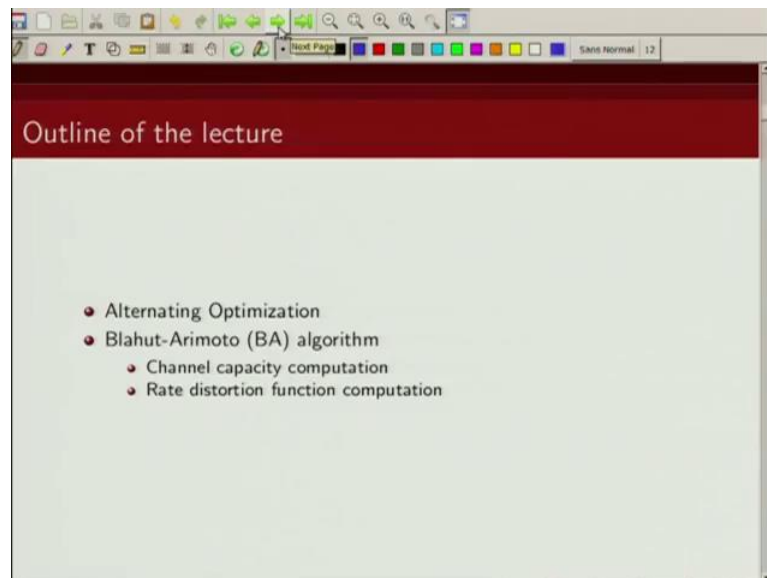


An Introduction to Information Theory
Prof. Adrish Banerjee
Department of Electronics and Communication Engineering
Indian Institute of Technology, Kanpur

Lecture – 14A
Blahut-Arimoto Algorithm

Welcome to the course on An Introduction to Information Theory. So, in this lecture, we are going to talk about an iterative algorithm to compute channel capacity and rate distortion function, and this is known as Blahut-Arimoto algorithm named after these two scientists who independently came up with this algorithm.

(Refer Slide Time: 00:54)



So, first we will talk about this algorithm in a general setting and then so will be talking about alternating optimization algorithm and then we will talk specific about this Blahut-Arimoto algorithm. In particular, we will talk about how we can use this algorithm to compute channel capacity, and how we can use this algorithm to compute rate distortion function. Now, for this lecture, we are going to use the book by Raymond Yeung on Information Theory and Network Coding, this is chapter 9 of that book.

(Refer Slide Time: 01:27)

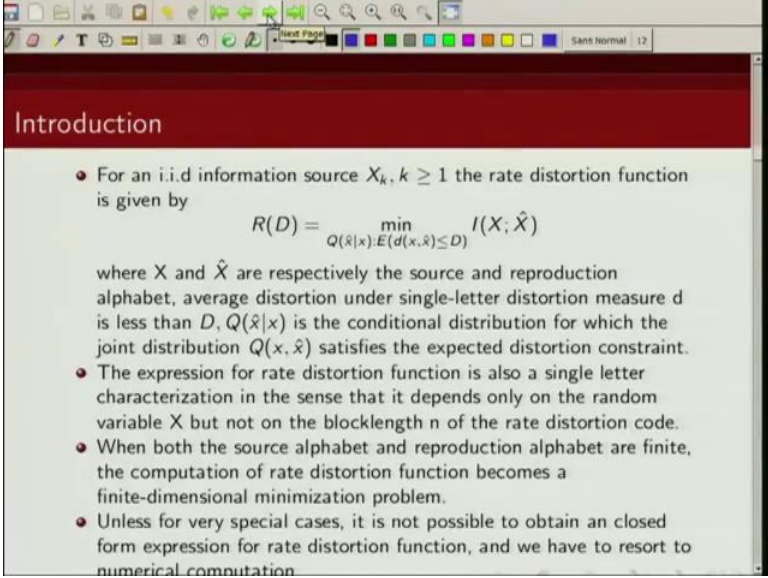
Introduction

- For a discrete memoryless channel, $p(y/x)$, the channel capacity is given by
$$C = \max_{r(x)} I(X; Y)$$
where X and Y are respectively the input and output of the channel, and $r(x)$ is the input distribution.
- The expression for channel capacity is called a single letter characterization in the sense that it depends only on the transition matrix of the channel but not on the blocklength n of the code.
- When both the input and output alphabet are finite, the computation of channel capacity becomes a finite-dimensional maximization problem.
- Unless for very special cases, it is not possible to obtain an closed form expression for channel capacity.

So, as we know for a discrete memoryless channel with transition matrix given by p of y given x , channel capacity is given as maximum mutual information between the input X and the output Y , and the maximization is taken over all input distribution. Here I am denoting the input distribution by r of x ; X is input to the channel and Y is the output to the channel. Now, as we know this expression for channel capacity is what we call as single letter characterization. Now, we know this channel capacity depends on the transition matrix of the channel; it does not depend on the block length of the code used.

And when the input alphabet size and the output alphabet size that is finite in that case this computation of this channel capacity is nothing but a finite-dimensional maximization problem. Now, we know that except for some simple specific cases, we do not have a close form expression for channel capacity. So, to compute channel capacity in general, we will have to resort to let us say numerical computation of something like that. Now Blahut-Arimoto algorithm is an algorithm to iteratively compute the channel capacity and rate distortion function.

(Refer Slide Time: 03:10)



Introduction

- For an i.i.d information source $X_k, k \geq 1$ the rate distortion function is given by
$$R(D) = \min_{Q(\hat{x}|x): E(d(x, \hat{x})) \leq D} I(X; \hat{X})$$
where X and $\hat{X}$ are respectively the source and reproduction alphabet, average distortion under single-letter distortion measure d is less than D , $Q(\hat{x}|x)$ is the conditional distribution for which the joint distribution $Q(x, \hat{x})$ satisfies the expected distortion constraint.
- The expression for rate distortion function is also a single letter characterization in the sense that it depends only on the random variable X but not on the blocklength n of the rate distortion code.
- When both the source alphabet and reproduction alphabet are finite, the computation of rate distortion function becomes a finite-dimensional minimization problem.
- Unless for very special cases, it is not possible to obtain a closed form expression for rate distortion function, and we have to resort to numerical computation.

Similarly, if we look at the rate distortion function this is defined as minimum mutual information between the source alphabet and the reproduce alphabet which is denoted by $\hat{X}$. And this minimization is taken over all conditional distribution of Q of $\hat{x}$ given x such that the expected single letter distortion is less than some quantity D . So, Q of x is the conditional distribution and of $\hat{x}$ given x and q of $x, \hat{x}$ denotes a joint distribution which satisfy the distortion constraint which is given by this here.

Now, here also this is the single letter characterization; it does not depend on the block length of the rate distortion code, it already depends on the random variable x . So, when the source alphabet and reproduce alphabet size is finite again this becomes a finite dimensional minimization problem. And again here in most of the cases, we do not have a close form expression for the rate distortion function. So, to compute rate distortion function then we will have to resort to some sort of numerical method. So, this gives a motivation to study an iterative algorithm to compute this rate distortion function and channel capacity.

