

نظریه اطلاعات کلاسیک - بخش دوم

وحیدکریمی پور- دانشکده فیزیک - دانشگاه صنعتی شریف

۲۴ اسفند ۱۳۹۳

۱ مقدمه

در فصل های قبل دیدیم که می توان پیام های یک منبع را که از کلمات $X = \{x_1, x_2, \dots, x_{2^n}\}$ تشکیل شده است را می توان با کدگذاری مناسب یعنی کدگذاری بلوکی چنان مخابره کرد که بجای n بیت به ازای هر کلمه، $nH(X) \leq n$ بیت مخابره کنیم. این نتایج البته در حد n های بزرگ دقیق هستند. تمام این حرف ها در صورتی بود که از نوفه داخل کانال صرف نظر کنیم و فرض کنیم که پیام بدون هیچ نوع تغییری به مقصد می رسد. مطالب بالا محتوی قضیه شانون را درباره کدگذاری بدون نوفه تشکیل می دهند. در این درس می خواهیم درباره کدگذاری در حضور نوفه بحث کنیم و نهایتا قضیه شانون درباره کدگذاری در حضور نوفه را بیان کنیم. در این درس هم چنین با مفهوم مهم ظرفیت کانال آشنا خواهیم شد و در حالت های خیلی ساده ای مقدار آن را محاسبه می کنیم.

۲ موازنه بین پایین آوردن خطا و بالا بردن نرخ انتقال اطلاعات

ازاین به بعد فرض می کنیم که کلمات منبع خود را به طریق بهینه ای کدگذاری کرده ایم و تنها می خواهیم کد کلمه هارا که رشته هایی از 0 و 1 هستند مخابره کنیم. فرض می کنیم که کانالی که پیام های خود را از طریق آن ارسال می کنیم تحت تاثیرنوفه قرار می گیرد و هرعلامت 0 یا 1 که ارسال می کنیم با احتمال q تبدیل به 1 یا 0 شده و با احتمال $1 - q$ دست نخورده باقی می ماند. می خواهیم کاری کنیم که گیرنده پیام ها همچنان بتواند پیام های صحیح را ازپیام های دریافت شده استخراج کند.

۱.۲ یک تذکرهمهم

آنچه که درمثال بالاگفته شد یک نوع کد موسوم به کد تکرار سه تایی برای تشخیص و تصحیح خطاست و باکدگذاری کلمات فرق می کند. ازاین به بعد دراین درس هروقت که مااز کلمه کدگذاری استفاده می کنیم منظورمان این نوع کدگذاری است. به عبارت دیگر فرض می کنیم که کلمات منبع X که می توانند حروف یک زبان مثل فارسی یاانگلیسی یاعلام ریاضی باشند به نحو بهینه ای تبدیل به رشته های 0 و 1 شده اند و هدف ما مخابره این رشته ها ازدرون کانال است. درانتهای کانال رشته های دریافت شده نخست برای تعیین و تصحیح خطاها بررسی شده وپس ازاصلاح به کلمات زبان موردنظر بازگشایی می شوند.

راهی که برای تعیین و تصحیح خطاباید پیش گرفت آن است عنصر تکرار را به نوعی وارد پیام های خود کنیم. به عنوان مثال می توانیم بجای مخابره یک 0 سه تا 0 و بجای یک 1 سه تا 1 مخابره کنیم واز گیرنده بخواهیم که با استفاده از قانون اکثریت تصمیم بگیرد که یک رشته سه تایی که دریافت کرده است درواقع چه رشته ای بوده است (جدول 1.2).

