

تعیین کران‌های واریانس توزیع‌های تک‌مدی دم‌سنگین به کمک آنتروپی توانی

فاطمه آشتاب^۱ محمدحسین دهقان^۲ منیژه صانعی طبس^۳

تاریخ دریافت: ۱۴۰۲/۰۷/۱۵

تاریخ پذیرش: ۱۴۰۲/۱۱/۱۵

چکیده:

واریانس و آنتروپی معیارهایی متمایز هستند که معمولاً برای اندازه‌گیری عدم قطعیت متغیرهای تصادفی استفاده می‌شوند. همان‌طور که واریانس هر متغیر تصادفی، پراکندگی آن حول میانگین را نشان می‌دهد، معیار آنتروپی عدم قطعیت یک رویکرد اطلاعاتی را اندازه‌گیری می‌کند به عبارت دیگر میانگین مقدار اطلاع یک متغیر تصادفی را اندازه‌گیری می‌کند. برای دو توزیع یکنواخت و نرمال واریانس نسبتی از آنتروپی توانی است. یافتن یک چنین رابطه یکنوا بین واریانس و آنتروپی برای یک رده بزرگ‌تر از این دو توزیع اهمیت و کاربرد زیادی در پردازش سیگنال، یادگیری ماشین، نظریه اطلاع و احتمال و آمار دارد. برای کم کردن خطاهای برآوردگرها مورد استفاده قرار می‌گیرد و یک راهبردی را انتخاب می‌کند که به‌طور متوسط بیشترین یا تقریباً بزرگ‌ترین کاهش در آنتروپی توزیع مکان هدف داشته باشد و اثربخشی این روش با استفاده از شبیه‌سازی‌ها با مدل‌های سنجش کاوی امتحان می‌شوند. در این مقاله کران بالای واریانس برای توزیع‌های تک‌مدی که دم آن‌ها سنگین‌تر از دم توزیع نمایی است به کمک آنتروپی توانی ایجاد می‌گردد.

واژه‌های کلیدی: آنتروپی توانی، کران‌های واریانس، توزیع‌های تک‌مدی، پیوستگی لیپ‌شیتز.

۱ مقدمه

ناهمگنی و بی‌نظمی یک سیستم تعریف می‌شود. در ترمودینامیک، آنتروپی نشان‌دهنده تعداد حالات میکروسکوپی که متناسب با یک حالت ماکروسکوپی (وضعیت قابل مشاهده) است؛ به عبارت دیگر، آنتروپی میزان عدم قطعیت در یک سیستم را نشان می‌دهد و به میزان ترتیب یا بی‌نظمی در سیستم نسبت می‌دهد. در نظریه اطلاع، آنتروپی به عنوان یک مفهوم برای اندازه‌گیری میزان اطلاعات موجود در یک مجموعه از داده‌ها استفاده می‌شود. آنتروپی در این حالت نشان‌دهنده میزان عدم قطعیت در پیام یا داده‌ها است؛ به عبارت دیگر، آنتروپی بیشتر به معنی وجود اطلاعات بیشتر و عدم قطعیت کمتر است؛ که بعد از گفته شدن آنتروپی شانون، دانشمندان با تعمیم دادن آن معیارهای دیگر نیز برای آنتروپی ارائه دادند [۲].

آنتروپی یک مفهوم مهم در حوزه اطلاعات و تئوری اطلاعات است. در اصل، آنتروپی یک میزان از عدم قطعیت و تصادفی بودن یک مجموعه از رویدادها را نشان می‌دهد. آنتروپی به عنوان یک معیار اطلاعاتی استفاده می‌شود و میزان اطلاعات موجود در یک سیستم را مشخص می‌کند. مقدار آنتروپی بیشتر، نشان‌دهنده بیشترین عدم قطعیت و تصادفی بودن سیستم است [۲].

تعریف ۱.۱. آنتروپی یک مفهوم کلیدی در حوزه‌های مختلف علمی و مهندسی، از جمله فیزیک، اطلاعاتی، نظریه اطلاع، تکنولوژی اطلاعات، آمار و احتمالات و رمزگذاری است. این مفهوم ابتدا توسط رودولف کلوژیوس در قرن نوزدهم مطرح شد و سپس توسط کلود شانون در دهه ۱۹۴۰ و ۱۹۵۰ به صورت دقیق‌تر مورد بررسی قرار گرفت. در فیزیک، آنتروپی به عنوان یک مفهوم مرتبط با نظریه ترمودینامیک استفاده می‌شود. آنتروپی معمولاً به عنوان یک معیار برای اندازه‌گیری

^۱ دانشجوی کارشناسی ارشد دانشگاه سیستان و بلوچستان

^۲ عضو هیئت علمی دانشگاه سیستان و بلوچستان

^۳ عضو هیئت علمی دانشگاه سیستان و بلوچستان (نویسنده مسئول: manijesanei@math.usb.ac.ir)

۱.۱ آنتروپی در حالت پیوسته

تعریف ۲.۱. اگر X یک متغیر تصادفی در حالت پیوسته با تابع چگالی احتمال $f_X(x)$ باشد آنتروپی آن عبارت است از:

$$H(X) = - \int_{-\infty}^{\infty} f(x) \log f_X(x) dx.$$

۲.۱ آنتروپی در حالت گسسته

تعریف ۳.۱. اگر X یک متغیر تصادفی گسسته با تابع احتمال $P_X(x)$ باشد، آنگاه

$$H(X) = - \sum_x P_X(x) \log P_X(x).$$

آنتروپی متغیرهای تصادفی گسسته خواهد بود.

تعریف ۴.۱. آنتروپی توانی مفهومی است که در حوزه تحلیل سیگنال و پردازش سیگنال مورد استفاده قرار می‌گیرد. آنتروپی توانی به صورت معمول در مطالعه توزیع‌های احتمالی پیوسته استفاده می‌شود. در واقع، آنتروپی توانی یک معیار از تنوع و پیچیدگی سیگنال به صورت زمانی است و مقدار آنتروپی توانی بیشتر، نشان‌دهنده وجود تغییرات پیچیده و تنوع بیشتر در سیگنال است [۲۹]. به طور خلاصه، آنتروپی مفهومی در حوزه اطلاعات و تئوری اطلاعات است که میزان عدم قطعیت و تصادفی بودن سیستم را اندازه‌گیری می‌کند، درحالی‌که آنتروپی توانی به عنوان یک مفهوم در حوزه پردازش سیگنال استفاده می‌شود و تنوع و پیچیدگی سیگنال را نشان می‌دهد.

آنتروپی توانی یک مفهوم در نظریه اطلاع است که برای اندازه‌گیری میزان تنوع و پیچیدگی یک سیگنال استفاده می‌شود. آنتروپی توانی یک سیگنال، مقداری است که نشان می‌دهد چقدر انرژی در سیگنال توزیع شده است. معمولاً آنتروپی توانی برای سیگنال‌های پیوسته محاسبه می‌شود. برای محاسبه آنتروپی توانی یک سیگنال پیوسته، ابتدا توان سیگنال را در طول زمان یا فضا محاسبه می‌کنیم و سپس میانگین لگاریتم طبیعی این توان‌ها را محاسبه می‌کنیم. این میانگین توان را می‌توان با استفاده از انتگرال‌گیری نیز محاسبه کرد. آنتروپی توانی در برخی موارد می‌تواند به عنوان یک معیار برای تحلیل سیگنال‌های پیچیده و تصادفی مورد استفاده قرار گیرد. می‌توانیم یک رابطه مستقیم بین واریانس و آنتروپی تفاضلی برای توزیع‌های یکنواخت متقارن و نامتقارن داشته باشیم. هم‌چنین در نظریه اطلاع نابرابری آنتروپی توانی به اصطلاح به قدرت آنتروپی متغیر تصادفی مربوط می‌شود. این نشان می‌دهد، آنتروپی مفهومی است که در حوزه‌های مختلف برای اندازه‌گیری

ناهمگنی، بی‌نظمی و عدم قطعیت استفاده می‌شود، اما اصول عمومی آنتروپی برای اندازه‌گیری وجود اطلاع و عدم قطعیت کاربرد دارد. تعیین کران‌های واریانس توزیع تک‌مدی با استفاده از آنتروپی توانی، در واقع یک روش برای تخمین زدن حداکثر و حداقل مقادیر واریانس در یک توزیع است. این روش به ما اطلاعاتی در مورد پراکندگی داده‌ها و نقطه تمرکز توزیع می‌دهد.

