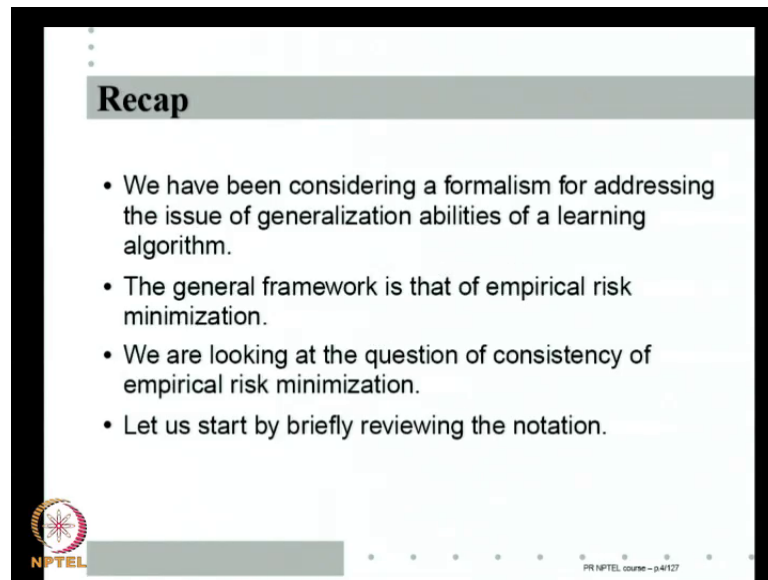


Pattern Recognition
Prof. P. S. Sastry
Department of Electronics and Communication Engineering
Indian Institute of Science, Bangalore

Lecture - 22
Consistency of Empirical Risk Minimization

(Refer Slide Time: 00:43)



Recap

- We have been considering a formalism for addressing the issue of generalization abilities of a learning algorithm.
- The general framework is that of empirical risk minimization.
- We are looking at the question of consistency of empirical risk minimization.
- Let us start by briefly reviewing the notation.

NPTEL PR NPTEL course - p4127

Hello and welcome to the next lecture in this pattern recognition course. We have been considering some basics of statistical learning theory specifically we are looking at a formalism that will allow us to address the issue of whether a learning algorithm learns correctly or what we call the generalisation abilities of a learning algorithm. So, the we have, we are basically considering a formalism for addressing the issue of generalisation abilities of a learning algorithm. What we have looking at is the general framework of empirical risk minimization.

We, are looking at approaches that define a suitable loss function, and then try to minimise the risk that is the general framework we are following. And the I, the question we are looking at is what I told you last class it is called the consistency of empirical risk minimization; that is that the minimizer of empirical risk, is it close in a proper sense to the minimizer of the true risk. So, let us briefly review the formalism, once again to get our notation, and then go ahead and see what is needed.

(Refer Slide Time: 01:40)



We are given

- $\mathcal{X}$ – input space; (as earlier, *Feature space*)
- $\mathcal{Y}$ – output space (as earlier, *Set of class labels*)
- $\mathcal{H}$ – hypothesis space (*family of classifiers*)

Each $h \in \mathcal{H}$ is a function, $h : \mathcal{X} \rightarrow \mathcal{A}$,
where $\mathcal{A}$ is called *action space*.

- Training data: $\{(X_i, y_i), i = 1, \dots, n\}$
drawn *iid* according to some distribution P_{xy} on
 $\mathcal{X} \times \mathcal{Y}$.

PR NPTEL course – p 10/127

So, this is the formalism that we started with the learning problem. We are given the following, we are given this script X which is the input space which for us is the feature space very often it is a d dimensional Euclidian space. All the input feature vectors come from X , but because it can be classification regression anything we, we use an abstract symbol X for the input space. Similarly, Y is the output space for the classification problem this sort of class labels; if I considering two class problem Y will be 0 1 or plus 1 minus 1, m class problem; it could be 1 to 3 m or it could be m unit vectors depending on what representation you want to use, for a regression problem it will real line and so on.

So, the target values are in the set Y that is called the output space. Then we had what is called the hypothesis space, this essentially is the family of classifiers or family of models or family of functions over which you are searching to find the one during the examples. As you know in, in our simpler earlier formalism we called it hypothesis space. And that allows only binary valued functions, but in general it can be a arbitrary set of functions from functions with domain X .

So, specifically we defined elements of this hypothesis space to be functions that map the input space X into some other set $\mathcal{A}$ which we call the action space. So $\mathcal{A}$ could for example, be Y if we think that each h is a classifier and say for example, Y could be 0 1. Then each h each little h which is a element of hypothesis space could be a binary valued

function on X or even when Y is equal to $\{0, 1\}$ as we saw A could be the real line then h assigns a real number to each feature vector so it is like a discriminant function.

So, you know this, this is a generic kind of formulation whereby we can accommodate many different learning situations. So, hypothesis space consists of functions that map the input space or the feature space to some other set called action space. Mostly it will be either be class labels or some real numbers. Then the training data is as usual sets X_i y_i via X_i in script X y_i in script Y which are drawn independently according to some distribution $P_{x,y}$ on $X \times Y$.

So, we are not talking of any target concept as I explained last class. So, there is no one particular hypothesis in the hypothesis space according to which all training data are classified, training data are simply drawn using some distribution $P_{x,y}$ on script $X \times$ script Y . So which in particular means as we explained last class given a particular X more than one Y can have positive probability. So, because we can always factorise $P_{x,y}$ as the marginal of X multiplied by the conditional of Y given X .

So, the conditional of Y given X is a proper distribution then same X can come from with different with different probabilities. This is like overlapping class conditional densities and so on. There is most often we have such situations of overlapping class conditional densities or other similar source of noise, this that the examples are drawn i.i.d according to distribution $X \times Y$ as I explained last class allows us to handle many realistic situations. But now that there is no target hypothesis as such, because given these examples I am asking which h is the best we have to define some goal for learning. So, we have used the loss functions for defining the goal of learning.

(Refer Slide Time: 05:31)



- Loss function: $L : \mathcal{Y} \times \mathcal{A} \rightarrow \mathbb{R}^+$.
- The risk function, $R : \mathcal{H} \rightarrow \mathbb{R}^+$, is given by

$$R(h) = E[L(y, h(X))] = \int L(y, h(X)) dP_{xy}$$

We assume that L is bounded so that the expectation always exists.

- Let $h^* = \arg \min_{h \in \mathcal{H}} R(h)$
- We define the goal of learning as finding h^* , the global minimizer of risk.

PR NPTEL course - p 15/127

So, a loss function maps $\mathcal{Y}$ cross $\mathcal{A}$ to $\mathbb{R}$ that is given a target value at a class label and a h X given a Y and a h X . It assigns a real number by a convention we will assume loss values are always non negative that is why we said $\mathbb{R}^+$. So, the idea is L of Y comma h X tells you on a random example X Y , if I use h to predict Y what is the loss of R . Then we define the risk function so on a random example Y X comma Y I of Y comma h X tells me the loss I suffer.

So, if I take the expectation of this with respect to the P_{xy} distribution. So, I have written the expectation as an integral with respect to P_{xy} distribution. So, if I take expectation of L of Y comma h X with respect to the P_{xy} distribution that is what we call the risk of the function h . So, risk of h is the expectation of loss, it tells you on the average, the average is defined by on samples drawn according to P_{xy} from X cross Y how well does h do. And doing well is in terms of the loss function which we specify we can decide how to measure the loss we have seen many different loss functions and so on last class.

So, this is the risk function and just for mathematical matchness, if you are defining this as the risk function for every h in my hypothesis space, this expectation integral should exist, L is anyway bounded from one side we are assuming it to be non negative. So, we will simply assume L is bounded on the other side also, it is the expectation always exists that is takes away any, any more technical questions on whether for every h R h X and so

on. Most of what we say quite a bit of it can also be said for a loss function that are not really bounded, but satisfy some other criterion, but it really, because we are only looking at a generic flavour of this statistical learning theory results. It does not take away anything if you make a slightly more simplifying assumption actually later on we will make even more simplifying assumptions.

So, we are simply assuming that L is bounded, so that the expectation always exists, we h^* is defined to be the global minimizer of h . So, it is the argument over the script h of R of h , so h^* is that function in my hypothesis state which achieves the global minimum value of risk. Of course, as we have seen h^* may not be unique that really does not matter we will compare different functions only with respect to their risk value. So, the idea is that we, our the goal of learning is to find h^* , essentially h^* is not unique by h^* we mean we want to find something whose risk is same as R of h^* that is the global minima of risk.

So, we are essentially interested in the global minima of risk, because we do not have a target concept now we, we, we specified a loss function that tells how to measure, how good the prediction by a particular function is and the expectation of loss is risk. And the goal of learning is to find a function that achieves the global minimum value of risk. The problem of course, is that we cannot directly minimise R given a particular h I cannot calculate R of h , because I do not know $P x$. So, as a matter of fact as we seen earlier our whole idea of the formalism is we should our algorithm should work no matter what $P x y$ is. So, given a particular h I cannot even calculate $R h$.

(Refer Slide Time: 09:35)



- However, we can not directly minimize $R(\cdot)$.
- The **empirical risk function**, $\hat{R}_n : \mathcal{H} \rightarrow \mathbb{R}^+$, is defined by
$$\hat{R}_n(h) = \frac{1}{n} \sum_{i=1}^n L(y_i, h(X_i))$$
- Let
$$\hat{h}_n^* = \arg \min_{h \in \mathcal{H}} \hat{R}_n(h)$$
- An algorithm learns $\hat{h}_n^*$ by minimizing empirical risk.

PR NPTEL course - p 19/127

So, minimising R directly is not feasible. So, what is our idea? We define an empirical risk function which is actually the sample mean estimator of R , R is the expectation of $L(Y, h(X))$. We have $X_1, y_1, \dots, X_n, y_n$ as i i d samples. So, instead of the expectation I take the sample mean that is $\frac{1}{n} \sum_{i=1}^n L(y_i, h(X_i))$ and that is what we call the empirical risk function.

Given our training data set $X_1, y_1, \dots, X_n, y_n$ I can calculate the empirical risk function for every h and we define the global minimizer of the empirical risk function as $\hat{h}_n^*$. So, this is an estimate of h^* that is why a hat and estimate obtained through n samples that is why the subscript n . So, we call this $\hat{h}_n^*$ and what does machine learning algorithm do? It learns $\hat{h}_n^*$ by minimising the empirical risk, this is our overall framework. So, we are interested; we specify some loss function; we are interested in minimizing risk. So, we want to find a h that h is global minimum risk we cannot minimise risk directly. So, we minimize empirical risk instead.

(Refer Slide Time: 10:53)



- Our objective is to find h^* , minimizer of risk $R(\cdot)$.
- Since we do not know R , we minimize $\hat{R}_n$ instead, and thus find $\hat{h}_n^*$.
- Hence we are interested in the question
$$R(\hat{h}_n^*) \rightarrow R(h^*)?$$
where the convergence is in probability.
- This is the issue of consistency of empirical risk minimization.

