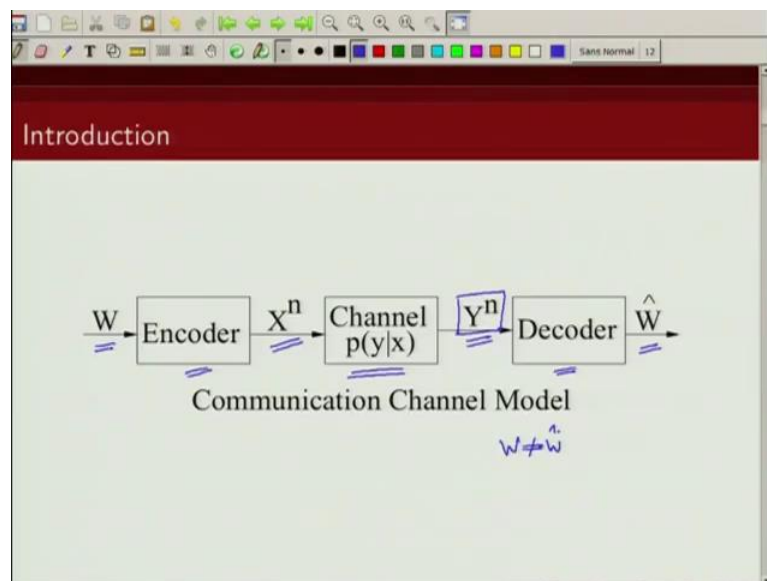


An Introduction to Information Theory
Prof. Adrish Banerjee
Department of Electronics and Communication Engineering
Indian Institute of Technology, Kanpur

Lecture - 10B
Noisy Channel Coding Theorem

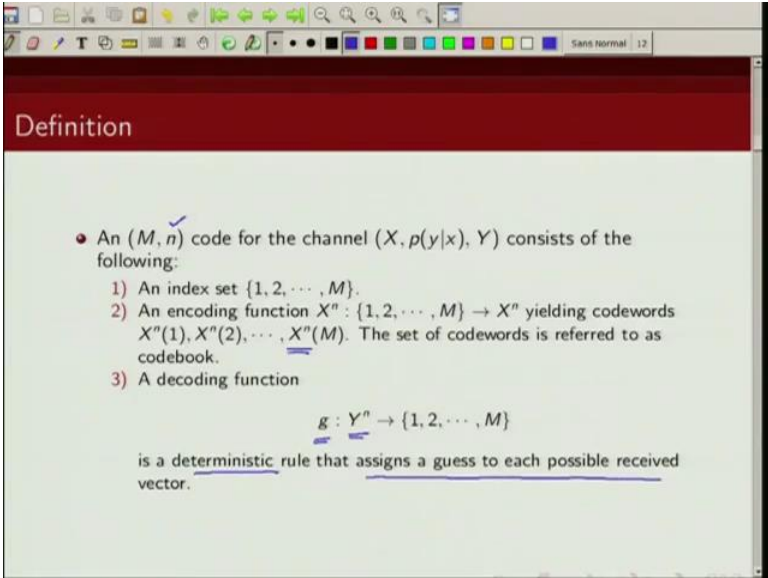
Welcome to the course on an Introduction to Information Theory. Now in this lecture we are going to talk about Shannon's noisy channel coding theorem. We are going to prove the achievability part of the result and we are also going to prove the converse of this theorem. In the last lecture we have already build up the background to prove our noisy channel coding theorem.

(Refer Slide Time: 00:46)



So, we start our lecture with some basic definition and model of the communication channel that we will consider. So, w is have a message that you want to send is encoded using this encoder and x of n is the code word, that is transmitted over a discrete memory less channel it is transition probability channel transit probabilities are given by probability of y given x , y^n are the received noisy code words and the job of the decoder is to estimate what was the message signal that was transmitted given received sequence y^n .

(Refer Slide Time: 01:49)



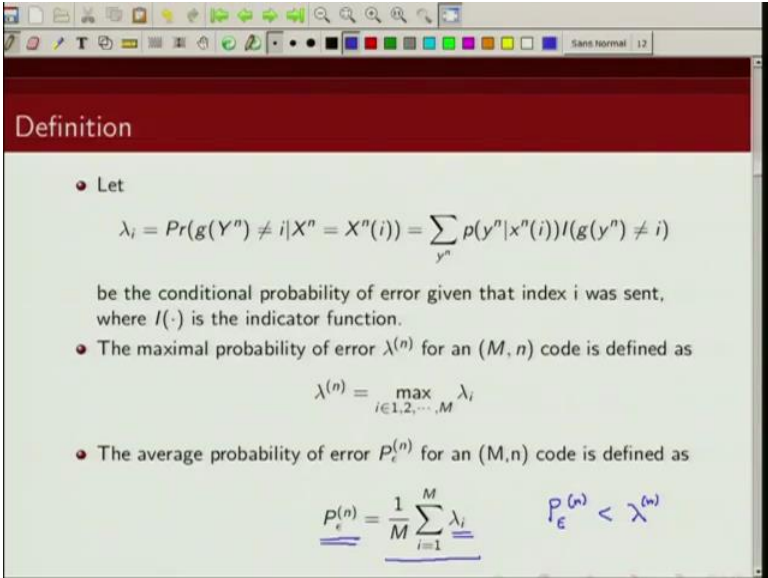
Definition

- An (M, n) code for the channel $(X, p(y|x), Y)$ consists of the following:
 - 1) An index set $\{1, 2, \dots, M\}$.
 - 2) An encoding function $X^n : \{1, 2, \dots, M\} \rightarrow X^n$ yielding codewords $X^n(1), X^n(2), \dots, X^n(M)$. The set of codewords is referred to as codebook.
 - 3) A decoding function
$$g : Y^n \rightarrow \{1, 2, \dots, M\}$$
is a deterministic rule that assigns a guess to each possible received vector.

