

AQIT

2025-26

① [Axioms of QM?] Lecture 1

Def. Quantum Channel.

A completely positive, trace preserving (CPTP) map $N: L(A) \rightarrow L(B)$.

Positive: $N(\rho) \geq 0$, if $\rho \geq 0$.

Note that $A \geq 0$ means $\langle \psi | A | \psi \rangle \geq 0 \ \forall |\psi\rangle$,
and $A \geq B$ means $A - B \geq 0$.

Completely Positive: $N_A \otimes I_B(\rho_{AB}) \geq 0$ if $\rho_{AB} \geq 0$,
for all B systems.

Trace Preserving: $\text{Tr } N(\rho) = \text{Tr } \rho \ \forall \rho$.

②

Elementary Elements.

• The identity map
 $I(X) = X$.

[AKA $\text{id}(x) = x$.]

Identity matrix: $\mathbb{1}$. Identity map I or id .
 $I(\rho) = \mathbb{1} \rho \mathbb{1} = \rho$.

• Isometry $V: A \rightarrow B$.

$$V^\dagger V = \mathbb{1}_A, \quad VV^\dagger = \Pi_{B(A)}.$$

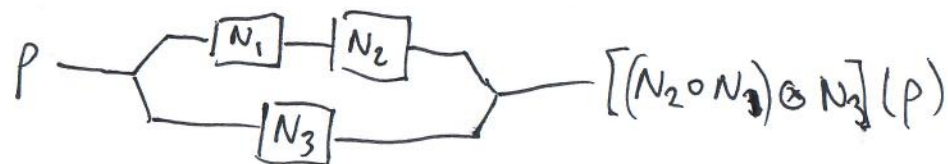
• Trace and Partial Trace.

$$\text{Tr}_A \rho_{AB} = \rho_B$$

• Compositions; Tensor Products

$N \circ M$

$N \otimes M$



• Convex Combos.

$$N = \sum_i p_i N_i.$$

Thm. The following are equivalent.

1. N is CPTP
2. N is TP and $I_{A'} \otimes N$ is positive where $N: A \rightarrow B$ and $A' \cong A$.

3. (Choi):
 $J(N) = (I_A \otimes N_{A \rightarrow B}) \Phi_{A'A}$
 is positive and $\text{tr}_B J(N) = \frac{1}{|A|} I_{A'}$

3. The Choi operator $J(N)^{AB}$ is positive $J(N) \geq 0$ and $\text{tr}_B J(N) = \frac{|A|}{|A|} I_{A'}$.

$$J(N) = I_A \otimes N_{A' \rightarrow B} \Phi_{AA'}$$

$$\Phi = |\Phi\rangle\langle\Phi| \text{ where } |\Phi\rangle = \frac{1}{\sqrt{|A|}} \sum_i |i\rangle_A \otimes |i\rangle_{A'}$$

$$\Phi = \frac{1}{|A|} \sum_{i,j} |i\rangle\langle j|_A \otimes |i\rangle\langle j|_{A'}$$

4. (Stinespring)
 $N(\rho) = \text{Tr}_E(V \rho V^\dagger)$ for some isometry V .

5. (Kraus).

$$N(\rho) = \sum_\alpha K_\alpha \rho K_\alpha^\dagger, \quad \sum_\alpha K_\alpha^\dagger K_\alpha = I.$$

PS.

~~QED~~

(Do proofs on board.)

[Cyclic proof from Andreas's lecture notes, see main lecture notes PDF.]

(5)

Examples.

- Identity ✓
- Constant/Preparation Channel

$$N_{\sigma}(p) = \text{tr}(p) \sigma$$

- Measure & Prepare (POVM $M = \{M_x\}_x$).

$$N(p) = \sum_x \text{tr}(M_x p) \sigma_x$$

- Depolarizing (Mixing w/ Noise)

$$N(p) = (1-\lambda)p + \lambda \frac{I}{d}$$

- Dephasing

$$N(p) = (1-\lambda)p + \lambda \sum_k |k\rangle\langle k| p |k\rangle\langle k|$$

- Qubit Pauli Error ✓

$$N(p) = p_0 p + p_x X p X + p_y Y p Y + p_z Z p Z$$

- Erasure channel.

$$N(p) = (1-\lambda)p + \lambda |\perp\rangle\langle\perp|$$

Classical Channels.

A classical channel is a map

$$N(p) = \sum_i \Lambda_{ji} |j\rangle\langle i| p_i,$$

where $\Lambda_{ji} \geq 0$ & $\sum_j \Lambda_{ji} = 1$.

- "Stochastic Matrix"
- Conditional Probability
- Maps prob to prob.
- Can embed classical prob. filtering as diagonal density matrices.

CQ channel: "Classical-Quantum"

$$N(p) = \sum_x \text{tr}(p |x\rangle\langle x|) \sigma_x$$

measures classical label x ($\equiv |x\rangle\langle x|$) to state p_x .

QC Channel:

$$N(p) = \sum_x \text{tr}(p |x\rangle\langle x|) |x\rangle\langle x|$$

⑥

After break: Adjoints.

IP extra time

- Classical / Comm
- EA Comm
- Quantum Comm
- 1-Shot vs. Asymptotic.



Lecture 2

Advanced Quantum Information Theory

- I promise we'll get very advanced and quantum. But not today. Today: Info Theory.
- Quantum theory is actually much older than Information Theory (1948). Even though "Quantum" sounds fancier.

Info theory was an amazing breakthrough! Before Shannon it was not clear that

- Info could be unambiguously quantified.
- Channels have a definite capacity for transmission.

Shannon (1948) showed there are fundamental, not just practical, bounds.

So the topic of today's lecture:

Shannon Entropy & Information.

① Shannon Entropy as uncertainty in a probability distribution.



High variance



Low variance

Same Entropy

Entropy doesn't care about outcome space only probabilities.



min Entropy



max Entropy

As a first try to capture uncertainty, consider

$$\langle p_i \rangle = \sum_i p_i^2$$

If $p_i = \delta_{ij}$ then $\langle p_i \rangle = 1$. [large expected prob.]

If $p_i = \frac{1}{N}$ then $\langle p_i \rangle = \frac{1}{N}$. [small expected prob.]

This is already a good uncertainty measure. ③

Why

$$H(p) = -\sum_i p_i \log p_i ?$$

We will see 3 reasons...

↳ Axiomatic Derivation.

↳ How many bits needed to represent outcomes?

↳ Asymptotic message rate.

By the way, $S(p) = -\log p = H(\text{e-val})$.

④ Axiomatic.

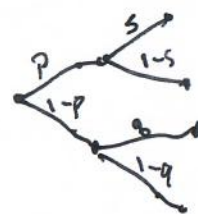
Shannon derives $H(p)$ from 3 axioms.

1. Monotonicity

$$H(\{ \frac{1}{n}, \dots, \frac{1}{n} \}) = A(n) \text{ monotone of } n.$$

2. Continuity in p .

3. Chain Rule



$$= H(\{p, 1-p\}) + p H(\{s, 1-s\}) + (1-p) H(\{q, 1-q\}).$$

[Thm. (Shannon) These axioms enforce $H(p) = -k \sum_i p_i \log p_i$.

Pf. Shannon's Paper.

(See also Renyi — giving up chain rule.)

This is nice, but the practical reasons are more powerful.

② How Many Bits?

- Shannon Entropy $H(p) = -\sum_i p_i \log p_i$
 " # bits to encode outcomes drawn from p with optimal encoding "
- Cross Entropy $H(p; q) = -\sum_i p_i \log q_i$
 " # of bits to encode from p when optimized for q "
- Relative Entropy $D(p||q) = H(p; q) - H(p)$
 "The cost of non-optimal encoding"
 # of extra bits.

Goal: Write down English text efficiently as binary strings.

26 Letters a, b, ..., x, y, z

Simplest Strategy — label using binary numbers $0 \rightarrow 26$ with $\lceil \log_2 26 \rceil = 5$ bit strings.

a → 00000	y → 10001	[z, a] → 11011
b → 00001	z → 11010	• → 11100
⋮	⋮	⋮

③

a → 00000
 letters → ~~code strings~~ codewords

But this is inefficient!

- Uses: 5 bits per letter.

Why?

- Some letters (a, e, t) are common, others (q, z) are rare.

Strategy... use short strings for common letters.

a → 001

q → 11101011

Let's be more precise.

Consider a league of 5 teams. (7)
 You will write down the winner each year, in binary, to store records.

1	1	1	1	1
A	B	C	D	E
.1	.4	.3	.1	.1

$H(p) \approx 2.05$ bits
 ← Win probabilities.

The naive strategy is to use $\lceil \log_2 5 \rceil = 3$ bits per letter.

Here's a better strategy

A → 00
 B → 01
 C → 10
 D → 110
 E → 111

This uses, on average,
 $\langle L \rangle = \sum p_i L_i = .8 \times 2 + .2 \times 3$
 $= 2.2$ bits/letter

That's nearly optimal.

Let's prove optimal length.

(8)

Lemma. (Kraft-McMillan). For any uniquely decodable code, with codeword lengths L_i , the lengths obey

$$\sum_i 2^{-L_i} \leq 1.$$

Pf. C.I.T. Thms 5.2.1, 5.5.1.

Thm. The expected codeword length
 $L = \langle L_i \rangle = \sum p_i L_i$
 is $\geq H(p)$, the Shannon Entropy.

Pf. $L = \sum p_i L_i = \sum p_i (-\log 2^{-L_i})$

$$= - \sum p_i \log \frac{2^{L_i}}{\sum_j 2^{L_j}} = - \sum p_i \log \frac{2^{L_i}}{\underbrace{\sum_j 2^{L_j} \leq 1}}$$

$$\geq - \sum p_i \log r_i \quad \leftarrow \text{Reference Prob Distribution}$$

$$= H(p) + D(p||r) \quad \leftarrow \text{Penalty}$$

$$\geq H(p).$$

Thm. The optimal codeword length is between

$$H(p) \leq L \leq H(p) + 1$$

Pf. C.T.

Now what if we look at long strings?

The source is

$$p^n(x^n) = \prod_k p(x_k),$$

with entropy

$$H(p^n) = nH(p).$$

If we do "block encoding", the # of bits per letter is

$$\frac{1}{n} H(p^n) \leq L \leq \frac{1}{n} (H(p^n) + 1)$$

$$H(p) \leq L \leq H(p) + \frac{1}{n},$$

with optimal scheme.

Using long blocks, expected length is

$$L \approx H(p)$$

bits per letter.

(9) (10)

(C) Asymptotic # of messages.

• Consider strings of n messages drawn from $p(x)$.

$$\hat{p}(x^n) = p(x_1) p(x_2) \dots p(x_n).$$

This is "IID".

• Law of large numbers says

$$\frac{1}{n} \sum_k x_k \rightarrow \langle x \rangle \text{ as } n \rightarrow \infty,$$

for any IID x_k .

$$(\text{Prob}(\text{diff} \geq \epsilon) \rightarrow 0).$$

• What is probability of a given string?

$$-\frac{1}{n} \log p^n(x^n) = \frac{1}{n} \sum_k (-\log p(x_k))$$

$$\xrightarrow{n \rightarrow \infty} \langle -\log p(x) \rangle = H(p).$$

The prob that a string obeys

$$-\log p^n(x^n) \approx_{\epsilon} nH(p)$$

goes to one.

Summary. Repeatedly drawing letters from $p(x)$ produces

$$2^{nH(p)}$$

equally likely "typical" messages, each with probability

$$2^{-nH(p)}$$

Any other message has negligible probability.

Pf. C&T (Asymptotic Equipartition).

Add ϵ^n , δ^n to make precise.

The Typical Set.

- Typical messages have to have the "right number" of each letter

- In English "A" is likely and "Z" not, but AAAA & ZZZZ are both atypical (highly unlikely).

- Typical: ALHENTTPA OOBTTVA

(11) (12)

So...

Having n copies of a source $p(x)$ is like having a uniform source of $2^{nH(p)}$ different messages.

An aside...

Back to writing down unknown term in binary.

- Suppose you initially know nothing. You simply have to guess all equally likely. $\log N$ bits per outcome.

- Now you're given distribution $p(x)$ (the unprobabilities).

This saves you

$$D(p||u) = \log N - H(p)$$

bits per letter.

- "Distraction is: worth $D(p||u)$ bits per letter."

- Now you get the actual outcome of the league, already knowing p_i . This gives you

$$H(p)$$

no more bits.

"Outcome worth $H(p)$ bits".

- What if you had wrong guess at p , but got to update it?

13

14

- ⑤ If time remaining, Review extra:

- Classical Channels
- Relative Entropy Props

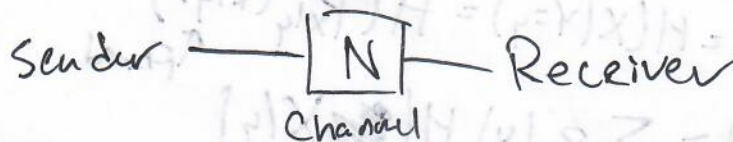
In prep for next lecture.

④ How Much Information is in a book?

- Depends on decoding mechanism.
- Go to projector? Show Shannon papers.

Lecture 3

① Intro



Goal: Create mutual information between Sender and receiver.

② Prob. Theory Notation

$P_X(x)$ ← Particular value
← Distribution
 X = Random Variable
= Function of outcome of $P(x)$.

Joint Distribution

$$P_{XY}(x, y)$$

Marginal Distribution

$$P_X(x) = \sum_y P_{XY}(x, y)$$

Conditional

$$P_X(x) P_{X|Y}(x|y) = P_Y(y) P_{Y|X}(y|x) = P_{XY}(x, y)$$

I will sometimes just write

P_{xy} , $P_{x|y}$, P_x , etc, for shorthand.

$$H(X) = H(P_X)$$

$$H(X|y) = H(X|Y=y) = H(P_{X|y}(x|y))$$

$$H(X|Y) = \sum P_Y(y) H(P_{X|y}(x|y))$$

$$H(X, Y) = H(P_{X,Y})$$

③ Conditional Probability.

A conditional prob dist

$$P_{X|Y}(x|y)$$

is a P.D. over x for each y .

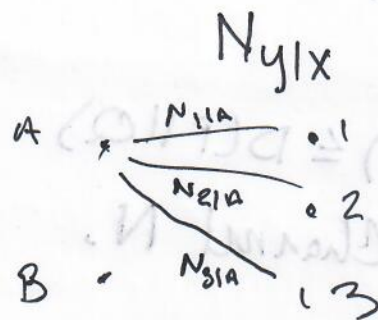
"Prob of x , given y ."

Thus $P(x|y) \geq 0 \forall x, y$.

$$\sum_x P(x|y) = 1, \forall y.$$

Also known as a "Stochastic Matrix"
or "Stochastic Channel".

④ A Classical Channel is defined by a conditional probability distribution from inputs to outputs.



A map from P.D.'s to P.D.'s:

$$g = N(p)$$

means

$$g_y = \sum_x N_{y|x} p_x.$$

It is a positive ^{sum} preserving map.

⑤ Markov Chain

$$X \rightarrow Y \rightarrow Z$$

is a Markov chain if

$$P(x, y, z) = P(x) P(y|x) P(z|y)$$

Fact. This implies also $Z \rightarrow Y \rightarrow X$.

Pf. C&T.

⑥ RE & Monotonicity

Relative Entropy:

$$D(P||Q) = \sum_x P(x) \log \frac{P(x)}{Q(x)}$$

"How different ARE P & Q ?"

Thm. (Monotonicity)

$$D(N(P)||N(Q)) \leq D(P||Q)$$

For Any classical Channel N .

⑦ Mutual Information

Joint distribution $P_{XY}(x,y)$. $(p)1 = p$

$$I(X:Y) \equiv D(P_{XY}||P_X P_Y)$$

Fact.

$$I(X:Y) = H(X) + H(Y) - H(X,Y)$$

$$= H(X) - H(X|Y)$$

$$= H(Y) - H(Y|X)$$

"How different is P_{XY} from its marginals"

"How much info do X, Y tell about each other?"

⑧ Data Processing Inequality

⑤

Thm. If $X \rightarrow Y \rightarrow Z$ is a Markov chain,
then $I(X:Z) \leq I(X:Y)$.

Cor. Also, $I(X:Z) \leq I(Y:Z)$.

~~The Cor. holds b/c also $Z \rightarrow Y \rightarrow X$.~~

~~Pf. Consider the channel $N_{Y \rightarrow Z}$~~
 ~~$N(y) = \sum_z N_{zy}$~~

Pf. Consider channel $N_{Y \rightarrow Z}$ defined by $N_{zy} = P_{zy}$ from the chain.

$$P_{xyz} = P_x P_{y|x} P_{z|y}$$

$$P_{xz} = \sum_y P_{xyz} = \sum_y P_{z|y} (P_x P_{y|x})$$

$$P_{xy} = \sum_z P_{xyz} = P_x P_{y|x}$$

$$N(P_{xy}) = \sum_y P_{z|y} P_x P_{y|x} = \sum_y P_{xyz} = P_{xz}$$

Similarly, $N(P_x P_y) = P_x P_z$.

