always survive under multiple faults; it depends on where the faults are, but if it is survivable after multiple faults, it will take a global view and say that ok, still I can give the services. So restoration will take care of that, whereas in protection it may not be possible to handle multiple faults. Implementation: protection is simple and restoration is naturally more complex.

(Refer slide time: 35:56)



For optical networks, we not only talk about physical links we talk about virtual wavelength paths. So in light paths also routing can be centrally controlled or distributed; resource reservation, forward reservation, as well as backward reservation are done as we do in optical networks. We will talk more about this later.

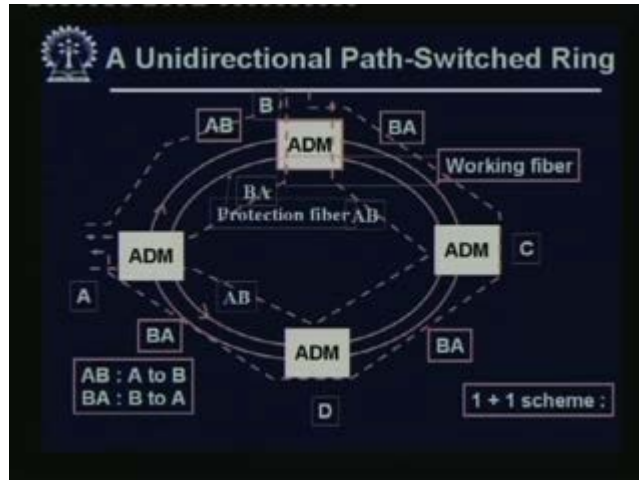
(Refer slide time: 36:22)



Now let us come to this all-important topic of fault management in ring networks. As I said, ring networks are very ubiquitous in WAN. All the telecom people will love these things because naturally it allows them to give very high level of service. So we have so many rings; these SONET rings are very common, and rings are a very common kind of topology. One of the chief

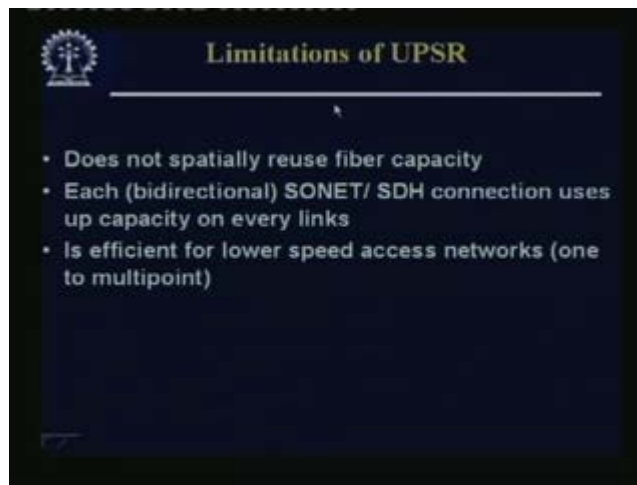
attractions of the ring topology is its capability to allow some kind of protection and restoration with some redundancy in-built, as we will see. We will look at two different cases, unidirectional path switch rings and bidirectional line switch rings or UPSR and BLSR.

(Refer slide time: 37:16)



