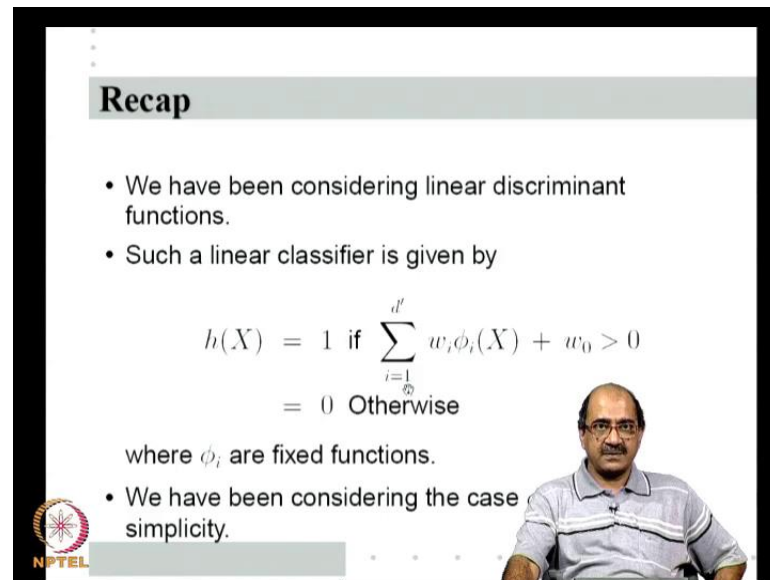


Pattern Recognition
Prof. P. S. Sastry
Department of Electronics and Communication Engineering
Indian Institute of Science, Bangalore

Lecture - 15
Linear Least Squares Regression; LMS Algorithm

(Refer Slide Time: 00:40)



Recap

- We have been considering linear discriminant functions.
- Such a linear classifier is given by
$$h(X) = \begin{cases} 1 & \text{if } \sum_{i=1}^{d'} w_i \phi_i(X) + w_0 > 0 \\ 0 & \text{Otherwise} \end{cases}$$
where ϕ_i are fixed functions.
- We have been considering the case simplicity.

NPTEL

Hello and welcome to the next lecture on pattern recognition. We will, we have been discussing linear discriminant functions last class. We have done with looking at Bayes classifier. We are looking at other ways of learning classifiers and regressors. So, we started with the linear discriminant functions. Basically, a classifier using linear discriminant functions has the following structure. The classifier will say one, if sum of $w_i \phi_i(X)$ plus w_0 is greater than 0 for some prefixed functions ϕ_i , right? This is what we have seen last time, otherwise 0. As I said in the beginning of last lecture, for now we are only looking at two classes. At the end of maybe next lecture or the lecture after that, we will consider how to generalize this to multi class problems?

(Refer Slide Time: 01:38)

Perceptron

- Perceptron is the earliest such classifier.
- Assuming augmented feature vector,
 $h(X) = \text{sgn}(W^T X)$.
- 'find weighted sum and threshold'

The diagram shows a perceptron model. On the left, there are input features $x_1, \dots, x_d$. Each feature is connected to a central circular node by a weighted line labeled $w_1, \dots, w_d$. The output of the node is a single value y . The NPTEL logo is visible in the bottom left corner of the slide.

So, this is the structure of the classifier that we are considering. We have considered last class. Of course, we have taken $\phi_i X$ to be X_i for simplicity. So, this simply becomes $w_i x_i$ and d' becomes d which is the dimension of x . As we said whether we use fixed functions ϕ_i or whether we use $w_i x_i$, the techniques involved in learning them are the same. So, all of them we will look at as linear classifiers. Now, the linear classifiers algorithm that we considered last time in the last lecture is perceptron, which are the earliest such classifiers.

Essentially as we have seen by augmenting the feature vector, we can write any linear classifier as $W^T X$. Now, both w and x will be plus one dimensional vector because of augmentation. So, this is nothing but you find a weighted sum. $w^T x$ is simply weighted sum of the feature values. You sum $w_i x_i$. Then, this sign function simply threshold set. So, we seen that is that kind of a structure where there are n input signals which are the features. Each one has its associated weight. You find this unit finds the weighted sum and then thresholds and that is the output. Okay?

(Refer Slide Time: 02:32)

Perceptron Learning Algorithm

- A simple iterative algorithm.
- Each iteration, we locally try to correct errors.

Let $\Delta W(k) = W(k+1) - W(k)$. Then

$$\Delta W(k) = \begin{cases} 0 & \text{if } W(k)^T X(k) > 0 \text{ \& } y(k) = 1, \text{ or } \\ & W(k)^T X(k) < 0 \text{ \& } y(k) = 0 \\ X(k) & \text{if } W(k)^T X(k) \leq 0 \text{ \& } y(k) = 1 \\ -X(k) & \text{if } W(k)^T X(k) \geq 0 \text{ \& } y(k) = 0 \end{cases}$$

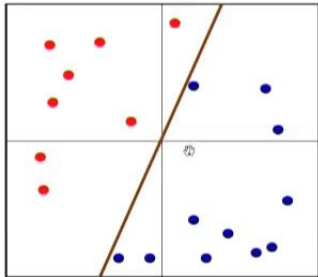

PR NPTEL course - p 10/17

Such a device is called perceptron. Perceptron learning algorithm, as we see in last class is a very simple and iterative algorithm. It tries to learn the weights iteratively. At each iteration, we take one feature vector trained. Ask whether the current weight vector correctly classifies it. If it does not, then we locally correct the errors, right? To recall, this was the algorithm. So, if I take ΔW to be $W_{k+1} - W_k$, this is the correction that is made. If W_k correctly classifies the next training part and that is X_k ; X_k is the training pattern picked up at iteration k .

(Refer Slide Time: 03:32)

Perceptron: Geometric view

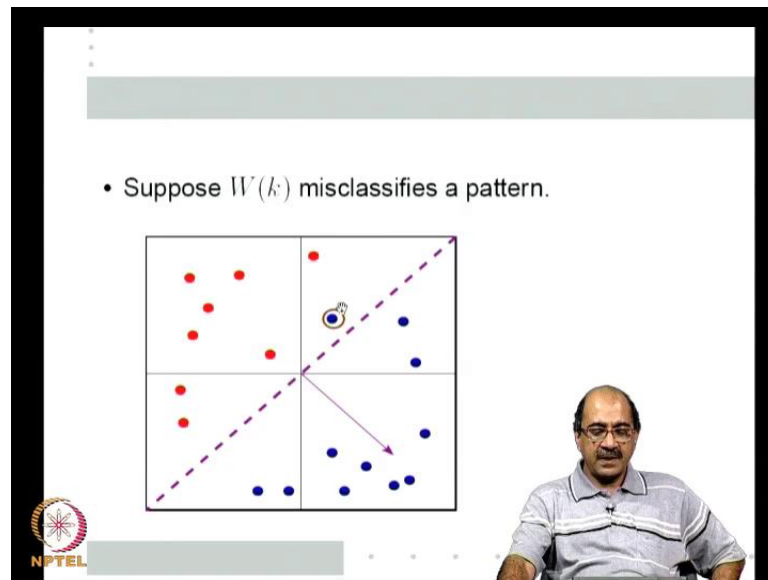
The algorithm has a simple geometric view. Consider the following data set.





So, if $W^T X_k$ is greater than 0 and y_k is 1, that is correct or if y_k is equal to 0; $W^T X_k$ is less than 0. Then, I do no correction. In the other two cases of errors, I either add or subtract X_k . As we have seen, this is a very nice method as to recall. It has a very simple geometric interpretation.

(Refer Slide Time: 03:44)



So, here is a simple linear least separable class that is that hyperplane could separate them. But, let us suppose we started with some hyperplane, which makes another one point is on the wrong side. This is the W vector. So, all the blues are the, so to say positive class. Reds are negative class in the sense for all the blues $W^T X$ is positive. For all the reds, $W^T X$ is negative. $W^T X$ is simple, the dot product. As you can see in the augmented feature space, all the hyperplanes pass through the origin; not the hyperplane equation is $W^T X$ is equal to 0.

(Refer Slide Time: 04:19)

• Now the correction made to $W(k)$ can be seen as

So, now we have this wrongly classified thing, which is on the wrong side of the hyper plane. As we seen what it means is that, I take $W(k)$, add $X(k)$ to it. So, that becomes a new normal vector, which is essentially same as moving, the rotating the hyperplane in the right direction whether we adding or subtracting; essentially a matter of rotating the hyperplane in one direction or the other. Of course, when we do a local correction like that in the previous hyper plane, this point is correctly classified. But now, after I took care of correction for this point, now this point becomes wrongly classified.

(Refer Slide Time: 05:02)

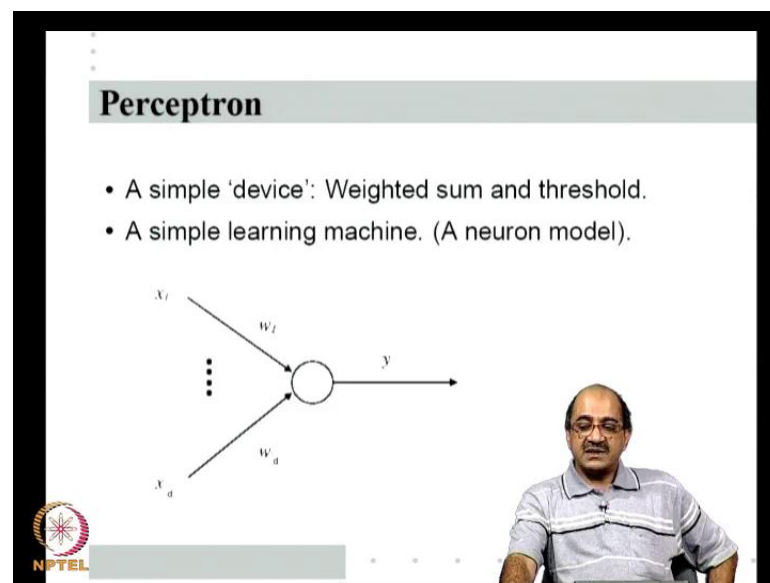
• We showed that: if the training set is linearly separable, the the algorithm would find a separating hyperplane in finitely many iterations.