PR NPTEL course - p.23/127

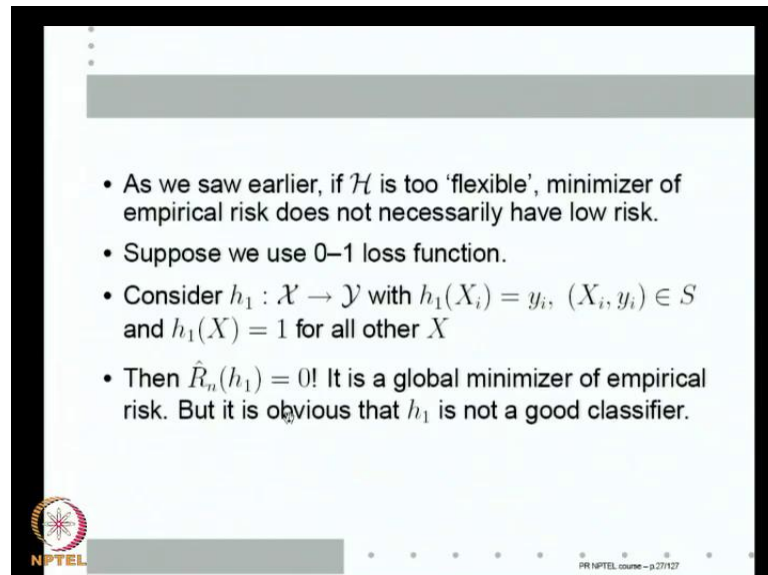
So, let us sum this up, our objective is to find h^* which is the global minimiser of risk. Since we do not know R , we cannot minimize R we minimize $\hat{R}_n$ you, you cannot do this so, whatever we can do, we do and that is find h^* $\hat{h}_n^*$. So, the question is is $\hat{h}_n^*$ a good, good approximation to h^* . And as I said we compared 2 functions only in terms of their risk, otherwise the, the whether the functions are same or not does not matter to us. If risk of 2 functions are same as far as we are concerned the functions are same.

So, essentially we are interested in the true risk of $\hat{h}_n^*$ that is R of $\hat{h}_n^*$ will be same as R of h^* , same in the sense of as number of examples goes to infinity. So, we are asking does R of $\hat{h}_n^*$ converts to R of h^* as n goes to infinity. If it does given sufficiently many samples I can be sure that what I learn has a true risk which is not true for away from the global minimum minimizer of risk. As we have already seen $\hat{h}_n^*$ depends on the random sample. So, it is a random variable so this is sequence of random variable.

So, I have to decide what this sense of this convergence is and the convergence is in probability. So, this is the issue of consistency of empirical risk minimization that is what we want to address now. So, the issue of consistency of empirical risk minimization is I am learning $\hat{h}_n^*$ while I, I am interested in h^* . Of course, functions are distinguishable for me only, based on there the values of their risk. So, is the true risk of

$\hat{h}$ is close to the global minimum of risk which is R of h^* , this is what I want to ask.

(Refer Slide Time: 12:46)



- As we saw earlier, if $\mathcal{H}$ is too 'flexible', minimizer of empirical risk does not necessarily have low risk.
- Suppose we use 0–1 loss function.
- Consider $h_1 : \mathcal{X} \rightarrow \mathcal{Y}$ with $h_1(X_i) = y_i$, $(X_i, y_i) \in S$ and $h_1(X) = 1$ for all other X
- Then $\hat{R}_n(h_1) = 0$! It is a global minimizer of empirical risk. But it is obvious that h_1 is not a good classifier.

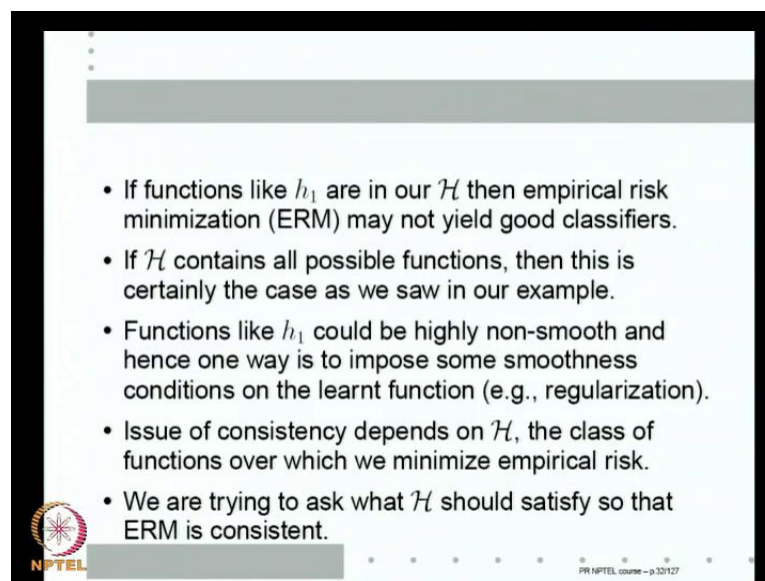
As we have seen earlier the reason the question needs to be asked is that is not always obvious, we can have $h \in \mathcal{H}$ such that if I minimise empirical risk I do not get a good function, we have seen an example where the minimizer of empirical risk gives me nothing. So, as we saw essentially if the hypothesis space is too flexible it has all kinds of functions. Then the minimizer of the empirical risk does not necessarily have a low true risk as a matter of fact. The example we considered we of course, made a particular R_2 classification problem under all that that is to make an example more interesting, but the essence of the example can be stated in just 2 lines.

Let us say we are considering the 0 1 loss function then consider a particular function h_1 that maps $\mathcal{X} \rightarrow \mathcal{Y}$, let us assume $\mathcal{Y}$ is equal to $\{0, 1\}$. So, all our hypothesis in the hypothesis space are functions that map $\mathcal{X} \rightarrow \mathcal{Y}$. So, let us consider a particular element of hypothesis space which is defined as $h_1(X_i) = y_i$, where X_i, y_i are samples that is on each of the training examples $h_1(X_i)$ takes the proper value the y_i value. And $h_1(X)$ is equal to 1 for all other X , what kind of a classifier is h_1 ? I just memorise blindly all the examples I have seen. If I see exactly the same example again I will recall it is class one, every other example I will simply close my eyes and say it is class zero, that is what h_1 is doing $h_1(X_i) = y_i$ for all the training examples X_i, y_i for all other X $h_1(X) = 1$.

1 of X is 1; obviously, a useless classifier it learns nothing it does not generalise it cannot do anything about unseen term, but its empirical risk is 0.

And hence it is a global minimum of empirical risk while it is obvious that h_1 is not a good classifier is as a matter of fact not a classifier at all that is not what we want. But the empirical risk of h_1 is 0 and hence is a global minimizer of empirical risk. This was the problem we had in our example, we just created an entire two dimensional pattern recognition problem where this happens. But essentially this is the issue if I have functions like h_1 then there, there will be global minimizers of empirical risk there might be other global minimizers of empirical risk. But how can I stop my algorithm from picking functions h_1 , because they are also global minimizers of empirical risk there will be many global minimizers of empirical risk. And the algorithm picks up something based on some general (()) it is difficult to ensure that it will not pick up this kind of things.

(Refer Slide Time: 15:40)



- If functions like h_1 are in our $\mathcal{H}$ then empirical risk minimization (ERM) may not yield good classifiers.
- If $\mathcal{H}$ contains all possible functions, then this is certainly the case as we saw in our example.
- Functions like h_1 could be highly non-smooth and hence one way is to impose some smoothness conditions on the learnt function (e.g., regularization).
- Issue of consistency depends on $\mathcal{H}$, the class of functions over which we minimize empirical risk.
- We are trying to ask what $\mathcal{H}$ should satisfy so that ERM is consistent.

So, the issue is if function like h_1 or in our hypothesis space then empirical risk minimization may not yield good classifiers we may say take simple functions. But as we have seen in our examples, the function like h_1 is in our hypothesis space and that happens to be simple by our definition of simple function. A differential simple function meaning; the smallest set, look like a very nice idea. It is a nice idea if I chose at the H , but if I chosen a two flexible a H as we saw it is a bad choice.

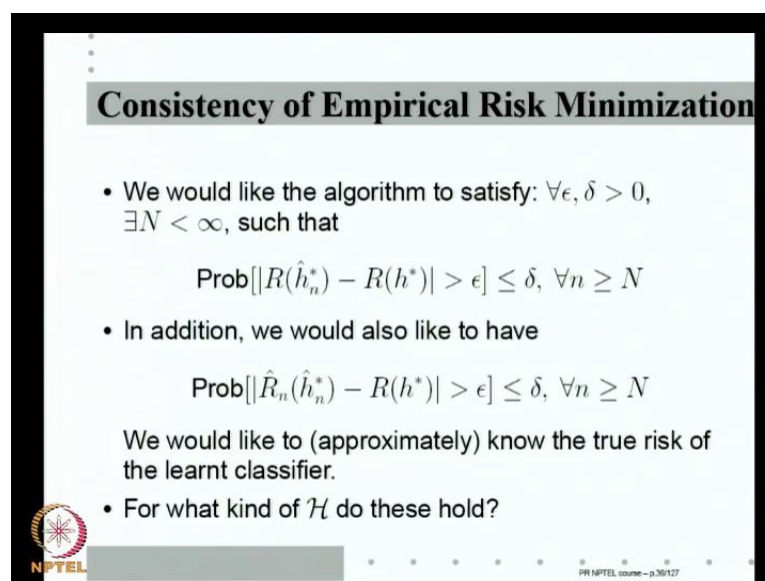
So, it is in general difficult to avoid empirical risk minimization to pull out to convert to functions like h_1 , if such functions are there in our H . Obviously, if H contains all possible functions say Y is equal to 0 comma 1, the two class classifier. And if I take H to be 2^X all possible two class classifiers on X . Then obviously h_1 will be in it and that is what we saw in our previous example as to why we can never learn. Of course, functions like h_1 could be highly non smooth, because at specific points it has to take proper y_i values everywhere else it has to take one. So, there might be many discontinuities the function may look very non smooth.

So, one way of imposing conditions is somehow look for nice functions. Of course, it is not always easy, we in our example, we put some criteria for niceness which did not turn out to be good. But we can certainly think of looking for some nice criteria to say the learn function should have some sufficiently smoothness properties that is what regularization was in the in the linear least squares algorithms. We saw regularization that is an algorithmic issue, we will we will come back to that many algorithm that we are going to consider from now on will impose something other than just empirical risk. We minimise empirical risk plus something which will ensure that highly non non smooth functions even though they have very low empirical risk will come over with a very high we are minimizing.

So that is one way to avoid this, but at this point for the next few lectures, we are not looking at how to go about algorithmically ensuring; we get very good functions; we are asking more theoretical questions. We are asking when would empirical risk minimization be consistent. So, we now know that consistency depends on the class of functions over which we are minimizing the empirical risk. If we choose proper class of functions it is all we choose a bad class of functions it may not be all.

