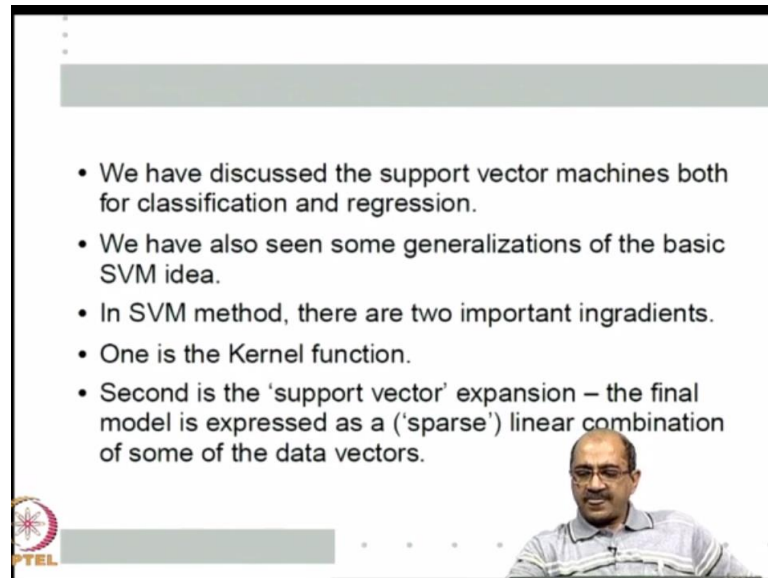


Pattern Recognition
Prof. P. S. Sastry
Department of Electronics and Communication Engineering
Indian Institute of Science, Bangalore

Lecture - 37
Positive Definite kernels; RKHS; Representer Theorem

(Refer Slide Time: 00:33)



The slide contains the following bullet points:

- We have discussed the support vector machines both for classification and regression.
- We have also seen some generalizations of the basic SVM idea.
- In SVM method, there are two important ingredients.
- One is the Kernel function.
- Second is the 'support vector' expansion – the final model is expressed as a ('sparse') linear combination of some of the data vectors.

In the bottom right corner, there is a small video inset showing Prof. P. S. Sastry speaking. The IIT Bombay logo is visible in the bottom left corner of the slide.

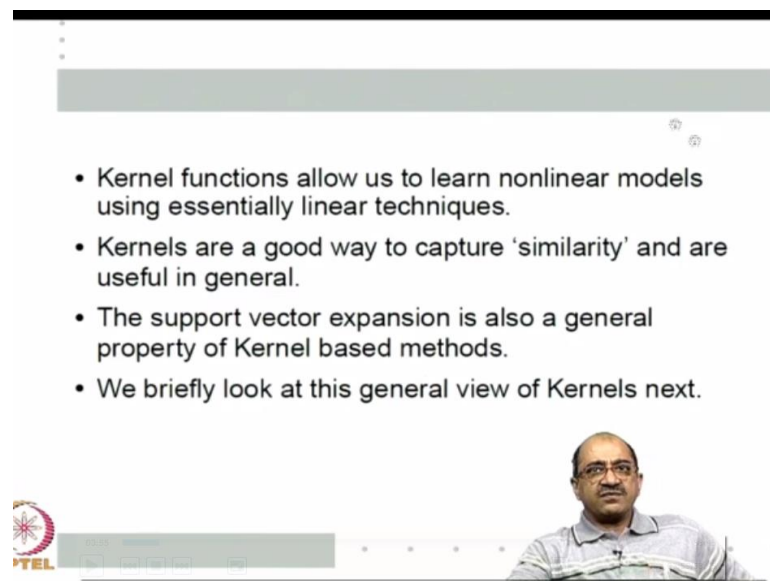
Hello and welcome to this next lecture in the pattern recognition course. We have been looking at support vector machines idea. This lecture, we will close this; we will look at kernels in general. Just to recall, we have looked at the support vector machines both for classification and regression. The basics of electro machine that glance optical hyperplane for a two class classification, and also seen how we can do support regression using epsilon in sensitive loss function. And we also seen some generalization of the basic SVM idea. For example, look at the new SVM the way to add b square to the primal objective functions of that the dual becomes very simple to solve. So, which is called successive over relaxation SVM.

So, we have seen some of the generalization of the basic SVM idea. So, to ultimately take some very broad look at what the SVM idea is added, there essentially two important ingredients in the SVM idea. So, if I am looking at only classification, then they told you the, the power of the algorithm comes from the fact that we are learning optimal hyperplane, which is what allows us to learn a classifier with very low true risk. And by using kernels to do inner products, we are able to learn a nonlinear classifier with

this thing, but if I look at everything all the SVM type ideas. So, essentially kernels is one important idea. And equally important idea, which is what allowed us to use the kernel idea in the first place is that the final function admits a representation in terms of kernel functions.

So, we seen for example, the in the SVM the final w is a linear combination of data vectors. And hence the final classification function of the discriminant function can be expressed as the linear combination of kernels of kernel function values. So, we can call this as support vector expansion. So, representing your final classifier say finite classifier function $f(X)$ where $\sin$ of $f(X)$ will be your classifier, representing that function $f(X)$ as a nice kernel expansion. So, we will call that support vector expansion, so these two are the two ingredients that allowed us to solve nonlinear problem using linear techniques.

(Refer Slide Time: 03:15)



The slide contains the following bullet points:

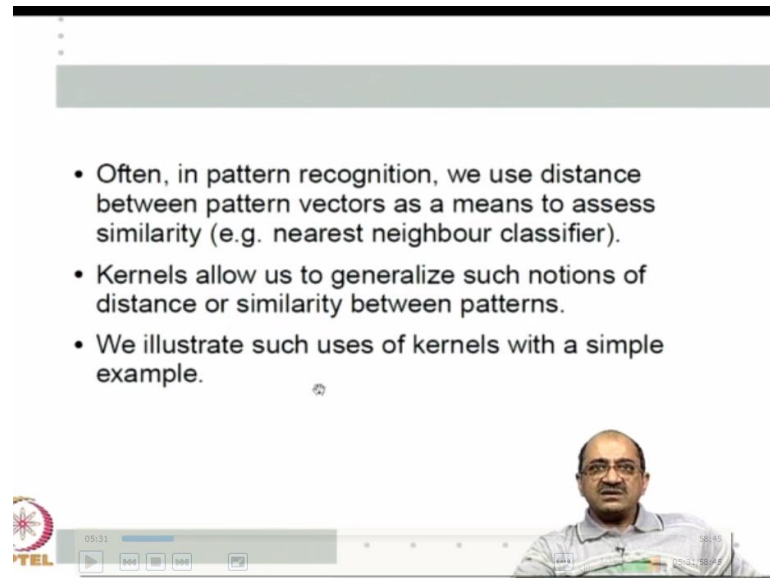
- Kernel functions allow us to learn nonlinear models using essentially linear techniques.
- Kernels are a good way to capture 'similarity' and are useful in general.
- The support vector expansion is also a general property of Kernel based methods.
- We briefly look at this general view of Kernels next.

A small video inset in the bottom right corner shows a man with glasses and a striped shirt speaking.

Now, both these are much general, much more general than what we have seen in the simple SVM method. Kernel functions in general allow us to learn nonlinear models using linear techniques and also a very good way to capture similarity that is useful in a general way. So, kernel functions can be thought of as capturing similarity, and dissimilarity between feature vectors. In the same way support vector expansion is also a very general property of kernel based methods. So, what we will do in this class is to look at this general overview of kernels. So, what I will do is I will first give you a very brief interaction or brief idea of what we mean by kernels at good way to capture

similarity. And then we look at positive different kernels in a little more detail to understand the kernel idea in a slightly more general fraction then what you seen so far.

(Refer Slide Time: 04:16)



The slide contains the following text:

- Often, in pattern recognition, we use distance between pattern vectors as a means to assess similarity (e.g. nearest neighbour classifier).
- Kernels allow us to generalize such notions of distance or similarity between patterns.
- We illustrate such uses of kernels with a simple example.

At the bottom of the slide, there is a video inset showing a man with glasses and a light blue shirt speaking. To the left of the video is a small logo with the word 'TEL' and a red star-like symbol. Below the video is a control bar with a play button, a progress bar, and other standard video controls.

In pattern recognition, we always use distance as a means to access similarity that is one of the ways which you can do this. You know at the beginning of the course, we looked at nearest neighbor classifier. Nearest neighbor classifier I stored some prototypes and when you give me new pattern I find the distance the equilibrium distance of the new pattern to each of the prototypes whichever prototype is closest to I will put that in the class, we have seen that as an interesting very simple classifier. And if I have sufficiently many prototypes it is worst case probability of error is at worst twice that of the basic error.

So, in that sense by simple classifier it does remarkably well. Now, kernel allow us to generalize such notions. So, the basic idea there is 2 patterns are closed to each other in a distance sense then they are also similar patterns. So, the question is what kind of distance is a good distance? So, we may have some similarity function, which we can think of as a kernel. And then we can use the kernel in place of the usual distance. We will, we will look at the very simple example to see how kernels can be used in a nearest neighbor classifier.

(Refer Slide Time: 05:38)



- Consider a 2-class classification problem with training data
 $\{(X_i, y_i), i = 1, \dots, n\}, X_i \in \mathbb{R}^m, y_i \in \{+1, -1\}$
- Suppose we implement a nearest neighbour classifier, by computing distance of a new pattern to a set of prototypes.
- Keeping with the viewpoint of SVM, suppose we want to transform the patterns into a new space using ϕ and find the distances there.

PR NPTEL course - p.15/133

So, for the example let us take as usual two class classification problem, let us say X_i, y_i that is X_1, Y_1, X_n, y_n are the samples and the classes are plus 1 minus 1 and the feature vectors are in $\mathbb{R}^m$. Now, suppose we implement in a nearest neighbor classifier by computing distance from a set of prototypes. And the whole idea is that what are these SVM idea I have amount to that we do not do it is the original feature space. But we map the feature vectors to some other high dimensional space and do distance as this. So, the same way we using some ϕ to map X s and we are actually finding Euclidian distance $\mathbb{R}$ distance using a inner product in the (ϕ) space of ϕ .

(Refer Slide Time: 06:30)



- Suppose we use two prototypes given by

$$C_+ = \frac{1}{n_+} \sum_{i: y_i=+1} \phi(X_i) \text{ and } C_- = \frac{1}{n_-} \sum_{i: y_i=-1} \phi(X_i)$$
 where n_+ is the number of examples in class +1 and n_- is that in class -1.
- The prototypes are 'centers' of the two classes.
- Given a new X , we would put it in class +1 if

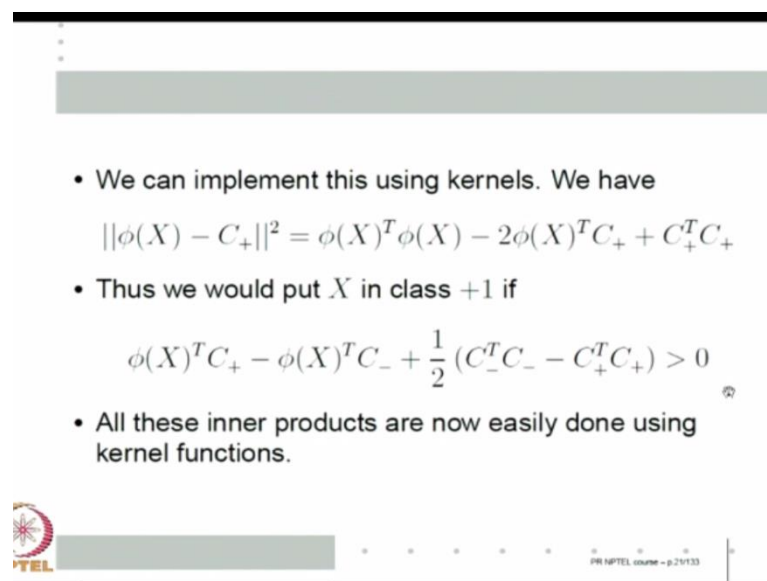
$$\|\phi(X) - C_+\|^2 < \|\phi(X) - C_-\|^2$$

PR NPTEL course - p.15/133

So, let us say we are using 2 prototypes, we call them as C plus and C minus, C plus is $\frac{1}{n_+} \sum \phi(X_i)$ where n_+ is the number of examples of class plus 1 similarly, n_- is the number of examples of class minus 1. So, this is nothing but the average of all the patterns of class one, not in the original space. But in the real space of ϕ because I have mapped the original feature vectors using ϕ to some other space, in that space C plus is the centre of all the class plus 1 examples.

Similarly, C minus is the centre of all the class minus 1 example. So, these are the 2 centers of the 2 classes. So, we use them as the prototypes, the idea is that if you give me a $\phi(X)$ will go to the real space of ϕ . So, I will find the distance between $\phi(X)$ and C plus this is $\|\phi(X) - C_+\|^2$. If that is less than $\|\phi(X) - C_-\|^2$ then I put it in class plus 1 otherwise I will put it in class minus 1 this. So, this is a simple nearest neighbor classifier, but distance has not found in the original feature space but in the ϕ space.

