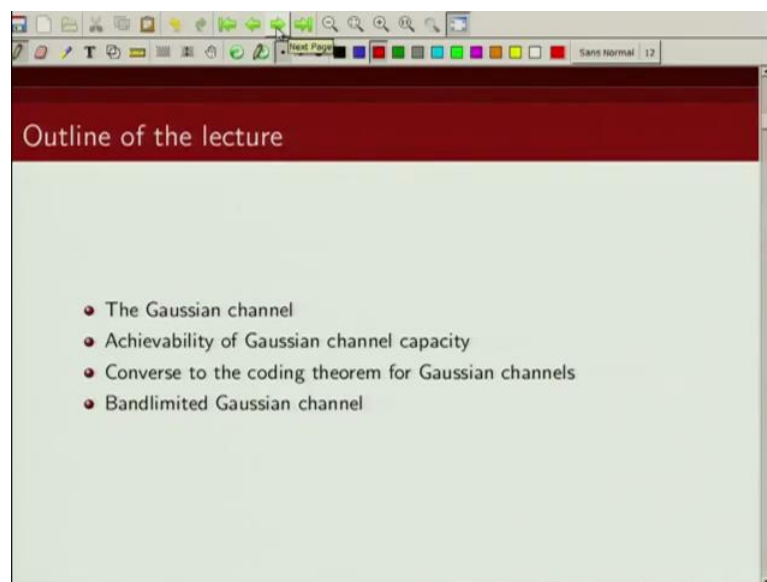


An Introduction to Information Theory
Prof. Adrish Banerjee
Department of Electronics and Communication Engineering
Indian Institute of Technology, Kanpur

Lecture – 12A
Gaussian Channel

Welcome to the course on an introduction to information theory. So, today we are going to talk about Gaussian channel in the last class we talked about entropy for continuous random variable one of the important channels continuous values channel is this Gaussian channel.

(Refer Slide Time: 00:47)



And today we are going to talk about capacity of a Gaussian channel, we will first talk about capacity of a Gaussian channel and we are going to prove the achievability of Gaussian channel capacity subsequently. We are going to show the converse to the coding theorem for Gaussian channel and finally, in this lecture we are going to talk about band limited Gaussian channel and we will compute the capacity of band limited Gaussian channel.

(Refer Slide Time: 01:12)

The Gaussian Channel

- The capacity of the Gaussian channel with power constraint P and noise variance N is given by

$$C = \max_{EX^2 \leq P} I(X; Y) = \frac{1}{2} \log \left(1 + \frac{P}{N} \right)$$

where maximum is attained when $X \sim N(0, K)$.

Proof:

- Capacity is given by

$$C = \max_{p(x): EX^2 \leq P} I(X; Y)$$

So, as I mentioned Gaussian channel are very important continuous value channel. If we write its corresponding discrete motion, we can write it. Let us see the input is x_i and noise is additive which is Gaussian distribution denoted by z_i what we get as output is y_i this is a discrete motion of a continuous value Gaussian channel. So, first result that we are going to show you is the capacity of the Gaussian channel, with power constraint p and noise variance n is given by this expression. Where the maximum is attained when, x is Gaussian distributed with 0 mean and variance k . Now let us first look at what is the capacity. If we do not have these constraints power constraints or if you do not have noise it is very easy to see, if the noise is 0 then, whatever we send will be reliably received. So, the channel capacity is infinite in that particular case similarly if there is no power constraint.

We can make other input arbitrarily large such that it is received correctly at the receiver. So, if we do not have this power constraint then also the capacity of this channel is infinite. Now let us look at the capacity of this Gaussian channel, when we do have a power constraint and we do have a noise variance given by n . So, from the definition of channel capacity its maximum information between the input x and output y and this maximization is taken over all input distribution p of x . Remember, we have a power constraint here it is as an average power constraint. So, expected value basically of x

square is less than p . So, we need to compute maximum mutual information under all input distribution given this power constraint, from the definition of the mutual information.

(Refer Slide Time: 03:46)

The Gaussian Channel

- We can write $I(X;Y)$ as

$$\begin{aligned}
 I(X;Y) &= h(Y) - h(Y|X) \\
 &= h(Y) - h(X+Z|X) \\
 &= h(Y) - h(Z|X) \\
 &= h(Y) - h(Z)
 \end{aligned}$$
 $Z \sim N(0, N)$
- We have $h(Z) = \frac{1}{2} \log 2\pi e N$. And we also have

$$\underline{EY^2} = \underline{E(X+Z)^2} = \underline{EX^2} + \cancel{2EXEZ} + \underline{EZ^2} = \underline{P} + \underline{N}$$

We can write mutual information between x and y as differential entropy of y minus differential entropy of y given x . Now what as I said y is the output of the channel which is nothing, but input x plus additive noise additive y additive Gaussian noise, which is z . So, i can write this as differential entropy of y minus differential entropy of x plus z given x . Now what is the uncertainty also what is the entropy of x plus z given x that only depends on z given x . So, we can write this expression as differential entropy of y minus differential entropy of z given x , now z is which is noise is independent of the signal x . So, the differential entropy of z given x is equal to differential entropy of z because z the noise is independent of the signal x .

So, we can write this mutual information between x and y as differential entropy of y minus differential entropy of z . Now we know that z is Gaussian distributed and if we have a Gaussian distributed random variable. We know what is its differential entropy is. So, if we assume that z is Gaussian distributed its mean. Let us say 0 and variance n then, the differential entropy of z is given by half log of $2\pi e$ times n , we can compute the

expected value of y square y is nothing, but x plus z . So, this can be written as expected value of x plus z square which is nothing, but expected value of x square plus expected value of z square plus 2 times expected value of x and expected value of z . Now since the expected value of z is 0 this term will be 0. So, what we would get is expected value of y nothing, but expected value of x square plus expected value of z square now expected value of x square is p and expected value of z square is n . So, expected value of y square is p plus n .

(Refer Slide Time: 06:52)

The Gaussian Channel

- Thus we have

$$I(X; Y) = h(Y) - h(Z)$$

$$\leq \frac{1}{2} \log 2\pi e(P + N) - \frac{1}{2} \log 2\pi eN$$

$$= \frac{1}{2} \log \left(1 + \frac{P}{N} \right)$$

- Maximum is attained when $X \sim N(0, P)$.

