

مقاله کامل

مجادله در مبانی آمار*

بردلی افرون*

ترجمه علی عمیدی

۱. مقدمه

به نظر می‌رسد که آمار موضوعی دشوار برای ریاضیدانان است، شاید به دلیل خصوصیت مبهم‌گونه و دامنه گسترده‌اش که باعث می‌شود در عرضه مطالب آن از روشن سنتی قضیه‌اثبات عدول شود. لذا ممکن است تسلی خاطری باشد که بگوییم آمار برای آماردانان نیز موضوعی مشکل است. ما اینکه، گذشت تقریباً دو سده از شروع مجادله‌ای بر سر مبانی آمار را گرامی می‌داریم. دو گروه اصلی در این ستیز فلسفی، یعنی بیزگراها و فراوانی‌گراها، به تناوب هر کدام برای مدت زمانی بر دیگری تفوق داشته‌اند، ولی در حال حاضر، فراوانی‌گراها از برتری مختصری برخوردارند؛ گروه سوم کوچکتري که بهتر است آنها را فیشرها بنامیم به هر دو گروه دیگر به شدت ایراد می‌گیرند. آمار، بنا به تعریف، به حالت‌های خاص توجهی ندارد. متوسطها خوراک آماردانان‌اند. «متوسط» در اینجا به مفهوم کلی یعنی به معنای هر حکم خلاصه‌ای درباره جامعه بزرگی از اشیاء است. حکم «متوسط IQی دانشجویان تازه وارد یک دانشگاه برابر ۱۰۹ است» و همچنین حکم «احتمال اینکه سکه‌ای بیفتد و برآمد آن شیر باشد $\frac{1}{2}$ است.» از این گونه حکمها هستند مجادله‌هایی که تقسیم‌کننده دنیای آمارند حول نکته اساسی زیرند: دقیقاً کدام متوسطها برای به دست آوردن استنباط از روی داده‌ها مناسب‌ترند؟ فراوانی‌گراها، بیزگراها، و فیشرها پاسخهای اساساً مختلفی به این سؤال می‌دهند. این مقاله به جای تشریح اصل موضوعی یا تاریخی دیدگاههای مختلف، با تعدادی مثال به پیش می‌رود. مثالها به خاطر درک عموم، به صورتی ساختگی

ساده هستند، ولی خوانندگان مطمئن باشند که این اختلاف دیدگاهها برای داده‌های واقعی هم به همین اندازه صادق‌اند. ذکر این نکته هم مناسب است: این اختلاف دیدگاهها، چه از لحاظ نظری و چه از لحاظ کاربردی اطمینان به آمار ندهند، و در واقع، سهمی در شادابی آن دارند. دستاوردهای جدید مهمی، به ویژه روشهای بیزی تجربی که در بخش ۸ به آنها اشاره می‌شود، مستقیماً از جدال بین دیدگاههای بیزگرا و فراوانی‌گرا حاصل شده‌اند.

۲. توزیع نرمال

همه مثالهای ما مربوط به توزیع نرمال‌اند، که به دلایل مختلف نقشی اساسی در آمار نظری و کاربردی ایفاء می‌کند. یک متغیر تصادفی نرمال یا گاوسی x ، کمیتی است که می‌تواند هر مقداری را روی محور حقیقی اختیار کند، ولی نه با احتمال یکسان. احتمال آنکه x در بازه $[a, b]$ قرار گیرد با مساحت سطح زیر خم معروف زنگدیس گاوس معین می‌شود، یعنی

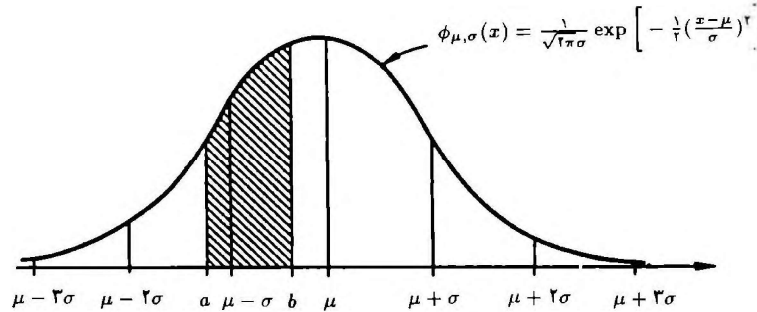
$$\text{Prob}\{a \leq x \leq b\} = \int_a^b \phi_{\mu, \sigma}(x) dx \quad (1.2)$$

که در آن

$$\phi_{\mu, \sigma}(x) = \frac{1}{\sigma\sqrt{2\pi}} \exp\left[-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2\right] \quad (2.2)$$

گرفت. زمینه علائق او بیشتر بخشهای آمار نظری و کاربردی — به خصوص کاربرد روشهای هندسی در مسائل آماری — را در بر می‌گیرد.

بردلی افرون (Bradley Efron) درجه دکتري خود را در آمار در سال ۱۹۶۴ با راهنمایی روبرت میلر (Rupert Miller) از دانشگاه استفرد



شکل ۱ توزیع نرمال. کمیت تصادفی $x \sim \mathcal{N}(\mu, \sigma^2)$ با احتمالی برابر با مساحت سطح هاشورزده، در $[a, b]$ رخ می‌دهد. ۶۸٪ احتمال در بازه $[\mu - \sigma, \mu + \sigma]$ ، ۹۵٪ در بازه $[\mu - 2\sigma, \mu + 2\sigma]$ ، و ۹۹٫۷٪ در بازه $[\mu - 3\sigma, \mu + 3\sigma]$ است.

برای آسانی، چنین متغیر تصادفی را با نماد

$$x \sim \mathcal{N}(\mu, \sigma^2) \quad (3.2)$$

نشان می‌دهیم، که بنا به قرارداد، σ^2 به جای σ ، شناسه دوم است.

شکل ۱، نمودار توزیع نرمال را نشان می‌دهد. قله $\phi_{\mu,\sigma}(x)$ به ازای $x = \mu$ به دست می‌آید. خم به ازای $|x - \mu| > \sigma$ سریعاً به محور x ‌ها نزدیک می‌شود. بیشترین مقدار احتمال، ۹۹٫۷٪، احتمال قرارگرفتن x در بازه $[\mu - 3\sigma, \mu + 3\sigma]$ است. می‌توانیم $x \sim \mathcal{N}(\mu, \sigma^2)$ را به صورت $x = \mu + \epsilon$ بنویسیم، که در آن $\epsilon \sim \mathcal{N}(0, \sigma^2)$ ؛ اضافه کردن مقدار ثابت μ ، صرفاً $\epsilon \sim \mathcal{N}(0, \sigma^2)$ را، واحد انتقال می‌دهد.

پارامتر μ ، «میانگین» یا «امید» کمیت تصادفی x است. اگر از نماد « E » برای نمایش امید استفاده کنیم^۱، نتیجه می‌شود که

$$\mu = E\{x\} \equiv \int_{-\infty}^{+\infty} x \phi_{\mu,\sigma}(x) dx \quad (4.2)$$

ممکن است خواننده مایل باشد $E\{g(x)\}$ را برای تابع دلخواه $g(x)$ صرفاً به عنوان نماد دیگری از انتگرال $g(x)$ نسبت به $\phi_{\mu,\sigma}(x)dx$ در نظر بگیرد

$$E\{g(x)\} \equiv \int_{-\infty}^{+\infty} g(x) \phi_{\mu,\sigma}(x) dx \quad (5.2)$$

از لحاظ شهودی، $E\{g(x)\}$ ، متوسط موزون مقادیر ممکن $g(x)$ است، که برحسب احتمالات $\phi_{\mu,\sigma}(x)dx$ برای بازه‌های بینهایت کوچک $[x, x+dx]$ ، و زنده شده‌اند. به عبارت دیگر، $E\{g(x)\}$ ، متوسط نظری یک جامعه نامتناهی از مقادیر $g(x)$ است، که به نسبت $\phi_{\mu,\sigma}(x)$ رخ می‌دهد.

بنابر تقارن، به آسانی دیده می‌شود که μ درواقع متوسط نظری خود x است وقتی که $x \sim \mathcal{N}(\mu, \sigma^2)$. محاسباتی مشکوکه (هر چند برای کسانی که با تابع گاما کار کرده‌اند آسان است) امید $g(x) = (x - \mu)^2$ را به دست می‌دهد

$$E\{(x - \mu)^2\} = \int_{-\infty}^{+\infty} (x - \mu)^2 \phi_{\mu,\sigma}(x) dx = \sigma^2 \quad (6.2)$$

۱. حرف E اول کلمه Expectation به معنی امید یا میانگین است. - م

پارامتر σ ، که «انحراف معیار» نامیده می‌شود، تغییرپذیری x حول مقدار مرکزی μ را، به صورتی که شکل ۱ نشان می‌دهد، مقیاس‌بندی می‌کند. متغیر تصادفی $\mathcal{N}(1, 10^{-6})$ ، تحت امتحانهای مکرر، تقریباً تغییرپذیری محسوسی نخواهد داشت، از ۱۰۰۰ تکرار، ۹۹۷ تا در $[0.997, 1.003]$ رخ می‌دهند، زیرا $\sigma = 10^{-3}$. یک متغیر تصادفی $\mathcal{N}(1, 10^6)$ ، اگر به زبان گویای نظریهٔ مخابرات صحبت کنیم، تقریباً تماماً نوبه است و نه سدیگال.

توزیع نرمال، ویژگی بسیار سودمندی دارد که باعث می‌شود کار کردن با تعداد زیادی از مشاهدات، به اندازه کار کردن با یک مشاهده تک آسان باشد. فرض کنید $x_1, x_2, \dots, x_n$ ، مشاهده مستقل باشند که هر یک $\mathcal{N}(\mu, \sigma^2)$ است، μ و σ برای همه n تکراریکی هستند. استقلال بدین معناست که مثلاً مقدار x_1 بر مقادیر دیگر اثری ندارد: مشاهده $x_1 > x_2$ ، احتمال ۳۴٪ رخداد $x_2 \in [\mu, \mu + \sigma]$ و غیره را افزایش نمی‌دهد. یک مثال آشنای (غیرنرمال) متغیرهای مستقل $x_1, x_2, \dots$ ، مشاهدات متوالی ریختن تاسی همگن است.

فرض کنید

$$\bar{x} \equiv \sum_{i=1}^n x_i / n \quad (7.2)$$

متوسط مشاهده شده n متغیر مستقل $\mathcal{N}(\mu, \sigma^2)$ باشد. به آسانی می‌توان نشان داد که

$$\bar{x} \sim \mathcal{N}(\mu, \sigma^2/n) \quad (8.2)$$

توزیع $\bar{x}$ نظیر توزیع x_i ی تنهاست، بجز آنکه پارامتر مقیاس آن از σ به $\sigma/\sqrt{n}$ تقلیل یافته است. می‌توانیم، با اختیار n ای به قدر کافی بزرگ، تغییرپذیری $\bar{x}$ حول μ را به سطحی به دلخواه کوچک تقلیل دهیم، وای البته در مسائل واقعی، n محدود است و $\bar{x}$ دارای مؤلفه‌ای غیرقابل تقلیل از تغییرپذیری تصادفی است.