Now, an error happens if w is not same as w hat otherwise there is no error. So, this is the communication channel model that we have. That we are considering was defined few terms before i go in to the noisy channel coding theorem. So, we define a code by these parameters m and n i will explain these terms in a little while and discrete memory less channel is described by this input x output y and channel transient probability which are given by probability of y given x . So, this index set will be you basically used to denote will be correspond to each of these messages, and what we are going to do encoder is a mapping of this in index set to a n bit sequence and this we are denoting by x of n .

So, we will have m code words this is x_{n1} x_{n2} denoted by x_{n3} , x_{nm} . So, these are our set of code words, and we will refer to a set of code words as code book. Now as you saw once you send these x ns at the channel output at the receiver input you are receiver you will received y of n . So, once you receive y of n then you essentially need to find out which of these are messages you have transmitted. So, decoding is essentially to be operation. So, it is a deterministic rule that assigns a guess to each possible received vector n is the length of the codeword and m are number of code words that you have.

(Refer Slide Time: 03:41)



Definition

- Let

$$\lambda_i = \Pr(g(Y^n) \neq i | X^n = X^n(i)) = \sum_{y^n} p(y^n | x^n(i)) I(g(y^n) \neq i)$$
 be the conditional probability of error given that index i was sent, where $I(\cdot)$ is the indicator function.
- The maximal probability of error $\lambda^{(n)}$ for an (M, n) code is defined as

$$\lambda^{(n)} = \max_{i \in 1, 2, \dots, M} \lambda_i$$
- The average probability of error $P_e^{(n)}$ for an (M, n) code is defined as

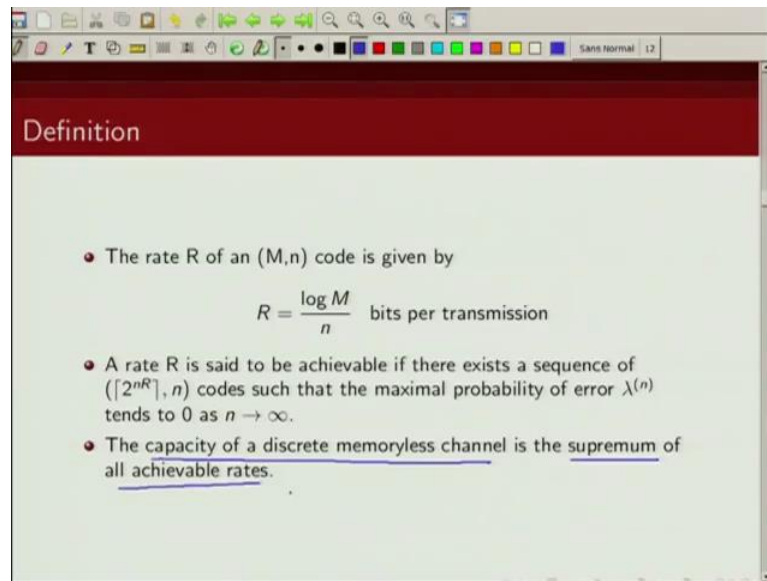
$$P_e^{(n)} = \frac{1}{M} \sum_{i=1}^M \lambda_i$$

Handwritten note: $P_e^{(n)} < \lambda^{(n)}$

Lambda i is nothing, but conditional probability of error given that index i was sent. So, i sent index i what is a probability of error. Now when will errors happen if the decoded. So, decoder function is denoted by g . So, if g of y of n is not same as i and when you transmit i th index. So, when x of n is x of n i when you transmit the i th index, that your decode are does not decoded as i th index then the error happens and that is given by this you know this is the indicator function indicating, whether error has happened or not and this probability y^n given x of n i . So, this is. So, lambda i is a conditional probability of error given that you transmitted the i th index. Now you define maximal probability of error as maximum of lambda i over all possible indexes.

So, index i can go for 1 to m . So, maxima of lambda i that is a maximal probability of error similarly we can define average probability of error. So, lambda i is the error, when i th index is sent. So, when you serve overall lambda i s and divide by m you get average probability of error is not very difficult to show that this average probability of error will be less than this maximal probability of error. Now we are going to show you later in this lecture that as long as a transmission rate is below channel capacity our average probability of error is about, is very small. It goes is to 0 as block length codeword length increases and we will show if the average probability of error goes to 0 then the maximal probability of error also it goes to 0 as then becomes large.

(Refer Slide Time: 06:09)



Definition

- The rate R of an (M,n) code is given by
$$R = \frac{\log M}{n} \text{ bits per transmission}$$
- A rate R is said to be achievable if there exists a sequence of $([2^{nR}], n)$ codes such that the maximal probability of error $\lambda^{(n)}$ tends to 0 as $n \rightarrow \infty$.
- The capacity of a discrete memoryless channel is the supremum of all achievable rates.

So, as I said that total m indexes that we are sending. So, that number of information bit's or $\log$ of m and the codeword length is n . So, the rate of the code is denoted by $\log$ of m divided by n .

