

Lecture - 27
Decision Trees for Big Data Analytics

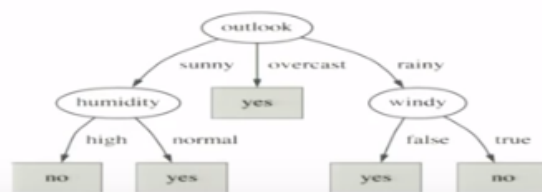
Decision trees for big data analytics.

Refer Slide Time (0:18)

Preface

Content of this Lecture:

- In this lecture, we will discuss Decision Trees for Big Data Analytics and also discuss a case study of medical application using a Decision Tree in Spark ML



Pre-phase. Content of this lecture in this lecture we will discuss decision trees for big data analytics.

We will also cover a case study of a medical application using decision trees in SPARC ML.

Refer Slide Time (0:37)

Decision Trees

Before, that let us understand about the basics of decision trees.

Refer Slide Time (0:42)

Predict if Sachin will play cricket

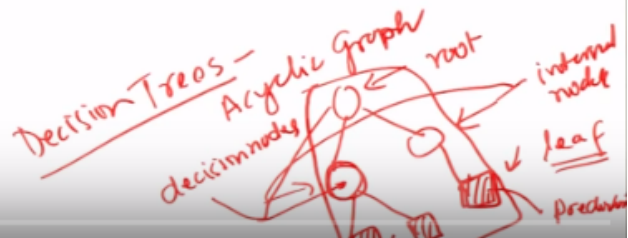
Consider the following data set
our task is to predict if Sachin is going to play cricket on a given day. We have observed Sachin over a number of days and recorded various things that might influence his decision to play cricket so we looked at what kind of weather it is. Is it sunny or is it raining humidity is it high as normal is it windy So this is our training set and these are the examples that we are going to build a classifier from.

Training examples: 9 yes / 5 no

Day	Outlook	Humidity	Wind	Play
D1	Sunny	High	Weak	No
D2	Sunny	High	Strong	No
D3	Overcast	High	Weak	Yes
D4	Rain	High	Weak	Yes
D5	Rain	Normal	Weak	Yes
D6	Rain	Normal	Strong	No
D7	Overcast	Normal	Strong	Yes
D8	Sunny	High	Weak	No
D9	Sunny	Normal	Weak	Yes
D10	Rain	Normal	Weak	Yes
D11	Sunny	Normal	Strong	Yes
D12	Overcast	High	Strong	Yes
D13	Overcast	Normal	Weak	Yes
D14	Rain	High	Strong	No

Refer Slide Time (0:46)

Decision Trees



So, decision tree is a tree, which is nothing but an acyclic graph, which has the root node and other nodes. This is called the root node of a tree and these are all called internal nodes of a tree and this is called the leaf. So, the square boxes represent the node, which is called a leaf node in the decision tree. The predictions are there at the leaf nodes, and the internal nodes and the root nodes they are called the decision nodes. Now, given a data set we are going to construct this particular decision tree, using the decision tree algorithm, and traversing the tree we can predict for a given query. So, whenever a new data set, we get we will be able to handle this particular data set for doing the predictions. So, details we are going to see in this discussion and, then how this decision tree is going to be useful for doing the big data analytics, also we will cover through an example.

Refer Slide Time (2:39)

Predict if Sachin will play cricket

So on day 15 if it's raining the humidity is high it's not very windy so is Sachin going to play or not and just by looking at the data it's kind of hard to it's it's kind of hard to decide right because it's you know some days when it's raining Sachin is playing other days Sachin is not playing sometimes he plays with strong winds sometimes he doesn't play with strong winds and with weak winds life so what do you do how do you predict it and the basic idea behind decision trees is try to at some level understand why Sachin plays. so this is the only classifier that that will cover that tries to predict Sachin playing in this way.

Training examples: 9 yes / 5 no

Day	Outlook	Humidity	Wind	Play
D1	Sunny	High	Weak	No
D2	Sunny	High	Strong	No
D3	Overcast	High	Weak	Yes
D4	Rain	High	Weak	Yes
D5	Rain	Normal	Weak	Yes
D6	Rain	Normal	Strong	No
D7	Overcast	Normal	Strong	Yes
D8	Sunny	High	Weak	No
D9	Sunny	Normal	Weak	Yes
D10	Rain	Normal	Weak	Yes
D11	Sunny	Normal	Strong	Yes
D12	Overcast	High	Strong	Yes
D13	Overcast	Normal	Weak	Yes
D14	Rain	High	Strong	No
New data:				
D15	Rain	High	Weak	?

Let us see, simple data set and let us see what kind of predictions this decision tree will be able to make. So, to understand in a very simple manner. Let us consider this data set add a query. So, the queries, that we want to predict, if searching will play the cricket or not, to answer this particular query, that means to get the prediction whether searching will play the cricket or not. We are going to be have observed the details of Sachin playing and also the weather conditions over 15 days. So, day 1 to day 50 14 or 14days. Wherein we have considered, that what is the outlook of the weather if it is sunny overcast raining and the humidity of every day, that is whether it is high normal or the wind speed whether it is weak or a strong and we have also seen, that in these conditions whether, that player such in highest played or not. So, on day one it has not played he has not played day three he has played well. The weather was outlook was overcast humidity was high and wind was weak and so on. So, we have observed the 14 days of these different parameters and also, we have seen whether the player Sachin has played under these conditions or not. Now, using this particular data set, can we predict whether Sachin will play to the days, which are a new day or a next day or not so decision tree. Let us see, how it is going to be useful? So, consider this data set and our task is to predict if the Sachin is going to play the cricket on a given day. So, we have observed the playing the game playing of Sachin our number of days. In this example we have monitored, or we have observed 14 days and recorded various conditions that might influence his decision to play the cricket. So, under under different weather conditions. So, is it sunny, or is it raining, or it is having humidity, and then in that case or it is having a normal windy temperature what conditions are required and necessary and to predict whether certain will play or not? Even for this is a small or a simple data set this becomes quite complex. So, let us see how we are going to simplify using decision tree for this particular purpose. Now, one thing is, that this particular data set is basically observation of the last 14 days. So,

that becomes the training set and using these examples can be predict this particular decision of a search in whether he will play the game or not? So, that means the query will be that on day 15th if outlook is, raining humidity is high and wind is weak, whether such an will play or not So, from this data set also it seems, that it's not easy to get or understand this entire data set and come out with the decision to predict or not? But, let us understand this particular data set, because decision tree is the technique, which will understand these different parts of the data set or the behaviour of that player Sachin it will understand and based on, that it will be just the predictions. So, on day 15 if it is raining and the humidity is high and the wind speed is not windy slow or a week wind is capacity, then weather Sachin is going to play or not, and by looking to this particular data and this kind it will become very hard to decide. The right kind of predictions, whether certain is playing in other days on one of these parameters and how this kind of condition, whether he will be playing or not? It's very difficult to predict.

Refer Slide Time (7:44)

Predict if Sachin will play cricket

- Hard to guess ✓
- Try to understand when Sachin plays ✓
- Divide & conquer:
 - Split into subsets
 - Are they pure ? (all yes of all no)
 - If yes: stop
 - If not: repeat
- See which subset new data falls into

Training examples: 9 yes / 5 no

Day	Outlook	Humidity	Wind	Play
D1	Sunny	High	Weak	No
D2	Sunny	High	Strong	No
D3	Overcast	High	Weak	Yes
D4	Rain	High	Weak	Yes
D5	Rain	Normal	Weak	Yes
D6	Rain	Normal	Strong	No
D7	Overcast	Normal	Strong	Yes
D8	Sunny	High	Weak	No
D9	Sunny	Normal	Weak	Yes
D10	Rain	Normal	Weak	Yes
D11	Sunny	Normal	Strong	Yes
D12	Overcast	High	Strong	Yes
D13	Overcast	Normal	Weak	Yes
D14	Rain	High	Strong	No

New data:
D15 Rain High Weak ? ✓

Let us see how in decision tree we are going to do this kind of predictions? So, as we have seen it is very hard to guess, whether Sachin will play on day 15 and but for this we have to understand, when such in play? What are the condition when certain will be when Sachin is playing? So, there are 19 different samples when Sachin has played and under different conditions 9 and these 9 different conditions, when Sachin has played is required as, the training to understand when he has played and where there are 5 different condition, when he has not played in those weather conditions. Now, we will to understand this behaviour of a player Sachin out of this data set, which is an observation of last 14 days. We can do this by splitting into the subset, into subsets of different classes and then we can extract this information and the the decision information we can then basically use this information for predicting a new data set.