So, this is a very important insight that is not just the criteria that you are minimising that is important, but over what class of functions you are minimising that criteria is equally important in deciding whether or not you learn well. So, what we are going to do now is we are going to ask what condition should H satisfy? So, that empirical risk minimization over H would be consistent, so request your addressing under what is known as statistical learning theory. As I said already we are we are only looking at a simple version of this, so to get a flavour of how such results look like.

(Refer Slide Time: 19:08)



Consistency of Empirical Risk Minimization

- We would like the algorithm to satisfy: $\forall \epsilon, \delta > 0$, $\exists N < \infty$, such that
$$\text{Prob}[|R(\hat{h}_n^*) - R(h^*)| > \epsilon] \leq \delta, \forall n \geq N$$
- In addition, we would also like to have
$$\text{Prob}[|\hat{R}_n(\hat{h}_n^*) - R(h^*)| > \epsilon] \leq \delta, \forall n \geq N$$

We would like to (approximately) know the true risk of the learnt classifier.

- For what kind of $\mathcal{H}$ do these hold?

 PR NPTEL course - p.36/127

So, let us state more formally what is that we want, this is the requirement of consistency of empirical risk minimization. We want our algorithm to satisfy the following, we wanted R of $\hat{h}_n^*$ to converge to R of h^* in probability what does that mean? You give me any epsilon delta positive, strictly positive epsilon delta. Then I can give you a number capital N such that probability that R of $\hat{h}_n^*$ minus R of h^* greater than epsilon is less than delta if I see if the number of examples I have seen is greater than capital N .

So, no matter what accuracy and confidence with which you want me to learn if you give me sufficiently many examples I can learn. So, this is this means that R of $\hat{h}_n^*$ converges to R of h^* in probability. So, essentially I want given any epsilon delta to find a N as you already seen N is often N will be a function of epsilon delta and is often called the sample complexity. So the N should satisfy that probability the difference between R of $\hat{h}_n^*$ and R of h^* being greater than epsilon is less than delta if I have seen at least n examples.

We may also like for under consistency one more thing, not only R of $\hat{h}_n^*$ should be close to R of h^* . But we may want $\hat{R}_n$ of $\hat{h}_n^*$ also to be close to R of h^* why, why would I need this? Because I, we, we I am learning $\hat{h}_n^*$ and h^* is our symbol for what I want to learn, and any two functions have to be distinguished only based on their risk values. So, this is what I want to be less why should I want $\hat{R}_n$

n of $h_{\text{hat star } n}$ to be close to R of h_{star} the reason is I cannot calculate, so $h_{\text{hat star } n}$ is what I have learnt. Now, I want to know how well $h_{\text{R star } n}$ will perform how well $h_{\text{R star } n}$ will perform is given by R of $h_{\text{hat star } n}$. But I cannot calculate R of $h_{\text{hat star } n}$ but I can calculate $R_{\text{hat } n}$ of $h_{\text{hat star } n}$. So, if I actually look at the error that I am getting on my training data with $h_{\text{hat star } n}$ does it tell me how well $h_{\text{hat star } n}$ in general will perform; obviously, we like to have some knowledge. So, we like to at least approximately know the true risk of what we learnt then we can estimate how well our classifier will do.

So, in addition to satisfying this which is the strict consistency requirement for empirical risk minimization we would also like to satisfy this as it turns out. If I can satisfy this I will be satisfying this also, I am not really asking much we will we will see at least some simple proofs for this. So, this is what we want, so we are asking for what kind of families of functions h do these conditions hold. Now, we are not considering algorithmic issues, we are somehow assuming that given some loss function I have an algorithm that can minimise empirical risk. That can actually find the global minimization of empirical risk which is not a simple issue, we have already seen that the 2 issues in learning; one is optimization and one is statistics.

So, the optimisation part is how for a given loss function how do I minimise empirical risk, we still need to get clever algorithms, we sometimes we may not be able to go to global minimizer of risk and all that, but we will come to that later on. Now, we are assuming that we are somehow finding the global minimizer of empirical risk that is what $h_{\text{hat star } n}$ is I am asking is that good enough if I can find the global minimizer of empirical risk. Then will, will the global minimizer of empirical risk will have risk that is close to the global minimum risk that is achievable.

(Refer Slide Time: 22:56)



- As we already saw, the law of large numbers (that $\hat{R}_n(h) \rightarrow R(h), \forall h$) is not enough.
- As it turns out, what we need is that the convergence under law of large numbers be **uniform** over $\mathcal{H}$.
- Such uniform convergence is necessary and sufficient for consistency of empirical risk minimization.
- We now explain what is meant by uniform convergence.
- Law of large numbers says that sample mean converges to expectation of the random variable.

PR NPTEL course - p.41/127

We have already seen that the law of large numbers is not enough for this though our intuitive idea for minimizing empirical risk the intuitive idea of a minimizing empirical risk is that for any h $\hat{R}_n(h)$ converges to $R(h)$ if there are sufficiently many examples. So, for every hypothesis the empirical risk is a good approximation to the true risk and hence minimizer of empirical risk should be able to approximate to minimizer of true risk that is the intuition.

But we have already seen in our example in our example problem earlier that this is not enough. As it turns out what is enough is this convergence law of large numbers ensure assures us that for any h $\hat{R}_n(h)$ converges to $R(h)$. It simply says that the sample mean estimator converges to the expected value that is what this means. What we need is that this convergence should be uniform over the family of classifiers $\mathcal{H}$.

The family of classifiers or family of functions we are considering should be such that the convergence assured by a law of large number should be uniform over $\mathcal{H}$. Some of you may not know what uniform convergence means we are going to explain that shortly. The uniform convergence is both necessary and sufficient for consistency of empirical risk minimization that is a very very strong result that if a family of classifiers is such that there is uniform convergence then empirical risk minimization is consistent. And if empirical risk minimization over a family of functions is consistent that means

that this convergence will be uniform over that H . So, uniform convergence is both necessary and sufficient for consistency of empirical risk.

(Refer Slide Time: 25:07)



- In the context of empirical risk this means:
- Given any $h, \forall \epsilon, \delta > 0, \exists N < \infty$ such that

$$\text{Prob}[|\hat{R}_n(h) - R(h)| > \epsilon] \leq \delta, \forall n \geq N$$
- The N that exists can depend on ϵ, δ and also on h .
- The convergence is said to be uniform if the N depends only on ϵ, δ and not on h .
- That is, for a given ϵ, δ the same $N(\epsilon, \delta)$ works for all $h \in \mathcal{H}$.

PR NPTEL course - p.45/127

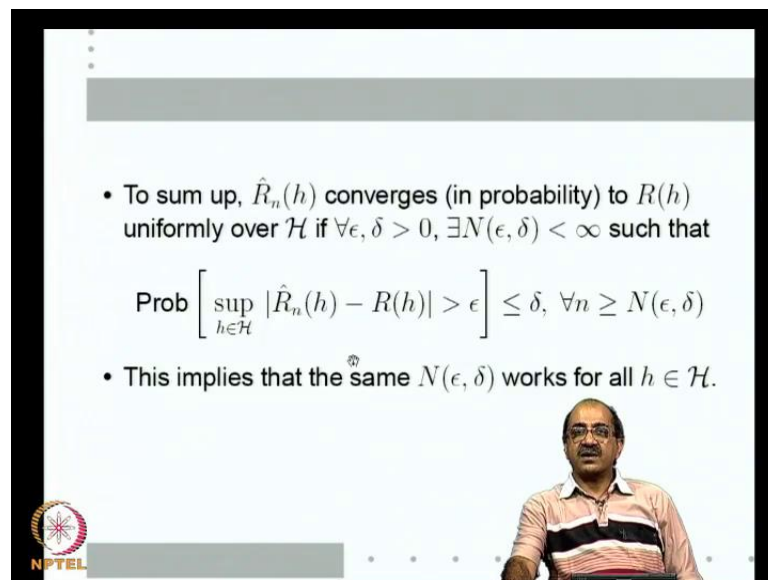
So, first we will explain what is meant by this uniform convergence? What does law of large numbers? Say in words the sample mean converges to expectation of the random variable that is what law of large number says converges in probability converges in almost truly and so on, but we only want weak law. So, we know sample mean converges to expectation of a random variable in probability which in the context of empirical risk. What does that mean? $\hat{R}_n(h)$ is the sample mean estimator of the expectation R of h for any given h . So $\hat{R}_n(h) - R(h)$ greater than ϵ will be a less than δ if I, if I have sufficiently many samples. So, in the context of empirical risk minimization law of large number says that for any for any specific h given any ϵ, δ greater than 0. There exists a N less than infinity such that probability that the difference between $\hat{R}_n(h) - R(h)$ is greater than ϵ is less than δ .

So, you give me any ϵ, δ I will tell you a sample size such that if I have seen that many samples that many i.i.d samples. Then the, the difference between the sample mean and the true mean be greater than ϵ would be less than δ this is what convergence and probability means this is what weak law of large numbers is. Because when we write ϵ, δ definitions, we sometimes do not pay attention to all the necessities even though we may implicitly know them. And we say this n exists the idea

is I fixed a h then you give me an epsilon delta I give you a n . So, what does that mean? The n that exists can depend on epsilon delta and can also depend on h for a particular h for a given epsilon and delta. No matter what epsilon delta you give me I have to exhibit such a N . But you are telling me for this h I want to know whether this converges in probability or not. So to exhibit that I have to say give me any epsilon delta I will find you a N once I can show that. Then the convergence is there which means the n that I have to give you can depend on epsilon delta and also on H .

The convergence is said to be uniform if the N that exists depends on epsilon delta, but not on h . So, for any h and any epsilon delta the N is a function of only epsilon delta naught h , what does that mean? That if I you give some epsilon delta and I give you an a epsilon delta and now this n epsilon N of epsilon delta I wrote n of epsilon delta to explicitly remember that N is a function of the epsilon and delta that we are giving the same number of samples works for all h that is you give me an epsilon delta Then I will give you a number n of epsilon delta if I see that many samples. Then the sample mean estimate of any h will be close to its true expectation that is what uniform convergence means.

(Refer Slide Time: 28:03)



• To sum up, $\hat{R}_n(h)$ converges (in probability) to $R(h)$ uniformly over $\mathcal{H}$ if $\forall \epsilon, \delta > 0, \exists N(\epsilon, \delta) < \infty$ such that

$$\text{Prob} \left[\sup_{h \in \mathcal{H}} |\hat{R}_n(h) - R(h)| > \epsilon \right] \leq \delta, \forall n \geq N(\epsilon, \delta)$$

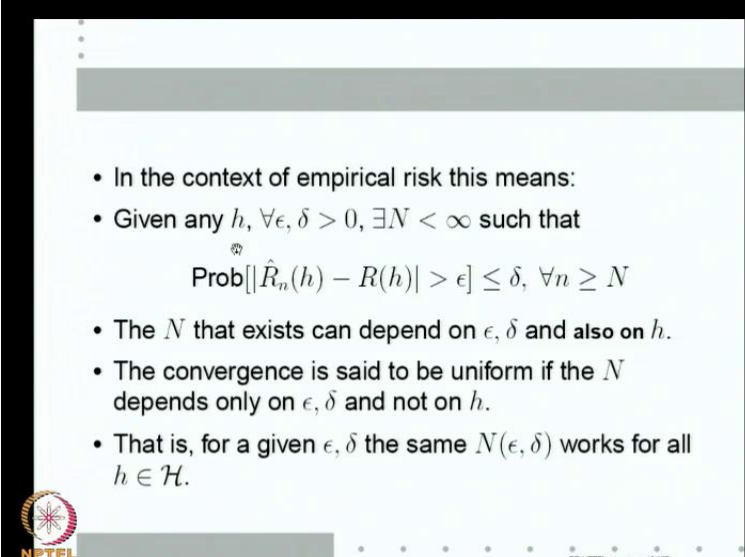
• This implies that the same $N(\epsilon, \delta)$ works for all $h \in \mathcal{H}$.