اهمیت تعیین کران‌های واریانس توسط آنتروپی توانی به دلایل زیر است [۲۷]:

۱. ارزیابی کیفیت و تنوع داده‌ها: واریانس یک معیار بسیار مهم برای اندازه‌گیری پراکندگی داده‌ها است. با تعیین کران‌های واریانس، می‌توانیم محدوده‌ای را مشخص کنیم که بیانگر پراکندگی مجموعه داده است. این محدوده ما را قادر می‌سازد تا به طور کیفی ارزیابی کنیم که آیا داده‌ها واحدی و یکنواخت هستند یا نه. اگر واریانس بسیار کم باشد، نشان می‌دهد که داده‌ها به طور کلی متمرکز و کم پراکنده هستند، درحالی‌که واریانس بزرگ نشانگر پراکندگی بیشتر و تنوع بیشتر است.
۲. تحلیل و پیش‌بینی رفتار توزیع: با تعیین کران‌های واریانس، می‌توانیم رفتار توزیع را بررسی کنیم و درک بهتری از شکل و خصوصیات آن پیدا کنیم. مثلاً، اگر واریانس بسیار کوچک باشد، نشان‌دهنده این است که توزیع تقریباً یکنواخت است و بیشتری تمرکز بر روی مقدار میانگین دارد؛ به عبارت دیگر، نقطه تمرکز توزیع به شدت تأثیرگذار است. از طرف دیگر، واریانس بزرگ نشان می‌دهد که داده‌ها پراکنده‌تر هستند و نقطه تمرکز توزیع کمتری اهمیت دارد.
۳. استفاده در تحلیل آماری: واریانس در بسیاری از روش‌ها و تحلیل‌های آماری به عنوان یک معیار مهم استفاده می‌شود. با تعیین کران‌های واریانس، می‌توانیم مطمئن شویم که تخمین‌های آماری ما دقیق و قابل اعتماد هستند.

بنابراین، تعیین کران‌های واریانس با استفاده از آنتروپی توانی، ابزاری قدرتمند برای ارزیابی و تحلیل داده‌هاست. این روش به ما کمک می‌کند تا به طور کمی و کیفی بهتری از پراکندگی و تنوع داده‌ها دست پیدا کنیم و به صورت کمی رفتار توزیع را تحلیل کنیم. نابرابری آنتروپی توانی در سال ۱۹۴۸ توسط کلود شانون [۶] در مقاله اصلی خود او که به نام نظریه ارتباطات ریاضی بود، اثبات شد که نابرابری آنتروپی توانی را در ادامه بیان خواهیم کرد.

تعریف ۵.۱. اگر X یک بردار تصادفی در $\mathbb{R}^n$ با چگالی f و آنتروپی $h(X)$ باشد، آنتروپی توانی آن به صورت زیر تعریف می‌گردد:

$$N(X) = \frac{1}{\pi e} e^{\frac{1}{n} h(X)} \quad (1)$$

تابعی از یک متغیر تصادفی که شرایط نظم را برآورده می‌کند، ارائه داده است. همه این کران شامل دو گشتاور اول مشتقات یا تفاوت‌های تابع است. واریانس و آنتروپی معیارهایی متمایز هستند که معمولاً برای اندازه‌گیری عدم قطعیت متغیرهای تصادفی استفاده می‌شوند. درحالی‌که واریانس نشان می‌دهد که چگونه یک متغیر تصادفی بیشتر از حد انتظارش گسترش می‌یابد، معیار آنتروپی رویکرد متفاوتی دارد. آنتروپی عدم قطعیت یک رویکرد اطلاعاتی را اندازه‌گیری می‌کند به عبارت دیگر میانگین مقدار اطلاع یک متغیر تصادفی را اندازه‌گیری می‌کند. اگرچه واریانس و آنتروپی به‌طور کلی یکسان نیستند، یک مورد وجود دارد که واریانس و آنتروپی معیارهای معادل هستند. هنگامی که توزیع نرمال در نظر گرفته می‌شود، حداکثر آنتروپی برابر واریانس به‌دست‌آمده است. این نشان می‌دهد که واریانس و آنتروپی تا حدی مرتبط هستند هنگامی که توزیع احتمال اساساً چندوجهی باشد به‌گونه‌ای که برای هر آزمایش بیش از دو نتیجه، واریانس به‌عنوان یک معیار عدم قطعیت ناکارآمد می‌شود. به‌عنوان مثال، ترکیبی از دو توزیع نرمال تک متغیره هم‌وزن با واریانس واحد و میانگین‌های متفاوت را نظر بگیرید. هنگامی که اختلاف میانگین دو متغیر بینهایت شود واریانس کل توزیع بی‌نهایت می‌شود ولی آنتروپی در این موقعیت پایدار باقی می‌ماند؛ بنابراین، آنتروپی به‌عنوان معیار مناسب‌تری برای عدم قطعیت متغیرهای تصادفی چندوجهی در نظر گرفته می‌شود. این مطلب به کمک نابرابری معروفی بنام نابرابری فانو نشان داده می‌شود که به آن اشاره خواهیم کرد.

تعریف ۷.۱. برای هر برآوردگر $\hat{X}$ به‌طوری‌که $X \rightarrow Y \rightarrow \hat{X}$ با $X \neq \hat{X}$ و با در نظر گرفتن احتمال خطا به شکل $P_e = P(X \neq \hat{X})$ نابرابری فانو به‌صورت زیر بیان می‌شود [۲۱]:

$$H(P_e) + P_e \log |X| \geq H(X|\hat{X}) \geq H(X|Y) \quad (۳)$$

توجه داشته باشید که در رابطه (۳) $P_e = ۰$ به این معنی است که $H(X|Y) = ۰$

هم‌چنین فرض کنید با شناخت یک متغیر تصادفی Y می‌خواهیم مقدار یک متغیر تصادفی همبسته X را حدس بزنیم. نابرابری فانو احتمال خطا در حدس زدن متغیر تصادفی X را به آنتروپی شرطی آن $H(X|Y)$ مرتبط می‌کند.

در حالت خاص وقتی X دارای توزیع نرمال با ماتریس کوواریانس K باشد، داریم:

$$N(X) = |K|^{\frac{1}{n}}$$

یک نابرابری مهم و پرکاربرد در رابطه با آنتروپی توانی معرفی شده است که به نابرابری آنتروپی توانی معروف است (EPI) در ادامه این نابرابری ارائه می‌گردد.

تعریف ۶.۱. اگر X و Y دو بردار تصادفی مستقل در $\mathbb{R}^n$ با تابع چگالی مربوطه و گشتاور مرتبه دوم متناهی باشند، نابرابری زیر برقرار است [۶]:

$$N(X+Y) \geq N(X) + N(Y) \quad (۲)$$

برای مطالعه بیشتر و کامل‌تر در خصوص آنتروپی و سایر مفاهیم نظریه اطلاع به کتاب نظریه اطلاع تألیف کاور و توماس [۲] مراجعه شود.