Now, we see a rate is achievable if there exists a sequence of codes that parameter m given by $\log$ of 2 raised to power nR of length n such that the maximal probability of error goes to 0 as n goes to infinity. So, we say a rate is achievable if for that particular code with parameter 2 raised to nR the maximal probability of error goes to 0 as the codeword length increases and we defined the capacity of a discrete memory less channel as the supremum of all achievable rates. So, so supremum is the smallest it is the upper bound of a set.

(Refer Slide Time: 07:36)

Noisy channel coding theorem

- All rates below capacity C are achievable. Specifically, for every rate $R < C$, there exists a sequence of $(2^{nR}, n)$ codes with maximum probability of error $\lambda^{(n)} \rightarrow 0$

Proof:

- We fix $p(x)$ and generate independently 2^{nR} codewords according to the distribution $p(x^n) = \prod_{i=1}^n p(x_i)$.
- We exhibit the 2^{nR} codewords as rows of the matrix

$$C = \begin{bmatrix} x_1(1) & x_2(1) & \cdots & x_n(1) \\ \vdots & \vdots & \ddots & \vdots \\ x_1(2^{nR}) & x_2(2^{nR}) & \cdots & x_n(2^{nR}) \end{bmatrix}$$

- Probability that we generate a particular code C is

$$Pr(C) = \prod_{w=1}^{2^{nR}} \prod_{i=1}^n p(x_i(w))$$

So, capacity is basically supremum of all achievable rates, now for a discrete memory less channel Shannon's noisy channel coding theorem says that all rates below channel capacity are achievable. That means, if we transmit at rate below capacity then our maximal probability of error will go to 0 as the codeword length increases to infinity. So, so specifically for every rate which is less than capacity there exists a code of these parameters. So, if that maximal error probability goes to 0.

Now, we are going to prove that cumulative part of this theorem first. First we are going to show that average probability of error goes to 0 as n become large and, then we will show that if the average probability. Whenever goes to 0 then the maximal probability of error is also bounded and it is it goes to 0. So, first we do is we fix our input distribution and generate 2^{nR} code words according to this distribution i had a distribution and we considered a rows matrix, where each row is basically randomly generated code words. So, this x_1 and x_2 and x_n is codeword similarly these. So, we generate 2^{nR} code words and they are start up in a matrix like this. This is randomly generated.

So, probability that we generate a particular code C is given by this probability. So, probability of $x_i w$, where is w can go for 1 to 2^{nR} and i goes from 1 to

n.

(Refer Slide Time: 09:40)

Noisy channel coding theorem

- Randomly generated code is then revealed to both sender and receiver. Channel transition matrix $p(y/x)$ is known to both sender and receiver.
- A message W is chosen according to a uniform distribution
$$Pr(W = w) = 2^{-nR}, \quad w = 1, 2, \dots, 2^{nR}$$
- The w th codeword $X^n(w)$ corresponding to the w th row of C is sent over the channel.
- The receiver receives a sequence Y^n according to the distribution
$$P(y^n | x^n(w)) = \prod_{i=1}^n p(y_i | x_i(w))$$

Now, once we randomly generate a code we reveal this code to both the sender and the receiver now we assume the channel transition matrix probabilities are both known to the sender and receiver. And we chose a message w according to a uniform distribution. So, probability that we chose a particular index is given by $1/2^{nR}$. So, that is this probability and the w th codeword corresponds to the w th row of this code matrix that we talked about. So, when we select this index w we are actually sending this w th row of this codeword matrix, that we talked about that we are randomly generating now what the receiver receives is y of n and probability y of n given a particular x of n of w was sent this probability for a discrete memory less channel with the feedback is given by this expression.

(Refer Slide Time: 11:02)

Noisy channel coding theorem

- We consider jointly typical decoding.
- The receiver declares that the index $\tilde{W}$ was sent if the following conditions are satisfied: $(X^n(\tilde{W}), Y^n)$ is jointly typical and there is no other index k , such that $(X^n(k), Y^n) \in A_{\epsilon}^{(n)}$.
- If no such $\tilde{W}$ exists or if there is more than one such, an error is declared. If $\tilde{W} \neq W$, there is a decoding error.

Now, once we receive y^n at the receiver we are going to do, what we call jointly typical decoding. So, we are going to do a typical set decoding now this is a suboptimal decoding technique, but asymptotically this is basically optimal. So, how does a typical set decoding works it works as follows. So, the receiver declares that an index w tilde was sent if the following conditions are met and what are those following conditions if x^n w tilde and y^n which is a received sequence. If this is jointly typical and there is no other index k such that x^n k and receives sequence y of n are jointly typical. So, we you say that we successfully decode and index if x^n w and that index and y of n they are jointly to become and there is no other index for which x^n k and y^n are jointly typical and that is basically your typical set decoding.

So, when will an error happen an error will happen when, you send w index and you decide in favor of some other index which is not w . So, an error can happen in a typical set decoding if there is no w tilde; that means, there does not exist any index for which this pair of that x^n w tilde y it is jointly typical that is 1 condition or the second condition under which an error can happen is as follows. If there is more than what 1 such index for which x^n k and y^n is jointly typical, so, there are 2 possible cases of error 1 if there does not exist any index w tilde such that x^n w tilde and y^n are jointly typical that is 1 condition second condition is there exist more than 1 index for which x^n k and

y of n are jointly typical then also you will make an error decoding error.