NPTEL

So, we can write it like that $\hat{R}_n(h)$ converges to $R(h)$ uniformly over the family h if for any epsilon delta you give me any epsilon delta. Then I can give you N of epsilon delta such that probability of this difference $\hat{R}_n(h) - R(h)$ greater than epsilon. But this

difference is $\hat{R}_n(h)$ minus $R(h)$ supremum over all h at the among all the elements of the hypothesis space the maximum difference between $\hat{R}_n(h)$ minus $R(h)$ being greater than epsilon is less than delta. That means if I see $N(\epsilon, \delta)$ samples then for every single h probability is that the reference is greater than epsilon less than delta, because the probability of the maximum difference is greater than epsilon it is less than delta. So, this definition implies that the same $N(\epsilon, \delta)$ works for all h belonging to $\mathcal{H}$. So, this is the definition of uniform convergence.

(Refer Slide Time: 29:13)



• In the context of empirical risk this means:

• Given any $h, \forall \epsilon, \delta > 0, \exists N < \infty$ such that

$$\text{Prob}[|\hat{R}_n(h) - R(h)| > \epsilon] \leq \delta, \forall n \geq N$$

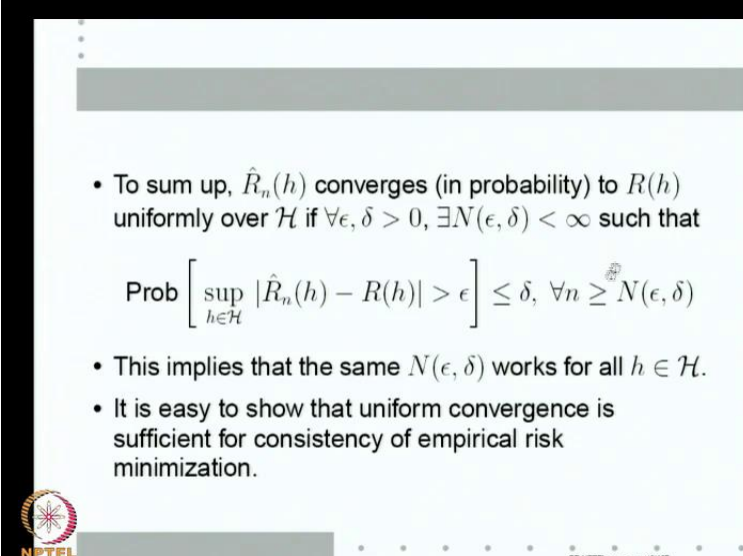
• The N that exists can depend on ϵ, δ and also on h .

• The convergence is said to be uniform if the N depends only on ϵ, δ and not on h .

• That is, for a given ϵ, δ the same $N(\epsilon, \delta)$ works for all $h \in \mathcal{H}$.

NPTEL PR NPTEL course - p-45127

(Refer Slide Time: 29:26)



• To sum up, $\hat{R}_n(h)$ converges (in probability) to $R(h)$ uniformly over $\mathcal{H}$ if $\forall \epsilon, \delta > 0, \exists N(\epsilon, \delta) < \infty$ such that

$$\text{Prob} \left[\sup_{h \in \mathcal{H}} |\hat{R}_n(h) - R(h)| > \epsilon \right] \leq \delta, \forall n \geq N(\epsilon, \delta)$$

• This implies that the same $N(\epsilon, \delta)$ works for all $h \in \mathcal{H}$.

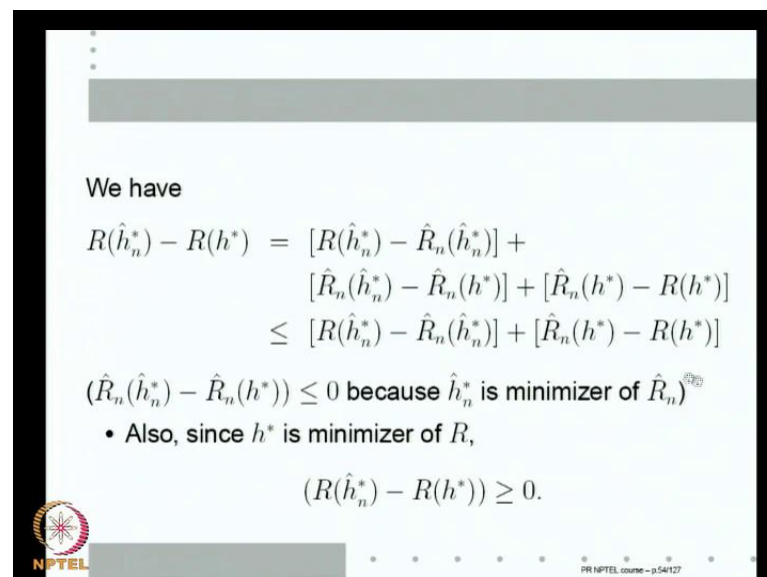
• It is easy to show that uniform convergence is sufficient for consistency of empirical risk minimization.

NPTEL PR NPTEL course - p-45127

So, convergence given by the law of large numbers only means this given a δ epsilon δ there can be N which can be a function of epsilon δ to satisfy this for that H . Whereas, uniform means the supremum over all elements of the hypothesis space of the difference between $R_n(h)$ minus $R(h)$ that supremum value being greater than epsilon should be less than δ . We have to put supremum greater than maximum, because h may be infinite if H is infinite maximum may not be defined generally maximum you, you, you say something is maximum if it is attained in the set.

So, for infinite sets maximum may or may not be attained that is why we call it supremum. So, it is easy to, so if we what we said earlier is that uniform convergence is necessary and sufficient for consistency of empirical risk minimization. So, what, what we will first do? We show that if uniform convergence is there then empirical risk minimization is consistent it is because showing the sufficiency of uniform convergence is much easier. So, let us show that at least then we understand where we are using the uniform convergence. So, we will show that if the convergence given by law of large numbers is uniform then empirical risk minimization would be consistent.

(Refer Slide Time: 30:45)



We have

$$\begin{aligned}
 R(\hat{h}_n^*) - R(h^*) &= [R(\hat{h}_n^*) - \hat{R}_n(\hat{h}_n^*)] + [\hat{R}_n(\hat{h}_n^*) - \hat{R}_n(h^*)] + [\hat{R}_n(h^*) - R(h^*)] \\
 &\leq [R(\hat{h}_n^*) - \hat{R}_n(\hat{h}_n^*)] + [\hat{R}_n(h^*) - R(h^*)] \\
 (\hat{R}_n(\hat{h}_n^*) - \hat{R}_n(h^*)) &\leq 0 \text{ because } \hat{h}_n^* \text{ is minimizer of } \hat{R}_n
 \end{aligned}$$

- Also, since h^* is minimizer of R ,

$$(R(\hat{h}_n^*) - R(h^*)) \geq 0.$$

So, what we have to do consistent of empirical risk minimization? We have to essentially show that $R(h) - R(h^*)$ can be controlled. Of course, the absolute value can be controlled means if you by taking N sufficiently large the difference can be made as small as I want as larger probability as I want. So, let us first handle this

number $R(h^*) - R(h)$ I can write it like this, what did I do? I added a $R(h)$ and I subtracted $R(h)$ and added a $R(h)$.

Once again subtracted a $R(h)$ and added a $R(h)$. So, this I just added and subtracted 2 numbers; one is $R(h)$ another is $R(h)$ it makes no difference. Then I grouped them differently what the, what is the idea? See in the first bracketed term I have the same element of h something that is called h^* . And I have R of that minus R of that this can be controlled using my law of large numbers. Similarly, here I have $R(h)$ and R on the same element of h some element that is called h^* . So, the first and last terms are easily bounded so to say by using law of large numbers.

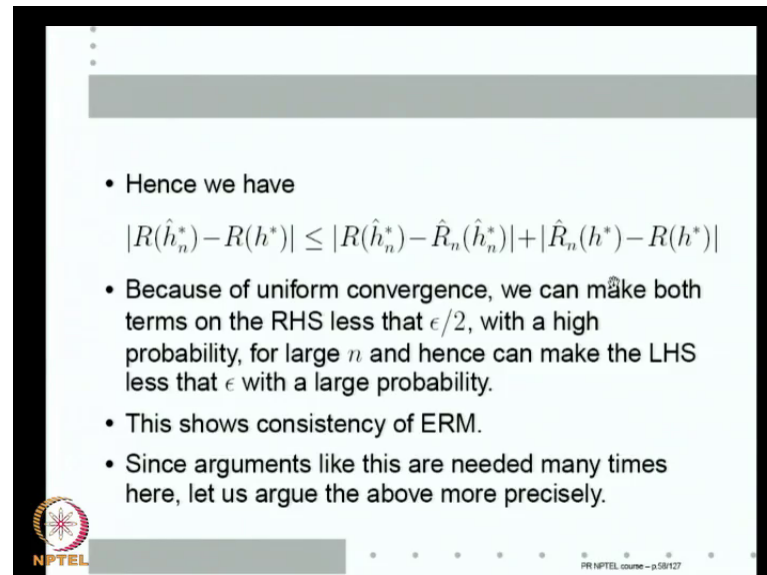
What about the middle term? I can get rid of the middle term, why can I get rid of the middle term? The middle term $R(h) - R(h^*)$ is always negative, it is non positive, why is it non positive? Because by definition h^* is the global minimiser of $R(h)$, so $R(h^*)$ is less than or equal to $R(h)$ of anything. And hence in $R(h) - R(h^*)$, because h^* is the global minimizer of $R(h)$ $R(h) - R(h^*)$ is less than or equal to 0, because this term is negative if I drop this term the $R(h)$ value can only increase.

So, this is less than or equal to this, now I got it in a form I want, because both the terms can now be controlled using law of large numbers. Of course, I still need to put absolute values then if I put absolute values on both sides of inequality the inequality may turn around. But in this case it does not, because this is also I have drawn negative quantity. This is non negative quantity why? Once again h^* by a definition is one that h is globally minimum of risk.

