

Randomized Algorithms
Prof. Benny George Kenkireth
Department of Computer Science & Engineering
Indian Institute of Technology, Guwahati

Lecture – 30
DNF counting

In today's lecture we will see how we can approximate the number of satisfying assignments for a DNF formula. So, the DNF formula is a formula in Disjunctive Normal Form which means you have clauses $C_1, C_2, \dots, C_m$ and each of these clauses is a disjunction of literals.

(Refer Slide Time: 00:41)

Approximating the # of satisfying assignments of a DNF

$$\phi = (l_{1,1} \wedge l_{1,2} \wedge l_{1,3}) \vee (l_{2,1} \wedge l_{2,2} \wedge l_{2,3} \wedge l_{2,4}) \vee \dots \vee (l_{m,1} \wedge l_{m,2} \wedge l_{m,3} \wedge l_{m,4})$$

$C_1 \qquad C_2 \qquad \qquad C_m$

$\phi = x_i \text{ or } \neg x_i$
 $n \text{ variables } x_1, \dots, x_n$

Polynomial Approximation Scheme.

I : instance of a problem $\text{NP} \leftarrow \text{AC}^0$

$A(I)$: # of satisfying assignments (or Hamiltonian Path, Matching)

ϵ : Real number $(0, 1)$ $\ln\left(\frac{1}{\epsilon}\right)$ $\left(\frac{1}{\epsilon}\right)$

$$(1 - \epsilon) A(I) \leq Y \leq (1 + \epsilon) A(I)$$

So, each formula is of each clauses are the from l_1 and l_2 and l_3 . So, on where each l_i is a literal; that means, it is a formula it is a of the form x_i or not x_i . So, let us say l any general literal is of the form its x_i or not x_i for some i where x_i is the variable. Now, what we are interested in is counting the number of assignments. So, suppose there were n variables x_1 to x_n there are 2^n different assignments possible, we need to count the number of assignments which are satisfying this particular DNF in other words the number of assignments which causes the formula to evaluate to 1 ok.

So, first of all we need to understand what is the meaning of approximating in this context. The first notion that we will look at is polynomial approximation scheme. So, suppose I is an instance of a problem; that means, one particular formula and we are

interested in $A(I)$ that is. So, here you could think of it as number of satisfying assignments or let us say Hamiltonian path matching etcetera; number of Hamiltonian paths in a given graph a number of matches in the given graph. So, $A(I)$ denotes the number of assignments which can satisfy the particular problem or some parameter associated with each of these instance. For each of these instance if you think of it as an a problem belonging to the class non-deterministically NP, then you can look at $A(I)$ as the number of certificates or number of y 's which will satisfy the formula corresponding to the non-deterministic problem.

So, we often write it as $A(x, y)$ we write a predicate associated with each NP problem instance $A(x, y)$ where y is what we call as a certificate. So, $A(I)$ you can look at as the number of certificates and when we say that we have a polynomial approximation scheme; that means, for any instance of the problem and for a given number epsilon. So, epsilon we can think of it is a real number, let us say between 0 comma 1; we can output something let us call it as Y this lies between. So, the actual value was $A(I)$, then our output will lie between $1 - \epsilon \times A(I)$ and $1 + \epsilon \times A(I)$.

So, if by $A(I)$ if we denote the number of satisfying assignments and your Y the output of your algorithm must lie between $1 - \epsilon \times \text{actual count}$ and $1 + \epsilon \times \text{actual count}$. And, if you can have such an algorithm the time taken should be polynomial in the input size and there is an epsilon here. So, if you look at epsilon as a number it will take let us say $\log \frac{1}{\epsilon}$ by epsilon we are taking $\frac{1}{\epsilon}$ by because it is a number less than between 0 and 1. So, $\log \frac{1}{\epsilon}$ is the amount of space required to present this number or the input size is $\log \frac{1}{\epsilon}$.

So, we could have said that the algorithm should run in polynomial time it being polynomial both in the input size and in $\log \frac{1}{\epsilon}$, but in this case we actually allow a little more leeway say happy if it runs in $\text{poly} \frac{1}{\epsilon}$.