(Refer Slide Time: 04:58)

Alternating Optimization

- Consider the double supremum

$$\sup_{u_1 \in A_1} \sup_{u_2 \in A_2} f(u_1, u_2)$$
 where A_i is a convex subset of $\mathbb{R}^n$ for $i = 1, 2$ and f is a real function defined on $A_1 \times A_2$.
- The function f is bounded from above, and is continuous and has continuous partial derivatives on $A_1 \times A_2$.
- Assume for all $u_2 \in A_2$, there exists a unique $c_1(u_2) \in A_1$ such that

$$f(c_1(u_2), u_2) = \max_{u'_1 \in A_1} f(u'_1, u_2)$$
- Assume for all $u_1 \in A_1$, there exists a unique $c_2(u_1) \in A_2$ such that

$$f(u_1, c_2(u_1)) = \max_{u'_2 \in A_2} f(u_1, u'_2)$$
- Let $u = (u_1, u_2)$ and $A_1 \times A_2$. Then the optimization problem is

$$\sup f(u)$$

So, first let me talk about this algorithm in a general setting. So, I am going to talk about alternating optimization algorithm and then we will come to the specifics of Blahut-Arimoto algorithm. So, let us see you want to compute double supremum, so you have function f over u_1 and u_2 ; u_1 belongs to a convex set A_1 , and u_2 belongs to convex set A_2 . Now you want to compute double supremum of this function f , which is function of u_1 and u_2 ; and this function f is a real function defined over $A_1 \times A_2$. The function is also bounded from above. So, this is function is upper bounded by some quantity finite quantity, and it is continuous, also it is partial derivatives, it has continuous partial derivatives defined. So, this is the conditions on the function f .

Now, assume for all u_2 belonging to this convex set A_2 , there exist mapping $c_1 u_2$ which belongs to this convex set A_1 such that f of $c_1 u_2$. And u_2 is equal to maximization of this function f of u_1 dash u_2 where u_1 dash belongs to this convex set A_1 . Similarly, let us assume that for all u_1 belonging to this convex set A_1 there exist. A unique mapping $c_2 u_1$ which belongs to this convex set A_2 such that this condition is satisfied that f of u_1 and $c_2 u_1$ is equal to maximizing of this function f of u_1 and u_2 dash where u_2 dash belongs to A_2 . So, then we can write this double supremum problem as we are computing a supremum of f of u where u is this and we are computing this over u which belongs to Cartesian to that of A_1 times A_2 .

(Refer Slide Time: 07:55)

Alternating Optimization

- Let $\mathbf{u}^{(k)} = (\mathbf{u}_1^{(k)}, \mathbf{u}_2^{(k)})$ for $k > 0$.
- Let $\mathbf{u}_1^{(0)}$ be any arbitrary chosen vector in A_1 , and let $\mathbf{u}_2^{(0)} = c_1(\mathbf{u}_1^{(0)})$
- For $k \geq 1$, $\mathbf{u}^{(k)}$ is defined as

$$\begin{aligned} \mathbf{u}_1^{(k)} &= c_1(\mathbf{u}_2^{(k-1)}) \\ \mathbf{u}_2^{(k)} &= c_2(\mathbf{u}_1^{(k)}) \end{aligned}$$
- Let the function f at k^{th} iteration $f^{(k)} = f(\mathbf{u}^{(k)})$, then

$$\begin{aligned} f^k &= f(u^k) = f(u_1^k, u_2^k) \\ &\geq f(u_1^k, u_2^{k-1}) \\ &\geq f(u_1^{k-1}, u_2^{k-1}) = f^{k-1} \end{aligned}$$

for $k \geq 1$.

So, let us take the alternating optimization algorithm. So, let us see at some time k , $\mathbf{u}^k$ is given by $\mathbf{u}_1^k$ and $\mathbf{u}_2^k$. So, we start up with some initial time k is equal to 0, it start an initial value of $\mathbf{u}_1$ which I am denoting by $\mathbf{u}_1^0$. So, let this been arbitrary chosen vector, and it belongs to this convex set A_1 . Then $\mathbf{u}_2^0$ is given by $c_1(\mathbf{u}_1^0)$, and this belongs to this convex set A_2 . In general, for k greater than equal to 1, we can define this $\mathbf{u}^k$ by $\mathbf{u}_1^k$ and $\mathbf{u}_2^k$ where $\mathbf{u}_1^k$ is nothing but c_1 of $\mathbf{u}_2^{k-1}$ and $\mathbf{u}_2^k$ is c_2 of $\mathbf{u}_1^k$.

Now, let this function f , which is basically this at k^{th} iteration, so we denoting this value by this is nothing but function f at k^{th} iteration is f of $\mathbf{u}^k$, which is f of $\mathbf{u}_1^k$ and $\mathbf{u}_2^k$. Now, this is greater than equal to f of $\mathbf{u}_1^k$ and $\mathbf{u}_2^{k-1}$, which in turn is greater than f of $\mathbf{u}_1^{k-1}$ and $\mathbf{u}_2^{k-1}$. Now, these two relations follow from these two properties So, we have this property and this property. So, for all $\mathbf{u}_2$ belonging to a A_2 there exist a unique $c_1(\mathbf{u}_2)$ will belongs to this convex set A_1 such that f of $c_1(\mathbf{u}_2)$ is nothing but maximizing this function f of $\mathbf{u}_1$ dash $\mathbf{u}_2$ when $\mathbf{u}_1$ dash belongs to A_1 . And similarly there exist for all $\mathbf{u}_1$ belonging to A_1 there exist a unique $c_2(\mathbf{u}_1)$ such that this relation holds. So, because of these two properties we know that these two hold. Now so what we have shown then is f of k , f as a function of k as A is basically a non decreasing function. So, f is a non decreasing function.

(Refer Slide Time: 11:05)

Alternating Optimization

- Since f^k is non decreasing and f is bounded from above, f^k must converge.
- Replacing f by $-f$, the optimization criterion becomes

$$\inf_{u_1 \in A_1} \inf_{u_2 \in A_2} f(u_1, u_2) \leftarrow R(D)$$

where f is bounded from below.

- The double infimum can be computed by the same alternating optimization algorithm.

So, f^k is a non decreasing function, and remember it is bounded from above that was one of the property of f^k . If you go back, this function f is bounded from above. So, f^k is non-increasing, but it is bounded from above. So, what does it mean it means that f^k must converge must have a limit as k goes to very high. So, f^k converges to some finite quantity.

Now, if you replace this function f by minus f then in this problem if we replace this f by minus f what we get is a double infimum of this function f , where we have similar conditions that f has to be continuous partial derivatives has to be continuous. f is bounded from below then this double infimum can be computed also using this alternating optimization algorithm. Now, this form we are going to use for computing the rate distortion function

(Refer Slide Time: 12:34)

Alternating Optimization

- Consider the double supremum

$$\sup_{\mathbf{u}_1 \in A_1} \sup_{\mathbf{u}_2 \in A_2} f(\mathbf{u}_1, \mathbf{u}_2) \leftarrow C$$
 where A_i is a convex subset of $\mathbb{R}^n$ for $i = 1, 2$ and f is a real function defined on $A_1 \times A_2$.
- The function f is bounded from above, and is continuous and has continuous partial derivatives on $A_1 \times A_2$.
- Assume for all $\mathbf{u}_2 \in A_2$, there exists a unique $\mathbf{c}_1(\mathbf{u}_2) \in A_1$ such that

$$f(\mathbf{c}_1(\mathbf{u}_2), \mathbf{u}_2) = \max_{\mathbf{u}_1' \in A_1} f(\mathbf{u}_1', \mathbf{u}_2)$$
- Assume for all $\mathbf{u}_1 \in A_1$, there exists a unique $\mathbf{c}_2(\mathbf{u}_1) \in A_2$ such that

$$f(\mathbf{u}_1, \mathbf{c}_2(\mathbf{u}_1)) = \max_{\mathbf{u}_2' \in A_2} f(\mathbf{u}_1, \mathbf{u}_2')$$
- Let $\mathbf{u} = (\mathbf{u}_1, \mathbf{u}_2)$ and $A_1 \times A_2$. Then the optimization problem is

$$\sup_{\mathbf{u}} f(\mathbf{u})$$

Handwritten notes: A red arrow points from the expression $\sup_{\mathbf{u}_1 \in A_1} \sup_{\mathbf{u}_2 \in A_2} f(\mathbf{u}_1, \mathbf{u}_2)$ to a red 'C'. Another red arrow points from the expression $\sup_{\mathbf{u}_2' \in A_2} f(\mathbf{u}_1, \mathbf{u}_2')$ to the handwritten note $\sup_{\mathbf{u} \in A_1 \times A_2} f(\mathbf{u})$.