(Refer Slide Time: 07:46)



• We can implement this using kernels. We have

$$\|\phi(X) - C_+\|^2 = \phi(X)^T \phi(X) - 2\phi(X)^T C_+ + C_+^T C_+$$

• Thus we would put X in class +1 if

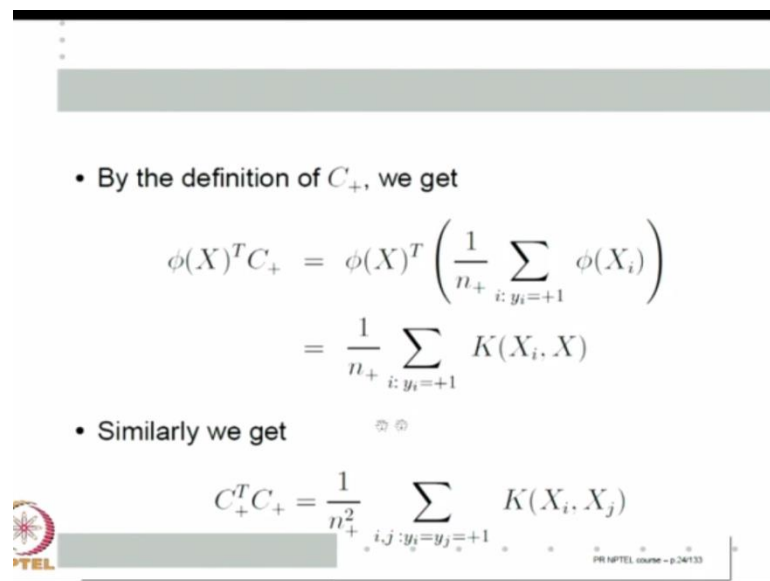
$$\phi(X)^T C_+ - \phi(X)^T C_- + \frac{1}{2} (C_-^T C_- - C_+^T C_+) > 0$$

• All these inner products are now easily done using kernel functions.

Now, we can do all this using kernels as follows, so basically we need to look at this $\phi(X) - C_+$ whole square that is the distance. By expanding $\phi(X) - C_+$ whole square becomes $\phi(X)^T \phi(X) - 2\phi(X)^T C_+ + C_+^T C_+$. Now, I want $\phi(X) - C_+$ whole square greater than $\phi(X) - C_-$ whole square. So, do this simple algebra this constant term cancel both

sides. So, finally, it means we will put X in class plus 1 if this quantity is greater than 0, $\phi(X)^T C_+ - \phi(X)^T C_- + \text{some constant}$ which is $C_+ - C_-$. So, basically I have to do all these inner products to be able to implement my nearest neighbor classifier by first transforming all the patterns using the function ϕ . But the idea is that all these inner products can now be kernelised.

(Refer Slide Time: 08:45)



• By the definition of C_+ , we get

$$\begin{aligned}\phi(X)^T C_+ &= \phi(X)^T \left(\frac{1}{n_+} \sum_{i: y_i = +1} \phi(X_i) \right) \\ &= \frac{1}{n_+} \sum_{i: y_i = +1} K(X_i, X)\end{aligned}$$

• Similarly we get

$$C_+^T C_+ = \frac{1}{n_+^2} \sum_{i, j: y_i = y_j = +1} K(X_i, X_j)$$

NPTEL

Very easy to see that they can be kernelised, what is $\phi(X)^T C_+$? $\phi(X)^T$ transpose, this is my C_+ . Now, push $\phi(X)$ inside then it becomes $\phi(X)^T \phi(X_i)$ which is nothing but $K(X_i, X)$. In a similar way, what is my C_- ? My C_- is this. I want $C_+^T C_+$, what I will ultimately I get is $\phi(X_i)^T \phi(X_j)$ for all i, j pairs at both y_i and y_j are plus 1 and $\phi(X_i)^T \phi(X_j)$ is $K(X_i, X_j)$. So, I get $C_+^T C_+$ like this similarly, I can write $C_-^T C_-$.

(Refer Slide Time: 09:25)

• Thus, our classifier is $\text{sgn}(h(X))$ where

$$h(X) = \frac{1}{n_+} \sum_{i: y_i=+1} K(X_i, X) - \frac{1}{n_-} \sum_{i: y_i=-1} K(X_i, X) + b$$

where

$$b = \frac{1}{2} \left(\frac{1}{n_-^2} \sum_{y_i, y_j=-1} K(X_i, X_j) - \frac{1}{n_+^2} \sum_{y_i, y_j=+1} K(X_i, X_j) \right)$$

PTEL PPTTEL course - p.25/133

Which means finally, my nearest neighbor classifier is $\text{sgn}(h(X))$ where $h(X)$ can be obtained like this. This is $\phi(X)^T C$ plus this is $\phi(X)^T C$ minus and the constant b is this plus $C^T C$ plus $C^T C$ minus $C^T C$ minus. So, everything can be kernelised. So, what it meant is that I can for example, use ϕ to transform my features into some other high dimensional space and finding the actual distances there. So, if I have a suitable kernel then I can actually implement nearest neighbor classifier by implicitly transforming features into high dimensional space. Now, another way of looking at it is essentially the kernel captures. So, if I think $\phi(X)^T \phi(y)$ is a good matrix for distance between pattern vectors X and y then correspondingly $K(X, y)$ is a good similarity measure.

(Refer Slide Time: 10:30)



- Thus we can implement such nearest neighbour classifiers by implicitly transforming the feature space and using kernel function for the inner product in the transformed space.
- The kernel function allows us to formulate the right kind of similarity measure in the original space.

PR NPTEL course - p.26/133

(Refer Slide Time: 10:55)



- Define

$$P_+(X) = \frac{1}{n_+} \sum_{i: y_i=+1} K(X_i, X),$$

$$P_-(X) = \frac{1}{n_-} \sum_{i: y_i=-1} K(X_i, X)$$

- With a proper normalization, these are essentially non-parametric estimators for the class conditional densities – the kernel density estimates.

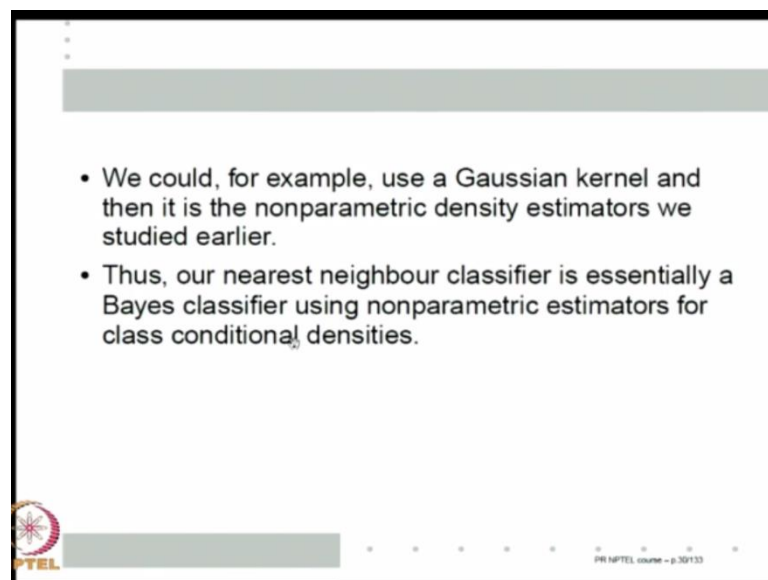
PR NPTEL course - p.26/133

So, that is we can implement nearest neighbor classifiers by implicitly transform the feature space. And the kernel function allows us to formulate the kind of similarity measure in the original space. For example, if I am using a Gaussian kernel so, which means, that I am measure similarity over some spatial extent by $((\sigma^2))$, my sigma parameter and so on. Another way of looking at is suppose I take this what do we call $\phi^T X^T C + \phi^T X^T C^-$ I call them $P_+ X$ and $P_- X$. Now, these are very familiar expressions, this is summed over all the examples of 1 class some kernel functions. So, there is some function defined here it is value at X is summed over

all examples of class 1, the kernel functions at each of the example points this X is same as this X that is the argument.

So, if for example, this was a a proper window function; this is nothing but a non parametric estimate of density. For example, if I take this to be Gaussian that means I am putting a Gaussian centered at each of the X_i 's and some sum up and sum up all their contribution, this X that is my estimated class conditional density at X . So, with a proper normalization these are nothing but non parametric estimate for class conditional densities the kernel density estimates. So, we have seen the kernel density estimates and my final classifier is if $P \text{ plus } X \text{ minus } P \text{ minus } X$ is greater than something.

(Refer Slide Time: 12:16)

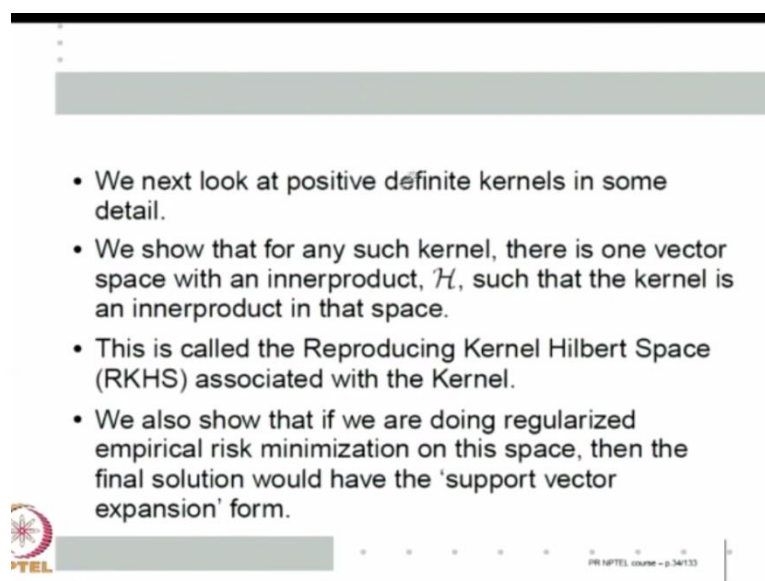


- We could, for example, use a Gaussian kernel and then it is the nonparametric density estimators we studied earlier.
- Thus, our nearest neighbour classifier is essentially a Bayes classifier using nonparametric estimators for class conditional densities.

NPTEL

NPTEL course - p30113

(Refer Slide Time: 12:52)



- We next look at positive definite kernels in some detail.
- We show that for any such kernel, there is one vector space with an innerproduct, $\mathcal{H}$, such that the kernel is an innerproduct in that space.
- This is called the Reproducing Kernel Hilbert Space (RKHS) associated with the Kernel.
- We also show that if we are doing regularized empirical risk minimization on this space, then the final solution would have the 'support vector expansion' form.

NPTEL course - p.34/133

So, so all it means is that no these are nothing but non parametric density estimators that we have seen earlier these are the kernel density estimators. And in that sense what I call the nearest neighbor classifier using kernels is like a Bayes classifier, using a non parametric density estimator for class conditional densities using the kernel a popularly normalized kernel with the density estimate. So, in this sense once again you can see kernels give us some kind of a similarity metric all. Now, let us, let us look at a few more general theoretical details of kernels. We defined positive definite kernels earlier we, we actually looked at two different kinds of definition of kernels.

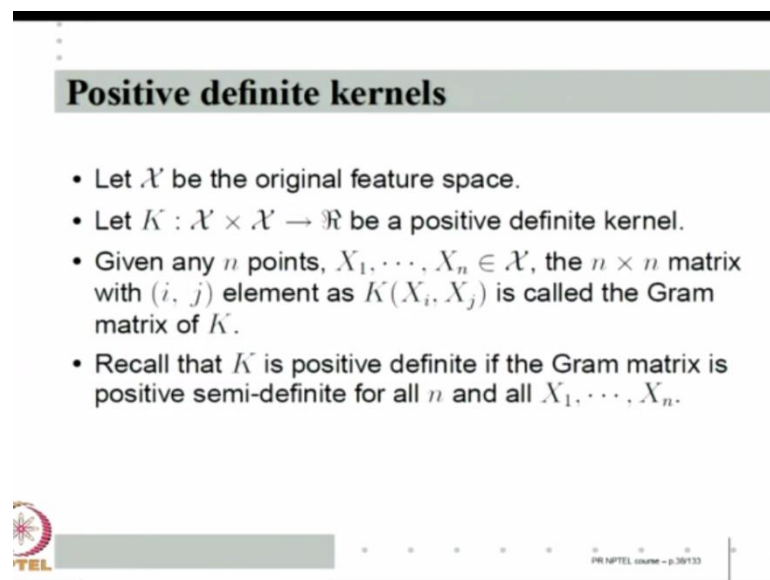
One of positive definite kernels, other kernel that satisfy Mercer's theorem, we just that anything that satisfies Mercer's theorem is such that there is always exist a ϕ and $K(X, X')$ will be $\phi(X)^T \phi(X')$. We did it say, we did not prove this theorem, but for positive definite kernels in this class, we will show that there are always exist such a space and we actually construct this space. So, we look at positive definite kernel in some detail, because these are one of the most important kernel pattern recognition today. So, we show that for any such kernel there is one vector space, we call that a $\mathcal{H}$ which has an inner product on it.