(Refer Slide Time: 05:49)

ϵ : Real number $(0,1)$

$$(1-\epsilon)A(I) \leq Y \leq (1+\epsilon)A(I)$$

$\text{Poly}(n, \frac{1}{\epsilon})$
 Fully Poly. App. Scheme (FPAS)

δ : Failure Probability.

$$\Pr((1-\epsilon)A(I) \leq Y \leq (1+\epsilon)A(I)) \geq 1-\delta$$

$\ln(\frac{1}{\delta})$ $\text{Poly}(n, \frac{1}{\epsilon}, \ln(\frac{1}{\delta}))$

Fully polynomial Randomized Approximation Scheme:

So, when we say polynomial approximation scheme with approximation as let us say ϵ mean epsilon we will say that it has to run in poly, if you think of the input size as n then should be poly n and $1/\epsilon$. And if we could find such an algorithm for every possible epsilon then we will say that is a Fully Polynomial Approximation Scheme and we call it as FPAS. Because that is one notion of approximation here we insist that the algorithm has to be deterministic.

So, deterministically an output has to be generated which lies very close to the actual answer, it deviates with the actual answer by an amount epsilon. This is the first notion regarding approximation. The second notion that we will have is we will allow our answers to be let us say randomized. So, we will think of randomized approximation schemes. So, everything remains pretty much same as polynomial approximation scheme, but we will insist that this output which should lie between these they need not always lie there, but it lies with a certain probability delta.

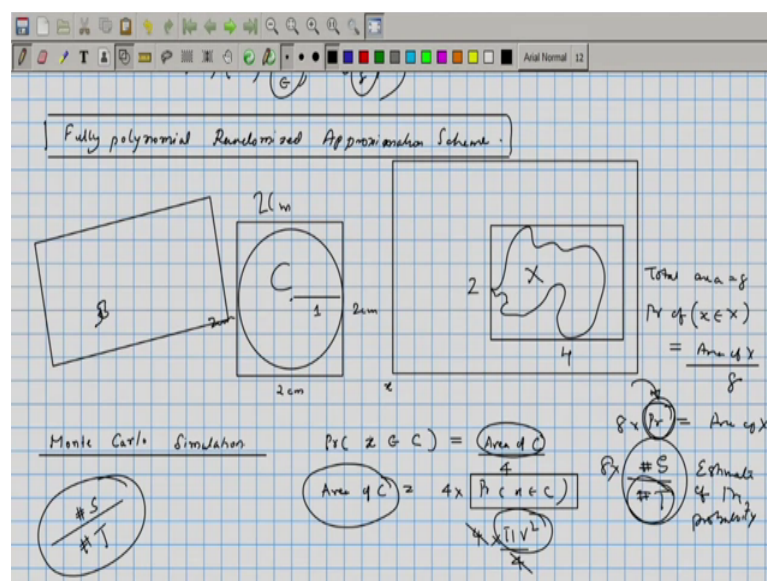
So, the additional parameter is delta if we say that this is the failure probability when there is the chance that the algorithm fails. So, we will say that the probability that the output lies between $1 - \epsilon$ times $A(I)$ and $1 + \epsilon$ times $A(I)$ this is going to be greater than $1 - \delta$; means the success probability is at least $1 - \delta$. And if we could do this for every delta, delta again being a real number between 0 and 1, the input size is $\log 1/\delta$. So, we will insist that the algorithm runs in poly in

polynomial time polynomial n ; polynomial n let us say $1 - \epsilon$ and $\log 1/\delta$. And, if we could do this for every ϵ and every δ then we will say that we have a fully polynomial randomized approximation scheme.

So, this would be, so, when we look at any approximation problem this would be the best case scenario that is if we can find an algorithm which satisfies these conditions. And, if we are unable to do that we will try and see if we can do it for just mean to say certain ϵ s, but δ we still want the failure probability we want it to be arbitrary low that is for any failure probability we should essentially have an algorithm. And its running time should be bounded by polynomial in $\log 1/\delta$, but regarding the approximation factor where a little less stringent ok.