So, let us try to calculate what is the average probability of error and we are doing calculating this average probability of error average over all possible code words and averaged over all possible code books. So, this a probability of error for particular code book and is a probability of this codeword C we average it over all possible code books.

(Refer Slide Time: 14:02)

Noisy channel coding theorem

- We will calculate the average probability of error, averaged over all codewords in the codebook and average over all codebooks.

$$\begin{aligned}
 Pr(E) &= \sum_C P(C) P_e^{(n)}(C) \\
 &= \sum_C P(C) \frac{1}{2^{nR}} \sum_{w=1}^{2^{nR}} \lambda_w(C) \\
 &= \frac{1}{2^{nR}} \sum_{w=1}^{2^{nR}} \sum_C P(C) \lambda_w(C)
 \end{aligned}$$

- Due to symmetry of the code construction, the average probability of error averaged over all codes does not depend on particular index that was sent.

So, this probability of error from definition this is given by this as you recall 1 by summation w 1 to n lambda w C right. So, this can be written like this now please note that this term does not depend on index w why, because there we are constructing our code due to symmetry of the code construction the average probability of error average over all possible codes is not a function of a particular index. Now what is that mean it means that then, we can calculate our probability of error given a particular index has been send since it does not depend on what index probability of error does not depend on what index have been sent.

So, we can calculate a probability of error assuming a particular index i was sent and therefore, we are going to do.

(Refer Slide Time: 15:32)

Noisy channel coding theorem

- We assume without loss of generality that the message $W = 1$ was sent. Average probability of error is given by

$$\begin{aligned} Pr(E) &= \frac{1}{2^{nR}} \sum_{w=1}^{2^{nR}} \sum_C P(C) \lambda_w(C) \\ &= \sum_C P(C) \lambda_1(C) \\ &= Pr(E|W = 1) \end{aligned}$$
- Let E_i be the event that the i th codeword and Y^n are jointly typical.

$$E_i = \left\{ (X^n(i), Y^n) \text{ is in } A_\epsilon^{(n)} \right\}, \quad i \in \{1, 2, \dots, 2^{nR}\}$$

So, without loss of generality we will assume that, let us say the first index was sent and we are going to calculate probability of error given the first index was sent. So, the probability of error can be given by probability of error given the first index was sent and that is probability of error given index1 was sent. So, let us compute this probability now let us denote e_i to be the event that i th codeword and received sequence y of n are jointly typical. So, e_i is the event that $x^n(i)$ and y^n belongs to jointly typical set now as I said probability of error does not depend on what index has been transmitted. So, we do analysis assuming; let us say index1 was sent. So, what is the probability of error, when index 1 was sent?

(Refer Slide Time: 16:29)

Noisy channel coding theorem

- We have $Pr(E) = Pr(E|W = 1)$ given by

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

$$\leq P(E_1^c | W = 1) + \sum_{i=2}^{2^{nR}} P(E_i | W = 1)$$
- By joint AEP, we have $P(E_1^c | W = 1) \leq \epsilon$ for large n .
- $X^n(1)$ and $X^n(i)$ are independent, so are Y^n and $X^n(i)$, hence the probability that Y^n and $X^n(i)$ are jointly typical is given by

$$\leq 2^{-n(I(X;Y) - 3\epsilon)}$$

So, this is given by. So, as I said this is given by probability that what is E_1^c ? E_1^c corresponds now if index 1 was sent the decoder will not make an error if x_{n1} and y of n are jointly typical and that corresponds to event e_1 . So, if e_1 complement happens then that is an error because if index 1 has been sent the decoder will not make a mistake if event e_1 occurs. So, an error will happen if the complement of the e_1 event even has happened or error can happen if any of the other index, which was not sent is jointly typical with the received sequence y of n . So, for example, if y_n is with x_{n2} or x_{n3} then there is an error because index 1 was sent. So, you will. So, I can write upper bound probability by i can write it as probability of e_1 complement union of e_2 event union of e_3 event union of up to $e_{2^{nR}}$ given, that first index was sent and usually union bound. I can upper bound it then as probability of e_1^c complement given index 1 was transmitted plus probability of e_i given w index 1 was transmitted and i go from 2 to 2^{nR} .

So, this corresponds to the event that y_n is not jointly typical with it. If you go back and see this corresponds to event e_1 corresponds to the e event that x and 1 and y_n are jointly typical. So, this corresponds should even event that x_{n1} and y_n are not jointly typical and this corresponds to the event that any other index. Let us say x_{n2} x_{n3} x_{n4} they are jointly typical with y of n now. So, then the error probability is upper bounded

by this probability and this probability, now from the property of joint a e p we know that if n is very large then x of $n-1$ and y of n are likely to be jointly typical. We have shown property of jointly typical sequence that probability that x of $n-1$ and y of n belongs to a typical jointly typical set that probability goes to 1 as n goes to infinity. So, then this probability that event e_1 complement happens given that w is 1 this probability is going to be very, very small less than epsilon.

Now, let us look at this now we have also shown that if, now if x_{n-1} and x_n where i is different from 1 they are independent. So, y_n and x_n are also independent with probability that y_n and x_n are jointly typical this probability is upper bounded by $2^{-n(I(X;Y) - 3\epsilon)}$ this also follows from the property of joint typical sequence this we have proved in the previous lecture. So, when x_{n-1} and y_n are independent probability that they are jointly typical this upper bounded by this.