• We also saw the 'batch' version of the algorithm. It is shown to be a gradient descent on a reasonable cost function.

So, correction now could have made something else that was correct earlier to become non correct. But, not withstanding that we know that this algorithm converges. We showed that if the training set is linearly separable, these series of corrections will stop at some finite time. We will always find a separating hyperplane after finitely many iterations. So, this algorithm is of course incremental in the sense I look at only one feature vector at a time.

So, in principle if I am just getting a stream of feature vectors, I can run the algorithm with that without needing to store all the feature vectors in the main memory. We also looked at a batch version where using the current W_k , we classify each one of the training vectors. Then, looking at the correctness of all of them together we make all the corrections together to the W_k . We shown that the batch version is essentially a gradient descent and a reasonable cost function and hence that also converges.

(Refer Slide Time: 05:58)



So, to sum up, perceptron is a simple device, essentially weighted sum and threshold. It is also a simple learning machine, one of the first learning machines, one of the first learning algorithms of pattern recognition. It is also can be thought of as a neuron model. As I told you, a simple model of a neuron is that it has many input ports which are called dendrites, on which inputs come at places called synapses. Each synapse can have a weight in the sense how much that synapse has fixed the on the surface of the neuron.

So, I can think of a neuron, a simple neuron for a mathematical model to be essentially doing a weighted sum of all the inputs and then thresholding it. Right? In that sense, is a simple neuron model and the weights can be adopted. As a point, we have to last class because the adaptation is simply based on the kind of error that is known at this point, you add W to X . That means for each component, you add the corresponding component.

So, to update this weight, all the information needed is available on this link itself because it just needs to know the value of X_i there. So, it is a very nice local algorithm. I did not really explain it in detail last class. But, perceptron which was proposed by Rosenblatt was proposed not just as an abstract pattern classification algorithm. But, this is proposed as one possible way in which our brain can learn. The idea is that, by then people knew that there are neurons like this.

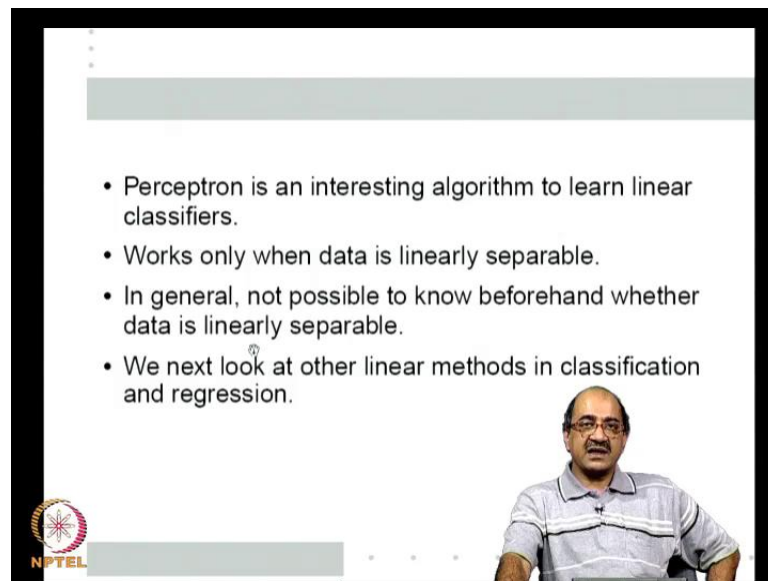
The outputs of other neurons will come and synapse at the inputs of some other neuron. That is how they interact. Essentially, all the currents passing through all the synapses will increase the potential on the surface of the neuron. Once the potential goes below some above some value, it triggers an action potential which will be then input to other neurons and so on. So, given this, any change in the information knowledge problem solving ability of the brain can come about only by changing synaptic weights.

So, everybody is looking at algorithms of how synaptic weights can be changed, so that non-trivial processing can be done. So, Rosenblatt proposed this as a model of how we may learn visual categorization. So, I am shown some images and say these are dogs. I am shown some other objects, say these are cats. We can think of these features as measuring some quantities on the image. As I have seen already, if X_i is replaced by ϕ_i of X , nothing changes in the algorithm. So, I can think of each of these as ϕ_i of the pattern that is put in front of the neuron which could be in front of your retina, for example.

So, given some pattern, there could be some measurements that we are bound with. So, these ϕ_i 's are fixed. So meaning, we have born with some kind of feature detection we can do on the image that is in front of us. Given that these things are fixed and somebody wants to teach us, let us say teach us to distinguish between Maruti cars and let us say Ambassador cars. Obviously, it is too much to expect that we have born with an ability to distinguish them.

So, using some generic measurements that we have, we somehow learn. Given enough supervised examples being pointed out that this is a Maruti car and that is an Ambassador car. Even though we may not be able to very precisely describe the difference in shapes between Maruti cars and Ambassador cars, we learn to distinguish between the shapes and can do the classification correctly. This could be one example of a method by which such visual categorizations could have been learnt.

(Refer Slide Time: 10:44)



- Perceptron is an interesting algorithm to learn linear classifiers.
- Works only when data is linearly separable.
- In general, not possible to know beforehand whether data is linearly separable.
- We next look at other linear methods in classification and regression.

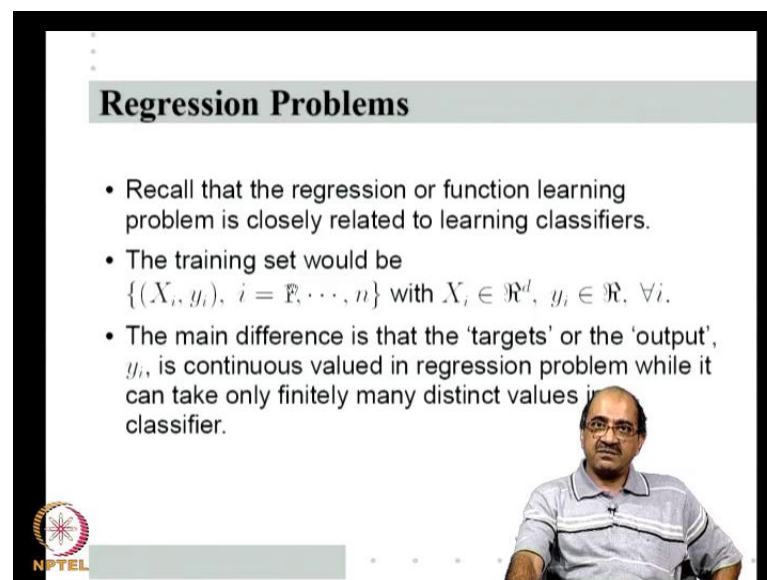
It is in that sense that perceptron was originally proposed. For its time, it is a very interesting algorithm with a very interesting convergence proof. As a matter of fact, it was actually built using electronic hardware those days. These are pre-IC days using mostly op-amps. A matter of fact, a model of the perceptron is even kept in the Smithsonian museum. So, it was in the development of a in pattern recognition. Perceptron is a historically very important step.

Anyway as you seen, it is an interesting and simple learning machine. It can learn linear classifiers. Its main drawback is that it works only when the data is linearly separable. Now, for a particular given task say, Maruti cars versus Ambassador cars. I do not even know what features, detectors I have. How can I at all say whether the data is linearly separable or not? In general, even if you are given some vectors in a d -dimensional space, it is not possible to easily deduce whether or not the data is nearly

separable. That brings us as we have seen in last class, if the training data is not linearly separable, then the perceptron algorithm can potentially go into an infinite loop.

Even though we have a convergence proof of finite iterations, we do not know how many iterations it actually takes. We do not know whether you know we are going towards convergence or we are going to an infinite loop. So, that is one real problem with the perceptron algorithm. So, we next look at other methods for learning linear classification. When you are looking at this, we look at two class classification along with the regression problem. If you remember right in the beginning, we looked at both classification as well as function learning problems and said both of them are quite similar. The function learning problem is what the regression problem is about.

(Refer Slide Time: 12:31)



Regression Problems

- Recall that the regression or function learning problem is closely related to learning classifiers.
- The training set would be $\{(X_i, y_i), i = 1, \dots, n\}$ with $X_i \in \mathbb{R}^d$, $y_i \in \mathbb{R}$, $\forall i$.
- The main difference is that the 'targets' or the 'output', y_i , is continuous valued in regression problem while it can take only finitely many distinct values in classifier.

NPTEL

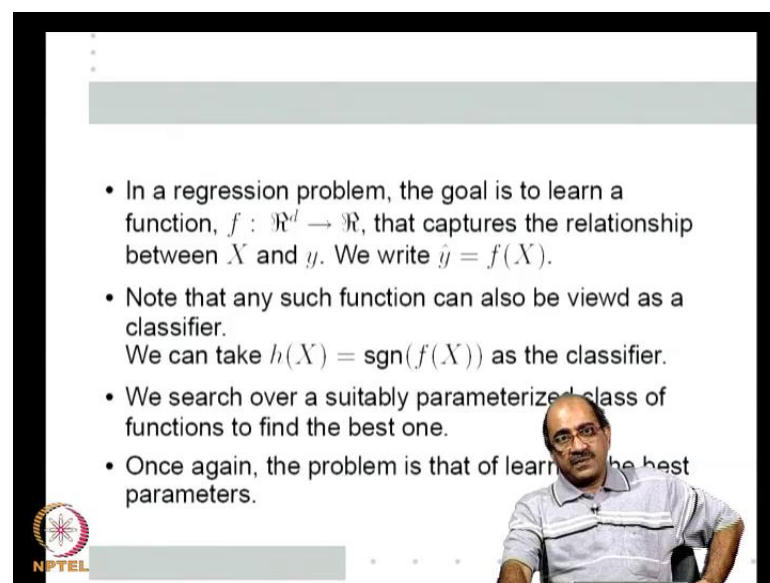
So, the next techniques we look at learning linear classifiers are also techniques for learning linear models of a function. So, we look at them together. So, to briefly recall, the regression or function learning problem is a closely related problem; a problem closely related to learning classifiers. Essentially, when I am learning a function, when I am asked to learn a function; from examples, the training data is X_i, y_i . Again, i going to 1 to n . X_i is still a d vector, but y_i can now be any real number, right?