So, this is what we mean when we say we have an approximation algorithm a fully polynomial randomized approximation scheme for a particular problem. Now, let us keep this definition inside for a minute and think about a simple problem.

(Refer Slide Time: 10:03)



So, let us say we have a circle whose radius is 1 ok, we need to estimate the area of this could have been any other geometric figure how do we estimate the area? So, there is a technique called as Monte Carlo simulation ok. What does the basic idea? Fit this circle inside the square ok. So, this is a square of size 2 centimeter. So, if we randomly choose a point from the square the probability that. So, let us call the randomly chosen point let

us say x , probability that x belongs to C is equal to area of C divided by the area of this square which is 4 ok.

Now, this will also mean that area of C is equal to 4 into probability that x belongs to C ok. So, if we could estimate this probability exactly that would mean that we can compute the area of c . So, now, probability that x belongs to C is going to be nothing, but πr^2 square by 4 and this multiplied by 4 gives us πr^2 square in some sense we already knew this ok. So, we just plugged in that to compute the probability, but let us say it is a more irregular figure ok.

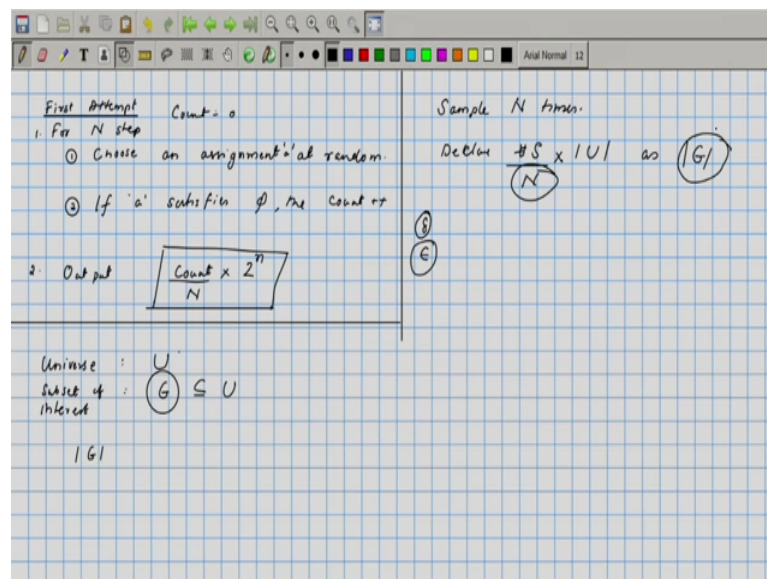
So, we enclose it in a in a known area let us say this is 2 and this is 4, the total area is 8 probability of x belonging to let us say x that region let us call it as x is equal to area of x divided by 8. So, 8 times that probability is going to be area of x . Now, how do we estimate this probability? One thing we can do is we can just repeatedly toss ok. So, let us say independently we toss the 100 coins and out of these 100 points some of them will lie inside and some of them may lie outside. Should have a method to figure out if the randomly chosen point lies inside or outside the region of interest. If we could do that the number of points which falls inside let us say that is success.

So, number of success by number of trials this is a good estimate of the probability you can argue that no more theoretical fashion, but you can intuitively feel that the fraction of points which lie inside is a good indicator of the probability that the point chosen lies inside ok, 8 multiplied by that gives us the area. So, that is a way to compute the area this is this same method can be extended to counting or approximate counting. So, this gives an approximate area if we could compute the probability exactly then the area computations would have been perfect it would have been exact, but in we cannot compute the probability exactly since we are unable to do that we estimate it and in the estimation there can be errors.

The estimation by central limit theorem certainly is going to be better and better as the number of trials becomes very large. Also you can see that if we had chosen a much larger area to contain this particular area of interest then the number of trials will have to increase. If you were to get let us a reasonable approximation to increase the accuracy it is best that the area which is holding this particular region is as close to the actual area as possible ok.

So, that will give better and better bounds, on the number of trials that we have to do if this was a tiny region inside a large let us say box then we do not expect to hit this region too often. But the quality of the answer that is if you look at the number of success that is points which are chosen and are lying inside the area of interest divided by the total number of trials. This answers accuracy depends on the number of trials in the number of trials will have to be larger and larger if this area was large in comparison to the area of interest ok.