Whereas, this form we are going to use to compute channel capacity.

(Refer Slide Time: 12:42)

Blahut-Arimoto Algorithm: Channel Capacity Computation

- Let $r(x)p(y|x)$ be a given joint distribution on $X \times Y$ such that $r > 0$, and let q be a transition matrix from Y to X . Then

$$\max_q \sum_x \sum_y r(x)p(y|x) \log \frac{q(x|y)}{r(x)} = \sum_x \sum_y r(x)p(y|x) \log \frac{q^*(x|y)}{r(x)}$$
- where the maximization is taken over all q such that $q(x|y) = 0$ if and only if $p(y|x) = 0$ and

$$q^*(x|y) = \frac{r(x)p(y|x)}{\sum_{x'} r(x')p(y|x')}$$
- i.e. the maximizing q is the one which corresponds to the input distribution r and the transition matrix $p(y|x)$.

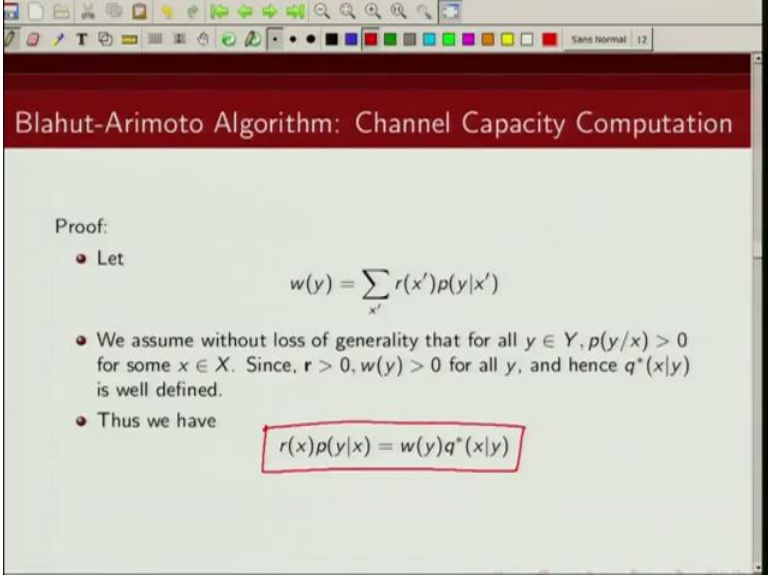
Handwritten notes: A red box highlights the expression $q^*(x|y) = \frac{r(x)p(y|x)}{\sum_{x'} r(x')p(y|x')}$. A red arrow points from the expression $\max_q \sum_x \sum_y r(x)p(y|x) \log \frac{q(x|y)}{r(x)}$ to the boxed expression.

So, let us first prove a lemma that we are going to use and then we will state our Blahut-Arimoto algorithm for computation of channel capacity. So, r of x is my input distribution and p of y given x is my transition probability. Now, let r of x times p of y given x be the joint distribution on X and Y such that r of x is greater than intersects r of x is strictly positive. And let q be the transition matrix from Y to X . Then what this

lemma says is as follows. $r(x)$ multiplied by $p(y|x)$ given x log of $q(x|y)$ divide by $r(x)$.

If you sum it over all x and y , and you maximize over all q then this is nothing but summation over all x and y of $r(x)p(y|x)$ log of $q^*(x|y)$ divide by $r(x)$, where this $q^*(x|y)$ is given by this expression. And this maximization is over all q such that $q(x|y) > 0$, if and only if $p(y|x) > 0$. So, what we have seen is if you try to maximize this over all q then this is basically given by this expression where expression of q^* is given here. So, we are going to use the divergence inequality to prove this result.

(Refer Slide Time: 14:56)



Blahut-Arimoto Algorithm: Channel Capacity Computation

Proof:

- Let $w(y) = \sum_{x'} r(x') p(y|x')$
- We assume without loss of generality that for all $y \in Y$, $p(y|x) > 0$ for some $x \in X$. Since, $r > 0$, $w(y) > 0$ for all y , and hence $q^*(x|y)$ is well defined.
- Thus we have $r(x)p(y|x) = w(y)q^*(x|y)$

So, let us see how we are going to prove this. Let $w(y)$ is given by this expression $r(x)$ prime $p(y|x)$ prime summation over all x prime. So, this is nothing but this one, this term let us denote this by $w(y)$. Now, without loss of generality, we will assume that for all Y $p(y|x)$ is greater than equal to 0 for some x and since we will consider a strictly positive r , so r greater than 0. So, then $w(y)$ will be greater than 0 for all y . And if $w(y)$ is greater than 0 you can go back and see if this is greater than 0 and these are greater than 0 then this will be also greater than 0.

(Refer Slide Time: 16:17)

Blahut-Arimoto Algorithm: Channel Capacity Computation

- Let $r(x)p(y|x)$ be a given joint distribution on $X \times Y$ such that $r > 0$, and let q be a transition matrix from Y to X . Then

$$\max_q \sum_x \sum_y r(x)p(y|x) \log \frac{q(x|y)}{r(x)} = \sum_x \sum_y r(x)p(y|x) \log \frac{q^*(x|y)}{r(x)}$$

where the maximization is taken over all q such that $q(x|y) = 0$ if and only if $p(y|x) = 0$ and

$$q^*(x|y) = \frac{r(x)p(y|x)}{\sum_{x'} r(x')p(y|x')} = w(y)$$

i.e. the maximizing q is the one which corresponds to the input distribution r and the transition matrix $p(y|x)$.

Now, from this, this is my $w(y)$. So, I can write $w(y)$ times $q^*(x|y)$ to be equal to $r(x)p(y|x)$.

(Refer Slide Time: 16:39)

Blahut-Arimoto Algorithm: Channel Capacity Computation

Proof (Contd.):

- Consider

$$\begin{aligned} & \sum_x \sum_y r(x)p(y|x) \log \frac{q^*(x|y)}{r(x)} - \sum_x \sum_y r(x)p(y|x) \log \frac{q(x|y)}{r(x)} \\ &= \sum_x \sum_y r(x)p(y|x) \log \frac{q^*(x|y)}{q(x|y)} \\ &= \sum_y \sum_x w(y)q^*(x|y) \log \frac{q^*(x|y)}{q(x|y)} \\ &= \sum_y w(y) \sum_x q^*(x|y) \log \frac{q^*(x|y)}{q(x|y)} \\ &= \sum_y w(y) D(q^*(x|y) || q(x|y)) \\ &\geq 0 \end{aligned}$$

$D(p||q) \geq 0$

So, this is what I am writing here. Next, you just recall we want to show that if we maximize this over all q , this is given by this expression where q^* is this. So, the way we are going to show that this is a maximum value of this function maximize over q . We are going to prove this by showing that this minus value of this function for some other q that is always greater than equal to 0. If we can show that then we have proved that q^*

x given y is the 1 that maximizes this function. So, what we are going to show is this function evaluated at q star minus this function you have evaluated at q . Now, so if we take this is a common term, if we take this out this is common term here. If we take this out, we get $\log$ of q star x given y divided by r of x minus $\log$ of q of x given y divide by r of x . So, this can be written as $\log$ of q star x given y by q of x given y . Now we just now showed that this term is nothing but $w y$ times q star. So, we plug in that value here, we get this expression, now this term does not depend on x . So, let us bring it out.