So, a vector space endow with inner product that we can construct such that the kernel is an inner product in that space. So, this particular space is called the reproducing kernel Hilbert space RKHS associated with this kernel K . And we will we will actually show

how to construct this space. And then we also show that if you are doing regularized empirical risk minimization under almost any loss function. Then the final solution would have a support vector expansion form. The, the, the rest of this lecture assumes some level of mathematical sophistication.

Specifically I am going to assume that everybody knows general vector space as norms basis orthogonal vectors, orthogonal complements of subspaces and so on and so forth. So, if any of you do not any of vector spaces may be it will be a little difficult to further part of the lecture. The, the, and the other, other hand the rest of this lecture is independent of the course. So, if you do not understand rest of the lecture can just skipped it.

(Refer Slide Time: 15:04)

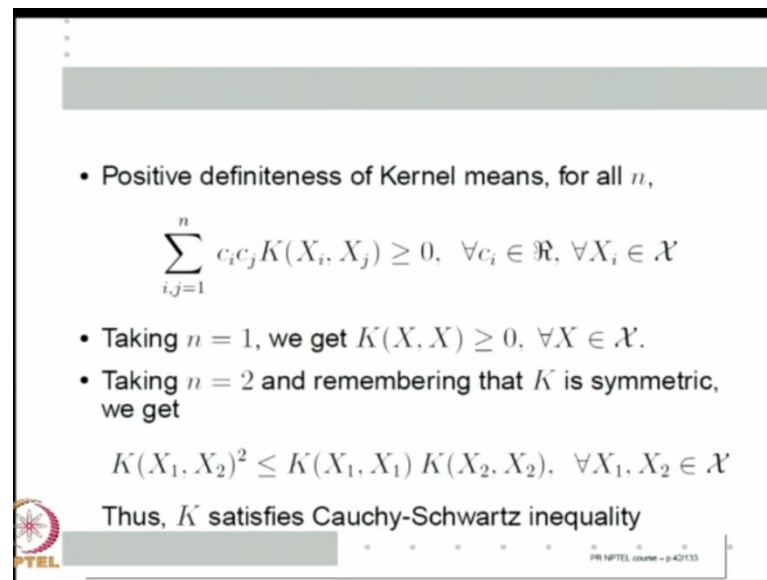


Positive definite kernels

- Let $\mathcal{X}$ be the original feature space.
- Let $K : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$ be a positive definite kernel.
- Given any n points, $X_1, \dots, X_n \in \mathcal{X}$, the $n \times n$ matrix with (i, j) element as $K(X_i, X_j)$ is called the Gram matrix of K .
- Recall that K is positive definite if the Gram matrix is positive semi-definite for all n and all $X_1, \dots, X_n$.

    PR NPTEL course - p.39133

(Refer Slide Time: 15:58)



- Positive definiteness of Kernel means, for all n ,

$$\sum_{i,j=1}^n c_i c_j K(X_i, X_j) \geq 0, \quad \forall c_i \in \mathbb{R}, \forall X_i \in \mathcal{X}$$

- Taking $n = 1$, we get $K(X, X) \geq 0, \forall X \in \mathcal{X}$.
- Taking $n = 2$ and remembering that K is symmetric, we get

$$K(X_1, X_2)^2 \leq K(X_1, X_1) K(X_2, X_2), \quad \forall X_1, X_2 \in \mathcal{X}$$

Thus, K satisfies Cauchy-Schwartz inequality

So, let us start looking at positive definite kernels as earlier, let X be the original feature space. And K is a positive definite kernel that is K is a symmetric function on X class X real value symmetric function X class X which, which satisfies some conditions just recall that given any n points. Let us say X_1 given any n and n points $X_1, X_2, \dots, X_n$ in my feature space. The n by n matrix whose i th element is $K(X_i, X_j)$ is called the Gram matrix of K . So, if I take all the X_i, X_j 's and arrange them as n by n matrix would be a symmetric matrix and recall that K is a positive definite kernel. If the gram matrix is positive semi definite that is its quadratic form is always greater than or equal to 0 for all n and all points X_1 to X_n , this is what we defined earlier.

So, specifically because the quadratic to be greater than equal to 0, a positive definite kernel if K is a positive definite kernel, then, for every n and every X_1 to X_n , the summation over i, j going from 1 to n $c_i c_j K(X_i, X_j)$. This is the quadratic form of the matrix is greater than equal to 0 for all scalars X_i . So, given any scalars even $c_1, c_2, \dots, c_n$ and any element from my feature space $X_1, X_2, \dots, X_n$ this double summation i, j going from 1 to n $c_i c_j K(X_i, X_j)$ has to be greater than or equal to 0 if K has to be a positive definite kernel. For this talk, we are going to confine ourselves to all our scalars being real even though everything we say can be extended to these scalars coming from the complex field, we will restrict ourselves to scalars being all real numbers.

So, what does this mean? Once again this is nothing but the Gram matrix n by n . Gram matrix is positive if I take n is equal to 1 of course the matrix has only one element. So, which means $K(x, x)$ should be greater than equal to 0 positive. Similarly, different if I take n is equal to 2 it will be a 2 by 2 matrix. The first row being $K(x_1, x_1) \quad K(x_1, x_2)$; second row will be $K(x_2, x_1) \quad K(x_2, x_2)$. We want that to be positive semi definite, which for example, mean that the determinant has to be greater than or equal to 0 if n is equal to 2. And I take the determinant, determinant will be K the main diagonal is $K(x_1, x_1) \quad K(x_2, x_2)$ of diagonal is $K(x_1, x_2) \quad K(x_2, x_1)$, but K is symmetric.

So, the determinant has to be positive so, $K(x_1, x_2)^2$ whole square should be less than equal to $K(x_1, x_1) \quad K(x_2, x_2)$, this is true for every pair of features x_1, x_2 . So, if K is a positive definite kernel then in particular it satisfies this given any 2 feature vectors x_1, x_2 $K(x_1, x_2)^2$ whole square is less than or equal to $K(x_1, x_1) \quad K(x_2, x_2)$. I hope all of you recognize this structure of this inequality this is nothing but Cauchy Schwartz inequality, because we think of K as a inner product if K was an indeed an inner product. Obviously, $\phi(x_1)^T \phi(x_2)$ whole square is less than or equal to $\phi(x_1)^T \phi(x_1) \quad \phi(x_2)^T \phi(x_2)$.

But we have not yet shown, that K is actually inner product, we just defined that a symmetric function with positive definite kernel if it satisfies this. But it satisfying this means that this kind of Cauchy Schwartz inequality is satisfied by the kernel. Of course, the Cauchy Schwartz inequality does not mean that there is a space in, in and a function ϕ is at $K(x, y)$ is $\phi(x)^T \phi(y)$. But certainly the function K by virtue of it being a positive definite kernel satisfies the Cauchy Schwartz inequality. This is going to be very important for us later on in the in, in studying positive definite kernels.

(Refer Slide Time: 19:15)



• Suppose $K(X, X') = \phi(X)^T \phi(X')$.

• Then K is a positive definite kernel:

$$\sum_{i,j} c_i c_j \phi(X_i)^T \phi(X_j) = \left(\sum_i c_i \phi(X_i) \right)^T \left(\sum_j c_j \phi(X_j) \right) = \left\| \sum_i c_i \phi(X_i) \right\|^2 \geq 0$$

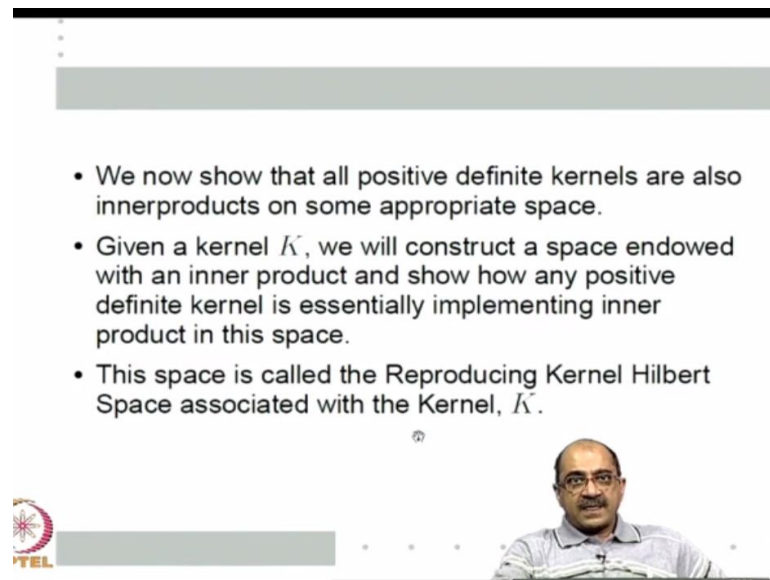
• Thus, if K satisfies Mercer theorem, then it is a positive definite kernel.

NPTEL course - p.45/133

If in fact, the kernel is obtained as a inner product in some other case that is if script X can be mapped to some other space using phi and the real space of phi has an inner product here I represent it the transpose. So, in, in, in fact, if $K(X, X') = \phi(X)^T \phi(X')$ then that such a K is certainly positive definite why, because for K to be positive definite kernel. I want to show that this $\sum_{i,j} c_i c_j K(X_i, X_j)$ should be greater than equal to 0. $K(X_i, X_j)$ is nothing but $\phi(X_i)^T \phi(X_j)$.

So, I need to show $\sum_{i,j} c_i c_j \phi(X_i)^T \phi(X_j) = \left(\sum_i c_i \phi(X_i) \right)^T \left(\sum_j c_j \phi(X_j) \right)$ is greater than or equal to 0. So, this can always be written as summation over i $c_i \phi(X_i)^T$ summation over j $c_j \phi(X_j)$ these two are same element. So, this is nothing but norm of summation over i $c_i \phi(X_i)$ norm square of that. So, this is always positive. So, if in fact, K happens to be an inner product then K will always be positive definite. But we only ask this for positive definiteness and this is of course, much easier to verify than your Mercer's theorem. So, where if K satisfies Mercer's theorem K satisfies Mercer's theorem. Then we know that there is some phi and K can be written like this. So, if K satisfies Mercer's theorem then it is a positive definite kernel.

(Refer Slide Time: 20:49)



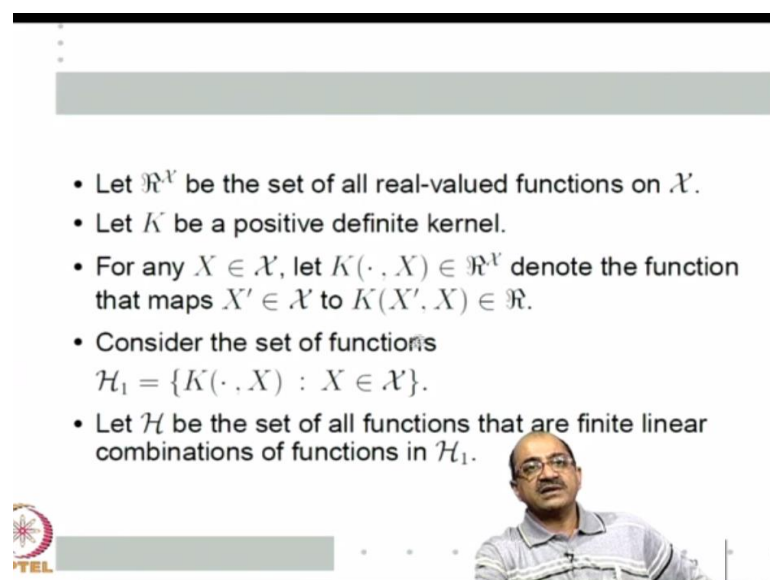
• We now show that all positive definite kernels are also inner products on some appropriate space.

• Given a kernel K , we will construct a space endowed with an inner product and show how any positive definite kernel is essentially implementing inner product in this space.

• This space is called the Reproducing Kernel Hilbert Space associated with the Kernel, K .

Now, we show that all positive definite kernels are also inner products on some appropriate space as I said we show this by saying given a kernel K we will construct a space endowed with the inner product. And show how positive definite kernel is essentially implementing inner product in this space, as I said this space is called the reproducing kernel Hilbert space.

(Refer Slide Time: 21:13)



• Let $\mathbb{R}^{\mathcal{X}}$ be the set of all real-valued functions on $\mathcal{X}$.

• Let K be a positive definite kernel.

• For any $X \in \mathcal{X}$, let $K(\cdot, X) \in \mathbb{R}^{\mathcal{X}}$ denote the function that maps $X' \in \mathcal{X}$ to $K(X', X) \in \mathbb{R}$.

• Consider the set of functions $\mathcal{H}_1 = \{K(\cdot, X) : X \in \mathcal{X}\}$.

• Let $\mathcal{H}$ be the set of all functions that are finite linear combinations of functions in $\mathcal{H}_1$.