(Refer Slide Time: 21:37)

Noisy channel coding theorem

Hence, we have

$$\begin{aligned}
 \underline{P(E)} = \underline{P(E|W=1)} &\leq \underline{P(E_1^c|W=1)} + \sum_{i=2}^{2^{nR}} \underline{P(E_i|W=1)} \\
 &\leq \epsilon + \sum_{i=2}^{2^{nR}} 2^{-n(I(X;Y) - 3\epsilon)} \\
 &= \epsilon + (2^{nR} - 1) 2^{-n(I(X;Y) - 3\epsilon)} \\
 &= \epsilon + 2^{3n\epsilon} 2^{-n(I(X;Y) - R)} \\
 &\leq 2\epsilon \quad \text{if } n \text{ is sufficiently large and } \boxed{R < I(X;Y) - 3\epsilon}
 \end{aligned}$$

Handwritten notes on the slide include a blue box around the final inequality and a blue arrow pointing to the condition $R < I(X;Y) - 3\epsilon$.

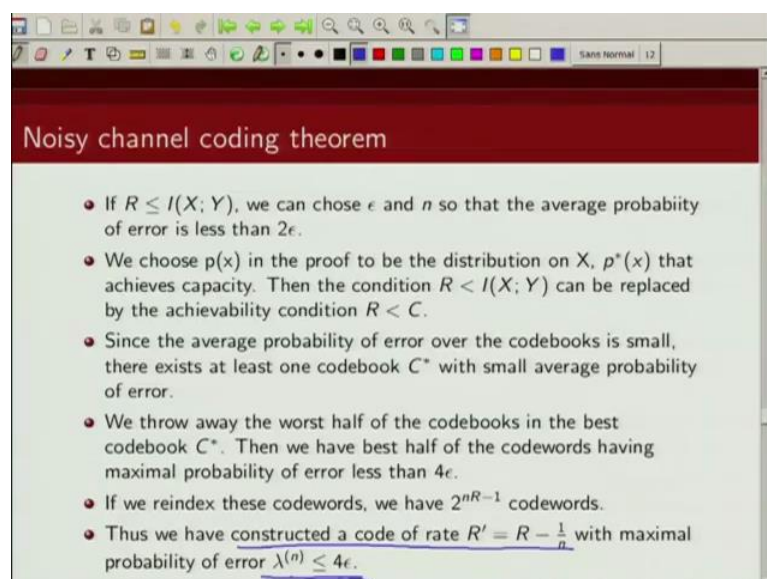
So, then we can compute the probability of error as we said it is independent of what index we have sent that is because the symmetry of code construction this is upper bounded we have shown by probability of the complement of event e_1 happening, given index 1 was sent plus this probability now this from the property of joint typicality is a very

small quantity you will call it epsilon this we know is upper bounded by this probability.

So, when we sum it above from i equal to 2, 2 raised to power nR what we get is equal to epsilon plus 2 raised to power nR minus 1 times this quantity now we just collect terms. We collect terms containing epsilon. So, this will become epsilon plus 2 raised to power $3n$ epsilon times 2 raised to power minus n mutual information minus r . Now as real as our transmission rate are is less than this quantity mutual information this particular term is going to be greater than 0 and for large n . So, this is this is this term is greater than equal to 0 if a large n 2 raised to power minus n this term will go towards 0.

So, if n is sufficiently large and our weight is less than see this we can combine as 2 raised to power minus n minus 3 epsilon minus R this is what you will get. So, as long as R is less than this quantity epsilon is very small this term will be positive for large n this whole term will go towards 0. So, some epsilon we calling it, so this will be less than some small quantity epsilon. So, what we have shown here is probability of error that is probability of decoding making a decoding error is less than 2 times epsilon provided our transmission rate is below this mutual information in x and y minus 3 epsilon and provided this length n is very large if n goes to infinity.

(Refer Slide Time: 24:34)



Noisy channel coding theorem

- If $R \leq I(X; Y)$, we can choose ϵ and n so that the average probability of error is less than 2ϵ .
- We choose $p(x)$ in the proof to be the distribution on X , $p^*(x)$ that achieves capacity. Then the condition $R < I(X; Y)$ can be replaced by the achievability condition $R < C$.
- Since the average probability of error over the codebooks is small, there exists at least one codebook C^* with small average probability of error.
- We throw away the worst half of the codebooks in the best codebook C^* . Then we have best half of the codewords having maximal probability of error less than 4ϵ .
- If we reindex these codewords, we have 2^{nR-1} codewords.
- Thus we have constructed a code of rate $R' = R - \frac{1}{n}$ with maximal probability of error $\lambda^{(n)} \leq 4\epsilon$.

Now, we can tighten this a bit. So, if you take care to be less equal to mutual information and we can choose our ϵ and n such that probability of error is less than 2ϵ . Now we can now choose our p of x remember we are generating our code words using this distribution p of x , now we can choose our p of x such that we choose our distribution that we maximize mutual information. So, we can choose a distribution i am calling it $p^* x$ that will achieve it is capacity then this condition R is less than mutual information can be replaced by the condition that R is less than capacity. So, what I have shown you. So, far is average probability of error over all possible code words is very small it is basically less than equal to 2ϵ now, if the average error probability is small; that means, there exist at least 1 codeword whose probability of error maybe less than the average probability of error right and remember in the beginning we said that if average probability of error goes to 0 we can also show that maximal error probability will also go to 0.