۳.۱ کران‌دار کردن واریانس

کران بالای واریانس نشان‌دهنده یک هدف مهم است، زیرا در بسیاری از زمینه‌های آماری مانند تخمین واریانس، هم‌چنین به‌عنوان یک حداکثر برای فرآیند تصادفی، پرکاربرد است. اولین بار موضوع کران واریانس برای متغیر تصادفی که دارای کران بالا است به‌صورت زیر مطرح گردید: فرض کنید C یک عدد مثبت ثابت باشد. فرض کنید X یک متغیر تصادفی باشد که فقط مقادیری بین ۰ و C می‌گیرد. این دلالت بر $P(0 \leq X \leq C) = ۱$ دارد. پس نابرابری زیر در مورد واریانس $V(X)$ ثابت می‌گردد:

$$V(X) \leq \frac{C^2}{۴}$$

نابرابری‌های پوانکاره (هم‌سوپی‌متری) که کران‌های بالایی را برای واریانس تابعی از یک متغیر تصادفی، به دست می‌دهند دارای تاریخچه طولانی و غنی است که با کار چرنوف (۱۹۸۱) [۴] شروع می‌شود. چرنوف با استفاده از چندجمله‌ای هرمیت، کران بالایی را برای واریانس تابعی از یک متغیر تصادفی معمولی به دست آورد. چن (۱۹۸۲) [۳] با استفاده از نابرابری کوشی-شوارتز، اثبات متفاوتی ارائه کرد و نابرابری را به حالت عادی چند متغیره تعمیم داد. کاکلوس (۱۹۸۲) [۵] نشان داده است که چگونه می‌توان کران‌های بالایی مشابه را برای توزیع‌های دیگر، از جمله موارد گسسته، به دست آورد. علاوه بر این، با استفاده از نابرابری کرامر-رائو^۴، کران‌های پایینی مشابهی برای واریانس

^۴Cramer-Rao

۲ کران‌های واریانس برای توزیع‌های تک‌مدی

۱۰۲ توزیع دم‌سنگین

توزیع دم‌سنگین^۵ یک مفهوم در آمار و احتمال است که به توزیعی اشاره دارد که دارای دم‌های بزرگ و احتمالات نادر در این دم‌ها است؛ به عبارت دیگر، در یک توزیع دم‌سنگین، احتمال وقوع رویدادهای نادر و خارج از محدوده معمول بزرگ است.

این مفهوم معمولاً در زمینه تحلیل داده‌ها و مدل‌سازی استفاده می‌شود. توزیع دم‌سنگین می‌تواند در مواردی مانند توزیع درآمد، سائز فایل‌ها در سیستم عامل، تعداد دنبال‌کنندگان در شبکه‌های اجتماعی و غیره مشاهده شود.

توزیع دم‌سنگین معمولاً به معنای وجود انحرافات شدید در داده‌ها است که از سنت‌های نرمال یا یک توزیع دم متوسط که دارای دمی سنگین‌تر است، جدا می‌شود. این انحرافات می‌توانند نشان‌دهنده رویدادهای نادر یا غیرمعمول باشند که در تحلیل داده‌ها بسیار مهم هستند.

مهم‌ترین خصوصیت توزیع دم‌سنگین، وجود مقادیر نادر (اعداد بزرگ) است که می‌توانند تأثیر بسیار زیادی بر نتایج تحلیل داشته باشند. این موضوع ممکن است به تحلیل‌های آماری سنتی که بر اساس فرضیاتی مانند توزیع نرمال بر مبنای داده‌های معمول توسعه یافته‌اند، چالش بیاورد؛ بنابراین، در تجزیه و تحلیل داده‌ها و استنتاج‌های آماری بر مبنای توزیع دم‌سنگین، نیاز به روش‌های خاص و مناسبی وجود دارد [۲۸].

۲۰۲ کران‌های واریانس چگالی‌های مخلوط تک‌مدی متقارن

ابتدا چگالی‌های مخلوط خطی تک‌مدی متقارن به فرم زیر را که مخلوطی از توزیع‌های نزولی نمایی، $p_i(x) \propto e^{-\beta_i|x-m|^{|\theta_i|}}$ برای هر $i = 1, 2, \dots, n$ است را در نظر می‌گیریم [۱]:

$$p(x) = \sum_{i=1}^n \alpha_i p_i(x), \quad \alpha_i > 0, \sum_{i=1}^n \alpha_i = 1. \quad (4)$$

که $p_i(x) \propto e^{-\beta_i|x-m|^{|\theta_i|}}$ آمیخته‌ای از توزیع‌های نزولی نمایی برای هر $\theta_i > 0$ و $\beta_i > 0$ است.

برای هر قضیه ۶۰۲ کران بالای واریانس چگالی نزولی نمایی مخلوط خطی تک‌مدی را تحت این فرض که نسبت مابین واریانس‌های

تعریف ۱۰۲. در توزیع‌های تک‌مدی مقادیر در ابتدا افزایش می‌یابند و به یک قله منفرد می‌رسند و سپس از پرتکرارترین عدد در یک حالت کاهش می‌یابند. به‌طور کلی اگر در یک توزیع یک مد وجود داشته باشد به آن توزیع تک‌مدی گفته می‌شود و این حالت هم در توزیع‌های گسسته و هم پیوسته تعریف شده است که می‌توانیم به چند توزیع مانند توزیع گاما، برنولی، هندسی، دوجمله‌ای منفی، فوق هندسی، پواسون، نمایی، خی‌دو، نرمال و غیره اشاره کرد. تعریف رسمی خاصیت تک‌مدی به صورت زیر است [۲۲].

تعریف ۲۰۲. یک متغیر تصادفی Z با تابع توزیع F را تک‌مدی حول (a) می‌گوییم، اگر F در فاصله‌ی $(-\infty, a)$ محدب (تقعر رو به بالا) و در فاصله‌ی (a, ∞) مقعر (تقعر رو به پایین) باشد.

تعریف ۳۰۲. توزیع تک‌مدی پیوسته: اگر تابع چگالی یک متغیر تصادفی پیوسته مطلق روی فاصله $(-\infty, a)$ صعودی و در فاصله (a, ∞) نزولی باشد، آنگاه آن توزیع تک‌مدی پیوسته حول a خوانده می‌شود.

تعریف ۴۰۲. توزیع تک‌مدی گسسته: یک توزیع گسسته با تابع جرم احتمال (p_n) را روی مجموعه‌ی اعداد صحیح، تک‌مدی حول a می‌نامند اگر در رابطه‌ی زیر صدق کند:

$$\begin{cases} p_n \geq p_{n-1}, & n \leq a \\ p_n \geq p_{n+1}, & n \geq a. \end{cases}$$

برای توزیع‌های تک‌مدی گسسته توزیع‌هایی چون برنولی، دوجمله‌ای، هندسی، پواسون و ... را می‌توان نام برد.

لم ۵۰۲. پیوستگی لیپشیتز: تحلیل ریاضی، پیوستگی لیپشیتز که به نام ریاضیدان رودولف لیپشیتز نام‌گذاری شده است، یک شکل قوی از پیوستگی یکنواخت برای توابع است. به‌طور شهودی یک تابع پیوسته لیپشیتز سرعت تغییر آن محدود است. فرض کنیم $f: (X, d_1) \rightarrow (Y, d_2)$ یک تابع باشد. گوئیم f در شرط لیپشیتز با ثابت M صدق می‌کند هرگاه عددی ثابت مانند $M > 0$ موجود باشد به قسمی که به ازای هر $x, y \in X$ داشته باشیم

$$d_2(f(x), f(y)) \leq M d_1(x, y).$$

⁵heavy-tailed distribution

با استفاده از این لم ۷.۲، می‌توانیم کران بالای واریانس توزیع مخلوط را ساده کنیم هرگاه همه θ_i ها مؤلفه مخلوط $p_i(x) = \frac{1}{Z_i(\theta_i, \beta_i)} e^{-\beta_i |x-m|^{\theta_i}}$ برای $i = 1, \dots, n$ کران‌دار پایینی به ثابت مثبت باشند. برای مثال، هرگاه برای هر i ، $\theta_i \geq 1$ باشد، بیشترین نسبت بین کران بالایی و کران پایینی برابر با $\frac{2\pi e}{12} M(r)$ می‌شود ازاین‌رو

$$\prod_{i=1}^n (1/A(\theta_i))^{\alpha_i} \leq 1/12.$$

هرگاه برای هر i ، $\theta_i \geq 1/2$ باشد، آنگاه $\prod_{i=1}^n (1/A(\theta_i))^{\alpha_i} \leq (2e^4)/15$ ، بیشترین نسبت بین کران بالایی و کران پایینی برابر $\frac{2\pi e^5}{15} M(r)$ می‌شود. توجه داشته باشید هرگاه برای هر i ، $\theta_i \geq 1/2$ باشد، مؤلفه‌های آمیخته $p_i(x)$ می‌توانند توزیع‌های دُم-سنگین باشند اما واریانس $p(x)$ هنوز از بالا به تابع یکنوا با آنتروپی توانی $e^{2h(p)}$ کران‌دار است [۱].