So, to start on this let $\mathbb{R}^{\mathcal{X}}$ be the set of all real valued functions on my feature space that is if I have any function g that maps my feature space

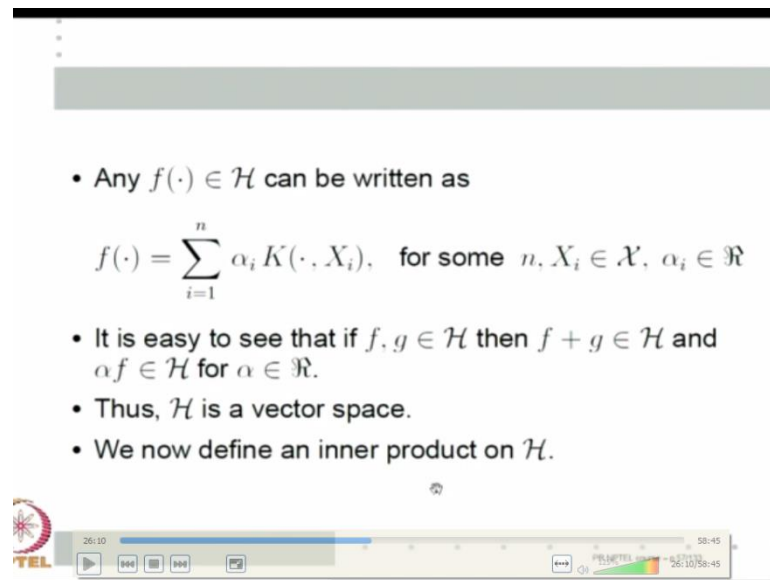
to real numbers then that g will be an element of this set. So, this set consists of all possible real valued functions on my feature space X . Now, K be the given positive definite kernel. So, given any element in my script X any feature vector X , let us represent by $K \cdot X$. This is general representation for functions the dot is where the functions argument is.

So, essentially this is a function because you put anything in place of this dot you get a real number. And think that you can put in place of this dot or elements of script X . So, $K \cdot X$ is nothing but a function that maps script X to real numbers. So, this belongs to $\mathbb{R}^X$. So, $K \cdot X$ is some real valued function. Let us say this denotes the function that maps any X' to $K(X', X)$. So, that is almost so simply dot is the notation for the dummy argument of the function.

So, $K \cdot X$ is the real valued function on X whose value at an argument X' is $K(X', X)$. Now, consider the set of functions H_1 which consist of all such $K \cdot X$'s for $X \in X$ for every X in my feature space actually we can think of $K \cdot X$ at the kernel functions centered at X . It is actually this norm is actually function its value at any argument value is the value of the kernel function X' comma X so, this we will call this a kernel centered X .

So, H_1 is the set of all the set of kernels centered at all possible feature vectors all possible elements of X . Let us script H be the set of function that are finite linear combinations of functions in H_1 , we will you take H_1 and make finite linear combinations. Let us say $\alpha_1 K \cdot X_1 + \alpha_2 K \cdot X_2$ like that so all possible finite linear combinations of functions H_1 . If you take all possible finite linear combinations of functions of H_1 , let us call that set of functions as script H . So, script H is the set of all functions that are finite linear combinations of functions in H_1 .

(Refer Slide Time: 23:53)



- Any $f(\cdot) \in \mathcal{H}$ can be written as

$$f(\cdot) = \sum_{i=1}^n \alpha_i K(\cdot, X_i), \text{ for some } n, X_i \in \mathcal{X}, \alpha_i \in \mathbb{R}$$

- It is easy to see that if $f, g \in \mathcal{H}$ then $f + g \in \mathcal{H}$ and $\alpha f \in \mathcal{H}$ for $\alpha \in \mathbb{R}$.
- Thus, $\mathcal{H}$ is a vector space.
- We now define an inner product on $\mathcal{H}$.

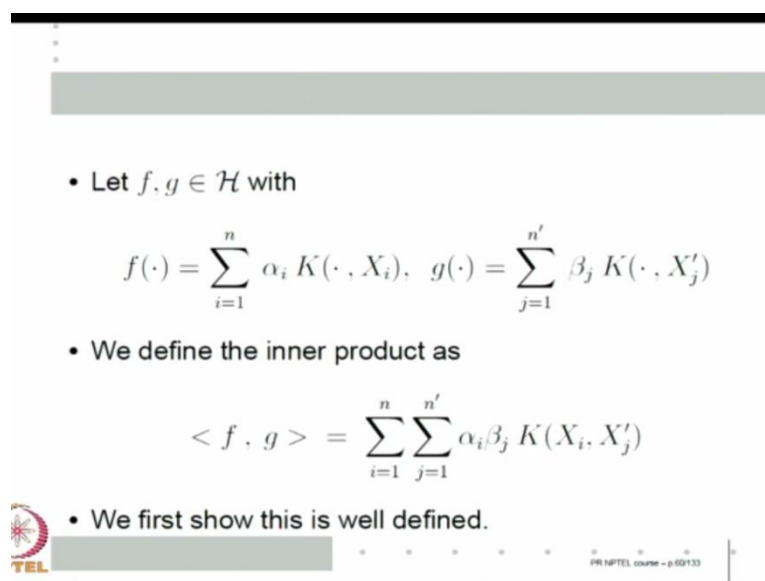
So, any f in $\mathcal{H}$ would be a linear combination of functions from $\mathcal{H}_1$ and the only functions in $\mathcal{H}_1$ are $K \cdot X$ type where X is an element of script $\mathcal{X}$. So, any f in my script $\mathcal{H}$ can be written as $\alpha_i K \cdot X_i$ for some n . So, it is a finite linear combination so, we do not know how many term they will be, but there will be some in $(())$ such that there will be n terms in the linear combination.

So, f will be i is equal to 1 to n for some n of $\alpha_i K \cdot X_i$ for some real numbers α_i and some X_i . So, every function in this script $\mathcal{H}$ can be written like this. For some X_1 to X_n elements of script $\mathcal{X}$ some α_1 to α_n real numbers. And some n f can be written as i is equal to 1 to f in $\alpha_i K X$ which means if I have 2 functions f and g in $\mathcal{H}$ is very easy to say f plus g is also in $\mathcal{H}$. Because f plus g would once again be some finite linear combination of $K \cdot X$'s. And similarly, α times f will also be in $\mathcal{H}$ for any real number α which means $\mathcal{H}$ is a vector space over the field of real s , because the rest of the vector space axioms are easy to verify, essentially we have a vector addition and scalar multiplication while define.

So, this, this set $\mathcal{H}$ know is a vector, vector space under the usual addition of functions and scalar multiplication of functions a vector space over the field of real numbers. As I said everything as a in this lecture can be extended to the field of complex numbers with minor modifications, but for simplicity we will stick to all scalars being layer. So, what we showed is that if we take finite linear combinations of kernels entire that any points in

X , if we call that as the space H that H is the is a vector space under the usual function addition and scalar multiplication. So, now what we are going to do is that we define an inner product on a space H inner product. And then show that essentially there can be a ϕ that maps script X to this H . So, that $K(X, X')$ will be $\phi(X) \phi(X')$ inner product that we are going to define.

(Refer Slide Time: 26:21)



• Let $f, g \in \mathcal{H}$ with

$$f(\cdot) = \sum_{i=1}^n \alpha_i K(\cdot, X_i), \quad g(\cdot) = \sum_{j=1}^{n'} \beta_j K(\cdot, X'_j)$$

• We define the inner product as

$$\langle f, g \rangle = \sum_{i=1}^n \sum_{j=1}^{n'} \alpha_i \beta_j K(X_i, X'_j)$$

• We first show this is well defined.

TEL PR NPTEL course - p-00133

So, first I have to define an inner product for any two elements of H . So, let us take any 2 elements of H and g , every element of H is a finite linear combination of kernel functions entered at X_i . So, f is some summation i is equal to 1 to n $\alpha_i K(\cdot, X_i)$ for some real α_i some elements X_i . So, g is $\beta_j K(\cdot, X'_j)$ so, $X_1, X_2, \dots, X_n, X'_1, X'_2, \dots, X'_{n'}$. These are all arbitrary I took n and n' also different, because the 2 functions may have different number of terms in there linear when they represented as linear combinations of K 's.

So, given 2 functions f and j like this, we will define inner product I am going to represent inner product like this in general the standard notation for inner product between the 2 angular brackets f comma g within the angular bracket. So, this is the inner product of f and g if f is this and g is this. We will represent the inner product of f comma g as i is equal to 1 to n j is equal to 1 to n' summation $\alpha_i \beta_j K(X_i, X'_j)$.

So, if f is this and g is this then the inner product is defined to be this. Of course, we have to show that this is an inner product it has satisfy all the properties of inner product before we can go for there to show that this is an inner product. We have to first show that this is well defined, what do you mean by well defined when I say elements of H R finite linear combinations of elements of H 1 which are nothing but kernel functions centered at point x .

So, this function f given different arguments say f of X is $\alpha_i K(x, X_i)$. Now, a given a function a of H there may be more than one way of writing it as a linear combination of some kernel function when we say we all possible linear combinations there is no guarantee that 2 different linear combinations do not give you the same function which on other words a given function may not be representable uniquely as a linear combination which means these α_i 's are not $(())$.

The function of course, is $(())$ because function is an element in H , but the same function may be representable with some α_i primes of and some $K(y, y_i)$'s who knows. So, essentially what we have to show is that it is not dependent on this α_i 's and β_j 's, because they can be different kinds of representations. So, the, the inner product should not be dependent on the expansion coefficients. So, first you have to show that it does not depend on the expansion coefficients.

(Refer Slide Time: 29:21)



- We note that

$$\langle f, g \rangle = \sum_{i=1}^n \alpha_i \sum_{j=1}^{n'} \beta_j K(X_i, X'_j) = \sum_{i=1}^n \alpha_i g(X_i)$$
- Similarly we have

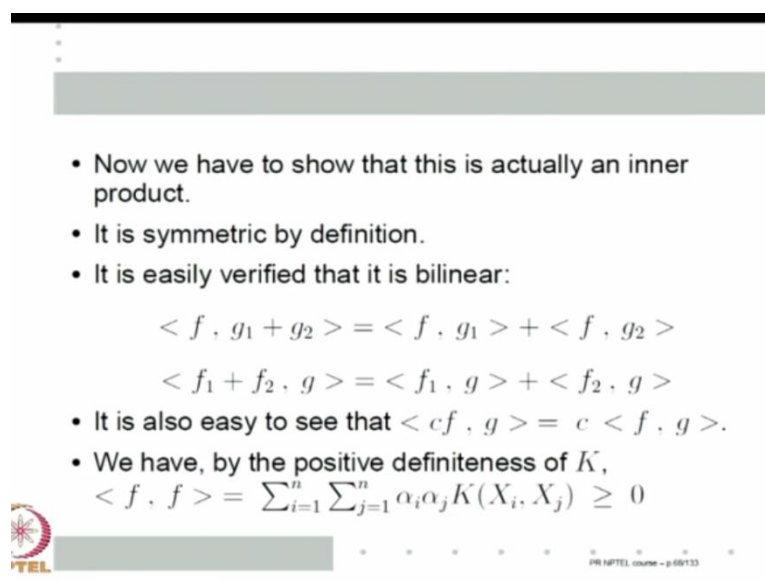
$$\langle f, g \rangle = \sum_{j=1}^{n'} \beta_j \sum_{i=1}^n \alpha_i K(X_i, X'_j) = \sum_{j=1}^{n'} \beta_j f(X'_j)$$
- Thus our inner product does not depend on the α_i and β_j and hence is well defined.

PR NPTEL course - p-63133

The other product as we defined I can take alpha i outside, what I have is j is equal to 1 to n prime beta j K X i X j prime g dot is j is equal to 1 to n prime beta is a K dot X j prime. So, g of X i will be beta j X i X j prime. So, if I take alpha i out, what I have is beta j K X i X j prime which is nothing but g f X i. So, I can write the inner product as i is equal to 1 to n alpha i g of X I, which means while it depends on the function g it does not depend on either beta j or X j prime i g can be written in different ways. It does not matter as long is the same function it only depends on the values of g at X i. So, i g has a different expansion it does not matter. Similarly, by taking beta j by interchanging the 2 summations and taking beta j out I can write as j is equal to 1 to n prime beta j i is equal to 1 to n alpha i K i X i X j prime, what is this? f dot is this.

So, f of X j prime will be i is equal to 1 to n alpha i K i of X j prime X i and your K is symmetric. So, whichever way I write it as does not matter. So, this entire summation nothing but alpha of X j i mean f of X j prime. So, I can also write f the inner product between f and g as summation j is equal to 1 to n b prime beta i j f of X j prime. So, it depends on the value of the function f at X j prime where does not does not specifically depend on the specific vectors f i or alpha I, which shows that the inner product does not depend on the specific expansion coefficients and hence it is well defined. So, this f well defined inner product. So, next you have to ask is this indeed an inner product I just said is an inner product, we have to show that this is an inner product.

(Refer Slide Time: 30:19)



- Now we have to show that this is actually an inner product.
- It is symmetric by definition.
- It is easily verified that it is bilinear:

$$\langle f, g_1 + g_2 \rangle = \langle f, g_1 \rangle + \langle f, g_2 \rangle$$

$$\langle f_1 + f_2, g \rangle = \langle f_1, g \rangle + \langle f_2, g \rangle$$
- It is also easy to see that $\langle cf, g \rangle = c \langle f, g \rangle$.
- We have, by the positive definiteness of K ,

$$\langle f, f \rangle = \sum_{i=1}^n \sum_{j=1}^n \alpha_i \alpha_j K(X_i, X_j) \geq 0$$

So, what is an inner product satisfies? An inner product vector space h takes space of vectors and gives you a real number. So, it is a function h class $h^2 l$, what does this function satisfy? It has to be symmetric, it has to be bilinear inner product of any, any X any element with itself is always greater than or equal to 0. An inner product is 0 if and only if the element is 0 inner product of some element with itself is 0 if and only if element is 0 these are the properties of inner product.