درواقع احتمال خطا که قبلاً برابر بود با q اینک کمتر شده است، زیرا احتمال وقوع خطا اینک برابر است با احتمال وقوع دو برگردان و یا سه برگردان درکد سه تایی که اولی برابر است با $3q^2(1 - q)$ ودومی برابر است با q^3 . درنتیجه احتمال وقوع خطا برابر خواهد بود با $q^3 + 3q^2(1 - q)$ که برای q های کوچک ازمرتبه q^2 است. البته بهایی برای این کاهش خطا پرداخت کرده ایم و آن این است که نرخ مخابره اطلاعات را پایین آورده ایم و از سه بیت برای مخابره یک بیت استفاده کرده ایم. دراین حالت نرخ مخابره اطلاعات یعنی R برابر است با $\frac{1}{3}$. درحالت کلی اگر از n بیت برای مخابره 2^k پیام استفاده کنیم (که درنمود نوفه می توانست

برای مخابره 2^n پیام مورد استفاده قرارگیرد) می‌گوییم که نرخ مخابره اطلاعات برابر است

$$R := \frac{k}{n}. \quad (1)$$

مسلم است که در این نوع کدگذاری می‌توان بازهم احتمال خطا را کاهش داد و در حد ایده آل آن را به صفر نزدیک کرد ولی در این صورت نرخ R نیز به سمت صفر میل خواهد کرد.

■ تمرین: فرض کنید که کد تکرار n تایی را به کار می‌بریم که در آن n یک عدد فرد است. نرخ مخابره اطلاعات برابر است با $R = \frac{1}{n}$. در این کد احتمال خطا را بر حسب n حساب کنید. سپس منحنی نرخ بر حسب احتمال خطا را رسم کنید.

■ تمرین: در کد هامینگ که با مشخصات $[7, 4, 3]$ تعیین می‌شود نرخ مخابره اطلاعات و هم چنین احتمال خطا را تعیین کنید.

آنچه که مهم است این است که ما نمی‌خواهیم اطلاعات را با احتمال خطای کم مخابره یا در جایی ذخیره کنیم. هدف ما آن است که این کار را با نرخ خطای صفر انجام دهیم یا به عبارت دقیق‌تر مطلوب ما آن است که این نرخ خطا را هر چقدر که می‌خواهیم بتوانیم به صفر نزدیک کنیم. حال سوالی که با آن مواجه هستیم آن است که آیا اصولاً می‌توان از یک کانال نوفه دار با نرخ محدود و غیرصفر اطلاعاتی را عبور داد به نحوی که احتمال وقوع خطا لااقل در حد مجانبی (برای رشته‌های بزرگ) به سمت صفر میل کند؟ آیا یک نوع کدگذاری وجود دارد که این امکان را برای ما فراهم آورد؟ آیا نرخ بیشینه‌ای وجود دارد که هر نوع مخابره اطلاعات با بیش از آن نرخ بدون خطا غیرممکن باشد؟ این‌ها سوالاتی است که در قضیه دوم شانون پاسخ داده شده است. قبل از بیان قضیه کلی سعی می‌کنیم که دو مثال ساده را بررسی کنیم.

نخست حالت ساده‌ای را در نظر می‌گیریم که منبع X پیام‌هایی را ارسال می‌کند که در آنها احتمال وقوع 0 و 1 یکسان باشد. به عبارت دیگر آنتروپی منبع در این مثال برابر است با 1. فرض کنید که پیام خود را در رشته‌های n بیتی ارسال کنیم. مجموعه تمام رشته‌های n بیتی دارای 2^n رشته است. نقاط این فضا نقاط یک شبکه ابرمکعبی n بعدی را پر کرده‌اند. اگر همه این نقاط را به عنوان کدکلمه‌های خود به کار ببریم به آسانی در طول عبور از کانال یک کد کلمه به یکی از کد کلمه‌های همسایه اش یا کد کلمه‌های نزدیکش تبدیل می‌شود و گیرنده واقعاً نمی‌فهمد که چه کد کلمه‌ای برایش ارسال شده است. بنابراین راه مقابله با خطا آن است که سعی کنیم که کدکلمه‌ها را با فاصله کافی از یکدیگر انتخاب کنیم تا احتمال وقوع خطا پایین بیاید.