در همهٔ مثالهای ما، σ برای آماردان معلوم فرض خواهد شد. پارامتر مجهول μ ، پارامتر مورد توجه است، و هدف، به دست آوردن استنباطهایی دربارهٔ مقدار μ بر پایهٔ داده‌های $x_1, x_2, \dots, x_n$ است. سرانجام بیشتر در ۱۹۲۵ به این نتیجهٔ اساسی رسید که در این وضعیت، میانگین $\bar{x}$ حاوی همهٔ اطلاعات ممکن دربارهٔ μ است. برای هر مسئلهٔ استنباط دربارهٔ μ ، دانستن

است که نه تنها یک «برآوردکننده نقطه‌ای» بلکه دامنه‌ای از مقادیر موجه برای μ را که سازگار با داده‌ها باشد تعیین کند. از (۸.۲) و شکل ۱ درمی‌یابیم که

$$\text{Prob}\{|\bar{x} - \mu| \leq 2\sigma/\sqrt{n}\} = 0.95 \quad (۴.۳)$$

که هم‌ارز حکم

$$\text{Prob}\{\bar{x} - 2\sigma/\sqrt{n} \leq \mu \leq \bar{x} + 2\sigma/\sqrt{n}\} = 0.95 \quad (۵.۳)$$

است. بازه $[\bar{x} - 2\sigma/\sqrt{n}, \bar{x} + 2\sigma/\sqrt{n}]$ را «بازه اطمینان ۹۵٪» برای μ می‌نامند. نظریه بازه‌های اطمینان را نیمین^۱ در اوایل دهه ۱۹۳۰ مطرح کرد. به عنوان مثال، با فرض $n = 4$ ، $\sigma = 1$ ، داده‌های $x_1 = 1.2$ ، $x_2 = 0.3$ ، $x_3 = 0.7$ ، $x_4 = 0.2$ را به دست آورده‌ایم. در این صورت $\bar{x} = 0.6$ و بازه اطمینان ۹۵٪ برای μ بازه $[-0.4, 1.6]$ است.

همه اینها آن قدر بی‌ضرر و سرراست به نظر می‌رسند که خواننده ممکن است تعجب کند که پس مجادله بر سر چیست. واقعیت این است که تمام نتیجه‌هایی که تا اینجا ارائه شدند از لحاظ ماهیت، «فراوانی‌گرایانه» هستند. یعنی وابسته‌اند به متوسط‌های نظری نسبت به توزیع $\mathcal{N}(\mu, \sigma^2/n)$ متغیر $\bar{x}$ ، با μ ایی که فرض می‌شود یک مقدار واقعی ثابت دارد. نااریبی هم مفهومی فراوانی‌گرایانه است؛ متوسط نظری $\hat{\mu}$ ، یعنی $E(\hat{\mu})$ ، با فرض ثابت بودن μ ، برابر با μ است. نتایج (۳.۳) و (۵.۳)، قضیه کرامر-رائو، احکامی فراوانی‌گرایانه هستند. مثلاً، تعبیر مناسب (۵.۳) آن است که بازه $[\bar{x} - 2\sigma/\sqrt{n}, \bar{x} + 2\sigma/\sqrt{n}]$ مقدار واقعی μ را، در دنباله‌ای طولانی از تکرارهای مستقل $\bar{x} \sim \mathcal{N}(\mu, \sigma^2/n)$ با فراوانی ۹۵٪ دربرمی‌گیرد.

کسی شک ندارد که این نتایج درست‌اند. سؤالی که بزرگراه و طرفداران فیشتر مطرح کرده‌اند این است که آیا متوسط‌های فراوانی‌گرایانه واقعاً به فرایندی استنباطی که دانشمندان در رسیدن از داده‌های خام به مدل‌های ریاضی زیربنایی به کار می‌برند مربوط‌اند یا نه. بعداً به نقطه‌نظر بیزیا برمی‌گردیم.

۴. برآورد بیزی میانگین

تا اینجا μ را ثابت، هر چند مجهول، در نظر گرفتیم. فرض کنید که μ خود متغیری تصادفی است که می‌دانیم دارای توزیع نرمال با میانگین m و انحراف معیار s است،

$$\mu \sim \mathcal{N}(m, s^2) \quad (۱.۴)$$

m و s ثابتی هستند که برای آماردان معلوم‌اند. مثلاً، اگر μ ، مقدار IQ فردی باشد که به تصادف از جامعه ایالات متحده انتخاب شده است، حکم (۱.۴) با $m = 100$ و $s = 15$ (تقریباً) برقرار است. حدود ۶۸٪ از IQها بین ۸۵ و ۱۱۵، حدود ۹۵٪ بین ۷۰ و ۱۳۰، و الخ، هستند. اطلاعاتی نظیر (۱.۴)، که به زبان بیزیا، «توزیع پیشین μ » است، طبیعت فرایند برآورد را تغییر می‌دهد. آزمون‌های استاندارد IQ طوری طرح می‌شوند که اگر فردی را که به تصادف انتخاب کرده‌ایم برای کشف مقدار خاص μ او آزمون کنیم، نمره*

1. J. Neyman

* نمادهای $\bar{x}$ برای نمره آزمون و $\sigma/\sqrt{n}$ برای انحراف معیار آن، برای هماهنگی با نماندگذاری قبلی انتخاب شده‌اند، گرچه نمره‌های IQ واقعی، عملاً متوسط n مورد مستقل آزمون نیستند. همان‌طور که در (۲.۴) بیان شد، نرمال بودن کامل، چیزی جز یک حالت ایدئالی نیست که نمره‌های آزمون واقعی، تقریب آن هستند.

$\bar{x}$ دقیقاً به اندازه آگاهی از مجموعه همه داده‌های $x_1, x_2, \dots, x_n$ مفید است. به زبان امروزی، $\bar{x}$ «آماره بسنده» برای پارامتر مجهول μ است.

نشان دادن بسندگی در این مورد خاص بسیار آسان است. اگر مقدار مشاهده‌شده $\bar{x}$ معلوم باشد، محاسبه متعارف احتمال نشان می‌دهد که کمیت‌های تصادفی $x_1 - \bar{x}, x_2 - \bar{x}, \dots, x_n - \bar{x}$ توزیع توأمی دارند که به هیچ وجه به پارامتر مجهول μ بستگی ندارد. به عبارت دیگر، آنچه بعد از آگاهی آمار دادن از $\bar{x}$ ، در داده‌ها باقی می‌ماند حاوی اطلاعاتی درباره μ نیست. (این اصلی ظاهراً ساده از چشم گاوس و لاپلاس دور مانده بود.)

۳. برآورد فراوانی‌گرایانه از میانگین

آماردان ممکن است بخواهد پارامتر مشاهده‌ناپذیر μ را بر پایه داده‌های مشاهده‌شده $x_1, x_2, \dots, x_n$ برآورد کند. «برآورد معمولاً به معنای مطرح کردن حدس $\hat{\mu}(x_1, x_2, x_3, \dots, x_n)$ وابسته به $x_1, x_2, x_3, \dots, x_n$ است، با در نظر گرفتن اینکه میزان توانایی که برای حدس خود خواهید پرداخت تابع صعودی همواری از خطای برآورد $|\hat{\mu} - \mu|$ است». تابع تاوان معمول، که ما نیز در اینجا به کار می‌بریم، $(\hat{\mu} - \mu)^2$ ، یعنی تابع زیان توان دوم خطاست که ابتدا گاوس آن را پیشنهاد کرد.

اصلی بسندگی فیشتر می‌گوید که ما فقط نیاز به در نظر گرفتن قاعده‌های برآوردی داریم که تابع $\bar{x}$ هستند. بدیهی‌ترین نامزد، خود $\bar{x}$ است

$$\hat{\mu}(x_1, x_2, \dots, x_n) = \bar{x} \quad (۱.۳)$$

این قاعده برآورد، برای μ «نااریب» است: مقدار واقعی μ هر چه باشد،

$$E\bar{x} = \mu \quad (۲.۳)$$

نااریب بودن، همان‌طور که کمی بعد خواهیم دید به هیچ وجه شرطی لازم برای خوب بودن قاعده برآورد نیست، اما این تضمین را می‌دهد که آماردان سعی نکرده فرایند برآورد را به سمت مقداری خاص از μ میل دهد و این تضمین از لحاظ شهودی جاذبه‌ای زیاد دارد.

در صورت استفاده از $\bar{x} = \hat{\mu}$ ، امید تاوان، بنابر (۶.۲) و (۸.۲) عبارت است از

$$E(\hat{\mu} - \mu)^2 = \sigma^2/n \quad (۳.۳)$$

گاوس نشان داد که بین همه قاعده‌های برآورد نااریب $\hat{\mu}(x_1, x_2, \dots, x_n)$ که برحسب $x_1, x_2, x_3, \dots, x_n$ خطی‌اند، قاعده $\hat{\mu} = \bar{x}$ به صورت یکنواخت، به ازای هر مقدار μ ، مقدار $E(\hat{\mu} - \mu)^2$ را مینیمم می‌کند. در اوایل دهه ۱۹۴۰، این حکم تعمیم یافت تا هر برآوردکننده نااریب خطی یا غیرخطی را شامل شود. برهان این مطلب که وابسته به ایده‌هایی است که فیشتر در دهه ۱۹۲۰ ارائه کرد، جداگانه به وسیله کرامر^۱ در سوئد و رائو^۲ در هند عرضه شد. اگر موافق تأکید بر ملاک نااریبی باشیم و زیان توان دوم خطا را به کار ببریم، به نظر می‌رسد که $\bar{x}$ بهترین برآوردکننده μ است. برای آماردان مفید

1. H. Cramér 2. C. R. Rao

اگر به تصادف تعداد زیادی از افراد را انتخاب می‌کردیم و از هر کدام، یک آزمون IQ به عمل می‌آوردیم، و زیرمجموعه‌ای از آنها را که نمره ۱۶۰ دارند، در نظر می‌گرفتیم باید متوسط واقعی IQ این زیرمجموعه ۱۴۸ می‌بود؛ یعنی ۶۸٪ از IQهای واقعی باید در بازه $[۱۴۸ - ۶۷, ۱۴۸ + ۶۷]$ قرار می‌گرفتند.