Now, our inner product is symmetric by definition, because this is how we define this. So, you given this and this there is an dependant on you know which is f and which is, because these $\alpha_i \beta_j K(X_i, X_j)$ and K is symmetric. So, in a product of f and g is same as product of g and f . I would like to remind you once again that will considering field of real numbers, this our vector space is vector space over reals that is why we are only looking at symmetry otherwise we have to worry about complex conjugateness. But we are not allowing complex scalars now for simplicity of exposition I am just taking everything to be real.

Similarly, it has to be bilinear that is inner product of $f, g_1 + g_2$ should be inner product of f, g_1 plus f, g_2 . Similarly, the other way which is also very straight forward, because of the definition of the inner product it is very easy to verify that it is bilinear. So, if I put $f_1 + f_2$ here. So, I will get some $\alpha_i K \cdot X_i$ plus

some γ_j $\sum_{j=1}^p \gamma_j \langle f, f_j \rangle = 0$ so, accordingly they come in the summation.

So, $\langle f, f \rangle$ will be $\langle f, f \rangle + \langle f, f \rangle + \dots$ and so on. So, by definition, it is symmetric and it is bilinear. And similarly, you can show that if one of the, are given multiply by a constant, the inner product gets multiplied by that constant. Next, what you have to show? You have to show that inner product of f comma f is greater than or equal to 0 while end product of f comma f if f is $\sum_{i=1}^p \alpha_i \phi(x_i)$ inner product of f comma f is $\sum_{i,j=1}^p \alpha_i \alpha_j K(x_i, x_j)$ So, this is nothing but the quadratic form of the gram matrix and because K is positive definite kernel, this is greater than or equal to 0. So, we also shown that inner product of f comma f is greater than or equal to 0.

(Refer Slide Time: 34:17)

- Finally, we have to show $\langle f, f \rangle = 0 \Rightarrow f = 0$.
- Let $f_1, \dots, f_p \in \mathcal{H}$ and let $\gamma_1, \dots, \gamma_p \in \mathbb{R}$.
Let $g_1 = \sum_{i=1}^p \gamma_i f_i \in \mathcal{H}$.
- Now we get

$$\begin{aligned} \sum_{i,j=1}^p \gamma_i \gamma_j \langle f_i, f_j \rangle &= \left\langle \sum_{i=1}^p \gamma_i f_i, \sum_{j=1}^p \gamma_j f_j \right\rangle \\ &= \langle g_1, g_1 \rangle \\ &\geq 0 \end{aligned}$$

So, the only thing that remains to be shown now is that if the inner product is 0 f is 0. Mind you, if I had assumed, if I defined positive definite kernels as a function K such that the Gram matrix is positive definite rather than positive semi definite then I am done. If I am asking for this to be positive definite then except when $\alpha_i \alpha_j$ is 0 which means function is 0 this quadratic form has to be strictly greater than 0. But if I want positive definite functions only there will be difficult for, for us to get kernels. Because we want kernels should represent similarity in general as it turns out in many

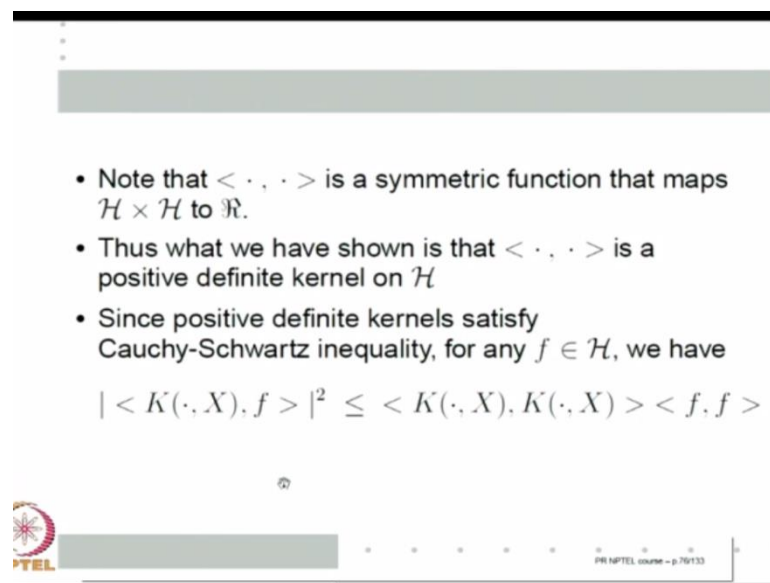
applications. While it is easy enough to show Gram matrix is positive semi-definite, it is much more difficult to show that Gram matrix is positive definite.

So, we do not want to assume positive definiteness, the kernel, because we do not know if we need it. I just want you to understand that because we assume only positive semi-definiteness of the Gram matrix, I only get $\langle f, f \rangle$ inner product of f and f is greater than or equal to 0 and it is not if I still have to separately show that if inner product is 0, f is equal to 0. So, let us go ahead and choose this separately. Show this, let us take any, any some P functions from H and some P scalars γ_1 to γ_P and let g_1 is $\sum \gamma_i f_i$, because H is the vector space g will also be in H .

Now, if I look at $\sum_{i,j} \gamma_i \gamma_j \langle f_i, f_j \rangle$ inner product of f_i and f_j summed over i, j is equal to 1 to n . By the bilinearity of inner product, it is same as inner product of $\sum_{i=1}^P \gamma_i f_i$ and $\sum_{j=1}^P \gamma_j f_j$. Because of the bilinearity of the inner product, we have seen that the inner product of $f_1 + f_2$ and g is inner product of f_1 and g plus f_2 and g . And similarly, the other way and hence this summation double summation can be written as inner product between these two.

Now, this is nothing but what we call the function g_1 . So, this is nothing but the inner product between g_1 and g_1 , because these two are the same summation i and j are dummy variables after all. And this we already show n to be greater than or equal to 0. So, what it means is given any P functions from H $f_1, f_2, \dots, f_P$ and any P scalars $\gamma_1, \gamma_2, \dots, \gamma_P$ $\sum_{i,j} \gamma_i \gamma_j \langle f_i, f_j \rangle$ summed over i, j equal to 1 to P is greater than or equal to 0. What is this mean? This is exactly what we wanted from a positive definite kernel, earlier positive definite kernel meant as we seen $C_{i,j}$ is C is a $K(X \times X)$ is greater than or equal to 0. So, I can think of the inner product itself as kernel on H our original K the the kernel I am considering K is a kernel on script X or feature space, but now it looks like $\langle f_i, f_j \rangle$ can.

(Refer Slide Time: 37:29)



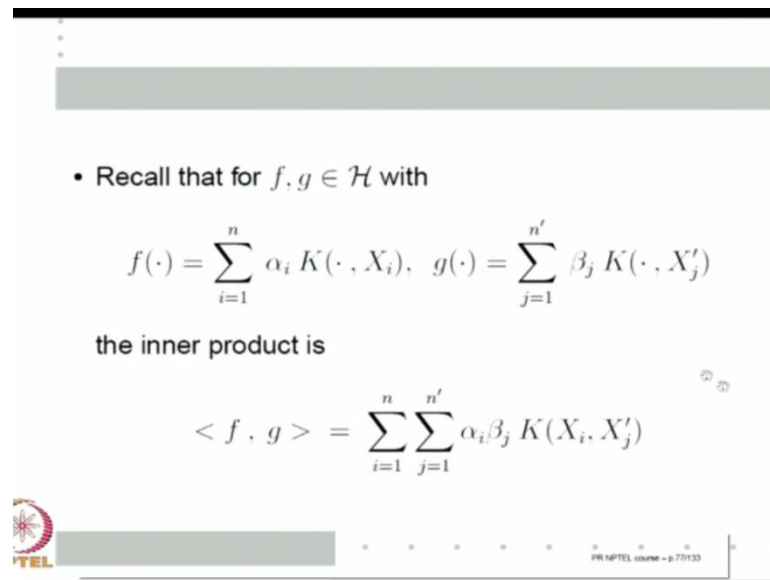
- Note that $\langle \cdot, \cdot \rangle$ is a symmetric function that maps $\mathcal{H} \times \mathcal{H}$ to $\mathbb{R}$.
- Thus what we have shown is that $\langle \cdot, \cdot \rangle$ is a positive definite kernel on $\mathcal{H}$
- Since positive definite kernels satisfy Cauchy-Schwartz inequality, for any $f \in \mathcal{H}$, we have

$$|\langle K(\cdot, X), f \rangle|^2 \leq \langle K(\cdot, X), K(\cdot, X) \rangle \langle f, f \rangle$$

What is $\langle f, f \rangle$? This inner product is a symmetric function that maps $\mathcal{H} \times \mathcal{H}$ to $\mathbb{R}$, the inner product is given any 2 elements of $\mathcal{H}$, it maps it to real number and the function is symmetric and any such symmetric function by virtue that it satisfies this mean that it is a kernel on $\mathcal{H}$. So, that is what we have shown is that this function; this function linear product function which maps $\mathcal{H} \times \mathcal{H}$ to $\mathbb{R}$ is in fact, a positive definite kernel on $\mathcal{H}$. Because it is a positive definite kernel on $\mathcal{H}$, we know that positive definite kernel satisfy Cauchy Schwartz inequality. What does, what did the Cauchy Schwartz inequality mean? The kernel value at X_1 comma X_2 whole square is less than or equal to kernel valued function X_1 comma X_1 into kernel value at X_2 comma X_2 .

So, if I take any 2 element from $\mathcal{H}$ their inner product square. So, this is element 1; this is element square is less than or equal to inner product of element 1 comma element 1 into inner product of element 2. So, in particular I know that if I choose f as 1 element and $K \cdot X$ as another then because of the because this satisfies Cauchy Schwartz inequality. We know have inner product of $K \cdot X$ comma f whole square is less than or equal to inner product of $K \cdot X$ comma $K \cdot X$ and inner product of f comma f [s]. So, let us calculate each of these inner products from our definition of inner product.

(Refer Slide Time: 39:04)



• Recall that for $f, g \in \mathcal{H}$ with

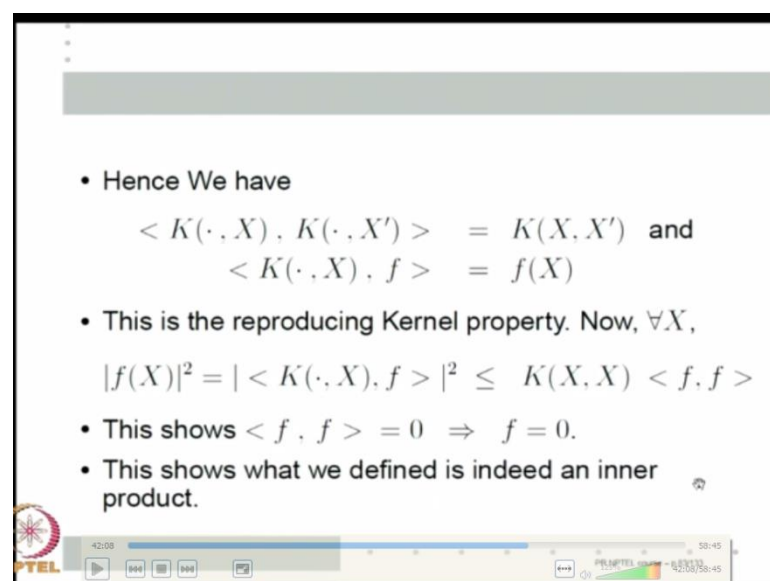
$$f(\cdot) = \sum_{i=1}^n \alpha_i K(\cdot, X_i), \quad g(\cdot) = \sum_{j=1}^{n'} \beta_j K(\cdot, X'_j)$$

the inner product is

$$\langle f, g \rangle = \sum_{i=1}^n \sum_{j=1}^{n'} \alpha_i \beta_j K(X_i, X'_j)$$

NPTEL course - p.7/133

(Refer Slide Time: 39:12)



• Hence We have

$$\begin{aligned} \langle K(\cdot, X), K(\cdot, X') \rangle &= K(X, X') \quad \text{and} \\ \langle K(\cdot, X), f \rangle &= f(X) \end{aligned}$$

• This is the reproducing Kernel property. Now, $\forall X$,

$$|f(X)|^2 = |\langle K(\cdot, X), f \rangle|^2 \leq K(X, X) \langle f, f \rangle$$

• This shows $\langle f, f \rangle = 0 \Rightarrow f = 0$.

• This shows what we defined is indeed an inner product.

NPTEL course - p.8/133

Recall that f is equal to $\alpha_i K(\cdot, X_i)$ g is equal to $\beta_j K(\cdot, X'_j)$ then inner product of f comma g is this. So, what it means? For example, if I want $K(\cdot, X)$ comma $K(\cdot, X')$ then both of these sums are only one element sums both these expansions are one element expansion. So, my f comma g will be nothing but $K(X, X')$ comma $K(X, X')$ so, $K(X, X)$ comma $K(X, X')$ is $K(X, X')$.