0	◦ ◦ ◦
1	۱ ۱ ۱

جدول ۱: کد تکرار سه تایی

درواقع وقتی که یک کد کلمه را ارسال می کنیم حول آن یک کره به اندازه کافی بزرگ رسم می کنیم و هر وقت نقطه ای درون این کره دریافت کنیم آن را به عنوان کد کلمه ای که در مرکز آن قرار دارد تعبیر و خطاهای بوجود آمده را تصحیح می کنیم. این کار به طور طبیعی نرخ را پایین می آورد زیرا همه نقاط شبکه استفاده نمی کنیم. حال می پرسیم که کد کلمه هارا با چه فاصله ای انتخاب کنیم. در این جا می بایست بین دو عامل متضاد موازنه ایجاد کنیم. اگر بخواهیم خطارا پایین بیاوریم می بایست فاصله کد کلمه هارا از یکدیگر زیاد کنیم. از طرفی این کار نرخ را پایین می آورد. خوب بیایید ببینیم به طور متوسط یک کد کلمه به چند کلمه نزدیک دیگر ممکن است تبدیل شود؟

در یک کلمه n حرفی هر حرف با احتمال q ممکن است که برگردانده شود. بنابراین تعداد حرف های برگردانده شده به طور متوسط برابر است با nq . این تعداد حرف می تواند به $2^{nH(q)}$ طریق در رشته n حرفی قرار گیرند. در نتیجه تعداد رشته های نزدیک به این کلمه که ممکن است از خطاهای بوجود آمده در این کلمه ایجاد شوند برابر است با $2^{nH(q)}$. در این جا نیز ایده اصلی این است که اگر چه کل خطاهای ممکن در یک رشته n تایی برابر است با 2^n اما تعداد خطاهای متعارف یا نمونه از این مقدار کمتر است. در حقیقت اگر وقوع یک خطا را با بیت 1 و عدم وقوع آن را با بیت 0 نشان دهیم آنگاه تمام خطاهای ممکن تبدیل می شوند به رشته هایی از صفر و یک که تعدادشان برابر است با 2^n ولی نکته در این است که ما تنها می بایست به خطاهای نمونه یا *typical* فکر کنیم که تعدادشان برابر است با $2^{nH(q)}$ که خیلی کمتر از مقدار کل است. در اینجا تابع $H(q)$ همان آنتروپی شانون برای خطا ست که برابر است با

$$H(q) = -q \log_2 q - (1 - q) \log_2 (1 - q). \quad (2)$$

تمامی قضایایی که در درس گذشته در باره تعداد و احتمال رشته های نمونه یاد گرفتیم در مورد خطاهای نمونه نیز صادق است.

در نتیجه یک نحوه کد گذاری خوب آن است که حول هر کلمه ناحیه ای در نظر بگیریم که به طور متوسط تعداد کلمات داخل آن در حدود $2^{nH(q)}$ باشد. حال اگر نرخ مبادله برابر با $R = \frac{k}{n}$ باشد معنایش این است که تعداد کل کد کلمه ها برابر است با $2^k = 2^{nR}$. چون تعداد کل کلمات ممکن برابر است با 2^n به این نتیجه می رسیم که می بایست شرط زیر برقرار باشد:

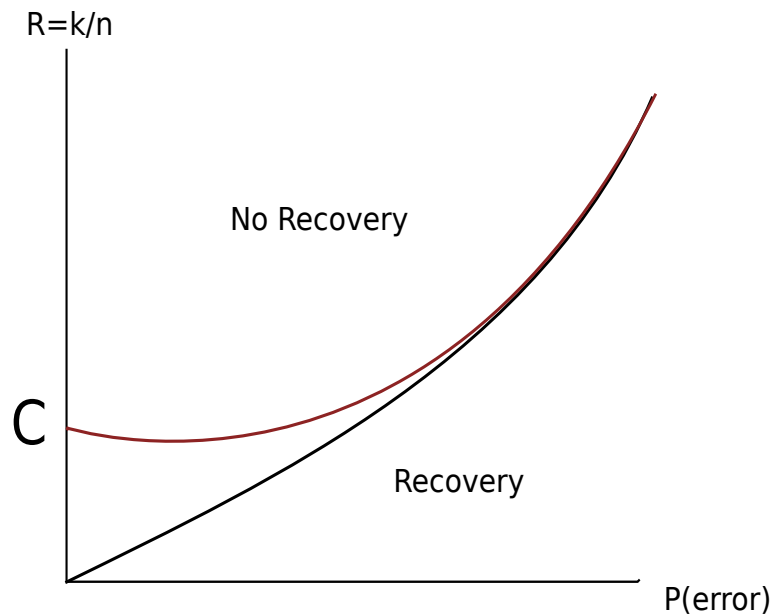
$$2^{nR} 2^{nH} \leq 2^n, \quad (3)$$

یا

$$R \leq 1 - H(q) =: C. \quad (4)$$

کمیت $C := 1 - H(q)$ ظرفیت این کانال نامیده می شود و این رابطه بیان می کند که برای پرهیز از خطا نرخ مبادله اطلاعات حتما باید از ظرفیت کانال کمتر باشد. در واقع آنچه که نشان داده ایم این است بر خلاف تصور اولیه می توان منحنی نرخ بر حسب خطا را از سقوط به نقطه صفر نجات داد و با توجه به این نکته که خطاهای نمونه خیلی کمتر از تعداد کل رشته های ممکن است می توانیم یک نوع کد کردن به کار ببریم که ضمن آنکه نرخ خطا را در حد مجانبی به صفر می رساند ولی نرخ مبادله اطلاعات را به سمت یک مقدار حدی میل می دهد. هر گونه مبادله اطلاعات با کمتر از این نرخ و بدون خطا ممکن است اما هر گونه تجاوز از این نرخ باعث می شود که هیچ گونه اطلاعاتی قابل بازیابی نباشد.

حال بایک سوال اساسی مواجه می شویم و آن اینکه آیا واقعا نوعی کد گذاری وجود دارد که بتوان با استفاده از آن نرخ مبادله اطلاعات را در حد مجانبی $n \rightarrow \infty$ به حداکثر ممکن یعنی به C رساند؟ پاسخ این سوال مثبت است. در واقع نشان می دهیم که حتما یک کد وجود دارد که می تواند نرخ مبادله بدون خطای اطلاعات را هر چقدر که می خواهیم به ظرفیت C نزدیک کند. به عبارت دیگر نشان می دهیم که حتما یک کد وجود دارد که با به کار بردن آن می توانیم نرخ مبادله خطا را هر چقدر که می خواهیم به C نزدیک کرده و نرخ خطا را هر چقدر که می خواهیم کوچک کنیم. دقت کنید که در اینجا مسأله ما ساختن صریح این کد نیست بلکه می خواهیم نشان دهیم که یک چنین کدی وجود دارد. برای اثبات به ترتیب زیر پیش می رویم. فضای تمام کلمه های n بیتی را در نظر می گیریم. این فضا دارای 2^n نقطه است. حال به طور تصادفی 2^k نقطه را درون این فضا در نظر می گیریم. این حرف به این معناست که با احتمال یکنواخت این نقاط را انتخاب می کنیم. بنابراین نحوه کد کردن ما به



شکل ۱: رابطه نرخ مبادله اطلاعات و نرخ خطا و مفهوم ظرفیت.

همین سادگی است، یعنی به جای تعریف صریح کد (مثلاً نوشتن کد خطی با یک ماتریس مولد یا ماتریس پارینه و نظایر آن) تنها 2^k نقطه را به طور تصادفی در این فضا اختیار می‌کنیم و این نقاط را به عنوان کد کلمه‌های خود به کار می‌بریم. به این ترتیب می‌توانیم k بیت را در n بیت کد کنیم. این کد را یک کد تصادفی می‌نامیم و آن را با C_1 نشان می‌دهیم. دقت کنید که ممکن است بعضی از کلمات کد به هم نزدیک انتخاب شوند که در این صورت احتمال رخ دادن خطا برای این کد کلمه‌ها زیاد است. در عوض بعضی فاصله یک کد کلمه می‌تواند از بقیه خیلی زیاد باشد که در این صورت احتمال رخ دادن خطا یعنی تبدیل این کد کلمه به یک کد کلمه دیگر خیلی کم است. هر بار که این نقاط را انتخاب می‌کنیم یک کد جدید به طور تصادفی مشخص می‌شود. این کدها را با C_1 ، C_2 و نظایر آن نشان می‌دهیم. حال یک کد مشخص تصادفی را در نظر می‌گیریم. پس از این نحوه کد کردن می‌بایست به نحوه تصحیح خطا بپردازیم. نحوه اصلاح خطا نیز خیلی ساده است. هر کلمه‌ای که دریافت می‌کنیم به دورش یک کره به شعاع $(H(q) + \delta)$ رسم می‌کنیم. اصطلاحاً این به این معناست که مجموعه تمام نقاطی را در نظر می‌گیریم که دارای $2^{n(H(q)+\delta)}$ در بر داشته باشند. هرگاه یک کد کلمه درون این کره قرار داشت می‌گوییم کلمه‌ای که ارسال شده است همان نقطه بوده است و کلمه دریافت شده را به همان کلمه اصلاح می‌کنیم، شکل ??