So $R(h^*)$ is less than or equal to $R(h)$ of anything, so $R(h) - R(h^*)$ has to be greater than or equal to 0. That means this term is greater than or equal to 0, because it is less than or equal to term this term is also greater than or equal to 0 and hence if I put absolute value on both sides it will not change. Now, if I put absolute value on both sides this side is absolute value of what I want $R(h) - R(h^*)$. This side I have absolute value of sum of 2 numbers absolute value of a plus B as you

know is by triangular inequality less than or equal to absolute value of a plus absolute value of b.

(Refer Slide Time: 34:13)



- Hence we have

$$|R(\hat{h}_n) - R(h^*)| \leq |R(\hat{h}_n) - \hat{R}_n(\hat{h}_n)| + |\hat{R}_n(h^*) - R(h^*)|$$
- Because of uniform convergence, we can make both terms on the RHS less than $\epsilon/2$, with a high probability, for large n and hence can make the LHS less than ϵ with a large probability.
- This shows consistency of ERM.
- Since arguments like this are needed many times here, let us argue the above more precisely.

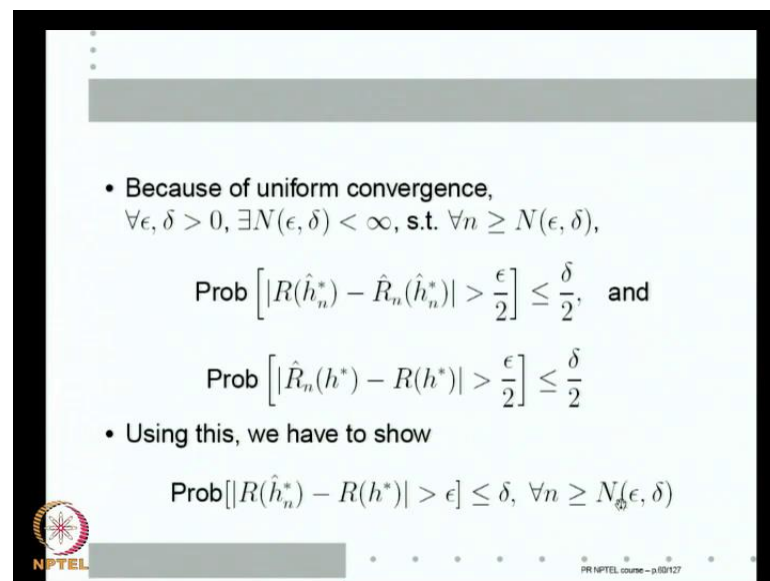
So, what have we got $R(\hat{h}_n) - R(h^*)$ absolute value is less than or equal to absolute the difference between $R(\hat{h}_n) - \hat{R}_n(\hat{h}_n)$ and $\hat{R}_n(h^*) - R(h^*)$. Now, because I have uniform convergence given a particular n the sample mean estimator of every element of h is close to its true value with the same level of accuracy. What does that mean? Because it is uniform convergence I can find a N such that if I have that many samples then both terms on the RHS can be made less than $\epsilon/2$ with a high probability.

Because this is of course, the law of large numbers for $R(\hat{h}_n)$, this is law of large numbers of h^* . But because I have uniform convergence the same n will work for any 2 elements of h . So, with the same number of I can find a number of sample such that both the terms are less than $\epsilon/2$, because both of them are less than $\epsilon/2$. With the same this will also; this will be less than ϵ that is the basic idea of the algorithm, because with the high probability I can make both the terms less than $\epsilon/2$ with a high probability I can make this term less than ϵ that shows consistency of empirical risk minimization.

In while for the rest of this series of talks and series of lectures and statistical learning theory, we will be using this kind of argument many times, we will have a inequality like

this among random variables. And I say if we can control the hand side you can control the left hand side, if you can make this substitution this small with the high probability. You can make this substitution small since people who come across argument for the first time you know may not really see it immediately, so for one time we will we will make this argument very precise. So, so that we also know how to actually argue it using epsilon deltas.

(Refer Slide Time: 36:30)



• Because of uniform convergence,
 $\forall \epsilon, \delta > 0, \exists N(\epsilon, \delta) < \infty, \text{ s.t. } \forall n \geq N(\epsilon, \delta),$

$$\text{Prob} \left[|R(\hat{h}_n^*) - \hat{R}_n(\hat{h}_n^*)| > \frac{\epsilon}{2} \right] \leq \frac{\delta}{2}, \text{ and}$$

$$\text{Prob} \left[|\hat{R}_n(h^*) - R(h^*)| > \frac{\epsilon}{2} \right] \leq \frac{\delta}{2}$$

• Using this, we have to show

$$\text{Prob}[|R(\hat{h}_n^*) - R(h^*)| > \epsilon] \leq \delta, \forall n \geq N(\epsilon, \delta)$$

NPTEL PPI NPTEL course - 080127

What is that we are saying, because of uniform convergence, you give me any epsilon delta there exists a N epsilon delta such that if I have seen N epsilon delta examples. So, if the number of examples n is greater than capital N of epsilon delta. Then for both h R star and n h star which are 2 arbitrary elements of h the, the, the uniform convergence of law of large numbers whole. So, probability absolute value of R h R star n minus R hat n h R star and greater than epsilon by 2 is less than delta by 2. And similarly, of course, I can do it for any epsilon deltas.

So, I can you give me epsilon delta I can put an epsilon two by delta by 2 and find corresponding n here. So, I can achieve this, the uniform, because we are assuming uniform convergence, uniform convergence guarantees that given any epsilon delta. There exists a n to satisfy this, this is what is given, given this, this is what we want to show, for the same N epsilon delta R h R star n minus R h star greater than epsilon

probability is less than delta, this is what you have to show given this we want to show this.

(Refer Slide Time: 37:45)



• Define events A, B, C by

$$A = [|R(\hat{h}_n^*) - R(h^*)| \leq \epsilon], \quad B = [|R(\hat{h}_n^*) - \hat{R}_n(\hat{h}_n^*)| \leq \frac{\epsilon}{2}],$$

$$C = [|\hat{R}_n(h^*) - R(h^*)| \leq \frac{\epsilon}{2}]$$

• Since

$$|R(\hat{h}_n^*) - R(h^*)| \leq |R(\hat{h}_n^*) - \hat{R}_n(\hat{h}_n^*)| + |\hat{R}_n(h^*) - R(h^*)|,$$

we have $A \supset (B \cap C)$ and hence $A^c \subset (B^c \cup C^c)$

PR NPTEL course - p/04127

So, to do this let us define 3 events A, B, C A is $|R(\hat{h}_n^*) - R(h^*)| \leq \epsilon$ this is what we want, B and C are the 2 events that we already know. We can control using law of large numbers this is $|R(\hat{h}_n^*) - \hat{R}_n(\hat{h}_n^*)| \leq \epsilon/2$. This is $|\hat{R}_n(h^*) - R(h^*)| \leq \epsilon/2$; this is the inequality that we already proved, what is that mean? If both B and C are true, if the event both B and C happened, if both B and C are true then this is less than $\epsilon/2 + \epsilon/2$ this is less than ϵ .

So, the sum is less than ϵ , so if both B and C happened then A is also true, what does that tell me given the events like this what it tells me is A is a super set of $B \cap C$. If B and C occur then A occurs of course, A and $B \cap C$ occurs without B and C , because it can be less than ϵ even though one of them is greater than $\epsilon/2$. If other is sufficiently small, so A may be bigger than $B \cap C$, but if both B and C happen then A will happen, so $B \cap C$ is a subset of A . If I take complements on both sides which means A^c is a subset of $B^c \cup C^c$.

(Refer Slide Time: 39:16)



- This gives us

$$\text{Prob}[A^c] \leq \text{Prob}[B^c \cup C^c] \leq \text{Prob}[B^c] + \text{Prob}[C^c]$$
- By uniform convergence, probability of both B^c and C^c are less than $\delta/2$. Hence,

$$\text{Prob}[A^c] \Rightarrow \text{Prob}[|R(\hat{h}_n^*) - R(h^*)| > \epsilon] \leq \delta$$



Now, as we know if A is a subset of B, then the probability of A is less than or equal to the probability of B. So, we know the probability of A complement is less than or equal to the probability of B complement union C complement. And we know the probability of the union of 2 sets is less than or equal to the sum of the probabilities, this is called the union bound. It is a very useful bound in statistical learning theory, we use it again, and again, the probability of the union is always less than or equal to the sum of the probabilities.

(Refer Slide Time: 37:45)



- Define events A, B, C by

$$A = [|R(\hat{h}_n^*) - R(h^*)| \leq \epsilon], \quad B = [|R(\hat{h}_n^*) - \hat{R}_n(\hat{h}_n^*)| \leq \frac{\epsilon}{2}],$$

$$C = [|\hat{R}_n(h^*) - R(h^*)| \leq \frac{\epsilon}{2}]$$
- Since

$$|R(\hat{h}_n^*) - R(h^*)| \leq |R(\hat{h}_n^*) - \hat{R}_n(\hat{h}_n^*)| + |\hat{R}_n(h^*) - R(h^*)|,$$
 we have $A \supset (B \cap C)$ and hence $A^c \subset (B^c \cup C^c)$

PR NPTEL course - p.54/127

(Refer Slide Time: 39:54)



• Because of uniform convergence,
 $\forall \epsilon, \delta > 0, \exists N(\epsilon, \delta) < \infty, \text{ s.t. } \forall n \geq N(\epsilon, \delta),$

$$\text{Prob} \left[|R(\hat{h}_n^*) - \hat{R}_n(\hat{h}_n^*)| > \frac{\epsilon}{2} \right] \leq \frac{\delta}{2}, \text{ and}$$
$$\text{Prob} \left[|\hat{R}_n(h^*) - R(h^*)| > \frac{\epsilon}{2} \right] \leq \frac{\delta}{2}$$

• Using this, we have to show

$$\text{Prob}[|R(\hat{h}_n^*) - R(h^*)| > \epsilon] \leq \delta, \forall n \geq N(\epsilon, \delta)$$

PR NPTEL course - p.03/127

Now, we know probability B complement and probability C complement both are less than delta by 2, what is B complement? This is greater than epsilon one by 2 C complement which is greater than epsilon by 2, that I know is less than delta by 2. So, we know that both B complement C complement of probability less than delta by 2. So, A complement as probability less than delta and A complement is nothing but what we want.

(Refer Slide Time: 40:03)



• This gives us

$$\text{Prob}[A^c] \leq \text{Prob}[B^c \cup C^c] \leq \text{Prob}[B^c] + \text{Prob}[C^c]$$

• By uniform convergence, probability of both B^c and C^c are less than $\delta/2$. Hence,

$$\text{Prob}[A^c] = \text{Prob}[|R(\hat{h}_n^*) - R(h^*)| > \epsilon] \leq \delta$$

• That shown consistency of empirical risk minimization.