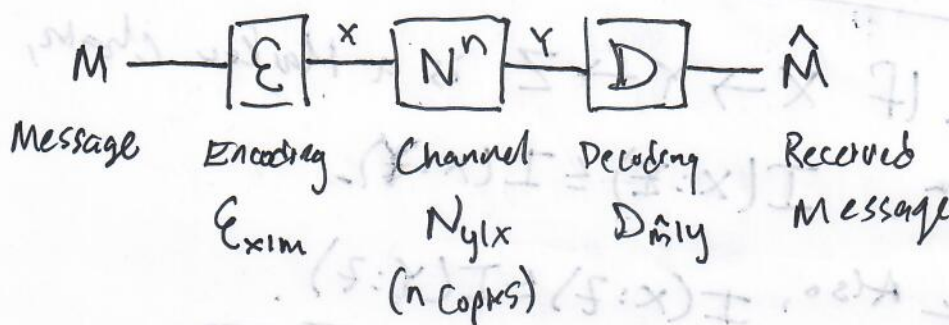
Thus $I(X:Z) = D(N(P_{xy}) || N(P_x P_y))$
 $\leq D(P_{xy} || P_x P_y) = I(X:Y)$

Cor. b/c $Z \rightarrow Y \rightarrow X$

* RE. Monotonicity

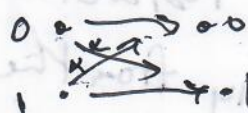
① The Communication Problem.

⑥

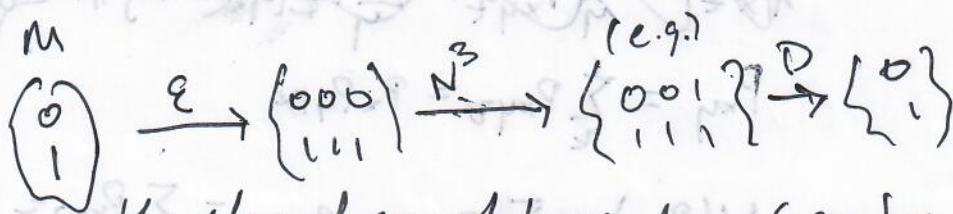


Question. At what Rate can we send messages Reliably using the channel N .

Example. Consider Bit Channel ~~perfect~~ with errors:



No single bit can be sent reliably in a single shot! But...



Using the channel several times, I can correct errors. Thus is a rate of

$$\frac{\log |M|}{n} = \frac{\log 2}{3} = \frac{1}{3} \text{ bits/signal}$$

Best possible Rate with negligible errors?

② Channel Capacity

⑦

Def. The capacity of channel N is the maximum amount of mutual information between the channel input & channel output,

$$C_N = \sup_{P_X} I(X; N(X)).$$

② Channel Capacity.

⑧

- Consider the mutual information between sent message and received message

$$I(M; \hat{M}) = \underbrace{H(M)}_{\text{Input Rate}} - \underbrace{H(M|\hat{M})}_{\text{Equivocation}}.$$

- The equivocation relates to Probability of error via Fano's Inequality:

$$H(M|\hat{M}) \leq 1 + P_e \log |M|,$$

$$\text{where } P_e = \sum_m P_m P_e(m) = \sum_m P_m (1 - P(\hat{m}=m)).$$

~~where $P_e = \sum_m P_m P_e(m)$~~

- By D.P.I., since $M \rightarrow X^n \rightarrow Y^n \rightarrow \hat{M}$, clearly

$$\underbrace{I(M; \hat{M})}_{\text{Depends on encoding \& decoding}} \leq \underbrace{I(X^n; Y^n)}_{\text{Depends on encoding \& decoding}}.$$

- But for any encoding/decoding,

$$I(X^n; Y^n) \leq n C_N,$$

where $C_N = \sup_{P_X} I(X; N(X))$ across a single use of the channel. This is the channel capacity.

- So for any $\epsilon(D)$,

$$nC_N \geq I(M; \hat{M}) = H(M) - H(M|\hat{M}).$$

- Suppose $P_M(m)$ is uniform over $|M| = 2^{nR}$ messages [Rate of R ^{message} bits per use of channel].

The worst case probability of error

$$P_e^* = \max_m P_e(m) \geq P_e^{\text{uniform}} = \sum_m \frac{1}{2^{nR}} P_e(m).$$

We want an encoding ensuring all messages get through successfully, so P_e^* is small.

- Combining above with Fano, using uniform input

$$nC_N \geq H(M) - H(\hat{M}|\hat{M})$$

~~$$\geq \log |M| - (1 + P_e^{\text{uniform}}) \log |M|$$~~

$$\geq \log |M| - (1 + P_e^{\text{uniform}}) \log |M|$$

$$\geq nR - 1 - P_e^{\text{uniform}} nR$$

$$\geq nR - 1 - P_e^* nR$$

or ...

$$nR P_e^* \geq nR - nC - 1$$

$$\boxed{P_e^* \geq 1 - \frac{C}{R} - \frac{1}{n}}$$

*For $R > C$,
Small P_e^* becomes
impossible!

Thm. (Shannon Noisy Channel Coding)

(Achievability)

- For any rate $R < C$, it is possible to transmit across the channel with negligible probability of error.

(Converse)

- For any rate $R > C$, ~~it is not possible to~~ the probability of errors is at least $P_e^* \geq 1 - \frac{C}{R} - \epsilon$, $\forall \epsilon > 0$.

- (Strong Converse). In fact, if $R > C$ the probability that the whole message is correct $\rightarrow 0$.

Pf. We already proved converse, we won't prove strong converse. But suppose you transmit $R_0 = C$ successfully and throw out $R_1 = R - C$. Then

$$P_{\text{succ}} = \frac{2^{nC}}{2^{nR}} = \frac{1}{2^{n(R-C)}}$$

which $\rightarrow 0$ exponentially with n .

The hard part is achievability.

for $R > C$
small is bound
(impossible)

$$\frac{1}{2^n} - \frac{C}{R} - 1 \leq \frac{1}{2^n}$$

③ Achievability

②

• The capacity is defined as

$$C_N = \sup_{P_X} I(X; N(X))$$

which is really

$$\sup_{P_X} D(P_X N_{Y|X} \| P_X \sum_X N_{Y|X} P_X)$$

$$= \sup_{P_X} D(P_X N_{Y|X} \| P_X Q_Y), Q = N(P)$$

Joint Dist
our input
& output

~~• Need to find~~
~~• Strategy: Define a code that achieves capacity.~~

~~• Strategy: Random Encoding~~

~~Jointly Typical Decoding.~~

• Need to find a code that achieves the capacity.

• Strategy :-

• Random Coding

• Jointly Typical Decoding

• Joint Typicality

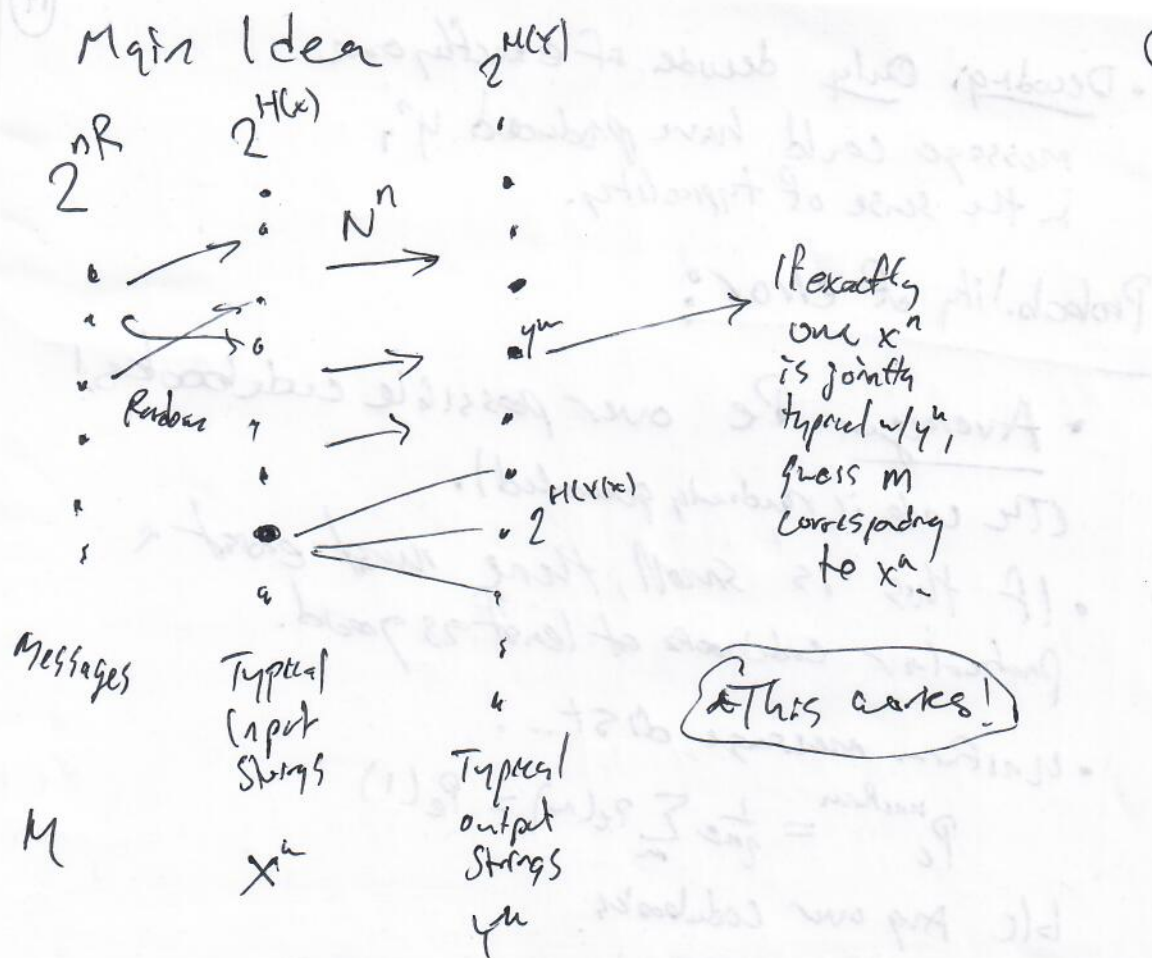
— x^n, y^n are jointly typical if x^n typical of $p(x)$, y^n typical of $p(y)$, and x^n, y^n typical of $p(x, y)$.

— Drawing IID from $p(x, y)$, almost always gives jointly typical strings.

— The probability that two independent marginally typical x^n, y^n are jointly typical is $2^{-nH(x, y)}$.

• We need to find a code that achieves the capacity.

~~Strategy: Random Coding~~
~~Jointly typical Decoding.~~



Encoding: For each m , generate codeword $C(m)$ from $P(x)$, IID, where $P(x) \ni$ capacity-achieving input rate.

Transmission: Codeword x^n translates to y^n with probability $P(y^n|x^n) = P(y_1|x_1) \dots P(y_n|x_n)$.
So together x^n, y^n drawn from $P(x^n, y^n) = P(x^n)P(y^n|x^n)$.

- Decoding: Only decide if exactly one message could have produced y^n , in the sense of typicality.

Probability of error:

- Average P_e over possible codebooks! (The code is randomly generated).
- If this is small, there must exist a particular codebook at least as good.
- Uniform message dist...

$$P_e^{\text{uniform}} = \frac{1}{q^n} \sum_m P_e(m) = P_e(1)$$

b/c Avg over codebooks.

- Calculate $P_e(1)$.

$$P_{\text{decode}} = P_{x^n(m) \text{ is atypically } w/y^n} - P_{\text{Another } x^n \text{ is too}}$$

$$\approx 1 - 2^{n(R-1)} \cdot 2^{-nI} - \epsilon$$

$$\approx 1 - 2^{-n(C-R)}$$

$$P_e^{\text{uniform}} \leq 2^{-n(C-R)}$$

which $\xrightarrow{n \rightarrow \infty} 0$ if $R < C$.

- Expurgate half! $P_e^E \leq 2 P_e^{\text{uniform}}$, $R' = R - \frac{1}{n}$

99 Relative Entropy Review

15

• Classical RE. (KL Divergence).

$$D(p \parallel q) = \sum_i p_i \log \frac{p_i}{q_i}.$$

Key Props.

① Non-Negativity

$$D(p \parallel q) \geq 0,$$

with equality iff $p = q$.

② Joint Convexity

$$D(\sum_k \lambda_k p_k \parallel \sum_k \lambda_k q_k) \leq \sum_k \lambda_k D(p_k \parallel q_k)$$

③ Monotonicity

$$D(N(p) \parallel N(q)) \leq D(p \parallel q)$$

for any classical channel N .

④ Chain Rule

$$D(p_{xy} \parallel q_{xy}) = D(p_x \parallel q_x) + \sum_x p_x D(p_{y|x} \parallel q_{y|x}).$$

• Quantum RE

$$D(\rho \parallel \sigma) = \text{tr}(\rho \log \rho - \rho \log \sigma)$$

Key Props.

① Non-negativity

$$D(\rho \parallel \sigma) \geq 0$$

Equality iff $\rho = \sigma$.

② Joint Convexity

$$D(\sum_k \alpha_k \rho_k \parallel \sum_k \alpha_k \sigma_k) \leq \sum_k \alpha_k D(\rho_k \parallel \sigma_k)$$

③ Monotonicity

$$D(N(\rho) \parallel N(\sigma)) \leq D(\rho \parallel \sigma)$$

For any CPTP map N .

No chain Rule.

These Props are your new best friend.

Other divergences: Rényi entropy.

(17)

Notice that

$$D(p||q) = \langle -\log \left(\frac{p}{q} \right) \rangle_p$$

You could similarly try

$$-\log \left\langle \left(\frac{p}{q} \right)^s \right\rangle_p.$$

More generally what about

$$-\log \left\langle \left(\frac{p}{q} \right)^s \right\rangle_p^{1/s}?$$

This is the Rényi α -Divergence

$$D_\alpha(p||q) = -\log \left\langle \left(\frac{p}{q} \right)^s \right\rangle_p^{1/s}$$

where $\alpha = s-1$.

• Everything is Relative entropies.

~~Entropy~~

$$H(p) = -D(p \parallel \mathbb{Q})$$

$$I(A:B)_p = D(p_{AB} \parallel p_A \otimes p_B)$$

$$H(A|B)_p = -D(p_{AB} \parallel p_A \otimes \mathbb{Q}_B)$$

And so on...

• Later we'll also see Generalized Divergences.
which have similar properties.

Lecture 4

①

Goals; Understand:

- Definition of one-shot comm. theory (plain and E.A.).
- ~~one-shot~~ Rate, Error Probability, codes.
- Basic Properties; Examples of one-shot rates.
- Superdense Coding
- Propagating info through the diagram.

① One Shot Communication

- You are sending messages using the channel N only once! (Not Asymptotic $N \rightarrow \infty$).
- You want to send one from a set of K messages.

$$\text{Rate: } R = \log K$$

(K messages $\leftrightarrow \log K$ ^{messages} bits).

- Is it possible?



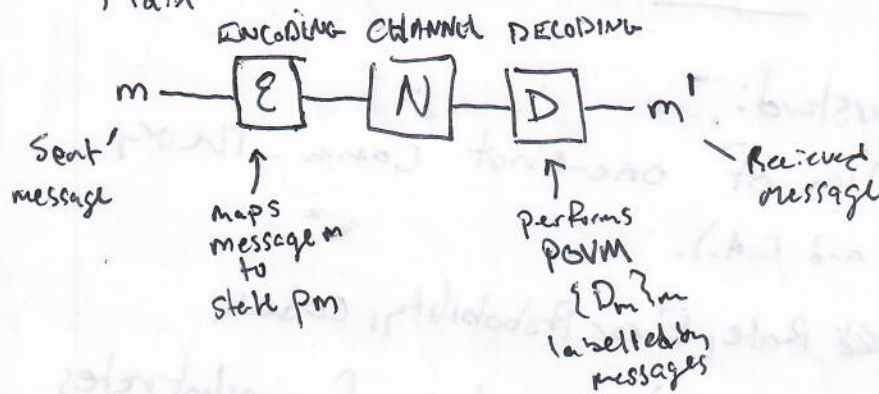
No! At least, not without error.
(unlike $N \rightarrow \infty$).

- The question: What rate R can you send messages at error level P_e ?

• Two Protocols: (like Asymptotic case)

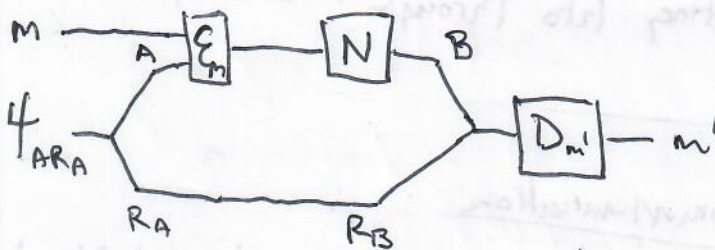
②

Plain



Entanglement Assisted

Now E_m is modulation.



- Different ways to define E, D are around, but you can always just treat both as channels (E is $CQ \rightarrow Q$, D is QC).

We'll treat both in detail.