نتیجه ۸.۲. چگالی تک‌مدی متقارن به فرم $p(x) = \sum_{i=1}^n \alpha_i p_i(x)$ برای $\alpha_i > 0$ ، $\sum_{i=1}^n \alpha_i = 1$ در نظر بگیرید که $p_i(x) = \frac{1}{Z_i(\theta_i, \beta_i)} e^{-\beta_i |x-m|^{\theta_i}}$ با ثابت نرمال شده $Z_i(\theta_i, \beta_i)$ و $\beta_i > 0$ است. هرگاه برای هر i ، مرتبه $\theta_i \geq 1$ باشد، واریانس به‌صورت زیر کران‌دار است [۱]:

$$\frac{e^{2h(p)}}{2\pi e} \leq \text{Var}(X) \leq \frac{M(r)}{12} e^{2h(p)}. \quad (9)$$

هرگاه برای هر i ، $\theta_i \geq 1/2$ ،

$$\frac{e^{2h(p)}}{2\pi e} \leq \text{Var}(X) \leq \frac{2e^4 M(r)}{15} e^{2h(p)}, \quad (10)$$

که $M(r)$ در رابطه (۸) داده شده است.

نتیجه ۹.۲. فرض کنید که $p(x)$ چگالی تک‌مدی متقارن با تکیه‌گاه کران‌دار به فرم $p(x) = \sum_{i=1}^n \alpha_i p_i(x)$ باشد که $p_i(x) = \frac{1}{Z_i(\theta_i, \beta_i)} e^{-\beta_i |x-m|^{\theta_i}}$ با ثابت نرمال شده $Z_i(\theta_i, \beta_i)$ و $\beta_i > 0$ برای $\text{unif}\left(m - \frac{1}{2\epsilon_i}, m + \frac{1}{2\epsilon_i}\right)$ همچنین، فرض کنید $r = \max_{i,j \in \{1, \dots, n\}} \left\{ \frac{\epsilon_i}{\epsilon_j} \right\}$ کران‌دار باشد. در این صورت واریانس به‌صورت زیر کران‌دار است [۱]:

$$\frac{e^{2h(p)}}{2\pi e} \leq \text{Var}(X) \leq \frac{M(r)}{12} e^{2h(p)}, \quad (11)$$

که $M(r)$ در رابطه (۸) داده شده است.

در قضیه ۶.۲ و نتیجه ۸.۲ کران‌داری نسبت r بین واریانس ماکزیم و واریانس مینیم مؤلفه‌های مخلوط را فرض کردیم؛ یعنی، $\max_{i,j \in \{1, \dots, n\}} \left\{ \sigma_i^2 / \sigma_j^2 \right\}$ در قضیه ۶.۲ و $\max_{i,j \in \{1, \dots, n\}} \left\{ \epsilon_i / \epsilon_j \right\}$

ماکزیم و مینیم اجزای مخلوط $p_i(x)$ کران‌دار است را ثابت می‌کند.

قضیه ۶.۲. فرض کنید $p(x)$ چگالی مخلوط خطی تک‌مدی متقارن به فرم (۴) با مؤلفه $p_i(x) = \frac{1}{Z_i(\theta_i, \beta_i)} e^{-\beta_i |x-m|^{\theta_i}}$ باشد که $Z_i(\theta_i, \beta_i)$ ثابت نرمال شده است و $\theta_i, \beta_i > 0$. فرض کنید σ_i^2 / σ_j^2 نشان‌دهنده واریانس $p_i(x)$ و نسبت مؤلفه واریانس‌های σ_i^2 / σ_j^2 برای همه $i \neq j$ کران‌دار باشد، به عبارت دیگر، فرض کنید $r = \max_{i,j \in \{1, 2, \dots, n\}} \left\{ \frac{\sigma_i^2}{\sigma_j^2} \right\}$. در این صورت واریانس چگالی $p(x)$ در رابطه زیر صدق می‌کند [۱]:

$$\frac{e^{2h(p)}}{2\pi e} \leq \text{Var}(X) \leq B(\theta, r) e^{2h(p)}. \quad (5)$$

در اینجا،

$$B(\theta, r) = M(r) \cdot \prod_{i=1}^n \left(\frac{1}{A(\theta_i)} \right)^{\alpha_i}, \quad (6)$$

برای $\theta = (\theta_1, \dots, \theta_n)^T$ که

$$A(\theta) = 2\theta^{-2} \frac{(\Gamma(1/\theta))^2}{\Gamma(2/\theta)} e^{2/\theta}, \quad (7)$$

با تابع گاما $\Gamma(t) = \int_0^\infty x^{t-1} e^{-x} dx$ که برای همه $t > 0$ تعریف می‌شود و

$$M(r) = \frac{(r-1)r^{\frac{1}{r-1}}}{e \log r}, \quad r \geq 1. \quad (8)$$

برابری در واریانس کران پایین در رابطه (۵) به دست می‌آید اگر و تنها اگر $p(x)$ توزیع گاوسی باشد. برابری در واریانس کران پایین زمانی برآورده می‌شود که همه $p_i(x)$ ها دارای یک توزیع باشند؛ به عبارت دیگر، برای هر i ، $p(x) = p_i(x)$. در نتیجه، کران‌های بالا و پایین هم‌ارز هستند هرگاه همه $p_i(x)$ ها دارای توزیع گاوسی باشند. برای یک توزیع مخلوط به فرم (۴)، احتمالاً $p_i(x) \neq p_j(x)$ برای هر جفت (i, j) ، اسکالر ثابت $B(\theta, r)$ در رابطه (۶) بزرگ‌تر یا مساوی $1/(2\pi e)$ است و ازاین‌رو میانگین هندسی $1/A(\theta_i) \geq 1/(2\pi e)$ برای هر θ_i ، با وزن $\{\alpha_i\}$ ، به صورت

$$\prod_{i=1}^n (1/A(\theta_i))^{\alpha_i} \geq 1/(2\pi e),$$

داده می‌شود و $1 \leq M(r)$ با $\lim_{r \rightarrow 1} M(r) = 1$.

لم ۷.۲. تابع $1/A(\theta)$ با $A(\theta)$ داده شده در (۷) روی $[0, 2]$ نزولی و روی $[2, +\infty)$ صعودی است. مینیم $1/A(\theta)$ در $\theta = 2$ با مقدار $1/A(2) = 1/(2\pi e)$ رخ می‌دهد. هرگاه $\theta \rightarrow 0$ باشد [۱]، $\lim_{\theta \rightarrow \infty} 1/A(\theta) = 1/12$ و $\lim_{\theta \rightarrow 0} 1/A(\theta) = \infty$.