So, the main difference between a classification and a functional learning problem is that, the targets we can call, or the outputs; namely y_i can be continuous value in a

function learning problem whereas they take only finitely many values in a classification problem. In a two class classification problem, they take only two values. In a multiclass classification problem, may be they will take k distinguishable values. But in general, in a classification problem, because it is a categorization problem, the targets or the outputs; they take only finitely many values whereas, in regression problem the output is continuous.

That is the only difference between learning a function or instead of learning a classifier otherwise both of them are functions on $\mathbb{R}^d$. Only thing is one could be a general real valued function. The other is a boundary valued function; a function that takes only finitely many values. This being so, it is no surprise that similar technique should work for both.

(Refer Slide Time: 13:58)



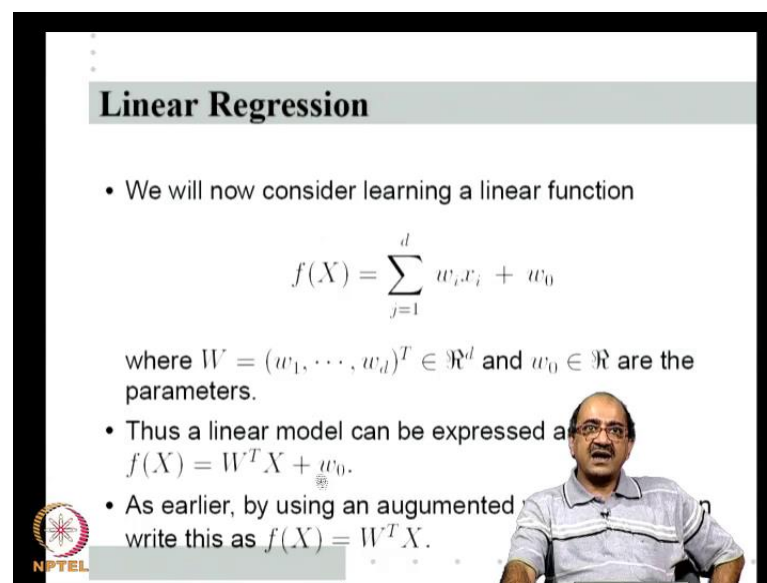
- In a regression problem, the goal is to learn a function, $f : \mathbb{R}^d \rightarrow \mathbb{R}$, that captures the relationship between X and y . We write $\hat{y} = f(X)$.
- Note that any such function can also be viewed as a classifier. We can take $h(X) = \text{sgn}(f(X))$ as the classifier.
- We search over a suitably parameterized class of functions to find the best one.
- Once again, the problem is that of learning the best parameters.

So, the next technique, we are going to look at, can we learn either classifiers or functions with the same algorithm. In a regression problem what is the goal? The goal is to learn a function from $\mathbb{R}^d$ to $\mathbb{R}$ because our target values are in $\mathbb{R}$. So, the idea is that, we are capturing the relationship between x and y using a function. So, we can write $\hat{y}$ as some function f of X . The function f is what you want to learn. So, given $X \in \mathbb{R}^d$, we can abstract the functional relationship between X and y , so that now you give me a new X . I can tell you what is y , namely, $\hat{y}$ is equal to f of X .

That is how I am going to make my prediction of the targets for any new X . In the classifier, this will be a binding function, but otherwise it can be a general; a function that takes any arbitrary real value. So, any such function can also be viewed as a classifier. For example, I can take $\text{sgn}(f(X))$ as a classifier, if I am any two class problem. So, learning a function R_d to R_c can also be viewed as a classification problem, two class classification problems simply by taking the classifier to be $\text{sgn}(f(X))$.

So, essentially the idea now we use to learn this function f given the data. So, if you have to learn a function, what we do is we start with a family of parameter as a functions. Then, we find the best function among them which boils down to once again learning the best set of parameters from a parameter space base. Given, a parameter class of functions and the training data using which we want to learn to get through the best function.

(Refer Slide Time: 15:41)



Linear Regression

- We will now consider learning a linear function

$$f(X) = \sum_{j=1}^d w_j x_j + w_0$$

where $W = (w_1, \dots, w_d)^T \in \mathbb{R}^d$ and $w_0 \in \mathbb{R}$ are the parameters.

- Thus a linear model can be expressed as $f(X) = W^T X + w_0$.
- As earlier, by using an augmented write this as $f(X) = W^T X$.

NPTEL

Now, in in this class, in this lecture and the next few lectures, we will be mostly be considering linear classifiers and linear regressors. Hence, we are looking at functions of the following general form. They are essentially called affine functions because of the constant. As you know, if there is a constant the function is not linear, but we already seen it earlier in the perceptron case that we can easily make it linear that is W transpose X by simply augmenting.

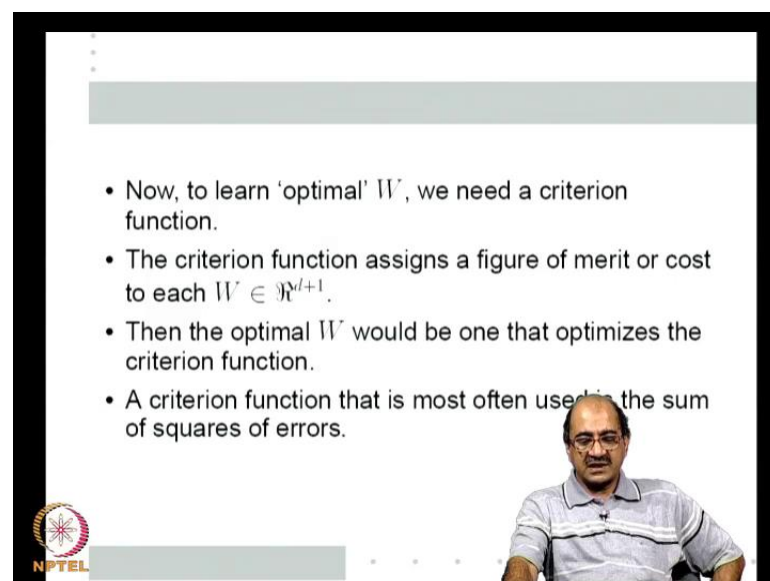
So, linear and affine functions are same for us, so looking at learning a general linear function of this form $f(X)$ is J . It is equal to 1 to d $W_i X_i$ plus W_0 where W_1 to W_d .

We call that vector W , is an $R^{d \times n}$. w_0 is a scalar. These are the parameters of the function. That is the linear model, that is, we are representing the relationship between y and x through a linear model. Namely, $W^T X + w_0$.

Just like in the classifier case, I can simply augment x with an extra feature which is always 1, extra component which is always 1. Assume that w_0 is substituted to w . So, the w now becomes $d+1$ dimensional vector. Then, I can write that as $W^T X$ equals to w_0 . So if you remember, last time when we did this, we for a while we used $\tilde{x}$ for the augmented x and $\tilde{W}$ for the augmented W . But now that we know this trick, we either write it as $W^T X + w_0$ or $W^T X$ without changing the notation.

Here also we use the same w . When I write like this, we mean X is $d+1$ dimensional w is $d+1$ dimensional. When we write like this, x is d dimensional and w is d dimensional. Right? As long as the difference is clear from context, we can easily take this. Hence, you know, without changing our notation; when we want we use augmented vector notation, when we want we use the full expanded notation.

(Refer Slide Time: 18:04)

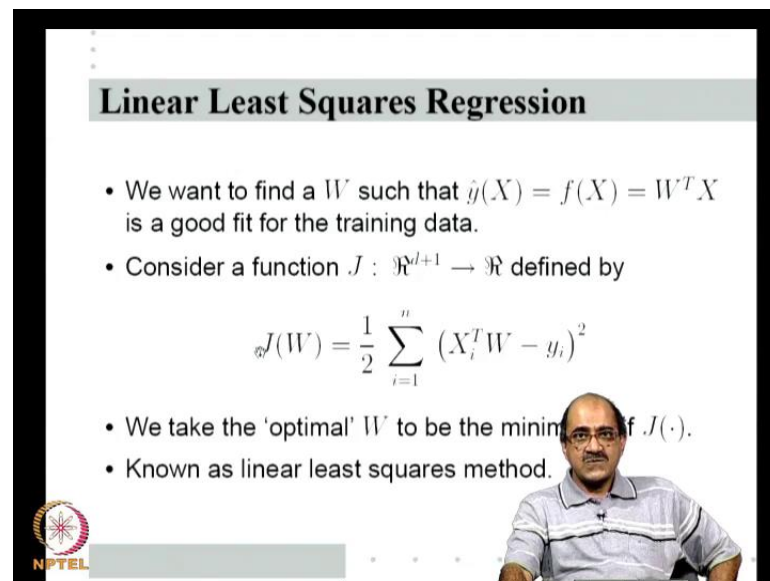


- Now, to learn 'optimal' W , we need a criterion function.
- The criterion function assigns a figure of merit or cost to each $W \in \mathbb{R}^{d+1}$.
- Then the optimal W would be one that optimizes the criterion function.
- A criterion function that is most often used is the sum of squares of errors.

So, now our task is to learn this W vector given all the training data. So, to learn the optimal W we need a criterion function. What does a criterion function do? We want

some function that assigns either a figure of merit or a cost to each W belonging to $\mathbb{R}^d$ plus 1 dimension space. So, essentially the function is telling how good W is, we want to maximize the function. If the function is telling me how bad W is, we want to minimize the function. Any case is called a criterion function. We use the criterion function to define what is our optimal W . So, the optimal W would be that value of W which will optimize the criterion function. A criterion function that is most often used is the so called sum of squares of errors.