Handwritten notes in red:

$$h(Y) \leq \frac{1}{2} \log 2\pi e(P + N)$$

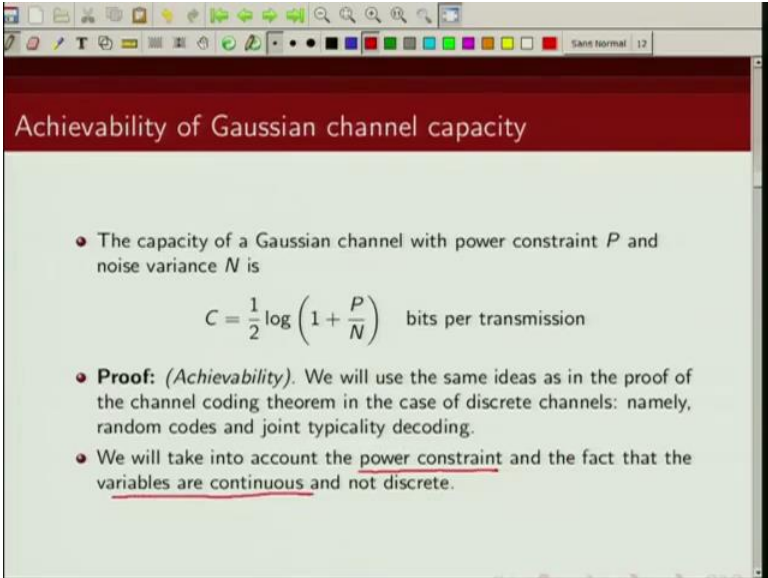
$$Y = X + Z$$

Now we have just shown that the mutual information between x and y can be written as differential entropy of y minus differential entropy of z . Now we know, what is the differential entropy of z that is because, z is a normal distributed with 0 mean and variance n , the differential entropy of z is given by this. We also know if y is a random variable which has variance p plus n then the differential entropy of y is upper bounded by the differential entropy of Gaussian random variable. So, we can upper bound the differential entropy of y by the differential entropy of a Gaussian random variable with same variance and in this case expected value of y square is p plus n . So, we can upper bound this differential entropy of y by differential entropy of a Gaussian random variable with variance given by p plus n , now variance given by p plus n .

So, now, basically what we have said this results follows from a theorem, that we have proved for differential entropy that is y is the random variable with variance p plus n then its differential entropy is upper bounded by differential entropy of a Gaussian random variable with same variance. So, clearly this is equal when y is also Gaussian and y is Gaussian, when x is also Gaussian because y is x plus z we know z is Gaussian distributed, If x is also Gaussian distributed then y will be sum of 2 Gaussian distributed random variable will also be a Gaussian distributed, in that case y will also be a Gaussian distributed random variable and in that case this would have been equality.

So, combining the terms, this is $\log e$ minus $\log v$ kind of term which can be combined. So, that becomes this. So, what we have shown is mutual information is upper bounded by half log of one plus p times n now remember the maximum is attained. When x is also Gaussian distributed with mean 0 and variance p in that case as i said y is also going to be Gaussian and this inequality would be equality. So, what we have shown is capacity of a Gaussian channel with power constraint p and noise variance n is given by this expression. Where the maximum is attained when input x is also Gaussian distributed with mean 0 and variance k .

(Refer Slide Time: 10:20)



Achievability of Gaussian channel capacity

- The capacity of a Gaussian channel with power constraint P and noise variance N is

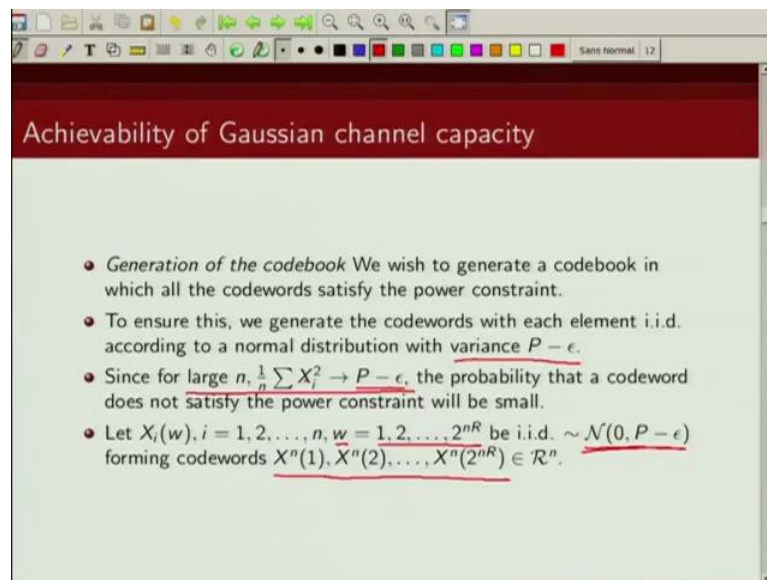
$$C = \frac{1}{2} \log \left(1 + \frac{P}{N} \right) \text{ bits per transmission}$$
- **Proof: (Achievability).** We will use the same ideas as in the proof of the channel coding theorem in the case of discrete channels: namely, random codes and joint typicality decoding.
- We will take into account the power constraint and the fact that the variables are continuous and not discrete.

Now, we are going to show the achievability of this. So, all as long as we choose rate less

than this capacity we are going to get reliable communication. So, let us prove the achievability part of this capacity theorem. So, we are going to use the same ideas that we have used earlier to prove channels, noisy channel, coding theorem the difference here is, now in that case we had discrete random variables.

Here we have continuous random variables this difference here we also have this average power constraint which we did not have let me prove a channel noisy channel coding theorem. So, we are going to use a same ideas as the as we have used including channel coding theorem for discrete channels and what was those ideas we are going to use random randomly. We are going to generate this code word and we are going the receiver, we are going to use typical set coding we are going to use this ideas of random codes and joint typical decoding as i said, we are going to take into account that now there is a constraint on powers is that is violated basically that is the error and the variable are now continuous not discreet. So, that is the difference from the proof that we had done earlier.

(Refer Slide Time: 12:03)



Achievability of Gaussian channel capacity

- *Generation of the codebook* We wish to generate a codebook in which all the codewords satisfy the power constraint.
- To ensure this, we generate the codewords with each element i.i.d. according to a normal distribution with variance $P - \epsilon$.
- Since for large n , $\frac{1}{n} \sum X_i^2 \rightarrow P - \epsilon$, the probability that a codeword does not satisfy the power constraint will be small.
- Let $X_i(w)$, $i = 1, 2, \dots, n$, $w = 1, 2, \dots, 2^{nR}$ be i.i.d. $\sim \mathcal{N}(0, P - \epsilon)$ forming codewords $X^n(1), X^n(2), \dots, X^n(2^{nR}) \in \mathcal{R}^n$.