So, we have this term summation over x and this term summation over y . Now you can see this particular term can be written as divergence of q star and q . Now, we know that divergence between two distribution p and q is greater than is equal to 0, and since $w y$ is also greater than 0. So, this whole thing will be greater than equal to 0. So, then what we have shown is this function evaluated at q star that has the maximum value. So, we have proved this result, this lemma which says if you try to maximize this over all q , and then this is a maximum value where q star is given by this expression.

(Refer Slide Time: 20:00)

Blahut-Arimoto Algorithm: Channel Capacity Computation

- For a discrete memoryless channel $p(y|x)$

$$C = \sup_{r>0} \max_{\mathbf{q}} \sum_x \sum_y r(x)p(y|x) \log \frac{q(x|y)}{r(x)}$$

where the maximization is taken over all $\mathbf{q}$ such that $q(x|y) = 0$ if and only if $p(y|x) = 0$.

- Proof:** Let $I(\mathbf{r}, \mathbf{p})$ denote the mutual information $I(X; Y)$ when $\mathbf{r}$ is the input distribution for a channel with transition probability $p(y|x)$. Then

$$C = \max_{r>0} I(\mathbf{r}, \mathbf{p})$$

We state that for a discrete memoryless channel the capacity is given by this expression, where maximization is taken over all q such that q of x given y is 0 if and only if p of y given x is 0. So, let us prove this. So, we will first consider a strictly positive distribution of r let mutual information between x and y , let us denote by I of r p where r is input

distribution and p is my channel transition probability. Now, capacity can be given as maximize mutual information and maximization over all input distribution r .

(Refer Slide Time: 20:54)

Blahut-Arimoto Algorithm: Channel Capacity Computation

- Let $\mathbf{r}^*$ achieves C . If $\mathbf{r}^* > 0$, then

$$\begin{aligned}
 C &= \max_{\mathbf{r} \geq 0} I(\mathbf{r}, \mathbf{p}) \\
 &= \max_{\mathbf{r} > 0} I(\mathbf{r}, \mathbf{p}) \\
 &= \max_{\mathbf{r} > 0} \max_{\mathbf{q}} \sum_x \sum_y r(x) p(y|x) \log \frac{q(x|y)}{r(x)} \\
 &= \sup_{\mathbf{r} > 0} \max_{\mathbf{q}} \sum_x \sum_y r(x) p(y|x) \log \frac{q(x|y)}{r(x)}
 \end{aligned}$$

Now, let $\mathbf{r}^*$ achieves capacity. And if $\mathbf{r}^*$ is greater than 0 then we can write this as maximizing mutual information to $\mathbf{r}$ greater than equal to 0, this becomes maximizing mutual information when $\mathbf{r}$ is greater than 0. And from the definition of mutual information we can write this as in this particular way and this is nothing but supremum $\mathbf{r}$ greater than 0 maximizing over $\mathbf{q}$ of this function.

(Refer Slide Time: 21:42)

Blahut-Arimoto Algorithm: Channel Capacity Computation

- Next we consider the case when $\mathbf{r}^* \geq 0$. Since $I(\mathbf{r}, \mathbf{p})$ is continuous in $\mathbf{r}$, for any $\epsilon > 0$, there exists $\delta > 0$, such that if $\|\mathbf{r} - \mathbf{r}^*\| < \delta$, then

$$C - I(\mathbf{r}, \mathbf{p}) < \epsilon$$

where $\|\mathbf{r} - \mathbf{r}^*\|$ denotes the Euclidean distance between $\mathbf{r}$ and $\mathbf{r}^*$.

- In particular, there exists $\tilde{\mathbf{r}} > 0$ that satisfies the above equation, then

$$\begin{aligned}
 C &= \max_{\mathbf{r} \geq 0} I(\mathbf{r}, \mathbf{p}) \\
 &\geq \sup_{\mathbf{r} > 0} I(\mathbf{r}, \mathbf{p}) \\
 &\geq I(\tilde{\mathbf{r}}, \mathbf{p}) \\
 &> C - \epsilon
 \end{aligned}$$

Now, in case r^* is not strictly positive then since mutual information is continuous in r , so for any ϵ greater than 0 there exist a δ such that the Euclidean distance between r and r^* is less than δ then this relation over the C capacity minus mutual information is less than ϵ . So, in particular, there exist an input distribution r greater than 0 such that this relation holds. So, capacity is given by maximizing mutual information over all input distribution r this can be written as this is greater than equal to 2 supremum of r p a supreme over all r greater than equal to 0. And since r tilde r tilde this, this, this can be written as greater than is equal to mutual information between some distribution r tilde p and some distribution r tilde p and since r tilde satisfies this relation a satisfy this relation satisfy this condition. So, we can write that this mutual information between r tilde p will be C minus ϵ this last things follows from the fact that there exist in r tilde which satisfies this equation.

(Refer Slide Time: 23:24)

Blahut-Arimoto Algorithm: Channel Capacity Computation

- Thus we have $C - \epsilon < \sup_{r>0} I(r, \mathbf{p}) \leq C$
- By letting $\epsilon \rightarrow 0$, we conclude $C = \sup_{r>0} I(r, \mathbf{p}) = \sup_{r>0} \max_q \sum_x \sum_y r(x)p(y|x) \log \frac{q(x|y)}{r(x)}$

So, we have this, now if we let go ϵ to 0, we get the expression of capacity as this.

(Refer Slide Time: 23:36)

Blahut-Arimoto Algorithm: Channel Capacity Computation

- We use alternating optimization algorithm to compute capacity.
- We arbitrary choose a strictly positive input distribution in A_1 and let it be $\mathbf{r}^{(0)}$. We define $\mathbf{q}^{(0)}$ and in general $\mathbf{q}^{(k)}$ for $k \geq 0$.

$$q^{(k)}(x|y) = \frac{r^{(k)}(x)p(y|x)}{\sum_{x'} r^{(k)}(x')p(y|x')}$$

- In order to define $\mathbf{r}^{(k)}$ for $k \geq 1$, we need to find $\mathbf{r} \in A_1$ that maximizes the function for a given $\mathbf{q} \in A_2$, where the constraints on $\mathbf{r}$ are $\sum_x r(x) = 1$ and $r(x) > 0$ for all $x \in X$.