نسبت بین واریانس‌های مؤلفه‌های مخلوط بی‌کران باشد، لزوماً کران بالا روی واریانس توزیع تک‌مدی متقارن وجود ندارد یعنی یک ثابت آنتروپی توانی مدرج شده است.

در نتیجه ۹۰۲ کران‌دار هستند. در اینجا، با یک مثال نقض نشان می‌دهیم که هرگاه این فرض نقض شود، واریانس چگالی مخلوط تک‌مدی متقارن لزوماً محدود در تابع یکنوای آنتروپی توانی نیست.

مثال ۱۰۰۲. توزیع تک‌مدی متقارن تشکیل شده از دو چگالی یکنواخت به فرم زیر را در نظر بگیرید:

$$p(x) = \sum_{i=1}^2 \alpha_i \cdot \text{unif} \left(-\frac{1}{2\epsilon_i}, \frac{1}{2\epsilon_i} \right) \quad (12)$$

که $\alpha_i > 0$ و برای $\epsilon_1 > \epsilon_2 > 0$ ، $\sum_{i=1}^2 \alpha_i = 1$ واریانس این توزیع برابر است با:

$$\text{Var}(X) = \frac{1}{12} \left(\alpha_1 \frac{1}{\epsilon_1^3} + \alpha_2 \frac{1}{\epsilon_2^3} \right), \quad (13)$$

و آنتروپی تفاضلی این توزیع به صورت زیر است:

$$h(p) = -\frac{1}{\epsilon_1} (\alpha_1 \epsilon_1 + \alpha_2 \epsilon_2) \log(\alpha_1 \epsilon_1 + \alpha_2 \epsilon_2) - \left(\frac{1}{\epsilon_2} - \frac{1}{\epsilon_1} \right) \alpha_2 \epsilon_2 \log(\alpha_2 \epsilon_2). \quad (14)$$

هرگاه $\epsilon_1/\epsilon_2 \rightarrow \infty$ باشد، حد آنتروپی تفاضلی به صورت رابطه زیر درمی‌آید:

$$\lim_{\epsilon_1/\epsilon_2 \rightarrow \infty} h(p) = -\alpha_1 \log \epsilon_1 - \alpha_2 \log \epsilon_2 + H_B(\alpha_1) \quad (15)$$

که در این صورت $H_B(\alpha_1) = -\alpha_1 \log \alpha_1 - (1 - \alpha_1) \log(1 - \alpha_1)$ آنتروپی توانی به صورت زیر درمی‌آید:

$$\lim_{\epsilon_1/\epsilon_2 \rightarrow \infty} e^{\gamma h(p)} = e^{H_B(\alpha_1)} \left(\frac{1}{\epsilon_1} \right)^{\alpha_1} \left(\frac{1}{\epsilon_2} \right)^{\alpha_2}. \quad (16)$$

برای یافتن یک کران بالا روی واریانس (۱۳) با آنتروپی توانی (۱۶)، نیاز به یک کران بالا روی میانگین حسابی $\left(\alpha_1 \frac{1}{\epsilon_1} + \alpha_2 \frac{1}{\epsilon_2} \right)$ در جملات میانگین هندسی $\left(\frac{1}{\epsilon_1} \right)^{\alpha_1}$ و $\left(\frac{1}{\epsilon_2} \right)^{\alpha_2}$ داریم. با این وجود، اگر $\epsilon_2 \rightarrow 0$ برای ثابت ϵ_1 ، بنا به این که میانگین حسابی روی مرتبه $\Theta(1/\epsilon_2^2)$ افزایش می‌یابد درحالی که میانگین هندسی روی مرتبه $\Theta(1/\epsilon_2^{2\alpha_2})$ برای $\alpha_2 < 1$ افزایش واریانس سریع‌تر از $e^{\gamma h(p)}$ است برای این که نمی‌تواند از بالا برای هر ثابت اسکالر آنتروپی توانی کران‌دار شود. از طرف دیگر، اگر $\epsilon_1 \rightarrow \infty$ برای ثابت ϵ_2 ، میانگین حسابی که متناسب با واریانس است، به‌طور تقریبی α_2/ϵ_1 است که یک ثابت است اما میانگین هندسی که متناسب با $e^{\gamma h(p)}$ است برای $\alpha_1 < 1$ به $\Theta(1/\epsilon_1^{2\alpha_1})$ میل می‌کند. در نتیجه، برای هر حالت که در ویژگی $\epsilon_1/\epsilon_2 \rightarrow \infty$ صدق می‌کند، واریانس $p(x)$ نمی‌تواند از بالا به ثابت اسکالر $e^{\gamma h(p)}$ کران‌دار شود. این مثال نشان می‌دهد که هرگاه

۳۰۲ کران‌های واریانس چگالی‌های تک‌مدی پیوسته-لیپ‌شیتز با تکیه‌گاه کران‌دار

با مؤلفه‌های مخلوط به اندازه کافی بزرگ، هر چگالی تک‌مدی متقارن می‌تواند به دلخواه همانند یک مخلوط خطی به فرم (۲) تقریب زده شود. این برای حالت خاص چگالی‌های یکنواخت $p_i(x)$ اثبات شده است [۲۳]. ما در اینجا، این نتیجه را تعمیم می‌دهیم. فرض کنید $p(x)$ چگالی تک‌مدی متقارن با تکیه‌گاه کران‌دار $[m-s, m+s]$ باشد. فرض کنید که $p(x)$ در شرط پیوستگی-لیپ‌شیتز با ثابت $c_s > 0$ صدق کند، به عبارت دیگر

$$|p(x+y) - p(x)| \leq c_s |y| \quad (17)$$

برای هر $x, y \in [m-s, m+s]$. برای چنین $p(x)$ ، چگالی مخلوط خطی $\bar{p}_n(x)$ را می‌سازیم که تقریب‌های $p(x)$ به قسمی که تمایز بین توان‌های آنتروپی $p(x)$ و $\bar{p}_n(x)$ و واریانس‌های $p(x)$ و $\bar{p}_n(x)$ می‌توانند به شکل زیر کران‌دار شوند [۱]:

$$e^{\gamma h(\bar{p}_n)} \leq e^{\gamma h(p)} (1 + c_1 n^{-1} \log n), \quad \left| \int_{m-s}^{m+s} (x-m)^{\gamma} (p(x) - \bar{p}_n(x)) dx \right| \leq c_2 n^{-1} \quad (18)$$

همچنین، یک کران بالای واریانس از چنین چگالی مخلوط خطی $\bar{p}_n(x)$ را ثابت می‌کنیم و با نماد $\text{Var}(\bar{p}_n)$ نشان می‌دهیم به قسمی که

$$\text{Var}(\bar{p}_n) \leq \frac{c_s s^{\gamma} e^{c_s s^{\gamma}}}{\gamma^{\gamma}} e^{\gamma h(\bar{p}_n)} (1 + c_2 n^{-1}) \quad (19)$$

با ترکیب روابط (۱۸) و (۱۹) و با فرض $n \rightarrow \infty$ ، کران بالای واریانس هر چگالی تک‌مدی متقارن پیوسته-لیپ‌شیتز با تکیه‌گاه کران‌دار را در آنتروپی توانی ثابت می‌کنیم.

قضیه ۱۱۰۲. برای هر چگالی تک‌مدی متقارن $p(x)$ با تکیه‌گاه کران‌دار $[m-s, m+s]$ ، هرگاه $p(x)$ در شرط لیپ‌شیتز (۱۷) با ثابت $c_s > 0$ صدق کند، واریانس $\text{Var}(X)$ می‌تواند از بالا و پایین به ثابت اسکالر آنتروپی توانی به صورت زیر است:

$$\frac{e^{\gamma h(p)}}{\gamma^{\gamma} e} \leq \text{Var}(X) \leq \frac{c_s s^{\gamma} e^{c_s s^{\gamma}}}{\gamma^{\gamma}} e^{\gamma h(p)}. \quad (20)$$

برای چگالی مخلوط خطی $p(x)$ از (۲۳) کران بالا روی واریانس در جملات آنتروپی توانی همانند چیزی که در قضیه ۶.۲ بیان شد به دست می‌آید که در اینجا اثبات را ارائه می‌دهیم [۱].