So, first step is generation of the codebook. So, we are going to generate a codebook in which all code words are going to satisfy the power constraint. So, to ensure this we are going to generate code words with each element identically or independently generated

according to Gaussian distribution, with variance given by $p - \epsilon$. Where $p - \epsilon$ is a small positive number now by choosing noise variance $p - \epsilon$ we are going to ensure that the average power constraint is not violated. So, we can see from last large numbers.

If we generate each code word normal distributed variance given by $p - \epsilon$ by law of large numbers the average value will basically tend to $p - \epsilon$ and hence the probability that a code word does not satisfy this average power constraint will be small. Let x_i where w are these 2^{nR} code words be i.i.d distributed and as we said we are generating this code words with variance given by $p - \epsilon$. So, let these code words be $x_1, x_2, \dots, x_{2^{nR}}$. So, these are n length code words which are randomly generated with Gaussian distributed with variance $p - \epsilon$.

(Refer Slide Time: 14:05)

Achievability of Gaussian channel capacity

- *Encoding.* After the generation of the codebook, the codebook is revealed to both the sender and the receiver.
- To send the message index w , the transmitter sends the w th codeword $X^n(w)$ in the codebook.
- *Decoding.* The receiver looks down the list of codewords $\{X^n(w)\}$ and searches for one that is jointly typical with the received vector.
- If there is one and only one such codeword $X^n(w)$, the receiver declares $\hat{W} = w$ to be the transmitted codeword. Otherwise, the receiver declares an error.
- The receiver also declares an error if the chosen codeword does not satisfy the power constraint.

So, after the generation of the codebook the codebook is revealed to both the sender and receiver. Now if we want to send message w , what the sender does it sends the w th code word $x_n w$ in the codebook. Now at the decoder we are going to do typical set decoding. So, how does typical set decoding works. So, we are going to see which code word is jointly typical with a received codebook and, if there is only 1 code, code word which is

jointly typical with this and there is know the code word which is jointly typical with a received codebook then, the decoder will correctly decode otherwise there is a decoding error also if the power constraint is valid then also there is an error.

So, the receiver is going to look down the least of code words which are just denoted by this x and w and it searches for this one code word which is jointly typical with this received vector y^n , it tries to find that one code word which is jointly typical with this received vector. Now if there is one and only one note this if there is one and only one such code word then, the receiver and let see that code word is x and w then, the receiver is going to declare that the transmitted code word was w otherwise the receiver will make an error since. There either more than, 1 code word which is jointly typical with the receive sequence or there is no code word which is jointly typical with the receive sequence then, the decoder will make an error also when power constraint is valid it then also receiver will make error. So, as I said the receiver also declares an error if the chosen code word does not satisfy this average power constraint.

(Refer Slide Time: 16:36)

Achievability of Gaussian channel capacity

- *Probability of error.* Without loss of generality, assume that codeword 1 was sent. Thus, $Y^n = X^n(1) + Z^n$. We define the following events:

$$E_0 = \left\{ \frac{1}{n} \sum_{j=1}^n X_j^2(1) > P \right\}$$

and

$$E_i = \{(X^n(i), Y^n) \text{ is in } A_e^{(n)}\}$$

- Then an error occurs if E_0 occurs (the power constraint is violated) or E_1^c occurs (the transmitted codeword and the received sequence are not jointly typical) or $E_2 \cup E_3 \cup \dots \cup E_{2^{nR}}$ occurs (some wrong codeword is jointly typical with the received sequence).

Now, let us compute what is the probability of error. So, when the decoder is doing joint typical set decoding. Let us compute what is the probability of error when, we are sending this with code, word over this Gaussian channel. So, without loss of generality

we will assume let say the code word one was sent. So, we are going to assume code word 1, was sent without loss of generality then the received code word y_n is going to be x_{n-1} plus z_n this is a noise this is the code word that was transmitted and this is what we have received.

Now, let us define the error events the first error event which i am defining by e_0 is the error event corresponding to the violation of average power constraint. So, e_0 is if the average power constraint is violated then, the event e_0 happens. So, this corresponds to violation of power constraint the average power constraint that we had let e_i be the event where, x_{n-i} and y_n belongs to this jointly typical set. So, e_i corresponds to the event when x_{n-i} and y_n are jointly typical. So, clearly an error happens if event e_0 occurs or if e_1 complement occurs why because, we have sent code word 1. So, there will be no error if y_n is jointly typical with x_{n-1} ; however, event e_1 complement will occur. If x_{n-1} and y_n are not jointly typical and that is this event e_1 complement.

So, if the transmitted code word and the receive sequence are not jointly typical that corresponds to the event e_1 complement then also error happens or if any of these events occur e_2, e_3, e_4, e_2 raise power n ; that means, if any other codeword other than the code word number one is jointly typical with the receive sequence y_n then there is an error and that that is denoted by this union of events e_2, e_3, e_4, e_2 raise power n . So, this corresponds to a wrong code word which was not sent to be jointly typical with the receive sequence.

(Refer Slide Time: 20:02)

Achievability of Gaussian channel capacity

- Let $\mathcal{E}$ denote the event $\hat{W} \neq W$ and let P denote the conditional probability given that $W = 1$. Hence,

$$Pr(\mathcal{E}|W = 1) = P(\mathcal{E}) = P(E_0 \cup E_1^c \cup E_2 \cup E_3 \cup \dots \cup E_{2^{nR}})$$

$$\leq \underbrace{P(E_0)}_{\leq \epsilon} + \underbrace{P(E_1^c)}_{\leq \epsilon} + \underbrace{\sum_{i=2}^{2^{nR}} P(E_i)}_{\leq \epsilon}$$
 by the union of events bound for probabilities.
- By the law of large numbers, $P(E_0) \rightarrow 0$ as $n \rightarrow \infty$.
- Now, by the joint AEP (which can be proved using the same argument as that used in the discrete case), $P(E_1^c) \rightarrow 0$, and hence