(Refer Slide Time: 18:57)



Linear Least Squares Regression

- We want to find a W such that $\hat{y}(X) = f(X) = W^T X$ is a good fit for the training data.
- Consider a function $J : \mathbb{R}^{d+1} \rightarrow \mathbb{R}$ defined by

$$J(W) = \frac{1}{2} \sum_{i=1}^n (X_i^T W - y_i)^2$$

- We take the 'optimal' W to be the minimum of $J(\cdot)$.
- Known as linear least squares method.

So, we will look at this criterion function next. So, basically before we write the criterion function let us ask what is our objective. See, the criterion function should be such that our objective is well captured. We want to approximate the relationship between X and y through a function f . So, we are writing $\hat{y}$ of X is f of X is $W^T X$. So, we want $W^T X$ to be a good guess at y corresponding to X . So, in particular, we want this to be a good fit for the training data that is, if I put X to be any training data, $X_i^T W$ should be close to the corresponding y_i .

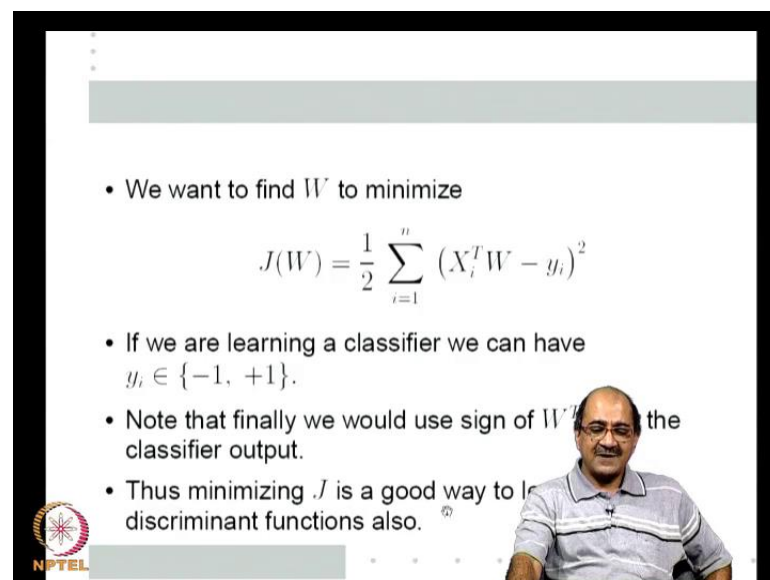
So, we can look at a criterion function like this. This is the function defined from $\mathbb{R}^d$ to $\mathbb{R}$ so it assigns a cost to every vector W . The cost is if W is my final solution for the function with that function w , with that w if you give me X_i , I would have predicted $X_i^T W$ where as the true value is y_i . So, $X_i^T W - y_i$ is the error. I square the errors because some errors may be positive some amount of errors

may be negative. So, if I simply add even though I make lot of errors, I may artificially think that I have no error.

So, we square out the errors and add them. So, this JW is always positive. It will be 0 if y_i exactly matches $X_i^T W$, otherwise it will not be 0. If JW is low then it means that in you know, most of the time using W , I am able to predict the corresponding y_i for a given X_i . If you give me X_i , my prediction is $X_i^T W$. The true value is y_i . I am taking the sum of squares of errors. If sum of squares of errors is small, that means my prediction is not true for half.

So, it is good to look at optimal W as the minimizer of J . If I can find a W that finds a global minimum of J , then that is the best error I can get. This is often called the linear least squares method or sometimes mean square method. So, we are taking the sum of squares of errors and asking which W will minimize. This is called the linear least squares method. So, we want a least square solution is that value W which minimizes this function. This function is very reasonable because essentially each W is being rated on how well it predicts the correct target value for each of the training samples.

(Refer Slide Time: 21:39)



- We want to find W to minimize

$$J(W) = \frac{1}{2} \sum_{i=1}^n (X_i^T W - y_i)^2$$

- If we are learning a classifier we can have $y_i \in \{-1, +1\}$.
- Note that finally we would use sign of $W^T X_i$ as the classifier output.
- Thus minimizing J is a good way to learn discriminant functions also.

NPTEL

So, our objective is to find a W to minimize this. As we seen for a function learning problem, this is a this is a this is a good criterion to have because essentially I am learning the function f such that $X_i^T W$ is equal to y_i . So, if you give me X_i as the input, I would have predicted $X_i^T W$ and the true value is y_i . So, if that W minimizes this then,

most of the time I am getting my predictions right. This can also be used for classifier.

Suppose, we want to learn classifiers. Then, let us say y_i are in the range minus 1 plus 1 instead of 0/1. Let us say, the two class error represent as minus 1 plus 1, that means if I once again put it in this, I will get low error if for all X_i , which are in class plus 1; $W^T X_i$ is closer to plus 1. For all X_i that are in class minus 1, $W^T X_i$ is closer to minus 1, that is when I will get lower value of J . So, because we would be using sgn of $W^T X$ as our final classifier output.

So, if $W^T X$ is positive, we will put in plus 1 class. $W^T X$ is negative, we will put in minus 1 class so that essentially, if I am learning a W , so the $W^T X_i$ is close to plus 1 when our X_i is in plus 1 class. $W^T X_i$ is close to minus 1 when our X_i is in minus 1 class. Then, such a W will do well as a classifier too. So, I can simply take y_i is to be plus 1 minus 1. Find a W that minimizes this. With that W , I can implement a classifier is simple, say if $X_i^T W$ is positive then, I will put in plus 1 class. $X_i^T W$ is negative, I put in minus 1 class.

(Refer Slide Time: 24:02)

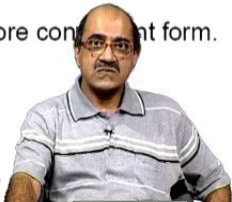


• We want to find minimizer of

$$J(W) = \frac{1}{2} \sum_{i=1}^n (X_i^T W - y_i)^2$$

• This is a quadratic function and we can analytically find the minimizer.

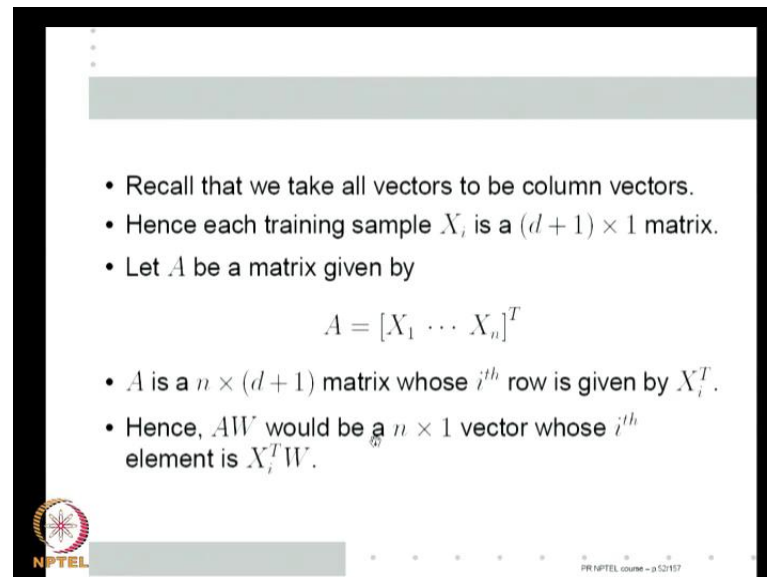
• For this we rewrite $J(W)$ into a more compact form.



So, minimizing the same criterion function, I can either learn a linear regression function or a linear classifier. Thus, minimizing J is a good way to learn linear discriminant functions. Right? So, minimizing J is a good way to learn regression functions as well as a good way to learn linear discriminant functions. Now, so let us let us look at

the computational problem. Now, I am given training data $X_i y_i$. I have this function J of W defined for every W in $\mathbb{R}^{d+1}$. So, I want to set your $\mathbb{R}^{d+1}$ to find a minimizer of this function.

(Refer Slide Time: 25:03)



- Recall that we take all vectors to be column vectors.
- Hence each training sample X_i is a $(d+1) \times 1$ matrix.
- Let A be a matrix given by

$$A = [X_1 \cdots X_n]^T$$
- A is a $n \times (d+1)$ matrix whose i^{th} row is given by X_i^T .
- Hence, AW would be a $n \times 1$ vector whose i^{th} element is $X_i^T W$.

NPTEL

Now, this is a quadratic function, right? So, because this is a quadratic function, it is no surprise that we can analytically find the minimizer. So, what we will do next is to derive the minimizer by using calculus. To do that, we will have to first put this expression into a slightly different form where it is easier to differentiate it and find the minimizer. So, what we are going to do is we are going to rewrite J of W into a more convenient form. Recall that we put it in the in this form, we need to do some vector algebra. We are going to represent this expression in in in in a vector matrix notation.

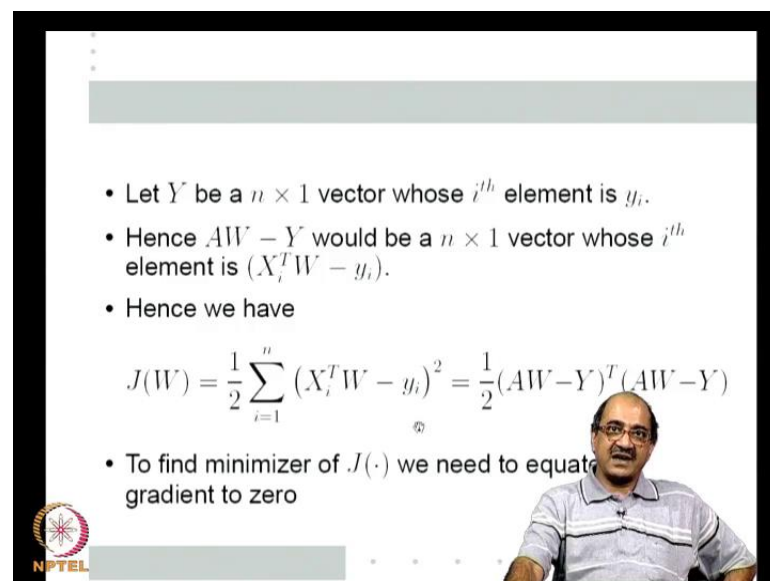
So, so far even though all our feature vectors are vectors, we never particularly bother about that representation. But right in the beginning, I told you that in this course, all vectors are always taken to be column vectors. So, we did not affect us much whether column vectors or row vectors. But now that we need a specific representation; let us recall that all vectors are column vectors. What does that mean? Each training sample X_i is a column vector. A column vector means it is a one column matrix.