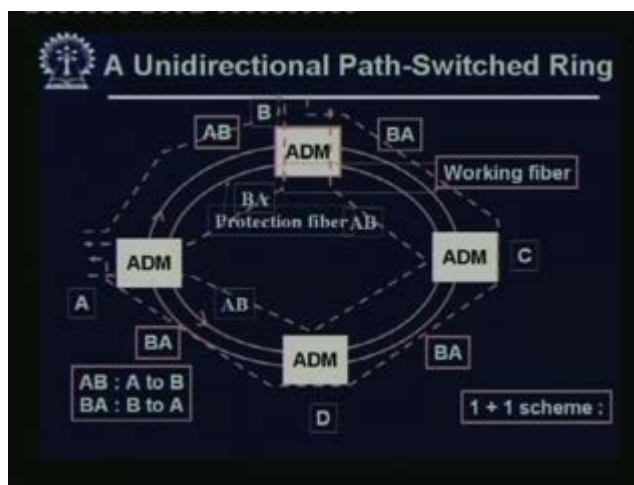
This is a diagram of UPSR. We have AB, means A to B and BA, means B to A. First of all you see that we have these two rings – we have this working fiber as well as a protection fiber. Working fiber is going in only one direction and protection fiber is going in the other direction, in the counter clockwise direction. So for connection from A to B and from B to A – this is A, and this is B – my working fiber is going in this direction. A to B is really going in this direction; B to A is coming in this direction. So A to B, B to A: this is the protection fiber part, A to B the working fiber is going in the other direction. A to B is going like this and B to A is coming all the way like this. Let us say the outer one is a working fiber, it is through the outer ring. For A to B, this is the path; and B to A, this is the path, whereas for protection purpose – and please note that this is a 1+1 scheme; that means for this the alternative is already provisioned and maybe it is something like a hot standby. So A to B is through the protection fiber and now this path from A to B is going via D and C, so A D C B: that is A to B and B to A. The protection path is from this through the protection fiber. So if there is a failure anywhere, this can still continue, so this is a unidirectional path switched ring.

(Refer slide time: 39:23)



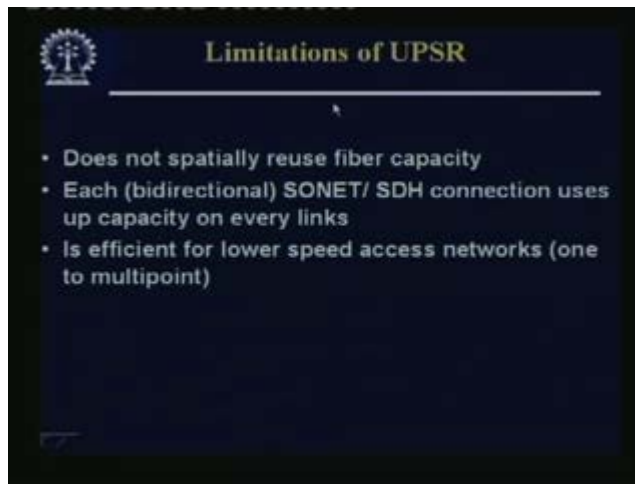
There are some limitations of UPSR. It does not spatially reuse the fiber capacity; so what is happening is that since this is unidirectional, what is happening is that even if two nodes are side by side, if they are in wrong direction then it has to come all the way through and naturally all those links – even if we are talking about particular λ s I mean WDM systems – in the ring are entirely sort of covered by this. Otherwise even if we are talking about a WDM system what is happening is that at least one λ around the whole ring is getting occupied because two side by side things want to communicate. So it does not spatially reuse fiber capacity and if there is some, you could use it for some other purpose.

(Refer slide time: 40:30)



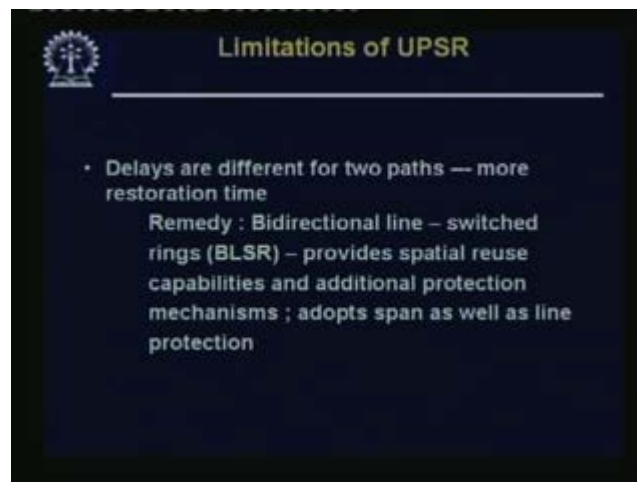
That is not possible in a unidirectional ring. Each bidirectional SONET SDH connection uses up capacity on every link. That is the thing; if you look at the previous picture we are just talking about a connection from A to B and B to A; that's all, so this entire ring now has been used up assuming only one wavelength. The entire ring now has been used up; these ADMs of course are the ADMs of SONET.