$$\underline{P(E_1^c)} \leq \epsilon \text{ for } n \text{ sufficiently large}$$

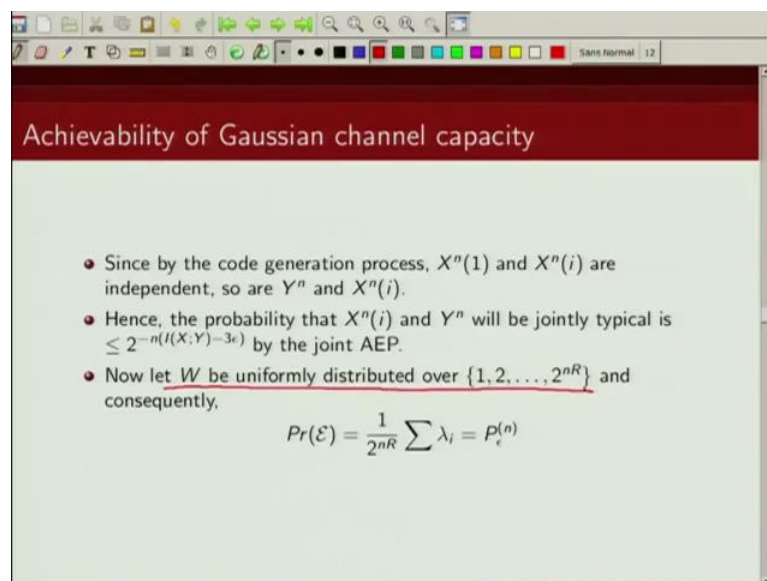
So, now that we have a numerated error events let us compute the probability of error given in the code word 1 was transmitted. So, let e denotes the event that what, which is an estimate of the code word that we sent is not same as w and let p be the conditional probability of error given w is 1. So, this is given by probability that event e_0 happens which corresponds to violation of power constraint union with event e_1 compliment which is the event that x and 1 and y_n are not jointly typical or union with e_2, e_3, e_4, e_2 raise power $n r$ which corresponds to the event that are wrong code word is jointly typical with the receive sequence y_n .

Now this can be upper bound bit using union bound using union bound, we can upper bound this probability of union of these events by. So, we can upper bound this using union bound this is probability of e_0 upper bounded by probability of e_0 plus probability of e_1 compliment plus probability of e_2 plus probability of e_3 up to plus probability of $e_{2^{nR}}$.

Fine next now, we are going to evaluate these probabilities now first let us, look at probability of e_0 now remember we were generating these code books with variance $p - \epsilon$ where ϵ is a positive number small number close to 0. So, if we are generating our code books like this by law of large number probability of

occurrence of this event e_0 is 0 as n is very large. So, this probability can be upper bounded by small quantity ϵ . Similarly, let us look at this probability of occurrence of the event that x_n is not jointly typical with y_n . Now we know from the property of joint ϵ p which we proved in the last class that probability of this event is also 0 as n goes to infinity. So, this can also be upper bounded by ϵ . So, we have upper bounded this by ϵ we are upper bounding this by ϵ .

(Refer Slide Time: 23:27)



Achievability of Gaussian channel capacity

- Since by the code generation process, $X^n(1)$ and $X^n(i)$ are independent, so are Y^n and $X^n(i)$.
- Hence, the probability that $X^n(i)$ and Y^n will be jointly typical is $\leq 2^{-n(I(X:Y)-3\epsilon)}$ by the joint AEP.
- Now let W be uniformly distributed over $\{1, 2, \dots, 2^{nR}\}$ and consequently,

$$Pr(\mathcal{E}) = \frac{1}{2^{nR}} \sum \lambda_i = P_e^{(n)}$$

Now, let us look at this error event now since the way we are generating our codebooks x_n and $x_n(i)$ where, I is the mother code, codeword they are independent. So, y_n and $x_n(i)$ are also going to be independent and we have shown that we have shown a result in the last class that what is the probability that if x_n , $x_n(i)$ and y_n which are chosen independently. What is the probability that they will be jointly typical? So, what is the probability that $x_n(i)$ and y_n are jointly typical this we have proved in the last class by the property of joint a e p this is upper bounded by 2 raise to power minus n times mutual information between x and y minus 3ϵ . So, if we let w to be uniformly distributed over all these 2 raise power nR code words then subsequently the probability of error is given by this average probability of error.

(Refer Slide Time: 24:43)

Achievability of Gaussian channel capacity

• Then

$$\begin{aligned}
 p_e^{(n)} &= Pr(\mathcal{E}) = Pr(\mathcal{E}|W=1) \\
 &\leq \underbrace{P(E_0)}_{\leq \epsilon} + \underbrace{P(E_1^c)}_{\leq \epsilon} + \sum_{i=2}^{2^{nR}} \underbrace{P(E_i)}_{\leq \epsilon} \\
 &\leq \epsilon + \epsilon + \sum_{i=2}^{2^{nR}} 2^{-n(I(X;Y)-3\epsilon)} \\
 &= 2\epsilon + \frac{(2^{nR} - 1)2^{-n(I(X;Y)-3\epsilon)}}{2^{-n(I(X;Y)-R-3\epsilon)}} \\
 &\leq 2\epsilon + \frac{2^{3nr} 2^{-n(I(X;Y)-R)}}{2^{-n(I(X;Y)-R-3\epsilon)}} \leq \epsilon \\
 &\leq 3\epsilon
 \end{aligned}$$

for n sufficiently large and $R < I(X;Y) - 3\epsilon$.

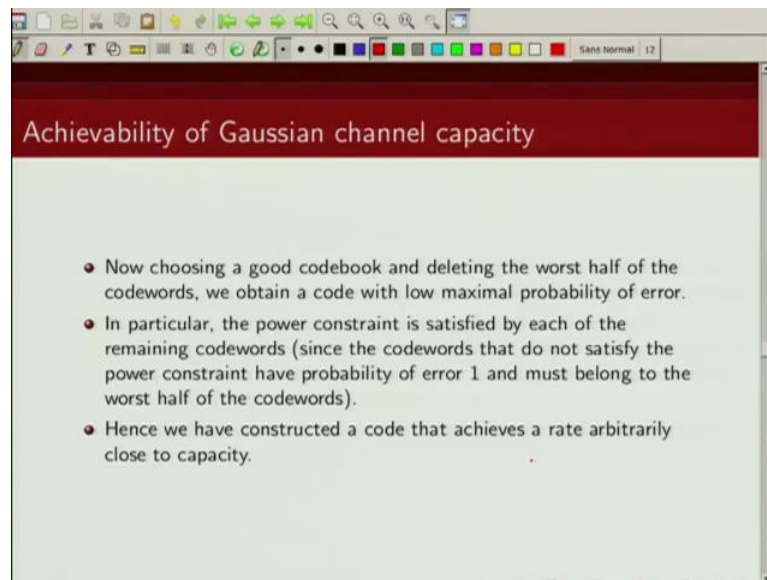
• This proves the existence of a good $(2^{nR}, n)$ code.