② Plain Classical Comm.

③



~~The encoder E maps the message space to the~~

- The encoder E maps messages to quantum states,
 $m \xrightarrow{E} \rho_m = E(m).$

You can think of this as a ^{Quantum} Channel, where the initial classical message is

$$m \rightarrow |m\rangle\langle m|.$$

But often defined as map from alphabet to states.

- The decoder is a POVM but you can think of it as a ^{ac} channel

$$D(p) = \sum_i \text{tr}(D_i p) |m'_i\rangle\langle m'_i|.$$

- The probability to receive m' given m is

$$P(m'|m) = \text{tr}(D_{m'} N(\rho_m)).$$

~~Given uniform input distribution $P(m) = \frac{1}{K}$, the success/failure probability is~~

- Given uniform input dist, $P_m = \frac{1}{K}$, the Probability of error (uniform, average) is

$$P_e = 1 - \frac{1}{K} \sum_m \text{tr}(D_m^* N(p_m))$$

- For a given input size K (rate $R = \log K$), the optimal error probability is

$$P_e(R) = \inf_{(E, D)} P_e$$

with message
size $K = 2^R$

More messages, higher error.

- We're interested in the complement of this, $C_e(N)$.

Highest rate possible given $P_e \leq \epsilon$.

This is the ϵ -one-shot capacity.

- Now formal definitions.

Def. A code for N is a pair $(\mathcal{E}, D)$.
 The rate of the code is $R = \log K$,
 with K size of message space. The
 (average, uniform) error probability is

$$P_e = 1 - \frac{1}{K} \sum_m \text{tr}(D_m N(p_m)).$$

Def. The ϵ -one-shot capacity of N is

$$C_\epsilon(N) = \sup_{\substack{(\mathcal{E}, D) \\ \text{s.t.} \\ P_e \leq \epsilon}} R(\mathcal{E}, D).$$

Largest Rate For fixed error probability.

• What is $C_0(N)$ for B.S.C.?

Examples.

• Identity channel I_A .

Let $K = d_A$. Let $\mathcal{E}(m) = |m\rangle\langle m| \equiv p_m$.

• Let $D_m = |m\rangle\langle m|$. $\rightarrow P_e = 0$. Thus

$$C_\epsilon(I_A) \geq C_0(I_A) \geq \log d_A.$$

On the other hand, suppose K messages and error $P_e \leq \epsilon$. Then for any such $(\mathcal{C}, D)$

$$\frac{1}{K} \sum_{m=1}^K \mathbb{E} \text{tr}(D_m p_m) \leq \frac{1}{K} \sum_{m=1}^K \mathbb{E} \text{tr}(D_m) = \frac{1}{K} \text{tr} \mathbb{I}_n = \frac{d_A}{K}$$

VI
 $1 - \epsilon$.

Thus $\frac{d_A}{K} \geq 1 - \epsilon$, so $R \leq \log d_A - \log(1 - \epsilon)$

So

$$\boxed{\log d_A \leq C_\epsilon(I_A) \leq \log d_A - \log(1 - \epsilon)}$$

• Constant Channel P_0 .

Intuitively, transmit nothing $\rightarrow$ is/are independent
But even broken clock is right 2x/day.

for any $(\mathcal{C}, D)$,

$$P_e = 1 - \frac{1}{K} \sum_{m=1}^K \text{tr}(D_m \sigma) = 1 - \frac{1}{K}$$

Thus ~~$R \leq \log \frac{1}{1 - \epsilon}$~~

$$R \leq \log \frac{1}{1 - \epsilon}$$

A code exists, such that whenever $R \leq \log \frac{1}{1 - \epsilon}$.

$$\boxed{C_\epsilon(P_0) = \log \left\lfloor \frac{1}{1 - \epsilon} \right\rfloor}$$

These examples are outliers... usually difficult^② to evaluate.

- Now an important "Data Processing" type result.

Prop. One shot capacity is monotonic under pre- and post-processing, and also under increasing ϵ . Namely, ~~$C_\epsilon(B \circ N \circ A) \leq C_\epsilon(N) \leq C_{\epsilon' \geq \epsilon}(N)$~~

$$C_\epsilon(B \circ N \circ A) \leq C_\epsilon(N) \leq C_{\epsilon' \geq \epsilon}(N).$$

Proof. Convert code for $B \circ N \circ A$ into a code for N with same Rate/ P_e .

Let $(\mathcal{E}, D)$ be code for $B \circ N \circ A$.

$$\mathcal{E}(m) = p_m, \quad D = \{D_m\}.$$

Then let $(\mathcal{E}', D')$ be a code for N .

$$\mathcal{E}'(m) = A(p_m) \leftarrow \text{Encoding map} \checkmark$$

$$D'_m = B^+(D_m) \leftarrow \text{postmap} \checkmark$$

But

$$\text{tr}[B^+(D_m) N(A(p_m))] = \text{tr}[D_m (B \circ N \circ A)(p_m)]$$

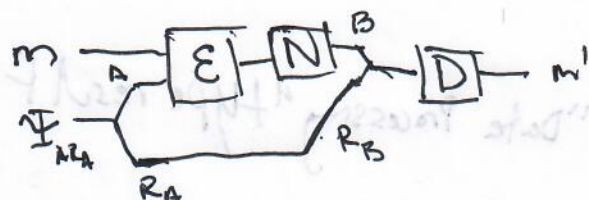
so $P'_e = P_e$. Thus $(\mathcal{E}', D')$ is an (R, ϵ)

code for N if $(\mathcal{E}, D)$ is an (R, ϵ) code for $B \circ N \circ A$.

Absorb into sup.

③ Entanglement Assisted.

⑧



- Alice can prepare arbitrary Ψ_{RA} and pass the reservoir part to Bob via $I_{R_A \rightarrow R_B}$. In other words Alice and Bob share state Ψ_{ARB} .

- The encoder is now a channel E_m that depends on message m . That is $\{A \rightarrow B\}$ is a channel for each m . Aka "modulation".

↳ Alternately, $E(|m\rangle\langle m| \otimes \rho_A) \rightarrow \rho_B$ is a channel $E_{MA \rightarrow B}$.

- This is clearly ~~at least as good as~~ ~~as good as~~ plain — it catches it!

Suppose $\Psi = \psi_A \otimes \psi_{RA}$ is a product state, and

$E_m = P_m$, and $D_m = D_m^B \otimes \mathbb{I}_{R_B}$. Then

$$P_{m|m} = \text{tr}(D_m^B N(p_m))$$

equivalent to an arbitrary plain code.

- Is it better? Super dense coding proves yes.

• Super Dense Coding

$N =$ qubit Identity channel I_2

$\Psi_{ABA} = \Phi^+ =$ Bell Pair $(|\Phi^+\rangle = \frac{1}{\sqrt{2}}(|00\rangle + |11\rangle))$.

$K = 4$ messages ($R = 2$)

$\mathcal{E}_m = \{I, X, Y, Z\}$

$D_m =$ Bell Basis Measurement

$P_e = 0$.

Thus, $C_e^{\text{EA}}(I_2) \geq C_o^{\text{EA}}(I_2) \geq 2$.

Def. ~~EA code for N is (E, D, \Phi)~~

~~EA~~ An EA code for N is $(\mathcal{E}, D, \Phi)$. The rate is still $R = \log K$. Now

$$P_{m|m} = \text{tr} [D_m ((N \circ \mathcal{E}_m) \circ I)(\Phi)],$$

$$P_e = 1 - \frac{1}{K} \sum_m P_{m|m}.$$

Def. The ϵ -one-shot EA capacity is

$$C_e^{\text{EA}}(N) = \sup_{\substack{(R, \epsilon) \text{ EA} \\ \text{codes}}} R.$$

Prop. ϵ -one-shot EA capacity is monotonic under Pre-/post-processing, and under merging: PP - same as plain case.

④ Aside...

Any questions about HW?

⑤ Overview of coming results.

Asymptotic

Plan Holevo $\chi(N) = \sup_{\substack{P_{XA} \\ \in \text{CCQ}}} I(X:B)_\omega$ [Regularize]

EA Mutual $I(N) = \sup_{P_{RA}} I(R:B)_\omega$

$\omega = N_{A \rightarrow B}(p)$

One-Shot [Not equal but near]

Plan $\chi_h^\epsilon(N)$

EA I_h^ϵ

Change to ~~CCQ~~

ϵ -Hypothesis Testing

Relative Entropy.

• Converse: Use DPI to remove D, then get M.T.

⑥ If necessary...

④

- Square Root Measurement
 - Hayashi/Nagaoka
 - $D_h^t(p||\sigma)$
-

Lecture 5

Today:

Goal: Prove Achievability bounds for plan. EA class. comm.

→ Define a code, check its (R, ϵ) .

① Preliminaries

- Hyp Test RE
- SQR Measurement
- H-N Lemma

② Plan: Random Coding [Wang/Renner 2012]

③ EA: Position-based Coding [Anshu et al 2019]

①

① Preliminaries

② Hypothesis Testing Relative Entropy

- Hypothesis Testing Problem:

Decide whether unknown state is σ or ρ .

- To decide, do a poven $\{M_\rho, \mathbb{I} - M_\rho\}$.
If result is 0, guess ρ . If 1, guess σ .

$$P_{\text{succ}} = \text{tr}(M_\rho \rho) \quad P_{\text{fail}} = \text{tr}[(\mathbb{I} - M_\rho) \sigma]$$

$$P_{\text{fail}} = \text{tr}[(\mathbb{I} - M_\rho) \rho] \quad P_{\text{succ}} = \text{tr}[M_\rho \sigma]$$

Type I Error Type II Error

- Given fixed bound on type I error,

$$\text{tr}(M_\rho) \geq 1 - \epsilon$$

Minimize the type II error:

$$\beta_\epsilon = \min_{M_\rho} \text{tr}(M_\sigma) \quad \text{tr}(M_\rho) \geq 1 - \epsilon$$

If small, very distinguishable.

If big, not distinguishable.

log of this
is like a
relative entropy

Def. ϵ -Hypothesis Testing Relative entropy

(3)

For two states ρ, σ on same system, $0 \leq \epsilon \leq 1$,

$$D_h^\epsilon(\rho || \sigma) = -\log \left(\min_{\substack{0 \leq M \leq I \\ \text{tr} M \rho \geq 1-\epsilon}} \text{tr}(M \sigma) \right).$$

- Big = Distinguishable. Small = Indistinguishable.
(Just like RE). [Recall $D(\rho || \sigma) = \text{tr}(\rho \log \rho - \rho \log \sigma)$]

Prop.

1. For any ρ, σ ,

$$D_h^\epsilon(\rho || \sigma) \geq D_h^\epsilon(\rho || \rho) = -\log(1-\epsilon) \geq 0.$$

2. For any ρ, σ , and any channel N ,

$$D_h^\epsilon(\rho || \sigma) \geq D_h^\epsilon(N(\rho) || N(\sigma)).$$

3. For $\tau = \frac{1}{d}$, $\epsilon = 1-\epsilon > 0$,

$$D_h^\epsilon(\tau || \tau) = \log d - \log(1-\epsilon).$$

Note: compare to

$$D(\rho || \sigma) \geq 0$$

$$D(\rho || \sigma) \geq D(N(\rho) || N(\sigma))$$

$$D(\tau || \tau) = \log d.$$

Proofs.

(2). Let M be optimal for $D_h^\epsilon(N(\rho) || N(\sigma))$.

Then $\text{tr}(M N(\rho)) \geq 1-\epsilon$, ~~$\text{tr}(M N(\sigma)) \leq 1-\epsilon$~~

But $N^\dagger(M)$ is a POVM element such that

$$\text{tr}(N^\dagger(M) \rho) = \text{tr}(M N(\rho)) \geq 1-\epsilon.$$

Thus

$$D_h^\epsilon(\rho || \sigma) = \sup_{\substack{M \\ \text{tr} M \rho \geq 1-\epsilon}} (-\log \text{tr}(M \sigma))$$

$$\geq -\log \text{tr}[N^\dagger(M) \sigma]$$

$$= -\log \text{tr}[M N(\sigma)]$$

$$= D_h^\epsilon(N(\rho) || N(\sigma)).$$

(1) First, $D_h^\epsilon(\rho || \sigma) \geq D_h^\epsilon(\rho || \rho) = D_h^\epsilon(\rho || \rho)$.

Then, $D_h^\epsilon(\tau || \tau) = \sup_{\substack{M \\ \text{tr} M \tau \geq 1-\epsilon}} (-\log \text{tr}(M \tau))$

$$= -\log \inf_{\substack{M \\ \text{tr} M \tau \geq 1-\epsilon}} \text{tr}(M \tau) = -\log(1-\epsilon).$$

(3) $\text{tr}(M \tau) = \frac{\text{tr} M}{d} \geq \frac{\text{tr}(M \tau)}{d} \geq \frac{1-\epsilon}{d}$. Equal for $M = (1-\epsilon)\tau$.

Relations to $D(p||\sigma)$

Thm. "Quantum Stein's Lemma".

$$\lim_{n \rightarrow \infty} \frac{1}{n} D_n^{\epsilon}(p||\sigma) = D(p||\sigma)$$

for all ϵ with $0 < \epsilon < 1$.

Prop. upper bound

$$D_n^{\epsilon}(p||\sigma) \leq \frac{D(p||\sigma) + H_2(\epsilon)}{1 - \epsilon}$$

with H_2 the binary entropy.

Pf. We'll prove these in later sections.

⑥ SQRT Measurement.

Def. For any set $\{S_x\}$ of ~~positive operators~~ positive operators,

$$D_x = \left(\sum_x S_x \right)^{-1/2} S_x \left(\sum_x S_x \right)^{-1/2}$$

are the elements of a POVM called the square root measurement.

- Used to convert unnormalized set of POVM elements to a POVM.

⑤

⑤ Hayashi-Nagaoka Lemma

- This is important because it relates the probabilities for D_x to S_x in the $\sqrt{\text{meas}}$.

Lemma. For $0 \leq S \leq \mathbb{1}$, $T \geq 0$, and any $c > 0$,

$$\left[(\mathbb{1} - S + T)^{-1/2} S (S + T)^{-1/2} \right] \leq (1+c)(\mathbb{1} - S) + (2+c+c^{-1})T.$$

Outline of proof.

- For any $c > 0$, $(A - cB)^{\dagger}(A - cB) \geq 0$.
- Let $A = \sqrt{T}$ and $B = \sqrt{T}((S+T)^{-1/2} - \mathbb{1})$.
- Use $T \leq (S+T)$.
- Use that $(\cdot)^{1/2}$ is operator monotone, $(S+T)^{1/2} \geq S^{1/2}$.
- Use $S^{1/2} \geq S$ b/c $0 \leq S \leq \mathbb{1}$.
- Bunch of algebra.

OK, now ready for achievability proofs is.

② Classical Communication: Random Coding.



- We want to find lower bounds on the capacity

$$C_\epsilon(N) = \sup_{\substack{(E,D) \\ P_e \leq \epsilon}} R$$

Do this by defining a code and calculating (R, F) .

Encoding:

- Choose any $P(x)$ and $\{p_x\}$.
- Set of messages is $m = 1, \dots, K$.
- Random Codebook: For each m , choose an x_m randomly w/prob $P(x)$.
 $\hookrightarrow$ Some $x_m = x_{m'}$ might be the same.

~~Let $w_x \equiv N(p_x)$~~

- Let $P_{XA} \equiv \sum_x P(x) |x\rangle\langle x| \otimes P_x$ and

$$\begin{aligned} w_{XB} &\equiv (I_X \otimes N_{A \rightarrow B})(P_{XA}) \\ &= \sum_x P(x) |x\rangle\langle x| \otimes N(p_x). \end{aligned}$$

Let $w_x \equiv N(p_x)$.

⑦

Decoding:

- Let S be a POVM for the hypothesis test

$$w_{XB} \text{ vs } w_X \otimes w_B$$

Suppose that $\text{tr}(S w_{XB}) = 1 - \epsilon'$.

- Let $S_x = \text{tr}_A(|x\rangle\langle x| \otimes \mathbb{1}) S$.

$$\hookrightarrow \text{tr}_B(w_x S_x) = \text{tr}(|x\rangle\langle x| \otimes w_x S).$$

~~Proof:~~

$$\text{tr}[(\mathbb{1} \otimes w_x)(|x\rangle\langle x| \otimes \mathbb{1}) S] = \text{tr}(w_x \text{tr}_A(S_x)).$$

- Let

$$D_m = \left(\sum_{n=1}^K S_{x_n} \right)^{-1/2} S_{x_m} \left(\sum_{n=1}^K S_{x_n} \right)^{-1/2}$$

be decoding POVM.

Now we calculate error for this scheme.

By definition,

$$P_e = 1 - \frac{1}{K} \sum_m \text{tr}(D_m N(p_{x_m})) = \frac{1}{K} \sum_m \text{tr}(w_{x_m} (\mathbb{1} - D_m))$$

But by H-N Lemma,

$$\mathbb{1} - D_m \leq (1 + c)(\mathbb{1} - S_{x_m}) + (2 + c + c^{-1}) \sum_{n \neq m} S_{x_n}.$$

So

$$P_e \leq \frac{1+c}{K} \sum_m \text{tr}(w_{x_m} \mathbb{1} - S_{x_m}) + \frac{2+c+c^{-1}}{K} \sum_m \sum_{m' \neq m} \text{tr}(w_{x_m} S_{x_{m'}}).$$

Now average over codes. Imagine we generate the random codebook many independent times.