در وضعیت بیزی چگونه باید μ را برآورد کنیم؟ استفاده از برآوردکننده $\mu^*(\bar{x})$ ، که امید شرطی $(\mu - \mu^*)^2$ را به شرط مقدار مشاهده شده $\bar{x}$ مینیمم می‌کند، طبیعی به نظر می‌رسد. به آسانی از (۳.۴) نتیجه می‌شود که این «برآوردکننده بیزی» عبارت است از

$$\mu^*(\bar{x}) = m + C(\bar{x} - m) \quad (۶.۴)$$

که میانگین توزیع پسین μ به فرض معلوم بودن $\bar{x}$ است. اگر $\bar{x}$ مشاهده شده برابر ۱۶۰ باشد، برآورد بیزی ۱۴۸ است و نه ۱۶۰. گرچه ما یک آزمون IQی ناریب را به کار برده‌ایم، تعداد IQهای واقعی کوچکتر از ۱۶۰ آنقدر بیشتر از IQهای بزرگتر از ۱۶۰ است که امید خطای برآورد را کاهش می‌دهد و نمره مشاهده شده را به سوی ۱۰۰ اریب می‌کند. شکل ۲ این وضعیت را نمایان می‌کند.

بازه‌های اطمینان، مشابه بیزی آشکاری دارند. از (۳.۴) داریم

$$\text{Prob}\{\mu^*(\bar{x}) - 2\sqrt{D} \leq \mu \leq \mu^*(\bar{x}) + 2\sqrt{D}|\bar{x}\} = ۰.۹۵ \quad (۷.۴)$$

نماد $\text{Prob}\{|\cdot|\bar{x}\}$ احتمال شرطی مبتنی بر مقدار مشاهده شده $\bar{x}$ را نشان می‌دهد. در مثال IQ،

$$\text{Prob}\{۱۳۴.۶ \leq \mu \leq ۱۶۱.۸|\bar{x} = ۱۶۰\} = ۰.۹۵$$

هیچ‌کس (بهر بگوئیم، تقریباً هیچ‌کس) با استفاده از روشهای بیزی در وضعیتهایی نظیر مسأله IQ، که برای آن، توزیع پیشین μ به روشنی تعریف

کلی آزمون، یعنی $\bar{x}$ ، نظیر بخش ۳، یک برآوردکننده ناریب با توزیع نرمال برای μ است

$$\bar{x}|\mu \sim \mathcal{N}(\mu, \sigma^2/n) \quad (۲.۴)$$

که $\sigma/\sqrt{n}$ حدود ۷.۵ است. می‌توانیم انتظار داشته باشیم که $\bar{x}$ در ۶۸٪ مواقع حداکثر به فاصله ۷.۵ واحد IQ از μ باشد، و نظایر آن. نماد $\langle \bar{x}|\mu \rangle$ تأکید می‌کند که توزیع $\mathcal{N}(\mu, \sigma^2/n)$ برای $\bar{x}$ ، روی مقدار خاصی که کمیت تصادفی μ اختیار کرده، شرطی است. دایل این تعویض نماد به زودی روشنتر خواهد شد.

قضیه بیز که آن را کشیش برجسته، تامس بیز، در حدود ۱۷۵۰ کشف کرد فرمولی ریاضی برای ترکیب (۱.۴) و (۲.۴) به منظور به دست آوردن توزیع شرطی μ به شرط معلوم بودن $\bar{x}$ است. در این حالت، فرمول بیز به دست می‌دهد

$$\mu|\bar{x} \sim \mathcal{N}(m + C(\bar{x} - m), D) \quad (۳.۴)$$

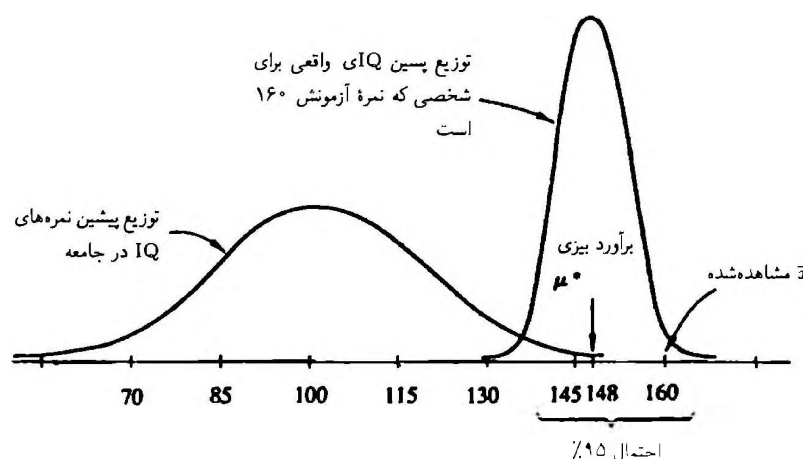
که در آن

$$D = \frac{1}{1/s^2 + n/\sigma^2} \quad \text{و} \quad C = \frac{n/\sigma^2}{1/s^2 + n/\sigma^2} \quad (۴.۴)$$

مثلاً اگر $\bar{x} = ۱۶۰$ (و $m = ۱۰۰, s = ۱۵, \sigma/\sqrt{n} = ۷.۵$)، آنگاه

$$\mu|\bar{x} \sim \mathcal{N}(۱۴۸, (۶.۷)^2) \quad (۵.۴)$$

عبارت (۵.۴) یا عبارت کلیتر (۳.۴)، «توزیع پسین μ به ازای مقدار مشاهده شده $\bar{x}$ » است. عرضه چنین حکمی در چارچوب بیزی امکان دارد، زیرا ما با این فرض شروع کردیم که خود μ متغیری تصادفی است. در چارچوب بیزی، فرایند متوسط‌گیری وارون می‌شود؛ داده $\bar{x}$ در مقدار مشاهده شده‌اش، ثابت فرض می‌شود، در حالی که این پارامتر μ است که تغییر می‌کند. مثلاً در (۵.۴) دیدیم که متوسط شرطی μ به شرط $\bar{x} = ۱۶۰$ برابر ۱۴۸ است.



شکل ۲ نمره‌های IQ در کل جامعه دارای توزیع $\mathcal{N}(۱۰۰, (۱۵)^2)$ هستند. برآورد می‌شود فردی که به تصادف انتخاب شده و در آزمون IQی ناریب نرمال با انحراف معیار ۷.۵ نمره دارای نمره ۱۶۰ است، IQی واقعی برابر با ۱۴۸ است. احتمال اینکه IQی واقعی شخص در بازه $[۱۳۴.۶, ۱۶۱.۸]$ باشد ۹۵٪ است.

ذهنی آماردان را از طریق یک رشته شرط‌بندی فرضی به پارامتر مجهول μ نسبت می‌دهد. این شرط‌بندیها به صورت «آیا مایل‌اید شرط بولی یک به یک روی $\mu > ۸۵$ در مقابل $\mu \leq ۸۵$ ببندید؟ آیا مایل‌اید شرط دو به یک روی $\mu < ۱۵$ در مقابل $\mu \geq ۱۵$ ببندید؟ ...» هستند. کارهای ساوج^۱ و دفتنی^۲ نشان می‌دهند که هر فرد کاملاً منطقی باید همواره بتواند با تعمق کافی به یک توزیع پیشین یکتا (برای خودش) از μ برسد.

رهیافت ذهن‌گرایانه، در مواردی که آماردان (البته معمولاً در همکاری با آزمایشگر) دارای نظرات پیشین هر چند میهمی درباره مقدار واقعی μ است و سعی در به هنگام کردن آنها بر پایه داده $\bar{x}$ مشاهده شده دارد می‌تواند بسیار نمر بخش باشد. چون این رهیافت ذهنی است، در مواردی که عینیت مهم است، مثلاً در انتشار نتایج علمی جدید مجادله‌انگیز، زیاد مورد استفاده قرار نمی‌گیرد.

خط دیگر تفکر بیزی، که می‌توان آن را «بیزگرایی عینی» نامید (که البته معمولاً چنین خوانده نمی‌شود) سعی دارد که در نبود شناخت پیشین، توزیعی پیشین برای μ به دست دهد که هر کسی قبول کند حاکی از نظر پیشین کاملاً بیطرفانه‌ای درباره μ است. در مسئله IQ چنین توزیع پیشین «همواری» باید صورت $\mu \sim \mathcal{N}(0, +\infty)$ را اختیار کند که به معنای $(s^2, \mu \sim \mathcal{N}(0, s^2))$ است، که در آن s^2 به پهنای می‌گراید، از (۳.۴) و (۴.۴) به دست می‌آوریم

$$\mu|\bar{x} \sim \mathcal{N}(\bar{x}, \sigma^2/n) \quad (۸.۴)$$

این نتیجه جاذبه زیادی دارد. برآوردکننده بیزی μ^* برابر با برآوردکننده فراوانی‌گرای $\bar{x} = \hat{\mu}$ است. بازه احتمال بیزی ۹۵٪ (۷.۴) همان بازه اطمینان فراوانی‌گرای ۹۵٪ (۵.۳) است. به علاوه، چون (۸.۴) حکمی بیزی است، نامه شرکت آزمون‌کننده IQ اثری بر آن ندارد. گویی که از مزایای دو جهان فراوانی‌گرا و بیزی‌گرا برخورداریم.

تلاش فراوانی برای تدوین دیدگاه بیزی عینی انجام شده است. خود بیز این دیدگاه را مطرح کرد (اذا ظاهراً با تردید بسیار — مقاله او پس از مرگش با کوششهای یک دوست علاقه‌مند منتشر شد) و لاپلاس آن را بی‌قید و شرط پذیرفت. اعتبار این رهیافت در اوایل ۱۹۰۰ دستخوش تزلزل شد ولی بعدها با کار هارولد جفریز^۳ تا حدی حیات تازه یافته است. یک مشکل این است که توزیع پیشین «هموار» برای μ ، ابتداً توزیعی هموار برای مثلاً μ^2 نیست، و اظهار بی‌اطلاعی ما از همواری توزیع پیشین، مربوط به انتخاب تابع موردنظر از پارامتر مجهول است. درباره مشکل زیان‌بارتری در بخش ۸ بحث شده است: در مسائلی که متضمن برآورد همزمان چندین پارامتر مجهول‌اند، آنچه، توزیع پیشین کاملاً بیطرفانه‌ای به نظر می‌رسد، مستلزم مفروضاتی نامطلوب درباره پارامترها خواهد بود.

۵. برآورد فیشری میانگین

رانلد فیشری یکی از پایه‌گذاران اصلی نظریه فراوانی‌گرایی بود. با این حال در تمام عمرش منتقد این نظریه بود و غالباً با حرارت زیاد از رهیافت فراوانی‌گرایی متعارف انتقاد می‌کرد. انتقادهای او از همان نوع انتقادهای بیزیه بود: چرا

شده و کاملاً معلوم است، مخالفتی ندارد. نظریه بیز، همان‌گونه که خواهیم دید، مزایای قابل توجهی از لحاظ وضوح و سازگاری دارد. این مزایا حاصل این واقعیت‌اند که متوسطهای بیزی تنها متضمن آن مقدار داده‌ای $\bar{x}$ ‌اند که عملاً دیده شده است و نه گردایه‌ای از مقادیر دیگر $\bar{x}$ که به طور نظری ممکن‌اند. مشکلات و مجادله‌ها به این دلیل رخ می‌دهند که آماردانان بیزی می‌خواهند روشهای بیزی را وقتی به کار برند که هیچ توزیع پیشین آشکاری برای μ وجود ندارد یا حتی فراتر از آن، وقتی که واضح است μ می‌مجهول یک ثابت تثبیت‌شده بدون سرشت تصادفی است. (مثلاً، وقتی μ ، ثابتی فیزیکی، نظیر سرعت نور باشد که به‌صورت تجربی برآورده شده است.) اما مسبب این نهضت بیزی، نوعی انحراف فکری نیست، بلکه پیشینه کاملاً مستند ناسازگاریهای ناخوشایند در رهیافت فراوانی‌گراست.

