

**Quantitative Finance**  
**Prof. Raghu Nandan Sengupta**  
**Department of Industrial and Management Engineering**  
**Indian Institute of Technology, Kanpur**

**Module – 06**

**Lecture - 35**

So, welcome back to this quantitative finance program. So, I am sure you have, till the last class, we discussed Black-Scholes model, the Ito's Lemma, the venna process and based on the assumption; different assumptions, you are able to understand so briefly, the concept of Black-Scholes model and what was the assumptions of Black-Scholes model; the concept of there is no brokerage cost, risk free interest rate, lending borrowing is fixed, then volatility concept and all those things; we did go into detail, but before that; obviously, we had considered different types of options; how options can be combined, the straddle, the butterfly, the bears spread, the bulls spread so on and so forth.

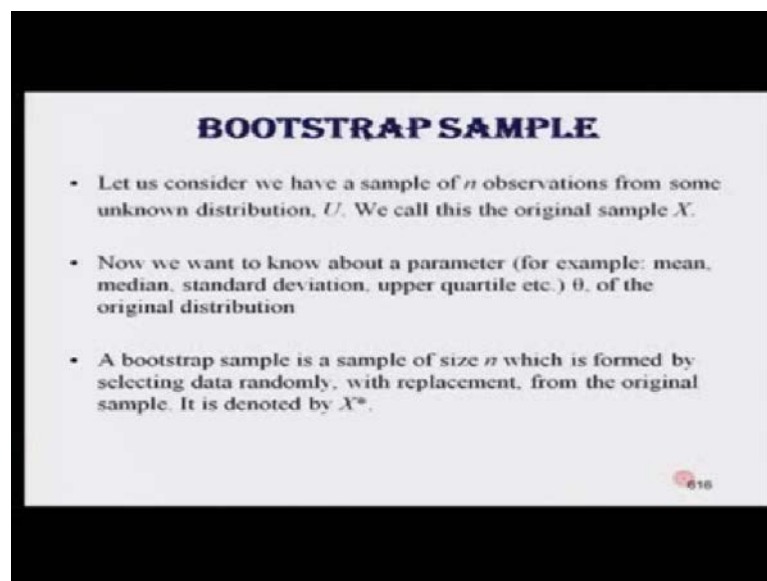
In the initial part if you remember; however, after going through the different averaging intending, before that, we considered the concept of portfolio management and how safety first principle, Mark Which model and all these things could be considered, or the ratios rank, the assets in the portfolio considering the excess return with respect to beta, excess return with respect to standard deviation; all these different models were considered. Now, for today and the next day, the class would be not any concept wise discussion. It will be, we will just deviate from our concept of how we have been tackling the course. Here, I will give you more of a very specific technique which is used in portfolio optimization, advance set of techniques, which is known as robust optimization. It is another concept of reliability optimization also. There is concept of stochastic programming also. So, the problems which we will discuss are the actual problems which we have solved and based on that, it would be much easier for me to explain the nuances of the optimization; what are the concepts it considers and so on and so forth.

Rather than trying to understand the nitty-gritty's, I am sure it will give you much better flavor that what are the changes in the model which occur, the moment we consider non deterministic nature of your of the parameters, which are considered in the optimization,

but before we start discussing the models, we will discuss some models. Before that we will discuss what are the parameters; what are the variables which are considered. Before that I will discuss a very important concept which is used in statistics and optimization, is the bootstrapping method. Now, as the name implies, bootstrapping means if somebody is wearing a very high, till the knee length boot and you have to tie the laces. So, you basically start tying the laces from the bottom most level, tighten up each and every lace, which is there in the eye of the shoe and then basically, continue doing it till you reach the highest level and you tie the lace.

So, bootstrapping exactly, like this; that we keep repeating some set of simulations; some set of mathematical procedures till we are able to understand the general characteristics of the distribution based on which, we are trying to do the study. It basically gives you some re-sampling techniques. Sampling means you pick up some set of observation from a population. Re-sampling means you keep repeating, picking up sample from the population and doing it in such a way that the characteristic of the samples and the resample and the average of the samples in the long run, gives you some information about what are the parameters of the population and they do give us to a high level of accuracy; what are the parameters of the population. So, what is the background of bootstrap? Let me consider in very simple words.