Refer Slide Time (9:07)



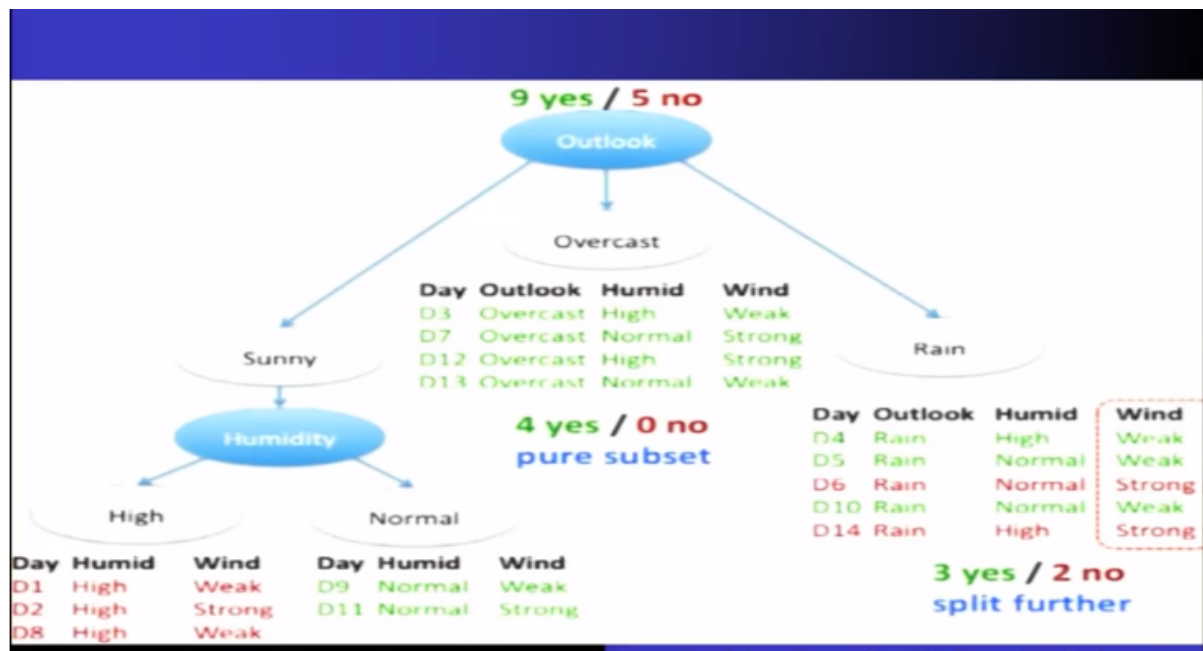
So, let us understand more details of this entire data set, which is given. So, if we classify this data set based on the outlook parameter, then there are 3 different outlook condition. One is sunny, than overcast and raining, and these subsets which we are going to get in these different conditions. So, we basically have obtained three different subsets. Now, out of these three different subsets we see, that in when the outlook is overcast, we can see, that here. There are 4 different entries, or 4 different days are there, when this becomes yes, and there is no entry which is having 0 no's. So, hence, this kind of subset is called a pure subset, which has all yes. When the outlook sunny, then we can see this particular subset, which has three no's and two yes. So, this kind of subset is called impure subset. That means it is mixed, whether to are saying yes when the outlook is sunny, then Sachin has played in two of these conditions, when the humid is normal and the event is a weak or strong and in three different cases, he has in played. So, this kind of subset is called impure subset. Similarly, when the outlook is rain, then there are three different cases, when Sachin has played and two different cases, when Sachin has not played. So, hence, this becomes an impure subset. Now, we can further split these impure subsets, and pure subsets need not to be splitted further.

Refer Slide Time (11:28)



So, now, we can see, that we can split so, that now, we have splitted this outlook sunny, based on the humidity. When the humidity is high or normal. So, what we can see here, that when the humidity is high, all three cases the Sachin has not played. So, four no's and zero yes. Hence, this is called a pure subset. Whereas when the humidity is normal, then there are two different in this subset, that means two yes is there, and there is no zero no's. Hence, this is also called a pure subset. So, when it is a pure subset, then decision can be easily made. For example, when the outlook is sunny and the humidity is normal, then Sachin has always played. So, whenever a data of this category comes, we can decide about this. So, here also when under this split humidity. We have obtained both the pure subsets of drop out of this split of sunny. Now, then let us see, the rain when the outlook is rain, then it has basically a mixed or impure subset which can be splitted further.

Refer Slide Time (13:00)

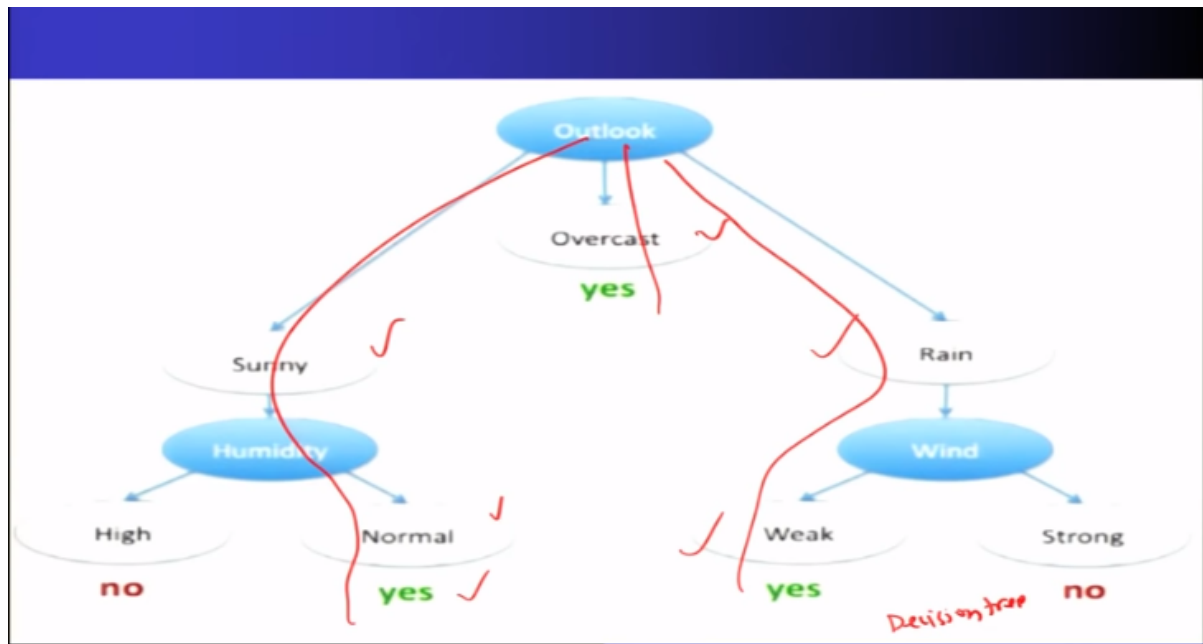


Refer Slide time (13:02)



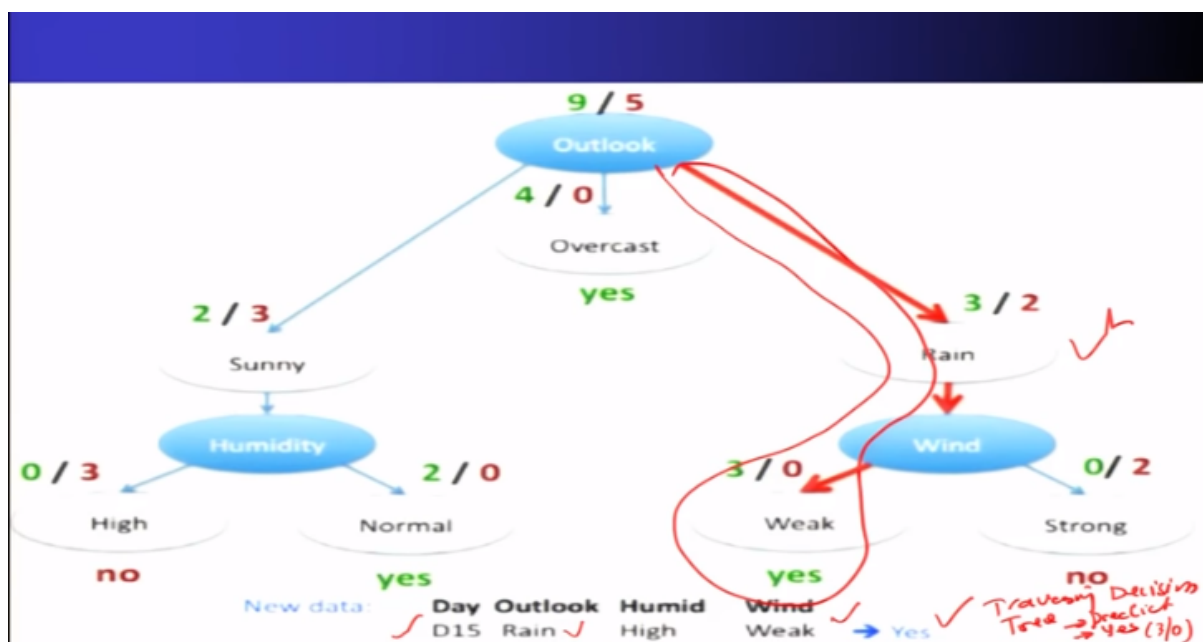
So, if we split further, on this wind what we obtain is, that when the outlook is rain, and the wind is weak, then we have got three different subsets. In this subset the three yes is there. Three times he has played and zero times it he has not played. So, this is also called a pure subset. Now, as far as when the wind is strong, then we have obtained a subset, where the zero yes is there, and two no's are there, and this is also called pure subsets. Now, there is no further division or split required and we have constructed decision tree of this use case.

Refer Slide Time (14:01)



Now, if we summarize this particular decision, that decision tree which is constructed out of, that dataset. We have obtained, that if when the humidity is normal, and when outlook is sunny, then always, that player Sachin hence, played where the outlook is overcast in Sachin has played and, when the wind is weak, and the outlook is rain, then also Sachin has played. So, there are three different cases, when Sachin has played, and which is shown over here. In all other cases Sachin has not played.

Refer Slide Time (14:56)



Now, not only this particular decision tree has yes and no cases. But, it also has how many cases which are having yes. So, it will add to the confidence on this particular decision using these

parameters. So, let us see, the query, which was given to us on day 15, that if the outlook is rain, and outlook is rain and humidity is high and wind is weak. When the outlook is rain and and the wind is weak. Obviously, it is going to play in this particular case. So, we are going to say yes. Hence, the decision is yes in this particular case using the traversing of the decision tree. So, by traversing the decision tree, we can predict, that Sachin will play, and the confidence with which we are going to predict is three by zero. So, that means with the high confidence this particular prediction can be made.