(Refer Slide Time: 15:33)



So, with these insights we will try and have an approximation algorithm for DNF. So, first attempt this will be our algorithm choose an assignment at random to do this for n steps if. So, let us call this assignment a if a satisfies phi then a count plus plus. So, initially we will have count equals 0 and then the last step would be number of. So, we are interested in the count. So, we will output its a count divided by n this is a fraction of satisfying assignments multiplied by total number of assignments we had 2 raise to n assignments ok.

So, we can round this off to an integer value and that is going to be our output. This is the algorithm any good what are the probability, what are the issues with this what are the I mean what can we say about the success probability and the quality of answer. So, let us look at a more general problem; let us say we had a universe which we call as U and let us say we have the set of interest which we call as G.

So, G is a subset of U , we are interested in computing size G ok. So, let us call this U size of U to be. So, this is ok. So, now, if we follow a similar algorithm Monte Carlo algorithm to compute the size of G then we would have sampled n times. So, the algorithm is simple sample n times declare number of success divided by N times size of U as let us say size of G ok.

So, when we do this what all could go wrong. How many times will we have to run the algorithm n is what we need to determine in order to get a reasonable estimate. So, suppose we say that we want let us say an algorithm which has error probability of at most δ and let us say approximation factor as ϵ .

Then how many times will we have to run this algorithm in order to get these parameters. The problem that could be there is G could be first of all very large or very small if G is very small compared to the size of the universe then we expect that this is not going to be a good algorithm and we will see the quantitative dependency between these.

(Refer Slide Time: 19:53)

$$\begin{aligned}
 Y &= \text{\# of success in } N \text{ trials} \\
 &= Y_1 + Y_2 + \dots + Y_N \quad (Y_i \text{ are i.i.d}) \\
 \tilde{G} &= \frac{Y}{N} \times |U| \quad E[Y] = N \cdot E[Y_i] = N \cdot p \quad (p = \frac{|G|}{|U|}) \\
 P_Y \left((1-\epsilon)|G| \leq \tilde{G} \leq (1+\epsilon)|G| \right) \\
 &= P_Y \left((1-\epsilon)p|U| \leq \tilde{G} \leq (1+\epsilon)p|U| \right) \\
 &= P_Y \left((1-\epsilon)p|U| \leq \frac{|U| \times \sum Y_i}{N} \leq (1+\epsilon)p|U| \right) \\
 &= P_Y \left((1-\epsilon)pN \leq \sum Y_i \leq (1+\epsilon)pN \right)
 \end{aligned}$$

So, we will define few random variables Y is equal to number of success in n trials. So, it could be any number between 0 to n and we can think of this is Y_1 plus Y_2 plus Y_N where each these are i.i.d random variables ok. The answer that we would have declared is. So, size of G we will declare it to be equal to Y by N into size of U ok, estimate of G that we will give is Y by N times U we want to know. So, note that we can calculate expectation of Y that is going to be equal to N times expectation of each Y_i and each Y_i

we can say that its expectation is ρ , where ρ is equal to $\frac{G}{U}$ size of G by size of U

The probability that let us call the estimate as $\tilde{G}$; so, $\tilde{G}$ lies between $1 - \epsilon$ times G and $1 + \epsilon$ times G , this is a probability that we need to estimate ok. And then we want this probability to be let us say greater than $1 - \delta$; that means, the failure probability is going to be δ and for that what should be the N that is what we will compute.

So, this probability it is going to be equal to probability that $1 - \epsilon$ and size G is going to be ρ into U and this is equal to probability that $1 - \epsilon$ times ρU lies between $\tilde{G}$ is nothing, but U divided by N into summation Y_i . So this probability is nothing, but the random variable Y which is summation of Y_i 's lies between $1 - \epsilon$ times ρN and $1 + \epsilon$ times ρN and ρN is the expectation of this random variable. So, we can simply write this as.