So, because X_i has $d+1$ dimensions, each X_i is essentially a $d+1$ by one matrix. Now, let A be a matrix given by, so you put each of the X_1 to X_n as columns of a matrix and take its transpose. That means the first row of A will be a X_1 transpose. It

has to be X^T because X is a column vector. So, the first row of A will be the first training sample arranged as the row vector. Second row of A will be the second training data vector and so on. So see, each X_i is $d+1$ dimensional.

So, without the transpose, this will be $d+1$ by n matrix. With the transpose, it becomes n by $d+1$ matrix and whose i th row is X_i^T . You have to write it as X_i^T because you have to make it a row vector. By definition, X_i is a column vector. So, the i th row of the matrix is given by X_i^T . Now, if I do AW where A is a n by $d+1$ matrix. W is a $d+1$ vector. So, AW will be a n by 1 vector. A 's rows are represented as X_1^T, X_2^T , so on. So, in this n vector, when I take this matrix vector multiplication, the i th element will become $X_i^T W$. So, AW will be a n vector whose i th component is $X_i^T W$.

(Refer Slide Time: 27:50)



- Let Y be a $n \times 1$ vector whose i^{th} element is y_i .
- Hence $AW - Y$ would be a $n \times 1$ vector whose i^{th} element is $(X_i^T W - y_i)$.
- Hence we have

$$J(W) = \frac{1}{2} \sum_{i=1}^n (X_i^T W - y_i)^2 = \frac{1}{2} (AW - Y)^T (AW - Y)$$

- To find minimizer of $J(\cdot)$ we need to equate gradient to zero

Now, let the capital Y denotes also a n vector whose i th element is y_i . See, I have n training samples $X_i Y_i$. So, I take all the targets $y_1, y_2, \dots, y_n$. Arrange them as a vector. Then, call that y . Now, if I make a new n by 1 vector; see, A is a n by $d+1$ matrix. W is a $d+1$ by 1 vector. So, AW is a n by 1 vector. y is also an n by 1 vector.

So, AW minus y will be a n by 1 vector whose i th element will be what? The i th element here is $X_i^T W - y_i$. So, i th element of $AW - y$ is $X_i^T W - y_i$.

minus y_i . Now, we are ready to rewrite the function J . J is half of $\sum_i (X_i^T W - y_i)^2$.

Now, I have a vector here, $AW - y$ whose i -th component is exactly this; $X_i^T W - y_i$. So, this sum is actually, you take each component of this vector, square it and sum. That is nothing, but the norm square of this vector. So, the norm of that vector is simply $\|AW - y\|^2$. So, I can write J as half of $\|AW - y\|^2$. Just a small point which I should have mentioned in beginning; the only reason we put a half in this J , it is just a convention.

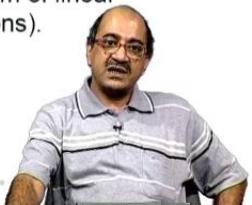
When you differentiate with respect to W because the square here, they will be 2 that comes out, so that we would not have any constants in your derivative that two will be cancelling with this half otherwise this half makes no difference. I am minimizing J . So, whether I put this half or do not put this half, the minimizer will be the same. It is just that putting that half is tradition so that this two from this square, if I differentiate, we will cancel this half. That is the only reason the half is there.

So, anyway now we have written J in terms of vector matrix notation as half $\|AW - y\|^2$ where A is a matrix, whose rows are the augmented feature vectors X_1 and X_2 , all the way upto X_n . y is all the targets. $y_1, y_2, \dots, y_n$ arranged as a n vector. W is the $d+1$ vector that I want to find by minimizing this. So, to find minimizer of J , we need to equate its gradient to 0. J is given by this. So, I need to find the gradient of this with respect to W . This is you know, $\frac{\partial}{\partial W} \frac{1}{2} \|AW - y\|^2$. So, if I differentiate with respect to W , I get $A^T(AW - y)$ and multiplied by the coefficient of W . So, of course, we have to properly take care of either A comes or A^T comes because it is a vector vector matrix differentiation.

(Refer Slide Time: 31:01)



- We have
$$\nabla J(W) = A^T(AW - Y)$$
- Equating the gradient to zero, we get
$$(A^T A)W = A^T Y$$
- The optimal W satisfies this system of linear equations. (Called normal equations).



If we do all that, we get the gradient as a transpose times AW minus y . That has to be because AW minus y is A n by 1 vector. So, I can only multiply with a transpose and not with A . So, I know that when I differentiate, of course, the half will go away. When I differentiate, I get a transpose AW minus y . So, the gradient turns out to be a transpose into AW minus y . So, if I equate the gradient to 0 , I get a transpose AW is equal to a transpose y . So, what we now know? Because we want to minimize J to find the optimal W , this is the gradient of J with respect to W .

(Refer Slide Time: 32:39)



- $A^T A$ is a $(d + 1) \times (d + 1)$ matrix.
- $A^T A$ is invertible if A has linearly independent columns. (This is because null space of A is same as null space of $A^T A$).
- Rows of A are the training samples X_i .
- Hence j^{th} column of A would give the values of j^{th} feature in all the examples.



The minimizer of $J(W)$ has to satisfy this equation. Now, this is some matrix, right? This is also, me given vector. So, this is a linear system of equations. So, what we can say is that the optimal w satisfies this system of linear equations. These are often called the normal equations. If you solve this system of linear equations, we can get W . A matter of fact, if $A^T A$ is invertible, I can pre-multiply both sides with $A^T A$. I can write an expression for W .

Now, what does a transpose A being invertible mean? A transpose J is A^T is $d+1$ by $d+1$ matrix. Recall that a transpose is A^T is n by $d+1$ matrix. So, a transpose is $d+1$ by n matrix. So, a transpose by A is $A^T A$ is $d+1$ by $d+1$ matrix. A transpose A is invertible. If A has linearly independent columns, if A has full column rank then, $A^T A$ would be invertible. This is easy to show essentially by showing that the null space of A would be same as the null space of a transpose $A^T A$. Anyway, if even, if you do not understand null space, does not matter.

$A^T A$ is invertible if and only if A has linearly independent column; that is A has full column rank. Now, let us ask when will A have full column rank. Rows of A are the training samples X_i . So, what will be the i th or j th column of A ? So, in the first row, j th column would be the j th value in the first training sample. So, that is the value of the j th feature in the first training sample. The second row j th column will have the value of the j th feature in the second training sample. The j th column of A would give us values of the j th feature in all the examples.

(Refer Slide Time: 34:29)



- Hence columns of A are linearly independent if no feature can be obtained as a linear combination of other features.
- If we assume features are linearly independent then A would have linearly independent columns and hence $A^T A$ would be invertible.
- This is a reasonable assumption.



So, each column of the matrix A will give me the values of one particular feature across all. Example, if I think of X_i 's as feature vectors and there are d features or $d + 1$ feature say, the augmented case. So, each column of A will essentially put together values of one particular feature. In all the examples, what does columns of A being linearly independent mean? No one column can be retained as a linear combination of the other columns.

Here, columns give you feature values. So, the columns will be linearly independent. If no feature can be obtained as a linear combination of other features, if one of the feature values can be predicted as a linear function of the other features, only then, the columns of A will not be linearly independent. But, that looks ridiculous because when we are looking at a linear classifier. If any one feature is a linear function of the other features, that feature is useless.

(Refer Slide Time: 35:40)



• The optimal W is a solution of $(A^T A)W = A^T Y$.

• When $A^T A$ is invertible, we get the optimal W as

$$\hat{W}^* = (A^T A)^{-1} A^T Y = A^\dagger Y$$

where $A^\dagger = (A^T A)^{-1} A^T$, is called the generalized inverse of A .

• The above W^* is the linear least squares solution for our regression (or classification) problem.



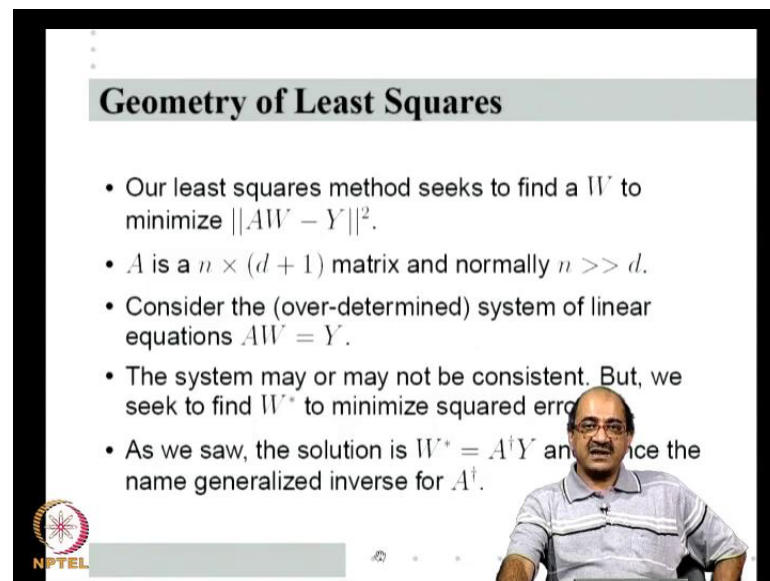
So, it is aaaa, it is reasonable to assume that our features are linearly independent. If I assume that the features are linearly independent, then A would have linearly independent columns. Hence, $A^T A$ would be invertible. As I just now said, this is a reasonable assumption, that is assumption, that $A^T A$ would be invertible because that simply means that my features are linearly independent. That is a very very reasonable assumption in either pattern recognition or regression problem because I am learning a linear model.