(Refer slide time: 40:45)



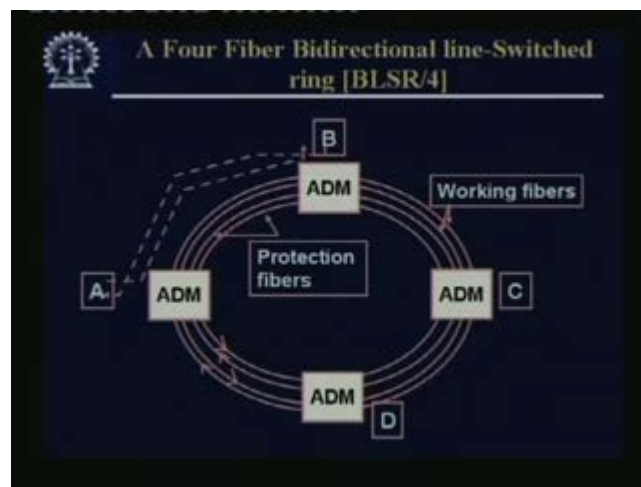
So it is efficient for lower-speed access networks, one to multipoint only.

(Refer slide time: 40:51)



And the other point, which may become a problem, is that delays are different for the two paths because one of them is small and the other one is quite large, all the way around the link. The delays in the two paths are quite different. A remedy could be bidirectional lines; that means, bidirectional line switched rings or BLSR. So we were in UPSR, now we will be talking about BLSR, bidirectional line switch rings. This provides spatial reuse capabilities and additional protection mechanisms and adopts span as well as line protection. We will have a look at both of these.

(Refer slide time: 41:28)



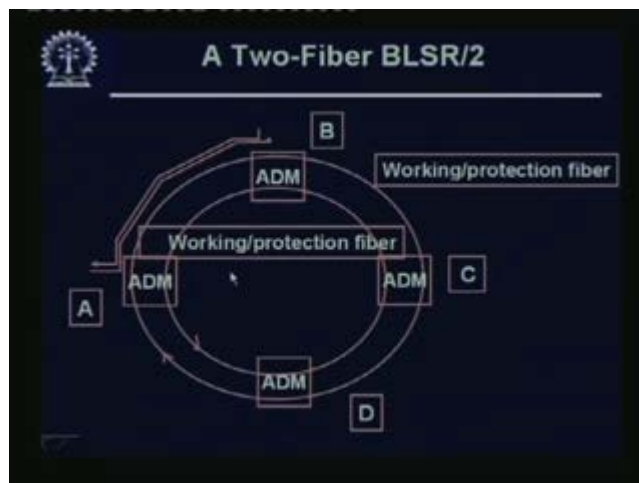
This is a four-fiber bidirectional line switched ring; so once again we have two working fibers, say, the outer two ones and two protection fibers or the inner ones. So we actually have four fibers. This is a four-fiber system and it is going, so we are using two such working fibers. One of them is going in one direction the other is going in the other direction, so this is bidirectional. We can communicate both ways, so if A to B and B to A both can communicate on the working fiber in this part at all, the rest of the space can be used for other communication. So spatial reuse is much better for bidirectional line switched rings, and the protection fibers are there because if one of them fails then the protection fiber may take up.

(Refer slide time: 42:27)



This is a span protection so we are talking about bidirectional traffic supports, maybe 16 nodes or some distances, etc., from A to B. You may pre-compute a span that if this span goes down what is the alternative span. If this is the case of line protection, if this line fails what is the alternative line? Now the alternative line of course may go in the other direction, depending on what the failure is. So there are many different possibilities with this BLSR.