Now as we said average probability of error without loss of generality can be written as probability of error given code word 1, was send using a union bound this was upper bounded by sum of these error events we have already shown, this is upper bounded by epsilon this is upper bounded by epsilon and this is upper bounded by this quantity this follows from the properties of joint a e p which we have proved in the last class. So, then this term does not depend on i. So, I can sum over all i going from p to 2 raise power n r. So, i get this epsilon will be 2 epsilon and this summation over i from 2 to 2 raise power n r this will be this term and then, we have this one which is this now combining the terms containing n i get this this can be written as 2 raise power 3 n epsilon into 2 raise to power minus n mutually is a minus r and this quantity.

If we choose our r to be less than mutual information being x and y minus 3 epsilon see this term, we can write as 2 raise to the power minus n mutual information between x and y minus r minus 3 epsilon can write in this way. So, if we choose our r to be less than mutual information between x and y minus 3 epsilon then, this term will be positive and. So, for large n this will also go to 0. So, this whole term can be upper bounded by epsilon. So, in other words this probability of error can then be upper bounded by 3 epsilon provided my rate is less than, mutual information between x and y minus 3 epsilon. So, this shows that there exists good codes with rate given by this and in that

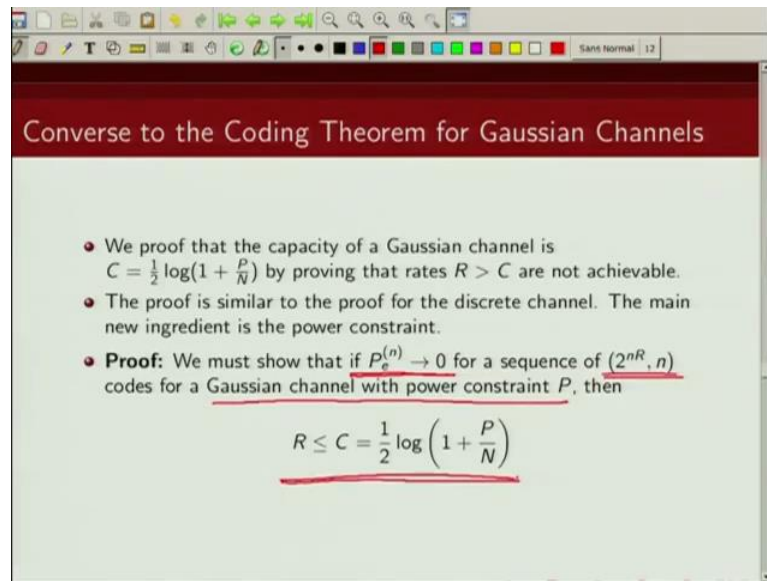
case the probability of error will go towards 0 as n goes to infinity.

(Refer Slide Time: 27:52)



Now, if we delete half of this words code words, we obtain a code with low maximum probability of error in particular the power constraint will be satisfied by this remaining code words because we are already throwing of these. So, called gauss code words and hence we are constructed a code that achieves a rate arbitrarily goes to capacity and gauss probability of error which is which goes probability of error which goes to 0 as n goes to infinity. So, this proves the achievability of Gaussian channel capacity.

(Refer Slide Time: 28:40)



Converse to the Coding Theorem for Gaussian Channels

- We prove that the capacity of a Gaussian channel is $C = \frac{1}{2} \log\left(1 + \frac{P}{N}\right)$ by proving that rates $R > C$ are not achievable.
- The proof is similar to the proof for the discrete channel. The main new ingredient is the power constraint.
- **Proof:** We must show that if $P_e^{(n)} \rightarrow 0$ for a sequence of $(2^{nR}, n)$ codes for a Gaussian channel with power constraint P , then

$$R \leq C = \frac{1}{2} \log\left(1 + \frac{P}{N}\right)$$

Next we are going to show the converse of the channel coding theorem for Gaussian channel. So, we are going to show that we prove that a channel capacity of Gaussian channel is given by this and we prove that rate greater than capacity is not achievable or we are going to show that; if probability of error goes to 0 then the rate is less than, channel capacity. Now the converse proof is very similar to the proof of discrete channel there is one additional constraint that we have here now, which is the average power constraint. So, we are going to show you now that this probability of error goes to 0 for a sequence of code given by these parameters for a Gaussian channel, with power constraint p as long as rate is less than capacity which is given by this expression.

(Refer Slide Time: 30:03)

Converse to the Coding Theorem for Gaussian Channels

- Consider any $(2^{nR}, n)$ code that satisfies the power constraint, that is,

$$\frac{1}{n} \sum_{i=1}^n x_i^2(w) \leq P$$
 for $w = 1, 2, \dots, 2^{nR}$.
- Proceeding as in the converse for the discrete case, let W be distributed uniformly over $\{1, 2, \dots, 2^{nR}\}$.
- The uniform distribution over the index set $W \in \{1, 2, \dots, 2^{nR}\}$ induces a distribution on the input codewords, which in turn induces a distribution over the input alphabet.
- This specifies a joint distribution on $\check{W} \rightarrow X^n(W) \rightarrow Y^n \rightarrow \check{W}$.

$\underline{I(W; \check{W})} \leq \underline{I(X^n(W); Y^n)}$

So, we are considering a 2^{nR} code which satisfies average power constraint which is given by this expression and this is satisfied by all code words w going from one to 2^{nR} . Let w be uniformly distributed now this uniform distribution over this index set induces a distribution of input code words which in turn induces distribution on the input alphabet and we can show, that $w \rightarrow x^n(w) \rightarrow y^n \rightarrow \hat{w}$ they form a Markov chain. So, w is my code word index x^n is my encoder encoded codebook y^n is the received codebook and $\hat{w}$ is an estimate of w based on what I receive which is y^n .