So, that is what we are going to show now or what we have shown. So, far is the average probability of error provided the transmission rate R is less than capacity is less than equal to 2ϵ , now what we do next is as follows we throw away half of the code words and what are these half of the code words we threw away worst half code words. So, we throw the code words which causes larger error. So, we throw half of them now we are left with other half which consists of good code words and the good I mean the one which will cause less probability of error now it can be shown that for these half code words the maximal error probability is less than 4ϵ . So, for these best half code words the maximal error probability is less than 4ϵ the reason being if it is more than 4ϵ then, the average probability of error overall bad and good code words cannot be less than 2ϵ .

So, by throwing away half of the worst code words we have shown with the half the number of code words left we have shown that the maximal probability of error is now bounded by 4ϵ now since we have thrown away half of the code words. So, we can reindex these code words and now we have $2^n R$ divided by 2 that is $2^{n-1} R$ minus 1 code words; so now, we have constructed a code with rate I am calling the rate R' which is R minus $1/n$ whose maximal probability of error is less than 4ϵ . So, if n is very large. So, this hardly any rate loss and

you can see the maximal probability of error goes to 0 as that side is large because maximal error probability is bounded by upper bounded by 4 times epsilon.

So, this proves that when our rate R is less than capacity then, we can reliably communicate because we have shown that the maximal probability of error is bounded by 4 epsilon next we are going to show the converse of this theorem now what is the converse of this theorem. So, we are going to show that if our transmission rate is above channel capacity; that means our R is greater than C then probability of error is bounded away from 0 so.

(Refer Slide Time: 29:27)

Converse to noisy channel coding theorem

- If information bits from a binary symmetric source (BSS) are sent at a rate R via a DMC of capacity C without feedback, then the bit error probability at the destination satisfies

$$P_b \geq H^{-1}\left(1 - \frac{C}{R}\right), \text{ if } R > C$$

where H^{-1} denotes the inverse binary entropy function defined by $H^{-1}(x) = \min\{p : H(p) = x\}$.

The slide also includes a graph of the binary entropy function $H(p)$ and a block diagram of the communication system. The graph shows a symmetric curve peaking at 1 when $p=0.5$. The block diagram shows the flow from Information Destination to Channel Decoder, then to DMC, then to Channel Encoder, then to BSS, and finally back to Information Destination via a feedback loop.

So, we will show that if transmission rate R is more than channel capacity then, probability of error is non 0. So, if information bit's. So, we consider a binary symmetric source if information bit's from binary symmetric source are sent at a rate R via discrete memory less channel whose capacity is C then probability of error is never bounded by h inverse 1 minus C by R where the transmission rate is our capacity.

So, if the transmission rate is above channel capacity then probability of error is lower bounded by quantity which is non0 and this h of p is nothing but our binary entropy function if you recall binary entropy function looks like this. So, here 2 1 at 0.5 this is

one. So, to proof this result let us look at the block diagram. So, we have a binary symmetric source output of that are u is $u_1, u_2, u_3, \dots, u_k$ and these bit's are sent to a channel encoder that generates these code words $x_1, x_2, x_3, \dots, x_n$. So, the rate of the code is k by n and these code words are sent over a discrete memory less channel what we receives is y_i . So, you see $y_1, y_2, y_3, \dots, y_n$. Now these are sent to a channel decoder which will try to estimate what were the information bit's that we have sent. So, the channel decoder will try to estimate u_1, u_2, u_3 . So, these estimates I am denoting by $\hat{u}_1, \hat{u}_2, \dots, \hat{u}_k$ and this is my. So, let us now show that if the transmission rate is above channel capacity then the probability of error is lower bounded by a non 0 quantity.

(Refer Slide Time: 31:44)

Converse to noisy channel coding theorem

Proof:

- Let us consider BSS, i.e. DMS with $P_U(0) = P_U(1) = 1/2$. Therefore $H(U) = 1$ bit.
- Also for DMC without feedback,

$$U \rightarrow X \rightarrow Y \rightarrow \hat{U}$$

$$P(y_1, \dots, y_n | x_1, \dots, x_n) = \prod_{i=1}^N P(y_i | x_i)$$

$$\Rightarrow H(Y_1 \dots Y_n | X_1 \dots X_n) = \sum_{i=1}^N H(Y_i | X_i)$$

- Rate of transmission, R is given by $R = K/N$ bits/use.
- Applying data processing lemma, we get

$$I(U_1 \dots U_K; \hat{U}_1 \dots \hat{U}_K) \leq I(X_1 \dots X_N; \hat{U}_1 \dots \hat{U}_K)$$

So, since we are considering a binary symmetric source. So, probability of 0 and probability of 1 is half it is a symmetric source and it is a binary source. So, it is sum is 0s and ones if probability has. So, we can write uncertainty u as 1 bit. Now for a discrete memory less channel without feedback probability of $y_1, y_2, y_3, \dots, y_n$ given $x_1, x_2, x_3, \dots, x_n$ can be given by probability of y_i given x_i and product taken from i go from 1 to n . So, we can write the uncertainty in $y_1, y_2, y_3, \dots, y_n$ given $x_1, x_2, x_3, \dots, x_n$ by summation of uncertainty of y_i given x_i and remember our transmission rate is k by n bit's per use.

So, we are going to first apply data processing lemma, what does data processing lemma

says that further processing of data does not increase information. So, what you had was we had u, x, y and $\hat{u}$ follow mark of chain. So, mutual information between u is and $\hat{u}$ is this is going to be less than mutual information between x is and $\hat{u}$ is. So, that is what i am writing here next again we will apply data processing lemma. Let we write down the mark of chain. So, u, x, y and $\hat{u}$ follows a mark of chain.