(Refer Slide Time: 04:01)



**BOOTSTRAP SAMPLE**

- Let us consider we have a sample of  $n$  observations from some unknown distribution,  $U$ . We call this the original sample  $X$ .
- Now we want to know about a parameter (for example: mean, median, standard deviation, upper quartile etc.)  $\theta$ , of the original distribution
- A bootstrap sample is a sample of size  $n$  which is formed by selecting data randomly, with replacement, from the original sample. It is denoted by  $X^*$ .

618

Let us consider we have a sample of  $n$  observation. Now, this small  $n$  which we are considering, so this small  $n$  will be considered, is not the whole population. Say for example, if we have all the number of students who have taken the CAT examination; the common admission test examination for the MBA program or say for example, the GATE examination; graduate aptitude test related to MTECH and PHD program in India or say for example, somebody is interested in taking the GRE examination or the GMAT examination or the JEEE examinations, either the main or the advanced; whatever you consider, all number of students who take, we consider as a population and that population in the practical sense or the theoretical sense, we will consider is to be infinite.

So, whatever characteristics we get; the mean, standard deviation, mode, variance, all these things are values, which we are interested to find out, because that will give us some information of what the distribution is. Like in the normal case, if you note, we always write mean and the variance. So, these are the basic parameters of the normal distribution which you want to know, because that will give us a lot of information or all information about the normal distribution. Now, as we do not have the population what we do? We have the big sum and small  $n$  number of observations. So, this  $n$  is the small  $n$  of observations which is a subset of the capital, so called infinite population which we have.

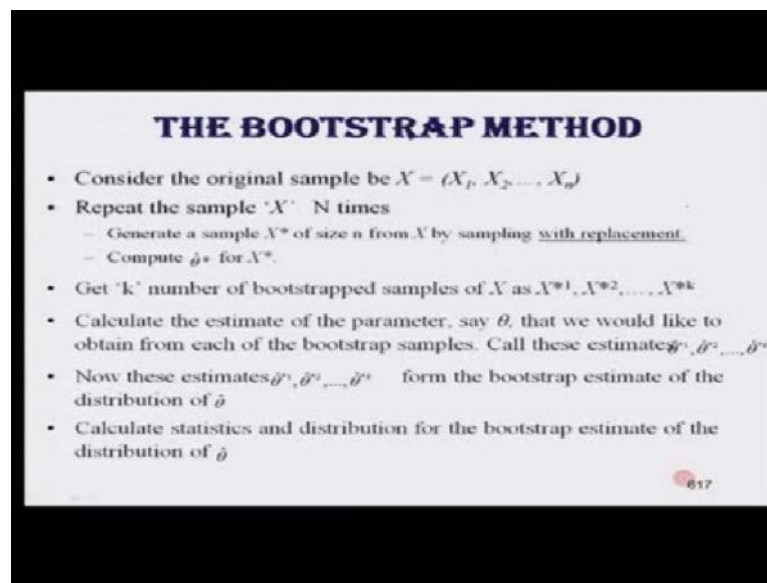
From consider this  $n$  number of observation has been picked up from any unknown observation; that means, we do not know what is the underlying distribution of the observation, from where we are picking up and consider that it to be  $u$ ;  $u$  is not uniform; It is any arbitrary distribution. We call this as the original sample  $x$ . So, let us pick up  $n$  observation. So, we will basically, have observations marked as  $x_1$  to  $x_n$ , where  $x_1$  is the first observation you pick up; first set of observation you pick up from the sample of  $n$  size;  $x_2$  is the second observation you pick up so on and so forth. Now, we want to know, as I mentioned what is that about the parameters for the example; the mean value, the median, the standard deviation, the upper quartile, the skewness courtesies of the original distribution, because till now, we have picked up a sample of size  $n$  from an unknown distribution.

Now a bootstrap sample; what we need is a sample of size  $n$  which you have picked up, which is forming by randomly selecting some set of data with replacement or without

replacement also it can be; without replacement can only be done if the sample size is very huge; can be done such that from the original population we want to pick up such that it will be denoted by  $x^*$ . So, what we actually want to do is that, consider the population is there. We pick up a small sample or a sample  $n$  and we want to pick up observations technically, from this end in such a way by that repeating the set of observations time and again, not the same set, like we pick up again, then note down the characteristics again, put it back.