Refer Slide Time (16:23)

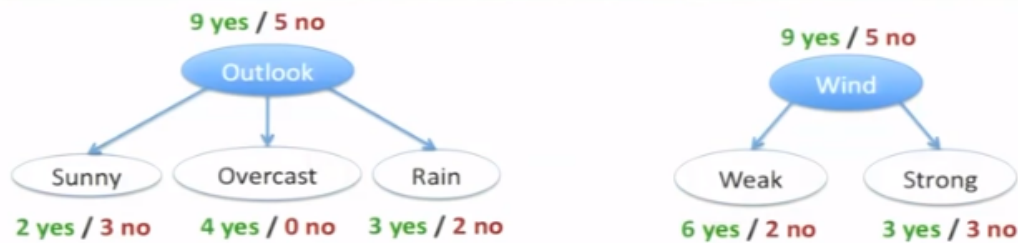
ID3 Algorithm
for Decision Tree

- In decision tree learning, ID3 (Iterative Dichotomiser 3) is an algorithm invented by Ross Quinlan used to generate a decision tree from a dataset. ID3 is the precursor to the C4.5 algorithm, and is typically used in the machine learning and natural language processing domains.
- Split (node, {examples}):
 1. $A \leftarrow$ the best attribute for splitting the {examples}
 2. Decision attribute for this node $\leftarrow A$
 3. For each value of A, create new child node
 4. Split training {examples} to child nodes
 5. For each child node / subset:
 - if subset is pure: STOP
 - else: Split (child_node, {subset})
- Ross Quinlan (ID3: 1986), (C4.5:1993)
- Breimanetal (CaRT:1984) from statistics

Now, as far as the decision trees are concerned, we are going to discuss one such algorithm which is called ID3algorithm for decision tree. So, this algorithm was invented by Ross Quinlan used to generate the decision tree from the data set ID3 is the precursor to see four point five algorithm and is typically used in the machine learning domain. So, let us see, how the split of a node is done in ID3. Now, let us consider A, which is assigned to the best attribute for the splitting, and the decision attribute for this particular node is assigned from A. For each value of A let us create a new child node. Now, split the training examples to the child node and for each child node or a subset if the subset is pure, then stop splitting further else we have to split the node child into the subsets further. Now, Ross - Lin queuing plan Ross Quinlan which has given this ID3 in 1986 and also C 4.5 in 1993 and Brien man Brein man natal and Briemanatal has given CaRT in 1984 from statistics about this DC entry algorithm.

Refer Slide Time (18:17)

Which attribute to split on ?



- Want to measure "purity" of the split
 - More certain about Yes/No after the split
 - pure set (4 yes / 0 no) $\Rightarrow$ completely certain (100%)
 - impure (3 yes / 3 no) $\Rightarrow$ completely uncertain (50%)
 - Can't use $P(\text{"yes"} \mid \text{set})$:
 - must be symmetric: 4 yes / 0 no as pure as 0 yes / 4 no
- pure subsets (100%) correct*

Now, the question is which attribute to split on. So, let us see, that we have seen, that there are three different attributes in that dataset. So, which one we are going to use to split the attributes on, that particular question is going to be, that is the split decision which attribute is split on is a decision which we are going to see how decision tree is going to take here in this case. Now, here let us see, that it wants to measure, the purity of the split, and this is the way it is going to take the decision, of which attribute to split on. So, based on the purity of the split the decision is to be evaluated. So, more certain about yes or no after the split. So, for example if we have a pure set where for yes is there, and zero no's are there, then we can say with hundred percentage certainty, that this is a pure and pure set and the decision is very certain. Similarly, if the set is impure, where three yes is there and three no's are there, then the decision which we are going to take is having completely uncertain, that is fifty percentage chance of making the decision correct, is basically is there. So, this way we are going to measure using how many yes is there whether the set is pure or impure and this is going to be the measure to go for the split. Now, so, this particular must be the use of this probability must be the symmetric, that is four yes and zero knows as pure as, zero yes and for no's. So, that means whether it is all yes or whether it is all no's are called pure subsets. Wherein the decision is hundred percentage certain to be correct. Whereas if the set is impure then it is uncertain, that the decision is correct. So, we are going to evaluate this measure of purity to split on.

Refer Slide Time (21:05)

Entropy

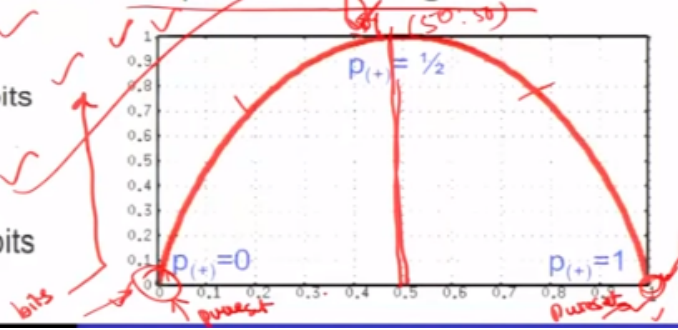
- Entropy: $H(S) = -p_{(+)} \log_2 p_{(+)} - p_{(-)} \log_2 p_{(-)}$ bits
 - S ... subset of training examples
 - $p_{(+)} / p_{(-)}$... % of positive/negative examples in S
- Interpretation: assume item X belongs to S
 - How many bits need to tell if X positive or negative

- Impure (3 yes / 3 no):

$$H(S) = -\frac{3}{6} \log_2 \frac{3}{6} - \frac{3}{6} \log_2 \frac{3}{6} = 1 \text{ bits}$$

- Pure set (4 yes / 0 no):

$$H(S) = -\frac{4}{4} \log_2 \frac{4}{4} - \frac{0}{4} \log_2 \frac{0}{4} = 0 \text{ bits}$$



Now, this decision tree uses different measures for make the decision whether to split or not? One such way is called entropy, that we are going to discuss. How using entropy, we can decide whether to split further or to not to split? So, entropy is represented by HS here and there are two variables. Let us say P plus and P minus. So, that is minus P plus log to the base 2 P plus minus P minus log to the base P minus I this is nothing but the number of bits. So, if we calculate the entropy of a particular set of training examples, then we will get the bits. So, this the interpolation of the bits we are going to use to infer, whether to split or not how this is all done we are going to see now. Now P plus and P minus percentage of the positive and negative examples in s we are going to evaluate. Now interpolation is let us assume an item X belongs to a set S and we have to ensure or basically estimate how many bits are needed to tell if X is positive or the negative? So, for example if it is impure set of three three yes and three no's, then if we calculate the entropy of this impure set S, then it comes out to be what bit? So, one bit if it is the entropy is one bit, then it is highly uncertain for the decision, whether it is yes or no. So, the number of bits which basically will represent it here it is 1. So, 1 means it is highly uncertain in this particular case. Now, if it is a pure set, let us hear, if it is four years and zero no's and we want to find out the entropy of that set. So, entropy will be zero in this case. So, if the set is pure the entropy will be zero bits. So, zero bits means the decision is very certain, and if it's uncertain, that is highly uncertain, then the bit will be one. Now, and rest of the decisions. are in between zero and one, that is what we have seen. Now, so, this particular diagram shows, that if the number of bits is one if the number of bit is one, here in this case if it is, so, if if it is one then the the decision will be half fifty percent chance, that the decision is going to be correct fifty percent chance. So, it will be a fifty fifty, that means fifty fifty and if all are positive, or all are negative, then the number of bits which is required is zero here in this case. For example, here in the pure set, all are all are yes, and zero are no's. Similarly, another pure set, when all are all are yes, all are no's, and zero yes, then also it is a pure set. In both the cases the number of bits here this represents the number of bits, in both the

cases the number of bits are zero, hence, this is going to be a pure set, and this is also going to be a pure set, and this is the impure set condition and so, basically either it is pure set, then it will take zero bits or it is highly uncertain then it will take one bit. Whereas all other conditions are lying on either on the ellipse. On the left side or on the right side, where the number of bits is varying between 1 to 0.5.

Refer Slide Time (26:00)

Information Gain

- Want many items in pure sets
- Expected drop in entropy after split:

$$Gain(S, A) = H(S) - \sum_{V \in Values(A)} \frac{|S_V|}{|S|} H(S_V)$$

V ... possible values of A
 S ... set of examples $\{X\}$
 S_V ... subset where $X_A = V$

- Mutual Information**
 - between attribute A and class labels of S

Gain (S, Wind)

$$\begin{aligned}
 &= H(S) - 8/14 * H(S_{weak}) - 6/14 * H(S_{strong}) \\
 &= 0.94 - 8/14 * 0.81 - 6/14 * 1.0 \\
 &= 0.049
 \end{aligned}$$

$- \frac{6}{14} \log_2 \frac{6}{14} - \frac{2}{14} \log_2 \frac{2}{14}$
 $H(S) = 0.94$
9 yes / 5 no
Weak
6 yes / 2 no
 $- \frac{6}{14} \log_2 \frac{6}{14} - \frac{2}{14} \log_2 \frac{2}{14}$
 $H(S_{weak}) = 0.81$
Strong
3 yes / 3 no
 $- \frac{3}{14} \log_2 \frac{3}{14} - \frac{3}{14} \log_2 \frac{3}{14}$
 $H(S_{strong}) = 1.0$

So, with this we have understood, that entropy can be useful to estimate, whether the whether we can go for the split or not? So, if it is a pure set, then entropy the size of the entropy is zero, that means we are not going to a split. If it is a highly uncertain, that is fifty percentage are saying yes and fifty percentage are saying no. Then the value of entropy will become half. So, then trophy will become one, that is highly uncertain, and all other in between which are lying in between. So, their values are between zero and one, that we have seen in the previous example.