■ تمرین: فرض کنید که $n = 20$ ، و $k = 10$ است. احتمال اینکه در یک کد تصادفی یک کد کلمه معین انتخاب شود چقدر است؟ احتمال اینکه دو کد کلمه که فاصله هامینگ آنها یک باشد انتخاب شوند چقدر است؟ فرض کنید که $q = 1/4$ باشد. در این صورت یک کره به شعاع $H(q)$ را در نظر بگیرید. تعداد نقاط درون این کره چند تاست. احتمال اینکه هیچ نقطه ای درون این کره نباشد چقدر است؟ احتمال این که فقط یک نقطه درون این کره باشد چقدر است؟ احتمال اینکه فقط دو نقطه درون این کره باشد چقدر است؟

همانطور که مثال بالا نشان می دهد احتمال اینکه این کره شامل هیچ کد کلمه ای نباشد خیلی کم است. به طور معمول این کره شامل یک نقطه است که همان کد کلمه اولیه ای است که مخابره شده است خطا وقتی رخ می دهد که این کره شامل یک کد کلمه دیگر باشد. احتمال این خط را می توانیم محاسبه کنیم. احتمال این که یک نقطه دیگر به طور اتفاقی درون این کره قرار گرفته باشد برابر است با نسبت حجم این کره به حجم کل فضا یعنی

$$P_1 = \frac{2^{n(H(q)+\delta)}}{2^n}. \quad (5)$$

از آنجا که تعداد نقاط برابر است با $2^k = 2^{nR}$ ، احتمال اینکه یک نقطه درون این کره قرار بگیرد برابر است با

$$P_{error} = 2^{nR} P_1 = 2^{nR} \frac{2^{n(H(q)+\delta)}}{2^n} = 2^{n(R-C+\delta)}. \quad (6)$$

این رابطه نشان می دهد که با انتخاب مناسب n می توانیم نرخ مبادله اطلاعات یعنی R را هرچقدر که می توانیم به C نزدیک کنیم و درعین حال خطا را هرچقدر که می خواهیم کم کنیم.

■ تمرین: اگر بخواهیم که میزان خطا کمتر از 0.001 شود و $R \leq C - 0.01$ باشد طول کد کلمه ها حداقل چقدر باید باشد؟ اگر بخواهیم با همین نرخ اطلاعات را مبادله کنیم ولی احتمال خطا را به زیر 0.00001 برسانیم حداقل طول کد کلمه ها چقدر خواهد بود؟

استدلال بالا نشان می دهد که احتمال خطا را به طور متوسط می توان از هر مقداری که بخواهیم کمتر کنیم. ولی این خطا در واقع یک خطای متوسط دوگانه است به این معنا که یک بار روی تمامی کد کلمه های درون یک کد متوسط گرفته ایم و بار دیگر روی تمام کدهای تصادفی $C_1, C_2, \dots$ متوسط گرفته ایم. در واقع عبارت خطای P_{error} را می بایست با $\langle \overline{P_{error}} \rangle$ عوض کنیم. به طور دقیق تر ثابت کرده ایم که می توان با انتخاب طول مناسب برای کد کلمه ها نرخ مبادله اطلاعات را هرچقدر که می خواهیم به ظرفیت کانال نزدیک کنیم آنچنان که خطا از هر مقداری که می خواهیم کمتر باشد یعنی

$$\langle \overline{P_{error}} \rangle \leq \epsilon, \quad (7)$$

که در آن

$$\langle O \rangle = \frac{1}{Z} \sum_C O(C), \quad (8)$$

یک متوسط روی تمام کدهای تصادفی C است. در این جا Z تعداد کدهای تصادفی است. هم چنین