Now, let us see how we are going to use alternating optimization algorithm to compute the channel capacity. So, first we are going to choose a strictly positive input distribution which I am denoting by $\mathbf{r}^{(0)}$. Of course, this belongs to this convex set A_1 . And we define $\mathbf{q}^{(0)}$ or in general $\mathbf{q}^{(k)}$ by this relation, why, if you recall we have proved the lemma earlier and this lemma is this, what does this lemma says the maximum value of this function maximizing over all $\mathbf{q}$ is given by this where $\mathbf{q}$ is given by this expression. So, given an $\mathbf{r}$, we should find $\mathbf{q}$ in this particular fashion because this will maximize my - this function over all $\mathbf{q}$. So, this is how $\mathbf{r}$ and $\mathbf{q}$, they are related, so that is why we start off with some initial arbitrary input distribution which is strictly positive $\mathbf{r}^{(0)}$. And then we can we have to find $\mathbf{q}^{(0)}$ for that we are going to use this.

Now once you know $\mathbf{q}^{(0)}$ you need to find $\mathbf{r}^{(1)}$ and how do we find $\mathbf{r}^{(1)}$ or in general for any k greater than equal to 1 we need to find $\mathbf{r}$ at time k . So, we need to find an $\mathbf{r}$ which belongs to this convex set A_1 that maximizes the functions for a given $\mathbf{q}$. Now, remember in addition to maximizing this function we also have some constraints on $\mathbf{r}$ and what are these constraints first one is sum of probabilities should be 1, and the probability is are basically greater than greater than 0. So, these are the two constraints we have. So, we want to find the optimal value of $\mathbf{r}$ given these two constraints. So, how do we proceed?

(Refer Slide Time: 26:20)

Blahut-Arimoto Algorithm: Channel Capacity Computation

- We use the method of Lagrange multipliers to find the best r . Ignoring temporarily the positivity constraints on r , we get

$$J = \sum_x \sum_y r(x) p(y|x) \log \frac{q(x|y)}{r(x)} - \lambda \sum_x r(x)$$
- Differentiating with respect to $r(x)$, we get

$$\frac{\partial J}{\partial r(x)} = \sum_y p(y|x) \log q(x|y) - \log r(x) - 1 - \lambda$$
- Equating $\frac{\partial J}{\partial r(x)}$ to zero, we get

$$\log r(x) = \sum_y p(y|x) \log q(x|y) - 1 - \lambda$$

or

$$r(x) = e^{-(\lambda+1)} \prod_y q(x|y)^{p(y|x)}$$

So, we will take help of this method of Lagrange multiplier. So, we will define this, this is my objective function, this is my constraint related to the fact the summation of this r of x is 1. Now, at the moment, I am ignoring these positivity constraints that because later on we will see when we compute the value of r of x that are already taken care of. So, this is a Lagrangian. Now, the next step is to find up to value I differentiate it with respect to r of x . So, when I differentiate with respect to r of x I get this summation over y p of y given x $\log$ of q given y minus $\log$ of r of x minus 1 minus λ . Now, when we equate this to 0 a unique collect terms what we get is $\log$ of r x is equal to summation over y q of y given x $\log$ of q x given y minus 1 minus λ or in other words I can write r of x . In this particular way, now we need to find λ right now we know that summation of r x of x over all x that is equal to 1.

(Refer Slide Time: 28:10)

Blahut-Arimoto Algorithm: Channel Capacity Computation

- We know that $\sum r(x) = 1$, hence

$$r(x) = \frac{\prod_y q(x|y)^{p(y|x)}}{\sum_{x'} \prod_y q(x'|y)^{p(y|x')}}$$

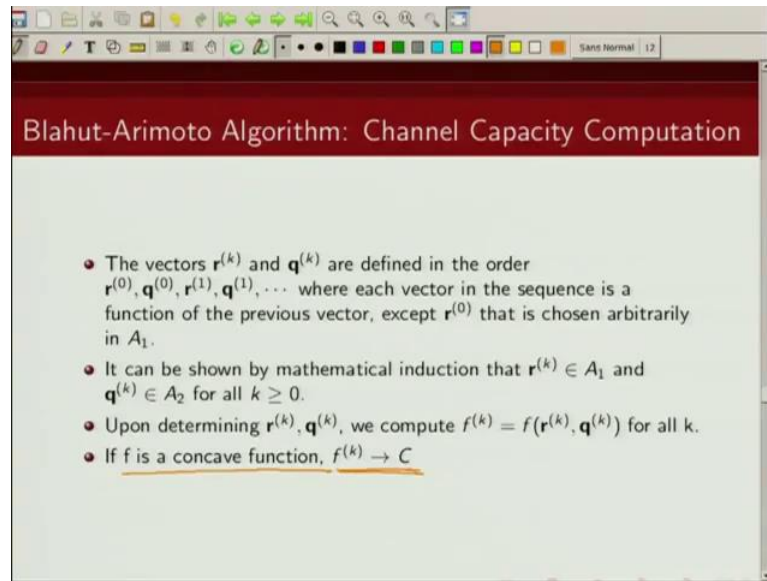
- The product is over all y such that $p(y|x) > 0$ and $q(x|y) > 0$ for all such y .
- This implies that both numerator and denominator terms on the right hand side above are positive and hence, $r(x) > 0$.
- Hence we define $r^{(k)}$ for $k \geq 1$

$$r^{(k)}(x) = \frac{\prod_y q^{(k-1)}(x|y)^{p(y|x)}}{\sum_{x'} \prod_y q^{(k-1)}(x'|y)^{p(y|x')}}$$

So, if we put in that constraint, we get the value of r of x to be this. Now, note that here we are taking product over all y here we are taking product over all y we have this term which is greater than 0 we have this term which is greater than zero. So, r of x is going to be greater than 0. So, earlier we have not taken this positivity constraints, but you can see that that r of x which comes out is constant is greater than 0, so that is implicitly taken care of. Now, So, this is the optimal value of r of x in terms of q of x .

Earlier, we have computed the optimal value of q of x in terms of r of x . So, at k th iteration basically I can write r of k x as product over all y q for k minus 1 time x given y raise to power p of y given x to address summation of this. So, note now we started with an arbitrary strictly positive distribution of r , r_0 we plug that in to get q of 0. Once, we get q of 0, it plugs the value of q of 0 here to that r of 1. Now, once we get r of 1 we are going to make use of once we know r of 1 we are going to make use of this expression to get q of 1. Now, once we get q of 1 we will make use of this to get r of 2, so that is how we are that is how we are proceeding.

(Refer Slide Time: 30:27)



Blahut-Arimoto Algorithm: Channel Capacity Computation

- The vectors $\mathbf{r}^{(k)}$ and $\mathbf{q}^{(k)}$ are defined in the order $\mathbf{r}^{(0)}, \mathbf{q}^{(0)}, \mathbf{r}^{(1)}, \mathbf{q}^{(1)}, \dots$ where each vector in the sequence is a function of the previous vector, except $\mathbf{r}^{(0)}$ that is chosen arbitrarily in A_1 .
- It can be shown by mathematical induction that $\mathbf{r}^{(k)} \in A_1$ and $\mathbf{q}^{(k)} \in A_2$ for all $k \geq 0$.
- Upon determining $\mathbf{r}^{(k)}, \mathbf{q}^{(k)}$, we compute $f^{(k)} = f(\mathbf{r}^{(k)}, \mathbf{q}^{(k)})$ for all k .
- If f is a concave function, $f^{(k)} \rightarrow C$

So, this you can see this is an iterative way we are, so these vectors $\mathbf{r}^{(k)}$ and $\mathbf{q}^{(k)}$ are defined in order we first start in arbitrary $\mathbf{r}^{(0)}$ which belongs to this convex set A_1 then we use $\mathbf{r}^{(0)}$ to compute $\mathbf{q}^{(0)}$ then we use $\mathbf{q}^{(0)}$ compute $\mathbf{r}^{(1)}$ and so on. So, you can see that each vector in this sequence is a function of previous vector except $\mathbf{r}^{(0)}$ initial values. Now, I am not showing you the proof, but it can be shown that $\mathbf{r}^{(k)}$ belongs to this convex set A_1 and $\mathbf{q}^{(k)}$ belongs to this convex set A_2 this can be proved using mathematical induction. Now, once we define and once we determine $\mathbf{r}^{(k)}$ and $\mathbf{q}^{(k)}$ we evaluate the function f at k th iteration, which is given by this. Now, the next obvious question is this function guaranteed to converge and when we will converge. So, when f is the concave function this function will converge to expression for capacity.