PR NPTEL course - p.07/127

(Refer Slide Time: 37:45)



- Define events A, B, C by

$$A = [|R(\hat{h}_n^*) - R(h^*)| \leq \epsilon], \quad B = [|R(\hat{h}_n^*) - \hat{R}_n(\hat{h}_n^*)| \leq \frac{\epsilon}{2}],$$

$$C = [|\hat{R}_n(h^*) - R(h^*)| \leq \frac{\epsilon}{2}]$$
- Since

$$|R(\hat{h}_n^*) - R(h^*)| \leq |R(\hat{h}_n^*) - \hat{R}_n(\hat{h}_n^*)| + |\hat{R}_n(h^*) - R(h^*)|,$$
 we have $A \supset (B \cap C)$ and hence $A^c \subset (B^c \cup C^c)$

PR NPTEL course - p.64/127

(Refer Slide Time: 40:03)



- This gives us

$$\text{Prob}[A^c] \leq \text{Prob}[B^c \cup C^c] \leq \text{Prob}[B^c] + \text{Prob}[C^c]$$
- By uniform convergence, probability of both B^c and C^c are less than $\delta/2$. Hence,

$$\text{Prob}[A^c] = \text{Prob}[|R(\hat{h}_n^*) - R(h^*)| > \epsilon] \leq \delta$$
- That shown consistency of empirical risk minimization.

PR NPTEL course - p.67/127

So, that is how we argued, so basically we should understand that given any inequality like this. If I know that I can make this sufficiently small with high probability and sufficiently small with high probability then I can make this also sufficiently small with high probability any such inequality implies this and this is the, we have showing this. So, that that shows the consistency of empirical risk minimization.

(Refer Slide Time: 40:41)



Consistency of Empirical Risk Minimization

- For consistency, the algorithm should satisfy:
 $\forall \epsilon, \delta > 0, \exists N < \infty$, such that

$$\text{Prob}[|R(\hat{h}_n^*) - R(h^*)| > \epsilon] \leq \delta, \forall n \geq N$$
- We have shown this. In addition, we wanted

$$\text{Prob}[|\hat{R}_n(\hat{h}_n^*) - R(h^*)| > \epsilon] \leq \delta, \forall n \geq N$$
- We can show this also as follows (using the uniform convergence)

PR NPTEL course - p.71/127

So, to sum up consistency the algorithm should satisfy this R of h R star n minus R of h star greater than epsilon will be less than delta for all greater than n given any epsilon delta you should be able to find N and we have shown this. We also said that we like to have this. Now, let us see whether uniform convergence will give me this also it, it is possible to show that this also follows uniform convergence, this is actually a little simple than the other one, what is it that we want to show? Probability of R hat n of h R star n minus R of h star that difference will no, will not be too much large probability, so R hat n h R star n minus R of h star I add and subtract R of h R star n .

(Refer Slide Time: 40:30)



- We have

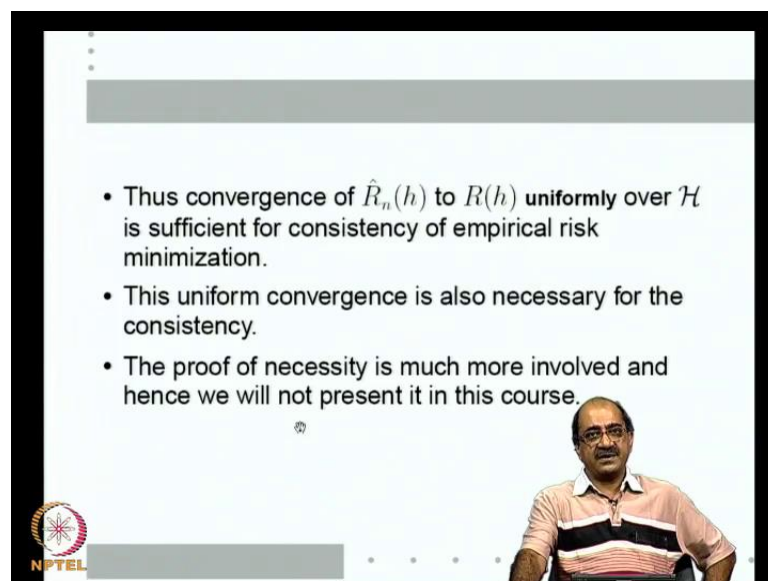
$$|\hat{R}_n(\hat{h}_n^*) - R(h^*)| \leq |\hat{R}_n(\hat{h}_n^*) - R(\hat{h}_n^*)| + |R(\hat{h}_n^*) - R(h^*)|$$
 by triangular inequality.
- By uniform convergence, for sufficiently large n , we can make both terms on the RHS smaller than $\epsilon/2$ with a large probability. Hence we can make the LHS smaller than ϵ with a large probability (for large n).
- This gives us the result we want.

PR NPTEL course - p.74/127

So, I can write this as $\hat{R}_n(h) - R(h) + R(h) - R(h)$. And by triangle inequality absolute value of this will be S_n is equal to sum of the absolute values. Now, the rest is easy this is controlled by law of large numbers, and this is controlled by what we just proved. So, by uniform convergence for sufficiently large N , we can make both terms in the $R(h) - R(h)$ smaller than ϵ , but the same argument we are using we already know that we can make this as small as we want that is given any ϵ there is this n . So, at this is true and this is also true by the uniform convergence and hence, because both of them can be made small with a high probability. This can be made small with a high probability that gives the result we want.

So, we get consistency as well as the, the empirical risk we find of the minimize of the empirical risk that is the global minimizer of the empirical risk where actually be closed to the global minimum of true risk. So, if we just minimise empirical risk we do find something that also minimises globe true risk that is what we want. And the actual empirical we calculate for the global minimizer that the global minimizer of the empirical risk is a good estimate of the true risk of the classifier, say everything is fine if we have uniform convergence. So, if, if the, if $\hat{R}_n(h)$ converges to $R(h)$ uniformly over the class of function h over then it consistency is true.

(Refer Slide Time: 43:13)

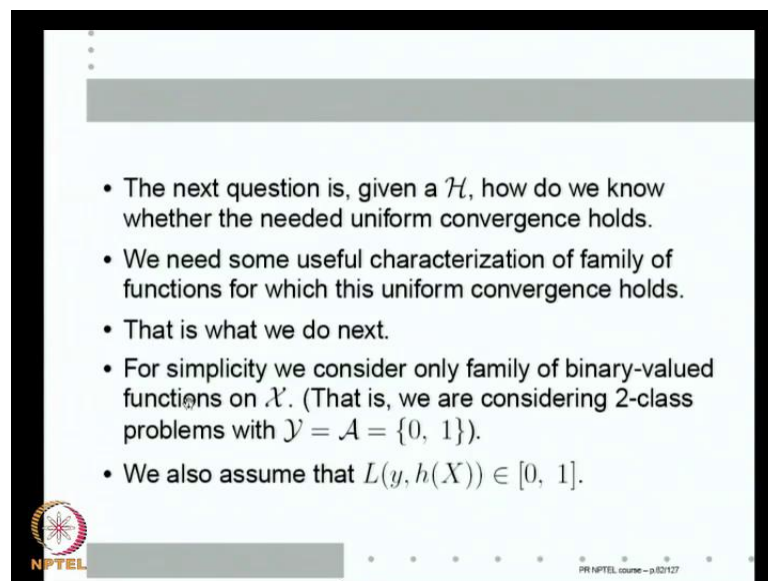


- Thus convergence of $\hat{R}_n(h)$ to $R(h)$ **uniformly** over $\mathcal{H}$ is sufficient for consistency of empirical risk minimization.
- This uniform convergence is also necessary for the consistency.
- The proof of necessity is much more involved and hence we will not present it in this course.

NPTEL

Of course, uniform convergence is also necessary for consistency. And unfortunately the, the, the proof of necessity is much more involved and hence it cannot be presented it is a little beyond the mathematics that we are pegging this course on. So, we would not prove the necessity, but we will just take that uniform convergence is both necessary and sufficient, we have seen the sufficiency proof in sufficiency proof we just take for granted.

(Refer Slide Time: 43:58)



- The next question is, given a $\mathcal{H}$, how do we know whether the needed uniform convergence holds.
- We need some useful characterization of family of functions for which this uniform convergence holds.
- That is what we do next.
- For simplicity we consider only family of binary-valued functions on $\mathcal{X}$. (That is, we are considering 2-class problems with $\mathcal{Y} = \mathcal{A} = \{0, 1\}$).
- We also assume that $L(y, h(X)) \in [0, 1]$.

The next question is given a $\mathcal{H}$ how do we know whether the needed uniform convergence holds or not. So, I wanted consistency you said consistency holds if law of large numbers convergence is uniform over $\mathcal{H}$. Now, how do I check given a $\mathcal{H}$ whether the, the convergence is uniform over $\mathcal{H}$, or not I need some useful and easily calculable characterisation of family of functions for which this holds. So, given a family there must be some simple way for me to check whether or not uniform convergence holds this is what you are going to do next. For this we, we consider only family of binary valued functions on $\mathcal{X}$, we do not consider any hypothesis space that is we are considering $\mathcal{Y}$ is equal to $\mathcal{A}$ is equal to $\{0, 1\}$.

So, we will only do this for two class classifiers and assuming that I am searching over only classifiers so not any real valued function, because during this for real valued functions is this is this will be complicate enough during it for real valued functions is

much more complicated. So, we will consider only this class and we will also assume for simplicity that we already anyway assumed that loss is bounded.

(Refer Slide Time: 45:27)



- First we note that if $\mathcal{H}$ is finite then the uniform convergence always holds.
- Suppose $\mathcal{H} = \{h_1, h_2, \dots, h_M\}$.
- By law of large numbers, for any h_i , given any $\epsilon, \delta > 0$, there would be a $N_i(\epsilon, \delta)$ such that

$$\text{Prob}[|\hat{R}_n(h_i) - R(h_i)| > \epsilon] \leq \delta, \forall n > N_i(\epsilon, \delta)$$
- Take $N(\epsilon, \delta) = \max_i N_i(\epsilon, \delta)$.
- This N would work for all h_i and hence we have uniform convergence.

PR NPTEL course - p.07127

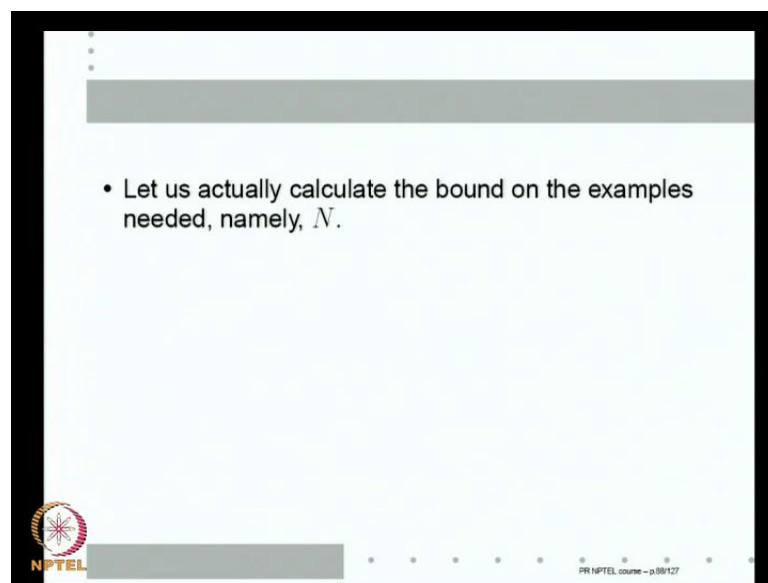
So, we assume is bounded between 0 onward and is bound on other side making it between 0 and 1 is not particularly more distinctive, but this makes the the rest of the mathematical statement little bit simpler. So, now, the question is for what kind of H is the uniform convergence holds, let us look at some simple cases. So, first we know that if H is finite then the uniform convergence always holds by suppose H is some m function $h_1, h_2, \dots, h_M$. It is for some finite number m the law of large numbers any ways say for each h the convergence holds.