(Refer Slide Time: 23:07)

Handwritten mathematical derivation on a grid background:

$$\begin{aligned}
 & \Pr \left((1-\epsilon)G \leq \tilde{G} \leq (1+\epsilon)G \right) \\
 &= \Pr \left((1-\epsilon)\rho U \leq \tilde{G} \leq (1+\epsilon)\rho U \right) \\
 &= \Pr \left((1-\epsilon)\rho U \leq \frac{U}{N} \sum Y_i \leq (1+\epsilon)\rho U \right) \\
 &= \Pr \left((1-\epsilon)\rho N \leq \sum Y_i \leq (1+\epsilon)\rho N \right) \\
 &= \Pr \left(\left| Y - \rho N \right| \leq \epsilon \rho N \right) \\
 & \Pr(Y - \rho N) \geq \epsilon \rho N \\
 & \Pr \left(\left| Y - \rho N \right| \geq \epsilon \rho N \right) \leq 2 \cdot e^{-\frac{\epsilon^2 \rho^2 N}{3}}
 \end{aligned}$$

Chernoff Bound:

$$\Pr(|X - \mu| \geq \delta \mu) \leq 2e^{-\frac{\delta^2 \mu}{3}}$$

Pr(Failure) = δ

$$\delta = 2e^{-\frac{\epsilon^2 \rho^2 N}{3}}$$

$$\frac{\epsilon^2 \rho^2 N}{3} = \ln\left(\frac{2}{\delta}\right)$$

$$N = \frac{3 \ln\left(\frac{2}{\delta}\right)}{\epsilon^2 \rho^2}$$

The probability that Y minus expectation of Y ρN the modulus of that, is less than ϵ times ρN . So, we can straight away apply Chernoff bound ok, Chernoff bound applies to the complement of this event. So, let us if you look at the complimentary event so, probability that Y minus ρN the absolute value is greater than or equal to ϵ times ρN .

Chernoff bound says if you sample Y_i 's independently with ρ then there some x will deviate from its expectation by an amount more than δ . U is going to be less than $2e$ raised to minus say δ^2 by 3 times μ ok. So, this is the form of Chernoff bound that we can just straightaway use here and that will give us the probability that Y minus ρN is greater than ϵ times ρ times N .

So, this is μ and this is going to be δ . So, this is going to be less than 2 into e raised to minus say ϵ^2 by 2 into sorry ϵ^2 by 3 into your μ is going to be ρN ok; so, that is that is what we obtained. So, we will look at the failure probability the failure probability is less than this. So, if we can tolerate the phase the probability of failure if it is equal to δ at most δ and we can equate these two quantities to determine what should N be ok. So, equating those if you said δ this is a failure probability this is equal to $2e$ raised to minus ϵ^2 by $3\rho N$, we can just take log.

So, ϵ^2 by $3\rho N$ is equal to $\log 2$ by δ . So, N will be equal to 3 sorry by this is ϵ^2 by $\log 2$ by δ by ρ . So, the number of trials that you have to do to get a probability of failure no more than δ is going to be this much.

(Refer Slide Time: 26:19)

$$P(|Y - \rho N| \geq \epsilon \rho N) \leq 2e^{-\frac{\epsilon^2 \rho N}{3}}$$

$$N = \frac{3 \ln(2/\delta)}{\epsilon^2 \rho}$$

$$\# \text{ of trials reqd (error } \leq \delta) = \frac{3 \ln(2/\delta) \times (1/\rho)}{\epsilon^2}$$

① ②ⁿ

So, number of trials required will be equal to 3 by. So, this is for error less than δ this is 3 by ϵ^2 by $\log 2$ by δ into 1 by ρ . So, if this ρ was very small, ρ was

the fraction of inputs that we had tried which were good, rho is size of G by size of U, G was the set that we were interested in and U was the universe. If rho was very small then this is not a good algorithm. So, in DNF there could be a DNF whose there is just 1 satisfying assignment amongst let us say 2^n possible assignments and then the number of trials record could be large ok.