به‌عنوان مثالی از نوع مشکلائی که فراوانی‌گرایان با آن مواجه می‌شوند، مسأله برآورد IQ را دوباره، ولی بدون فرض آگاهی از توزیع پیشین (۱.۴) برای μ در نظر بگیرید. به عبارت دیگر، فرض کنید که تنها مشاهده $\bar{x} \sim \mathcal{N}(\mu, \sigma^2/n)$ ، $\bar{x} = ۷۵$ و $\sigma/\sqrt{n} = ۷$ را داریم و می‌خواهیم μ را برآورد کنیم. با مشاهده $۱۶۰ = \bar{x}$ نتایج بخش ۳ به ما می‌گویند که μ را با $\hat{\mu} = ۱۶۰$ و با بازه اطمینان ۹۵ درصد $[۱۴۵, ۱۷۵]$ برآورد کنیم. حال فرض کنید که شخص فراوانی‌گرا نامهای از شرکتی که آزمون IQ را برگزار می‌کند دریافت کند: «در روزی که نمره $\bar{x} = ۱۶۰$ گزارش شد، ماشین نمردهی آزمون بد عمل کرده است هر نمره $\bar{x}$ کمتر از ۱۰۰ ، برابر ۱۰۰ گزارش شده است. ماشین برای نمره‌های $\bar{x}$ بعد از ۱۰۰ درست عمل کرده است.»

ممکن است به نظر آید که فرد فراوانی‌گرا دلیلی برای نگرانی در این مورد ندارد، زیرا نمره‌ای که دریافت کرده است، یعنی $\bar{x} = ۱۶۰$ ، صحیح گزارش شده است. اما، دلیل استفاده او از $\hat{\mu} = \bar{x}$ برای برآورد μ آن است که $\bar{x}$ بهترین برآوردکننده نااریب است. بد عمل کردن ماشین نمردهی به این معنی است که $\hat{\mu}$ حتی دیگر نااریب هم نیست!

اگر مقدار واقعی μ برابر ۱۰۰ باشد، بد عمل کردن ماشین، به نحوی که در نامه توصیف شده است، $E\bar{x} = ۱۰۳$ را با اریبی $+۳$ امتیاز تولید می‌کند. برای باز یافت نااریبی، آماردان فراوانی‌گرا باید به جای قاعده برآورد $\hat{\mu} = \bar{x}$ قاعده $\hat{\mu}' = \bar{x} - \Delta(\bar{x})$ را به کارگیرد، که در آن تابع $\Delta(\bar{x})$ برای حذف اریبی حاصل از بد عمل کردن ماشین، انتخاب شده است.

جمله تصحیح‌کننده $\Delta(\bar{x})$ برای $\bar{x} = ۱۶۰$ خیلی کوچک خواهد بود، اما اینکه اصلاً تغییری لازم باشد ناراحت‌کننده است. نامه شرکت نمردهی، حاوی اطلاع جدیدی درباره نمره‌ای که عملاً گزارش شده، یا درباره IQ‌ها به طور کلی، نیست. این نامه فقط مربوط به اتفاق بدی است که ممکن بود رخ داده باشد ولی رخ نداده است. چرا باید استنباطمان را درباره مقدار واقعی μ عوض کنیم؟ روشهای بیزی از این نقص بری هستند؛ استنباطهایی که از این روشها به دست می‌آیند تنها به مقدار داده‌ای $\bar{x}$ ای که عملاً مشاهده شده است بستگی دارد، زیرا متوسطهای بیزی نظیر (۶.۴) و (۷.۴) روی $\bar{x}$ مشاهده‌شده شرطی هستند.

روش تحلیل آماردان بیزی، اگر اطلاعی از توزیع پیشینی مانند (۱.۴) نداشته باشد چگونه است؟ در این مورد از دو رهیافت مختلف استفاده می‌شود. شاخه «ذهن‌گرای» آمار بیزی، قبل از گردآوری داده‌ها، توزیع احتمال

همه، این موضوع را یا صورتی از بزرگاری عینی به شمار می‌آورند یا یک اشتباه می‌دانند. کاربرد استدلال فیدوسیالی در برآورد همزمان چندین پارامتر، همان‌طور که در بخش ۸ نشان خواهیم داد ممکن است مشکل بزرگی به بار آورد.

برای اینکه خواننده زیاد برای فیشر متأثر نشود باید گفت دو اندیشه بدیع دیگرش در مورد متوسط‌گیری، استنباط شرطی و تصادفی‌سازی، هنوز بسیار مورد توجه‌اند و موضوعهای دو بخش آینده‌اند.

۶. استنباط شرطی

به دیدگاه فراوانی‌گرا برمی‌گردیم، اما این بار از مفهوم «شرطی‌سازی» که فیشر در ۱۹۳۴ معرفی کرد استفاده می‌کنیم. استنباط شرطی منشأ ابهام عمده دیگری در مجموعه روشهای فراوانی‌گرایانه است که عبارت است از روش انتخاب گردابه‌ای از مقادیر داده‌های نظراً ممکن که برای به دست آوردن استنباطی فراوانی‌گرایانه، متوسط آنها محاسبه می‌شود.

دوباره فرض کنید که متغیرهای نرمال مستقل $x_1, \dots, x_n$ را داریم که هر x_i دارای توزیع $\mathcal{N}(\mu, \sigma^2)$ است، اما قبل از انجام مشاهده، n را با انداختن سکه‌ای همگن به تصادف انتخاب کرده‌ایم

$$n = \begin{cases} 10 & \text{با احتمال } \frac{1}{2} \\ 100 & \text{با احتمال } \frac{1}{2} \end{cases} \quad (1.6)$$

هنوز مایلیم μ را بر پایه داده‌های $x_1, x_2, \dots, x_n$ و n ، با σ ای که مثل قبل ثابتی معلوم است، برآورد کنیم.

توزیع شرطی $\bar{x}$ وقتی مقدار مشاهده‌شده n معلوم است، مثل (۸.۲) عبارت است از

$$\bar{x}|n \sim \mathcal{N}(\mu, \sigma^2/n) \quad (2.6)$$

متوسط مشاهده‌شده، $\bar{x}$ ، در این وضعیت، آماره‌ای بسنده نیست ما هنوز نیاز داریم که بدانیم n برابر با ۱۰ است یا ۱۰۰. بدون این آگاهی باز برآوردکننده ناریب μ ، یعنی $\hat{\mu} = \bar{x}$ را داریم، اما انحراف معیار $\hat{\mu}$ را نمی‌دانیم

در چنین وضعیتی امید مربع خطای $\hat{\mu} = \bar{x}$ چیست؟ اگر متوسط (۳.۳) را به ازای دو مقدار n به دست آوریم، داریم

$$E(\hat{\mu} - \mu)^2 = \frac{1}{2} \frac{\sigma^2}{10} + \frac{1}{2} \frac{\sigma^2}{100} \quad (3.6)$$

فیشر تذکر داده که این محاسبه مضحک است. محاسبه درستی $\hat{\mu}$ به شرط داشتن مقدار مشاهده‌شده واقعی n به وضوح مناسبتر است، یعنی

$$E\{(\hat{\mu} - \mu)^2 | n\} = \begin{cases} \sigma^2/10 & n = 10 \\ \sigma^2/100 & n = 100 \end{cases} \quad (4.6)$$

چیز نادرستی در (۳.۶) نیست، بجز آنکه متوسط مربع خطایی که در آن محاسبه می‌شود به هر مقدار خاص n و $\bar{x}$ ای که واقعاً مشاهده‌شده نامربوط

اگر تعداد بینهایت مقدار $\bar{x}$ به تصادف از $\mathcal{N}(\mu, \sigma^2/n)$ با μ ثابت تولید شود باید برای استنباط درباره آنچه رخ می‌دهد به متوسطهای نظری توجه کنیم؟ ما تنها یک مقدار مشاهده‌شده $\bar{x}$ در هر مسأله استنباط داریم، و فرایند استنباط باید درست بر این مقدار مشاهده‌شده متمرکز شود.

فیشر با روش بیزی هم مخالف بود شاید بدین سبب که آن نوع از مسائل تحلیل داده‌ها که او در کار ژنتیک و کشاورزی خود با آنها مواجه بود چندان برای نسبت دادن توزیعهای پیشین به آنها مناسب نبودند. او با قریحه و ذکاوت خاص خود صورت دیگری از استنباط را به وجود آورد، که نه بیزی بود و نه فراوانی‌گرا.

رابطه $\bar{x} \sim \mathcal{N}(\mu, \sigma^2/n)$ را می‌توان به صورت زیر نوشت

$$\bar{x} = \mu + \epsilon, \quad \epsilon \sim \mathcal{N}(0, \sigma^2/n) \quad (1.5)$$

ما مشاهده $\bar{x}$ را با اضافه کردن نوفه نرمال $\epsilon \sim \mathcal{N}(0, \sigma^2/n)$ ، به میانگین μ غیرقابل مشاهده به دست می‌آوریم. عبارت (۱.۵) را می‌توان به صورت

$$\mu = \bar{x} - \epsilon \quad (2.5)$$

نیز نوشت. آشکار است، یا حداقل برای فیشر آشکار بود، که در وضعیتی که ما چیزی از پیش درباره μ نمی‌دانیم، مشاهده $\bar{x}$ به ما چیزی درباره ϵ نمی‌گوید. فیشر می‌گفت: در واقع اگر ما بتوانیم از روی $\bar{x}$ اطلاعی درباره ϵ به دست آوریم، آنگاه مدل (۱.۵) به صورت تنها، جنبه مهمی از وضعیت آماری را نادیده می‌گیرد. ما این بحث را باز هم، به صورت ملموستر، در بخش بعد پیگیری می‌کنیم.

اگر $\epsilon \sim \mathcal{N}(0, \sigma^2/n)$ ، آنگاه به دلیل تقارن خم زنگدیس حول نقطه مرکزیش $\epsilon \sim \mathcal{N}(0, \sigma^2/n)$ ، تغییر فیشر از (۲.۵) عبارت بود از

$$\mu|\bar{x} \sim \mathcal{N}(\bar{x}, \sigma^2/n) \quad (3.5)$$

این شبیه حکم بیزی عینی‌گرایانه (۸.۴) به نظر می‌رسد، ولی بدون توسل به توزیع پیشین μ به دست آمده است. حکم بازه‌ای حاصل از (۳.۳) عبارت است از

$$\text{Prob}\{\bar{x} - 2\sigma/\sqrt{n} \leq \mu \leq \bar{x} + 2\sigma/\sqrt{n}|\bar{x}\} = 0.95 \quad (4.5)$$

به زبان فیشر این یک حکم احتمالاتی «فیدوسیالی» است. در بحث «فیدوسیالی»، تصادفی بودن، نه در داده $\bar{x}$ ، نظیر محاسبات فراوانی‌گراها، جایی دارد و نه در μ ، نظیر محاسبات بیزیها، بلکه در مکانیسمی که μ ی غیرقابل مشاهده را به $\bar{x}$ مشاهده‌شده تبدیل می‌کند جای می‌گیرد. (در بحث حاضر این مکانیسم، افزودن $\epsilon \sim \mathcal{N}(0, \sigma^2/n)$ به μ است) احکام فیدوسیالی نظیر (۴.۵) به صورت متوسطهایی روی مکانیسم تبدیل تصادفی به دست می‌آیند.