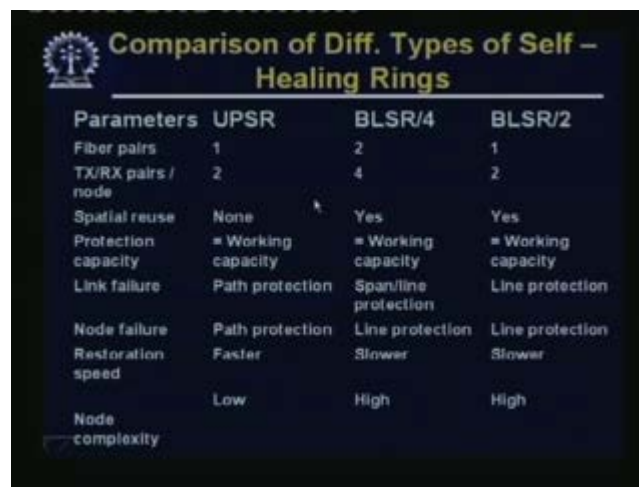
(Refer slide time: 43:11)



There is also two-fiber version of BLSR, namely BLSR/2. Here both the fibers are working as well as protection fibers. That means if one of them fails the other one will give the protection in the other direction of the ring. If you note, the rings are going in the opposite directions and there

is a line failure in both of these working as well as protection fiber. That means A to B is communicating in this direction, B to A is communicating in this direction. But if one of them fails, the other one will give the protection by going in the other direction. In that case we have to go all the way round, then you will have to reserve the resource and if there is something already going on here, you may be blocked or this may have to be dropped and so on. So naturally you have two fibers, we have less flexibility. But with two fibers we can get this bidirectional thing going on.

(Refer slide time: 44:25)

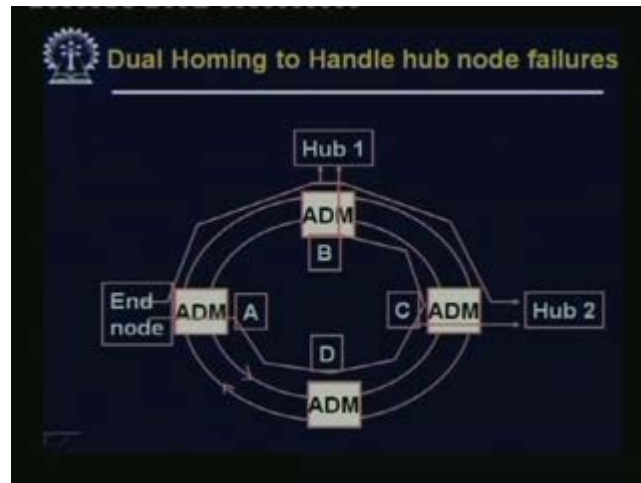


Parameters	UPSR	BLSR/4	BLSR/2
Fiber pairs	1	2	1
TX/RX pairs / node	2	4	2
Spatial reuse	None	Yes	Yes
Protection capacity	= Working capacity	= Working capacity	= Working capacity
Link failure	Path protection	Span/line protection	Line protection
Node failure	Path protection	Line protection	Line protection
Restoration speed	Faster	Slower	Slower
Node complexity	Low	High	High

Comparison of different types: these are all self healing rings. What we mean by self healing is that the nodes, these ADMs and the SONET, etc., are programmed in such a fashion that as soon as they sense a failure they know what local action to take and how to adjust the switch internally, so that it automatically switches from the working fiber or from the working span or working fiber or whatever through the protection side. That is why these are self healing rings, so they heal automatically. And as I said these SONET rings really do this in less than 15 milliseconds; the entire thing will again be up. So we can compare the three: UPSR, BLSR/4, and BLSR /2. For example, fiber pairs 1, 2 and transmission receiver pair 2 4 and 2 spatial reuse in UPSR it is none, in BLSR it is there, and in BLSR/4 it is there. Production capacity is equal to the working capacity. Link failure is path protection in the case of UPSR or span or line protection in the case of BLSR/4. In case of BLSR/2, it is only line protection because it will have to go in the other direction. Node failure: it is path protection, line protection and line