(Refer Slide Time: 32:00)

BA Algorithm: Convergence

- In order to show that Blahut Arimoto algorithm for computing channel capacity converges, we need to show that the function
$$f(\mathbf{r}, \mathbf{q}) = \sum_x \sum_y r(x)p(y/x) \log \frac{q(x|y)}{r(x)}$$
is concave.
- Let us consider two ordered pairs $(\mathbf{r}_1, \mathbf{q}_1)$ and $(\mathbf{r}_2, \mathbf{q}_2)$. For any $0 \leq \lambda \leq 1$, we have using log sum inequality

$$\begin{aligned} & (\lambda r_1(x) + (1-\lambda)r_2(x)) \log \frac{\lambda r_1(x) + (1-\lambda)r_2(x)}{\lambda q_1(x|y) + (1-\lambda)q_2(x|y)} \\ & \leq \lambda r_1(x) \log \frac{r_1(x)}{q_1(x|y)} + (1-\lambda)r_2(x) \log \frac{r_2(x)}{q_2(x|y)} \end{aligned}$$

So, next we are try we will show that this expression for channel capacity this function is a concave function. So, to show that this algorithm converges to channel capacity we are going to show that this function that we evaluated this is the concave function of $\mathbf{r}$ and $\mathbf{q}$. So, let us consider two ordered pair $\mathbf{r}_1, \mathbf{q}_1$ and $\mathbf{r}_2, \mathbf{q}_2$. And let λ be and there will be between 0 and 1. Now, using log sum inequality, we can write λ times $r_1(x)$ plus 1 minus λ times $r_2(x)$ log of λ times $r_1(x)$ plus 1 minus λ times $r_2(x)$ divided by λ times $q_1(x|y)$ plus 1 minus λ times $q_2(x|y)$. This is less than equal to λ times $r_1(x)$ log of $r_1(x)$ by $q_1(x|y)$ plus 1 minus λ times $r_2(x)$ log of $r_2(x)$ divided by $q_2(x|y)$. So, this follows from log sum inequality. Now we are going to take the reciprocal, so reciprocal of this log. So, what you will notice is, so what I did was I just took a reciprocal of these log terms. If you take reciprocal of this log time is less than equal to term will become greater than equal to and that is what happened here.

(Refer Slide Time: 34:04)

BA Algorithm: Convergence

- Taking the reciprocal of the logarithms, we get

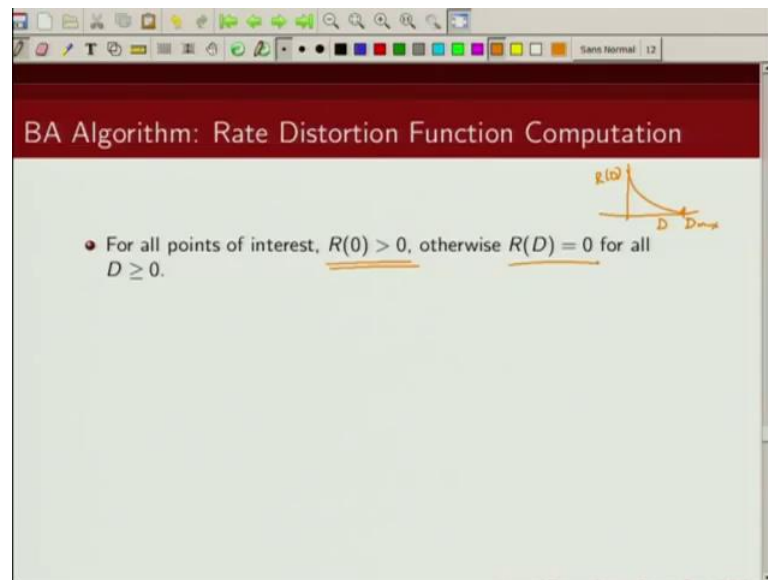
$$\begin{aligned}
 & (\lambda r_1(x) + (1 - \lambda)r_2(x)) \log \frac{\lambda q_1(x|y) + (1 - \lambda)q_2(x|y)}{\lambda r_1(x) + (1 - \lambda)r_2(x)} \\
 & \geq \lambda r_1(x) \log \frac{q_1(x|y)}{r_1(x)} + (1 - \lambda)r_2(x) \log \frac{q_2(x|y)}{r_2(x)}
 \end{aligned}$$
- Multiplying both sides by $p(y|x)$ and summing over all x and y , we get

$$f(\lambda r_1 + (1 - \lambda)r_2, \lambda q_1 + (1 - \lambda)q_2) \geq \lambda f(r_1, q_1) + (1 - \lambda)f(r_2, q_2)$$

So, I took the reciprocals of the logarithm. So, then this log of this by this is now greater than equal to lambda times $r_1(x) \log \frac{q_1(x|y)}{r_1(x)}$ plus $(1 - \lambda)r_2(x) \log \frac{q_2(x|y)}{r_2(x)}$. So, I essentially took reciprocals of this log terms and this term which was earlier less than equal to is now greater than equal to. Next, I multiply both sides by $p(y|x)$ and sum over all x and y . So, if I do that what I get on the left hand side is function evaluated at $\lambda r_1 + (1 - \lambda)r_2$ and $\lambda q_1 + (1 - \lambda)q_2$ is greater than equal to $\lambda f(r_1, q_1) + (1 - \lambda)f(r_2, q_2)$. And we know from the definition of concavity that if function evaluated at this is greater than equal to expected value of the function basically this is the condition for concavity.

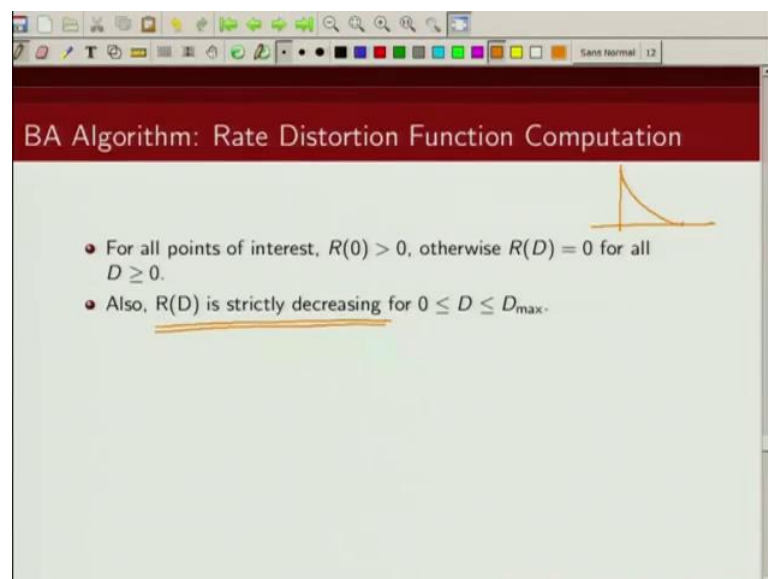
So, this shows that the function f is a concave function. So, hence we have shown that when we iteratively compute R and Q as k increases basically f evaluated at k th equation will tends towards the expression for channel capacity and this is Blahut-Arimoto algorithm for computation of channel capacity. Now Blahut-Arimoto algorithm for rate distortion function is very, very is a similar we will this quickly browse over the results.