اوج شکوفایی استدلال فیدوسیالی در دهه ۱۹۴۰ بود و از آن زمان به بعد به تدریج از گردونه خارج شده است. اکثر آماردانان معاصر، البته نه

نشان دهیم. در این صورت $\hat{\theta}$ برآوردکننده بدیهی θ است. این برآوردکننده، نالاریب است، $E\hat{\theta} = \theta$ ، و امید مربع خطا عبارت است از

$$E(\hat{\theta} - \theta)^2 = 0.12 \quad (۸.۶)$$

(که با انتگرالگیری عددی به دست آمده است؛ در (۸.۶) قرارداد می‌کنیم که $\hat{\theta} - \theta$ به ازای هر مقدار θ ، بردی از $-\pi$ تا π دارد، بزرگترین خطای برآورد ممکن وقتی رخ می‌دهد که $(\bar{x}_1, \bar{x}_2)$ نسبت به (μ_1, μ_2) متقاطع باشد. این قرارداد مهم نیست زیرا احتمال $|\hat{\theta} - \theta| > \frac{\pi}{4}$ تنها 0.04 است.) واقعیت ناآشنای که فیشر متذکر شد آن است که r همان نقشی را دارد که n در مثالهای (۱.۶) تا (۴.۶) داشت.

(الف) توزیع r به مقدار واقعی θ بستگی ندارد. (برای خوانندگان آشنا با چگالی نرمال دو متغیره، این مطلب از تقارن دایره‌ای توزیع (۵.۶) متعلق به جفت $(\bar{x}_1, \bar{x}_2)$ حول (μ_1, μ_2) نتیجه می‌شود.)

(ب) اگر r کوچک باشد، آنگاه $\hat{\theta}$ درستی کمتر از آنچه (۸.۶) نشان می‌دهد دارد. در حالی که اگر r بزرگ باشد، $\hat{\theta}$ درستی بزرگتر از آنچه (۸.۶) نشان می‌دهد دارد. جدول ۱، امید مربع خطای شرطی $E\{(\hat{\theta} - \theta)^2 | r\}$ را به صورت تابعی از r ارائه می‌کند.

در اصطلاح فیشر، r یک آماره «کمکی» است. به دلیل ویژگی (الف) این آماره مستقیماً اطلاعی درباره θ نمی‌دهد، ولی مقدارش، درستی $\hat{\theta}$ را تعیین می‌کند. اینک به نظر واضح می‌آید که باید ارزیابی خودمان از درستی $\hat{\theta}$ را روی مقدار مشاهده‌شده r شرطی کنیم. اگر $r = 2$ ، نظیر شکل ۳، آنگاه $E\{(\hat{\theta} - \theta)^2 | r\} = 0.18$ بیشتر از امید غیرشرطی $E(\hat{\theta} - \theta)^2 = 0.12$ به دقت $\hat{\theta}$ مرتبط است.

بسیاری از مسائل آماری واقعی دارای این ویژگی‌اند که برخی از مقادیر داده‌ها به وضوح آگاهی‌بخشتر از بقیه هستند. در این‌گونه مسائل، شرطی

است! اگر $n = 100$ ، آنگاه نتیجه (۳.۶) در مورد درستی $\hat{\mu}$ خیلی بدبینانه است. در حالی که اگر $n = 10$ ، خیلی خوش‌بینانه است.

اینها همه ممکن است آن قدر بدیهی به نظر آید که به گفتنش نیرزد. کار شگفت‌فیشر این بود که نشان داد دقیقاً همین وضعیت، با پیچیدگی و ظرافت بیشتر در مسأله‌های دیگر استنباط آماری نیز پیش می‌آید. ما این مطلب را با مثالی شرح می‌دهیم که متضمن برآورد میانگینهای دو توزیع نرمال مختلف، مثلاً μ_1 و μ_2 است. برآورد بر پایه برآوردهای نرمال نالاریب مستقل هر یک از آنهاست

$$\bar{x}_1 \sim \mathcal{N}(\mu_1, 1), \quad \bar{x}_2 \sim \mathcal{N}(\mu_2, 1) \quad (۵.۶)$$

$\bar{x}_1$ و $\bar{x}_2$ از هم مستقل‌اند. (برای سادگی، فرض کرده‌ایم که σ^2/n برای هر دو برابر ۱ باشد) بردار داده‌ای دوبعدی $(\bar{x}_1, \bar{x}_2)$ می‌تواند هر جای صفحه باشد، اما با احتمال زیاد، چندان دورتر از بردار میانگینها، (μ_1, μ_2) ، قرار نمی‌گیرد.

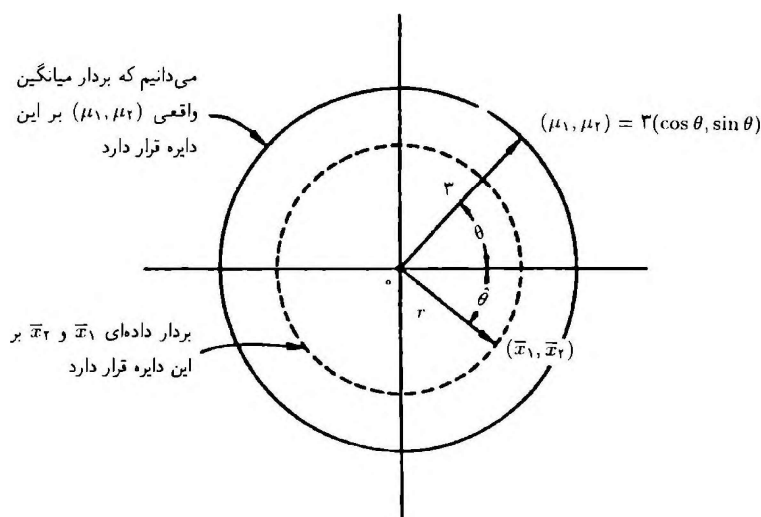
اگر اطلاع دیگری در دست نباشد، احتمالاً (μ_1, μ_2) را با $(\bar{x}_1, \bar{x}_2)$ برآورد می‌کنیم. (لکن بخش ۸ را ببینید!) اما، اینک این فرض را اضافه می‌کنیم که می‌دانیم (μ_1, μ_2) بر دایره به شعاع ۳ و به مرکز مبدا قرار دارد.

$$(\mu_1, \mu_2) = 3(\cos \theta, \sin \theta) \quad -\pi \leq \theta \leq \pi \quad (۶.۶)$$

مسأله آماری، همان‌طور که در شکل ۳ نشان داده‌ایم، برآورد پارامتر مجهول θ بر پایه $(\bar{x}_1, \bar{x}_2)$ است.

فرض کنید مختصات قطبی $(\bar{x}_1, \bar{x}_2)$ را با

$$\hat{\theta} \equiv \arctan(\bar{x}_2/\bar{x}_1), \quad r \equiv \sqrt{\bar{x}_1^2 + \bar{x}_2^2} \quad (۷.۶)$$



جدول ۱ امید مربع فضای شرطی برآورد در مسئله دایره، یعنی $E\{(\hat{\theta} - \theta)^2 | r\}$ به صورت تابعی از آماره کمکی $x = \sqrt{x_1^2 + x_2^2}$ درستی $\hat{\theta}$ وقتی r افزایش می‌یابد صلاح می‌شود. فیشتر استدلال کرد که $E\{(\hat{\theta} - \theta)^2 | r\}$ معیار مناسب‌تری از درستی $\hat{\theta}$ است تا امید غیرشرطی $E(\hat{\theta} - \theta)^2$.

مقدار غیرشرطی

r	۱٫۵	۲	۲٫۵	۳	۳٫۵	۴	۴٫۵	۵	$E(\hat{\theta} - \theta)^2$
$E\{(\hat{\theta} - \theta)^2 r\}$	۰٫۲۶	۰٫۱۸	۰٫۱۴	۰٫۱۲	۰٫۱۰	۰٫۰۹	۰٫۰۸	۰٫۰۷	۰٫۱۲

با مقدار معلوم σ و مقادیر مجهول μ_1 و μ_2 هستند. می‌خواهیم این «فرض صفر» $\mu_1 = \mu_2$ را در برابر «فرض مقابل» $\mu_2 > \mu_1$ که اغلب به صورت

$$A: \mu_2 > \mu_1 \quad H: \mu_2 = \mu_1 \quad (2.7)$$

نوشته می‌شود، آزمون کنیم. (برای مقاصد ما، $\mu_2 < \mu_1$ غیرممکن فرض می‌شود.)

در آزمون فرض، فرض صفر H معمولاً نقش «مخالف‌خوان مصلحتی» را دارد که آزمایشگر سعی می‌کند نظر او را رد کند. مثلاً، H ممکن است پاسخهای مربوط به یک دزوی قدیم و H ها پاسخهای مربوط به یک دزوی جدید باشند که آزمایشگر امید به بهتر بودن آن دارد. چون انگیزه قوی برای بی‌اعتبار کردن H وجود دارد، روشهای آماری محافظه‌کارانه‌ای ارائه شده‌اند که شواهدی در سطح نسبتاً بالا برای اعلام بی‌اعتباری H طلب می‌کنند. نظریه فراوانی‌گرا که در آزمون فرض، نظریه‌ای غالب است لازم می‌داند که احتمال به غلط رد کردن H به نفع A ، وقتی H درست است، همواره کمتر از سطحی پایین، معمولاً 0.05 باشد. آزمونی که این محک را در نظر می‌گیرد آزمون در «سطح 0.05 » برای آزمون کردن H در برابر A نامیده می‌شود. با داده‌هایی مثل (۱.۷)، به نظر طبیعی می‌رسد که $\bar{x} = \sum_{i=1}^n x_i/n$ و $\bar{y} = \sum_{i=1}^n y_i/n$ را حساب کنیم و H را به نفع A رد کنیم اگر

$$\bar{y} - \bar{x} > c \quad (3.7)$$

ثابت c به قسمی انتخاب می‌شود که اگر H راست باشد آنگاه $\text{Prob}\{\bar{y} - \bar{x} > c\} = 0.05$. محاسبات احتمالاتی متعارف نشان می‌دهد که انتخاب صحیح $c = 2.326\sigma/\sqrt{n}$ است. نظریه آزمون ایتیمال که نیم و بیرسن آن را در حدود 1930 ارائه دادند نشان می‌دهد که (۳.۷) در واقع بهترین آزمون فرض H در برابر A در سطح 0.05 است، بدین معنا که اگر A واقعاً راست باشد آنگاه احتمال رد H به نفع A ، ماکسیمم است.