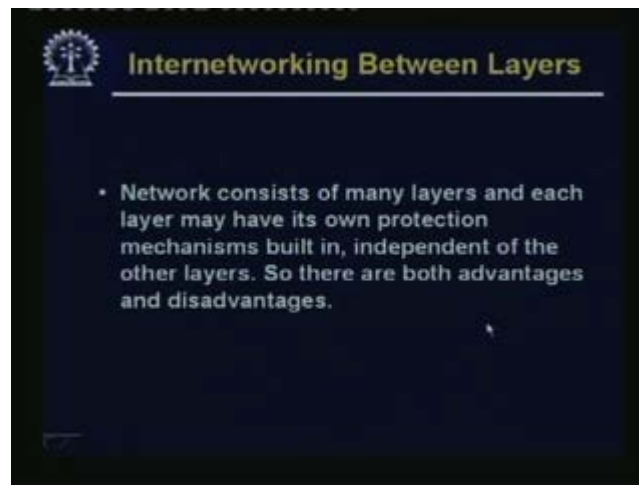
protection restoration is faster in UPSR; somewhat slower in BLSRs and restoration speed is this. Node complexity is low, high and high. Another thing we talked about is the dual homing.

(Refer slide time: 46:21)



By dual homing we mean that suppose we want to deploy your network in such a way that it is very mission critical and no failure is acceptable and we want a hot standby. So what you might to do is that you may get connected through two different hubs and what you want to do is that you want dual home to these two hubs and these two hubs take independent paths to your destination. This shows a dual homing to handle hub on node failure; so we have these four ADMs: once again four nodes, let's say A B C D. So the end node is A and these are the two hubs, B and C. What we want is that not only some link failure but even if one of the hubs fails, I should still be able to communicate. So what you do is on this ring you communicate with hub 1, let us say you are communicating with hub 1 through this A D C B and you are communicating with hub 2 through A B C. As we will note, even if one of the hubs fails, you can still communicate through other one. So this is another kind of protection.

(Refer slide time: 47:51-48:14)



And finally I have mentioned this point before – just a reminder that a network consists of many layers and each layer may have its own protection mechanism built in, independent of other layers. So there are both advantages and disadvantages to this. We have already talked about the advantages.

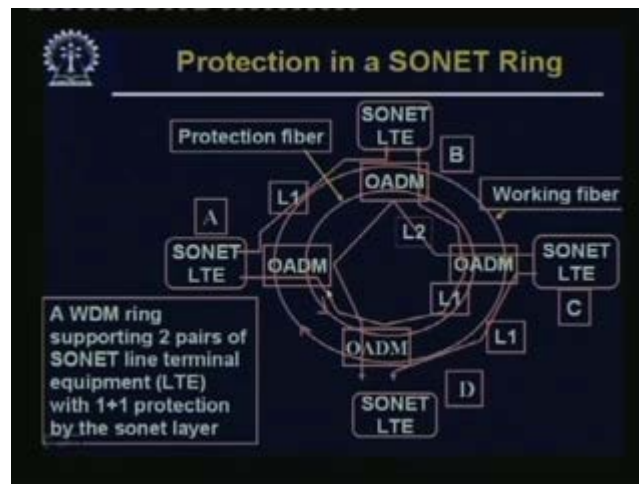
(Refer slide time: 48:15)



This is an example of a WDM link carrying SONET traffic. So there is a WDM link, so there is a SONET ADM. This is a working fiber pair and protection fiber pair. Please note that the pair has been shown as one line over here because usually as you know that fiber optic line is a simplex line; that means it goes in only one direction; there is a source in one side and the detector on the

other side. Usually these fibers always come in pairs. The other side is for the communication in other direction so we have a working fiber pair over there and protection fiber pair through this WDM link. Look at this 1:2 protected scheme. What is happening is that there is one protection fiber pair, there may be protection fiber pair and through the working fiber there may be multiple virtual links going through the working fiber using different wavelengths. A is a normal operation and B link is cut and the traffic is restored by the optical layer. That means you automatically assign new wavelengths and new paths through the fibers etc., to bring it up. So this you may do at the optical layer rather than at the electronic layer.