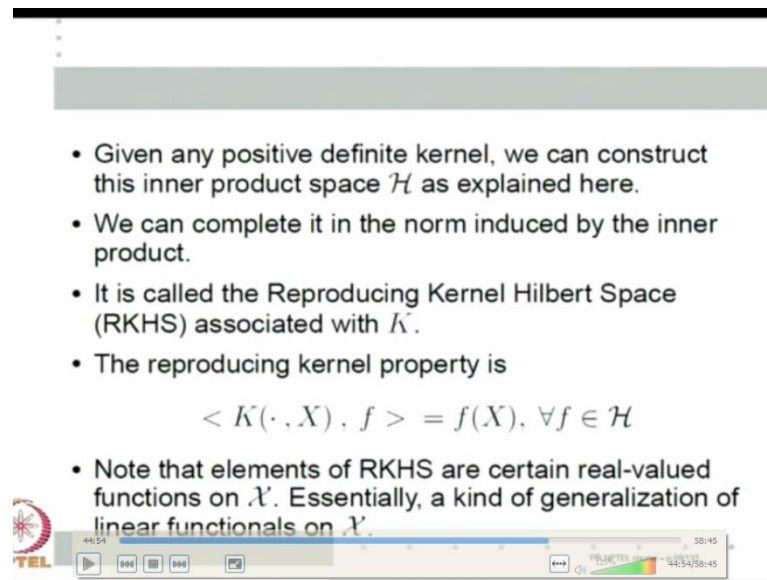
Similarly, if f is this, but g is simply $K(X, X')$ 1 element then what will be the inner product? $\alpha_i K(X, X')$ comma X_i . So, if g is simply $K(X, X')$ then the inner

product will be $\sum_{i=1}^n \alpha_i K(x, x_i)$ which is nothing but $f(x)$. I hope this is clear, $K(x, x)$ is $f(x)$. So, very, very important property, just for our definition of inner product shows this let us understand this again. Let us say f is $\sum_{i=1}^n \alpha_i K(x, x_i)$. And let us say g is simply a single function $K(x, x)$ then what will be this double summation be? This implies single summation i is equal to 1 to n $\alpha_i K(x, x_i)$.

Now, $\sum_{i=1}^n \alpha_i K(x, x_i)$ is nothing but $f(x)$. So, definition of inner product is such that $K(x, x)$ is $f(x)$. Now this is called the reproducing property of kernel. If I take a kernel centered at x and take its inner product with f what I get is the value of f at x . If I take a kernel centered at x and take its inner product with respect to f then what I get is the value of f at that x this is called the reproducing kernel property, because the reproducing property what we have is the following.

Now, $\|f\|_H^2$ which is nothing but the inner product of $K(x, x)$ with f whole square by Cauchy Schwartz inequality. This is inner product of $K(x, x)$ with $K(x, x)$ which is nothing but $K(x, x)$ into inner product of f with f . So, what is that mean? It means if $\|f\|_H^2$ inner product of f with f is 0 then for every x $f(x)^2$ whole square is 0 that means for every x $f(x)$ is equal to 0 that means f is equal to 0. So, we have shown that our inner product is such that if inner product of f with f is equal to 0 then f is identically equal to 0 using this so called reproducing kernel property. So, this shows that what we have defined indeed a proper inner product.

(Refer Slide Time: 42:11)



- Given any positive definite kernel, we can construct this inner product space $\mathcal{H}$ as explained here.
- We can complete it in the norm induced by the inner product.
- It is called the Reproducing Kernel Hilbert Space (RKHS) associated with K .
- The reproducing kernel property is
$$\langle K(\cdot, X), f \rangle = f(X), \forall f \in \mathcal{H}$$
- Note that elements of RKHS are certain real-valued functions on $\mathcal{X}$. Essentially, a kind of generalization of linear functionals on $\mathcal{X}$.

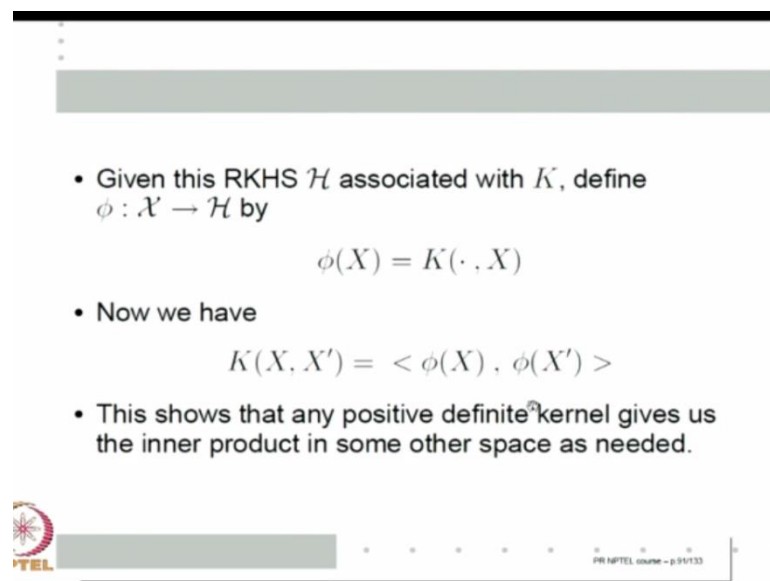
So, given a any positive definite kernel, we can construct the inner product space H as explained give me a f kernel K and have the feature space X . I will take kernel centered at each point on X which we called $K \cdot X$, we take the set of $K \cdot X$ for every i for X in my feature space. Then I make all possible finite linear combinations of these, these kernels $\alpha_i K \cdot X_i$ summed over i . And I can consider the set of all functions that is my space h , this space happens to be a vector space on which I can put an inner product. And you know once I put this inner product technically I want this space to be complete, what does what do I mean by complete? That a an inner product; obviously, gives me metric.

If I have an inner product between 2 elements f comma g then I can have a norm for this as inner product norm f square is inner product f comma f . And I will define distance between f and g as norm of f minus g . Under such a distance metric if any sequence is Cauchy that is a a sequence of that as $n \rightarrow \infty$ tends to infinity. The n th and m th element of the sequence comes close to each other in the distance metric. Then the sequence should also have a limit that is what is meant by complete. This space is not complete, we can always complete, it this just a technical diognisation if you do not understand it does not matter essentially, H is a vector space is the inner product. And we simply assume that under this inner product all convergent sequences have their limits in the space that is what is called completing the H .

This space is called a reproducing Kernel Hilbert space. Essentially Hilbert space is a vector space on which you define an inner product and in the metric induced by the inner product this space is complete. So, that's why it is called a Hilbert space, so this H , we constructed is called a reproducing kernel Hilbert space and the reproducing kernel properties. So, given a positive definite kernel K we constructed the H and that H there is an inner product and the inner product is such that inner product of $K \cdot X$ comma f will give me the value of f at X .

This is called the reproducing property of the kernel, this is true for every f in the space. Essentially the elements of this RKHS are real valued functions on H not all real valued functions of certain real valued functions at $f(X)$, which can essentially be written as linear combinations of kernel centered at X . So, this is a kind of generalization of linear functional on X . So, that is another way of looking at RKHS, which have essentially the special reproducing kernel property.

(Refer Slide Time: 45:14)



• Given this RKHS $\mathcal{H}$ associated with K , define $\phi : \mathcal{X} \rightarrow \mathcal{H}$ by

$$\phi(X) = K(\cdot, X)$$

• Now we have

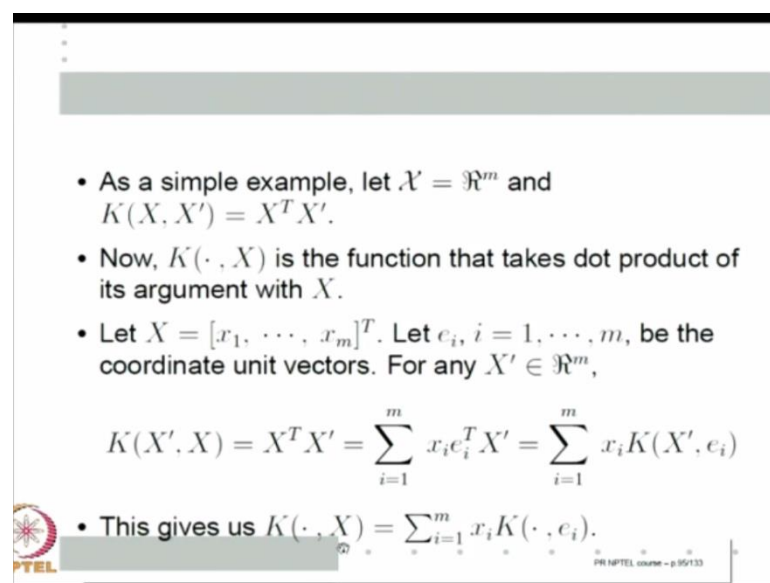
$$K(X, X') = \langle \phi(X), \phi(X') \rangle$$

• This shows that any positive definite kernel gives us the inner product in some other space as needed.

So, given this RKHS $\mathcal{H}$ be associated with K , now we can define ϕ that maps X to $\mathcal{H}$ that is it maps every element of $\mathcal{X}$ to every element of $\mathcal{H}$ namely it maps the element X in my feature space to the element $K \cdot X$ in $\mathcal{H}$ it maps X to the kernel functions centered at X . So, ϕ maps X to $K \cdot X$, if I take this ϕ then the inner product of $\phi(X)$ comma $\phi(X')$ is nothing but inner product of $K \cdot X$ comma $K \cdot X'$ which is nothing but $K(X, X')$.

So, $K(X, X')$ is nothing but inner product between $\phi(X)$ and $\phi(X')$ in H . So, H is the space you are looking for given a positive definite kernel K . Now, I have constructed a specific space H and I showed you a function ϕ that maps X to H such that $K(X, X')$ is nothing but inner product between $\phi(X)$ and $\phi(X')$. So, this is; this means that any positive definite kernel gives us an inner product in some other space as needed. Now, let us just get a little more idea of what this RKHS is a very familiar setting if you look at it then we make at some more idea about this RKHS is.

(Refer Slide Time: 46:41)



- As a simple example, let $\mathcal{X} = \mathbb{R}^m$ and $K(X, X') = X^T X'$.
- Now, $K(\cdot, X)$ is the function that takes dot product of its argument with X .
- Let $X = [x_1, \dots, x_m]^T$. Let $e_i, i = 1, \dots, m$, be the coordinate unit vectors. For any $X' \in \mathbb{R}^m$,

$$K(X', X) = X^T X' = \sum_{i=1}^m x_i e_i^T X' = \sum_{i=1}^m x_i K(X', e_i)$$
- This gives us $K(\cdot, X) = \sum_{i=1}^m x_i K(\cdot, e_i)$.

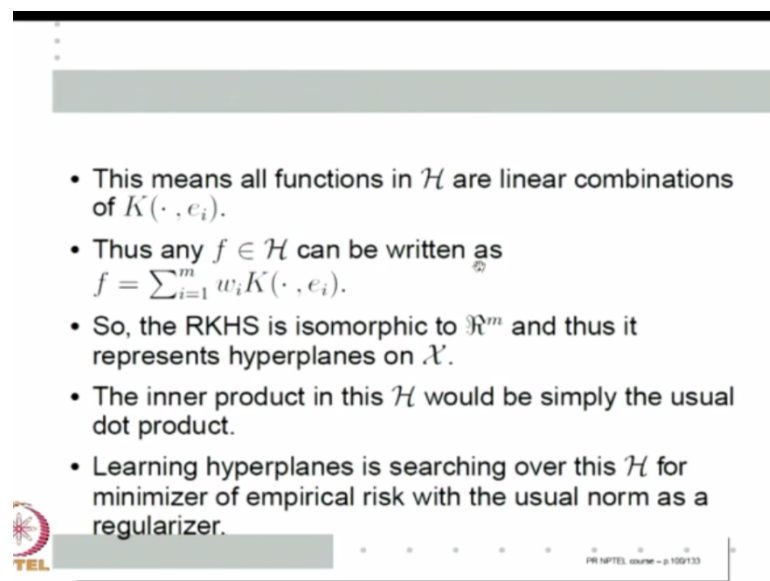
So, we look at a very simple example, let us assume X is $\mathbb{R}^m$, and let us take this simplest linear kernel $K(X, X')$ is transpose of X' . So, what will be $K(X, X')$ now $K(X, X')$ is a function that takes dot product of its argument with X for each X . I have function $K(X, \cdot)$ which is essentially taking inner dot product with X that is the name of the function.

So, you give X' as the argument of the function outcomes $X'^T X$. So, essentially $K(X, \cdot)$ is the function that takes dot product of its argument with X . So, let us take any, any vector X in $\mathbb{R}^m$ with components X_1 to X_m . And let us say e_i are the coordinate unit vectors that is e_1 is $[1, 0, 0, \dots]^T$, e_2 is $[0, 1, 0, \dots]^T$ and so on. If you give me any X' belonging to $\mathbb{R}^m$. Now, $K(X', X)$ is by definition $X'^T X$.

prime $X^T X$ prime is nothing but summation of over i of i th element of X and i -th element of X prime.