So, that is the reason why we cannot use the simple Monte Carlo algorithm to solve the approximation problem for DNF counting; so, what is the way out. We were now sampling from here should not expect that this is going to be a good sampling because they said that we had was very tiny compared to the region from where we are sampling. If you could find something which is small enough to encompass this then we could do a better approximation and that is what we will see in the next part ok.

So, instead of sampling from the set of all possible assignments we will try and sample from a small region ok. So, what does these regions going to be ok. So, let us look at this structure of DNF formulas little more carefully.

(Refer Slide Time: 28:29)

Handwritten notes on a grid background explaining DNF formulas and their components:

- Formula: $\phi = (l_{11} \wedge l_{12} \dots) \vee (l_{21} \wedge \dots) \vee \dots \vee (l_{m1} \dots \wedge l_{mk})$
- Definition: $H_i \triangleq \# \text{ of assignments satisfying the } i^{\text{th}} \text{ clause.}$
- Union: $\bigcup H_i$
- Clause example: $(x_1 \wedge x_2 \wedge x_k) \quad m-k$
- Count: 2^{m-k}
- Boxed note: $m \times \# \text{ of satisfying assignments}$
- Right side notes:
 - Multiset of H_i s.
 - $H_i = \{ (x_i, i) \mid \text{satisfies } C_i \}$
 - $\bigcup_{i=1}^m H_i = \text{Multiset of satisfying assignments.}$

So, if the formula was of the form let us say l_{11} or $l_{11} l_{12}$ so on or l_{21} and so on or l_{m1} and $l_{m1} l_{m2}$. Can we try and count how many satisfying assignments can be there they will split into parts. So, let H_i denote the number of assignments satisfying the i th clause ok.

And what we are really interested in this union of H_i , can we determine any of these if we find the H_i and we could maybe try inclusion exclusion formula, but since there are m clauses the inclusion exclusion formula would essentially have 2^m different terms. So, we cannot really go by that route that is going to be computationally very costly.

So, what can we really do here. Let us see if we can just compute each H_i . If you look at each H_i they are of the form let us say x_1 and x_2 and x_k maybe not and. So, how many satisfying assignments are there for this? Well, there is only one possible value of these variables that will satisfy the particular clauses H_i , but the other $n - k$ variables can be arbitrarily assigned.

So, each assigned count is extremely simple, the number of satisfying assignments for the i th clause can be simply computed as 2^{n-k} , where k is the number of variables that is appearing in the clause or number of literals appearing in that particular clause. Note that if any clause appears more than once we just throw them away that simplification can be done some variable appears both as positive and negative then of course, that clause is not satisfiable therefore, we can just throw it away.

So, 2^{n-k} under these assumption 2^{n-k} is the number of assignments which satisfies the i th clause. If we wanted to take 2 of them simultaneously then it becomes little more complicated, but that also you can compute the formula. But the real issue here is we cannot compute all of them there are too many of them 2^m of them. So, we cannot apply inclusion exclusion.

What we will do is instead of sampling from the entire when we run our Monte Carlo simulation or the Monte Carlo algorithm instead of sampling from 2^n possible assignments. We will sample from another sample space which whose size is not too large it is at most m times the number of satisfying assignments ok.

So, if we could find a space which contains at most m times the number of satisfying assignments then by this method ρ is going to be at most m sorry $1/m$ so, $1/m$ is m . So, that is going to work in $3 \log 2 / \epsilon^2 \log 2 / \delta$ ok.

(Refer Slide Time: 32:35)

$$= P_Y \left((1-\epsilon) p_N \leq Y \leq (1+\epsilon) p_N \right)$$

$$= P_Y \left(|Y - p_N| \leq \epsilon p_N \right)$$

$$P_Y(|Y - p_N| \geq \epsilon p_N) \leq 2e^{-\frac{\epsilon^2 p_N}{3}}$$

$$P_Y(|Y - p_N| \geq \epsilon p_N) \leq 2e^{-\frac{\epsilon^2 p_N}{3}}$$

$$\# \text{ of trials reqd (error } \leq \epsilon) = \frac{3}{\epsilon^2} \ln\left(\frac{2}{\delta}\right) \times \left(\frac{1}{p}\right)$$

$$N = \frac{3 \ln(2/\delta)}{\epsilon^2 p}$$

If $p \geq \frac{1}{m}$ then simple Monte Carlo works.