(Refer Slide Time: 36:31)



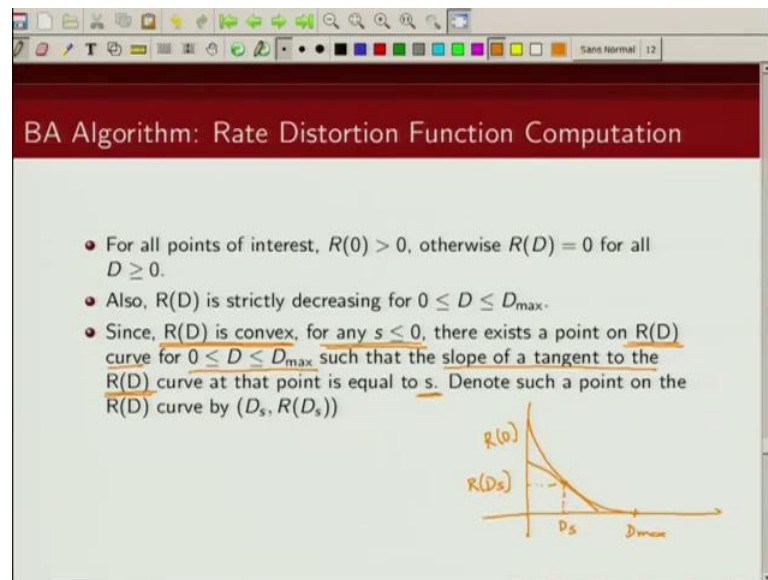
So, we know the rate distortion function looks like this typically something like this you have this is a rate this distortion. So, this is maximum distortion $D_{\max}$, $R(D) > 0$ typically greater than is equal to 0 or otherwise $R(D) = 0$ for D greater than 0.

(Refer Slide Time: 37:03)



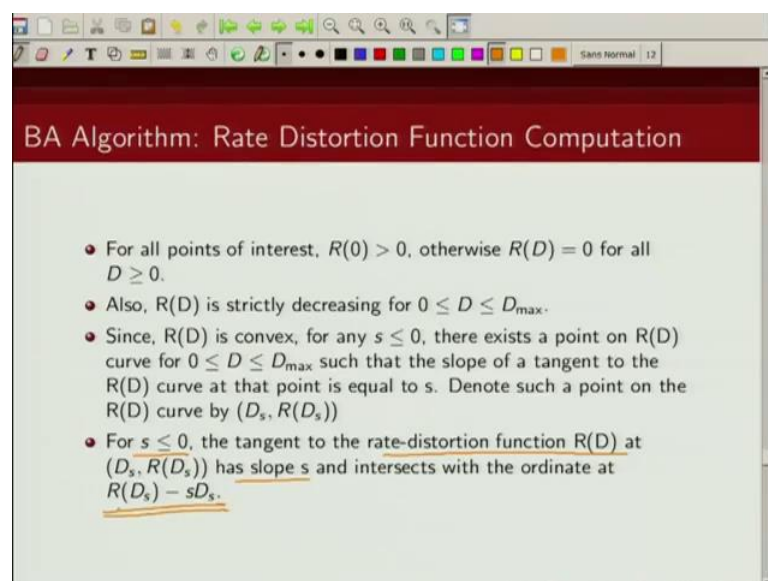
It is a typically a strictly decreasing function as we saw basically $R(D)$ typically is like this something like that.

(Refer Slide Time: 37:16)



Now, we also know that this rate distortion function is a convex function. So, then for any s which is negative there exist a point on this rate distortion curve for D between 0 to $D_{\max}$ such that the slope of the tangent to the rate distortion curve at that point is equal to this slope s . So, what I am saying is if you have a rate distortion function let us say something like this, this like $D_{\max}$. So, there exist a point and this rate distortion curve such that the slope of the tangent to this curve at that point is equal to s something s here. and let us denote this point s this is my D of s and this is my R of D of s .

(Refer Slide Time: 38:38)



So, s less than 0, the tangent to the rate distortion function have slope s , and its y intercept is given by $R(D) - sD$.

(Refer Slide Time: 38:59)

BA Algorithm: Rate Distortion Function Computation

- Let $I(\mathbf{p}, \mathbf{Q})$ denote the mutual information $I(X, \hat{X})$ and $D(\mathbf{p}, \mathbf{Q})$ denote the expected distortion $Ed(X, \hat{X})$ when $\mathbf{p}$ is the distribution for X and $\mathbf{Q}$ is the transition matrix from X to $\hat{X}$.
- Then for any $\mathbf{Q}$, $(I(\mathbf{p}, \mathbf{Q}), D(\mathbf{p}, \mathbf{Q}))$ is a point in the rate distortion region, and the line with slope s passing through $(I(\mathbf{p}, \mathbf{Q}), D(\mathbf{p}, \mathbf{Q}))$ intersects the ordinate at $I(\mathbf{p}, \mathbf{Q}) - sD(\mathbf{p}, \mathbf{Q})$.
- Since the $R(D)$ curve defines the boundary of the rate-distortion region, we see that

$$R(D_s) - sD_s = \min_{\mathbf{Q}} [I(\mathbf{p}, \mathbf{Q}) - sD(\mathbf{p}, \mathbf{Q})]$$
- For each $s \leq 0$, if we can find a $\mathbf{Q}_s$ that achieves the minimum, then the line passing through $(0, I(\mathbf{p}, \mathbf{Q}_s) - sD(\mathbf{p}, \mathbf{Q}_s))$ gives a tight lower bound on the $R(D)$ curve.
- In particular, if $(R(D_s), D_s)$ is unique, then $D_s = D(\mathbf{p}, \mathbf{Q}_s)$ and $R(D_s) = I(\mathbf{p}, \mathbf{Q}_s)$.
- By varying over all $s \leq 0$, we can trace out the whole $R(D)$ curve.

Now, let $I(\mathbf{p}, \mathbf{Q})$ denote the mutual information between X and $\hat{X}$ and $D(\mathbf{p}, \mathbf{Q})$ denotes the expected distortion, where $\mathbf{p}$ is the distribution for X and $\mathbf{Q}$ is the transition matrix from X to $\hat{X}$. For any $\mathbf{Q}$, then this is a point in this rate distortion region, and the line with slope s is going to pass through this point. So, this will give a y intercept of this. Now, since this rate distortion curve defines the boundary of the rate distortion region, we can write this as minimizing mutual information minus s times this distortion minimizing over $\mathbf{Q}$. So, for each s less than or equal to 0, we can find a $\mathbf{Q}$ of x such that this is minimized. So, we can find the $\mathbf{Q}_s$ that will achieve this minimum value. So, then a line passing through this point is going to give me a tight lower bound on the $R(D)$ curve and if I do this for all s 's, I essentially can trace out this rate distortion curve.