x ها و y هایی که مشاهده می‌کنیم در واقع اندازه‌های مربوط به نوعی از واحدهای آزمایشی، مثلاً دانشجویان تازه وارد یا موشهای سفید یا مبتلایان به سردرد، هستند. فرض کنید این واحدها را با $U_1, U_2, \dots, U_n$ نشان دهیم. موقعیت برای تصادفی‌سازی وقتی پیش می‌آید که آزمایشی داشته باشیم که در آن بتوانیم از قبل تصمیم بگیریم که کدام n واحد، x ها و کدام n واحد، y ها هستند. اگر دقیق نباشیم می‌توانیم درست اولین n واحدی را که به دست می‌آوریم تیمار x و n واحد آخر را تیمار y بگیریم. این آغاز

کردن، به طور شهودی راه صحیح حل مسئله است، اما وضعیتهای معدودی، به وضوح ساختاری مثل مسئله دایره را دارند. گاهی بیش از یک آماره کمکی وجود دارد، و یک مقدار داده، بسته به اینکه کدام آماره کمکی شرطی شود، برآوردهای درستی مختلفی را نتیجه می‌دهد. بیشتر اوقات آماره کمکی وجود ندارد، ولی آماره‌های کمکی تقریبی مختلف به ذهن القا می‌شوند. آنچه مثال دایره آشکار می‌کند آن است که حکمهای فراوانی‌گرایانه نظیر (۸.۶) ممکن است راست ولی نامربوط باشند. دیدگاه فیشتر این بود که متوسط نظری $(\hat{\theta} - \theta)$ نباید روی همه مقادیر داده‌ای ممکن محاسبه شود، بلکه باید تنها روی آن مقادیری محاسبه شود که حاوی مقدار یکسانی از اطلاعات درباره θ باشند. تاکنون صورت‌بندی این حکم به صورتی رضایتبخش میسر نشده است.

بیرگر قبول دارد که شرطی کردن عقیده‌ای درباره درستی $\hat{\theta}$ روی مقدار مشاهده‌شده x کاری صحیح است، اما می‌پرسد که چرا فراتر نرویم و شرطی کردن را روی مقدار مشاهده‌شده خود $(\bar{x}_1, \bar{x}_2)$ انجام ندهیم. این کار در چارچوب فراوانی‌گرایی غیرممکن است، زیرا اگر مجموعه‌ای را که متوسط‌گیری روی آن انجام می‌شود به یک نقطه داده‌ای تقلیل دهیم، برای میانگین‌گیری چیزی باقی نمی‌ماند. استنباطهای بیزی همیشه روی نقطه داده‌ای واقعاً مشاهده شده شرطی هستند. در مسئله دایره، توزیع پیشین هموار طبیعی، توزیع یک‌نواختی روی $\theta \in [-\pi, \pi]$ است. با این توزیع پیشین نتیجه می‌شود که $E\{(\hat{\theta} - \theta)^2 | (\bar{x}_1, \bar{x}_2)\}$ ، همان‌طور که در جدول ۱ نشان داده‌ایم برابر با $\sqrt{\bar{x}_1^2 + \bar{x}_2^2}$ است، پس در این حالت خاص، دیدگاههای بیرگرایی عینی و فراوانی‌گرایی شرطی همانندند. (توجه کنید که در اولین امید، «کمیتی تصادفی است، در حالی که در دومی، این $(\hat{\theta})$ است که تغییر می‌کند.)

۷. تصادفی‌سازی

تصادفی‌سازی هم صورت دیگری از متوسط‌گیری استنباطی است که فیشتر معرفی کرده است. برای اینکه بحث درباره این موضوع ساده شود، باید مسائل آماری را از چارچوب نظریه برآورد به چارچوب «آزمون فرض» منتقل کنیم. داده‌ها اینک، به صورت $2n$ مشاهده نرمال مستقل $x_1, \dots, x_n, y_1, \dots, y_n$ از توزیعهای

$$x_i \sim \mathcal{N}(\mu_1, \sigma^2), \quad y_i \sim \mathcal{N}(\mu_2, \sigma^2), \quad i = 1, 2, \dots, n \quad (1.7)$$

مصیبت است! اولین فرد مبتلا به سردرد ممکن است آنهایی باشند که سردردهای شدیدی دارند، اولین n موش ممکن است آنهایی باشند که با حیوانات سنگینتر در قفس هستند و الخ. در آزمایشی که بدون دقت انجام شده است ممکن است احتمال به غلط رد کردن فرض صفر، به دلیل چنین عاملهای کنترل نشده‌ای، خیلی بیش از $\alpha = 0.05$ باشد.

فیشر در اثر بسیار مهم و مؤثرش دربارهٔ طرح آزمایش، استدلال کرد که انتخاب واحدهای آزمایش باید با تصادفی‌سازی صورت گیرد. یعنی، تخصیص n واحد به گروه تیمار x و n واحد به گروه تیمار y برای هر یک از $(2n!)/(n!n!)$ تخصیص، با احتمال یکسان انجام شود. برای اجرای فرایند تصادفی‌سازی یک ابزار مولد اعداد تصادفی به کار می‌رود.

فیشر خاطرنشان کرد که بررسیهای مبتنی بر تصادفی‌سازی احتمالاً از آن نوع آریبی‌های آزمایشی که در بالا از آن بحث شد بری هستند. مثلاً فرض کنید که نوعی «همتغیر» وجود دارد که به واحدهای آزمایش مربوط است. منظور ما از این اصطلاح، کمیتی است که تصور می‌رود بر مشاهدهٔ مربوط به هر واحد، صرف‌نظر از اینکه چه تیماری داده شده است، اثر می‌گذارد. مثلاً وزن ممکن است همتغیری مهم برای موش سفید باشد. موش سنگین وزن ممکن است کمتر از موش سبک وزن به یک محرک پاسخ دهد. اگر n به صورتی معقول بزرگ، مثلاً برابر 10^4 باشد بسیار نامحتمل است که در آزمایش مبتنی بر تصادفی‌سازی، همهٔ موشهای سنگین در گروه x و همهٔ موشهای سبک در گروه y قرار گیرند. این حکم در مورد هر همتغیری صادق است. چه بدانیم و چه ندانیم که بر پاسخ اثر می‌گذارد، و حتی اگر از وجود آن آگاه نباشیم.

هیچکدام از این مطالب ارتباطی با متوسط‌گیری ندارد. ارتباط، از پیشنهاد بعدی فیشر حاصل می‌شود: که متوسطهای نظری را روی همهٔ توزیعهای نرمال مفروض محاسبه نکنیم، بلکه روی خود فرایند تصادفی‌سازی حساب کنیم. فرض کنید که اگر همهٔ $2n$ واحد آزمایش، تیمار x را دریافت کنند، مشاهدات، $X_1, X_2, \dots, X_{2n}$ باشند، که X_i مشاهدهٔ مربوط به واحد U_i است. حروف بزرگ نشان می‌دهند که اینها مشاهدات فرضی اند و الزاماً داده‌های مشاهده‌شده نیستند. تحت فرض صفر H_0 ، تیمار y همان تیمار x است و بنابراین می‌توانیم فرض کنیم که همهٔ $2n$ واحد، تیمار x را دریافت می‌کنند. در این حالت، داده‌های مشاهده‌شدهٔ $X_1, X_2, \dots, X_n, x_1, x_2, \dots, x_n, y_1, y_2, \dots, y_n$ با مقادیر نظری $X_1, X_2, \dots, X_{2n}$ یکی هستند. فرض کنید $\mathcal{S}(x)$ مجموعهٔ اندیسه‌های آن واحدهایی باشد که واقعاً به تیمار x تخصیص یافته‌اند و $\mathcal{S}(y)$ مجموعهٔ اندیسه‌هایی باشد که به تیمار y تخصیص یافته‌اند. در این صورت اگر H درست باشد،

$$\bar{x} = \sum_{i \in \mathcal{S}(x)} X_i/n, \quad \bar{y} = \sum_{i \in \mathcal{S}(y)} X_i/n \quad (4.7)$$

اگر بررسی بر اساس تصادفی‌سازی انجام شود، آنگاه $\bar{x}$ صرفاً میانگین n مقدار X است که به تصادف انتخاب شده‌اند و $\bar{y}$ متوسط بقیهٔ n تا X هاست. آزمون تصادفی‌سازی (یا «جایگشت») H مشابه با (۳.۷) به صورت زیر طرح می‌شود:

(الف) اگر داده‌های مشاهده‌شدهٔ $X_1, X_2, \dots, X_n, x_1, x_2, \dots, x_n, y_1, y_2, \dots, y_n$ مفروض باشند، تعریف می‌کنیم $x_1 \equiv X_1, u_1 \equiv x_2, u_2 \equiv x_3, \dots$

(ب) برای هر افزاز $\mathcal{P} \equiv \{\mathcal{S}_1, \mathcal{S}_2\}$ از $\{1, 2, \dots, 2n\}$ به دو زیرمجموعهٔ مجزا با اندازهٔ n ، مقدار زیر را حساب کنید

$$(\bar{y} - \bar{x})_{\mathcal{P}} \equiv \sum_{i \in \mathcal{S}_1} u_i/n - \sum_{i \in \mathcal{S}_2} u_i/n \quad (5.7)$$

(پ) همهٔ $(n!)/(2n!)$ مقدار $(\bar{y} - \bar{x})_{\mathcal{P}}$ را به ترتیب صعودی فهرست کنید.

(ت) فرض H را به نفع A رد کنید اگر مقدار $\bar{y} - \bar{x}$ y که عملاً مشاهده شده است در قسمت ۵٪ بالایی فهرست باشد.

در آزمون تصادفی‌سازی، ۵٪ احتمال دارد که H به غلط رد شود، و این احتمال $\alpha = 0.05$ اشاره بر متوسطی دارد که روی همهٔ $(2n!)/(n!n!)$ تخصیص تصادفی انواع تیمارها به واحدهای آزمایش محاسبه می‌شود. آزمون، باز به صورت H به نفع A رد می‌شود اگر $\bar{y} - \bar{x} > c$ است بجز آنکه c دیگر برابر با مقدار ثابت $2.326\sigma/\sqrt{n}$ نیست. در عوض c تابعی از مجموعهٔ مقادیر $\{u_1, u_2, \dots, u_{2n}\}$ است که در (الف) ساخته شد. برای هر مجموعهٔ $\{u_1, u_2, \dots, u_{2n}\}$ ، c به طوری انتخاب می‌شود که در (ت) صدق کند.