So, if rho is equal to 1 by m or if it is greater than or equal to 1 by m then simple Monte Carlo works. So, that sample space is what we will construct its a simple sample space to construct. Let us imagine the multi set of a H_i 's. So, H_i here we denoted it as a number we can overload it and say that H_i is a set containing all the assignments which satisfy the i th clause.

So, when you say multi set we will look at the assignment comma i ok. So, consider the pair v comma i ok. So, let us say we will again H_i itself H_i is equal to the set of all v comma i such that v satisfies the clause i ok. So, clearly if you take union of H_i s that is going to be multi set of you can put it i mean in 1 to 1 correspondence with the multi set of satisfying assignments that is for every satisfying assignment we are counting the number with its multiplicities.

If the assignment v satisfies both let us say clause 5 and 3 and 25 it is counted 3 times, this is once for each clause it satisfies. Now we could define another set which is basically so, this is going to be our universe. So, union H_i have which we will call as H .

(Refer Slide Time: 34:45)

$$\text{Universe} = \bigcup H_i = H$$

$$|H| = \sum |H_i|$$

$$\quad \quad \quad \uparrow$$

$$\quad \quad \quad 2^{n-k_i}$$

We will find a G s.t. $|G| = \# \text{ of satisfying assignments}$ &
 $G \subseteq H$.

Note: $|H| \leq m \cdot \# \text{ of satisfying assignments}$ where
 m is the # of clauses.

$$\frac{|G|}{|H|} \geq \frac{1}{m}$$

$$|H| \leq m \cdot |G|$$

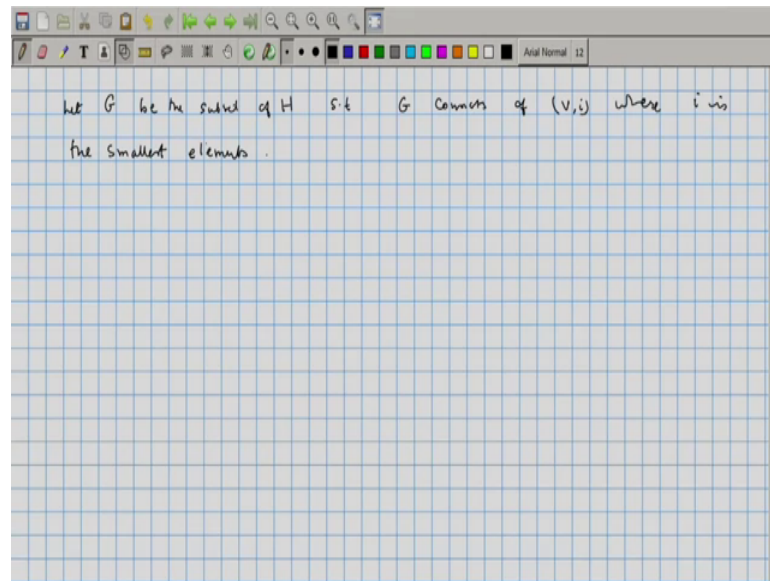
$$\frac{|G|}{|H|} \geq \frac{1}{m}$$

So, universe is going to be union H_i which we will call as H ok. Now in this universe when the count means obtaining the count of this universe is simple because that is just going to be sum over sizes of H_i s because these each H_i has a unique i in the second coordinate. So, you look at each H_i the elements inside that H_i s the first component is an assignment and the second component is a number indicating i . So, it is different for different classes. So, this summation is just 2 raise to n minus k_i where k_i is the number of literals appearing there. So, size H_i can be easily computed.

Now, we will define G to be the subset. So, find G such that size G is equal to number of satisfying assignments and additionally G is a subset of H ok. Now number of satisfying assignments and H what is the relationship? So, H is not going to be so, it surely will be less than m times number of satisfying assignments where m is the number of clauses. This is because you take any satisfying assignment it can be present in at most m clauses ok.