Again pick up, note down the characteristics and put it back. So, we keep repeating in such a way that the  $x^*$  stars, which you pick up at each and every go, in the long run, the characteristics of the  $x^*$  stars will be in such a way that it will actually, give you the characteristics of the population. Now, you may be asking that why we are not drawing all the information from  $x$ . The reason is that we do not have any other information apart from  $x$ , which was of size  $n$ , because if we had, we would keep repeating the set of observations of small  $n$  size and which would be denoted by  $x$ ;  $x$  is this one, such that whenever you keep picking up the samples, the characteristics of the sample in the long run would actually, be the population, but we cannot do that, because we have only one set of sample of  $n$  size. So, we resample from this  $x$  of size  $n$  such that the characteristics of this  $x^*$  in the long run, slowly mimics the properties of  $x$ , which in a way would definitely give you some or lot of information of the actual population and you want to basically, have some study about the mean, standard deviation, variance so on and so forth of the actual population.

(Refer Slide Time: 08:48)



### THE BOOTSTRAP METHOD

- Consider the original sample be  $X = (X_1, X_2, \dots, X_n)$
- Repeat the sample 'X' N times
  - Generate a sample  $X^*$  of size n from X by sampling with replacement.
  - Compute  $\hat{\theta}^*$  for  $X^*$ .
- Get 'k' number of bootstrapped samples of X as  $X^{*1}, X^{*2}, \dots, X^{*k}$
- Calculate the estimate of the parameter, say  $\theta$ , that we would like to obtain from each of the bootstrap samples. Call these estimates  $\hat{\theta}^1, \hat{\theta}^2, \dots, \hat{\theta}^k$
- Now these estimates  $\hat{\theta}^1, \hat{\theta}^2, \dots, \hat{\theta}^k$  form the bootstrap estimate of the distribution of  $\hat{\theta}$
- Calculate statistics and distribution for the bootstrap estimate of the distribution of  $\hat{\theta}$

So, how we do that? I will give you a simple pictorial and step by step concept; how you do that. Consider the original sample which you have picked up is  $x$  and it is known as  $x_1$  to  $x_n$ . Now, also consider that you are able to pick up such  $x$  number, capital  $n$  number of times. Now, your question would be; would those observations which you are picking up in the first small  $n$  number; would some or all of them would be repeated in the second set of observations we pick up? Answer is yes. It may be with repetition; it may not be with repetition, but some of the observations may be repeated. Now, once you repeat the sample of  $n$ 's times, generate a sample  $x^*$  of size  $n$  from  $x$  by sampling with replacement. Now, note the word with replacement. Reason is very simple. Say for example, you have a box of chip marked 1, 2, 3, 4 and consider theoretically, that is the size of the population.

So, if I ask you a question or let me make it a little bit more explicit. Consider there are two ones, one 2, one 3, one 4. So, total number of such chips is 5. If I ask you, if I pick up one what is the probability? I only pick up one chip at a time. So, the probability of getting 1 is 2 by 5. Probability of getting 2 or 3 or 4 is 1 by 5. Now, see what happens. Consider you pick up two observations; one after the other and you do it with replacement. So, the probability of getting a 1 for the first observation will be 2 by 5. You note it down; keep it back in the box. So, again the total number of size of the orchids in the box is still 5.

Again, if you pick an observation then the probability of getting a 1 still remains 2 by 5 or else if you pick up the first observation as 2, note it down; the probability is 1 by 5, again replacing in the box, again pick up and consider 2 number also comes again. Hence, the probability again, remains as or rather than using the word probability let me use the word of relative frequency. It remains as 1 by 5. Now, if you keep repeating for any set of observations with replacement, the corresponding probability remains the same, but consider that if you are doing without replacement; that means, you pick up one observation, note it down.

Consider for the time being it is 1. So, its probability would have been 2 by 5. Note down 2 by 5 and keep that chip aside. So, now, if you think what are the number of chips which are left in the box; there are only four now; it is 1 1, 1 2, 1 3 and 1 4. So, if you pick up observations, any one of them in the second picking, now the corresponding probability of 1 getting is now, not 2 by 5. It is now 1 by 4. Similarly, the probability of getting 2, which was actually with replacement coming out to be, say for example, 1 by 5. It is not 1 by 5 now. It is now 1 by 4 or say for example, consider that we are doing with replacement of the first chip we pick up is say for example, 2. So, the probability is 1 by 5.