On average,

$$P_e \leq \frac{1+c}{K} \sum_m \sum_x P(x) \text{tr}(w_x (\mathbb{1} - S_x)) + \frac{2+c+c^{-1}}{K} \sum_m \sum_{m' \neq m} \sum_{x, x'} P(x) P(x') \text{tr}(w_x S_{x'})$$

After averaging, ~~all~~ all terms m -independent.

So $\sum_m \rightarrow K$, $\sum_{m' \neq m} \rightarrow (K(K-1))$.

$$P_e \leq \frac{(1+c)}{K} \sum_x P(x) \text{tr}(w_x (\mathbb{1} - S_x)) + \frac{(2+c+c^{-1})(K-1)}{K} \sum_{x, x'} P(x) P(x') \text{tr}(w_x S_{x'})$$

$$= (1+c) \left[1 - \sum_x P(x) \text{tr}(w_x S_x) \right] +$$

Now observe that

$$\sum_x P(x) \text{tr}(w_x S_x) = \sum_x P(x) \text{tr}(|x\rangle\langle x| w_x S) = \text{tr}(w S).$$

(9)

And,

$$\begin{aligned} \text{tr}(w_x w_B S) &= \text{tr} \left[\sum_x P(x) |x\rangle\langle x| w_x S \right] \\ &= \sum_{x, x'} P(x) P(x') \text{tr} \left[(|x\rangle\langle x| w_x) (|x'\rangle\langle x'| w_{x'} S) \right] \\ &= \sum_{x, x'} P(x) P(x') \text{tr} [w_x S_{x'}]. \end{aligned}$$

Plugging back in,

$$P_e \leq (1+c) \underbrace{\left[1 - \text{tr}(w S) \right]}_{1-\epsilon'} + (2+c+c^{-1})(K-1) \underbrace{\left[\text{tr}(w_x w_B S) \right]}_{\text{Choose optimal } S \text{ to minimize}}$$

If we take the optimal S for the Hyp. Test.

$$P_e \leq (1+c)\epsilon' + (2+c+c^{-1}) 2^R 2^{-D_h^{\epsilon'}(w_x w_B \| w_x w_B)}$$

used $K-1 \leq K \leq 2^R$.

~~we used $K-1 \leq K \leq 2^R$. Let $P_e = \epsilon$ (relabel). Rearranging~~

$$2^{R-D_h^{\epsilon'}} \geq \frac{\epsilon - (1+c)\epsilon'}{2+c+c^{-1}}$$

$$R \geq D_h^{\epsilon'} - \log \frac{2+c+c^{-1}}{\epsilon - (1+c)\epsilon'}$$

Optimizing over c gives $c = \frac{\epsilon - \epsilon'}{\epsilon + \epsilon'}$, so

$$R \geq D_h^{\epsilon'} - \log \left(\frac{4\epsilon}{(\epsilon - \epsilon')^2} \right)$$

• This is average over many codes. It follows, at least one code in the class has P_e below the average.

Suppose we want $P_e \leq \epsilon$ for fixed ϵ . (11)

$$P_e \leq (1+c)\epsilon' + (2+c+c^{-1})2^R 2^{-D_h^{\epsilon'}} \leq \epsilon.$$

Possible only if $\epsilon' < \epsilon$. If so, we need

$$2^{R-D_h^{\epsilon'}} \leq \frac{\epsilon - (1+c)\epsilon'}{2+c+c^{-1}}$$

for some $c > 0$. Optimizing gives $c = \frac{\epsilon - \epsilon'}{\epsilon + \epsilon'}$,

$$R \leq D_h^{\epsilon'} - \log\left(\frac{4\epsilon}{(\epsilon - \epsilon')^2}\right).$$

This is average rate over many random codes. At least one code exists better than average.

We have proved:

Thm.

Fix $\epsilon > 0$, $\epsilon' < \epsilon$, $P(x)$, $\{P_z\}$. If

$$R \leq D_h^{\epsilon'}(W_{XB} \| W_{X \otimes B}) - \log\left(\frac{4\epsilon}{(\epsilon - \epsilon')^2}\right)$$

then $\exists$ code with $P_e \leq \epsilon$.

Thus $\forall \epsilon' < \epsilon$,

$$C_\epsilon(N) \geq \sup_{P_{XA}} D_h^{\epsilon'}(W_{XB} \| W_{X \otimes B}) - \log(2).$$

What is this $D_h^{\epsilon'}(W \| W_{\otimes})$?

$$X(N) = \sup_{P_{XA}} I(X:B)_{W=N(p)}$$

$$X_\epsilon^t(N) = \sup_{P_{XA}} I_h^t(X:B)_{W=N(p)}$$

$$D(W \| W_{\otimes}) \quad \text{vs} \quad D_h^t(W \| W_{\otimes})$$

Generalized Mutual Information.

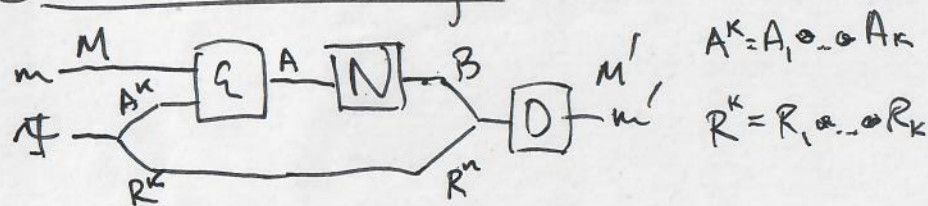
Lecture 6

Today:

(... Finish Previous Lecture: Random Coding)

- EPR Position Based Coding
- Converse Bounds
- Asymptotic Limit?

① Position Based Coding



Setup:

- Alice and Bob share K copies of PAR:

$$\mathbb{P}_{A^K R^K} = (\text{PAR})^{\otimes K}$$

- Messages $m = 1 \dots K$.

- Let (one-copy)

$$W_{BR} = N_{A \rightarrow B} \otimes I_R (\text{PAR})$$

so

$$W_B = \text{tr}_R W_{BR}, W_R = P_R.$$

Encoding:

- Alice sends with system to Bob through N ,
 $N \circ E_m = N_{A_m \rightarrow B}$.

Equivalently, $E_m = \text{tr}_{A^m/A_m}$ (trace out all $m' \neq m$)

Decoding:

- Again let S be hypothesis tester for W_{BR} vs $W_B \otimes W_R$.

Let $S_m = S^{BR_m}$.

- Again let

$$D_m = \left(\sum_n S_n \right)^{-1/2} S_m \left(\sum_n S_n \right)^{-1/2}$$

Analysis:

- What does Bob see? Bob's state is

$$N_{A_m \rightarrow B} (p_{A^R_1} \otimes \dots \otimes p_{A^R_K})$$

But can only access BR^K , so really $\text{tr}_{A^m/A_m}(m)$.

Thus,

$$\begin{aligned}\sigma_m^{BR^n} &= p^{R_1} \otimes \dots \otimes N_{A_n \rightarrow B}(p^{R_n}) \otimes \dots \otimes p^{R_K} \\ &= p^{R_1} \otimes \dots \otimes w^{R_n B} \otimes \dots \otimes p^{R_K}.\end{aligned}$$

The 2-Party marginals of this are

$$\sigma_m^{BR_m} = w^{R_m B}, \quad m=m$$

$$\begin{aligned}\sigma_m^{BR_{m'}} &= w^B \otimes p^{R_{m'}} \\ &= w^B \otimes w^{R_{m'}}, \quad m' \neq m.\end{aligned}$$

Alice

Bob

p^{R_1}		$w^B \otimes p^R$	] Hyp. Test. to identify state.
p^{R_2}	$\xrightarrow{m=2}$	w^{BR}	
p^{R_3}		$w^B \otimes p^R$	
p^{R_4}		$w^B \otimes p^R$	

Error Prob:

• Like before

$$P_e = \frac{1}{K} \sum_m \text{tr}(\sigma_m (I - D_m))$$

$$I - D_m \leq (1+\epsilon)(I - S_m) + (2\epsilon c + \epsilon^{-1}) \sum_{k \neq m} S_m$$

③

• Plugging this in with optimal S ,

$$\begin{aligned}\text{tr}(S_m \sigma) &= \text{tr}(S^{BR_m} \sigma) \\ &= \text{tr}(S \sigma^{BR_m})\end{aligned}$$

$$\begin{aligned}\downarrow \\ \text{tr}(S w^{RB}) &= 1 - \epsilon' \\ (m=m)\end{aligned}$$

$$\begin{aligned}\downarrow \\ \text{tr}(S w^{R \otimes w^B}) &= 2^{-D_n^H} \\ (m' \neq m)\end{aligned}$$

~~So yet again, this~~

~~$P_e \leq \epsilon$~~

• This is the same for every term in sum, so again $\sum \rightarrow K$, $\sum \sum \rightarrow K(K-1)$.

So yet again

$$P_e \leq (1+\epsilon)\epsilon' + (2\epsilon c + \epsilon^{-1}) 2^{R - D_n^H} \leq \epsilon.$$

Solving we have ~~this~~ then

Lemma. Fix $\epsilon > 0$, $\epsilon' < \epsilon$, PAR-

IF

$$R \leq D_n^H - \log\left(\frac{4\epsilon}{(1+\epsilon)^2}\right),$$

there exists an (R, ϵ) EA code.

Thm. For any $0 \leq \epsilon' < \epsilon$,

$$C_{\epsilon}^{\text{EA}}(N) \geq \sup_{\text{PRA}} D_n^{\epsilon'}(W_{RB} \| U_{R \otimes B}) - \log \frac{4\epsilon}{(\epsilon - \epsilon')^2}$$

Additionally, the sup can be limited to pure states with $R \sim A$.

Proof. We just showed the existence of such a code for any PRA. For the sup restriction:

- Purify PRA to Ψ_{RAE} . This changes nothing in the code, which acts only on RA. Thus for any PRA code, there also exists a Ψ_{RAE} code.

- Now considering sup Ψ_{RAE} (w.l.g.) over pure states, any pure Ψ_{RAE} , in the Schmidt Basis, looks like a state on RAA' . Can't delete the rest of the space without changing anything.

• We can even use this lemma to prove the plain capacity theorem!

- Any Prod-Assisted Code = A plain code.

• Note that P_{ϵ} is convex in $\mathcal{F}$,

$$P_{\epsilon} = \frac{1}{n} \sum_{\mathcal{F}} \text{tr}(D_n(N \otimes E_n(\mathcal{F})))$$

Thus if $\mathcal{F} = \sum p_i \mathcal{F}_i$, at least one $\mathcal{F}_i$ must do better than $\mathcal{F}$.

- Thus, starting from any CQ state

$$P_{XA} = \sum_x p(x) |x\rangle\langle x| \otimes P_x$$

we can convert P.B.C. into plain code.

Thm (Again). For any $0 \leq \epsilon' < \epsilon$,

$$C_{\epsilon}(N) \geq \sup_{\substack{P_{XA} \\ (\text{CQ})}} D_n^{\epsilon'} - \log(n)$$

To summarize so far; let us define

$$X_h^t(N) = \sup_{P_{XA}} I_h^t(X:B)_\omega$$

$$I_h^e(N) = \sup_{P_{RA}} I_h^t(R:B)_\omega,$$

with $\omega = N_{A \rightarrow B}(p)$.

• We've proved achievable rates, it turns out they're nearly maximal. Next we'll find, that

$$X_h^t - \log \frac{4\epsilon'}{(\epsilon - \epsilon')^2} \leq C_\epsilon \leq X_h^e$$

$$I_h^t - \log \frac{4\epsilon'}{(\epsilon - \epsilon')^2} \leq C_\epsilon^{EA} \leq I_h^{EF}$$

Achievability
 Bounds
 ✓

Converse
 Bounds
 (next.)

② Converse Bounds.

Generalized Divergence

Def. A function $D(p||\sigma)$ on states p, σ is said to be a divergence from the same system such that

- 1. D takes values in $\mathbb{R} \cup \{\infty\}$.
- 2. $D(p||\sigma) \geq 0$ for all states p, σ , with $D(p||\sigma) = 0$ iff $p = \sigma$.
- 3. $D(p||\sigma) \geq D(M_N(p)||M_N(\sigma))$ for CPTP N .

is a generalized divergence.

Examples.

- Kullback: $D(p||\sigma) = \text{tr}(p \log p - p \log \sigma)$
- Hyp-Test: $D_h^t(p||\sigma) = \sup_{\text{tr}(mp) = 1-t} \text{tr}(M_t \sigma)$
- Rényi: $D_\alpha(p||\sigma) = \frac{1}{\alpha-1} \log \text{tr}(p^\alpha \sigma^{1-\alpha})$, $\alpha \in (0, \infty)$
 [is divergent for $\alpha \in (0, 2]$]
- Sandwiched Rényi: $\tilde{D}_\alpha(p||\sigma) = \frac{1}{\alpha-1} \log \text{tr} \left(p^{\frac{1+\alpha}{2}} \sigma^{\frac{1-\alpha}{2}} p^{\frac{1+\alpha}{2}} \right)$
 [converges for $\alpha \in (\frac{1}{2}, \infty)$].

② Converse Bounds:

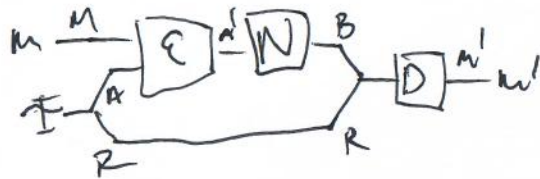
Thm.

$$C_e(N) \leq X_h^e(N)$$

$$C_e^{EA}(N) \leq I_h^e(N).$$

Proof.

EA First. Consider EA setup



Alice starts with state

$$\gamma_{MRA} = \frac{1}{K} \sum_m |m\rangle\langle m| \otimes \Psi_{AR}$$

This turns into

$$\begin{aligned} P_{mm'} &= D_{RB \rightarrow m'} \circ N_{A' \rightarrow B} \circ E_{MA \rightarrow A'm}(\gamma_{MRA}) \\ &= \sum_{mm'} \frac{1}{K} P_{m'}(m) |m\rangle\langle m| \otimes |m'\rangle\langle m'| \end{aligned}$$

(Alice "keeps a copy".)

8/9

By defn,

$$I_h^e(N) = \sup_{P_{RA}} D_h^e(W_{RB} \| W_R \otimes W_B)_{W=N_{A \rightarrow B}(P_{RA})}$$

Absorb Alice's message system into $\bar{R}$, letting

$$\begin{aligned} P_{\bar{R}A} &= E_{MA \rightarrow MA'}(\gamma_{MRA}) \\ &= \frac{1}{K} \sum_m |m\rangle\langle m| \otimes E_m^{A \rightarrow A'}(\Psi_{RA}). \end{aligned}$$

Then

$$\begin{aligned} W_{\bar{R}B} &= N \circ E(\gamma) \\ &= \frac{1}{K} \sum_m |m\rangle\langle m| \otimes N_{A' \rightarrow B} \circ E_m^{A \rightarrow A'}(\Psi_{RA}) \end{aligned}$$

And also

$$P_{mm'} = D_{RB \rightarrow m'}(W_{\bar{R}B}).$$

By defn of sup, in particular

$$I_h^e(N) \geq D_h^e(W_{\bar{R}B} \| W_{\bar{R}} \otimes W_B).$$

$$\geq \inf_{\sigma_B} D_h^e(W_{\bar{R}B} \| W_{\bar{R}} \otimes \sigma_B).$$

$$\text{But } W_{\bar{R}} = W_M \otimes W_R = P_M \otimes W_R = P_M \otimes \Psi_R.$$

10

So

$$I_h^f \geq \inf_{\sigma_B} D_h^f(w_{M \otimes B} \| P_M \otimes w_R \otimes \sigma_B).$$

Now apply decoding $D_{RB \rightarrow M'}$. By DPI,

$$\begin{aligned} I_h^f &\geq \inf_{\sigma_B} D_h^f(P_{MM'} \| P_M \otimes D(w_R \otimes \sigma_B)) \\ &= \inf_{\sigma_B} D_h^f(P_{MM'} \| P_M \otimes \sigma_{M'}). \end{aligned}$$

Now consider H.T. PPM element

$$\Pi = \sum_m |m\rangle\langle m| \otimes |m\rangle\langle m|. \quad (\text{Comparator})$$

This obeys $\text{tr}(\Pi P_{MM'}) = 1 - p_e \geq 1 - \epsilon$.

Also, $P_M = \frac{\mathbb{1}}{K}$, so

$$\begin{aligned} \text{tr}(\Pi P_M \otimes \sigma_{M'}) &= \sum_m \frac{1}{K} \text{tr}(|m\rangle\langle m| \otimes |m\rangle\langle m| \otimes \sigma_{M'}) \\ &= \sum_m \frac{1}{K} \text{tr}(|m\rangle\langle m| \otimes \sigma_{M'}) \\ &= \frac{1}{K} \text{tr}(\mathbb{1} \otimes \sigma_{M'}) = \frac{1}{K}. \end{aligned}$$

Thus $D_h^f \geq -\log \frac{1}{K} = \log K$.