(Refer Slide Time: 31:17)

Converse to the Coding Theorem for Gaussian Channels

• To relate probability of error and mutual information, we can apply Fano's inequality to obtain

$$H(W|\hat{W}) \leq 1 + nR P_e^{(n)} = n\epsilon_n$$

$H(W|\hat{W}) \leq H(P_e) + P_e \log_2 \frac{2^n - 1}{P_e}$
 $\leq 1 + P_e \log_2 2^n$
 $\leq 1 + nR P_e$

where $\epsilon_n \rightarrow 0$ as $P_e^{(n)} \rightarrow 0$. Hence,

$$nR = H(W) = I(W; \hat{W}) + H(W|\hat{W})$$

$$\leq I(W; \hat{W}) + n\epsilon_n$$

$$\leq I(X^n; Y^n) + n\epsilon_n$$

$$= h(Y^n) - h(Y^n|X^n) + n\epsilon_n$$

$$= h(Y^n) - h(Z^n) + n\epsilon_n$$

$I(W; \hat{W}) \leq I(X^n; Y^n)$
 $Y^n = X^n + Z^n$
 $h(X^n + Z^n|X^n) = h(Z^n|X^n)$
 $= h(Z^n)$

Now, we are going to take help of Fano's lemma to relate probability of error and to mutual information. Now we in the class has used this version of Fano's lemma which was given by now we did each of p_e plus $p_e \log$ of 1 minus 1 we use this portion and we said this can be weak end because binary entropy function is less than equal to 1 . So, this can be weak end also 1 and this can be we further we can $p_e \log$ of all the possibilities 1 . So, number of code words here is 2 raise to power $n r$. So, if we take log of that this will become less than equal to one plus n times $r n$ probability of error.

So, this is the version that we are using uncertainty in w given w hat is upper bounded by 1 plus $n r$ times probability of error and as we said here we are going to show that, if probability of error goes to 0 for large n then the rate has to be less than c . So, this is you know. So, uncertainty in w given w hat is less than equal to n times. So, small value of epsilon where epsilon goes to 0 and probability of error goes to 0 . Now w is uniformly distributed over this index said 1 2 2 raise power $n r$. So, uncertainty in w is given by log of 2 raise power $n r$ which is nothing, but $n r$ now from the definition of mutual information we can write uncertainty in w in terms of mutual information between w and w hat plus conditional entropy of w given w hat. Now if probability of error goes to 0 as, n goes to infinity we have shown that this term is upper bounded by n times.

So, I can write then uncertainty in w is upper bounded by mutual information between w and $\hat{w}$ plus n times ϵ_n . Now we just saw in the last slide that w , x , n , w , y , n and $\hat{w}$ form a Markov chain. So, then from data processing lemma we know that mutual information between w and $\hat{w}$ is going to be less than mutual information between x and y . So, this follows from the data processing lemma. So, then we can write this that the mutual information between w and $\hat{w}$ is less than equal to mutual information between x and y . So, that is what we are writing here and this comes as n times ϵ_n now following the definition of mutual information, we can write the mutual information between x and y as differential entropy of y minus differential entropy of y given x and this is n times ϵ_n .

Now what is y ? y is the output of the Gaussian channel which is x plus z where z is the Gaussian noise this 0 mean n variance n . Now uncertainty in y differential entropy of y given x can be written as differential entropy of x plus z given x and this is nothing, but uncertainty in z given x now note that the noise z and the signal x are independent and hence this can be written as uncertainty in differential entropy of z of n .

(Refer Slide Time: 36:55)

Converse to the Coding Theorem for Gaussian Channels

$$\begin{aligned}
 nR &\leq h(Y^n) - h(Z^n) + n\epsilon_n \\
 &\leq \sum_{i=1}^n h(Y_i) - h(Z^n) + n\epsilon_n \\
 &= \sum_{i=1}^n h(Y_i) - \sum_{i=1}^n h(Z_i) + n\epsilon_n \\
 &= \sum_{i=1}^n I(X_i; Y_i) + n\epsilon_n
 \end{aligned}$$

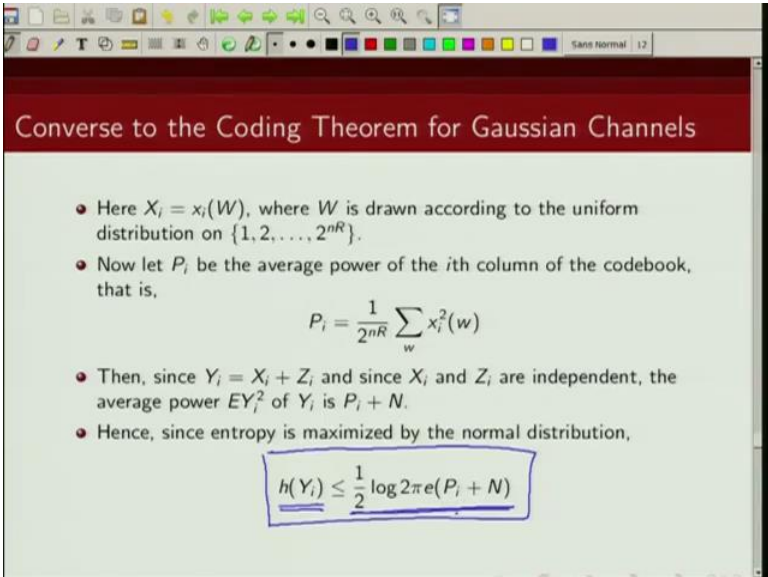
$h(Y^n) \leq \sum_{i=1}^n h(Y_i)$
 $h(Z^n) \leq \sum_{i=1}^n h(Z_i)$

$h(Z_i)$
 $= h(z_i/x_i)$
 $= h(x_i + z_i/x_i)$
 $= h(z_i/x_i)$

So, I can write differential entropy of y given x as differential entropy of z next. So,

what we have shown. So, far is $n r$ is actually less than equal to h of y_n minus differential entropy of z_n plus n times epsilon. Now this can be upper bounded by this again this is straight forward to prove we have prove this using chain rule we can write h of y_n in terms of $y_1 y_2 y_3 y_n$ conditioned on these y_i and we know that conditioning cannot increase entropy. So, from that we get this result similarly we can write h of z_n as less than equal to summation over i from 1 to n h of z_i . So, this is what we are writing here plus n times epsilon n this is this term here now, we have just shown h of z_i is same as h of z_i given x_i and this is same as h of x_i plus z_i given x_i which is same as h of y_i given x_i . So, this is differential entropy of y_i this is differential entropy of y_i given x_i . So, then this can be written as mutual information between x_i and y_i .

(Refer Slide Time: 39:17)



Converse to the Coding Theorem for Gaussian Channels

- Here $X_i = x_i(W)$, where W is drawn according to the uniform distribution on $\{1, 2, \dots, 2^{nR}\}$.
- Now let P_i be the average power of the i th column of the codebook, that is,
$$P_i = \frac{1}{2^{nR}} \sum_w x_i^2(w)$$
- Then, since $Y_i = X_i + Z_i$ and since X_i and Z_i are independent, the average power EY_i^2 of Y_i is $P_i + N$.
- Hence, since entropy is maximized by the normal distribution,
$$h(Y_i) \leq \frac{1}{2} \log 2\pi e(P_i + N)$$

So, from this result and this result we can write $n r$ is upper bounded by summation of mutual information between x_i and y_i where, summation is over i going from one to n plus n times epsilon now let x_i be w_i where w is drawn uniformly over this index set and let p_i be the average power of the i th column of this codebook we know the x_i signal and z_i the noise they are independent.