Second time, if I ask you the question; what is the probability of getting a 2? The 2 has already vanished, because it has been kept aside. Then the corresponding probability would of 2 is now 0. Hence the concept of with replacement would be very important. Now, if I ask you the question; can we do it without replacement? Answer would be yes, if and only if, the actual sample size and the number of observation is huge, because in that case, if one observation goes then the corresponding ratio which we find out, that is the total number in the numerator, the total number of events or total number of favorable cases by divided by the total number of such pickings; that ratio would not change much, because say for example, if you find out the ratio of 1000 to 1001, and then you, so, consider this; what I was saying is this. This and this, almost remain in the same. So, if you increase it too large; that means, the observations or pickings also increase very large; that means, numerator and denominator increase; then reducing one number from the numerator then denominator, does not change the corresponding probability.

So, in order to basically overcome that, we will always consider with replacement. Now, once you have done with the replacement what you have? So, you have picked up  $n$  and small  $n$ , and you basically, repeat it capital  $n$  number of times. So, what I am saying is that in the bootstrapping, you note down the first characteristics of the first sample  $n$ , small  $n$ , which we picked up. Characteristics may be mean, median, mode, whatever you want to find out. Note it down. Then you again, replace that in the population. Again, you close your eyes and again, you pick up a chunk.

Then again, remember the chunk is of size  $n$ . Again, you note down the characteristics. Say for example; mean, median, mode, whatever you want to find out. Note it down. Now, you keep repeating it; say for example, a huge number of times. Now, the means are, say for example; the means value which you have and each and every such pickings, you are trying to find out; note them down in a separate column. Now, if you have the mean, say for example, for small size  $n$ ;  $n$  is considered for the time being as 30 and you have found out the average of such 30 observations 200 number of times. So, first picking, you pick up is 30; note down the mean; keep it aside. Again, with replacement, again pick up a size of 30; note down the mean, keep it aside again.

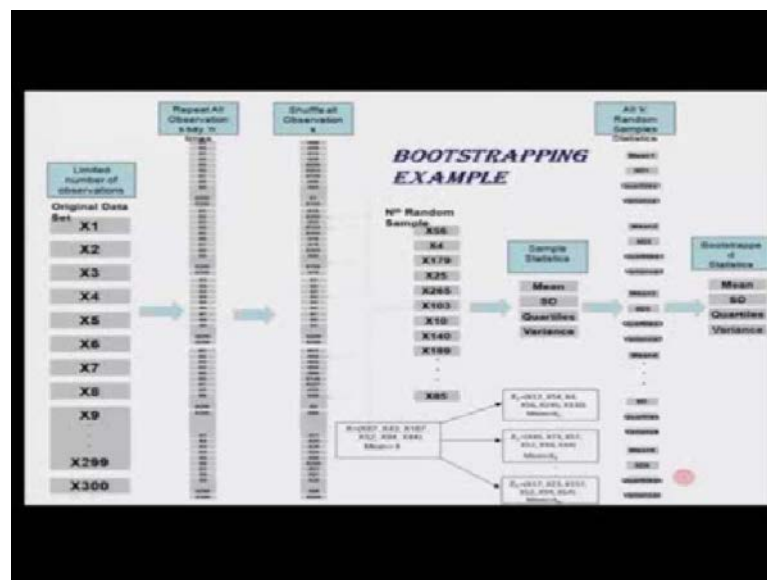
You pick up the third time; note down the mean; keep it aside. You such do replacements and picking up 200 number of times. So, actually what you have? You would basically have such 200 averages, where each and average is basically, the average of such 30 observations. So, technically if I ask you the question that what is the averages of these averages; this is actually the average of 30 into 200 such observations. Now, 30 into 200 is very large number, which means that if I am able to increase this number, not  $n$ , if I am able to increase that number 200 to a large number, then the averages of the averages of whatever characteristics I want to find out, in the long run, will slowly come out to be the population average and that comes out to be true, depending on some characteristics of fulfilling meant if we are able to do. So, you will follow this principle and basically, try to enumerate in a very simple qualitative sense; how this concept of bootstrapping can be done. Now, it says that get  $k$  number of bootstrap samples.