Refer Slide Time (26:50)

Entropy

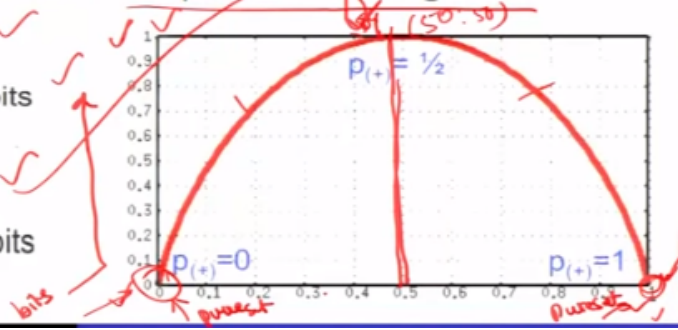
- Entropy: $H(S) = -p_{(+)} \log_2 p_{(+)} - p_{(-)} \log_2 p_{(-)}$ bits
 - S ... subset of training examples
 - $p_{(+)} / p_{(-)} \dots \%$ of positive/negative examples in S
- Interpretation: assume item X belongs to S
 - How many bits need to tell if X positive or negative

- Impure (3 yes / 3 no):

$$H(S) = -\frac{3}{6} \log_2 \frac{3}{6} - \frac{3}{6} \log_2 \frac{3}{6} = 1 \text{ bits}$$

- Pure set (4 yes / 0 no):

$$H(S) = -\frac{4}{4} \log_2 \frac{4}{4} - \frac{0}{4} \log_2 \frac{0}{4} = 0 \text{ bits}$$



Refer Slide Time (26:51)

Which attribute to split on ?



- Want to measure "purity" of the split
 - More certain about Yes/No after the split
 - pure set (4 yes / 0 no) → completely certain (100%)
 - impure (3 yes / 3 no) → completely uncertain (50%)
 - Can't use $P(\text{"yes"} | \text{set})$:
 - must be symmetric: 4 yes / 0 no as pure as 0 yes / 4 no
- any entropy decision split 6 nodes*
pure subsets (100%) correct

So, we are going to evaluate, this purity of the set using entropy, to decide the split of the nodes.

Refer Slide Time (27:09)

Entropy

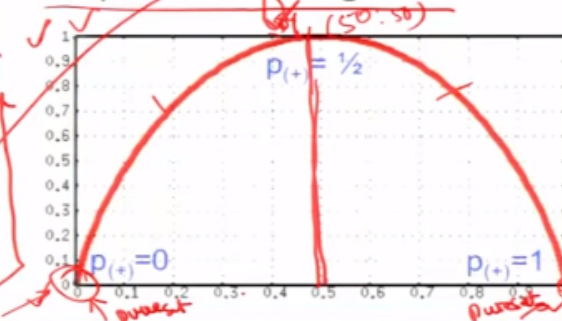
- Entropy: $H(S) = -p_{(+)} \log_2 p_{(+)} - p_{(-)} \log_2 p_{(-)}$ bits
 - S ... subset of training examples
 - $p_{(+)} / p_{(-)}$... % of positive/negative examples in S
- Interpretation: assume item X belongs to S
 - How many bits need to tell if X positive or negative

- Impure (3 yes / 3 no):

$$H(S) = -\frac{3}{6} \log_2 \frac{3}{6} - \frac{3}{6} \log_2 \frac{3}{6} = 1 \text{ bits}$$

- Pure set (4 yes / 0 no):

$$H(S) = -\frac{4}{4} \log_2 \frac{4}{4} - \frac{0}{4} \log_2 \frac{0}{4} = 0 \text{ bits}$$



Refer Slide Time (27:10)

Information Gain

- Want many items in pure sets
- Expected drop in entropy after split:

$$\text{Gain}(S, A) = H(S) - \sum_{V \in \text{Values}(A)} \frac{|S_V|}{|S|} H(S_V)$$

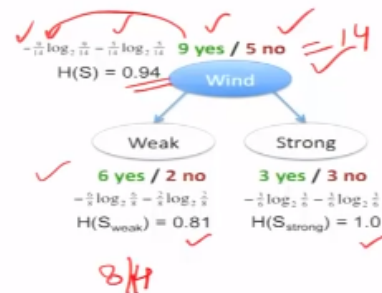
V ... possible values of A
S ... set of examples {X}
S_v ... subset where X_A = V

- Mutual Information

-between attribute A
and class labels of S

Gain (S, Wind)

$$\begin{aligned} &= H(S) - \frac{8}{14} * H(S_{\text{weak}}) - \frac{6}{14} * H(S_{\text{strong}}) \\ &= 0.94 - \frac{8}{14} * 0.81 - \frac{6}{14} * 1.0 \\ &= 0.049 \end{aligned}$$



Let us see, how this decision tree has used this concept? This concept is termed as, information gain. So, we want to see, how many items are there in that pure set, and how many and this particular expectation will drop after the split in, that entropy value, and this can be termed as, the gain and that is called information gain. So, information gain will be HS minus the summation of HS $H(S_v)$ and that is the summation of all the values of A and the set of all samples divided by, that means subsets of all the samples divided by the set of all examples. Let us see, how with this information gain, we are going to do this kind of decision, whether to split or not? So, let us see, here in this case of the wind, we are in the wind we have seen, that there are nine yes and five no's. So, the positive variable is nine,

which will be fed, and the total number of samples are fourteen. So, minus nine divided by fourteen log to the base nine divided by fourteen minus five or no five are negative. So, five divided by fourteen log to the base five divided by fourteen, that comes out to be point nine four. So, information game is point nine four here in the wind. How about calculating this same information gained, for weekend for the strong? For the week we obtained point eight one and for the strong we are obtained one. Now, we can see what is the information gained here in this case of wind, whether we can split or not? So, HS is basically the information gain of a wind. So, that is point nine four, and this particular, then for the weak case it is point eight one and it will be divided by eight upon fourteen. So, so, total number of samples if we count here. They are fourteen. So, that total number of fourteen and out of fourteen we are four we case we are taking only eight six two eight. eight divided by fourteen and minus this comes out to be six divided by fourteen and if we evaluate the gain comes out to be zero point zero four nine information gain. So, if we split, then the net gain will be the positive gain of point zero four nine. Hence, it is better to split.

Refer Slide Time (30:41)

Decision Trees for Regression

So, this way using information gain, we can estimate, or the decision tree will make the decision whether to split or not?

Refer Slide Time (30:51)

How to grow a decision tree

- The tree is built **greedily** from top to bottom
- Each split is selected to **maximize information gain (IG)**

$$IG = \underbrace{\text{Impurity}(Z)}_{\text{Error before split}} - \underbrace{\left(\frac{|Z_L|}{|Z|} \text{Impurity}(Z_L) + \frac{|Z_R|}{|Z|} \text{Impurity}(Z_R) \right)}_{\text{Error after split}}$$

Now, let us see, how to grow, how to build, this particular decision tree. So, the trees is built greedily from top to the bottom and each split is selected to maximize the information gained. So, what we have seen here is, that in the previous example, that after the split the gain is positive. So, we have to maximize after the split. Hence, we are going to split in this particular case.

Refer Slide Time (31:16)

Decision Tree for Regression

- Given a training set: $Z = \{(x_1, y_1), \dots, (x_n, y_n)\}$
 y_i -real values

- Goal is to find $f(x)$ (a tree) such that

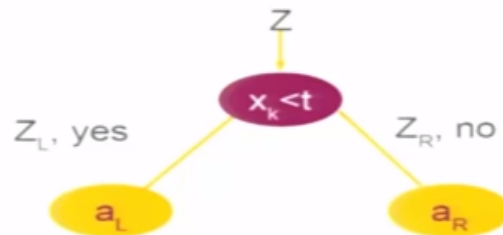
$$\min \sum_{i=1}^n (f(x_i) - y_i)^2$$

- How to grow a decision tree for regression ?

Refer Slide Time (31:20)

How to find the best split

- A tree is built from top to bottom. And at each step, you should find the best split in the internal node. You get a test dataset Z . And here is a splitting criteria x_k less than t or x_k is greater or equal than a threshold t . The problem is how to find k , the number of feature and the threshold t . And also you need to find values in leaves a_L and a_R .



Refer Slide Time (31:29)

What happens without a split?

- What happens without a split?**
- Without a split, the best you can do is to predict one number, 'a'. It makes sense to select such a number 'a' to minimize the squared error.
- It is very easy to prove that such variable a equals an average of all y_i .
- Thus you need to calculate the average value of all the targets.
- Here $\bar{a}$ ($\hat{a}$) denotes the average value.
- Impurity z equals the mean squared error if we use $\bar{a}$ ($\hat{a}$) as a prediction.

Without split: predict one number a

$$\hat{a} = \min_a \sum_{i \in Z} (a - y_i)^2$$

$$\hat{a} = \frac{1}{|Z|} \sum_{i \in Z} y_i$$

$|Z|$ - number of elements in Z

$$\text{Impurity}(Z) = \frac{1}{|Z|} \sum_{i \in Z} (\hat{a} - y_i)^2$$



Refer Slide Time (31:35)

Find the best split ($x_k < t$)

- What happens if you make some split by a condition $x_k < t$? For some part of training objects Z_L , you have $x_k < t$. For other part Z_R holds $x_k \geq t$.
- The error consists of two parts for Z_L and Z_R respectively. So here we need to find simultaneously k , t , a_L and a_R .
- We can calculate the optimal a_L and a_R exactly the same way as we did for the case without a split. We can easily prove that the optimal a_L equals an average of all targets which are the targets of objects which get to the leaf. And we will denote these values by $\hat{a}_L$ and $\hat{a}_R$ respectively.

$$\min_{k, t, a_L, a_R} \sum_{i \in Z_L} (a_L - y_i)^2 + \sum_{i \in Z_R} (a_R - y_i)^2$$

Values in leaves

$$\hat{a}_L = \frac{1}{|Z_L|} \sum_{i \in Z_L} y_i, \quad \hat{a}_R = \frac{1}{|Z_R|} \sum_{i \in Z_R} y_i$$

$|Z_L|$ - number of elements in Z_L ,

$|Z_R|$ - number of elements in Z_R

$$\min_{k, t} \sum_{i \in Z_L} (\hat{a}_L - y_i)^2 + \sum_{i \in Z_R} (\hat{a}_R - y_i)^2$$

So, we can see here, by this way we can find out the the splits, and we have already seen, that if we can do this kind of predictions. To find out this whether to split a node or not. So, whenever there will be information gained and to maximize the gain according to that maximizing the information gained this particular splitting is to be done. So, we have to find out the best split, and we have to do the decision based on that.