⑩

Putting this together, for any (R, ϵ) code, we find ⑪
 $R \leq I_h^f(N)$. ✓

For $C_f(N)$, same exact logic.

• Choose $P_{XX'} = \frac{1}{K} \sum_m |m\rangle\langle m| \otimes P_m$.

• Then $X_h^f(N) \geq D_h^f(w_{XB} \| w_X \otimes w_B)$

$$\geq \inf_{\sigma_B} D_h^f(w_{XB} \| w_X \otimes \sigma_B)$$

• Apply decoder,

$$\geq \inf_{\sigma_B} D_h^f(P_{XX'} \| P_X \otimes \sigma_{X'})$$

• Use Comparator $\Pi = \sum_m |m\rangle\langle m| \otimes |m\rangle\langle m|$,
 $\geq \frac{-\log}{\text{tr}}(\Pi P_X \otimes \sigma_{X'}) = \log K$. ✓

That's it.

Next time: Metaconverse / Gai / ZOC Divergence
 Asymptotic Limit.

Lecture 7

(1)

Today:

- Metaconverse
- Generalized Divergences
- Asymptotic Limit.

Generalized Divergence

Defn. A function $D(p||\sigma)$ states p and $\sigma \geq 0$ from the same system, is a generalized divergence if

1. D takes values in $\mathbb{R} \cup \{\infty\}$.
2. $D(p||\sigma) \geq d_0 \quad \forall p, \sigma \in \{\text{states}\}$,
where d_0 is system-independent constant.
3. $D(p||S) \geq D(N(p)||N(S))$
for all states p and operators $S \geq 0$.

(2)

Examples:

- $D(p||\sigma) = \text{tr}(p \log p - p \log \sigma)$ (Umegaki RE)
- $D_h^t(p||\sigma)$ (Hyp Test RE)
- $d(p, \sigma) = \frac{1}{2} \|p - \sigma\|_1$ (Trace Distance)
- $d(p, \sigma) = \sqrt{1 - F(p, \sigma)^2}$ (Purified Distance/Fidelity)
- $D_\alpha(p||\sigma) = \frac{1}{\alpha-1} \log \text{tr}(p^\alpha \sigma^{1-\alpha})$, $\alpha \in (0, 2)$ (Petz-Rényi RE)
- $\tilde{D}_\alpha(p||\sigma) = \frac{1}{\alpha-1} \log \text{tr}[(\sigma^{\frac{1-\alpha}{2\alpha}} p \sigma^{\frac{1-\alpha}{2\alpha}})^\alpha]$, $\alpha \in (\frac{1}{2}, \infty)$ (sand. Petz RE)
- $D_{\max}(p||\sigma) = \inf_{p \leq \lambda \sigma} \lambda$ (Max RE)

Some Preparatory Stuff

(3)

Fact. Let $0 < p_c < t < 1 - \frac{1}{k}$. Then

$$D(\{ \frac{p_c}{1-p_c} \} \| \{ \frac{1-\frac{1}{k}}{1/k} \}) \geq D(\{ \frac{t}{1-t} \} \| \{ \frac{1-\frac{1}{k}}{1/k} \}).$$

$$\begin{array}{c|c|c} 1 & 1 & 1 \\ \hline p_c & t & 1-\frac{1}{k} \end{array} \Rightarrow \text{New Dist is more similar.}$$

Proof. Exercise: Show $\exists$ classical channel Λ
s.t. $\{ \frac{p_c}{1-p_c} \} \xrightarrow{\Lambda} \{ \frac{t}{1-t} \}$ and $\{ \frac{1-\frac{1}{k}}{1/k} \} \rightarrow \{ \frac{1-\frac{1}{k}}{1/k} \}$.

Fact. $\sup_{p_{XB}} D(W_{XB} \| W_X \otimes W_B) \geq$

$$\sup_{p_{XA}} \inf_{\sigma_B} D(W_{XB} \| W_X \otimes \sigma_B).$$

Proof. This is obvious.

Fact. Let $\Pi = \sum_m |m\rangle\langle m| \otimes |a\rangle\langle a|$.

$$\text{Then } \text{tr}(\Pi \frac{\Pi}{k} \otimes \sigma) = \frac{1}{k}.$$

Proof. $\text{tr}(\Pi \frac{\Pi}{k} \otimes \sigma) = \frac{1}{k} \text{tr}(\sum_m |m\rangle\langle m| \otimes |a\rangle\langle a| \sigma) = \frac{1}{k} \text{tr}(\sigma|_a) = \frac{1}{k}.$

(4)

Defn. (Comparator Channel).

For a state $P_{MM'}$, the comparator channel is $Q(P) = \text{tr}(\Pi p) |0\rangle\langle 0| + \text{tr}(\mathbb{I} - \Pi) p |1\rangle\langle 1|$,
where $\Pi = \sum_m |m\rangle\langle m| \otimes |m\rangle\langle m|$.

Now we are ready for the Metconverse.

Lemma. (Metconverse). For any channel N and $0 < \epsilon < 1 - \frac{1}{K}$, if $\exists$ a code with $P_e \leq \epsilon$ for K messages, then

$$\sup_{P_X A} \inf_{\sigma_B} D(W_{XB} \| W_X \otimes \sigma_B)_{W=N_{A \rightarrow B}(P)} \\ \geq D\left(\left\{\frac{\epsilon}{1-\epsilon}\right\} \middle\| \left\{\frac{1-\frac{1}{K}}{1/K}\right\}\right)$$

for any given divergence.

Lemman. (EA Metconverse),

5

$$\sup_{P_{RA}} \inf_{\sigma_B} D(W_{RB} || W_R \otimes \sigma_B)_{W=N_{A \rightarrow B}(P)} \\ \geq D\left(\left\{\frac{t}{1-t}\right\} || \left\{\frac{1-\frac{1}{n}}{1/n}\right\}\right).$$

Proof, Plan First:

Let $P_{XA} = \frac{1}{K} \sum_m |m\rangle\langle m| \otimes P_m$, with P_m from $(\mathbb{C}, K)$ code.

Then after channel, $N_{A \rightarrow B}$,

$$W_{XB} = \frac{1}{K} \sum_m |m\rangle\langle m| \otimes N(P_m).$$

Now apply decoding $D_{B \rightarrow m'}$ and compression Q ,

$$D(W_{XB}) = \frac{1}{K} \sum_m |m\rangle\langle m| \otimes D(N(P_m))$$

$$= \frac{1}{K} \sum_{m, m'} P_{m'|m} |m\rangle\langle m| \otimes |m'\rangle\langle m'|$$

$$Q \circ D(W_{XB}) = (1 - P_e) |0\rangle\langle 0| + P_e |1\rangle\langle 1|$$

For the other state,

(6)

$$D(w_x \otimes \sigma_B) = w_x \otimes D(\sigma_B) = \frac{1}{K} \otimes D(\sigma_B)$$

$$Q \circ D(w_x \otimes \sigma_B) = Q\left(\frac{1}{K} \otimes \sigma_B\right) = \frac{1}{K} |0\rangle\langle 0| + \left(1 - \frac{1}{K}\right) |1\rangle\langle 1|$$

Thus by monotonicity,

~~$$D(w_x \otimes \sigma_B) \geq D(Q \circ D(w_x \otimes \sigma_B))$$~~

$$D(w_x \otimes \sigma_B) \geq D(Q \circ D(w_x \otimes \sigma_B))$$

$$= D\left(\left\{\frac{1-p_k}{K}\right\} \parallel \left\{\frac{1}{K}\right\}\right)$$

$$\geq D\left(\left\{\frac{1-\epsilon}{K}\right\} \parallel \left\{\frac{1}{K}\right\}\right) \checkmark$$

For EA, do the same thing but start from

$$P_{RA} = \frac{1}{K} \sum_m |m\rangle\langle m| \otimes E_m(\mathbb{I}_{R'A})$$

using existence of the code. (Absorb $R'M = R$.) $\checkmark$

②

Apply to different D

Relative Entropy

$$\begin{aligned} D\left(\frac{1-\epsilon}{K}\right) \parallel \left(\frac{1-\epsilon}{K}\right) &= -h_2(\epsilon) - \epsilon \log \frac{K-1}{K} - (1-\epsilon) \log \frac{1}{K} \\ &= -h_2(\epsilon) + \epsilon \log \frac{K}{K-1} + (1-\epsilon) \log K \\ &\geq (1-\epsilon) \log K - h_2(\epsilon). \end{aligned}$$

Thus,

$$(1-\epsilon) \log K - h_2(\epsilon) \leq \sup_{P \times A} I(X:A)_{\omega=X(N)}$$

$$R \leq \frac{X(N) + h_2(\epsilon)}{1-\epsilon} \quad \text{for any (R, \epsilon) code.}$$

$$C_\epsilon(N) \leq \frac{X(N) + h_2(\epsilon)}{1-\epsilon}$$

Similarly,

$$\left| C_\epsilon^{EA}(N) \leq \frac{I(N) + h_2(\epsilon)}{1-\epsilon} \right|.$$

(Only thing that changes is $\sup_{P \times A}$.)

Sanitized Rényi:

You end up with

$$\tilde{D}_\alpha(\{1-\epsilon\} \| \{1-\epsilon\}^{1-\frac{1}{\alpha}} \{ \frac{\epsilon}{n} \}^{\frac{1}{\alpha}}) = \log K + \frac{\alpha}{\alpha-1} \log(1-\epsilon)$$

$$C_\epsilon(N) \leq \chi_\alpha(N) - \frac{\alpha}{\alpha-1} \log(1-\epsilon)$$

($\alpha > 1$). Nice form.

~~scribbled out text~~

Note ... Since here dists are classical,
 $\tilde{D}_\alpha = D_\alpha$ which is usual Rényi divergence:

$$\log \langle \left(\frac{P}{Q} \right)^s \rangle^{1/s} = \frac{1}{s} \log \langle \left(\frac{P}{Q} \right)^s \rangle$$

$$= \frac{1}{s} \log \sum \frac{P^s}{Q^s}$$

$$= \frac{1}{\alpha-1} \log \sum \frac{P^\alpha}{Q^{\alpha-1}} \quad (s = \alpha-1)$$

$$= D_\alpha(P \| Q).$$

(9)

Hypothesis Testing

Choose $M = \{1\}$. Then $\text{tr}(M \{1-\epsilon\}) = 1-\epsilon$ is valid for

$$D_n^\epsilon(\{1-\epsilon\} \| \{1-\frac{1}{K}\}) = \sup_{\substack{M \\ \text{tr}(M\sigma) \geq 1-\epsilon}} -\log(\text{tr}(M\sigma)) \geq -\log(\text{tr}(\{1\} \{1-\frac{1}{K}\}))$$

$$= -\log \frac{1}{K} = \log K.$$

$$\Rightarrow C_\epsilon(W) \leq \chi_h^\epsilon(N)$$

$$C_\epsilon^\text{KA}(W) \leq I_h^\epsilon(N)$$

As we already had. ✓

Asymptotic Limit.

(10)

Converse.

We found from M.C.,

$$C_\epsilon(N) \leq \frac{\chi(N) + h_2(\epsilon)}{(1-\epsilon)}.$$

For $\epsilon = 0$,

$$C_0(N) \leq \chi(N) + h_2(\epsilon)$$

Multiple uses,

$$\frac{C_0(N^{(n)})}{n} \leq \frac{1}{n} \chi(N^{(n)}) + \frac{h_2(\epsilon)}{n}$$

Let $N \rightarrow \infty$,

$$C(N) \leq \lim_{n \rightarrow \infty} \frac{1}{n} \chi(N^{(n)}) \quad \checkmark$$

Achievability. (Wang-Pennec.)

(11)

Fix $\epsilon > \epsilon' > 0$. Choose a k and consider $N^{\otimes k}$.

For any k , we have a code where

$$\frac{R(k)}{k} \geq \frac{1}{k} \chi_h^{\epsilon'}(N^{\otimes k}) - \frac{1}{k} \log \frac{4\epsilon}{(\epsilon - \epsilon')^2}$$

Now, repeating $N^{\otimes k}$ many times, $(N^{\otimes k})^{\otimes n}$, we have

$$\lim_{n \rightarrow \infty} \frac{1}{n} \frac{R(k)}{k} \geq \lim_{n \rightarrow \infty} \frac{1}{n} \frac{1}{k} \chi_h^{\epsilon'}((N^{\otimes k})^{\otimes n}) - \frac{1}{n} \frac{1}{k} \log \frac{4\epsilon}{(\epsilon - \epsilon')^2}$$

$$\rightarrow \frac{1}{k} \lim_{n \rightarrow \infty} \frac{1}{n} \chi_h^{\epsilon'}((N^{\otimes k})^{\otimes n}).$$

But A.E.P. told us that $\frac{1}{n} D_h^{\epsilon}(p^{\otimes n} \| q^{\otimes n}) \rightarrow D(p \| q)$, for all ϵ . So

$$\frac{1}{n} \chi_h^{\epsilon}(N^{\otimes n}) = \frac{1}{n} \sup_{p_{XA^n}} D_h^{\epsilon}(I_{X \otimes N_A^{\otimes n}}(p_{XA^n}) \| I_{X \otimes N_A^{\otimes n}}(p_{X \otimes p_A^n}))$$

$$\geq \frac{1}{n} \sup_{(p_{XA})^{\otimes n}} [m]$$

$$= \frac{1}{n} \sup_{p_{XA}} D_h^{\epsilon}(I_{\otimes N}(p)^{\otimes n} \| I_{\otimes N}(p \otimes p)^{\otimes n})$$

$$\xrightarrow{n \rightarrow \infty} \sup_{p_{XA}} D^{\epsilon}(I_{\otimes N}(p) \| I_{\otimes N}(p \otimes p)) = \chi(N).$$

So going back,

(12)

$$\lim_{n \rightarrow \infty} \frac{1}{n} \frac{R(n)}{K} \geq \frac{1}{K} \lim_{n \rightarrow \infty} \frac{1}{n} \chi_n^{\epsilon'}((N \otimes k)^n) \\ \geq \frac{1}{K} \chi(N).$$

This holds for any ϵ, ϵ' . ~~We have~~ We have

For any $\epsilon > \epsilon'$, for any K ,

$$\lim_{n \rightarrow \infty} \frac{C_{\epsilon}((N \otimes k)^n)}{nK} \geq \frac{1}{K} \chi(N \otimes k).$$

Thus, taking the limit in this way,

$$\lim_{K \rightarrow \infty} \lim_{n \rightarrow \infty} \frac{C_{\epsilon}((N \otimes k)^n)}{nK} \geq \lim_{K \rightarrow \infty} \frac{1}{K} \chi(N \otimes k).$$

Taking the limit $\epsilon' \rightarrow 0$, then $\epsilon \rightarrow 0$,

$$C(N) \geq \lim_{K \rightarrow \infty} \frac{1}{K} \chi(N \otimes k).$$

Thus, $\boxed{C(N) = \lim_{K \rightarrow \infty} \frac{1}{K} \chi(N \otimes k)}.$

The HSW theorem.

Similarly, $C^{\text{EH}}(N) = I(N).$

Lecture 8

①

Today:

- Entropy Zoo

Recall: Generalized Divergences

$D(p||\sigma)$ s. h.

$D(p||\sigma) \in \mathbb{R} \cup \{\infty\}$.

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

$D(N(p)||N(\sigma)) \leq D(p||\sigma)$.

Examples

- GRE
- Hyp. Test. RE
- Trace Distance
- Purified distance (Fidelity)
- Petz-Renyi
- Sandwiched Renyi
- Max

②

RE as Parent Quantity

For QRE

$$D(\rho \parallel \sigma) = \text{tr}(\rho \log \rho - \rho \log \sigma)$$

we have

- Entropy

$$\begin{aligned} H(\rho) &= -D(\rho \parallel \mathbb{I}) \\ &= -\text{tr} \rho \log \rho \end{aligned}$$

- Conditional Entropy

$$\begin{aligned} H(A|B)_\rho &= -D(\rho_{AB} \parallel (\mathbb{I}_A \otimes \rho_B)) \\ &= H(A, B)_\rho - H(B)_\rho \end{aligned}$$

- Mutual Information

$$I(A:B)_\rho = D(\rho_{AB} \parallel \rho_A \otimes \rho_B).$$

(3)

But also,

$$I(A:B)_p = \inf_{\sigma_B} D(p_{AB} \| p_A \otimes \sigma_B)$$

$$H(A|B)_p = - \inf_{\sigma_B} D(p_{AB} \| \mathbb{I}_A \otimes \sigma_B)$$

Proof.

$$D(p_{AB} \| p_A \otimes \sigma_B) = -S(p_{AB}) + \text{tr } p_{AB} \log(p_A \otimes \sigma_B)$$

$$= -S(p_{AB}) + \text{tr } p_{AB} \log[(p_A \otimes \mathbb{I})(\mathbb{I} \otimes \sigma_B)]$$

$$= -S(p_{AB}) + \text{tr}(p_{AB} (\log(\mathbb{I} \otimes \mathbb{I}) + \log(\mathbb{I} \otimes \sigma_B)))$$

$$= -S(p_{AB}) + S(p_A) - \text{tr } p_B \log \sigma_B$$

$$= -S(p_{AB}) + S(p_A) + S(p_B)$$

$$= I(A:B)_p.$$

Because

$$D(p_B \| \sigma_B) = \text{tr } p_B \log p_B - \text{tr } p_B \log \sigma_B \geq 0.$$

Similarly for conditional.