$$\overline{O} = \frac{1}{2^{nR}} \sum_{i=1}^{2^{nR}} O(i), \quad (9)$$

که در آن متوسط روی تمام کد کلمه های درون یک کد محاسبه می شود. حال از رابطه ۷ نتیجه می شود که حتما یک کد تصادفی مثل C وجود دارد که برای آن رابطه زیر برقرار است:

$$\overline{P_{error}}(C) \leq \epsilon. \quad (10)$$

اما این عبارت متوسط خطا را روی تمام کد کلمه ها نشان می دهد. ممکن است برای بعضی از کد کلمه ها خطا بیشتر از ϵ و برای بعضی دیگر کمتر از ϵ باشد. این کد البته مطلوب ما نیست چرا که ما کدی می خواهیم که احتمال خطا برای تمام کد کلمه هایش از ϵ کمتر باشد. برای رسیدن به این کد تنها کافی است که کد را کمی خلوت کنیم و بعضی از کد کلمه ها را حذف کنیم. تنها یک نگرانی وجود دارد و آن اینکه با خلوت کردن کد و دور ریختن این نوع کلمات تعداد خیلی کمی کد کلمه در کد باقی بماند و نرخ به شدت پایین بیاید یا حتی به صفر برسد. نشان می دهیم که چنین نیست. در واقع نشان می دهیم که تنها با دور ریختن کمتر از نیمی از کد کلمه ها (یا رشته ها) می توانیم به کدی برسیم که خطا برای تک تک کد-کلمه ها از

2ϵ کمتر باشد و چون 2ϵ نیز یک مقدار بی نهایت کوچک است به هدف خود دست پیدا می کنیم. برای این کار باز هم از لم اول چبیشف استفاده می کنیم. خطای مربوط به کد کلمه اول را با E_1 ، خطای مربوط به کد کلمه دوم را با E_2 نشان می دهیم و همینطور تا آخر. در این صورت می دانیم که متوسط این خطا ها از ϵ کمتر است. یعنی

$$\langle E \rangle = \frac{1}{2^{nR}} \sum_{i=1}^{2^{nR}} E_i \leq \epsilon. \quad (11)$$

در این صورت بنا بر لم چبیشف می دانیم که احتمال اینکه یک خطا از 2ϵ بیشتر باشد در نامساوی زیر صدق می کند:

$$P(E \geq 2\epsilon) \leq \frac{\langle E \rangle}{2\epsilon} \leq \frac{\epsilon}{2\epsilon} = \frac{1}{2}. \quad (12)$$

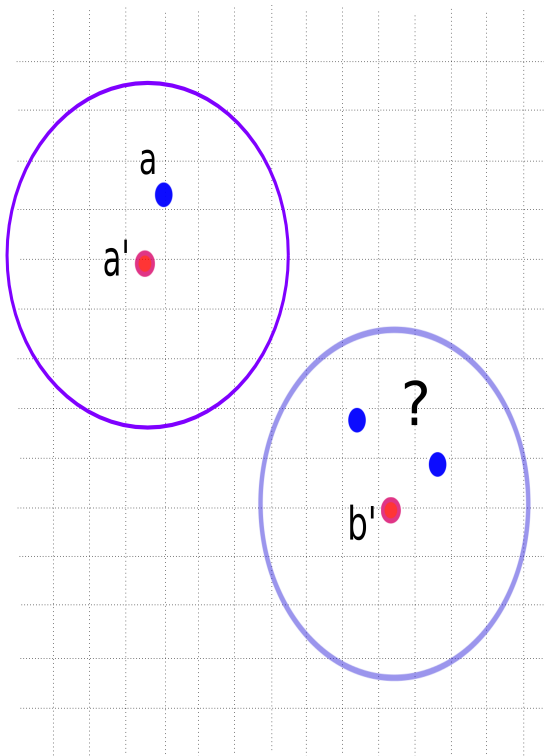
معنای این حرف این است که کمتر از نصف تعداد کل کد کلمه ها مقدار خطایشان 2ϵ بیشتر است و اگر این تعداد از کد کلمه ها را دور بریزیم بقیه همه خطایشان از 2ϵ کمتر است و این همان چیزی است که می خواستیم ثابت کنیم. بنابراین با دور ریختن کمتر از نیمی از کد کلمه ها به کدی می رسم که خطا برای تمام کد کلمه هایش از 2ϵ کمتر است. کد جدیدی که ساخته ایم تعداد کد کلمه هایش برابر است با:

$$\frac{2^{nR}}{2} = 2^{n(R - \frac{1}{n})} \quad (13)$$

و این به این معناست که نرخ جدید در حد n های بزرگ با نرخ قبلی برابر است. بنابراین ثابت کردیم که یک کد وجود دارد که نرخ مبادله اطلاعات را به حد شانون یعنی ظرفیت کانال می رساند و در عین حال خطا را هم هر چقدر که بخواهیم کم می کند.

۳ تعریف ظرفیت در حالت کلی

در بخش قبل خود را به منبع پیامی محدود کردیم که دارای آنتروپی بیشینه بود زیرا در این منبع حرف 0 و حرف 1 هر دو با احتمال $\frac{1}{2}$ رخ می دادند و بنابراین داشتیم $H(X) = H(\frac{1}{2}) = 1$. در این بخش حالت کلی را در نظر می گیریم که در آن آنتروپی منبع دلخواه است. فرض می کنیم که منبع حرف 0 را با احتمال p و حرف 1 را با احتمال $1-p$ تولید کند. در نتیجه آنتروپی منبع



شکل ۲: به دور هر رشته دریافتی یک کره با شعاع مناسب رسم می کنیم. هرگاه یک کدکلمه درون این کره وجود داشته باشد می گوییم همان کلمه ارسال شده و در اثر خطا به این رشته دریافتی تبدیل شده است. خطا وقتی رخ می دهد که وقتی چنین کره ای حول یک رشته دریافتی رسم می کنیم شامل یک کد کلمه دیگر باشد.

برابراست با

$$H(X) = p \log \frac{1}{p} + (1-p) \log \frac{1}{1-p}. \quad (۱۴)$$

باید دقت کنیم که در اینجا منظور ما از منبع همان رشته صفر و یک هایی است که پس از کد کردن الفبا بوجود آمده اند و این آنتروپی با آنتروپی الفبا فرق دارد. تمرین زیر این موضوع را روشن می کند.

در ابتدای بحث های خود وقتی که می خواستیم یک روش کد گذاری مشخص را برای کد کردن k بیت در n بیت تعریف می کنیم نقطه شروع ما رشته های k بیتی بود که آنها را به طریق معینی کد می کردیم. در این جا چون هدف ما پیدا کردن ظرفیت کانال و اثبات وجود یک کدگذاری است که این ظرفیت را اشباع کند روش استدلال ما اندکی متفاوت است اگر چه از

نظر مفهومی با استدلال هایی که قبلاً انجام داده ایم معادل است.

رشته های n حرفی از منبع $X = \{0, 1\}$ را در نظر می گیریم. هرگاه یک رشته n حرفی انتخاب کنیم با احتمال نزدیک به یک این رشته یک رشته نمونه خواهد بود و هرچه که n زیادتر باشد این احتمال نیز به یک نزدیک تر است. تعداد این رشته ها برابر است با $2^{nH(X)}$. می دانیم که کانال کلاسیک دارای یک نوفه است و بنابراین همه این رشته ها را نمی توانیم به سلامت از کانال عبور بدهیم زیرا این رشته ها خیلی نزدیک به هم هستند و براحتی باهم اشتباه می شوند. چاره کار آن است که تنها تعدادی از این رشته ها را انتخاب کرده و آنها را به عنوان رشته های کد شده از کانال عبور دهیم به نحوی که بتوانیم در مقصد خطای اعمال شده روی آنها را اصلاح کنیم و رشته های صحیح را بازیابی کنیم. هرگاه بتوانیم 2^k رشته را در این فضا انتخاب کنیم که دارای خصلت فوق باشند می گوئیم در این کد k بیت را در n بیت کد کرده ایم. بنابراین بازهم به ترتیب قبل پیش می رویم به این نحو که به طور تصادفی 2^k نقطه را انتخاب می کنیم و این کد ما را تعریف می کند. برای اصلاح خطا به دور هر کلمه یا رشته ای مثل y که دریافت می کنیم یک کره به شعاع $H(X|y) + \delta$ رسم می کنیم. این کره دارای $2^{n(H(X|Y) + \delta)}$ رشته یا کلمه را در بر می گیرد. این رشته ها محتمل ترین رشته هایی x ای هستند که می توانسته اند در اثر عبور از کانال و تحمل خطا به رشته y دریافت شده تبدیل شده باشند. رجوع کنید به شکل ۱.۲.