(Refer Slide Time: 33:24)

Converse to noisy channel coding theorem

- Applying data processing lemma we also get,

$$I(X_1 \dots X_N; \hat{U}_1 \dots \hat{U}_K) \leq I(X_1 \dots X_N; Y_1 \dots Y_N)$$
- From above two inequalities, we get

$$I(U_1 \dots U_K; \hat{U}_1 \dots \hat{U}_K) \leq I(X_1 \dots X_N; Y_1 \dots Y_N)$$
- We can write $I(X_1 \dots X_N; Y_1 \dots Y_N)$ as

$$I(X_1 \dots X_N; Y_1 \dots Y_N) = \frac{H(Y_1 \dots Y_N) - H(Y_1 \dots Y_N | X_1 \dots X_N)}{N}$$

$$= \frac{H(Y_1 \dots Y_N) - \sum_{i=1}^N H(Y_i | X_i)}{N}$$

$$\leq \frac{\sum_{i=1}^N [H(Y_i) - H(Y_i | X_i)]}{N}$$

$$\leq \sum_{i=1}^N I(X_i; Y_i) \leq NC$$

So, we have shown the mutual information between u is and $\hat{u}$ is less than mutual information between x is and $\hat{u}$ is. Now we can say that mutual information between x is and $\hat{u}$ is less than mutual information between x is and y is again that follows from data processing lemma. So, if we combine these 2 result let us call it a and call it b if i combine a and b what i get is this condition that mutual information between $u_1, u_2, \dots, u_k$ and $\hat{u}_1, \hat{u}_2, \dots, \hat{u}_k$ this is less than mutual information between $x_1, x_2, \dots, x_n$ and $y_1, y_2, \dots, y_n$ next.

Let us try to simplify this term mutual information between $x_1, x_2, x_3, \dots, x_n$ and $y_1, y_2, y_3, \dots, y_n$ now from the definition of mutual information i can write this as joint entropy of $y_1, y_2, y_3, \dots, y_n$ minus conditional entropy of $y_1, y_2, y_3, \dots, y_n$ given $x_1, x_2, x_3, \dots, x_n$ now since this is a discrete memory less channel without feedback. This term can be written as summation of uncertainty in y_i given x_i . So, i can write this term as joint entropy minus

the summation from i goes from 1 to n h of y_i given x_i . Now this joint entropy can be written using symbol as h of y_1 plus h of y_2 given y_1 plus h of y_3 given $y_1 y_2$ plus h of y_n given $y_1 y_2 y_n$ minus 1. Now this can be further written as this is upper bounded by h of y_1 plus h of y_2 plus h of y_3 plus h of y_n .

Now, here I used the fact that conditioning cannot increase entropy. So, if I do that then I can write that this joint entropy is upper bounded by summation of entropy of $h y_i$ where i goes from 1 to n . So, then from here I can write this expression as summation i goes from 1 to n h of y_i minus h of y_i given x_i and what is this term this term is mutual information between x_i and y_i . So, then I can write this as mutual information between x_i and y_i and mutual information between x_1, y_1, x_2, y_2 that is less than channel capacity.

(Refer Slide Time: 37:46)

Converse to noisy channel coding theorem

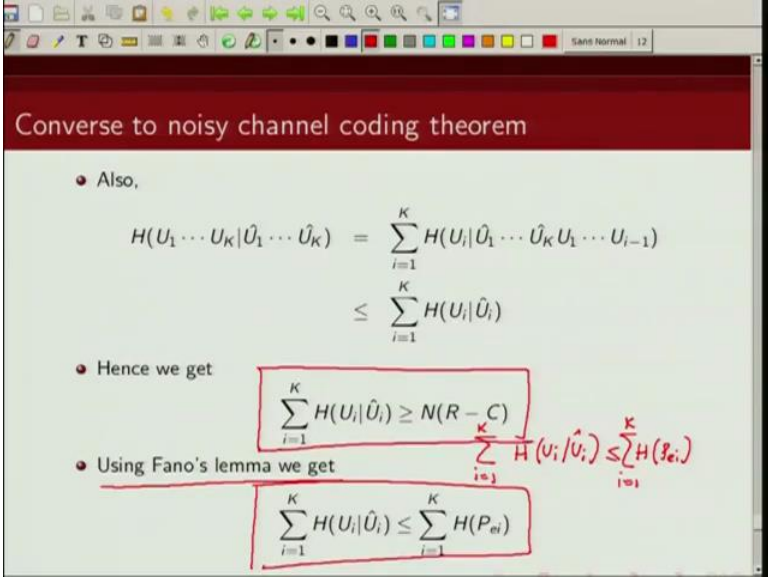
- This implies that $I(U_1 \dots U_K; \hat{U}_1 \dots \hat{U}_K) \leq NC$
- We define bit error probability as
$$P_b = \frac{1}{K} \sum_{i=1}^K P_{ei}$$
 where $P_{ei} = P(\hat{U}_i \neq U_i)$.
- We can $H(U_1 \dots U_K | \hat{U}_1 \dots \hat{U}_K)$ as
$$\begin{aligned} H(U_1 \dots U_K | \hat{U}_1 \dots \hat{U}_K) &= H(U_1 \dots U_K) - I(U_1 \dots U_K; \hat{U}_1 \dots \hat{U}_K) \\ &= K - I(U_1 \dots U_K; \hat{U}_1 \dots \hat{U}_K) \\ &\geq K - NC \\ &= N(R - C) \end{aligned}$$
 $\frac{K}{N} = R$

So, I can write then mutual information between x_1, x_2, x_3, x_n and y_1, y_2, y_3, y_n this is less than equal to n times C next. Now we know that from data processing lemma we know that this quantity is less than this quantity and this is less than n times c . So, mutual information between u_1, u_2, u_k and u_1 hat, u_2 hat, u_k hat this will be less than equal to n times c . So, combining these results we get this next.