④

So for generalized information quantities there are usually two versions, regular ($\uparrow$) and infimized ($\downarrow$). (Note K&W Book uses all $\downarrow$).

Smoothing.

Another way to modify info quantities is by smoothing.

Defn. Fidelity. (Wilde Books Defn).

$$F(p, \sigma) = \max_{\psi, \phi} |\langle \psi | \phi \rangle|^2$$

where ψ, ϕ are purifications of p, σ .

Fact. $F(p, \sigma) = \|\sqrt{p} \sqrt{\sigma}\|_1^2$

Defn. Purified Distance.

$$P(p, \sigma) = \sqrt{1 - F(p, \sigma)}.$$

Defn. Smoothed Divergence.

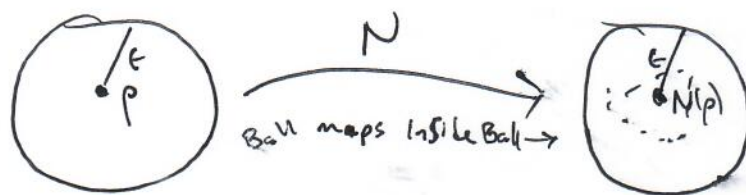
$$D^\epsilon(p, \sigma) = \inf_{\substack{p' \text{ st.} \\ P(p', \sigma) \leq \epsilon}} D(p' \| \sigma).$$

Lemma. If $\mathbb{D}$ is a divergence, then
the smoothed $\mathbb{D}^\varepsilon$ is fco.

(5)

Proof. We only need to use the fact that
the smoothing balls B_ε are defined wrt
a divergence. (We have specified $\mathbb{D}(p, \sigma)$
but it doesn't matter.) Only thing to check is DPI.

Let $\mathbb{D}(p \| \sigma) \neq \mathbb{D}(N(p) \| N(\sigma)) \forall p, \sigma$.



Then

$$\mathbb{D}^\varepsilon(N(p) \| N(\sigma)) = \inf_{p' \in B_\varepsilon(N(p))} \mathbb{D}(p' \| N(\sigma))$$

$$\leq \inf_{p' \in N(B_\varepsilon(p))} \mathbb{D}(p' \| N(\sigma))$$

$$= \inf_{p' \in B_\varepsilon(p)} \mathbb{D}(N(p') \| N(\sigma))$$

$$\leq \inf_{p' \in B_\varepsilon(p)} \mathbb{D}(p' \| \sigma) \\ = \mathbb{D}^\varepsilon(p \| \sigma).$$

So for any generalized divergence we have ⑥

• Entropy

$$H_D(p) = -D(p \| \mathbb{Q})$$

• Conditional Entropy

~~$H_D(A|B)$~~

$$H_D^{\uparrow}(A|B) = -D(p_{AB} \| \pi_A \otimes p_B)$$

$$H_D^{\downarrow}(A|B) = \cancel{-D(p_{AB} \| \pi_A \otimes p_B)} \\ = \inf_{\sigma_B} D(p_{AB} \| \pi_A \otimes \sigma_B)$$

• Mutual Info

$$I_D^{\uparrow}(A:B) = D(p_{AB} \| p_A \otimes p_B)$$

$$I_D^{\downarrow}(A:B) = \inf_{\sigma_B} D(p_{AB} \| p_A \otimes \sigma_B)$$

As well as ϵ -smoothed versions.

Why smooth and infimize?

(Note $\mu_{\text{train}} \leftrightarrow \text{Data}$)

Thm. Generalized Mutual Info is monotone under local channels,

⑦

$$I_D^{\uparrow}(A:B)_{\rho_{AB}} \geq I_D^{\uparrow}(A:B)_{(N_A \otimes N'_B) \rho_{AB}}$$

$$(\downarrow \geq \downarrow).$$

Proof. ~~Recall ρ_{AB} is a state on $\mathcal{H}_A \otimes \mathcal{H}_B$ including $\mathcal{H}_B = \mathcal{H}_B$~~

1P $\rho'_{AB} = N_A \otimes N'_B(\rho_{AB})$ then $\rho'_A = N_A(\rho_A)$, $\rho'_B = N'_B(\rho_B)$.

B, DPI

$$D(\rho'_{AB} \| \rho'_A \otimes \rho'_B) = D(N_A \otimes N'_B(\rho_{AB}) \| N_A \otimes N'_B(\rho_A \otimes \rho_B))$$

$$\leq D(\rho_{AB} \| \rho_A \otimes \rho_B) \checkmark$$

for minimized

$$\inf_{\sigma_B} D(\rho'_{AB} \| \rho'_A \otimes \sigma_B) \leq \inf_{\sigma_B} D(\rho'_{AB} \| \rho'_A \otimes N'_B(\rho_B))$$

$$= \inf_{\sigma_B} D(N_A \otimes N'_B(\rho_{AB}) \| N_A \otimes N'_B(\rho_A \otimes \sigma_B))$$

$$\leq \inf_{\sigma_B} D(\rho_{AB} \| \rho_A \otimes \sigma_B) \checkmark$$

Thm. Suppose V_A is a channel s.t. $V_A(\mathbb{I}_A) = \mathbb{I}_A$.

Then $H_D^{\uparrow}(A|B)_{V_A \otimes N_B(\rho)} \geq H_D^{\uparrow}(A|B)_{\rho}$.

Proof. Left as exercise.

Some particular relations

(8)

We define

$$D_\alpha(p||\sigma) = \frac{1}{\alpha-1} \log \text{tr } p^\alpha \sigma^{1-\alpha}$$

$$\tilde{D}_\alpha(p||\sigma) = \frac{1}{\alpha-1} \log \text{tr} \left(\left(\sigma^{\frac{1-\alpha}{2\alpha}} p \sigma^{\frac{1-\alpha}{2\alpha}} \right)^\alpha \right)$$

$$D(p||\sigma) = \text{tr}(p \log p - p \log \sigma)$$

$$D_{\max}(p||\sigma) = \inf_{p \leq \lambda \sigma} \lambda$$

Aside, think of $D_{\max}$ as

$$\sum p_i \log \frac{p_i}{q_i} \leq \log \max_i \frac{p_i}{q_i} = \log \max_i d_i$$

$$= \inf_{\lambda} \lambda$$

As defined, $D_\alpha, \tilde{D}_\alpha$ make sense for $\alpha \in (0,1) \cup (1,\infty)$.

But, Fact.

$$D_0(p||\sigma) = \lim_{\alpha \rightarrow 0} D_\alpha(p||\sigma) = D_h(p||\sigma)$$

$$\tilde{D}_\infty(p||\sigma) = \lim_{\alpha \rightarrow \infty} \tilde{D}_\alpha(p||\sigma) = D_{\max}(p||\sigma)$$

$$D_1 = \tilde{D}_1 = D \quad (\lim_{\alpha \rightarrow 1}).$$

Fact. D_α is G.D. for $\alpha \in [0, 2]$.

(9)

$\tilde{D}_\alpha$ is G.D. for $\alpha \in [\frac{1}{2}, \infty)$.

Lemma. $D_{\max}$ and $D_{\max}^\epsilon$ are G.D.'s.

Proof. Since $\lambda\sigma \geq p$, $\lambda \text{tr} \sigma \geq \text{tr} p \Rightarrow \lambda \geq 1$.
Thus $D_{\max} \geq 0$, $D_{\max} \in \mathbb{R} \cup \{\infty\}$.

Monotonicity:

Let $\lambda\sigma - p \geq 0$.

Then $N(\lambda\sigma - p) = \lambda N(\sigma) - N(p) \geq 0$ by CP.

Thus $\lambda N(\sigma) \geq N(p)$.

So $\inf_{p \leq \lambda\sigma} \log \lambda \geq \inf_{N(p) \leq \lambda N(\sigma)} \log \lambda$.

Fact. Araki-Lieb-Thirring inequality, for $A, B \geq 0$ states

$$\text{tr}[(ABA)^\alpha] \leq \text{tr}[A^\alpha B^\alpha A^\alpha] \quad (\alpha \geq 1)$$

$$\text{tr}[(ABA)^\alpha] \geq \text{tr}[A^\alpha B^\alpha A^\alpha] \quad (\alpha \leq 1)$$

Applying with $p=B$, $\sigma^{-\frac{1-\alpha}{\alpha}}=A$,

$$D_\alpha(p||\sigma) \leq \tilde{D}_\alpha(p||\sigma) \quad \forall \alpha.$$

• Show the "Entropy Zoo".

(12)

Semi-Definite Programming

• Khachi + Wilde 2.4:

- A class of linear optimization problems where you optimize over the cone of semi-definite matrices.
- Can be efficiently numerically solved.
- Arises in many pattern info problems.

Def. Given a Hermiticity preserving superoperator $\mathcal{I}$ and Hermitian operators A, B , an SDP consists of two optimization problems.

~~REDACTED~~

Primal SDP

(11)

$$\begin{array}{ll} \text{maximize} & \text{tr}(AX) \\ \text{subject to} & \Phi(X) \leq B \\ & X \geq 0. \end{array}$$

Dual SDP

$$\begin{array}{ll} \text{minimize} & \text{tr}(BY) \\ \text{subject to} & \Phi^*(Y) \geq A \\ & Y \geq 0 \end{array}$$

Let

$$S(\Phi, A, B) = \sup_{\substack{X \geq 0 \\ \Phi(X) \leq B}} \text{tr}(AX)$$

$$\hat{S}(\Phi, A, B) = \inf_{\substack{Y \geq 0 \\ \Phi^*(Y) \geq A}} \text{tr}(BY)$$

Prop. $S(\Phi, A, B) \leq \hat{S}(\Phi, A, B)$ (Weak Duality)

Strong Inequality: Often they are equal (see below).

(12)

Examples:

These are tricky examples, they're not obvious. ~~see~~
See notes:

- $D_n^t(p||\sigma)$
- $I_h^t(A:B)_p$
- $H_{\min}(A|B)_p$
- Trace Norm

Useful: Evaluate efficiently. Convert min $\leftrightarrow$ max.

Max-Min Inequality

$$\inf_x \sup_y F(x, y) \geq \sup_y \inf_x F(x, y)$$

Smallest grant $\geq$ biggest dwarf

(Sometimes equal - minimax theorems).

• Homework Q?

Lecture 9

①

Today

- Continue SDP
- AEP

(Comment: Why smooth & minimize?)

SDP

(See Previous Notes).

AEP (Asymptotic Equipartition Property).

Many divergences converge to $D(p||\sigma)$ in the large $N^{\text{asymptotic}}$ limit.

Theorem. (Quantum Stein's Lemma).

For any p, σ states, any $\epsilon \in (0, 1)$,

$$\lim_{n \rightarrow \infty} \frac{1}{n} D_n^{\epsilon}(p^{\otimes n} || \sigma^{\otimes n}) = D(p || \sigma).$$

Proof. Use R\'enyi bounds

②

$$D_{\alpha}^{\epsilon}(p||\sigma) \geq D_{\alpha}(p||\sigma) - \frac{\alpha}{\alpha-1} \log \epsilon \quad (\alpha \in (0,1))$$

$$D_{\alpha}^{\epsilon}(p||\sigma) \leq \tilde{D}_{\alpha}(p||\sigma) - \frac{\alpha}{\alpha-1} \log(1-\epsilon) \quad (\alpha \in (1,\infty))$$

First one, see KKL.

Second follows because of DPI.

$$D_{\alpha}(p||\sigma) \geq D_{\alpha}(N(p)||N(\sigma))$$

where $N(p) = \text{tr}(p) \text{diag}(c_1, \dots, c_n)$ is the channel of a hypothesis test.

Now use that

$$D_{\alpha}(p^{\otimes n}||\sigma^{\otimes n}) = \frac{1}{\alpha-1} \log \text{tr} \left(p^{\otimes n} (\sigma^{\otimes n})^{1-\alpha} \right)$$

$$= \text{tr} \left((p^{\alpha} \sigma^{1-\alpha})^{\otimes n} \right)$$

$$= \text{tr} \left((p^{\alpha} \sigma^{1-\alpha})^{\otimes n} \right)$$

$$= \left(\text{tr} (p^{\alpha} \sigma^{1-\alpha}) \right)^n$$

$$= n D_{\alpha}(p||\sigma) \quad (\text{additivity})$$

So

$$\lim_{n \rightarrow \infty} \frac{1}{n} D_{\alpha}^{\epsilon}(p||\sigma) \geq D_{\alpha}(p||\sigma) - \frac{1}{n} \frac{\alpha}{\alpha-1} \log \epsilon$$

Limit as $\epsilon \rightarrow 0$ gives proof.

Another Thm. For $t \geq 0$,

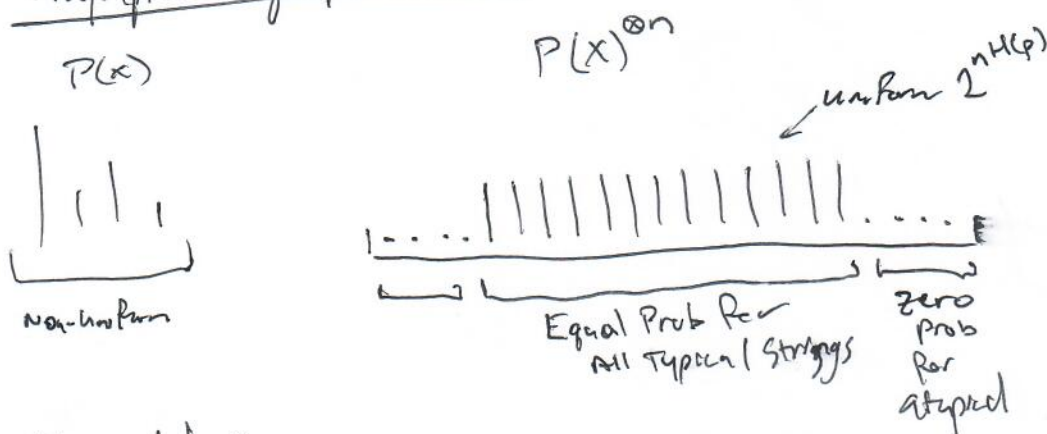
(3)

$$\lim_{n \rightarrow \infty} \frac{1}{n} D_{\max}(p^{\otimes n} \| \sigma^{\otimes n}) = D(p \| \sigma).$$

Pf. Another funny α -proof (see notes).

• But why does this happen?

"Asymptotic Equipartition"



Typicality:

Non uniform dist $\xrightarrow{\otimes n}$ Uniform Distribution over typical ~~strings~~ subspace.

For quantum state w/ eigenbasis

$$\rho = \sum_x p(x) |x\rangle \langle x|$$

$$\begin{aligned} \rho^{\otimes n} &= \sum_{x_1, x_2, \dots} p(x_1) p(x_2) \dots |x_1, x_2, \dots\rangle \langle x_1, x_2, \dots| \\ &= \sum_{x_n} p(x_n) |x_n\rangle \langle x_n| \quad (\text{mixture of strings}) \end{aligned}$$

Project into typical subspace

(4)

$$\Pi_n = \sum_{\text{typical } x_n} |x_n\rangle\langle x_n|$$

Then $\rho^{\otimes n} \approx \Pi_n \rho^{\otimes n} \Pi_n \approx \frac{\Pi_n \rho^{\otimes n} \Pi_n}{\text{tr}(\Pi_n \rho^{\otimes n} \Pi_n)}$ b/c

you've barely gotten rid of any probability.
(Gentle measurement lemma.)

So always

$$\rho^{\otimes n} \approx \sum_{\text{typical } x_n} \frac{|x_n\rangle\langle x_n|}{2^{nH(p)}}$$

is essentially uniform on typical subspace.

- So the question is, how do these entropies look to uniform distributions?

~~Suppose~~

Suppose that

$$p = \frac{\pi}{\text{tr} \pi} = \frac{\pi}{W} \quad \sigma = \frac{\pi'}{W'} = \frac{\pi'}{\text{tr} \pi'}.$$

Let's assume $\pi \leq \pi'$ (support), otherwise everything ∞ .