با استفاده از مقعر بودن آنتروپی تفاضلی $h(p)$ در توزیع $p(x)$ داریم

$$h(p) \geq \sum_{i=1}^n \alpha_i h(p_i), \quad (25)$$

و در نتیجه

$$e^{\gamma h(p)} \geq e^{\sum_{i=1}^n \gamma \alpha_i h(p_i)} = \prod_{i=1}^n \left(e^{\gamma h(p_i)} \right)^{\alpha_i}. \quad (26)$$

$$\text{برای } p_i(x) = \frac{1}{Z_i(\theta_i, \beta_i)} e^{-\beta_i |x-m|^{\theta_i}} \text{ داریم}$$

$$\sigma_i^{\gamma} = \frac{1}{A(\theta_i)} \cdot e^{\gamma h(p_i)} \quad (27)$$

و در نتیجه

$$e^{\gamma h(p)} \geq \left(\prod_{j=1}^n A(\theta_j)^{\alpha_j} \right) \cdot \left(\prod_{i=1}^n (\sigma_i^{\gamma})^{\alpha_i} \right). \quad (28)$$

با استفاده از معکوس نامساوی میانگین توانی داده شده در [۲۴، ۲۵] یک کران بالا روی میانگین هندسی $\{\sigma_i^{\gamma}\}$ با مرتبه‌های $\{\alpha_i\}$ در جملات میانگین حسابی $\{\sigma_i^{\gamma}\}$ با مرتبه‌های $\{\alpha_i\}$ به صورت زیر داده می‌شود:

$$\sum_{i=1}^n \alpha_i \sigma_i^{\gamma} \leq M(r) \prod_{i=1}^n (\sigma_i^{\gamma})^{\alpha_i} \quad (29)$$

که $M(r)$ در رابطه (۸) تعریف شده است و

$$r = \max_{i,j \in \{1, \dots, n\}} \left\{ \frac{\sigma_i^{\gamma}}{\sigma_j^{\gamma}} \right\}$$

بنا به این که واریانس توزیع مخلوط $p(x) = \sum_{i=1}^n \alpha_i p_i(x)$ برابر $Var(X) = \sum_{i=1}^n \alpha_i \sigma_i^{\gamma}$ است، با ترکیب روابط (۲۸) و (۲۹) داریم:

$$\begin{aligned} Var(X) &= \sum_{i=1}^n \alpha_i \sigma_i^{\gamma} \\ &\leq M(r) \prod_{i=1}^n (\sigma_i^{\gamma})^{\alpha_i} \\ &\leq M(r) \left(\prod_{i=1}^n \left(\frac{1}{A(\theta_i)} \right)^{\alpha_i} \right) e^{\gamma h(p)}. \quad (30) \end{aligned}$$

۳ مثال‌ها

۱. نمودار کران واریانس توزیع نمایی:

توزیع نمایی تنها توزیع پیوسته‌ای است که خاصیت بی‌حافظگی دارد و از این رو بیشتر در حل مسائل احتمال و نظریه صف به کار گرفته می‌شود. هم‌چنین از این توزیع برای مدل‌سازی کردن و آسان ساختن شیوهی حل مسائل واقعی استفاده می‌کنند. این ویژگی تابع را می‌توان این‌طور تفسیر کرد که رویدادهایی را که در گذشته اتفاق افتاده می‌توانیم در نظر نگیریم و از زمان حال

که c_s ثابت لیپ‌شیتز و s نصف اندازه تکیه‌گاه است. در ادامه، واریانس کران بالای را برای حالتی که چگالی‌ها تک‌مدی متقارن لیپ‌شیتز پیوسته باشند، تعمیم می‌دهیم. چگالی تک‌مدی $p(x)$ با تکیه‌گاه کران‌دار $[b - s_l, b + s_r]$ با $s_l, s_r > 0$ را در نظر بگیرید. تعریف کنید: $s = \max\{s_l, s_r\}$. فرض کنید حالت مد این چگالی در $x = b$ رخ دهد. فرض کنید m نشان‌دهنده میانگین این چگالی باشد، $m = \int_{-\infty}^{+\infty} xp(x)dx$. فرض کنید $p(x)$ لیپ‌شیتز پیوسته باشد؛ به عبارت دیگر،

$$|p(x+y) - p(x)| \leq c_s |y| \quad (21)$$

برای هر $x, y \in [b - s_l, b + s_r]$. در نتیجه، واریانس این چگالی $p(x)$ می‌تواند از بالا به منهای ثابت اسکال آنتروپی توانی $(m - b)^{\gamma}$ کران‌دار شود که ثابت یک تابع با ثابت لیپ‌شیتز c_s و اندازه ماکزیم $s = \max\{s_l, s_r\}$ از تکیه‌گاه یک‌بعدی از مد $x = b$ است [۱].

قضیه ۱۲.۲. برای هر چگالی تک‌مدی $p(x)$ روی تکیه‌گاه کران‌دار $[b - s_l, b + s_r]$ با مد $x = b$ و میانگین m که $p(x)$ در شرط لیپ‌شیتز (۲۱) صدق می‌کند، واریانس $Var(X)$ از بالا و پایین کران‌دار به جملات آنتروپی توانی به صورت زیر است:

$$\frac{e^{\gamma h(p)}}{\gamma \pi e} \leq Var(X) \leq \frac{c_s s^{\gamma} e^{c_s s^{\gamma}} M(128(c_s s^{\gamma})^{\gamma})}{\gamma} e^{\gamma h(p)} - (m - b)^{\gamma} \quad (22)$$

با $s = \max\{s_l, s_r\}$ و $M(r)$ که در رابطه (۸) داده شده است [۱].

۴.۲ کران بالای واریانس چگالی‌های مخلوط

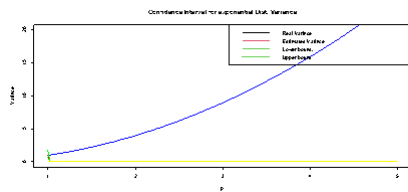
تک‌مدی متقارن

فرض کنید توزیع مخلوط $p(x)$ از تعداد متناهی توزیع‌های نزولی نمایی، همان میانگین m با تضمین تک‌مدی $p(x)$ است. توجه داشته باشید برای هر مؤلفه مخلوط $p_i(x)$ ، ثابت نرمال شده برابر $Z_i(\theta_i, \beta_i) = \beta_i^{-\gamma} \frac{\Gamma(\gamma/\theta_i)}{\Gamma(1/\theta_i)}$ و واریانس $\sigma_i^{\gamma} = \beta_i^{-\gamma} \frac{\Gamma(3/\theta_i)}{\Gamma(1/\theta_i)}$ است. واریانس چگالی مخلوط تک‌مدی متقارن ۲۳ به صورت زیر است:

$$p(x) = \sum_{i=1}^n \left(\frac{1}{Z_i(\theta_i, \beta_i)} e^{-\beta_i |x-m|^{\theta_i}} \right) \quad (23)$$

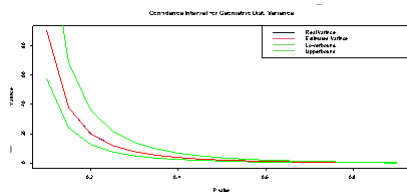
که $\alpha_i > 0$ و $\sum_{i=1}^n \alpha_i = 1$. همه مؤلفه‌های مخلوط $p_i(x)$ دارای همان میانگین m با تضمین تک‌مدی $p(x)$ است. توجه داشته باشید برای هر مؤلفه مخلوط $p_i(x)$ ، ثابت نرمال شده برابر $Z_i(\theta_i, \beta_i) = \beta_i^{-\gamma} \frac{\Gamma(\gamma/\theta_i)}{\Gamma(1/\theta_i)}$ و واریانس $\sigma_i^{\gamma} = \beta_i^{-\gamma} \frac{\Gamma(3/\theta_i)}{\Gamma(1/\theta_i)}$ است. واریانس چگالی مخلوط تک‌مدی متقارن ۲۳ به صورت زیر است:

$$Var(X) = \int_{-\infty}^{\infty} (x - m)^{\gamma} p(x) dx = \sum_{i=1}^n \alpha_i \sigma_i^{\gamma}. \quad (24)$$



شکل ۲: نمودار کران واریانس توزیع نرمال.