That is defining probability of bit error. So, basically it is this is some all errors divided by the block size that we give me probability of bit error and probability of error when does an error happen when, my transmitted date is not same as the decoded date. Now let us try to write uncertainty in u_1, u_2, u_3, u_k given $u_1 \text{ hat}, u_2 \text{ hat}, u_k \text{ hat}$ and these from the definition of mutual information can be written in this particular form what is this quantity the u is are the output of a binary symmetric source and this is independently identically distributed. So, h of u_1, u_2, u_3, u_k is a h of u_1 plus h of u_2 plus h of u_k and h of u_i was 1 bit because is a binary symmetric source. So, this particular term will be equal to k k minus this now we have proofed earlier that this mutual information is less than equal to n time's c . So, if we subtract the larger quantity this will become greater than equal to.

So, then this is greater than equal to k minus n times C and what is k by n is r . So, k is n times r . So, when we write k is n time R and we take n out what we get is uncertainty in u_1, u_2, u_k given $u_1 \text{ hat}, u_2 \text{ hat}, u_k \text{ hat}$ is greater than equal to n times R minus c .

(Refer Slide Time: 40:29)



Converse to noisy channel coding theorem

- Also,

$$H(U_1 \dots U_K | \hat{U}_1 \dots \hat{U}_K) = \sum_{i=1}^K H(U_i | \hat{U}_1 \dots \hat{U}_K U_1 \dots U_{i-1})$$

$$\leq \sum_{i=1}^K H(U_i | \hat{U}_i)$$
- Hence we get

$$\sum_{i=1}^K H(U_i | \hat{U}_i) \geq N(R - C)$$
- Using Fano's lemma we get

$$\sum_{i=1}^K H(U_i | \hat{U}_i) \leq \sum_{i=1}^K H(p_{ei})$$

Handwritten red notes:

$$\sum_{i=1}^K H(U_i | \hat{U}_i) \leq \sum_{i=1}^K H(p_{ei})$$

Now, let us further simplify this channel. So, using chain rule i can write this conditional entropy in this particular fashion and as we know, conditioning cannot increase entropy. So, this particular term can be upper bounded by uncertainty in u_i given $u_i \text{ hat}$. So, we

have shown that this particular term is greater than equal to $n C$ and earlier and we are showing that this particular term is greater than this now combining these 2 results we get that $h(u_i)$ given $\hat{u}_i$ sum over all i from 1 to k this is greater than equal to n times R minus c . Now what do we know from Fano's lemma from Fano's Lemma, we know that uncertainty in u_i given $\hat{u}_i$ is less than equal to $p_{ei} \log(l - 1)$ where this angular being u_i can take l possible values plus $h(p_{ei})$.

So, this is in this case l is 2 because of binary random variable. So, Fano's Lemma for this will be $h(u_i)$ given $\hat{u}_i$ is less than equal to $h(p_{ei})$ because, the other term $p_{ei} \log(l - 1)$ term that is 0. So, and if you take summation over all i from 1 to k both sides we get this. So, that is what I am saying using Fano's Lemma we get this relation right and. So, on what we proved earlier we showed that this term can be greater than equal to n times R minus c . So, in other words this particular term is lower bounded by n times R minus c .

(Refer Slide Time: 43:04)

Converse to noisy channel coding theorem

- Combining the previous results we get
$$\frac{1}{K} \sum_{i=1}^K H(P_{ei}) \geq \frac{N}{K} (R - C) = 1 - \frac{C}{R}$$
- Since $H(p)$ is a concave function, we get
$$\frac{1}{K} \sum_{i=1}^K H(P_{ei}) \leq H\left(\frac{1}{K} \sum_{i=1}^K P_{ei}\right) = H(P_b)$$
- Combining the above results we get
$$\underline{H(P_b)} \geq \underline{1 - \frac{C}{R}} \quad \text{Handwritten: } R > C$$

So, we can write summation of $h(p_{ei})$ is lower bounded by n times R minus c , and if you divide both sides by 1 by k . So, divide both side by 1 by k what you get is on the right hand side n by k times R minus C which is nothing, but 1 minus capacity divide by the rate next.

This binary entropy function is a concave function and from Jensen's inequality, what do we know about the concave function? So, Jensen's inequality says if the function f of x is concave then expected value of the function is less than equal to function evaluated at expected value. If f of x is concave, so since this binary entropy function is concave expected value of the function here is a binary entropy function should be less than equal to function evaluated at expected value, and what is this term this term is nothing, but average bit error probability.

So, by combining this result, with this result we have shown that binary entropy function of this probability of error is lower bounded by $1 - C/r$, and remember R is greater than C here. So, when R is greater than C this is a non 0 quantity. So, probability of error cannot be 0. So, if we try to transmit at rate above capacity our probability of error will be non 0. So, this proves the converse of the noisy channel coding theorem. So, with this we will conclude our discussion on noisy channel coding theorem.

Thank you.