(5)

QRE

$$D(p||\sigma) = \text{tr} \left(p \log \frac{\pi}{w} - p \log \frac{\pi'}{w'} \right)$$

Restrict to support π , everywhere else zero,
 On support π , $\pi' = \pi = \pi$.

$$D(p||\sigma) = \text{tr} \left(p \log \frac{\pi}{w} - p \log \frac{\pi}{w'} \right)$$

$$= -\text{tr} p \log w + \text{tr} p \log w'$$

$$\boxed{D(p||\sigma) = \log \frac{w'}{w}} \quad \text{For uniform distr.}$$

Renyi (The p, σ commute so either one.)

$$D_\alpha(p||\sigma) = \frac{1}{\alpha-1} \log \text{tr} (p^{\alpha} \sigma^{1-\alpha})$$

$$= \frac{1}{\alpha-1} \log \text{tr} \left(\frac{\pi^{\alpha}}{w^{\alpha}} \frac{\pi^{1-\alpha}}{w'^{1-\alpha}} \right) \quad [\text{Restrict to Support}]$$

$$= \frac{1}{\alpha-1} \log \frac{\text{tr} \pi}{w^{\alpha}} \frac{1}{w'^{1-\alpha}}$$

$$= \frac{1}{\alpha-1} \log \frac{w^{1-\alpha}}{w'^{1-\alpha}}$$

$$\boxed{D_\alpha(p||\sigma) = \log \left(\frac{w'}{w} \right)}$$

Hyp. Test

6

$$D_h^t(p||\sigma) = -\log \inf_{\text{tr}(M\frac{\Pi}{W}) \geq 1-\epsilon} \text{tr}(M\frac{\Pi'}{W'})$$

For feasible M , $\text{tr}(M\Pi) \geq W(1-\epsilon)$. But $\Pi' \geq \Pi$, so

$$\text{tr}\left(\frac{M\Pi'}{W'}\right) \geq \frac{\text{tr}(M\Pi)}{W'} \geq \frac{W}{W'}(1-\epsilon).$$

on the other hand, Π is feasible ($\text{tr}(\Pi) = 1$), and

$$\text{tr}(\Pi\sigma) = \text{tr}\left(\frac{\Pi\Pi'}{W'}\right) = \frac{\text{tr}(\Pi)}{W'} = \frac{W}{W'}.$$

so

$$\log \frac{W'}{W} \leq D_h^t(p||\sigma) \leq \log \frac{W'}{W} - \log(1-\epsilon)$$

$\nearrow$ scales like n $\uparrow$ const

In asymptotic setting, $\frac{W'_n}{W_n} = \left(\frac{W'}{W}\right)^n = \frac{\left(\frac{2^{H(\sigma)}}{2^{H(p)}}\right)^n}{\left(\frac{2^{H(\sigma)}}{2^{H(p)}}\right)^n}$

dividing by n gives

$$\log\left(\frac{2^{H(\sigma)}}{2^{H(p)}}\right) \leq \frac{1}{n} D_h^t(p||\sigma) \leq \log\left(\frac{2^{H(\sigma)}}{2^{H(p)}}\right) - \frac{1}{n} \log(1-\epsilon)$$

$$\frac{1}{n} D_h^t(p||\sigma) \rightarrow D(p||\sigma).$$

Max.

7

$$D_{\max}(p||\sigma) = \inf_{p \perp \sigma} \log \lambda.$$

Restrict to support (for other states all $\lambda \geq 0$ valid).

$$\frac{\pi}{w} \leq \lambda \frac{\pi}{w'} \Rightarrow \lambda \geq \frac{w'}{w}. \text{ Also } \lambda = \frac{w'}{w} \checkmark$$

$$\text{So } D_{\max}(p||\sigma) = \log \frac{w'}{w}$$

~~Fidelity~~ Fidelity.

$$F(p, \sigma) = \|\sqrt{p}\sqrt{\sigma}\|_1^2 = \left(\text{tr} \left| \frac{\pi}{w^{1/2}} \frac{\pi'}{w'^{1/2}} \right| \right)^2$$
$$= \left(\text{tr} \left| \frac{\pi}{w^{1/2}} \right| \right)^2 = \left(\frac{w^{1/2}}{w'^{1/2}} \right)^2 = \frac{w}{w'}$$

$$-\log F(p, \sigma) = \log \frac{w'}{w}$$

(note: This is $\tilde{D}_{\frac{1}{2}}(p||\sigma)$).

*Not rigorous we assumed π, π' commute, etc.

Summary. (Roughly Speaking)

• These all converge when applied to uniform p, σ , up to constants (independent of dimension of the inputs)!

• Asymptotic limit makes everything uniform!

~~scribbles~~

With extra time:

8

- Resource Theories
- Names & Distances
- Questions

Today

- Compression
- State Merging

Smooth Conditional Min/Max Entropy

~~Smooth Conditional Min/Max Entropy~~

$$H_{\min}(A|B)_\rho \equiv -\inf_{\sigma_B} \tilde{D}_\infty(p_{AB} \| \mathbb{1}_{A \otimes \sigma_B})$$

$$H_{\max}(A|B)_\rho \equiv -\inf_{\sigma_B} \tilde{D}_{1/2}(p_{AB} \| \mathbb{1}_{A \otimes \sigma_B})$$

$$H_{\min}^\epsilon(A|B)_\rho \equiv -\inf_{\sigma_B} D_\infty^\epsilon(p_{AB} \| \mathbb{1}_{A \otimes \sigma_B})$$

$$\nearrow H_{\max}^\epsilon(A|B)_\rho \equiv \inf_{\tilde{\rho} \in B_\epsilon(\rho)} H_{\max}(\tilde{\rho})$$

Breaks the pattern.

For any pure $|\psi_{ABR}\rangle$,

$$H_{\max}^\epsilon(A|B)_{\rho_{AB}} = -H_{\min}^\epsilon(A|R)_{\rho_{AR}}.$$

(Intuition...)

$$\begin{aligned} H(AB) - H(B) &= H(R) - H(AR) \\ &= -(H(AR) - H(R)). \end{aligned}$$

What are these? Trivial Bsystem...

(2)

$$H_{\min}(A)_{p_A} = -D_{\infty}(p_A \| \mathbb{I}_A)$$

$$= -\log \min_{\lambda} \lambda$$

$$p \leq \lambda \mathbb{I}$$

$$= -\log \|p\|_{\infty} \quad \leftarrow \text{Largest E.V.}$$

~~$$H_{\min}(A)_{p_A} = -\log \frac{1}{\|p\|_{\infty}}$$~~

$$H_{\min}^f(A)_{p_A} = -\inf_{\tilde{p} \in \mathcal{B}(p)} \log \|\tilde{p}\|_{\infty} \quad \leftarrow \text{Smallest largest E.V.}$$

$$= \sup_{\tilde{p} \in \mathcal{B}(p)} (-\log \|\tilde{p}\|_{\infty})$$

$$H_{\max}(A)_p = -\tilde{D}_{\chi^2}(p \| \mathbb{I})$$

$$= 2 \log F(p, \mathbb{I})$$

$$= 2 \log \text{tr} \sqrt{p}$$

$$\leq \log(\text{rank}(p))$$

$$= -D_0(p \| \mathbb{I}) \quad [\text{Alt. Property Def.}]$$

$$\text{tr} \sqrt{p} = \sum_{i=1}^r \sqrt{\lambda_i}$$

$$\leq \sqrt{\sum_{i=1}^r \lambda_i} \sqrt{\sum_{i=1}^r 1}$$

$$= \sqrt{r}$$

$$H_{\max}^f(A)_p = \inf_{\tilde{p} \in \mathcal{B}(p)} 2 \log \text{tr} \sqrt{\tilde{p}}$$

$$\leq \inf_{\tilde{p} \in \mathcal{B}(p)} \log \text{rank} \tilde{p} \quad \leftarrow \text{Compression}$$

Entanglement: Schmidt Coefficients

(3)

Bipartite pure entanglement is very simple.

• Every $|\psi_{AB}\rangle$ has Schmidt decomposition

$$|\psi_{AB}\rangle = \sum_{k=1}^r \sqrt{\lambda_k} |K\rangle_A |K\rangle_B$$

$\uparrow$ Schmidt coeffs!

• The Schmidt basis is also the reduced density matrix (ρ_A, ρ_B) eigenbasis,
and $\{\lambda_k\}$ are reduced density matrix E.V.s.

$$\rho_A = \sum_{k=1}^r \lambda_k |K\rangle_A \langle K|, \quad \rho_B = \sum_{k=1}^r \lambda_k |K\rangle_B \langle K|.$$

The Schmidt coeffs $\{\lambda_k\}$ completely
determine the entanglement structure.

→ Asymptotic Distillable Entanglement $H(\{\lambda_k\}) = S(\rho_A)$

→ Asymptotic Entanglement Cost $S(\rho_B)$

→ LOCC Interconvertibility.

④

Interconvertibility: (Nielsen Thm).

Thm. It is possible to ~~transform~~ transform $|\psi\rangle \rightarrow |\phi\rangle$ using LOCC if and only if the ψ schmidt coeffs $\{\lambda_k^\psi\}$ are majorized by the ϕ schmidt coeffs $\{\lambda_k^\phi\}$.

Let P majorize q . Then P is more concentrated.



Def. P majorizes q if

$$\sum_{i=1}^k p_i^\downarrow \geq \sum_{i=1}^k q_i^\downarrow \quad \forall k,$$

where $p_i^\downarrow, q_i^\downarrow$ are ordered (large to small).

- (if means: Fastest-growing cumulative distribution is always bigger.

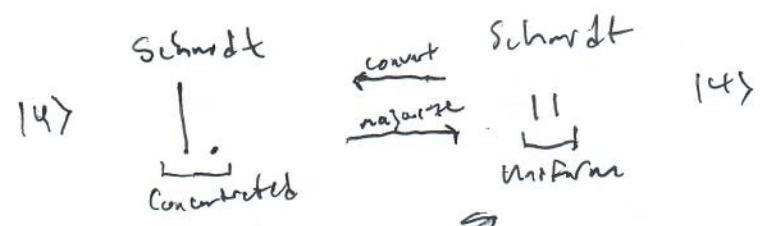
~~Order~~ (Original Order doesn't matter.)

→ Can be accomplished with "binary equating" operations!

5

Product State
 $| \psi_A \rangle \otimes | \psi_B \rangle$

Bell Pair
 $| \psi_{AB} \rangle = \frac{1}{\sqrt{2}} (| 00 \rangle + | 11 \rangle)$



Uniform coeffs $\Rightarrow$ More Entangled

~~LOCC~~ LOCC decreases Entanglement

New $| \psi \rangle \xrightarrow{\text{majorizes}} \text{Old } | \psi \rangle$

("odd majorized by new")

Cor. LOCC (Pure $\rightarrow$ Pure) cannot increase Schmidt rank.

PF. Schmidt Rank = # nonzero λ_k .

Let $| \psi \rangle$ have ~~rank~~ Schmidt rank r_A and $| \psi \rangle, r_B$.

If ~~$r_A > r_B$~~ then we need



$$\sum_{i=1}^k \lambda_i^4 \leq \sum_{i=1}^k \lambda_i^4 \quad \forall k.$$

But for $k = r_A < r_B$, this is

$$1 \leq \sum_{i=1}^{r_A} \lambda_i^4 < 1$$

which is false!

Cor. Let $|\psi_{AB}\rangle$ be pure state and $\textcircled{6}$
 let $|\Phi_K\rangle$ be maximally entangled state
 of Schmidt rank K . Then

$$|\psi_{AB}\rangle \xrightarrow{\text{LOCC}} |\Phi_K\rangle$$

if and only if $\lambda_{\max} \leq \frac{1}{K}$, where
 $\lambda_{\max}$ is biggest schmidt coeff.

That is, from $|\psi_{AB}\rangle$ you can 1-shot
 distill ~~perfect Bell pairs~~ $\log K = -\log \lambda_{\max} = H(p_A)$
 perfect Bell pairs.

Exercise. Use this to prove asymptotic distillation
 rate is $H(p_A)$.

Proof. For all l ,

$$\sum_{i=1}^l \lambda_i \leq l \lambda_{\max} \leq l \frac{1}{K} = \sum_{i=1}^l \frac{1}{K}$$

if $\lambda_{\max} \leq \frac{1}{K}$.

⑦

LOCC: Kraus Operators form
conditional chans of local ops.

$$K_{ij} = (\mathbb{1}_A \otimes K_{ji}^B) (K_i^A \otimes \mathbb{1}_B)$$

depending on
long LO/LOCC/
SEP/PPT
target

For each fixed i , $\{K_{ji}^B\}_j$ is a set
of Kraus operators. ("conditional
on measurement outcome from A")

Lemma. A product operator $A \otimes B$ acting on
a ~~separable~~ pure state $|\psi\rangle$ cannot
increase its Schmidt rank.

PF. Schmidt decomp is really SVD.

$$|\psi\rangle = \sum_{ij} Y_{ij} |i\rangle \otimes |j\rangle$$

$\nwarrow \lambda_k = \text{Sing. Vals.}$

$$r_\psi = \text{Rank } Y$$

But

$$A \otimes B |\psi\rangle = \sum_{\substack{i,j \\ i',j'}} Y_{ij} A_{ii'} \otimes B_{jj'} |i'\rangle \otimes |j'\rangle$$

$$= \sum_{ij} (AYB^T)_{ij} |i\rangle \otimes |j\rangle,$$

And $\text{rank}(AYB^T) \leq \text{rank}(Y)$ by lm. alg.

8

Cor. ~~Any state~~ If $\rho = \sum_i p_i |\psi_i\rangle\langle\psi_i|$ such that $p_i \in \mathbb{R}$, and Λ is LOCC, then

$$\Lambda(\rho) = \sum_j q_j |\psi_j\rangle\langle\psi_j|$$

where $q_j \in \mathbb{R}$.

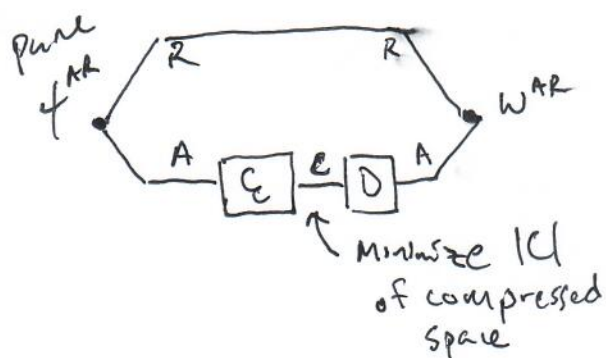
Def. Kraus Operators are strings of products.

Finally :).

Now on to compression.

Compression

9



One-Shot:
 • Allow error ϵ .
 Want $W_{AR} \approx \psi_{AR}$

Def. A compression scheme is a code $(\mathcal{C}^{A \rightarrow C}, D^{C \rightarrow A})$. It maps

$$W_{AR} = \bigotimes_{i=1}^R \mathcal{C}^{A \rightarrow C}(\psi_{AR}).$$

We say it has error ϵ if $P(w, \psi) \leq \epsilon$ in the purified distance, and rate $R = \log |\mathcal{C}|$. The ϵ -one-shot compressibility of ψ is

$$l_\epsilon(\psi) = \min_{\substack{(C, D) \\ \text{error} \leq \epsilon}} R.$$

We'll show the lower bound is the smooth max entropy.

Thm. (compressibility const):

$$L_\epsilon(\psi) \geq H_{\max}^\epsilon(p^A)$$

(Compare to Schumacher $L(\psi) = S(p^A)$).

Proof.

- The intermediate state $\rho_{CR} = \mathcal{E}(\psi)$ has E.R.

$$\rho_{CR} = \sum_i p_i |\psi_i\rangle \langle \psi_i|$$

where $|\psi_i\rangle$ has S.R. $\leq K$. Therefore

- So does

$$w_{AR} = \sum_j q_j |\xi_j\rangle \langle \xi_j|$$

with $|\xi_j\rangle$ S.R. $\leq |C|$, by earlier

lemma, since $D^{\text{RSA}}(\psi_{CR})$ is local.

- But since $P = \sqrt{1-F}$ we have

$$1-F^2 \leq F(\psi, w) = \text{tr}(fw) = \sum p_i \text{tr}(\psi \xi_i)$$

$$\leq \max_{\substack{\text{S.R.} \leq |C| \\ \psi}} F(\psi, \psi)$$

$$= \max_{\substack{\sigma^A \\ \text{rank } \sigma \leq |C|}} F(p^A, \sigma^A) \quad \left[\text{B/c Rank } \sigma_A = \text{S.R. } \psi \right]$$

~~Thus, $\exists$ a state σ with~~

Thus $\exists \sigma$, $\text{rank} \leq |C|$, with

$$P(p^A, \sigma) \leq \epsilon \quad (F(p^A, \sigma) \geq 1 - \epsilon^2).$$

• Thus,

$$H_{\max}^t(p) = \min_{\tilde{p} \in \mathcal{B}(p)} H_{\max}(\tilde{p}) \leq H_{\max}(\sigma)$$

And earlier we saw

$$H_{\max}(\sigma) = 2 \log \text{Ar} \sqrt{\sigma} \leq \log \text{rank } \sigma$$

$$\text{So } (\quad) \leq \dots \leq \log |C| \checkmark$$

~~Adversity~~
~~Let $H_{\max}^t(p) \leq \log |C|$~~

Achievability

12

Consider

$$M_H^s(p) = \log \min_{\text{tr } M \geq 1-s} \text{tr } M.$$

if M is s.p.h. $\text{tr } M \geq 1-s$, then also

$$M' = \sum_n |n\rangle \langle n| M |n\rangle \langle n|$$

$\nearrow$ P eigenvectors
"pinching"
"dephasing"
think about pinching?