So,  $x$  which you have picked up initially, from that actual in the seconds of  $x$  star. Consider that the first, second, third, fourth reading, which you do; you do  $k$  number of times and it is  $x$  star 1,  $x$  star 2. This 1, 2, 3, 4 are the picking numbers. So, each set which you pick up with  $x$  star,  $x$  star,  $x$  star and they are marked as  $x$  star 1,  $x$  star 2,  $x$

star 3 till x star k. Now, calculate the estimate of the parameters with theta that you will like to obtain and call these estimates as this. So, that means, theta star hat star 1 is the characteristics of that actual small sample of size n, which you have picked up the first time; that one basically knows it. Then if you, when you pick up the second time, it becomes theta star 2. So, this hats which basically means they are the estimated value, corresponding to the first sample which you are drawing, second sample which you are drawing, third sample which you are drawing so on and so forth.

So, you such find out, say for example, such groups of k in number; each group be basically having small n number of observations. Now this estimates from the bootstrap basically, estimate the distribution parameter which is basically, theta hat. Calculate the statistics on the distribution for the bootstrapping estimates which is corresponding to the distribution of theta of parameter which you want to find out. So, the averages, so called averages are estimates which you find out from each sub sample n and you repeat it k number of times. Then the averages of this average, depending on some characteristics would basically, give you the actual parameters values in the long run.

(Refer Slide Time: 17:37)



So, let us understand that with this slide. This slide may be little bit cluttered, but if you look carefully, when you are reading or when you are listening to the notes; obviously, the slide would also be there by the side and I am sure when you are for all the other lectures, till now, which you have studied; obviously, as you listen to my lectures, you

also referred to the notes or the slides. Now, as you are listening to whatever I am saying try concentrating on the slide also.

Now, if you see the left most set of column which is there, it basically gives you the limited number of observations and there are 300 in number. So, consider the 300 is the maximum such sample which I can ever pick up from observation of the population. Consider that I am doing some destructive testing. I am only giving a simple example or you are doing some marketing survey or you are doing some HR related issue survey or you are doing some geographical survey. You are trying to find out some concept of how frequently, earth quake happened or say for example, what is the rainfall, which is happening; whatever it is. Consider that you do not have actual observations and you have the observations. Actual observation means you do not have the whole list. You only have with you 300 number of observations. They are marked as  $x_1$  to  $x_{300}$ . Now, what you do is that you stack this  $x_{300}$  in the column format. So, the first and consider that you are doing it in an excel sheet, a column.

So, let  $x_1$  to  $x_{300}$ ; we placed from cell 1 to cell 300 and then you basically, put pickup that block and cut, copy, paste it repeatedly, say for example,  $k$  number of times;  $k$  can be 200, 300, whatever it is; that is immaterial. Now, what you do is that once you have that if you look at that column, it is basically, 300 is being repeated in blocks. Now, if I ask you to pick up set of observations what will you do? You will basically, pick up observations like if I tell you to do 30 number of times so, you will pick up the first 30, then 31st to 60 and repeat it. Now, once it is there if 300 are there so, 300 if you pick up 30; so, within 10 such pickings, it will be exhausted; the first block. Then again, if you go for the 11th block again, 11th block would basically be repeated by the first one. Then the 12th block would be repeated by the second one, which means the averages would come out to be same, which you do not want it to be. So, what you will do is that the blocks which you have; you jumble; shuffle them totally. So, if you shuffle them totally, how would it look like?

Let us pay attention to the second and the third column; the second column which you have is all those 300, which have been repeated. So, this is the first block which you have; the second block; the third block so on and so forth. So, you have they are n blocked. Now, shuffle all the observations as we noted down is basically, they are in kept in a box and you shuffle it rigorously, and then basically, take down in the column. So, if

you lay down in the column then randomly, this is just snapshot of random  $x$  40 observation is there; then the  $x$  100 and 98 is there; 674 is there; 58 is there so on and so forth.

Now, if you note down, if you now pick up a small block of say for example, 30 observations and you find out the mean lower in the case of if you pick up 30 observations, then the 60; 1st to the 31st to the 60; 61st to the 90. If you keep repeating it, there would not be any instance, where there may cases at or by 0 probability. Be rest assured there would not be any instances, where they would be repetition. So, now, what you have been able to do is that you have been able to jumble them in such a way that such repetition would not occur and if that does not occur, then your first major task is already of our trying to basically, characterize from the samples; some characteristics of the population such that you can proceed and try to find out the average characteristics in a very nice manner. So, now, I will take a pause, end this class. In the next class, I will basically repeat it in a much more greater details with an example, which will make things much clearer to you.

Thank you.