So, the average power of the received sequence y_i can be written as p_i plus n where p_i is the power of x_i n is the noise variance now again this result will have proved in the

previous class that if, y has y_i have the same second order movement then Gaussian random variable will have the maximum differential entropy. So, h of y_i is upper bounded by the differential entropy of an equivalent of differential entropy of Gaussian random variable with same variance. So, basically h of y_i is an upper bounded by half log of $2\pi e$ times p_i plus n . So, the differential entropy of y_i is upper bounded by the differential entropy of Gaussian random variable which has 0 mean and variance given by p_i plus n actually mean can be anything because translation does not change the differential entropy.

(Refer Slide Time: 41:13)

Converse to the Coding Theorem for Gaussian Channels

- Continuing with the inequalities of the converse, we obtain

$$\begin{aligned}
 nR &\leq \sum (h(Y_i) - h(Z_i)) + n\epsilon_n \\
 &\leq \sum \left(\frac{1}{2} \log(2\pi e(P_i + N)) - \frac{1}{2} \log 2\pi eN \right) + n\epsilon_n \\
 &= \sum \frac{1}{2} \log \left(1 + \frac{P_i}{N} \right) + n\epsilon_n
 \end{aligned}$$

$R \leq \frac{1}{n} \sum \frac{1}{2} \log \left(1 + \frac{P_i}{N} \right) + \epsilon_n$
- Since each of the codewords satisfies the power constraint, so does their average, and hence

$$\frac{1}{n} \sum P_i \leq P,$$

Now continuing with the inequalities that we had we just had the expression that n of r is upper bounded by summation over all i . So, $\sum_{i=1}^n h(y_i) - h(z_i)$, now we just have now showed you that differential entropy of y_i is upper bounded by this quantity and since z_i is Gaussian distributed with variance n is differential entropy is given by this expression. Now, combining these 2 log times this term and this term we get this expression the nR is upper bounded by summation from i goes to one to n half log of one plus p_i divided by n plus n times epsilon. Now the way we generated these codebooks code words they all satisfied average power constraint. So, they all satisfy power constraint because we were we are generating these code word with variance given by p minus epsilon. So, this code words satisfying power constraint. So, that average is also

going to satisfy the power constraint.

(Refer Slide Time: 42:53)

Converse to the Coding Theorem for Gaussian Channels

- Since $f(x) = \frac{1}{2} \log(1+x)$ is a concave function of x , we can apply Jensen's inequality to obtain

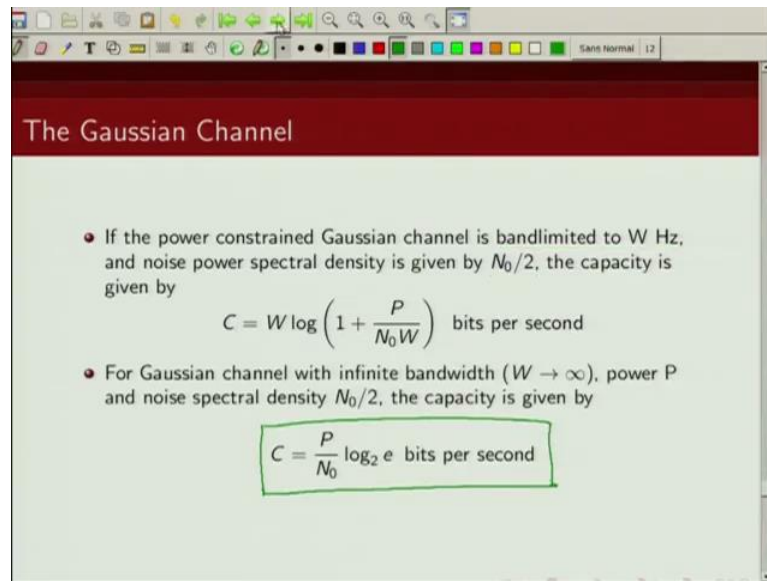
$$\frac{1}{n} \sum_{i=1}^n \frac{1}{2} \log \left(1 + \frac{P_i}{N} \right) \leq \frac{1}{2} \log \left(1 + \frac{1}{n} \sum_{i=1}^n \frac{P_i}{N} \right) \leq \frac{1}{2} \log \left(1 + \frac{P}{N} \right)$$

- Thus $R \leq \frac{1}{2} \log \left(1 + \frac{P}{N} \right) + \epsilon_n$, $\epsilon_n \rightarrow 0$, and we have the required converse.

Now we know that log is a concave function and if a function is concave, we know from Jensen's inequality that expected value of the function is less than or equal to function of expected value. So, this is what we know from Jensen's inequality. If f of x is a concave function in this case log is a concave function of x . So, then here we have this expression that nR is less than or equal to summation over all i of $\frac{1}{2} \log(1 + \frac{P_i}{N})$ plus ϵ_n . If I divide this by n , what I would get is this will go away, this will go away and I will have here $R \leq \frac{1}{2} \log(1 + \frac{P}{N}) + \epsilon_n$.

So, have I have this expression R is less than or equal to $\frac{1}{2} \log(1 + \frac{P}{N}) + \epsilon_n$. Now I know that $\log(1+x)$ is a concave function of x . Then expected value of log function will be upper bounded by log of expected value of this x . So, that is what we are doing here. So, this is expected value of the log function this is from Jensen's inequality upper bounded by function is log of the expected value alright now this one by n summation P_i is nothing, but average power constraint. So, what we have shown here now is that R is less than or equal to $\frac{1}{2} \log(1 + \frac{P}{N}) + \epsilon_n$. So, what we have shown is if probability of error goes to 0 then, R is also less than channel capacity. So, R is less than channel capacity if probability of error goes to 0 as n goes to infinity.