Refer Slide Time (32:00)

Stopping rule

- The node depth is equal to the **maxDepth** training parameter.
- No split candidate leads to an information gain greater than **mininfoGain**.
- No split candidate produces child nodes which have at least **minInstancesPerNode** training instances ($|Z_L|, |Z_R| < \text{minInstancesPerNode}$) each.

Now, there is another criteria which is called a stopping criteria, that the the growth of the spanning tree or the building of the spanning tree it has to be stopped. So, the node depth is equal to the maximum and depth of the training parameter, and no split candidate leads to an information getting

greater than minimum information gain, and no candidate produces the child, which have at least the minimum instances per node training instances.

Refer Slide Time (32:33)

Summary: Decision Trees

- **Automatically handling interactions of features:** The benefits of decision tree is that this algorithm can automatically handle interactions or features because it can combine several different features in a single decision tree. It can build complex functions involving multiple splitting criteria.
- **Computational scalability:** The second property is a computational scalability. There exists effect of algorithms for building decision trees for the very large data sets with many features.
- **Predictive power:** Single decision tree actually is not a very good predictor. The predictive power of a single tree is typically not so good.
- **Interpretability:** You can visualize the decision tree and analyze this splitting criteria in nodes, the values in leaves, and so one. Sometimes it might be helpful.

So, in summary what we have seen in the decision tree is, that here the automatic handling of interactions of the features is done, and computationally scalability also we have seen, and interpretability aspects also.

Refer Slide Time (32:47)

Building a tree using MapReduce

Now, let us see, how to build tree using MapReduce.

Refer Slide Time (32:53)

Problem: Building a tree

- Given a large dataset with hundreds of attributes ✓

- Build a decision tree! *Big data*

- General considerations:

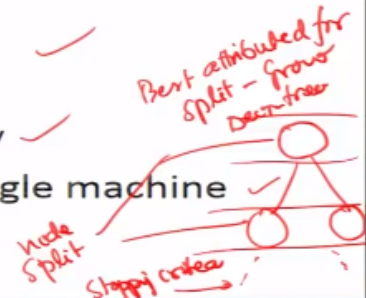
- Tree is small (can keep it memory):
 - Shallow (~10 levels) ✓
- Dataset too large to keep in memory ✓
- Dataset too big to scan over on a single machine ✓
- MapReduce to the rescue! ✓

Algorithm 1 FindBestSplit

Require: Node n , Data $D \subseteq D^*$

```

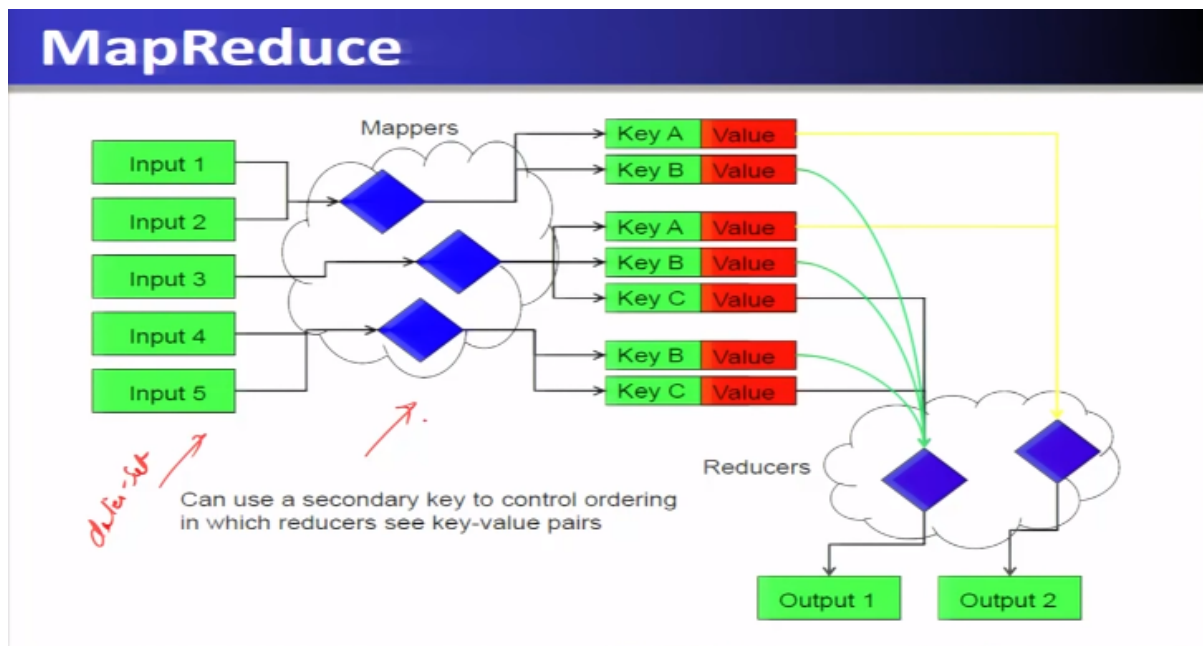
1: ( $n \rightarrow \text{split}, D_L, D_R$ ) = FindBestSplit( $D$ ) ✓
2: if StoppingCriteria( $D_L$ ) then ✓
3:    $n \rightarrow \text{left\_prediction}$  = FindPrediction( $D_L$ ) ✓
4: else ✓
5:   FindBestSplit( $n \rightarrow \text{left}, D_L$ ) ✓
6: if StoppingCriteria( $D_R$ ) then ✓
7:    $n \rightarrow \text{right\_prediction}$  = FindPrediction( $D_R$ ) ✓
8: else ✓
9:   FindBestSplit( $n \rightarrow \text{right}, D_R$ ) ✓
    
```



So, now MapReduce will be applied for the decision tree algorithm, and here we will see the approach and we will see more detail of it. How this MapReduce will be applied in the decision tree. So, that the big data analytics can be performed with the help of this parallel a decision tree **Clementessin**

Now, here we are given a large data sets, and which has a lot of attribute the size of the data set is also big, and the number of attributes are also more, and we want to build a decision tree and that is what is called the Big Data applying that is the big data analysis. So, here we can we have to see is, that the tree all the data is very big, but the tree is small and it can be stay in the memory and maximum levels of that particular tree, which is extracted out of big data is on off out of the ten levels, and also we have to keep the data set the very large data set into the memory and this data set is too big to be kept and one machine. So, basically MapReduce we are going to use which is going to compute over the clusters of machines. Now, here here in this particular technique which is called the decision tree. We have to find out the best the splits and, then we have to find out which attribute along with, which we are going to build the tree, that is the single tree which we are going to build. So, we have to first find out which is the best variable around, which the split has to be done or the tree has to be constructed. So, the best we have to find out the best attribute, for the split to grow the the decision tree. So, after having identified a particular attribute. Now, we have to identify the how the best split is to be done? So, that in the left side and the right side of the tree can basically grow and so on. So, these particular growth of the spanning of the decision tree which we are going to see level wise growth of the tree and also the stopping criteria. When we have to stop further growth of the tree. So, this particular decision is called when we are going to split. So, this is whether the node, what is the condition in which we are going to split the node, along whether to split or not and if it is then which particular attribute will be the best to split. So, all these things we are going to discuss which are going to be implemented using MapReduce in.

Refer Slide Time (36:17)



So, MapReduce we have already seen, that it is it will the input data set, is now divided into the into the chunks, and these chunks are given to the map or functions and the mapper function will generate the key value pair, which is further given to the reduce functions and they will generate the output.