If one of the features is obtainable as a linear function of the other features then, it should not be putting that feature at all in new model. Now, as you seen in the optimal W satisfies $A^T A W = A^T Y$. So, if $A^T A$ is invertible, I can write my optimal W , which I will write as W^* as $(A^T A)^{-1} A^T Y$. This $(A^T A)^{-1} A^T$, that matrix is often called $A^\dagger$.

This sign is called dagger. This is called $A^\dagger$. So, I can write W^* as $A^\dagger Y$. This $A^\dagger$ is often called generalized inverse of A . Because A is a rectangular matrix, so it does not have an inverse. But, this $A^\dagger$ is called generalized inverse of A . We will currently see why it is called the generalized inverse of A . But, any case starting from my normal equations; if $A^T A$ is invertible, I can write W^* as $(A^T A)^{-1} A^T Y$.

This W^* is the linear least square solution to a regression or classification problem. Essentially, if I am given my feature vectors, the my example training patterns, that is I am given $X_i Y_i$ then, I can form the matrix A . I can form the vector Y . Then, I can calculate $A^T A$ whole inverse $A^T A$ into Y . That will give me W^* .

(Refer Slide Time: 37:16)



Geometry of Least Squares

- Our least squares method seeks to find a W^* to minimize $\|AW^* - Y\|^2$.
- A is a $n \times (d + 1)$ matrix and normally $n \gg d$.
- Consider the (over-determined) system of linear equations $AW = Y$.
- The system may or may not be consistent. But, we seek to find W^* to minimize squared error.
- As we saw, the solution is $W^* = A^\dagger Y$ and hence the name generalized inverse for $A^\dagger$.

NPTEL

So, given the training data, I can use this expression to compute W^* . So, this W^* is your linear least square solution. That is your final regression function of classifier. So, let us look a little bit into what the solution means. What is our least squares method trying to do? As we have seen that that function J is simply the norm square of the vector AW minus Y , right? Now, what is AA^T , is A n by $d + 1$ matrix. W is a $d + 1$ vector. Y is A n vector. Normally see, d is the dimension of the feature vectors. n is the number of examples.

Normally, we have many more examples that the dimension of feature vector has otherwise is really meaningless. If you give me, you know one point in $\mathbb{R}^2$ and ask me to find a classifier, it is hardly meaningful. I should get many many more points; certainly more than two points, so really say whether linearly separable or not. So, in most normal problems of function learning or classifier learning, n will be greater than d . n is often much larger than d . Now, let us look at this. AW is equal to Y system. This is a system of

linear equations. W is the d vector that is the unknown. $d + 1$ vector, that is the unknown.

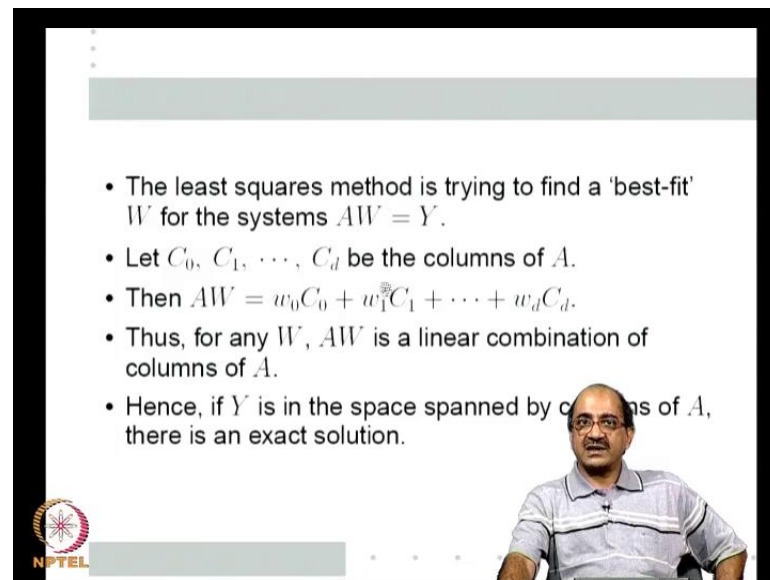
Now, A is a rectangular matrix, n by $d + 1$. So, in a linear system like this, n the rows of A , n could be the number of equations. The columns of A that is $d + 1$ same as the dimension W , those are the unknowns. Because n is much larger than d , this is an over determined system of equations. We have more equations than unknowns. Now, the system may be consistent, may not be consistent. As you know, when we given an over determined system of linear equations since the system is not consistent would not have a solution. Beforehand, we have no way of knowing their solution.

They may not even be because we are just fitting a function. The actual $X_i Y_i$ may not be on a line on a hyper plane. So, i may not be, i may not even be able to satisfy $W^T A X_i = Y_i$. X_i is equal to Y_i , for each i . So, the system may not be consistent. Whether it is consistent or not, we are asking find AW , but even if it does not exactly solve this; it is all set in the sense of minimizing errors over the training data. That is, find a least square solution of this.

So, if I cannot exactly satisfy AW minus Y equal to Y , what I am saying is at least finds me some W , so that the norm AW minus Y whole square is small. That is, the mean square error solution for this over determined system of equations. As we saw that solution is, one can write it as W^* is equal to $A^\dagger Y$. So essentially, I am solving this linear system with A as a rectangular matrix. So, really A does not have an inverse. But, if A had an inverse, I would have return W is equals to $A^{-1} y$.

Ultimately, when I did the least square solution of finding AW such that AW minus Y norm square is minimum, my solution is given by W^* is some matrix times A . That matrix is given the symbol $A^\dagger$. That is the reason why $A^\dagger$ is called the generalized inverse. It is like an inverse for a solving a linear system of equations, using the inverse of the coefficient matrix. We are essentially doing the same thing except in a mean square sense. The solution once again is W^* , is some matrix into the right hand side. So, that matrix is $A^\dagger$.

(Refer Slide Time: 40:56)



- The least squares method is trying to find a 'best-fit' W for the systems $AW = Y$.
- Let $C_0, C_1, \dots, C_d$ be the columns of A .
- Then $AW = w_0C_0 + w_1C_1 + \dots + w_dC_d$.
- Thus, for any W , AW is a linear combination of columns of A .
- Hence, if Y is in the space spanned by columns of A , there is an exact solution.

So, $A^\dagger$ is called the generalized inverse of A . Let us look a little more into the solution. So, the least squares method is trying to find the best fit W for the system of linear equations. AW is equal to Y . What does the best fit solution mean? Suppose columns of A are $C_0, C_1, \dots, C_d$. There are $d+1$ columns of A . Let us call them, $C_0, C_1, \dots, C_d$. Now, when I multiply any matrix on the right with a vector; essentially, what I get is another vector. It is a linear combination of the columns of the matrix whose linear combination coefficients; are the coefficient of the vectors.

So, any AW is simply, if W has components, $w_0, w_1, \dots, w_d$, then, AW will be written as $w_0C_0 + w_1C_1 + \dots + w_dC_d$, so the AW vector is simply $w_0C_0 + w_1C_1 + \dots + w_dC_d$ vector. So, every AW is essentially a linear combination of the column vectors of A . So, when I search over all possible W , all I can get out of AW are linear combinations of columns of A . So, for any W , AW will be a linear combination of columns of A .

So, when can I solve this system exactly? If Y is in the space spanned by columns of A ; space spanned by the columns of A is called the column space. What do we mean by space spanned by the columns of A ? If you consider this, the vector space obtained as all possible linear combinations of columns of A . All vector that can be obtained as linear combination of columns of A , that is called the column space of A .

(Refer Slide Time: 42:51)



- Otherwise, we want the projection of Y onto the column space of A .
- That is, we want to find a vector Z in the column space of A that is closest to Y .
- Any vector in the column space of A can be written as $Z = AW$ for some W .
- Hence we want to find Z to minimize $\|Z - Y\|^2$ subject to the constraint that $Z = AW$ for some W .
- That is the least squares solution.

PR NPTEL course - p 04/17

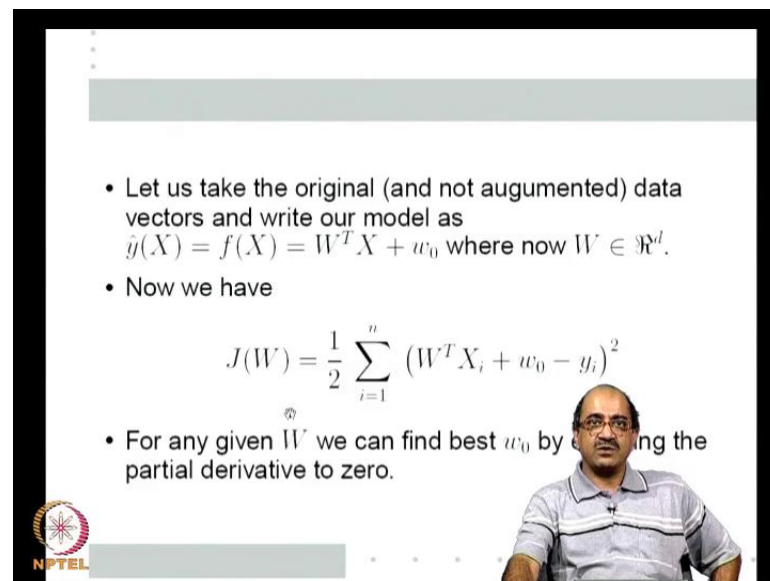
So, if the vector Y happens to be in column space of A , then we have an exact solution. There will be some linear combination of columns of A that will give me the Y . That linear combination is the W vector, but suppose, we do not Y is not in the column space of A , then what is the best we can do? Then, we are asking project Y onto the column space of A . What do you mean by project Y on to the column space of A ?