(Refer Slide Time: 45:48)

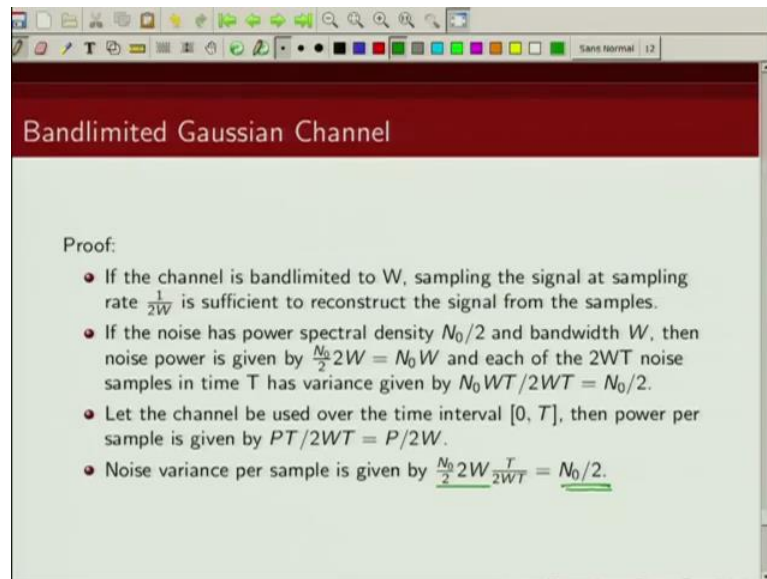


The Gaussian Channel

- If the power constrained Gaussian channel is bandlimited to W Hz, and noise power spectral density is given by $N_0/2$, the capacity is given by
$$C = W \log \left(1 + \frac{P}{N_0 W} \right) \text{ bits per second}$$
- For Gaussian channel with infinite bandwidth ($W \rightarrow \infty$), power P and noise spectral density $N_0/2$, the capacity is given by
$$C = \frac{P}{N_0} \log_2 e \text{ bits per second}$$

So, this proves the converse of the channel coding theorem for Gaussian channel next, we are going to talk about what is the capacity of a Gaussian channel. If the Gaussian channel is band limited to w hertz and the noise power spectral density is given by n naught by two. So, we are going to show you that if, the Gaussian channel is band limited to w hertz and its noise power spectral density is given by n naught by 2 then, the capacity of this band limited Gaussian channel is given, by $w \log$ of one plus p divided by n naught w this is the capacity of the band limited Gaussian channel and if we let w go to infinity then the capacity is given by this expression.

(Refer Slide Time: 46:50)



Bandlimited Gaussian Channel

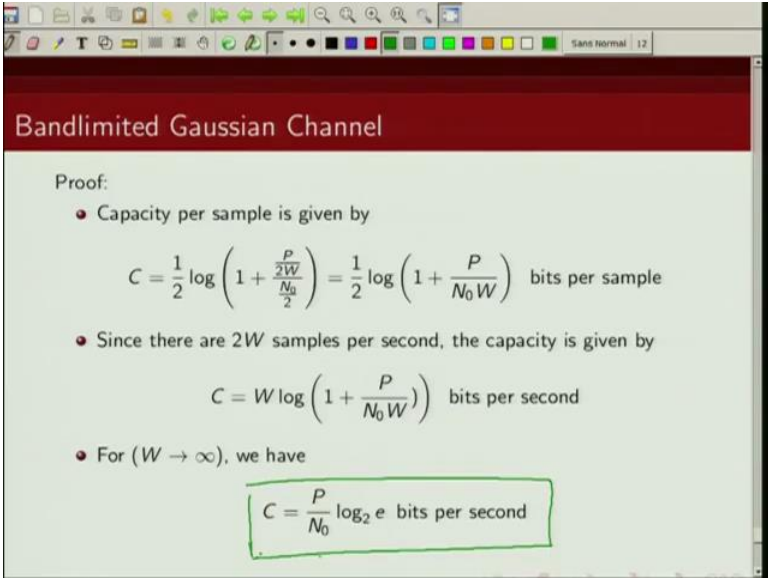
Proof:

- If the channel is bandlimited to W , sampling the signal at sampling rate $\frac{1}{2W}$ is sufficient to reconstruct the signal from the samples.
- If the noise has power spectral density $N_0/2$ and bandwidth W , then noise power is given by $\frac{N_0}{2} 2W = N_0 W$ and each of the $2WT$ noise samples in time T has variance given by $N_0 W T / 2WT = N_0/2$.
- Let the channel be used over the time interval $[0, T]$, then power per sample is given by $PT / 2WT = P/2W$.
- Noise variance per sample is given by $\frac{N_0}{2} 2W \frac{T}{2WT} = N_0/2$.

So, let us prove this. So, if the channel is band limited to w then we can sample this signal at rate 1 by $2w$ and that would be sufficient to reconstruct the signal from the samples this follows from the criteria. Now if the noise has power spectral density given by $N_0/2$ and the band width is w then the noise power is given by $N_0/2$ times w . So, you have basically noise power spectral density given by $N_0/2$ and you have band width of w to w . So, this noise power is given by $N_0/2$ into $2w$ times w which is $N_0 w$ and over of period from 0 to t have, many noise samples because we are sampling at rate one by $2w$.

So, each of these $2wt$ noise samples in time t will have variance given by $N_0 w$ times t divided by $2wt$ which is nothing, but $N_0/2$. Now let us use the channels over time 0 to t , power per sample will be p times t divide by total samples which is $2wt$ and that comes out to be p divided by 2 times w . So, similarly we can compute noise variance per sample is $N_0/2$ into $2wt$ divided by $2wt$ which is $N_0/2$.

(Refer Slide Time: 48:45)



Bandlimited Gaussian Channel

Proof:

- Capacity per sample is given by
$$C = \frac{1}{2} \log \left(1 + \frac{P}{\frac{N_0 W}{2}} \right) = \frac{1}{2} \log \left(1 + \frac{P}{N_0 W} \right) \text{ bits per sample}$$
- Since there are $2W$ samples per second, the capacity is given by
$$C = W \log \left(1 + \frac{P}{N_0 W} \right) \text{ bits per second}$$
- For $(W \rightarrow \infty)$, we have
$$C = \frac{P}{N_0} \log_2 e \text{ bits per second}$$

So, capacity per sample is given by half log one plus power per sample this is noise per sample if you go back here, noise variance per sample is n_0 by 2 and power per sample is given by p by 2 w . So, then the capacity is given by half log of one plus p by 2 w divided by n_0 by 2 and this is nothing, but half log of one plus p divided by n_0 w now this is capacity per samples and how many sample we have per second we have $2 w$ samples per second.

So, then the capacity is given by $w \log$ of 1 plus p divided by $n_0 w$ bits per second this is signal power this is noise power and we can similarly compute, if you let w go to infinity and take the limit the capacity comes out to be this expression p times p divided by n_0 into log of e bits per second. So, this is the expressions for capacity for band limited Gaussian channel. So, with this we will conclude our discussion on Gaussian channel. In the next lecture we will talk about parallel Gaussian channel.

Thank you.