آزمون تصادفی‌سازی یک مزیت عمده بر آزمون (۳.۷) دارد. احتمال $\alpha = 0.05$ رد کردن H به غلط، تحت هر فرض صفری که می‌گویید $2n$ مقدار x و y از یک توزیع احتمال، نرمال یا غیر آن، تولید شده‌اند معتبر باقی می‌ماند. درواقع اصلاً نیازی به فرض تصادفی بودن مشاهدات نیست. می‌توانیم فرض صفر را تنها این بگیریم که هر واحد U_i دارای پاسخ تثبیت‌شدهٔ X_i می‌باشد. به آن است، بدون توجه به آنکه تیمار x یا y منظور شده باشد. این حکم اخیر مجدداً تأکید می‌کند که آزمون تصادفی‌سازی باید متضمن یک صورت غیرفراوانی‌گرایانهٔ متوسط‌گیری باشد.

تصادفی‌سازی، یا حداقل استنباط مبتنی بر تصادفی‌سازی، به نظر آماردان بیزی مردود می‌آید. بیزگرای واقعی باید شرطی کردن را روی تخصیص $\{\mathcal{S}(x), \mathcal{S}(y)\}$ واحدها به تیمارهایی که عملاً به کار رفته است انجام دهد، زیرا این، بخشی از داده‌های موجود است، و نه متوسطی روی همهٔ افزازهای ممکن. (به نظر می‌رسد که استدلالهای فیشر دربارهٔ آمارهٔ کمکی دقیقاً اشاره‌ای در همین راستاست که مستقیماً مخالف با تصادفی‌سازی است!)

یک جنبهٔ تصادفی‌سازی، هم فراوانی‌گرایان و هم بیزگرایان را نگران می‌کند. فرض کنید تنها به دلیل بخت بد، در فرایند تصادفی‌سازی، همهٔ موشهای سنگین وزن به تیمار x و همهٔ موشهای سبک وزن به تیمار y تخصیص یابند. آیا باز هم می‌توانیم از آزمون تصادفی‌سازی در سطح $\alpha = 0.05$ برای رد H به نفع A استفاده کنیم؟ به نظر می‌رسد که جواب به وضوح منفی است، اما مشخص کردن راه اجتناب از چنین دامهایی مشکل است. با نگرش به مطلب از راهی دیگر، فرض کنید وزنه‌های $w_1, w_2, \dots, w_{2n}$ موشها را قبل از شروع آزمایش بدانیم. تحت پذیره‌های معقول فراوانی‌گرا، بهترین راه یکنای $\{\mathcal{S}(x), \mathcal{S}(y)\}$ برای تخصیص موشها به تیمارها به منظور آزمون تیمار x در برابر تیمار y وجود دارد که به صورت بهینه، تخصیصهای وزن به دو

جدول ۲ احتمال آنکه $||\bar{x}|| \geq ||\mu||$ همیشه بزرگتر از ۵۰ در صد است. برای حالت $k = 10$ ، احتمالها به ازای مقادیر نه چندان بزرگ و نه چندان کوچک $||\mu||$ خیلی بیش از ۵۰ در صد اند.

$ \mu ^2$	۰	۶	۱۲	۱۸	۲۴	۳۰	۴۰	۶۰
$\text{Prob}\{ \bar{x} > \mu \}$	۱/۰۰۰	۰/۹۶۷	۰/۹۰۴	۰/۸۵۷	۰/۸۲۲	۰/۷۹۵	۰/۷۶۲	۰/۷۱۹

جدول ۲، به نظر می‌رسد که احتمال $||\mu||^2 < 12$ بسیار زیاد است. برای $||\mu||^2$ در بازه $[0, 40]$ که اگر $||x||^2 = 12$ ، روی دادن آن تقریباً قطعی است، بیش از ۷۵٪ مواقع داریم $||\bar{x}|| > ||\mu||$. اما، این یک «هفتاد و پنج درصد» فراوانی‌گرایانه است که با μ ی تثبیت شده و $\bar{x}$ ی که به تصادف مطابق (۱.۸) تغییر می‌کند حساب شده است. استدلالی مشابه استدلال بیزگرایانه عینی که در بخش ۴ معرفی شد نتایج کاملاً متفاوتی به دست می‌دهد.

با فرض ناآگاهی کامل از توزیع پیشین بردار پارامتر μ ، استفاده از توزیع پیشین همواری به صورت $\mu_i \sim \mathcal{N}(0, \infty)$ (یعنی $\mu_i \sim \mathcal{N}(0, s^2)$ با $s^2 \rightarrow \infty$) که μ_i ها به ازای $i = 1, 2, \dots, k$ ، مستقل اند، طبیعی به نظر می‌رسد. این فرض، به توزیع پسین (۸.۴) برای هر پارامتر μ_i ، یعنی

$$\mu_i | \bar{x}_i \sim \mathcal{N}(\bar{x}_i, 1) \quad (4.8)$$

که به ازای $i = 1, 2, \dots, k$ ، μ_i مستقل از هم هستند، منجر می‌شود. البته این، حکمی بیزی با $\bar{x}_i$ های تثبیت شده در مقادیر مشاهده شده است و μ_i ها به طور تصادفی مطابق (۴.۸) تغییر می‌کنند. تعویض نامهای کمیتهای تثبیت شده و تصادفی در جدول ۲، نتیجه می‌دهد که

$$\text{Prob}\{||\mu|| > ||\bar{x}|| \mid ||\bar{x}||^2 = 12\} = 0.904 \quad (5.8)$$

اینکه به نظر می‌رسد که احتمال $||\bar{x}|| > ||\mu||$ زیاد است. در واقع برای هر بردار داده‌ای مشاهده شده $\bar{x}$

$$\text{Prob}\{||\mu|| > ||\bar{x}|| \mid \bar{x}\} > 0.50 \quad (6.8)$$

استدلال فیدوسیالی فیشتر در بخش ۵ هم به (۴.۸) تا (۶.۸) منجر می‌شود.

معادله‌های (۳.۸) و (۶.۸) تناقضی آشکار بین دیدگاههای فراوانی‌گرا و بیزگرا را نشان می‌دهند. کدام صحیح است؟ استدلالی بسیار شگفت‌آور و متقاعدکننده به نفع محاسبه فراوانی‌گرای (۳.۸) وجود دارد. این استدلال را چارلز استاین در اواسط دهه ۱۹۵۰ عرضه کرده است و به برآورد μ بر پایه بردار داده‌ای $\bar{x}$ (یا هم‌ارز آن، برآورد پارامترهای $\mu_1, \mu_2, \dots, \mu_k$ بر پایه $\bar{x}_1, \bar{x}_2, \dots, \bar{x}_k$) مربوط است.

برآوردکننده‌ای بدیهی عبارت است از

$$\hat{\mu}(\bar{x}) = \bar{x} \quad (7.8)$$

که هر μ_i را با $\bar{x}_i$ به صورتی که در (۱.۳) آمده، برآورد می‌کند. امید ریاضی مربع زیان خطای این برآورد برای هر بردار پارامتر μ برابر است با:

$$E||\hat{\mu} - \mu||^2 = k \quad (8.8)$$

گروه را یکسان می‌کند. آماردانانی که در مکتب فیشتری کار آموخته‌اند چنین «طرحهای آزمایش بهینه» را به دلیل حذف شدن عنصر تصادفی‌سازی، مشکل می‌پذیرند.

۸. پدیده استاین

خواننده ممکن است توجه کرده باشد که بحث و معادله‌ای که تا اینجا شرح داده شد، بیشتر علمی بود تا عملی. همه نجله‌های فلسفی توافق دارند که با عدم شناخت پیشین، بازه $[\bar{x} - 2\sigma/\sqrt{n}, \bar{x} + 2\sigma/\sqrt{n}]$ یک بازه ۹۵٪ برای μ است، عدم توافق، درباره معنای «۹۵٪» است. این وضعیت وقتی به صورت بدی درمی‌آید که برآورد همزمان پارامترهای زیادی را در نظر بگیریم. فرض کنید که باید چند میانگین نرمال $\mu_1, \mu_2, \dots, \mu_k$ را برآورد کنیم که برای هر یک از آنها یک برآورد نرمال نااریب مستقل را به صورت

$$\bar{x}_i \sim \mathcal{N}(\mu_i, 1) \quad i = 1, 2, \dots, k \quad (1.8)$$

مشاهده می‌کنیم. (برای راحتی، باز هم واریانس σ^2/n را برابر ۱ گرفته‌ایم.) وقتی چندین پارامتر برای برآورد کردن وجود دارند، مشابه طبیعی زیان مربع خطا، مربع فاصله اقلیدسی است. برای ساده کردن نمادگذاری، فرض کنید $\bar{x} = (\bar{x}_1, \bar{x}_2, \dots, \bar{x}_k)$ بردار متوسطهای مشاهده شده باشد، $\mu = (\mu_1, \mu_2, \dots, \mu_k)$ بردار واقعی میانگین‌هاست، و $\hat{\mu} = (\hat{\mu}_1, \hat{\mu}_2, \dots, \hat{\mu}_k)$ بردار برآوردها باشد. در این صورت مربع خطا، یعنی تاوان برآورد بد، عبارت است از

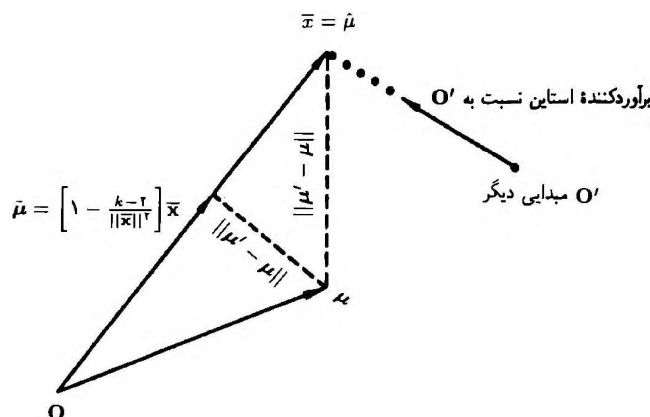
$$||\hat{\mu} - \mu||^2 = \sum_{i=1}^k (\hat{\mu}_i - \mu_i)^2 \quad (2.8)$$

قبل از ادامه مسأله برآورد کردن μ بر پایه $\bar{x}$ ، یک حکم مقدماتی ولی مهم را تذکر می‌دهیم. این حکم، که خوانندگان آشنا با توزیع نرمال چند متغیره می‌توانند آن را در یک سطر ثابت کنند، آن است که برای هر بردار پارامتر μ داریم

$$\text{Prob}\{||\bar{x}|| > ||\mu||\} > 0.50 \quad (3.8)$$

یعنی، بردار داده‌ای $\bar{x}$ گرایش به این دارد که بیشتر از بردار پارامتر μ ، هر چه باشد μ ، از مبدأ دور باشد. جدول ۲ نشان می‌دهد که به ازای $k = 10$ ، برای مقادیر نه چندان بزرگ و نه چندان کوچک $||\mu||$ ، احتمال خیلی بزرگتر از ۵۰ در صد است.