۳. نمودار کران واریانس توزیع هندسی:



شکل ۳: نمودار کران واریانس توزیع هندسی.

۴ بحث و نتیجه‌گیری

در این مقاله کران‌های بالا روی واریانس از زیررده‌های توزیع‌های تک‌مدی در جملات آنتروپی توانی را ثابت کردیم. ابتدا چگالی‌های مخلوط تک‌مدی متقارن از توزیع‌های نزولی نمایی و توزیع‌های یکنواخت را در نظر گرفتیم. تنگی کران‌های بالا روی واریانس روی نسبت بین ماکزیمم و مینیمم واریانس‌ها از مؤلفه‌های مخلوط مستقل است. با ارائه یک مثال نقض، نشان دادیم که هرگاه نسبت بین واریانس‌های مؤلفه‌های مخلوط بی‌کران است و لزوماً کران بالا روی واریانس چگالی‌های مخلوط تک‌مدی یکنواخت وجود ندارد که ثابت مدرج شده آنتروپی توانی باشد. همچنین نشان دادیم که واریانس هر چگالی تک‌مدی لیپ‌شیتز-پیوسته با تکیه‌گاه کران‌دار می‌تواند به جملات آنتروپی توانی کران‌دار شود. همه کران‌های بالا روی واریانس چگالی‌های تک‌مدی ارائه‌شده در اینجا مدرج-پایا هستند.

در پردازش سیگنال، سنجش تطبیقی و یادگیری ماشین، جانشین‌های نظری اطلاعات مانند واگرایی کولبک لیبلر، آنتروپی و اطلاعات فشر به‌طور گسترده به‌جای توابع هزینه‌های خاص کار مانند میانگین مجذور خطا یا احتمال خطای طبقه‌بندی مورد استفاده قرار گرفته‌اند. از آنجایی که چنین توابع هزینه‌های خاص کار اغلب غیرقابل حل هستند، جانشین‌های اطلاعاتی به‌عنوان اهداف طبیعی برای توسعه شکل‌های موج یا استراتژی‌های انتخاب حسگر برای جمع‌آوری و یا فیلتر کردن اطلاعات استفاده می‌شوند. نتایج گزارش‌شده در اینجا می‌تواند برای توجیه استفاده از آنتروپی دیفرانسیل به‌عنوان

به بعد را مبدأ زمان قرار بدهیم. مثلاً لامپی که طول عمرش ۱۰ ساعت است و تا ساعت ۶ هنوز نسوخته است را می‌توان مثل یک لامپ نو به حساب آورد.

$$f(x) = \lambda e^{-\lambda x}, \quad \forall x \geq 0.$$

سپس

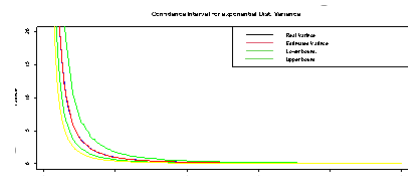
$$\begin{aligned} h(p) &= - \int_0^\infty \lambda e^{-\lambda x} \log(\lambda e^{-\lambda x}) dx h_e(X) \\ &= - \left(\int_0^\infty (\log \lambda) \lambda e^{-\lambda x} dx + \int_0^\infty (-\lambda x) \lambda e^{-\lambda x} dx \right) \\ &= - \log \lambda \int_0^\infty f(x) dx + \lambda E[X] \\ &= - \log \lambda + 1. \end{aligned}$$

به‌سادگی می‌توان نتیجه گرفت:

$$h(p) = - \log \lambda + 1 \Rightarrow \lambda = e^{1-h(p)}$$

واریانس توزیع نمایی:

$$\begin{aligned} v(x) &= \frac{1}{\lambda^2} = e^{2(h(p)-1)}, \\ \frac{e^{2h(p)}}{2\pi e} &\leq e^{2(h(p)-1)} = v(x) \end{aligned}$$



شکل ۱: نمودار کران واریانس توزیع نمایی.

۲. نمودار کران واریانس توزیع نرمال:

$$\begin{aligned} h(p) &= \int_{-\infty}^\infty f(x) \log \left(\frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(x-\mu)^2}{2\sigma^2}} \right) dx \\ &= \int_{-\infty}^\infty f(x) \log \frac{1}{\sqrt{2\pi}\sigma} dx + \\ &\quad \log(e) \int_{-\infty}^\infty f(x) \left(-\frac{(x-\mu)^2}{2\sigma^2} \right) dx \\ &= -\frac{1}{2} \log(2\pi\sigma^2) - \log(e) \frac{\sigma^2}{2\sigma^2} \\ &= -\frac{1}{2} (\log(2\pi\sigma^2) + \log(e)) \\ &= -\frac{1}{2} \log(2\pi e\sigma^2). \end{aligned}$$

از رابطه بالا واریانس توزیع نرمال نتیجه می‌شود که:

$$h(p) = \frac{1}{2} \log(2\pi e\sigma^2) \Rightarrow \sigma^2 = \frac{e^{2h(p)}}{2\pi e}.$$

ماتریس کوواریانس $k \times k$ ، Σ برابر $\Sigma = \prod_{i=1}^k \text{Var}(X_i)$ است. به علاوه، آنتروپی توانی $e^{\text{th}(X)}$ برابر $e^{\text{th}(X)} = \prod_{i=1}^k e^{\text{th}(p_i)}$ است. با استفاده از این حقایق، می‌توانیم داشته باشیم

$$|\Sigma| \leq c e^{\text{th}(X)} \quad (31)$$

برای ثابت مثبت $c > 0$. بنابراین برای این حالت توزیع چندمتغیره با متغیرهای تصادفی مستقل در شرط لیپ‌شیتز-پوسته صدق می‌کند و تک‌مدی است همان‌گونه که آنتروپی $h(X)$ از بردار تصادفی X به $-\infty$ میل می‌کند، تعیین‌کننده ماتریس کوواریانس است که از واریانس‌های توزیع‌های حاشیه‌ای تعیین می‌شود همگرا به ۰ است و توزیع‌های چندمتغیره تصادفی با کران از ماتریس کوواریانس تعیین می‌شوند؛ اما توجه به این نکته مهم است که تعیین‌کننده ماتریس کوواریانس همیشه مرتبط با توزیع k -بعدی نیست هرگاه متغیرهای تصادفی دوبه‌دو مستقل نباشند. برای مثال، هرگاه $X_1 = X_2 = \dots = X_k$ بیان‌کننده واریانس X_i باشد، تعیین‌کننده ماتریس کوواریانس برابر ۰ است. بنابراین، یک سؤال باز جالب پیدا کردن یک متریک مناسب است که یک توزیع چندمتغیره را نشان دهد و یک نابرابری مرتبط برای متریک برحسب توان آنتروپی برای یک رده مناسب از توزیع‌ها استخراج شود.

جایگزینی برای میانگین مربعات خطا در چنین کاربردهایی با ضمانت عملکرد قابل اثبات، زمانی که توزیع خلفی لیپ‌شیتز و تک‌وجهی با پشتیبانی محدود است یا زمانی که می‌تواند با مخلوطی خطی از چگالی‌های به‌طور نمایی کاهشی و چگالی‌های یکنواخت تقریبی شده است، مورد استفاده قرار گیرد.