So, give me any h_i and an epsilon and delta then I will have a N that is a function of h_i epsilon delta I will write it as $N_i(\epsilon, \delta)$ so give me h_i epsilon and delta then they will be N_i of epsilon delta. So, that probability $|\hat{R}_n(h_i) - R(h_i)| > \epsilon$ or less than delta if N is greater than $N_i(\epsilon, \delta)$. Now, I take $N(\epsilon, \delta)$ to be maximum over i $N_i(\epsilon, \delta)$ I can always do this because I am taking maximum over finite numbers give me any finitely many numbers the maximum always exist.

The will be finite, if we give me finitely many finite numbers each of these N_i 's are finite, if you give me finitely many finite numbers the maximum is always finite. But if you give me infinitely many finite numbers, it is it if you give me 1 comma 2 comma 3

so on. What is the maximum? For any finite subset of it there exist a maximum, but for the that does not exist a max, but because here I only finitely many functions the maximum exist. So, what does this mean? If you now give me enough epsilon delta examples this will hold for every h_i . So, this, this N will work for all h_i and hence uniform convergence holds, this is very straight forward there are only finitely many functions for each h there is an n . And if I take the maximum of all those n 's for all the h S and hence for finite h the uniform convergence always holds. Is of course, is not particularly great insight, but let us actually calculate the bound on the examples needed.

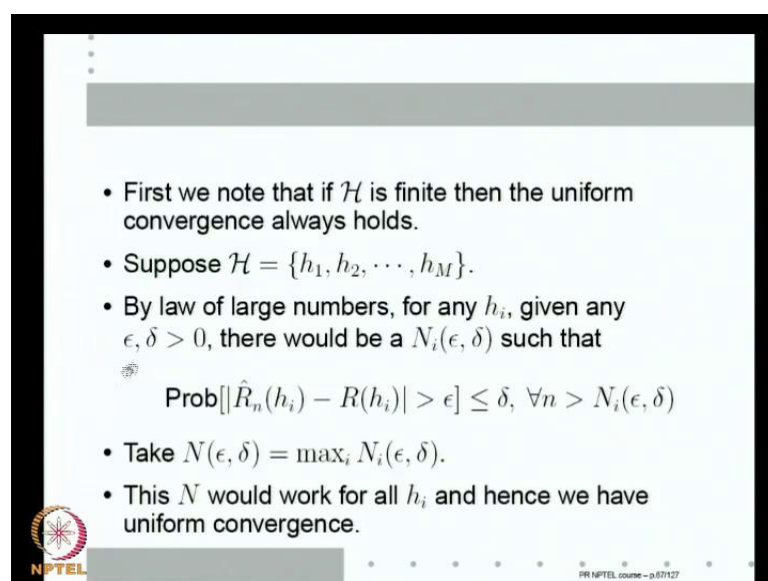
(Refer Slide Time: 47:26)



• Let us actually calculate the bound on the examples needed, namely, N .

NPTEL PR NPTEL course - p.58/127

(Refer Slide Time: 47:35)



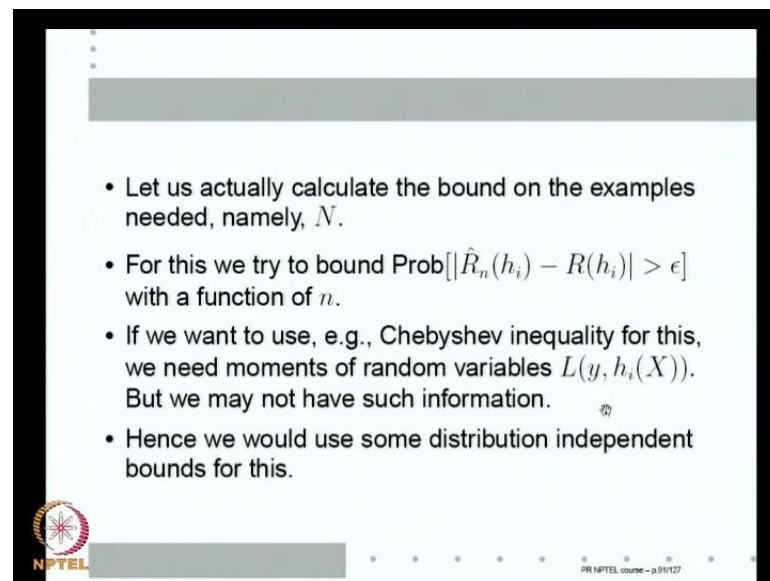
- First we note that if $\mathcal{H}$ is finite then the uniform convergence always holds.
- Suppose $\mathcal{H} = \{h_1, h_2, \dots, h_M\}$.
- By law of large numbers, for any h_i , given any $\epsilon, \delta > 0$, there would be a $N_i(\epsilon, \delta)$ such that

$$\text{Prob}[|\hat{R}_n(h_i) - R(h_i)| > \epsilon] \leq \delta, \forall n > N_i(\epsilon, \delta)$$
- Take $N(\epsilon, \delta) = \max_i N_i(\epsilon, \delta)$.
- This N would work for all h_i and hence we have uniform convergence.

NPTEL PR NPTEL course - p.57/127

Here we have not calculate an epsilon do you saying it axis let us ay can I calculate an epsilon delta. One way I can calculate is if I can bound this probability the LHS probability by a function that depends on little r if I can find it depending on little n. Then I can say that function should be less than delta and then do some algebra to say that function has to be less than delta how large N should be. So, basically the idea is I should bond this probability the probability at left hand side of this inequality by some function of n.

(Refer Slide Time: 48:22)



- Let us actually calculate the bound on the examples needed, namely, N .
- For this we try to bound $\text{Prob}[|\hat{R}_n(h_i) - R(h_i)| > \epsilon]$ with a function of n .
- If we want to use, e.g., Chebyshev inequality for this, we need moments of random variables $L(y, h_i(X))$. But we may not have such information.
- Hence we would use some distribution independent bounds for this.

Then this is a sample mean; this is the expectation I want probability of the difference being greater than epsilon then many bonds for example, I can put a Chebyshev, bond, but all such bonds. So, what is that I want to do? We want to bond this probability with a function of n. Of course, I can use many bonds, but if I use any bonds for example, if you want to use Chebyshev bond. Then you need any variance of this random variable which is like knowing moments of the random variable $L(Y, h(X))$ for X, Y random I got into $P \times Y$ as I do not know $P \times Y$ I cannot calculate those moments.

(Refer Slide Time: 49:07)



• Let Z_i be iid random variables taking values in $[a, b]$, with mean μ . Then the two sided Hoeffding inequality is

$$\text{Prob} \left[\left| \frac{1}{n} \sum_{i=1}^n Z_i - \mu \right| > \epsilon \right] \leq 2 \exp \left(-\frac{2n\epsilon^2}{(b-a)^2} \right)$$

• This gives a distribution independent bound and hence we can use this.



So, if when I want to bound it actually I need to find some proper in equal bonds, the standard inequality is which may need the moments of the random variables or not good enough for me. So, we essentially need some distribution independent bonds there are such distribution dependent bonds. Let us say Z_i is or a Z_i meaning $Z_1 Z_2$ some Z_n or iid random variables. All of them take values in some interval a comma b and let us say there mean is μ , because there iid all of them have same distribution so same mean.

So, $Z_1 Z_2 Z_n$ are independent random variables. So, identically distributed taking values in a comma B and having mean μ then there is what is called a two sided Hoeffding inequality which says that the difference between the sample mean and the expectation being greater than ϵ is bounded above by 2 times exponential minus 2 n epsilon square by b minus a whole square. Say it only depends on n the ranger the random variables are nothing less it does not depend on any moments.

(Refer Slide Time: 50:08)



- Take $Z_i = L(y_i, h(X_i))$. Then Z_i are *iid* random variables taking values in $[0, 1]$.
- Then $\frac{1}{n} \sum Z_i = \hat{R}_n(h)$ and $EZ_i = R(h)$.
- Hence, for any h , we have

$$\text{Prob} \left[|\hat{R}_n(h) - R(h)| > \epsilon \right] \leq 2 \exp(-2n\epsilon^2)$$


So, this is the very useful inequality to use as this is a distinguished independent bond and we can use this for our case two data bond in the finite h case. So, let us calculate the bond essentially what we are saying is Z_i to be L of y_i h X_i then Z_i R i i d random variables taking values in 0 1 is very nice for us b is 1 , a is 0 . So, b minus a is 1 , so even that term will go away that the reason why I have taken 0 1 . Then 1 by n summation Z_i is same as R hat n X and expected value of Z_i is nothing, but R h by definition.

(Refer Slide Time: 51:03)



- Recall that $\mathcal{H} = \{h_1, \dots, h_M\}$.
- In the probability space corresponding to drawing n *iid* samples according to P_{xy} , define the events

$$C_\epsilon^i = \left[|\hat{R}_n(h_i) - R(h_i)| > \epsilon \right], \quad i = 1, \dots, M$$

- We have just seen that

$$\text{Prob}(C_\epsilon^i) \leq 2 \exp(-2n\epsilon^2), \quad \forall i$$

PR NPTEL course - p 109/127

So, what Hoeffding inequality gives me is probability $\hat{R}_n(h) - R(h)$ is greater than ϵ is less than $2 \exp(-2n\epsilon^2)$. Of course, this is not the bound I want. I want to put a supremum inside here. I want to put a supremum over h of this. So, I have to find a bound for that, so that bound can also be obtained, so this is a bound that shows me that as n tends to infinity it goes to 0. So, I can make it as small as I want for example, smaller than δ by taking n sufficiently large, but before that let us first put the supremum inside this probability.

So, H is h_1 to h_M , so basically all our probabilities with respect to drawing antiple of samples. So, in the same probability space let us define an event $C_i(\epsilon)$ over drawing antiples of this, we $C_i(\epsilon)$ set $C_i(\epsilon)$ or $\hat{R}_n(h_i) - R(h_i)$ that is what i minus $R(h_i)$. It is expectation being greater than ϵ that is the event $C_i(\epsilon)$, $C_i(\epsilon)$ says that for the i th hypothesis the empirical risk, and the true risk differ by more than ϵ . Then what we have just shown is probability of the event $C_i(\epsilon)$ is less than $2 \exp(-2n\epsilon^2)$, this is true for every i .