فرض کنید که $k = 10$ ، بردار داده‌ای $\bar{x}$ را با مربع طول $||\bar{x}||^2 = 12$ مشاهده کنیم. فرض کنید که اطلاع قبلی نیز درباره μ نداریم. با ملاحظه



شکل ۴ برآورد استاتین $\hat{\mu}$ با انقباض برآورد بدیهی $\bar{x} = \mu$ به سمت O به دست آمده است. ضریب انقباض، هر چه $||\bar{x}||$ نزدیکتر به O باشد فرین تر است. استاتین و جیمز نشان دادند که برای هر μ ، $E||\hat{\mu} - \mu||^2 < E||\bar{\mu} - \mu||^2$ می‌توانیم هر مبدأ دیگر O' را انتخاب کنیم و برآورد استاتین دیگری به دست آوریم که بر $\hat{\mu}$ نیز ارجح باشد.

نتیجه استاتین برای فراوانی‌گراها و بیزیها مشکلات زیادی پدید آورده است که در اینجا نمی‌توانیم به آنها بپردازیم. استلزامهای این نتیجه برای بیزگرایان عینی و فیدوسالیها بسیار نگران‌کننده بوده است. توزیع پیشین ظاهراً همواری که به (۴.۸) منجر می‌شود ابتدا هموار نیست: این توزیع، بردار پارامتر را واهی دارد که نسبتاً دور از هر مبدأ از قبل انتخاب شده O' باشد. اگر نظریه رضایتبخشی از استنباط بیزگرای عینی وجود داشته باشد، برآوردکننده استاتین نشان می‌دهد که باید بسیار پیچیده‌تر از آنچه باشد که قبلاً انتظار می‌رفت.

دردسر مسأله برآورد چند پارامتری این نیست که سخت‌تر از برآورد کردن یک پارامتر است. بلکه آسانتر هم هست، به این معنا که هم‌زمان، با مسائل زیادی سر و کار دارد که ممکن است اطلاعاتی اضافی به دست دهند که در غیر این حالت در دسترس نیستند. دردسر، مربوط به یافتن و به‌کار بستن این اطلاعات اضافی است. مدل بیزی (۱.۴) را در نظر بگیرید. اگر تنها یک μ برای برآورد کردن داشته باشیم، پارامترهای این مدل را باید صرفاً بر اساس اعتقاد یا تجربه شخصی اختیار کرد. اما، اگر چندین میانگین $\mu_1, \mu_2, \dots, \mu_k$ را برای برآورد کردن داشته باشیم که هر یک مستقلاً از یک جامعه $N(m, s^2)$ استخراج شده باشد، داده‌های $\bar{x}_1, \bar{x}_2, \dots, \bar{x}_k$ به ما امکان می‌دهند که m و s^2 را برآورد کنیم، به جای اینکه مقادیری برای آنها مفروض بگیریم. با قراردادن برآوردها در (۶.۴)، یک «قاعده بیز تجربی» که خیلی شبیه قاعده استاتین (۱۱.۸) است به دست می‌آید. نظریه بیزی تجربی که ابتدا هربرت رابینز^۱ در اوایل دهه ۱۹۵۰ آن را مطرح کرد، نویدبخش مصالحه‌ای نسبی بین فراوانی‌گراها و بیزگراهاست.

۹. چند نکته پایانی

علی‌رغم مجادله بر سر مبانی آمار و تا حدی به سبب آن، رشد و شکوفایی آمار همچنان ادامه دارد. بدون مبالغه میلیونها تحلیل آماری در ۵۰ سال

آنچه استاتین نشان داد آن است که اگر k ، تعداد میانگینهایی که باید برآورد شوند، بزرگتر از ۳ یا برابر با آن باشد، آنگاه برآوردکننده

$$\hat{\mu}(\bar{x}) = \left[1 - \frac{k-2}{||\bar{x}||^2}\right] \bar{x} \quad (۹.۸)$$

برای هر μ دارای

$$E||\hat{\mu} - \mu||^2 < k \quad (۱۰.۸)$$

است! (این صورت خاص $\hat{\mu}$ با همکاری جیمز^۱ در ۱۹۶۰ به دست آمد.) از دیدگاه فراوانی‌گرا، $\hat{\mu}$ بردار μ را به طور یکنواخت بهتر از $\bar{\mu}$ برآورد می‌کند. این برآوردکننده از دیدگاه بیزگرا هم بهتر است: با فرض هر توزیع پیشین برای μ ، برآورد کردن این بردار با $\hat{\mu}$ به جای $\bar{\mu}$ ، باعث می‌شود امید مربع خطای برآورد کلی کمتر باشد (در این حال، متوسط‌گیری به شرط تصادفی بودن μ و تصادفی بودن $\bar{x}$ انجام می‌شود).

برآوردکننده استاتین مبتنی بر (۳.۸) است. چون $||\hat{\mu}|| = ||\bar{x}||$ با احتمال زیاد گرایش به بزرگتر بودن از $||\mu||$ دارد، یک ضریب انقباض $\left[1 - (k-2)/||\bar{x}||^2\right]$ برای به دست آوردن برآوردکننده‌ای که به μ نزدیکتر است به کار می‌رود. ضریب انقباض، وقتی $||\bar{x}||^2$ کوچک است مؤثرتر است. با $k=10$ ، $||\bar{x}||^2 = 12$ ، داریم $\bar{\mu} = [0.333]\bar{x}$ ، اگر به جای آن $||\bar{x}||^2 = 800$ ، آنگاه $\bar{\mu} = [0.99]\bar{x}$ ، شکل ۴ شمایی از این موضوع را نشان می‌دهد.

توجه کنید که مبدأ O نقشی خاص در ساختن $\hat{\mu}$ دارد، گرچه O در بیان مسأله برآورد نقشی ندارد. درواقع، می‌توانیم مبدأ را به هر نقطه دیگر فضای k بعدی، مثلاً O' ببریم، و برآورد استاتین دیگری به صورت

$$\hat{\mu}' = O' + \left[1 - \frac{k-2}{||\bar{x} - O'||^2}\right] (\bar{x} - O') \quad (۱۱.۸)$$

به دست آوریم که باز هم به طور یکنواخت بهتر از $\hat{\mu}$ است.

5. R. A. Fisher, Statistical Methods and Scientific Inference, Oliver and Boyd, London, 1956.

[آخرین رساله عمده فیشر، بحث استدلالهای فیدوسیالی و شرطی در حد قانع کننده ای پیشرفته اند. ولی این مطالب را باید با تردید و احتیاط مطالعه کرد!]

6. H. Jeffreys, Theory of Probability, 3rd Edition, Clarendon Press, Oxford, 1967.

[مهمترین رساله نوین درباره بیزگرایی عینی.]

7. D. V. Lindley, Bayesian Statistics – A Review. SIAM Monographs in Applied Mathematics, SIAM, Philadelphia, (1971).

[مرجع خوبی برای دیدگاه بیزی، هم عینی و هم ذهنی.]

8. ———, The future of statistics – a Bayesian 21st century. Proceedings of the Conference on Directions for Mathematical Statistics, (1974). Special supplement to Advances in Applied Probability, September 1975.

[مقاله های هوبر و رابینز هم به آینده آمار مربوط اند.]

9. L. J. Savage, The Foundations of Statistics. Wiley, New York, 1954.

[این کتاب توجه به دیدگاه بیزی ذهنی را دوباره برانگیخت.]

- Bradley Efron, "Controversies in the foundations of statistics", Amer. Math. Monthly, (4) 85 (1978) 231-246

* بردلی افرون، بخش آمار دانشگاه استفرد، آمریکا

گذشته اجرا شده اند که این تعداد تحلیل محققاً بر قابلیت اعتماد روشهای معمول آماری، وقتی با احتیاط به کار گرفته شوند، صحه می گذارد. من در کار مشاوره ای خود همواره توان روشهای استاندارد را در تحلیل و توضیح مجموعه های داده ای بزرگ از رشته های علمی مختلف می بینم. از لحاظی، مهمترین باور، صرف نظر از مجادله های فراوانی گرایان و بیزیها آن است که کم و بیش تکنیکهایی کلی برای استخراج اطلاعات از داده های نوفه ای وجود دارند که برای تقریباً هر زمینه از نیازها مناسب اند. به عبارت دیگر، آماردانان بر این باورند که آمار اصالتاً به عنوان یک نظام وجود دارد حتی اگر نتوانند در ماهیت دقیق آن هم نظر باشند.

چشم انداز آینده چیست؟ در یک همایش اخیر، دنیس لیندلی^۱ از دانشگاه لندن، سخنرانی تحت عنوان «آینده آمار – قرن بیست و یکم قرن بیزگرایی» ایراد کرد. احتمال ذهنی شخصی من برای وقوع این پیشامد ۱۵٪ است. مزیت عمده بیزگرایی ذهنی، که پروفیسور لیندلی به آن اشاره کرده است، سازگاری منطقی آن است. معمولاً فیلسوفانی که مبانی استنباط علمی را بررسی می کنند، سرانجام از فراوانی گرایی روی گردان می شوند و به استدلال بیزی روی می آورند.

اما، سازگاری کافی نیست. بیزگرایی ذهنی باید با چالش عینیت علمی مواجه شود. این سنگر نهایی دیدگاه فراوانی گراست. اگر سده بیست و یکم، سده بیزی باشد، گمان من آن است که ترکیبی از بیزگرایی ذهنی، عینی و تجربی است، که پیچیدگی و تناقض آن چندان کمتر از وضع موجود نخواهد بود. پیچیدگی مسائلی که بررسی آنها را از آماردانان می خواهند با سرعتی ترسناک رو به افزایش است. این روزها سر و کار داشتن با مجموعه های داده ای میلیونی، و مدلهایی با هزاران پارامتر، غیرمعمول نیست. همان طور که از بخش ۸ برمی آید، این روند، احتمالاً مشکلات ایجاد یک نظریه آمار منطقی سازگار را بیشتر می کند.

مراجع

1. V. Barnett. Comparative Statistical Inference, Wiley, New York, 1973.

[شامل بحث روشی درباره دیدگاه فراوانی گرا در مقایسه با روشهای بیزی.]

2. A. Birnbaum, On the foundations of statistical inference (with discussion). J. Amer. Statist. Assoc. 57 (1962) 269-326.

[بحث عمیقی درباره مجادلات بر سر مبانی. این بحث فی نفسه عالی است.]

3. B. DeFinetti, Foresight: Its logical laws, its subjective sources. Studies in Subjective Probability, ed. by M. Kyburg and H. Smokler, 93-158, Wiley, New York, 1964

[این اثر مهم را افراطی ترین، و (همراه با ساوج) پرنفوذترین، آماردان ذهنی گرای عصر ما در ۱۹۳۵ نوشته است. این جلد شامل مقاله هایی نیز از زون، بورل، رمزی، کویمن، و ساوج است.]

4. B. Efron, Biased versus unbiased estimation. Advances in Math., No. 3, 16 (1975) 259-277.

[برآوردکننده استاین در نظریه و عمل.]

1. Dennis Lindley