That is, we want to find a vector Z , which is in the column space of A and that is closest to Y . Because AW can only give me vectors in the column space of A , I cannot now satisfy AW is equal to Y anyway. So, the best I can do is to find a vector in the column space of A that is that has least distance to Y . That is the difference between that vector and Y is the smallest. So, that is what is meant by project Y on to the column space of A .

How do I find such a Z ? To find such as Z now, any vector in the column space of A can be written as AW for some W . Thus, that is what this definition of column space of A . So, finding a vector Z in the column space of A that is closest to Y is same as we want to find Z to minimize $\|Z - Y\|^2$. We want to find Z to minimize $\|Z - Y\|^2$. That is the distance between Z and Y under the constraint that Z has to be AW , for some W because, Z has to be in the column space. So essentially, projecting Y on to the column space of A is same as minimizing.

Now, because Z has to be AW , I can put this as AW . So, minimize AW minus Y norm square. So, minimizing AW minus Y norm square is same as projecting Y on to column space of A , that is the least square solution. A matter of fact, what it means is, if I have A vectors of space defined through some vectors X_1, X_2, X_n . So, if I, if there is a vector Y and I want to project it on to the space spanned by X_1, X_2, X_n , then, I can organize those X 's as rows of matrix A . Take A transpose, A whole inverse, A transpose and multiply post multiply with this. Then, that will be the projection.

(Refer Slide Time: 45:42)



- Let us take the original (and not augmented) data vectors and write our model as $\hat{y}(X) = f(X) = W^T X + w_0$ where now $W \in \mathbb{R}^d$.
- Now we have

$$J(W) = \frac{1}{2} \sum_{i=1}^n (W^T X_i + w_0 - y_i)^2$$

- For any given W we can find best w_0 by finding the partial derivative to zero.

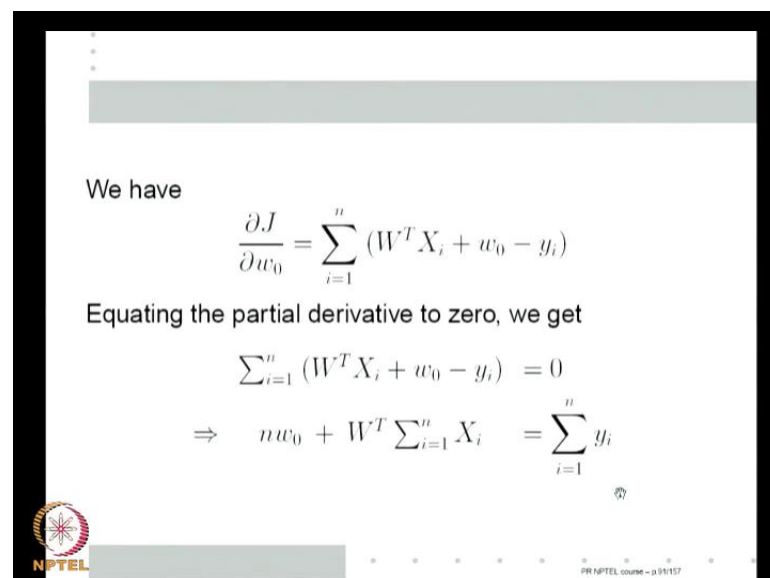
So, as a matter of fact, this A transpose, A whole inverse, A transpose is also called a projection matrix. It projects Y on to the column space of A . So, least square solution can be looked at as obtaining an approximate solution to an over determinant system of equations, which may not be consistent hence, essentially projecting this vector Y on to the column space of A . The projection matrix is given by A transpose, A whole inverse, A transpose. Now, let us look at a few more things.

We have been we have been saying that we can either use augmented representation or representation of the linear function as W transpose X . So, let us ask what the extra constant does. See, essentially if I have this $f(X)$ defined as W transpose X plus w_0 is not a linear function, in the sense $f(X_1 + X_2)$ is not equal to $f(X_1) + f(X_2)$ because of this w_0 . Of course, is called an affine function. It is just a constant plus a linear function. We artificially made into a linear

function, not artificially. We made it a linear function by simply working in the augmented space. But, suppose you want to work in the original space and ask what is this extra constant doing? Why did I need this constant?

Why couldn't I have simply done $W^T X$? To see that, let us go back to what we minimize by this time. Let us not work in the augmented space, but work in the original space. So, I will write $J = \frac{1}{2} \sum_{i=1}^n (W^T X_i + w_0 - y_i)^2$, so, minimizing this, minimizing this. Now, with respect to both w and w_0 , I should have put w_0 also here. But, anyway that does not matter.

(Refer Slide Time: 47:17)



We have

$$\frac{\partial J}{\partial w_0} = \sum_{i=1}^n (W^T X_i + w_0 - y_i)$$

Equating the partial derivative to zero, we get

$$\sum_{i=1}^n (W^T X_i + w_0 - y_i) = 0$$

$$\Rightarrow n w_0 + W^T \sum_{i=1}^n X_i = \sum_{i=1}^n y_i$$

So, let us ask if I fix AW vector, what is the best w_0 , I can get? For a fixed w , if I want to get the best w_0 then, I have to find the partial derivatives of J with respect to w_0 and equate it to 0. That will give me the best w_0 . So, let us look at that best w_0 . So, what is $\frac{\partial J}{\partial w_0}$? This is J . If I differentiate it with respect to w_0 , I get this $\frac{1}{2}$ cancel summation; this into $W^T X_i + w_0 - y_i$, because, w_0 has no other coefficient; that is the whole derivative.

So, $\frac{\partial J}{\partial w_0}$ is given by this. So, if I equate this to 0, what do we get? This is equal to 0. $\frac{1}{2} \sum_{i=1}^n (W^T X_i + w_0 - y_i) = 0$. Now, you take the summation inside. Get n times w_0 plus W^T times summation X_i . You take y_i on this side, is summation y_i .

(Refer Slide Time: 48:00)



This gives us

$$w_0 = \frac{1}{n} \sum_{i=1}^n y_i - W^T \left(\frac{1}{n} \sum_{i=1}^n X_i \right)$$

- Thus, w_0 accounts for difference in the average of $W^T X$ and average of y .
- So, w_0 is often called the bias term.



So, I can solve this to get w_0 . That gives me w_0 is $\frac{1}{n} \sum y_i - W^T \left(\frac{1}{n} \sum X_i \right)$. See this, n will come this side. If I divide this by n , is $\frac{1}{n} \sum y_i - W^T \left(\frac{1}{n} \sum X_i \right)$, which I can write it as $\frac{1}{n} \sum y_i - W^T \left(\frac{1}{n} \sum X_i \right)$. What is this $\frac{1}{n} \sum X_i$? This is the average of X_i . This is the average of y_i . So essentially, w_0 accounts the difference in the average of $W^T X$ and average of y .

If I ultimately wanted to represent y as $W^T X$, assuming essentially thinking of y and X as random, what it should mean? This expected value, y should be equal to W^T expected value of x . So, if expected value of y is the average of y is not equal to W^T expected value of X . Then, I have to account for this difference because that cannot be accounted for my linear function $W^T X$. That difference is accounted for by w_0 . That is the reason, why we need a constraint essentially.

(Refer Slide Time: 49:32)



- We have taken our linear model to be

$$\hat{y}(X) = f(X) = \sum_{j=0}^d w_j x_j$$

- As mentioned earlier, we could instead choose any fixed set of basis functions ϕ_i .
- Then the model would be

$$\hat{y}(X) = f(X) = \sum_{j=0}^{d'} w_j \phi_j(X)$$

PR 10809 Course - p 37/117

If some of these averages can be matched, otherwise then we do not need the constants. So, w_0 is often called the bias term because it is time to adjust between the averages of, average of y to the average of $W^T X$. So, that is why whenever we write a general linear model as $W^T X$ plus w_0 ; the w_0 is called the bias term. A few more things about our linear model. We are taking our linear model once again. Let us put w_0 back into this and take X_0 to be 1, so that, we are working back in the augmented space. So, our model is $\hat{y}(X) = \sum_{j=0}^d w_j x_j$.

Now, as mentioned earlier, we do not have to use x_j . We can use any fixed sets of basis functions ϕ_i . What does that mean? The model could have been $\hat{y}(X) = \sum_{j=0}^{d'} w_j \phi_j(X)$. It could be any number of ϕ functions; need not naturally be equal to the dimension of x_j is equal to 0 to d' . $w_j \phi_j(X)$ can be written as $\sum_{j=0}^{d'} w_j \phi_j(X)$, this d' becomes d' . So, this as we seen is also a linear model.

(Refer Slide Time: 50:53)



- We can use the same criterion of minimizing sum of squares of errors.

$$J(W) = \frac{1}{2} \sum_{i=1}^n (W^T \Phi(X_i) - y_i)^2$$

where $\Phi(X_i) = (\phi_0(X_i), \dots, \phi_d(X_i))^T$.

- We want the minimizer of $J(\cdot)$.

PR NPTEL course - p 09/107

A matter of fact, if I have chosen this as the model using any functions phi exactly the same method that we have seen so far will work, right? Why I can now take the same criterion minimizing sum of squares of zeroes?