So, $\sum_i X_i$ and I can get i th element of X prime as $e_i^T X$ prime, because e_i is the coordinate vector. Now, by definition of kernel $e_i^T X$ prime is nothing but $K(X, e_i)$ or $e_i^T X$ prime does not matter whichever way we write. So, if summation $X_i K(X, e_i)$, so $K(X, X)$ is nothing but summation $i=1$ to m $X_i K(X, e_i)$, what is that mean? $K(X, X)$ is equal to $\sum_{i=1}^m X_i K(X, e_i)$. So, for any arbitrary X if you want a kernel function centered at X that itself can be written as a linear combination of these m specific kernel functions $K(\cdot, e_i)$. That means every $K(X, \cdot)$ can be written as a linear combination of $K(\cdot, e_i)$ which means any linear combinations of $K(X, \cdot)$ for all X for different X 's can once again be rewritten as a linear combination of $K(\cdot, e_i)$.

(Refer Slide Time: 48:40)



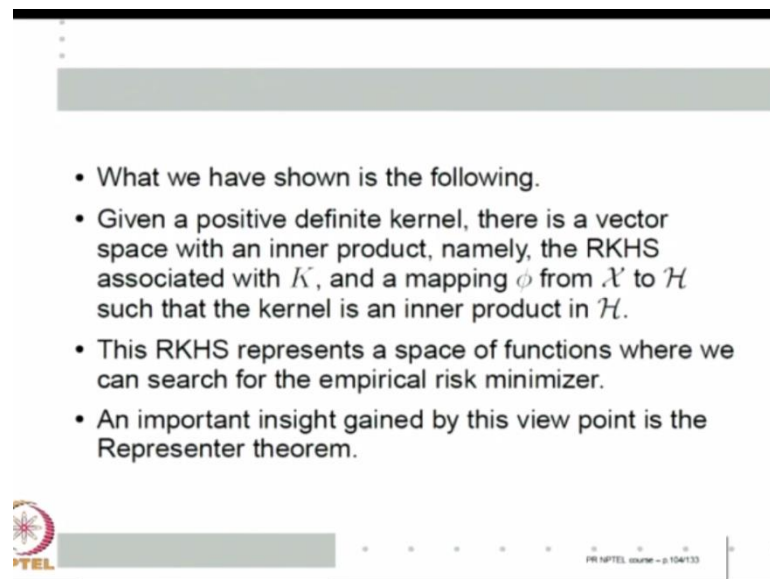
- This means all functions in $\mathcal{H}$ are linear combinations of $K(\cdot, e_i)$.
- Thus any $f \in \mathcal{H}$ can be written as $f = \sum_{i=1}^m w_i K(\cdot, e_i)$.
- So, the RKHS is isomorphic to $\mathbb{R}^m$ and thus it represents hyperplanes on $\mathcal{X}$.
- The inner product in this $\mathcal{H}$ would be simply the usual dot product.
- Learning hyperplanes is searching over this $\mathcal{H}$ for minimizer of empirical risk with the usual norm as a regularizer.

PRIPTEES course - p.100/133

Which means all functions in $\mathcal{H}$ are simply linear combinations of $K(\cdot, e_i)$. And for i is equal to 1 to m there only m such kernel functions. And every other function $\mathcal{H}$ can be written as linear combination of these, which means any f can be written as $\sum_{i=1}^m w_i K(\cdot, e_i)$. So, any f can be represented now by m w_i 's which means it can be represented by a vector in $\mathbb{R}^m$, which means every f can be uniquely defined by a m component vector which means my RKHS itself is isomorphic to $\mathbb{R}^m$. And that means

every element in H can be associated with a hyperplane on H . So, my for my linear kernel the reproducing kernel Hilbert space is nothing but the set of all the hyperplanes on X . That is what a linear classifier gives me essentially if my if I take the linear kernel extra X transpose X then what I am doing is I am searching over hyperplane over x . So, I have actually searching over the RKHS so, the inner product H now is simply the usual dot product in $\mathbb{R}^m$ and learning hyperplanes is nothing but searching over this H for a minimize of empirical risk.

(Refer Slide Time: 50:02)



- What we have shown is the following.
- Given a positive definite kernel, there is a vector space with an inner product, namely, the RKHS associated with K , and a mapping ϕ from $\mathcal{X}$ to $\mathcal{H}$ such that the kernel is an inner product in $\mathcal{H}$.
- This RKHS represents a space of functions where we can search for the empirical risk minimizer.
- An important insight gained by this view point is the Representer theorem.

PRIPITEEL course - p.104/133

(Refer Slide Time: 50:32)

Representer Theorem

- Let K be a positive definite Kernel and let $\mathcal{H}$ be the RKHS associated with it.
- Let $\{(X_i, y_i), i = 1, \dots, n\}$ be the training set.
- For any function f , the empirical risk, under any loss function can be represented as a function

$$\hat{R}_n(f) = C((X_i, y_i, f(X_i)), i = 1, \dots, n)$$

- We search over $\mathcal{H}$ for a minimizer of empirical risk.
- Let $\|f\|^2 = \langle f, f \rangle$ be the norm under our inner product.

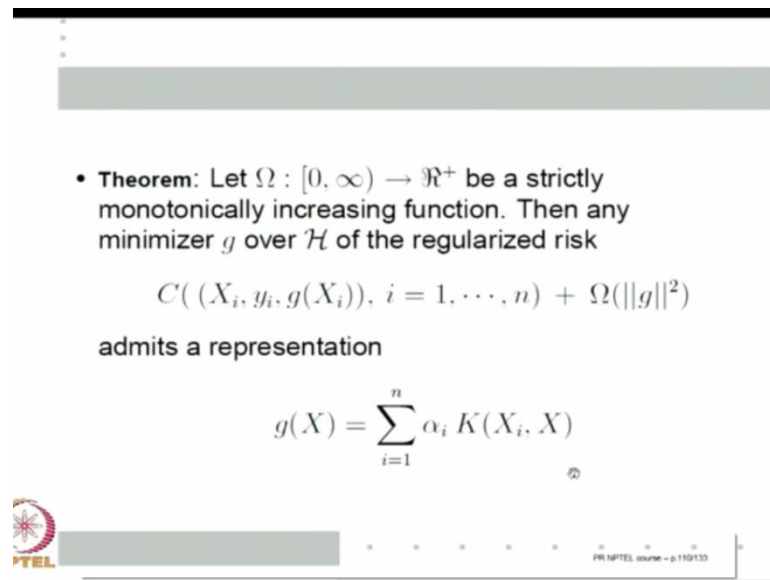


© NPTEL, 2019. All rights reserved. | Page 109/133

So, we shown the following given a positive definite kernel there is a vector space within a inner product namely the RKHS associated with K and a mapping phi from X to H. Such that the kernel is an inner product in H the RKHS represents a space of functions where you can search for the empirical risk minimize is a is is is essentially that that kind of functions. Now, given this, we will just do one very important insight which is called the Representer theorem. What is Representer theorem say? Let K be a positive definite kernel and H be the associated RKHS and X_i, y_i be the training data set as earlier. Now, any given any function f suppose, I want empirical risk under this training data set.

So, empirical risk depends on some loss function l of y_i comma $f(X_i)$ whatever so, no matter what is the loss function is, the empirical risk ultimately depends on what $X_i, y_i, f(X_i)$ for i is equal to 1 to n. It cannot depend on anything else, given this training examples and a function f. The empirical risk can depend only on $X_i, y_i, f(X_i)$ for i is equal to 1 to n. So, let us write empirical risk of any function f as some function C of $X_i, y_i, f(X_i)$ for i is equal to 1 to n. So, we do not have to worry about what our loss function is, let us say we want to minimize H i empirical risk. But we want to do a regularized empirical risk minimization that means we use a regularization term for which we use the norm in the H space. And let us say the inner product f comma f , we will represent as norm f square like this.

(Refer Slide Time: 51:38)



• **Theorem:** Let $\Omega : [0, \infty) \rightarrow \mathbb{R}^+$ be a strictly monotonically increasing function. Then any minimizer g over $\mathcal{H}$ of the regularized risk

$$C((X_i, y_i, g(X_i)), i = 1, \dots, n) + \Omega(\|g\|^2)$$

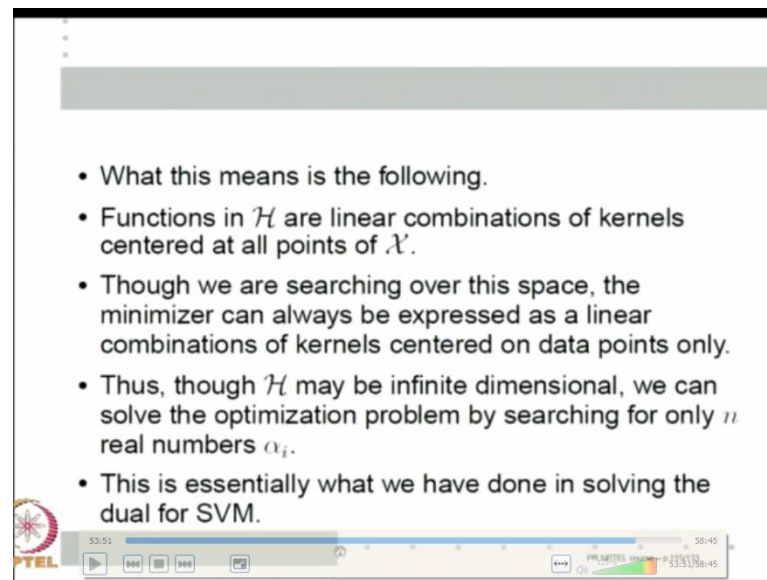
admits a representation

$$g(X) = \sum_{i=1}^n \alpha_i K(X_i, X)$$

PR NPTEL course - p.110/133

So, this is the theorem, let Ω be any strictly monotonically increasing function then if g is any minimizer minimize is I am searching over the RKHS. And g is the minimizer of the regularized risk there regularized risk is this empirical risk under any arbitrary loss function plus the regularization term Ω of which is a function of the norm of g in H . So, essentially the regularization term should be some increasing function of the norm of g . So, if, if I am using norm of g in H are the regularizing under not, not directly norm of g any increasing function of norm of g would do. Then any minimizer can be represented as the linear combination of kernels centered at X_i . So, g dot is $\alpha_i K(X_i, X)$ dot this. So, the the final minimizing function can be simply written as linear combination of kernel functions centered at data points X_i .

(Refer Slide Time: 52:45)

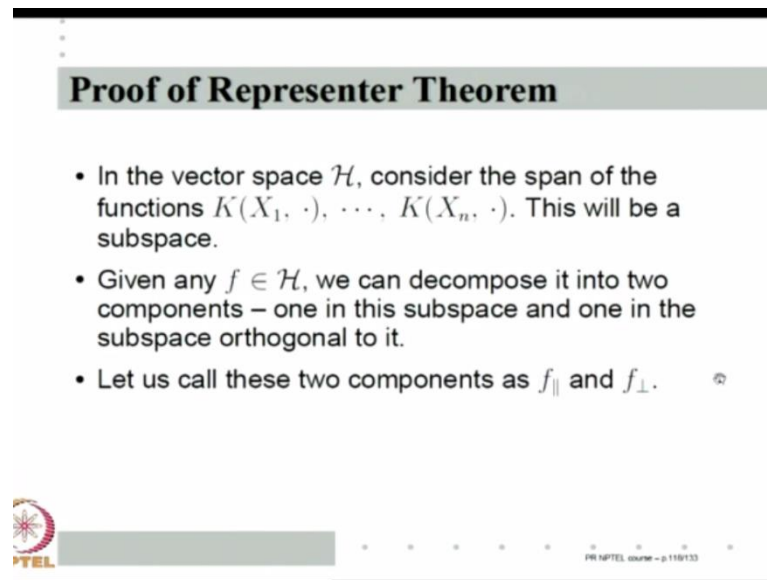


- What this means is the following.
- Functions in $\mathcal{H}$ are linear combinations of kernels centered at all points of $\mathcal{X}$.
- Though we are searching over this space, the minimizer can always be expressed as a linear combinations of kernels centered on data points only.
- Thus, though $\mathcal{H}$ may be infinite dimensional, we can solve the optimization problem by searching for only n real numbers α_i .
- This is essentially what we have done in solving the dual for SVM.

Let us see, what it means, functions in H are linear combination of kernel centered at all points of X there be, there may be uncountably infinitely many points in X . So that many functions are there in H though, we are searching over the space, the minimizer can always be expressed as a linear combination of kernel centered at data points only. And they are only finitely many data points, which means though H may be infinite dimensional, we can solve the optimization problem by searching for only n real numbers α_i .