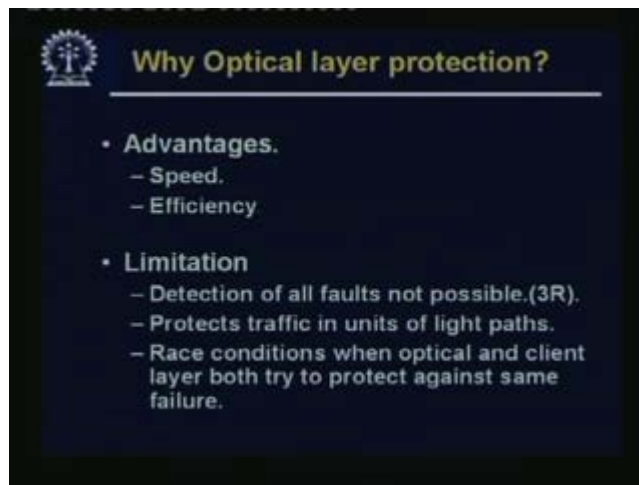
(Refer slide time: 49:52)



But the case we are talking about here is that the SONET is riding on the optical layer. So we have the optical layer at the bottom then you have a SONET on top of it. On top of SONET also there will be the data link layer and then so on all the other links. What we are saying is that each layer may have its own protection mechanism. So for example I mentioned that the SONET will also have a protection mechanism of its own. So if a SONET LTE, i.e., the SONET line terminal equipment, senses that it cannot communicate to the next LTE, this will automatically reroute the traffic and try to reroute the traffic in the other direction. Sometimes it is good to have multiple protections at multiple layers but what might also happen is that they might sort of cancel each other or they might go into a race condition. So these are the disadvantages of having protection at various layers. There could be some disadvantage also if they are not very well coordinated, which usually would not be because these layers are sort of independent of each other. They talk

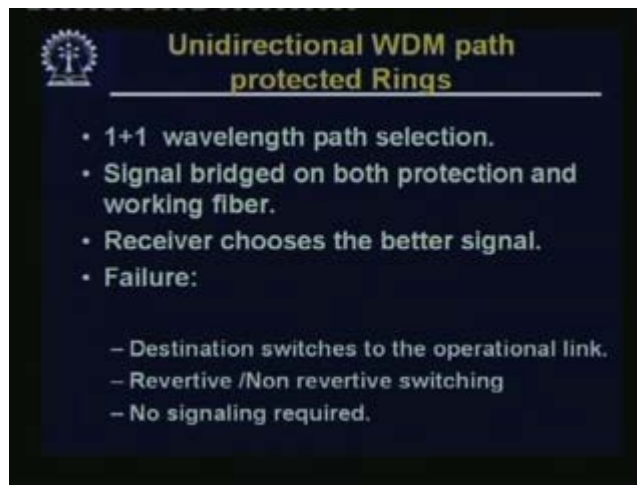
only to their peers and go through their own protocol and give protection. Apart from these of course the other disadvantage is that you may redo some part of the protection unnecessarily. You may be unnecessarily duplicating the work at various places. So these are the disadvantages of having protection at multiple layers.

(Refer slide time: 51:59)



What is the advantage of optical layer protection – speed and efficiency. Limitation would be detection of all faults may not be possible; protects traffic in units of light paths. So this is another problem. As I mentioned, in light path the granularity is very coarse. It may be 2.5 Gbps; so you are really giving the protection at a level of granularity, which may be quite high as I just now mentioned; it could lead to race conditions when optical and client layers both try to protect against the same failure.

(Refer slide time: 52:35)



Of course, on the optical layer you have one more dimension to play with, which is in the case of a WDM. That means you have different wavelengths, so instead of 1+1 link, you can talk about 1+1 wavelength path selection. You can sort of try to select two independent light paths and the signal is bridged on both protection and working fibers if you are doing 1+1 protection kind of thing; the receiver chooses the better signal. In case of a failure, the destination switches to the operational link that is operational light path; there is revertive or non revertive switching; that means, if the original link comes back, it may revert or it may not revert; and no signaling is required.