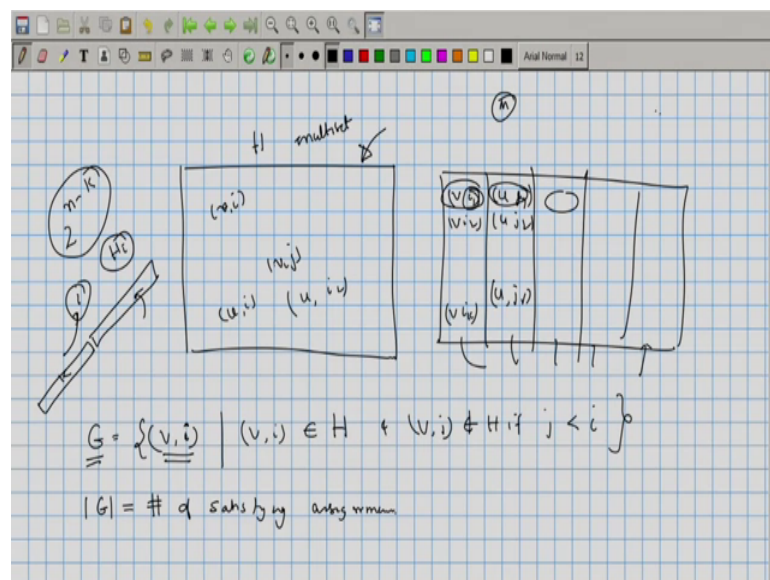
Therefore each quantity can each assignment can appear at most m times therefore, total number the size the multiset is bounded by m times the number of assignments. So, when you get these the set G which is equal to number of satisfying assignments and this is the subset of H size of G by size of H is going to be greater than 1 by m ok. So, G by H is going to be greater than 1 over m and that will take care of all the requirements ok. Now how do we construct this set? It is simple.

(Refer Slide Time: 37:45)



So, let G be the subset of H such that G consists of elements of the form $v i$ where i is the smallest element may be i will just rewrite it.

(Refer Slide Time: 38:47)



So, we have this set H which is the multi set. So, this set these various elements $v i v j u i u i 2$ etcetera. Now we can think of this as split into various assignments. So, let us say $v i 1 v i 2 v i k$. So, v is an assignment which appears k times and then let us say u appears say $j 1 u j 2 u j r$ and u appears r times and so on. So, we want to pick one element each

of these for that we will pick this smallest j ok. So, G will basically be the set of v_i such that v_i belongs to H and v_j does not belong to H if j is less than i ok.

So, pick the smallest elements on the basis of the second coordinate for each assignment. So, the number of elements that we pick is going to be equal to the number of satisfying assignments. So, this is easy G is equal to number of satisfying assignments ok. Now, how do we check if a particular assignment is in G ? Now let us sample one of the elements and we sample uniformly from H . Well, we can I mean the moment of thought would tell you that from this multi set if you wanted to sample that is not a difficult thing because you can first choose a particular clause uniformly at random and then from that.

So, if you choose let us say this the second coordinated random. For that particular second coordinate, all the satisfying assignments of that there are 2^{n-k} of them. They can be found by fixing the first k coordinates first k here means whatever of the variables it is appearing in that particular clause those who set to the value as prescribed by the clause. And for the other variables you just choose randomly ok; uniformly at random toss a coin and fix the values of the other variables and that will give you a uniform sampling from each of these H_i 's and that basically translates into a uniform sampling for the entire set.

So, you can uniformly sample from H . Further once you have picked one particular element, how do you check whether that belongs to the set G . Well if your i was some particular element, for every element smaller than i , there are at most m of them you can check whether they belong to any of these smaller elements. There are only m of them. So, you can actually do this computation that computation polynomial time.

So, this verification of whether a particular assignment α belongs to G can be done in polynomial time and therefore, we are done with the requirements we could sample from this set x uniformly at random. The size of H is known and we can in polynomial time figure out whether the sampled element belongs to G or not and therefore, that will give us a polynomial time approximation scheme for computing the number of satisfying assignment of DNF formula.

We will stop here and we will continue with other approximation algorithms in the next class.