(Refer Slide Time: 51:53)



Now we have

$$\text{Prob} \left[\sup_{h \in \mathcal{H}} |\hat{R}_n(h) - R(h)| > \epsilon \right] = \text{Prob} (C_\epsilon^1 \cup \dots \cup C_\epsilon^M)$$

$$\leq \sum_{i=1}^M \text{Prob}(C_\epsilon^i)$$

$$\leq 2M \exp(-2n\epsilon^2)$$

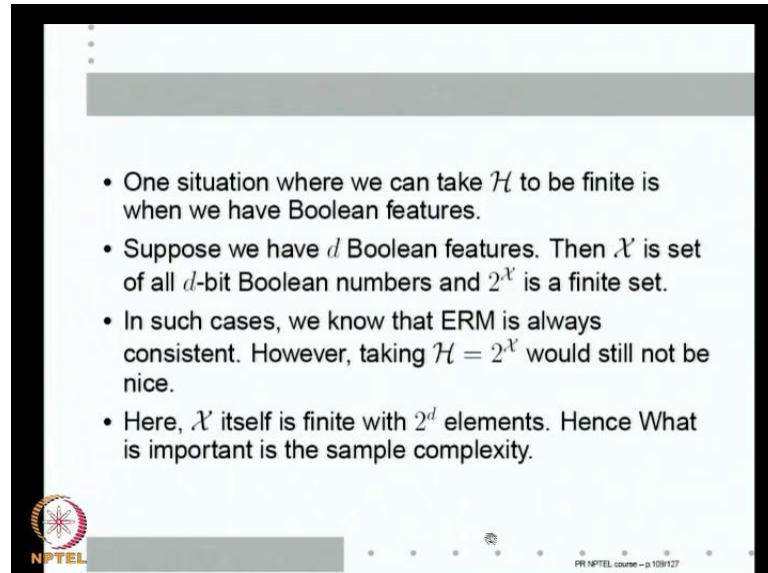
- Now we can find how large n should be so that the bound on the RHS is less than δ

PR: NPTEL course - p 105/127

Now, if I want to put a supremum here. So, I am saying for some what is the probability for some h or the other this difference is greater than ϵ that is same as probability of the union of all the $C_i(\epsilon)$ sets. $C_i(\epsilon)$ tells me that it is true for h_i i.e. asking at least one of the I have for some h_i or the other this is greater, but is simply given by the union of these events. Now, we have already seen union bounds probability of union is

less than or equal to sum of their probabilities and each individual probability we know what it is each individual probability $2^{-n\epsilon^2}$.

(Refer Slide Time: 52:53)



- One situation where we can take $\mathcal{H}$ to be finite is when we have Boolean features.
- Suppose we have d Boolean features. Then $\mathcal{X}$ is set of all d -bit Boolean numbers and $2^{\mathcal{X}}$ is a finite set.
- In such cases, we know that ERM is always consistent. However, taking $\mathcal{H} = 2^{\mathcal{X}}$ would still not be nice.
- Here, $\mathcal{X}$ itself is finite with 2^d elements. Hence What is important is the sample complexity.

NPTEL

PR NPTEL course - p 10/127

So, this is less than or equal to $2m$ times exponential minus $2\epsilon^2$. So, this tells me that this can be made as δ . So, we can now find n how large is n should be so that this can be made less than δ . So, this is how I can actually calculate n if you give me ϵ and δ . Before we go forward where can where is finite h i useful thing if you have Boolean features, if you have Boolean features. Then the $\mathcal{X}$ itself is set of d -bit Boolean numbers $\mathcal{S}$ itself is finite and hence $2^{\mathcal{X}}$ is also finite set.

If $2^{\mathcal{X}}$ is finite set every subset of $2^{\mathcal{X}}$ is also finite set. So, if I am learning Boolean features two class classifiers. Then obviously even multiclass classifier for that matter; obviously, every h I can take will be finite and for all finite h uniform convergence holds, because it does not mean even in finite case taking $2^{\mathcal{X}}$ should be a a nice choice. As we already seen two complete over all possible classifiers never learns well that we can still see why basically, because $\mathcal{X}$ itself is finite saying given sufficiently many examples I can learn means nothing. Because there are only $2^{\mathcal{X}}$ examples if you see all the $2^{\mathcal{X}}$ examples then what more do you want to see there is nothing more to see nothing more to predict.

(Refer Slide Time: 54:01)



- We have shown
$$\text{Prob} \left[\sup_{h \in \mathcal{H}} |\hat{R}_n(h) - R(h)| > \epsilon \right] \leq 2M \exp(-2n\epsilon^2)$$
- We can bound this probability by δ if
$$2M \exp(-2n\epsilon^2) \leq \delta \quad \text{or} \quad n \geq \frac{1}{2\epsilon^2} \ln \left(\frac{2M}{\delta} \right)$$
- The bound may be poor (and may be more than 2^d).

PR NPTEL course - p 112/127

So, what is important is given an epsilon delta how many samples to be made, it is this my M is a total number of elements. If X has 2^d elements and I am considering all boundary valued functions are on X M will be 2^{2^d} . So, this may not be a very nice bound for example, if I want is less than delta and is greater than $\frac{1}{2\epsilon^2} \ln \frac{2M}{\delta}$.

(Refer Slide Time: 54:50)



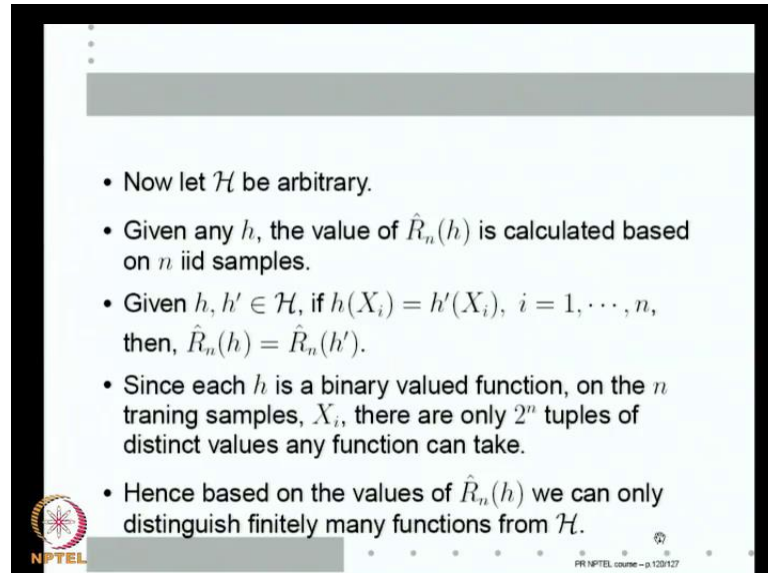
- For specific finite $\mathcal{H}$ we can get better bounds.
- There are classes of Boolean functions that can be learnt efficiently.
- But, for us, the reason for doing the finite $\mathcal{H}$ case is that it gives us ideas on how to tackle the general case.

PR NPTEL course - p 115/127

Now, M is 2^{2^d} , so this might be this bound might be even greater than 2^{2^d} . So, just because we found a bound it does not mean that it is a very nice way to

learn I mean even though it is consistent the number of examples I am asking means show me every single example then I can predict all of them.

(Refer Slide Time: 55:16)



- Now let $\mathcal{H}$ be arbitrary.
- Given any h , the value of $\hat{R}_n(h)$ is calculated based on n iid samples.
- Given $h, h' \in \mathcal{H}$, if $h(X_i) = h'(X_i)$, $i = 1, \dots, n$, then, $\hat{R}_n(h) = \hat{R}_n(h')$.
- Since each h is a binary valued function, on the n training samples, X_i , there are only 2^n tuples of distinct values any function can take.
- Hence based on the values of $\hat{R}_n(h)$ we can only distinguish finitely many functions from $\mathcal{H}$.

NPTEL
PR NPTEL course - p 123127

So, infinite sets actually we want better bounds we are much more interested in how fast is number of examples grows with epsilon and delta. There are class of Boolean function be learnt efficiently, we can show that we do not need exponentially many examples, but we can do with polynomially many examples, But that is not what we are going now I might come back to that. Later the reason why we actually did the finite case in detail is that it gives us ideas on how to tackle the general case the idea is the as follows.

Now, consider h is arbitrary it may have infinite uncountable infinite items, given any h the empirical risk is calculated based on n iid samples. What is that mean? If I have two functions h and h' in my hypothesis space there such that $h(X_i)$ is equal to $h'(X_i)$ on all the n training examples. Then the empirical risk of h on h' is same the h and h' may take many vastly different values on other axis. But let us say they just agree on the training samples I have with they agree on training samples I have then their empirical risk are same.

So I can only distinguish between function while minimising empirical risk only those functions use values differ on the training samples each h is bounded valued function. And I have n training samples given n training samples, how many different values can bounded valued functions taken n type training samples they are only 2^n some different

possible values that any function can take on which means at most I can distinguish through an possible elements 2 in h . Even if h contain uncountable infinite things given any n training examples I cannot distinguish between more than 2 power n different h 's hence based on the values of R hat n we can only distinguish between finitely many functions from h .

(Refer Slide Time: 56:40)



- We can sum-up this insight as follows.
- Given n training examples, as far as empirical risk is concerned, only finitely many (at most 2^n) functions from $\mathcal{H}$ can be distinguished.
- Hence we may be able to employ the argument we used for finite $\mathcal{H}$ case to tackle the general case.

PR NPTEL course - p 123/127

(Refer Slide Time: 57:03)



- Recall that in the finite case we had

$$\text{Prob} \left[\sup_{h \in \mathcal{H}} |\hat{R}_n(h) - R(h)| > \epsilon \right] \leq 2M \exp(-2n\epsilon^2)$$

- Using our insight we may be able to use this but then M would be a function of n . So, this depends on how the number of distinguishable functions grow with n , the number of examples.
- Next, we explore this intuitive idea in a more precise fashion.



So, what is the insight we got? Given N training examples as far as empirical risk is concerned only finitely many at most 2 power n functions form h can be distinguished

which means we may be able to employ the same argument that we have used in this finite h case to tackle the general case that is the reason why looked at it. So, how do I apply the finite h case, in the finite h case I have this where M is the total number of functions unfortunately. If I put M is equal to 2^n , this inequality is useless for me, because this goes exponential with n this falls exponential with n what do I get.

So, using our insight may be able to bond this, but then M would be a function of n . So, this whether or not I can bond this properly depends on how the number of distinguishable functions grow with n the number of examples. We know the given n examples and given a family of functions h they can at most be 2^n function that from h that can be distinguished based on the example, but they may be less.

So, for a particular h the question is how this number of functions grows with n . This is the insight as a matter of fact, we will explore this intuitive idea in more precise fashion next class. But essentially looking at this idea and knowing that this M may become a function of n , we want to know, what is the kind of growth function I should put here? And then I may be able to characterise the h for which rather the distinguishable functions grows with 2^n or not this is what we do next.

Thank you.