(Refer Slide Time: 40:53)

BA Algorithm: Rate Distortion Function Computation

- Let $p(x)Q(\hat{x}|x)$ be a given joint distribution on $X \times \hat{X}$ such that $Q > 0$, and let t be any distribution on $\hat{X}$ such that $t > 0$. Then

$$\min_{t>0} \sum_x \sum_{\hat{x}} p(x)Q(\hat{x}|x) \log \frac{Q(\hat{x}|x)}{t(\hat{x})} = \sum_x \sum_{\hat{x}} p(x)Q(\hat{x}|x) \log \frac{Q(\hat{x}|x)}{t^*(\hat{x})}$$

where $t^*(\hat{x}) = \sum_x p(x)Q(\hat{x}|x)$

- Applying the above lemma, we can write

$$R(D_s) - sD_s = \min_Q [I(\mathbf{p}, \mathbf{Q}) - sD(\mathbf{p}, \mathbf{Q})]$$

$$= \inf_{Q>0} \min_{t>0} \left[\sum_{x, \hat{x}} p(x)Q(\hat{x}|x) \log \frac{Q(\hat{x}|x)}{t(\hat{x})} - s \sum_{x, \hat{x}} p(x)Q(\hat{x}|x) d(x, \hat{x}) \right]$$

Now, let us talk about an iterative way to compute this thing. So, without proving, I am just stating the lemma that I am going to use the proof of this lemma is all on the similar lines as the proof that we did for lemma for computation of channel capacity. So, let p of x Q of x hat given x be the joint distribution on X and X hat such that Q is strictly positive and t is the distribution on x hat such that t is also strictly positive. Then minimizing this over all t is given by this expression where t star is this. Again this proof is very similar to we can use divergence inequality to prove this result, I am not stating it here again; again very similar to prove that we did for channel capacity, if now this is mutual information and this is my distortion, basically, R of D of minus s of D which is nothing, but minimizing this over all Q . Now, this mutual information is given by this and this is given by this channel. So, then I can write R of D s minus s of D s is basically this term.

(Refer Slide Time: 42:46)

BA Algorithm: Rate Distortion Function Computation

- We can now apply alternating optimization algorithm.
- For computation of rate-distortion function, we start with any strictly positive transition matrix $\mathbf{Q}^{(0)}$.
- Then we define $\mathbf{t}^{(0)}$ and in general $\mathbf{t}^{(k)}$ as

$$\underline{t^{(k)}}(\hat{x}) = \sum_x p(x) Q^{(k)}(\hat{x}|x)$$
- In order to define $\mathbf{Q}^{(1)}$, and in general $\mathbf{Q}^{(k)}$, we need to find $\mathbf{Q}$ that minimizes the function for a given $\mathbf{t}$, where the constraints on $\mathbf{Q}$ are

$Q(\hat{x}|x) > 0$

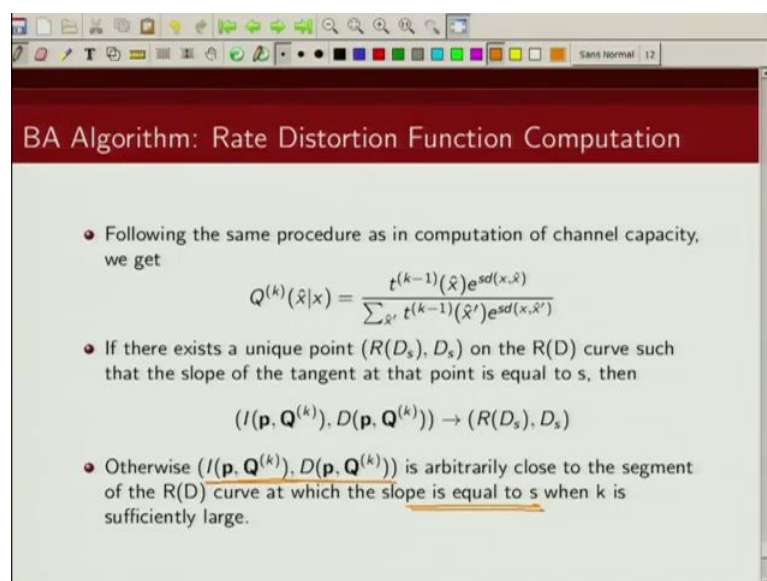
 for all $(x, \hat{x}) \in X \times \hat{X}$ and

$\sum_{\hat{x}} Q(\hat{x}|x) = 1$

 for all $x \in X$.

Now, we are going to use alternative minimization algorithm to compute this. So, we start with any strictly positive matrix $\mathbf{Q}$ of 0. The next step is we need to find $\mathbf{t}$ at time at iteration index k and this is given by this. Now, this follows from the lemma that minimum value of this minimum over all $\mathbf{t}$ is given by this, where the value of $\mathbf{t}^*$ is given by this expression. So, given a $\mathbf{Q}$, you know what is the next $\mathbf{t}$. Now, once you find $\mathbf{t}$, you need then next find $\mathbf{Q}$ of 1 at $\mathbf{Q}$ at iteration index 1. Now how do I find $\mathbf{Q}$ in general for k , again this process is very, very similar to what we did you want to maximize this with some constraints. And what are this constraints we have those positivity constraints $\mathbf{Q}$ of $\hat{x}$ given x is greater than 0 and we have this sum of probability adding up to 1. So, we will again follow the method of Lagrange. So, we will form the Lagrangian of objective function plus lambda times this constraint, we will initially ignore the positivity constraints and later on we will show that $\mathbf{Q}$ comes out to be positive, so that condition is expressively taken care.

(Refer Slide Time: 44:46)



BA Algorithm: Rate Distortion Function Computation

- Following the same procedure as in computation of channel capacity, we get
$$Q^{(k)}(\hat{x}|x) = \frac{t^{(k-1)}(\hat{x})e^{sd(x,\hat{x})}}{\sum_{\hat{x}'} t^{(k-1)}(\hat{x}')e^{sd(x,\hat{x}')}}$$
- If there exists a unique point $(R(D_s), D_s)$ on the $R(D)$ curve such that the slope of the tangent at that point is equal to s , then
$$(I(\mathbf{p}, \mathbf{Q}^{(k)}), D(\mathbf{p}, \mathbf{Q}^{(k)})) \rightarrow (R(D_s), D_s)$$
- Otherwise $(I(\mathbf{p}, \mathbf{Q}^{(k)}), D(\mathbf{p}, \mathbf{Q}^{(k)}))$ is arbitrarily close to the segment of the $R(D)$ curve at which the slope is equal to s when k is sufficiently large.

So, without showing you the details of this expression I am just directly writing the expression for Q for iteration index k and again this is exactly same procedure that we followed for when we compute computed R of at some iteration index x for computation of channel capacity. So, then if there exist a point R of D_s D_s on this rate distortion curve, such that the slope of the tangent at that point is equal to s then what we will see is as k increases, this converges to R of D_s and D_s , and if otherwise you will see that this term will be arbitrary close to a segment on this rate distortion curve, where slope is equal to s when k is sufficiently large. And the condition for convergence is so this function has to be convex function. So, I am not proving this, again very similar proof to what we did for the computation of the channel capacity.

So, to summarize, we have given to example one for the case of computation of channel capacity. We have shown how we can iteratively compute it. And similarly for the rate distortion function, as we said initially start off with some strictly positive value of Q of 0 and then we used that to compute t of 0 and then we use t of 0 to compute Q of 1, and we process that in a iterative fashion to compute the rate distortion function. So, with this, I will conclude this lecture.

Thank you.