Refer Slide Time (36:47)

PLANET

Parallel Learner for Assembling Numerous Ensemble Trees [Panda et al., VLDB '09]

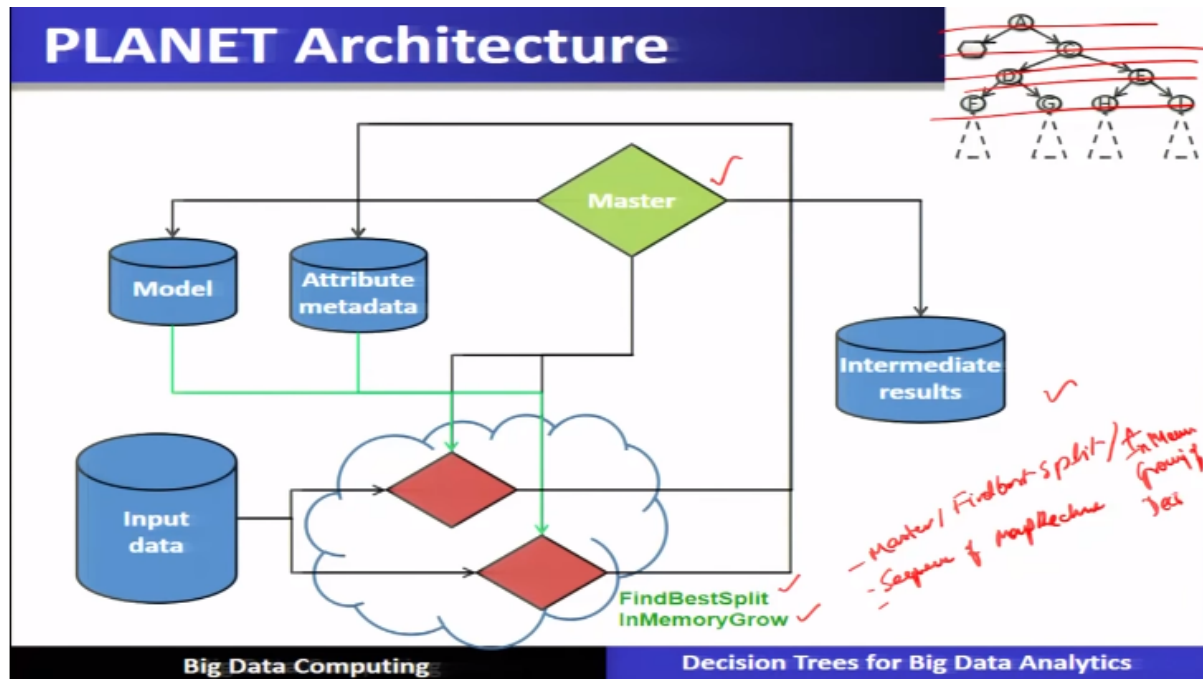
- A **sequence** of MapReduce jobs that build a decision tree ✓
- **Setting:**
 - Hundreds of numerical (discrete & continuous) attributes ✓
 - Target (class) is numerical: **Regression**
 - Splits are binary: $X_j < v$
 - Decision tree is small enough for each Mapper to keep it in memory ✓
 - Data too large to keep in memory

binary tree

Now, this particular implementation of a MapReduce for decision tree, is published in a paper which is called planet by the Google. So, parallel learner for assembling numerous n sampling n symbol trees which we are going to see here as, the MapReduce implementation of a decision tree. So, this is nothing but a sequence of MapReduce job, that will build the decision trees, and here we assume, that

the values which are given in the data set, they are basically hundreds of numerical attributes, and we are going to build the binary tree, which is the decision tree, which is a binary which is a binary tree and this particular tree which we assume, that is small enough to be for each mapper to be kept in the memory and data is too large.

Refer Slide Time (37:45)



Let us see the different constituents in this architecture. So, here we see, that there will be a master node, and this particular master node in turn will decide about finding the best split and also the in memory growth of that tree at the mapper end, and the intermediate results are also stored, and this all happens in the coordination of the master. So, the tree is built level wise using these different, MapReduce sequence of MapReduce operations. MapReduce operations are being coordinated by the master. There will be a master iron, which will decide, the best is split and also tries to decide about in memory growing of the decision tree. I had the mapper.

Refer Slide Time (39:08)

```

graph TD
    A --> B
    A --> C
    B --> D
    B --> E
    D --> F
    D --> G
    E --> H
    E --> I
    F --- F_group[ ]
    G --- G_group[ ]
    H --- H_group[ ]
    I --- I_group[ ]
    style F_group fill:none,stroke-dasharray: 5 5
    style G_group fill:none,stroke-dasharray: 5 5
    style H_group fill:none,stroke-dasharray: 5 5
    style I_group fill:none,stroke-dasharray: 5 5
  
```

-
- ```

graph TD
 A --> B
 A --> C
 B --> D
 B --> E
 D --> F
 D --> G
 E --> H
 E --> I
 F --- F_group[]
 G --- G_group[]
 H --- H_group[]
 I --- I_group[]
 style F_group fill:none,stroke-dasharray: 5 5
 style G_group fill:none,stroke-dasharray: 5 5
 style H_group fill:none,stroke-dasharray: 5 5
 style I_group fill:none,stroke-dasharray: 5 5

```

```

graph TD
 A --> B
 A --> C
 B --> D
 B --> E
 D --> F
 D --> G
 E --> H
 E --> I
 F --- F_group[]
 G --- G_group[]
 H --- H_group[]
 I --- I_group[]
 style F_group fill:none,stroke-dasharray: 5 5
 style G_group fill:none,stroke-dasharray: 5 5
 style H_group fill:none,stroke-dasharray: 5 5
 style I_group fill:none,stroke-dasharray: 5 5

```

```

graph TD
 A --> B
 A --> C
 B --> D
 B --> E
 D --> F
 D --> G
 E --> H
 E --> I
 F --- F_group[]
 G --- G_group[]
 H --- H_group[]
 I --- I_group[]
 style F_group fill:none,stroke-dasharray: 5 5
 style G_group fill:none,stroke-dasharray: 5 5
 style H_group fill:none,stroke-dasharray: 5 5
 style I_group fill:none,stroke-dasharray: 5 5

```

```

graph TD
 A((A)) --> B((B))
 A --> C((C))
 C --> D((D))
 C --> E((E))
 E --> F((F))
 E --> G((G))
 E --> H((H))
 E --> I((I))
 style F stroke:#00FF00,stroke-width:2px
 style G stroke:#00FF00,stroke-width:2px
 style H stroke:#00FF00,stroke-width:2px
 style I stroke:#00FF00,stroke-width:2px

```

- 
- ```

graph TD
    A((A)) --> B((B))
    A --> C((C))
    C --> D((D))
    C --> E((E))
    E --> F((F))
    E --> G((G))
    E --> H((H))
    E --> I((I))
    style F stroke:#00FF00,stroke-width:2px
    style G stroke:#00FF00,stroke-width:2px
    style H stroke:#00FF00,stroke-width:2px
    style I stroke:#00FF00,stroke-width:2px

```

```

graph TD
    A((A)) --> B((B))
    A --> C((C))
    C --> D((D))
    C --> E((E))
    E --> F((F))
    E --> G((G))
    E --> H((H))
    E --> I((I))
    style F stroke:#00FF00,stroke-width:2px
    style G stroke:#00FF00,stroke-width:2px
    style H stroke:#00FF00,stroke-width:2px
    style I stroke:#00FF00,stroke-width:2px

```

leaf node L, and for each node L it keeps the statistics about, the data reaching L and the data in the left and right subtree under the splits S. This way the mapper. So, each mapper will now, has to deal with a particular level of the nodes and has to decide, whether it has to be further split or not. If it is to be splitted further, then another iteration will go on with the next set of MapReduce job. Reducer aggregates the statistics of one and two and determines the best split at each node.

Refer Slide Time (41:00)

PLANET: Components

- **Master**
 - Monitors everything (runs multiple MapReduce jobs)
- **Three types of MapReduce jobs:** ✓
- (1) **MapReduce Initialization (run once first)**
 - For each attribute identify values to be considered for splits ✓
(root node)
- (2) **MapReduce FindBestSplit (run multiple times)** ✓ for each level of decision tree
 - MapReduce job to find best split when there is too much data to fit in memory
- (3) **MapReduce InMemoryBuild (run once last)**
 - Similar to FindBestSplit (but for small data)
 - Grows an entire sub-tree once the data fits in memory



So, overall different components in this MapReduce implementation of decision tree is like this. So, master will monitor everything and runs multiple MapReduce jobs. So, besides master there are three different type of MapReduce jobs, which will be executed. So, first is the initialization of MapReduce, that is it will be run once. So, for each attribute, it will identify the values to be considered, for the split. So, for example the root node so, for the root node decision, has to be done in during the initialization phase, and then it will be a second set of, MapReduce job to find out the best is split, and it will run multiple times for each level of decision tree growth. So, the second step will keep on increasing keep on growing the the the levels of the tree as they are being invoked. So, the number of time is it is invoked. So, that indicates, that it will be growing in that levels. So, that particular so MapReduce job is to find the best displayed, when there are too much of data to fit in the memory. So, this way this MapReduce will try to find out the best split and keep on splitting and also, it will be then the next step is finally the MapReduce will build their-memory build, that is sub trees are built at each MapReduce then they will be runs once. So, similar to find the best split for small data. It will grow and entire sub tree once the data fits in the memory. So, these three different MapReduce jobs one is for the root node, that is the initialization. The second one is at the level wise splitting of the nodes, and growing the decision tree, and finally it has to take this subtrees, which are in memory and built during each run. So, it will be done once.

Refer Slide Time (43:33)

Reference

- B. Panda, J. S. Herbach, S. Basu, and R. J. Bayardo. PLANET: Massively parallel learning of tree ensembles with MapReduce. *VLDB* 2009.
- J. Ye, J.-H. Chow, J. Chen, Z. Zheng. Stochastic Gradient Boosted Distributed Decision Trees. *CIKM* 2009.

These are the references for this discussion, which we have just seen.

Refer Slide Time (43:43)

Example: Medical Application using a Decision Tree in Spark ML

Example medical application using decision tree in spark machine learning.

Refer Slide Time (43:50)

Create SparkContext and SparkSession

- This example will be about using Spark ML for doing classification and regression with decision trees, and in samples of decision trees.
- First of all, you are going to create SparkContext and SparkSession, and here it is.

```
In [ ]: from pyspark import SparkContext
        from pyspark.sql import SparkSession
```

*Consider Medical application using
Decision trees for
Big data Analytics*

```
In [ ]: sc = SparkContext(appName = "module3_week4")
```

*1. Create
SparkContext*

```
In [ ]: ! echo $PYSPARK_SUBMIT_ARGS
```

```
In [ ]: spark = SparkSession.Builder().getOrCreate() # required for dataframes
```

*2. build
spark
session*

Now, we will consider an example of the medical application of the application of decision trees, for using spark a ML. So, in this example, we will see how the classification using the decision trees is done. So, let us see the steps, for doing this decision tree based analytics, for the medical application. So, here we are going to consider, the medical application using decision trees, for big data analytics. Now, here we have to see, that we have to give the name of the application and now, we have to create the spark context, using this application name. So, we call it as by spark context. We have to create a spark context. After creating the spark context, then we have to build a session that is called a spark session. This is step number one, and this is step number two, which is essential in all such applications, which runs under spark. First, we have to create the context, then a spark session and this is required to build the data frames.