(Refer Slide Time: 51:47)



- We can learn W using the same method as earlier.
- Thus, we will again have

$$W^* = (A^T A)^{-1} A^T Y$$

- The only difference is that now the i^{th} row of matrix A would be

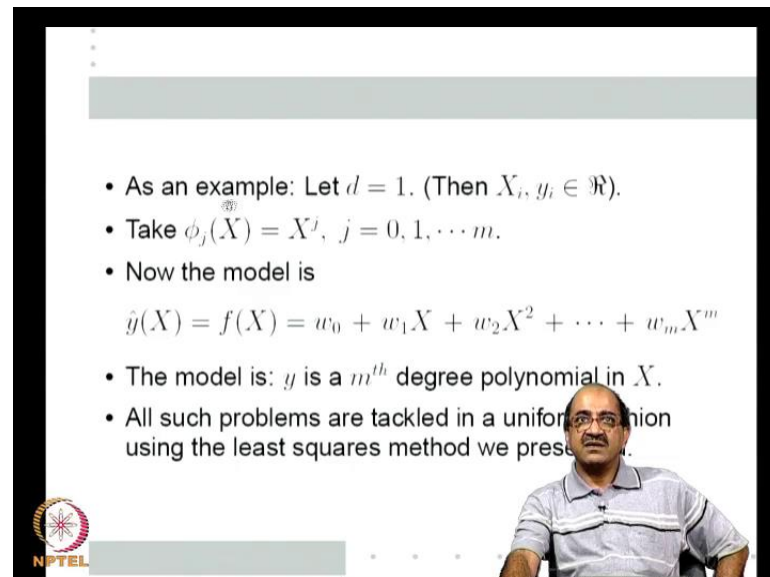
$$[\phi_0(X_i) \ \phi_1(X_i) \ \dots \ \phi_d(X_i)]$$

PR NPTEL course - p 102/107

That is half W transpose, say now, I call it ϕX . See, the the the the the sum is $w_j \phi_j x$ that is $w_0 \phi_0 x$, $w_1 \phi_1 x$, $w_2 \phi_2 x$ and so on. So, if I write it as w transpose ϕx I, this vector ϕx should have components $\phi_0 x$ $\phi_1 x$ $\phi_2 x$ and so on. So, I can write j as i is equal to 1 to n w transpose capital ϕ of x i minus y_i square. So, if I put capital

ϕ of X_i in place of x_i , then exactly the same algebra goes through. We want the minimizer of J . The minimizer of J can be learnt using the same method as earlier. We will once again get a transpose A , inverse A , transpose y . The only difference is that now, the rows of A earlier, the i -th row of A is x_i .

(Refer Slide Time: 52:19)



- As an example: Let $d = 1$. (Then $X_i, y_i \in \mathbb{R}$).
- Take $\phi_j(X) = X^j, j = 0, 1, \dots, m$.
- Now the model is

$$\hat{y}(X) = f(X) = w_0 + w_1 X + w_2 X^2 + \dots + w_m X^m$$
- The model is: y is a m^{th} degree polynomial in X .
- All such problems are tackled in a uniform fashion using the least squares method we presented.

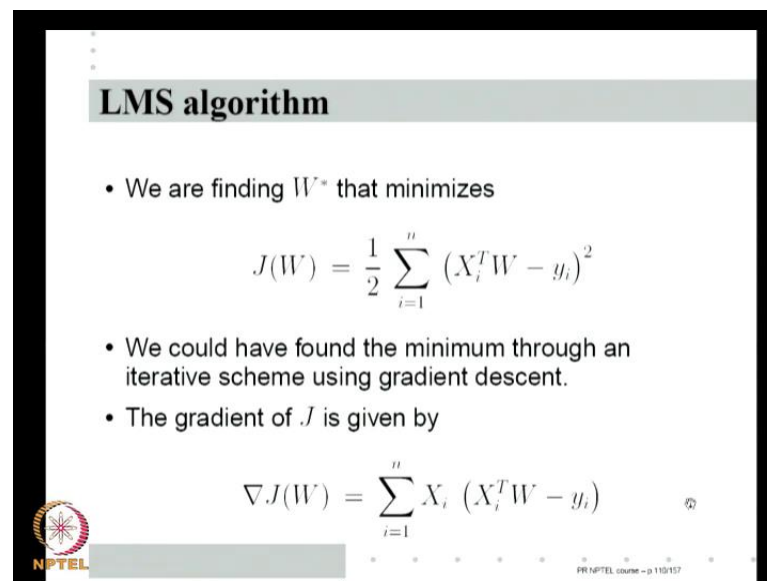
Now, instead of that, it will be capital phi of X_i , which is ϕ_0 of x_i , ϕ_1 of x_i , all the way upto $\phi_{d'} \phi_{d'}$ of x_i . So, the only difference is that i -th row of matrix A would be now the capital phi of x_i , except for this exactly same thing go. So, let us look at one example of this generality. Let us take d is equal to 1. 1 dimension case, so that means both x_i and y_i are belonging to $\mathbb{R}$. So, my training data said, this is simply real numbers pairs of real numbers, x and y , $x_1 y_1, x_2 y_2, x_n y_n$. You want to find a function relation with x and y .

What is this problem? This is the familiar curve fitting problem that all of you would have had. So, you want to put y is equal to $f(X)$. Then, you find the least squares curve fitting. The least squares curve fitting is you take $f(X_i) - y_i$ whole square, sum over i and minimize that. You, most of you would have done the linear fit that is, you do $y_i - m x_i - c$ whole square, differentiate it with a fine $(())$, which is equal to m and c . Then, find the optimal m and c . That is the least square's curve fitting. So, the problem that we have is a least square's curve fitting.

Let us say, we want a generic least square's curve fitting what kind of curve. Let us take ϕ_j of X to be X to the power J . J is equals to 0 to n . So, what does the model mean now? $\hat{y}_x$ is f of X , $w_0 \phi_0$ of x ϕ_0 of x is 1 and w_1 of ϕ_1 x ϕ_1 x is x to the power 1 and so on. So, $\hat{y}_x$ is a w_0 plus $w_1 x$ plus $w_2 x^2$ plus $w_m x^m$. So, what is this? This is nothing but a polynomial index. So, if I want to model y , as the m th degree polynomial on x , I can simply think of it as a linear least square's problem say, m th degree polynomial is settling at an linear curve.

But, I can think of it as a linear least squares problem by simply working with ϕ_j of x is equals to x^j . So, the model is, I want to model y as an n th degree polynomial on x . I want to find the best coefficients for the polynomial w_0, w_1, w_m . I can solve it using the same linear method that I have of this, the generalized inverse method.

(Refer Slide Time: 54:48)



LMS algorithm

- We are finding W^* that minimizes

$$J(W) = \frac{1}{2} \sum_{i=1}^n (X_i^T W - y_i)^2$$

- We could have found the minimum through an iterative scheme using gradient descent.
- The gradient of J is given by

$$\nabla J(W) = \sum_{i=1}^n X_i (X_i^T W - y_i)$$

NPTEL PRE NPTEL course - p 110107

So, the all such problems are tackled in a uniform fashion using the least squares method that we proposed here. So, this is what essentially the generality that you get with respect to ϕ means. We are finding. Let us just look at a little more structure in this. We introduce one more algorithm what is called the LMS algorithm. We are finding W^* to minimize this. To minimize any function, I can use gradient descent. If I want to minimize a function f of X where X is a vector then, I can do an iterative algorithm X of k plus is X of k minus θ times gradient of f of X k .

(Refer Slide Time: 55:43)



• The iterative gradient descent scheme would be

$$W(k+1) = W(k) - \eta \sum_{i=1}^n X_i (X_i^T W(k) - y_i)$$

• In analogy with what we saw in Perceptron algorithm, this can be viewed as a 'batch' version.

• We use the current W to find the errors on all training data and then do all the 'corrections' together.

• We can instead have an incremental version of this algorithm.

PR NPTEL course - p 114/117

That algorithm converges to the local minimum. So, I can find minimizer using gradient descent. So, we could have found the minimum through an iterative scheme using gradient descent. What is the gradient of J ? Gradient of J is once again the 2 cancels with half. So, I get X_i into $X_i^T W - y_i$. So, what will be an iterative gradient descent scheme for this? So, $W(k+1)$ is equal to $W(k)$ minus some step size η into my gradient is this. i is equal to one to n X_i into $X_i^T W - y_i$.

Now, if you look at it, in analogy with perceptron, it looks like a batch version of the algorithm. Given a current W , these are the errors I am making on each of the training samples. On i training sample, my error is $X_i^T W - y_i$. So, I calculate the gradient with respect to all the errors and then do correction together. So, that is a batch version. I could as well have done it incrementally. So, here we are first using the current W . We find errors on all the training data. Then, do all the corrections together instead of that we could do an incremental version of the algorithm.

(Refer Slide Time: 56:35)



- For the incremental version, at each iteration we pick one of the training samples. Call this $X(k)$.
- The error on this sample would be $\frac{1}{2}(X(k)^T W(k) - y(k))^2$.
- Using the gradient of only this term, we get the incremental version as
$$W(k+1) = W(k) - \eta X(k) (X(k)^T W(k) - y(k))$$
- This is called the LMS algorithm.

PR: NPTEL course - p 110/117

What will the incremental version mean? In the incremental version, at each iteration, we pick one of the training samples $(X(k), y(k))$ and correspondingly $Y(k)$. Then, the error on this sample will be $X(k)^T W(k) - y(k)$ whole square. So, using only that gradient in the gradient descent, I get $W(k+1) = W(k) - \eta X(k) (X(k)^T W(k) - y(k))$. That is the gradient. This is called the LMS or the least mean square algorithm.

(Refer Slide Time: 57:14)



- In the LMS algorithm, we iteratively update W as
$$W(k+1) = W(k) - \eta X(k) (X(k)^T W(k) - y(k))$$
- Here $(X(k), y(k))$ is the training example picked at iteration k and $W(k)$ is the weight vector at iteration k .
- We do not need to have all training examples together with us. We can learn W from a stream of examples without needing to store them.
- If η is sufficiently small this algorithm also converges to the minimizer of $J(W)$.

PR: NPTEL course - p 120/117

So, LMS algorithm is essentially the linear least squares method, where I minimize J using an iterative gradient descent but, using an incremental algorithm. So, that is the LMS algorithm. We update W iteratively like this. Here, $X_k Y_k$ is the training sample with data iteration k . W_k is the weight vector at iteration k . So, like in the perceptron case because it is incremental, we do not need to store all the examples. We can have examples coming in a stream.

We can do from there, if it has sufficiently small. This algorithm also converges to the minimizer of J (()). Such an algorithm is called the LMS algorithm. So, next class, we will start with the LMS algorithm and see. LMS algorithm was also as old as perceptron. It also came from a simple model called Adaline. So, we look at the LMS algorithm, various ramifications of LMS algorithm and then look at more structure in the linear least squares method.

Thank you.