I do not have to worry about what the g 's representation is my minimizer g always representation like this where X_i 's are always all given. So, my g dot is nothing but $\alpha_i K(X_i, \cdot)$. So, essentially I need only α_1 to α_n to find my g , my minimizer. Even though H , H contains linear combination kernel centered all points of X . And H may be α_1 to α_n , find my g my minimize even though H , H contains linear combination kernel centered all points of X . And H may be infinite dimensional I can find the minimizer by searching for only these n numbers α_i . This is essentially what we did when we solve the dual in SVM irrespective of the, as I said are the dimensions of the transform of the feature vector. The dual is a dimensionality equal to the number of examples, because this is a very generic property when we work with kernel. If you are doing a search over the RKHS for a minimize of the regularized risk under any loss function as long as regularizer uses the norm of the function in H .

(Refer Slide Time: 54:19)

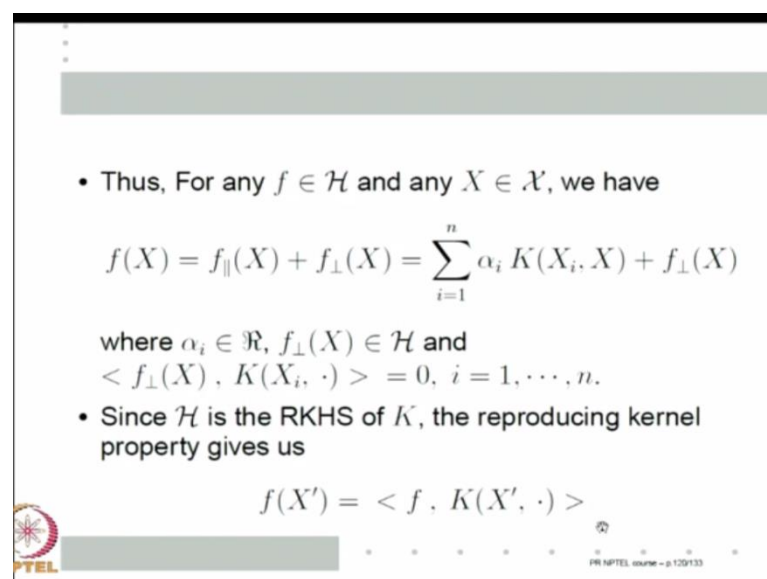


Proof of Representer Theorem

- In the vector space $\mathcal{H}$, consider the span of the functions $K(X_1, \cdot), \dots, K(X_n, \cdot)$. This will be a subspace.
- Given any $f \in \mathcal{H}$, we can decompose it into two components – one in this subspace and one in the subspace orthogonal to it.
- Let us call these two components as $f_{\parallel}$ and $f_{\perp}$.

NPTEL course - p.118/133

(Refer Slide Time: 54:47)



- Thus, For any $f \in \mathcal{H}$ and any $X \in \mathcal{X}$, we have

$$f(X) = f_{\parallel}(X) + f_{\perp}(X) = \sum_{i=1}^n \alpha_i K(X_i, X) + f_{\perp}(X)$$

where $\alpha_i \in \mathbb{R}$, $f_{\perp}(X) \in \mathcal{H}$ and $\langle f_{\perp}(X), K(X_i, \cdot) \rangle = 0, i = 1, \dots, n$.

- Since $\mathcal{H}$ is the RKHS of K , the reproducing kernel property gives us

$$f(X') = \langle f, K(X', \cdot) \rangle$$

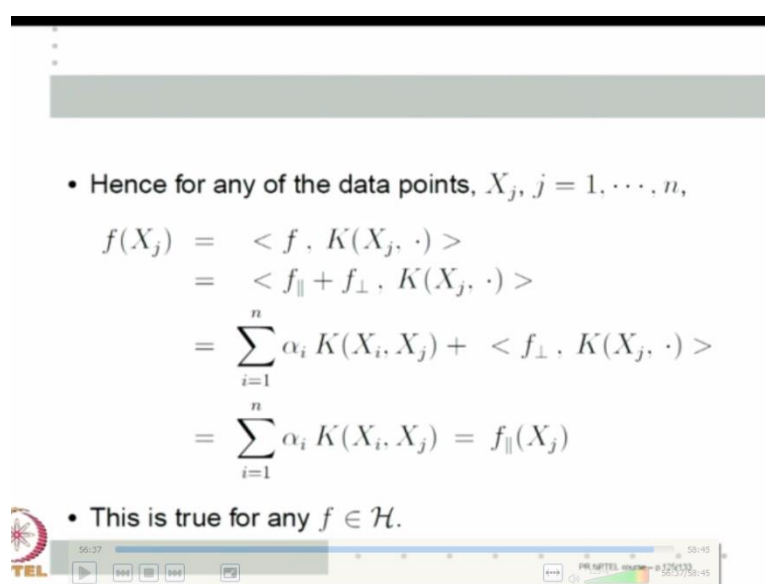
NPTEL course - p.120/133

So, let us quickly prove the Representer theorem in the vector space of $\mathcal{H}$, consider this span of the functions $K(X_1, \cdot), \dots, K(X_n, \cdot)$. So take the kernel centered at each of the n data points and take their linear span that is a subspace so, because these are subspace given any f in my RKHS that f can be resolved into that component which is in the subspace. And the component that is orthogonal to it, let us call these 2 components $f_{\parallel}$ and $f_{\perp}$ which means for any f in $\mathcal{H}$ and any X . Because f is f can be written as $f_{\parallel} + f_{\perp}$ $f(X) = f_{\parallel}(X) + f_{\perp}(X)$

what is that part of f which is in the linear span of X_i that means $f_{\parallel}$ can be written as a linear combination of $K(X_i, \cdot)$.

So, $f_{\parallel}$ is nothing but $\sum_{i=1}^n \alpha_i K(X_i, \cdot)$ for some α_i 's and $f_{\perp}$ is perpendicular to this, that means its inner product with this with $f_{\parallel}$ should be 0. And by bilinearity its inner product with each of the $K(X_i, \cdot)$ should be 0 for if X_i 's are the data points. Then $f_{\perp}$ is a function such that $f_{\perp} \text{ comma } K(X_i, \cdot) \text{ dot}$ is equal to 0 for $i = 1$ to n , this is true for every function. Given my data points $X_1, \dots, X_n$ I can resolve every function as the parallel and perpendicular components of this K of the linear span of $K(X_i, \cdot)$. And hence this kind of decomposition holds. Since H is RKHS reproducing kernel property gives me $f(X_i) = f \text{ comma } K(X_i, \cdot) \text{ dot}$, the inner product of f and $K(X_i, \cdot)$ is $f(X_i)$.

(Refer Slide Time: 56:02)



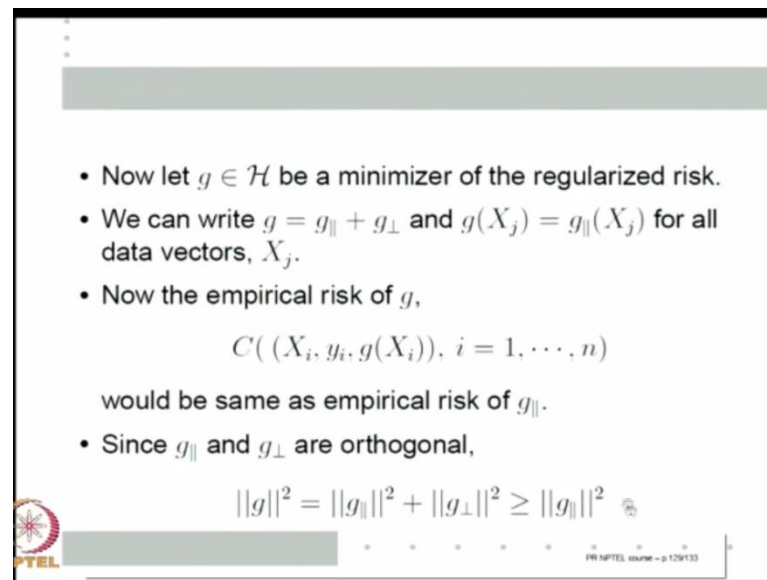
• Hence for any of the data points, $X_j, j = 1, \dots, n$,

$$\begin{aligned}
 f(X_j) &= \langle f, K(X_j, \cdot) \rangle \\
 &= \langle f_{\parallel} + f_{\perp}, K(X_j, \cdot) \rangle \\
 &= \sum_{i=1}^n \alpha_i K(X_i, X_j) + \langle f_{\perp}, K(X_j, \cdot) \rangle \\
 &= \sum_{i=1}^n \alpha_i K(X_i, X_j) = f_{\parallel}(X_j)
 \end{aligned}$$

• This is true for any $f \in \mathcal{H}$.

Which means if I take any data points X_j $f(X_j)$ is the inner product of f and $K(X_j, \cdot)$, f is $f_{\parallel}$ plus $f_{\perp}$ $K(X_j, \cdot)$. By bilinearity this is $f_{\parallel}$ inner product with $K(X_j, \cdot)$ and $f_{\perp}$ inner product with $K(X_j, \cdot)$. $f_{\perp}$ inner product with $K(X_j, \cdot)$ is 0, which means for each of the data points X_j $f(X_j)$ is same as $f_{\parallel}(X_j)$ this is true for every function f in H .

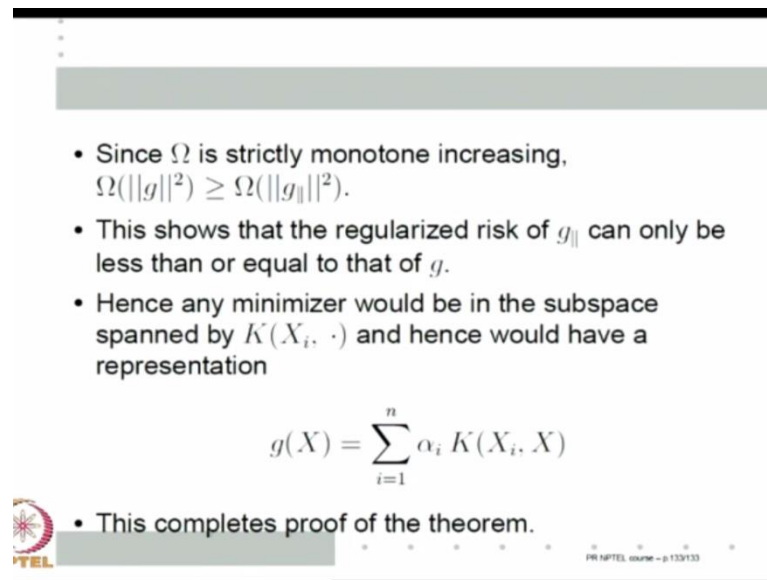
(Refer Slide Time: 56:39)



- Now let $g \in \mathcal{H}$ be a minimizer of the regularized risk.
- We can write $g = g_{\parallel} + g_{\perp}$ and $g(X_j) = g_{\parallel}(X_j)$ for all data vectors, X_j .
- Now the empirical risk of g ,
$$C((X_i, y_i, g(X_i)), i = 1, \dots, n)$$
would be same as empirical risk of $g_{\parallel}$.
- Since $g_{\parallel}$ and $g_{\perp}$ are orthogonal,
$$\|g\|^2 = \|g_{\parallel}\|^2 + \|g_{\perp}\|^2 \geq \|g_{\parallel}\|^2$$

Now, let g be any minimizers of the regularized risk we can write g as $g_{\parallel}$ plus $g_{\perp}$. And hence we know $g(X_j)$ is $g_{\parallel}(X_j)$ for all data vectors X_j . Now, the empirical risk only depends on $X_i, y_i, g(X_i)$. So, it depends on values of g 's only on the data points on the data points g and $g_{\parallel}$ are same. So, the empirical risk of g is same as empirical risk of $g_{\parallel}$. So, we know that empirical risk of g and $g_{\parallel}$ are same. Now, let us look at the regularizing term. So, we know $\|g\|^2$ is $\|g_{\parallel}\|^2 + \|g_{\perp}\|^2$ these two are orthogonal and so hence is greater than $\|g_{\parallel}\|^2$.

(Refer Slide Time: 57:22)



- Since Ω is strictly monotone increasing, $\Omega(\|g\|^2) \geq \Omega(\|g_{\parallel}\|^2)$.
- This shows that the regularized risk of $g_{\parallel}$ can only be less than or equal to that of g .
- Hence any minimizer would be in the subspace spanned by $K(X_i, \cdot)$ and hence would have a representation

$$g(X) = \sum_{i=1}^n \alpha_i K(X_i, X)$$

- This completes proof of the theorem.

PHILIPPS UNIVERSITÄT ERLANGEN-NÜRNBERG

So, which means $\Omega(g^2)$ is greater than $\Omega(g_{\parallel}^2)$, which means $g_{\parallel}$ cannot have a regularized risk anymore than that of g that means $g_{\parallel}$ is also minimizer of risk. And hence my minimizer always admits this representation. So, apart from the proof this particular theorem is very important, it shows that essentially my minimizer if I am doing empirical risk over empirical, risk minimizer any loss function over the RKHS using the norm of a the function under H as the regularizer. Then it has this nice support vector expansion, this is a very generic property of all kind of kernels as long as kernels are pointed out, that is the reason why one looks at positive definite kernels. So, we will stop about the kernel based methods here. So, next class, we will, we will just round up about what all we have done about learning non linear classifiers. And then move on to a few special topics to wind up the course.

Thank you.