یکی از مسیرهای جالب تحقیقات آینده، گسترش این کار به یک محیط چند متغیره است. نامساوی‌های مرتبط برای توابع تعداد زیادی از متغیرهای تصادفی یک موضوع تحقیقاتی فعال با کاربردهای متنوع در ابعاد بالا بوده است. آمار، یادگیری ماشین و نظریه اطلاعات. پیشرفت‌های اخیر در نابرابری‌های تمرکز و کاربردهای آن در نظریه اطلاعات را می‌توان در یک مونوگراف عالی یافت [۱۹]. در میان بسیاری از جنبه‌های جالب، ما به‌ویژه به یافتن شرایطی در توزیع چند متغیره برای استخراج نابرابری غلظت از نظر توان آنتروپی علاقه‌مندیم. هرگاه متغیرهای تصادفی با بردار تصادفی k -بعدی $X = [X_1, X_2, \dots, X_k]$ دوبه‌دو مجزا ترکیب و توزیع‌های حاشیه‌ای p_i از هر بردار تصادفی X_i تک‌مدی باشد و در شرایط بحث شده در اینجا صدق کند، واریانس $\text{Var}(X_i)$ هر متغیر تصادفی X_i از بالا به $\text{Var}(X_i) \leq c_i e^{\text{th}(p_i)}$ برای ثابت $c_i > 0$ کران‌دار و تعیین‌کننده

مراجع

- [1] Chung, H. W., Sadler, B. M., Hero, A. O. (2015, September). Bounds on variance for symmetric unimodal distributions. In *2015 53rd Annual Allerton Conference on Communication, Control, and Computing (Allerton)*, (pp. 1235-1240). IEEE.
- [2] Cover, T. M. (1999). *Elements of information theory*. John Wiley & Sons.
- [3] Chen, L. H. (1982). An inequality for the multivariate normal distribution. *Journal of Multivariate Analysis*, **12**(2), 306-315.
- [4] Chernoff, H. (1981). A note on an inequality involving the normal distribution. *The Annals of Probability*, **9**(3), 533-535.
- [5] Cacoullos, T. (1982). On upper and lower bounds for the variance of a function of a random variable. *The Annals of Probability*, **10**(3), 799-809.
- [6] Shannon, C. E. (1948). A mathematical theory of communication. *The Bell system technical journal*, **27**(3), 379-423.
- [7] Jedynek, B., Frazier, P. I., Sznitman, R. (2012). Twenty questions with noise: Bayes optimal policies for entropy loss. *Journal of Applied Probability*, **49**(1), 114-136.
- [8] Wang, H., Yao, K., Pottie, G., Estrin, D. (2004, April). Entropy-based sensor selection heuristic for target localization. In *Proceedings of the 3rd international symposium on Information processing in sensor networks*, (pp. 36-45).

- [9] Bakry, D., Barthe, F., Cattiaux, P., and Guillin, A. (2008). A simple proof of the Poincaré inequality for a large class of probability measures. *Electron. Commun. Probab.*, **13**, .60-66
- [10] Bobkov, S. G. (1999). Isoperimetric and analytic inequalities for log-concave probability measures. *The Annals of Probability*, **27(4)**, .1903-1921
- [11] Madiman, M., Melbourne, J., Xu, P. (2017). Forward and reverse entropy power inequalities in convex geometry. *In Convexity and concentration*, (pp. 427-485). New York, NY: Springer New York.
- [12] Bobkov, S., Madiman, M. (2011). The entropy per coordinate of a random vector is highly constrained under convexity conditions. *IEEE Transactions on Information Theory*, **57(8)**, .4940-4954
- [13] Hensley, D. (1980). Slicing convex bodies—bounds for slice area in terms of the body's covariance. *Proceedings of the American Mathematical Society*, **79(4)**, .619-625
- [14] Webb, S. P. (1996). *Central slices of the regular simplex*. University of London, University College London (United Kingdom).
- [15] Marsiglietti, A., Kostina, V. (2017, June). A lower bound on the differential entropy for log-concave random variables with applications to rate-distortion theory. *In 2017 IEEE International Symposium on Information Theory (ISIT)*, (pp. 46-50). IEEE.
- [16] Walther, G. (2009). Inference and modeling with log-concave distributions. *Statistical Science*, **24(3)**, .319-327
- [17] Bagnoli, M., Bergstrom, T. (2006). Log-concave probability and its applications. *In Rationality and Equilibrium: A Symposium in Honor of Marcel K. Richter*, (pp. 217-241). Springer Berlin Heidelberg.
- [18] Saumard, A., Wellner, J. A. (2014). Log-concavity and strong log-concavity: a review. *Statistics surveys*, **8**, .45
- [19] Raginsky, M., Sason, I. (2013). Concentration of measure inequalities in information theory, communications, and coding. *Foundations and Trends in Communications and Information Theory*, **10(1-2)**, .1-246
- [20] Principe, J. C. (2010). *Information theoretic learning: Renyi's entropy and kernel perspectives*. Springer Science Business Media.
- [21] Fano, R. M. (2008). Fano inequality. *Scholarpedia*, **3(10)**, .6648
- [22] Alimohammadi, M., Alamatsaz, M. H. (2011). Properties of Unimodal Distributions and Some of Their Applications. *Andisheye Amari*, **16(1)**, .47-62
- [23] Feller, W. (1971). *An introduction to probability theory and its applications*. John Wiley & Sons.
- [24] Specht, W. (1960). Zur theorie der elementaren mittel. *Mathematische Zeitschrift*, **74(1)**, .91-98
- [25] Mitrinovic, D. S., Vasic, P. M. (1970). *Analytic inequalities (Vol. 1)*. Berlin: Springer-verlag.
- [26] Tsiligkaridis, T., Sadler, B. M., Hero, A. O. (2014). Collaborative 20 questions for target localization. *IEEE Transactions on Information Theory*, **60(4)**, 2233-2252.
- [27] Travers, N. F. (2013). *Bounds on Convergence of Entropy Rate Approximations in Hidden Markov Processes*. University of California, Davis.

[28] Resnick, S. I. (2007). *Heavy-tail phenomena: probabilistic and statistical modeling*. Springer Science Business Media.

[29] Kwok, T. (2005). Power entropy and its applications. *IEEE Signal Processing Magazine*, **22(5)**, 151-15.

Determining the variance boundaries of single-mode distributions using power entropy

Fateme Ashtab¹ Dr. Mohammad Hossein Dehghan² Dr. Manizheh Sanei Tabas³

Abstract:

Variance and entropy are distinct metrics commonly used to measure the uncertainty of random variables. While variance shows how a random variable spreads more than expected, entropy measures the uncertainty of an information approach; in other words, it measures the average amount of information of a random variable. For both uniform and normal distributions, variance is a measure of power entropy. Finding such a monotonic relationship between variance and entropy for a larger class of these two distributions is very important and useful in signal processing, machine learning, information theory, probability, and statistics. For example, it is used to reduce the errors of estimators and choose a strategy that gives, on average, the greatest or nearly greatest reduction in the entropy of the distribution of the target location. The effectiveness of this method is tested using simulations with mining assay models. In this article, the upper bound of the variance for single-mode distributions, whose tails are heavier than the tails of exponential distributions, is created with the help of power entropy.

Keywords: Power entropy, Variance bounds, unimodal distributions, Lipschitz continuity.

¹Master student, University of Sistan and Baluchestan

²Assistant Professor, Dept. of statistics, Faculty of Math, University of Sistan and Baluchestan, Daneshgah Ave., Zahedan, Iran.

³Assistant Professor, Dept. of statistics, Faculty of Math, University of Sistan and Baluchestan, Daneshgah Ave., Zahedan, Iran.