Refer Slide Time (46:10)

Download a dataset ✓

- Now you are downloading a dataset which is about breast cancer diagnosis, here it is.

Download a dataset (breast cancer diagnosis) ✓

```
In [ ]: !wget https://archive.ics.uci.edu/ml/machine-learning-databases/breast-cancer-wisconsin/wdbc.data ✓

In [ ]: !head -n 3 wdbc.data ✓

In [ ]: !wc -l wdbc.data ✓

In [ ]: #1) ID number ✓
        #2) Diagnosis (M = malignant, B = benign) ✓
        #3) Features ✓

In [ ]: from pyspark.ml.linalg import Vectors
        from pyspark.ml.feature import StringIndexer

In [ ]: # Load a text file and convert each line to a Row.
        data = []
        with open("wdbc.data") as infile: ✓
```

m/B
Diagnosis
features
numeric

The next step would be to download the data set, which is used here in this example. So, we will download the data set, from some link, and let us, call it as, this data set which is being downloaded. Once the data set is available and is downloaded into the spark system. Then we have to see, that this particular data set has three different parts. That is in the form of key value pair. it has an ID number and, then it will have, the key as, whether it is having a medical problem yes or no means align or having benign and then different features of it. So, the data will be in this particular form. It will be a diagnosis, and then it will be a features. Now, this diagnosis if it is M or B it has to be converted into the numeric form. So, we have to convert, we have to load into file and we have to convert this field into the numeric field.

Refer Slide Time (47:36)

Exploring the Dataset

- Let's explore this dataset a little bit. The first column is ID of observation. The second column is the diagnosis, and other columns are features which are comma separated. So these features are results of some analysis and measurements. There are total 569 examples in this dataset. And if the second column is M, it means that it is cancer. If B, means that there is no cancer in a particular woman.

```
In [6]: !head -n 3 wdbc.data

842302,M,17.99,10.38,122.8,1001,0.1184,0.2776,0.3001,0.1471,0.2419,0.07871,1.09
5,0.9053,8.589,153.4,0.006399,0.04004,0.05373,0.01587,0.03003,0.006193,25.38,1
7.33,184.6,2019,0.1622,0.6656,0.7119,0.2654,0.4601,0.1189
842517,M,20.57,17.77,152.9,1326,0.08474,0.07864,0.0869,0.07017,0.1812,0.05667,
0.5435,0.7339,3.398,74.08,0.005225,0.01308,0.0186,0.0134,0.01389,0.003532,24.9
9,23.41,158.8,1956,0.1238,0.1866,0.2416,0.186,0.275,0.08902
84300903,M,19.69,21.25,130.1203,0.1096,0.1599,0.1974,0.1279,0.2069,0.05999,0.74
56,0.7869,4.585,94.03,0.00615,0.04006,0.03832,0.02058,0.0225,0.004571,23.57,25.
53,152.5,1709,0.1444,0.4245,0.4504,0.243,0.3613,0.08758

In [ ]: !wc -l wdbc.data

In [ ]: #1) ID number
        #2) Diagnosis (M = malignant, B = benign)
        #3) Features

In [ ]: from pyspark.ml.linalg import Vectors
        from pyspark.ml.feature import StringIndexer

In [ ]: # Load a text file and convert each line to a Row.
```

So, out of this particular data set, that is five sixty nine examples in a data set we have to convert in this particular field of diagnosis into the numeric values, and then we have to see this particular explore the data set.

Refer Slide Time (47:48)

Exploring the Dataset

- First of all you need to transform the label, which is either M or B from the second column. You should transform it from a string to a number. We use a StringIndexer object for this purpose, and first of all you need to load all these datasets. Then you create a Spark DataFrame, which is stored in the distributed manner on cluster.

```

In [ ]: from pyspark.ml.linalg import Vectors
        from pyspark.ml.feature import StringIndexer

In [ ]: # Load a text file and convert each line to a Row.
        data = []
        with open("wdbc.data") as infile:
            for line in infile:
                tokens = line.rstrip("\n").split(",")
                y = tokens[1]
                features = Vectors.dense([float(x) for x in tokens[2:]])
                data.append((y, features))

In [ ]: inputDF = spark.createDataFrame(data, ["label", "features"])

In [ ]: inputDF.show()

In [ ]: stringIndexer = StringIndexer(inputCol = "label", outputCol = "labelIndexed")
        si_model = stringIndexer.fit(inputDF)
        inputDF2 = si_model.transform(inputDF)
  
```

So, after having read this particular data set, which we have downloaded. Now, it will be creating the data frames with the label and the feature, and the label will be converted into into the numeric values, Refer Slide Time (48:19)

Exploring the Dataset

- inputDF DataFrame has two columns, label and features. Okay, you use an object vector for creating a vector column in this dataset, and then you can do string indexing. So Spark now enumerates all the possible labels in a string form, and transforms them to the label indexes. And now label M is equivalent 1, and B label is equivalent 0.

```

In [11]: inputDF = spark.createDataFrame(data, ["label", "features"])

In [12]: inputDF.show()
  
```

label	features
M	[17.99,10.38,122.]
M	[20.57,17.77,132.]
M	[19.69,21.25,130.]
M	[11.42,20.38,77.5]
M	[20.29,14.34,135.]
M	[12.45,15.7,82.57]
M	[18.25,19.98,119.]
M	[13.71,20.83,90.2]
M	[13.0,21.82,67.5]
M	[12.46,24.04,83.9]
M	[16.02,23.24,102.]
M	[15.78,17.89,103.]
M	[19.17,24.8,132.4]
M	[15.85,23.95,103.]
M	[13.73,22.61,93.6]
M	[14.54,27.54,96.7]
M	[14.68,20.13,94.7]

Refer Slide Time (48:24)

Train/Test Split

- We can start doing the machine learning right now. First of all, you make training test splitting in the proportion 70% to 30%. And the first model you are going to evaluate is one single decision tree.

train/test split

```
In [15]: (trainingData, testData) = inputDF2.randomSplit([0.7, 0.3], seed = 23)
```

Training Decision Tree

```
In [ ]: from pyspark.ml.classification import DecisionTreeClassifier
        from pyspark.ml.evaluation import MulticlassClassificationEvaluator

In [ ]: decisionTree = DecisionTreeClassifier(labelCol = "labelIndexed")

In [ ]: dtModel = decisionTree.fit(trainingData)

In [ ]: dtModel.numNodes

In [ ]: dtModel.depth

In [ ]: dtModel.featureImportances
```

Refer Slide Time (48:27)

Exploring the Dataset

- inputDF DataFrame has two columns, label and features. Okay, you use an object vector for creating a vector column in this dataset, and then you can do string indexing. So Spark now enumerates all the possible labels in a string form, and transforms them to the label indexes. And now label M is equivalent 1, and B label is equivalent 0.

```
In [11]: inputDF = spark.createDataFrame(data, ["label", "features"])

In [12]: inputDF.show()
```

```
-----+-----+
|label|features|
-----+-----+
M|[17.99,10.38,122....|
M|[20.57,17.77,132....|
M|[19.69,21.25,130....|
M|[11.42,20.38,77.5....|
M|[20.29,14.34,135....|
M|[12.45,15.7,82.57....|
M|[18.25,19.98,119....|
M|[13.71,20.83,90.2....|
M|[13.0,21.82,87.5....|
M|[12.46,24.04,83.9....|
M|[16.02,23.24,102....|
M|[15.78,17.89,103....|
M|[19.17,24.8,132.4....|
M|[15.85,23.95,103....|
M|[13.73,22.61,93.6....|
M|[14.54,27.54,96.7....|
M|[14.68,20.13,94.7....|
-----+-----+
-- -- -- -- --
```

and after that this particular data set, which will having this values, converted into so, these M fields will be converted into the numeric values, and then now after converting it, now it will be ready. So, the conversion will be M is equivalent to one and V is equivalent to zero. So, wherever one is there M is there. So, it will be one and wherever B is there. So, after this conversion now this particular data set will be now, ready for the use of machine learning.

Refer Slide Time (48:58)

Train/Test Split

- We can start doing the machine learning right now. First of all, you make training test splitting in the proportion 70% to 30%. And the first model you are going to evaluate is one single decision tree.

train/test split

```
In [15]: (trainingData, testData) = inputDF2.randomSplit([0.7, 0.3], seed = 23)
```

Training Decision Tree

```
In [ ]: from pyspark.ml.classification import DecisionTreeClassifier
        from pyspark.ml.evaluation import MulticlassClassificationEvaluator

In [ ]: decisionTree = DecisionTreeClassifier(labelCol = "labelIndexed")

In [ ]: dtModel = decisionTree.fit(trainingData)

In [ ]: dtModel.numNodes

In [ ]: dtModel.depth

In [ ]: dtModel.featureImportances
```

So, the first after preparing the data. Now, we have to use the data for the machine learning, for that we have to split this data into the training data set, and the test dataset, into the proportion of seventy and thirty, So, to split this data set into the training and test data splits. Now, we have to use this particular spark come on random split, that is seventy percent is of training data, and thirty percent is the test data. So, after this splitting of the training data and the test data. Now, we are going to apply this decision tree algorithm. So, we will now use the decision tree classifier and now we will use this label indexed, which we have seen in the previous slide, and then we will do this decision tree fit on the decision on the training data set, which we have already obtained from the the data set is split seventy percent of it. So, after this step we are ready with the the decision tree. Now, we can see out of this decision tree it's different parameters such as, how many nodes are there in the decision tree what is the depth of the decision tree? and also, the what are the features which are of importance? and this will build the model of a decision tree. So, decision tree model, is being built here after this execution.