~~For fixed M , optimal M' is~~

always $\text{tr}(M/p) = \text{tr}[M' \sum_n |n\rangle \langle n| P |n\rangle \langle n|] = \text{tr}(M' P).$

• So we only need to consider M' that are diagonal in P eig. basis.

• Let $P = \sum_{k=1}^r \lambda_k |k\rangle \langle k|$ with $\lambda_0 = \lambda_1 = \dots$

For fixed $\text{tr } M$, optimal M' is

$$M' = \sum \mu_i |i\rangle \langle i|, \quad \mu_i = \underbrace{\begin{matrix} | & | & | & | & | \\ \mu_1 & \mu_2 & \mu_3 & \mu_4 & \mu_5 \end{matrix}}_{r \text{ terms}} \dots$$

Let P project onto this subspace for optimal M .

$$\text{tr } M \leq \text{tr } P \leq \text{tr}(M) + \epsilon$$

Now compress to this subspace

$$C(P) = P P P + \sigma(I - P) P (I - P).$$

Decompress w/ identity isometry.

$$D \circ C(P) = \text{above}$$

(13)

• The output

$$w_{AR} = D_{OL}(\psi_{AR}) \geq \cancel{D_{AR}}(P \otimes Q) \psi_{AR}(P \otimes Q_P)$$

so

$$\begin{aligned} \text{tr} \psi_w &\geq \text{tr}(\psi(Q \otimes P) \psi(Q \otimes P)) \\ &= \cancel{\text{tr}(\psi(Q \otimes P) \psi(Q \otimes P))} |\langle \psi | P \otimes Q | \psi \rangle|^2 \\ &= |\text{tr}(\psi(P \otimes Q))|^2 \\ &= (\text{tr}(P \rho_A))^2 \\ &\geq (\text{tr}(M \rho_A))^2 \\ &\geq (1-s)^2. \end{aligned}$$

Let $1-e^2 = (1-s)^2$ to get purified distance bound.

Thus $\exists C$ with $|e| \leq \text{tr} P \leq \text{tr} w + 1$.

$$\boxed{l_e(\psi) \leq H_n^{1-\sqrt{1-e^2}}(\rho_A) + 1}.$$

Also a similar lower bound

$$H_n^{2(1-e^2)}(\rho_A) \leq l_e(\psi).$$

Lecture 11

①

Today:

- Compression
- State merging/QC

Compression:

↳ Finish previous lecture.

Compression $\rightarrow$ Teleportation.

- Compressibility $h_c(\rho)$ is also the number of ebits you need to teleport ρ from $A \rightarrow B$!

~~• If you have R ebits for a teleportation scheme available.~~

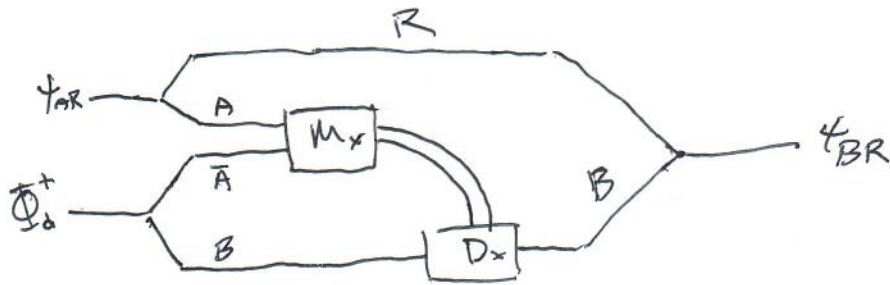
- Recall teleportation.



M_X = Bell basis measurement
 D_X = Pauli correction

- You can do the same for Qdts + Pcs...

(2)

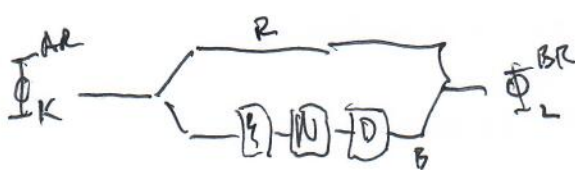


Schmidt rank (Kebits) of Φ_d^p is $= |A|$ for tws to work.

$$\psi_{AR} \xrightarrow[\text{to Bob}]{\text{compress}} \tilde{\psi} \xrightarrow[\text{Kebits}]{\text{Teleport}} \psi'_{AR}$$

- Quantum Communication =

Entanglement Distillation / Transmission using N + Teleportation (Perfect Channel)



[Picture Above]

Entanglement Transmission (Quantum Capacity)

(Alman) Put

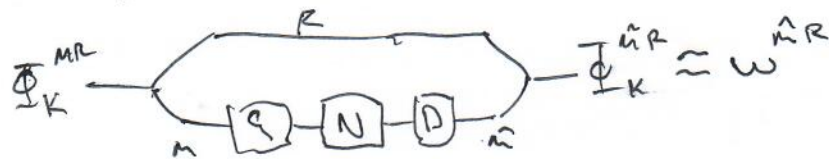
generates state ψ_{ABE}

State Merging
- Distillation
- Teleportation
LOCC

③

1. Q Comm (Entanglement Transmission)

Def. An ϵ -one-shot quantum error correcting code is a pair of encoding/decoding maps



Def. Error:

$$P_{err} = P(\Phi^{MR}, \omega^{\hat{M}R})$$

Rate: $R = \log K$

Capacity: $Q_\epsilon(N) = \max_{P \leq \epsilon} R$

* Many variants! See ~~book~~ K+L for comparisons.

~~Another Def~~

A general result:

Thm. $Q_\epsilon(N) \geq \max_{\phi_{AA'}} H_{\min}^{\epsilon/5}(A|E)_\phi - 2 \log \frac{5}{\epsilon^2}$

where

$$\phi_{A|BE} = V_N^{A \rightarrow BE}(\phi)$$

is Stinespring of N .

& Convergence = asymptotic value



Conditional min-entropy
 $\Rightarrow I(A \rangle B)$

asymptotically.

* But not 1-shot tight.

Proof?

- See notes: Uses postprocessing by S.M., then teleport.
- Also see Kew for other protocols.
- Notes proof uses

$$Q_E(N) \geq \max_{\Phi_{AA'}} (-\tilde{\mathcal{L}}(\Psi))$$

↑ state energy capacity
of Ψ_{ABE} .

Examples:

(1) Identity $\mathbb{I}_A$

- Clearly $R = \log |A|$ achieved w/ zero error.
 - Also $F(\mathbb{I}_A, \omega) \leq \frac{|A|}{K}$ by the S. Rank argument.
- $$\log |A| \leq Q_c(\mathbb{I}_A) \leq \log |A| - \log(1 - e^2)$$

(2) Entanglement Breaking (Separable output)

- Output ω is mixture of S. Rank 1.

$$F(\mathbb{I}_A, \omega) \leq \frac{1}{K}.$$

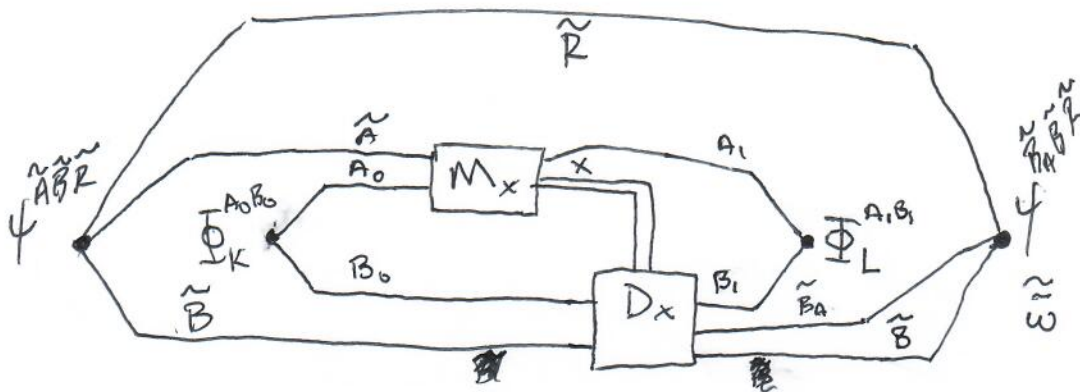
(Low overlap of reduced density matrices).

(3) Antidegradable

Stronger $\rightarrow$ No cloning $\rightarrow$ Low Capacity

Quantum State Merging

(5)



Def. A Q.S.M. Scheme for pure $\psi^{\tilde{A}\tilde{B}\tilde{R}}$ is a joint ^{LOCC} operation

$$\Lambda: \tilde{A}A_0 \otimes \tilde{B}B_0 \rightarrow A_1 \otimes B_1 \tilde{A}_A \tilde{B}$$

with error ϵ ,

$$F(W_{A_1B_1\tilde{B}_A\tilde{B}\tilde{R}}, \Phi_L^{A_1B_1} \otimes \psi^{\tilde{B}_A\tilde{B}\tilde{R}}) \geq 1 - \epsilon$$

where

$$W_{A_1B_1\tilde{B}_A\tilde{B}\tilde{R}} \equiv \Lambda(\Phi_K^{A_0B_0} \otimes \psi^{\tilde{B}_A\tilde{B}\tilde{R}})$$

with rate (entanglement cost)

$$R = \log K - \log L.$$

* What does cost mean if spaces change? ⑥

The smallest achievable rate

$$\tilde{L}_t(\psi) = \inf_{P_{E \leq t}} R$$

is the average cost -

* Different
cost
rates

• Gives meaning to negative conditional entropies.

↙ +Rate
Need to input
Entanglement
(Teleportation)

↘ -Rate
Merge +
get entanglement
as output
(distillation).

Examples:

- Note: By teleporting A using Φ_{AB} ,
generating nothing, $\tilde{L}_t(\psi) \leq \log |A|$.

Examples...

⑦

~~① ψ_{ABR}~~

(1) $\psi_{ABR} \approx \psi_{AR}$.

- Compress + Teleport!

$$\tilde{L}_E \leq L_E$$

- Schmidt Rank Argument for
same converse as L_E .

(Teleportation)

(2) $\psi_{ABR} = \psi_{AB}$

- Bob creates ψ_{AB} locally

- Use original ψ_{AB} to distill

- Entanglement generated from ψ_{AB}
is $H_{\min}^E(P_A)$ using smoothed
version of arguments from before

(Distillation)

(3) GHZ

$$|\psi_{ABR}\rangle = \frac{1}{\sqrt{2}}(|000\rangle + |111\rangle)$$

$$R \approx 0 \dots$$

Teleportation / Distillation cancel out.

The general picture:

(8)

Then.

x Note after G.

~~Then~~

$$H_{\max}^{\epsilon}(A|B) \leq \tilde{H}_{\epsilon}(\epsilon) \leq -H_{\max}^{\epsilon/2}(A|B)_{\epsilon} + 2 \log \frac{5}{\epsilon^2}$$

(Recall Minimax duality).

Proof?

All relies on decoupling:

$$\rho_{A, \tilde{R}}^{\otimes K} \approx \tau_{A,1} \otimes \rho_{\tilde{R}}$$

After Alice's measurement.

↳ Necessary & sufficient

↳ Achieved using Randomness.

Lecture 12

①

Today:

- Recap
- Misc

Recap

Intro to
Q.I.T.



Adv
Q.I.T.

"One-Shot" Quantum Shannon Theory

Lecture 2

- Classical Information Theory
 - Compression: $H(P) = \text{min bits}$
 - Noisy Channel Coding Thm
 - ↳ Fundamental Limits on Communication
 - Relative Entropy is King
 - ↳ Monotonicity, Non-Negativity, Joint Convexity, Chain Rule
 - Information Can be Quantified!

Entropy

- ↳ Asymptotic Derivation of $H(p) = -\sum p \log p$ ②
- ↳ How many bits? (Variable length coding). Kraft-McMillan.
- ↳ Uncertainty
- Typicality $\rightarrow$ Typical strings $-\log p = n H(p)$
w/ prob. 1.
- Info in a book?

Lecture 3

- Classical Channels & Prob. Theory Notation.
- Markov Chains, Data Processing
- Noisy Channel Coding Theorem
 - At what rate can channel reliably send messages!

=

Channel Capacity

• Proof used

↳ Random Coding

↳ Jointly Typical Decoding

* Lecture 1

- Quantum Channels

• Classical Comm over Quantum Channels.

2.5

• Plain Codes

• Achievability: ~~Random Coding~~

• Random Coding

• SQRT Measurement Decoding

↳ Hayashida/Nagaoka Lemma

• Hypothesis Testing RE (Wang/Renner)

↳ Proved ϵ -one-shot capacity
~ ϵ -hypothesis X quantity.
(Q mutual info across channel).

QA Codes

• Achievability:

• Position Based Coding

• SQRT meas / ~~Hayashida~~ H-N Lemma

• Hyp testing @ each position

↳ Proved ϵ -one-shot capacity
~ ϵ -hyp testing mutual info

• Courses = Proved general meta-courses
based on monotonicity & comparator,
using generalized divergences.

• Asymptotic Limit

(3)

↳ Asymptotic Equipartition Property
i.e. Quantum Typicality.

↳ Proved ϵ -one-shot capacities
for $n \rightarrow \infty$ become
 $X_{\text{reg}}(N)$, $I(N)$.

~~Using~~ Using AEP for Hyp Testing

• The Zoo of Entropies

• Generalized Divergences

↳ Entropy, Mutual Info, Conditional

↳ Monotonicity rules all.

↳ G -Smoothing

σ_B -minimizing

↳ α -Perry's $\hat{\epsilon}$ minimax

• SDP's / minmax / Duality

• Entropy Zoo

- Quantum Compression & Communication
 - Smallest ~~state~~ Hilbert space p can be compressed to with ϵ -error.
 - ↳ one-shot compressibility $\sim H_{\max}^{\epsilon}(p)$
 - ↳ Duality of ϵ -smooth min/max entropies
 - ↳ Asymptotically: AEP $\rightarrow H(p)$
= Schumacher Compression

- Bottlenecks based on the Schmidt ~~rank~~ rank
- Majorization & LOCC-interconvertibility (Entanglement Theory)
- What is LOCC?
- Q. Comm = Entanglement Distillation + teleportation

↳ one-shot Distillation $H_{\max}^{\epsilon}$

↳ Quantum State Merging $I(A:B)$
 $= -H(A|B)$

③
↳ One-shot Quantum Capacity
in terms of entangled transmission
↳ Achievability w/ Q.S.M. $\rightarrow I(A:B)$

↳ Method for Q. Comm:
• Keep info out of (Steeping) Environment!
• Decouple Alice from Env.
• Distribute Ent. Then use LOCC.

A few principles:

- Information can be quantified!
Which entropic quantity fits the task?
- Relative Entropy Manifordity!
↳ Know whether your proof is specific, or
for the particular quantity -
↳ What property are you using!
- Revisit Entropy Zoo PDF

Finally, a little bit of

- Resource Theory
- Entanglement / Correlations

Resource Theory:

- Free States ($\mathcal{F}$)
- Free Operations ($\mathcal{F}$)
 - ↳ Leave free states invariant
- Relative Entropy Measures

$$E_R(p) = \inf_{\sigma \in \mathcal{F}} D(p \| \sigma)$$
- Resource monotone!

Ex. Purity.

Free State: Maximally mixed $\tau = \frac{I}{d}$

Free Ops: Unital

$$\Phi(\tau) = \tau$$

$$\begin{aligned} RE \text{ of Purity} &= D(p \| \frac{I}{d}) \\ &= \log d - H(p). \end{aligned}$$

Ex.

Resource: Total Correlations

Free States: Product States
 $\rho_A \otimes \rho_B$

Free Operations: Local Operations
 $\Lambda_A \otimes \Lambda_B$

RE of Correlations:

$$E(\rho) = \inf_{\sigma_A, \sigma_B} D(\rho_{AB} \| \sigma_A \otimes \sigma_B)$$

$$= D(\rho_{AB} \| \rho_A \otimes \rho_B)$$

$$= \text{Mutual Info!}$$

Ex.

Resource: Entanglement

Free States: Separable States

Free Ops: LOCC

$$RE: \quad E(\rho) = \inf_{\sigma \in \text{SEP}} D(\rho \| \sigma)$$

RE of entanglement!

Entanglement & Quantum Correlations

(8)

- Product State

$$\rho = \rho_A \otimes \rho_B$$

- mutual info

- Classically Correlated State (CC)

$$\rho = \sum_{ij} P_{ij} |i\rangle\langle i| \otimes |j\rangle\langle j| \quad (\text{Local O.N. Bases})$$

(classical prob dist)

- RE of Quantumness / Discord

- Separable State

$$\rho = \sum_k P_k \rho_k^A \otimes \rho_k^B$$

- RE of Entanglement

$$\rightarrow \rho = \frac{1}{2} (|00\rangle\langle 00| + |11\rangle\langle 11|)$$

is SEP but not CC.

- Entangled State

$$\rho = \text{Not Separable}$$

- Bell Nonlocal State

- Discard:
Difference
or
mutual
info

→ Different Operational
Tasks!

For Comm we say,
Ent. is long.