این موضوع را بعداً نشان می دهیم. احتمال خطا این است که در درون این کره علاوه بر رشته اولیه x یک رشته دیگر نیز وجود داشته باشد. یعنی این که یک رشته دیگر در اثر خطای زیاد از جای خود حرکت کرده و درون این کره افتاده باشد. این خطا را براحتی می توان حساب کرد. از آنجا که کد تصادفی است احتمال این که یک رشته معین درون این کره قرار گرفته باشد برابر است با:

$$P_1 = \frac{2^{n(H(X|Y) + \delta)}}{2^{nH(X)}}. \quad (15)$$

این احتمال این است که یک رشته معین در اثر خطا درون این کره قرار گرفته باشد. احتمال این که هر کدام از رشته های کد درون این کره قرار گرفته باشند برابر است با:

$$P_{error} = 2^k \times P_1 = 2^{nR} \times \frac{2^{n(H(X|Y) + \delta)}}{2^{nH(X)}} = 2^{n(R - H(X) + H(X|Y) + \delta)}. \quad (16)$$

بنابراین هرگاه که

$$R > H(X) - H(X|Y) \quad (17)$$

باشد می توانیم با بزرگ تر کردن طول رشته ها خطا را هر قدر که می خواهیم کوچک کنیم. به این ترتیب به تعریف زیر برای ظرفیت یک کانال می رسم:

$$C = \max_X I(X; Y) = \max_X (H(X) - H(X|Y)). \quad (18)$$

و در نتیجه نشان دادیم که امکان مخابره اطلاعات با نرخ R که کمتر از ظرفیت کانال باشد امکان پذیر است.

■ تمرین: فرض کنید که $H(X) = 1/3$. کانال را نیز به صورت زیر تعریف می کنیم:

$$P(0|1) = P(1|0) = p. \quad (19)$$

الف: عبارت $H(X) - H(X|Y)$ را حساب کنید.

ب: اگر بخواهیم که نرخ مبادله اطلاعات به اندازه $\frac{1}{100}$ از این مقدار کمتر باشد و درعین حال خطا از 0.0001 کمتر باشد طول رشته ها از چه عددی می بایست بیشتر باشد؟

حال به آخرین قسمت استدلال می رسم و آن اینکه به طور متوسط تعداد رشته های x ای که می توانسته اند به یک y تبدیل شده باشند برابر است با $2^{nH(X|Y)}$. برای فهم این موضوع دقت می کنیم که تعداد رشته های نمونه از نوع x برابر است با $2^{nH(X)}$. تعداد جفت رشته های نمونه های (x, y) برابر است با $2^{nH(X, Y)}$. هم چنین تعداد رشته های نمونه از نوع y برابر است با $2^{nH(Y)}$. مطابق با قضایایی که در درس پیشین ثابت کردیم تمام این رشته ها با احتمال یکنواخت توزیع شده اند. بنابراین از تعداد $2^{nH(X, Y)}$ رشته نمونه با احتمال $2^{-nH(Y)}$ به یک رشته برمی خوریم که یک y معین داشته باشد. در نتیجه تعداد جفت رشته های نمونه که y آنها یک رشته معین باشد برابر است با

$$2^{-nH(Y)} \times 2^{nH(X, Y)} = 2^{nH(X|Y)}. \quad (20)$$

و این همان چیزی بود که می خواستیم ثابت کنیم. به این ترتیب اثبات رابطه اصلی در مورد ظرفیت کانال های کلاسیک و هم چنین درس مربوط به اطلاعات کلاسیک به پایان می رسد.