Refer Slide Time (51:14)

Train/Test Split

use data for ML

- We can start doing the machine learning right now. First of all, you make training test splitting in the proportion 70% to 30%. And the first model you are going to evaluate is one single decision tree.

```
train/test split ✓
In [15]: (trainingData, testData) = inputDF2.randomSplit([0.7, 0.3], seed = 23)
Training Decision Tree
In [ ]: from pyspark.ml.classification import DecisionTreeClassifier
        from pyspark.ml.evaluation import MulticlassClassificationEvaluator
In [ ]: decisionTree = DecisionTreeClassifier(labelCol = "labelIndexed")
In [ ]: dtModel = decisionTree.fit(trainingData) ✓
In [ ]: dtModel.numNodes ✓
In [ ]: dtModel.depth ✓
In [ ]: dtModel.featureImportances ✓
```

test data ✓
training data ✓
Decision Tree model ✓

Refer Slide Time (51:13)

Train/Test Split

- We are making import DecisionTreeClassifier object. We create a class which is responsible for training, and we call the method fit to the training data, and obtain a decision tree model.

```
In [17]: decisionTree = DecisionTreeClassifier(labelCol = "labelIndexed")
In [18]: dtModel = decisionTree.fit(trainingData)
In [19]: dtModel.numNodes ✓
Out[19]: 29 ✓
In [20]: dtModel.depth
Out[20]: 5 ✓
In [21]: dtModel.featureImportances ✓
Out[21]: SparseVector(30, (1: 0.0589, 6: 0.0037, 10: 0.0112, 13: 0.0117, 20: 0.0324, 21: 0.0302, 22: 0.7215, 24: 0.01, 26: 0.0191, 27: 0.1013))
In [22]: dtModel.numFeatures ✓
Out[22]: 30 ✓
In [ ]: print dtModel.toDebugString
```

Now, now we can see, that if we find out how many in this particular data set of around six hundred different samples. We have seen, that there are twenty nine and nodes are there, and the number of depth is only five, and different important features, which we have obtained in the form of the graph and number of features are thirty.

Refer Slide Time (51:46)

Train/Test Split

- Number of nodes and depths of decision tree, feature importances, total number of features used in this decision tree, and so on.
- We can even visualize this decision tree and explore it, and here is a structure of this decision tree. Here are splitting conditions If and Else, which predict values and leaves of our decision trees.

```
Out[21]: SparseVector(30, {1: 0.0589, 6: 0.0037, 10: 0.0112, 13: 0.0117, 20: 0.0324, 21: 0.0302, 22: 0.7215, 24: 0.01, 26: 0.0191, 27: 0.1013})

In [22]: dtModel.numFeatures
Out[22]: 30

In [ ]: print dtModel.toDebugString

In [ ]: predictions = dtModel.transform(testData)

In [ ]: predictions.select('label', 'labelIndexed', 'probability', 'prediction').show()

In [ ]: evaluator = MulticlassClassificationEvaluator(labelCol = "labelIndexed", predictionCol = "prediction", metricName = "accuracy")
accuracy = evaluator.evaluate(predictions)
print("Test Error = %g" % (1.0 - accuracy))
```

So, after constructing this particular decision tree. Now, we have to see that we have to evaluate this particular decision tree, whether it is the quality wise whether it is good or bad. How it is good for the predictions? So, evaluation is done. So, multiclass classifier multiclass classification evaluator is used for this particular purpose, to evaluate the predictions accuracy.

Refer Slide Time (52:20)

Train/Test Split

- Now we are applying a decision tree model to the test data, and obtain predictions. And you can explore these predictions. The predictions are in the last column.
- And in this particular case, our model always predicts zero class.

```
In [23]: print dtModel.toDebugString

DecisionTreeClassificationModel (uid=DecisionTreeClassifier_4653be4ce1bd9e589b6
4) of depth 5 with 29 nodes
  If (feature 22 <= 114.2)
    If (feature 27 <= 0.1613)
      If (feature 20 <= 16.57)
        If (feature 27 <= 0.1258)
          If (feature 10 <= 0.9289)
            Predict: 0.0
          Else (feature 10 > 0.9289)
            Predict: 1.0
        Else (feature 27 > 0.1258)
          If (feature 21 <= 32.85)
            Predict: 0.0
          Else (feature 21 > 32.85)
            Predict: 1.0
      Else (feature 20 > 16.57)
        If (feature 1 <= 16.54)
          Predict: 0.0
        Else (feature 1 > 16.54)
          If (feature 24 <= 0.1084)
            Predict: 0.0
          Else (feature 24 > 0.1084)
            Predict: 1.0
    Else (feature 22 > 114.2)
      Predict: 0.0
```

Refer Slide Time (52:23)

Predictions

```
In [ ]: predictions = dtModel.transform(testData)

In [ ]: predictions.select('label', 'labelIndexed', 'probability', 'prediction').show()

In [ ]: evaluator = MulticlassClassificationEvaluator(labelCol = "labelIndexed", predictionCol = "prediction", metricName = "accuracy")
accuracy = evaluator.evaluate(predictions)

print("Test Error = %g" % (1.0 - accuracy))
```

[illegible]

Refer Slide Time (52:41)

Predictions

[illegible]

Here we can see that these predictions are shown here in this case. So, this is the decision tree, which is generated which is nothing but here you can see all the predictions, that means predictions are there at the leaf nodes and all other nodes are there the decision nodes are there at the internal nodes

Refer Slide Time (52:53)

Train/Test Split

- Now we are applying a decision tree model to the test data, and obtain predictions. And you can explore these predictions. The predictions are in the last column.
- And in this particular case, our model always predicts zero class.

```
In [23]: print dtModel.toDebugString
```

```
DecisionTreeClassificationModel (uid=DecisionTreeClassifier_4653be4ce1bd9e589b6
4) of depth 5 with 29 nodes
  If (feature 22 <= 114.2)
    If (feature 27 <= 0.1613)
      If (feature 20 <= 16.57)
        If (feature 27 <= 0.1258)
          If (feature 10 <= 0.9289)
            Predict: 0.0
          Else (feature 10 > 0.9289)
            Predict: 1.0
        Else (feature 27 > 0.1258)
          If (feature 21 <= 32.85)
            Predict: 0.0
          Else (feature 21 > 32.85)
            Predict: 1.0
        Else (feature 20 > 16.57)
          If (feature 1 <= 16.54)
            Predict: 0.0
          Else (feature 1 > 16.54)
            If (feature 24 <= 0.1084)
              Predict: 0.0
            Else (feature 24 > 0.1084)
              Predict: 1.0
```

Decision

Decision

So, this is the prediction node this is also the prediction nodes, this is also the prediction node, and you see that this is the decision tree which is being calculated out of the training data set,

Refer Slide Time (53:32)

Predictions

```
In [ ]: predictions = dtModel.transform(testData)

In [ ]: predictions.select('label', 'labelIndexed', 'probability', 'prediction').show()

In [ ]: evaluator = MulticlassClassificationEvaluator(labelCol = "labelIndexed", predictionCol = "prediction", metricName = "accuracy")
accuracy = evaluator.evaluate(predictions)

print("Test Error = %g" % (1.0 - accuracy))
```

and you and using the test data set we can evaluate this particular model of a training of a decision tree which is being generated based on its prediction accuracy,

Refer Slide Time (53:45)

Accuracy

```
| 0 | 0.0 | 0.99086757990867... | 0.0 |
| 0 | 0.0 | 0.99086757990867... | 0.0 |
| 0 | 0.0 | 0.99086757990867... | 0.0 |
| 0 | 0.0 | 0.99086757990867... | 0.0 |
| 0 | 0.0 | 0.99086757990867... | 0.0 |
| 0 | 0.0 | 0.99086757990867... | 0.0 |
| 0 | 0.0 | 0.99086757990867... | 0.0 |
| 0 | 0.0 | 0.99086757990867... | 0.0 |
| 0 | 0.0 | 0.99086757990867... | 0.0 |
| 0 | 0.0 | 0.99086757990867... | 0.0 |
| 0 | 0.0 | 0.99086757990867... | 0.0 |
| 0 | 0.0 | 0.99086757990867... | 0.0 |
| 0 | 0.0 | 0.99086757990867... | 0.0 |
| 0 | 0.0 | 0.99086757990867... | 0.0 |
| 0 | 0.0 | 0.99086757990867... | 0.0 |
| 0 | 0.0 | 0.99086757990867... | 0.0 |
| 0 | 0.0 | 0.99086757990867... | 0.0 |
| 0 | 0.0 | 0.99086757990867... | 0.0 |
| 0 | 0.0 | 0.99086757990867... | 0.0 |
| 0 | 0.0 | 0.99086757990867... | 0.0 |
+-----+
only showing top 20 rows

In [26]: evaluator = MulticlassClassificationEvaluator(labelCol = "labelIndexed", predicti
accuracy = evaluator.evaluate(predictions)

print("Test Error = %g" % (1.0 - accuracy))
Test Error = 0.0451977 ✓ ab16
```

and we have seen that the prediction accuracy comes out to be ninety six percent which is fairly well.

Refer Slide Time (53:50)

Conclusion

- In this lecture, we have discussed Decision Trees for Big Data Analytics.
- We have also discussed a case study of Breast Cancer Diagnosis using a Decision Tree in Spark ML.

So, conclusion in this lecture we have discussed the decision trees for our bigdata analytics, and we have also seen one case study of a medical application. Wherein we have applied the decision tree based on machine learn is part machine learning. Thank you.