(Refer slide time: 53:29)



So that is the unidirectional light path. I have just shown you some of the schemes which are used for protection and restoration. And as I mentioned, the schemes may be partially deployed and some of the parts may be protected and some of the parts may not be protected and so on. But the essential idea is the same – that you build in some kind of redundancy and the redundancy may be in the form of entire fibers or the redundancy may be in the form of light paths or wavelengths. It may be pre-computed and pre-reserved like 1+1; it may be pre-computed but not pre-reserved like in a shared one or it may be computed as and when the failure occurs; that is, in the case of restoration. So there are various approaches to it. Depending on how critical the problem is or how critical the application is, what is the cost and how much extra provisioning you can do, you can choose your own way of protection and restoration. Thank you.

Preview of the next lecture

Lecture – 13

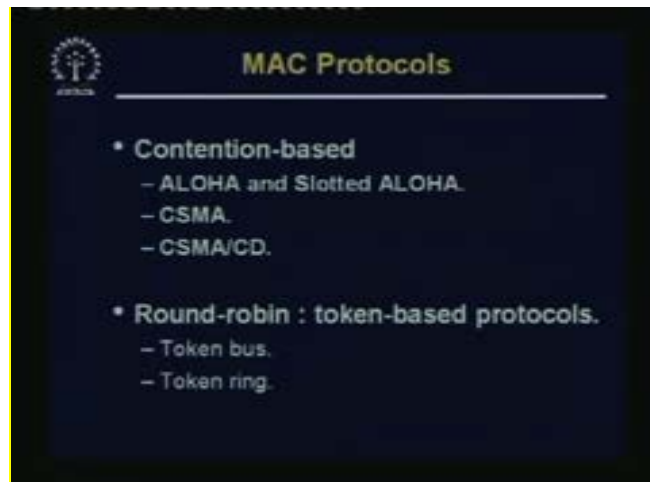
Multiple Access

Good day so today we will talk about multiple access ok now what is multiple access (Refer slide time: 54:59-57:54)



If you remember that we had seven layers in the so-called OSI stack, the top one being application, then we have presentation, session, transport, network, link, and physical. Ok, we had been mostly talking about the physical layer till now, although for optical networks we sort of ventured into

some of the higher layers but from this lecture onwards we want to concentrate on the link layer
ok now what medium access does is it coordinates competing request request for what for
medium that means that there is a medium which may be an object of contention meaning that I
mean several nodes may want to use it and this medium access control protocol has to do with
how to handle that so sharing of link and transport of data over the link that is an general the
description of what the data link layer does so when we share a link there is a question of
committing request and we have to have some way of reserving that and of course there is also
question of transport of data over a link link if you remember when we say a link I mean that
two nodes which are connected the nodes may be or computers routers switches etc they are two
networking nodes they are directly connected now the when I do like this it might mean cable
copper cable or a fiber or it might mean a shared medium like the free space ok so there is a but
there is some way of communicating directly between these two nodes that is what we mean by
link that means that is just one hop in the network that is what we are talking about in the data
link layer so there is a question of reliable transfer of data over these link and if this link like
when you have a free space transmission with so many nodes in the network now so many
people would like to transmit so there is a question of sharing this medium and there is a
question of who would access it when and just as I said free space could be a shared medium
similarly if you remember that if you have some kind of a bus if you remember our discussion
about topology of networks when we have some kind of a bus from which the number of nodes
are hanging and that bus may also be an object of contention so that bus is the medium through
which communication is taking place and there is an there is some kind of competition or
sharing between the of this shared medium between the nodes so we have to handle that that is
the other thing so examples of contention based or ALOHA and slotted ALOHA (Refer slide
time: 57:57-59:00)



these refers to some protocols which we are used in satellite communication so we will discuss these and we talk about satellite communication we have CSMA it stands for carrier sense multiple access or CSMA CD which is carrier sense multiple access with collision detection there are other variance of these like carrier science multiple access with collision avoidance and things like that so these ALOHA slotted ALOHA for satellite CSMA CSMA CD or specific specifically CSMA CD is used by Ethernet in many situations and then we have CSMA CA may be for cellular communication etc so these are contention based MAC and Round Robin there are these are token based protocols and so two very common ones are token bus and token ring so actually we will discuss these in the next in this lecture as well as the next we will be discussing these