

فلسفه آمار

(۱)

دنيس ليندلى *

ترجمه حميد پزشک

۱. مقدمه

این مقاله به جای بحث راجع به یک مسأله معین، چشم‌اندازی ارائه می‌دهد که از طریق آن می‌توان غالب موضوعات آماری را بررسی کرد. در عنوان مقاله، «فلسفه» به معنی «مطالعه اصول کلی شاخه خاصی از دانش، تجربه، یا فعالیت» [۲۲] آمده است. کلمه «فلسفه» این روزها بیشتر یادآور مفاهیم مجرد و جدا از واقعیت است. ولی قصد من در اینجا این است که از تجرید بیش از حد دوری جسته به موضوعات عملی مرتبط با آمار بپردازم. اگر آمارپیشه‌ای که این مقاله را می‌خواند آن را برای خود سودمند نیابد نوشته من به یکی از هدفهایش نرسیده است. در مقاله تلاش می‌شود روشی برای نگرش به آمار عرضه شود که ما آماردانان به کمک آن بهتر بتوانیم راه حل صحیح را برای مسائلی که به آنها برمی‌خوریم بیابیم. از بحثهای فنی اغلب پرهیز کرده‌ام نه به دلیل آنکه این بحثها مهم نیستند بلکه چون خواسته‌ام بیشتر توجه‌مان معطوف به روشن شدن این موضوع باشد که چگونه یک مسأله آماری را می‌توان بررسی کرد. در برخی جاها جزئیات را نیاورده‌ام تا ایده‌های اساسی مشخص‌تر شوند. برای مثال؛ چگالیهای احتمال بدون اشاره صریحی به اندازه کنترل‌کننده آنها مورد استفاده قرار گرفته‌اند.

مقاله با پذیرش این نظر آغاز می‌شود که مقولات آماری با عدم حتمیت سروکار دارند، و در ادامه نتیجه‌گیری می‌کند که عدم حتمیت را فقط می‌توان با احتمال اندازه‌گیری کرد. با این نتیجه‌گیری شرح نظام‌مند «استنباط» بر اساس حساب احتمال میسر می‌شود، که نشان خواهیم داد متفاوت با بعضی رویکردهای متداول است. اصل درستمایی از نقشی پایه‌ای که احتمال ایفا می‌کند نتیجه می‌شود. سپس نقش تحلیل داده‌ها در ساختن الگوهای احتمال و طبیعت الگوها مورد بحث قرار می‌گیرد. این بحث به روشی برای تصمیم‌گیری می‌انجامد و طبیعت مخاطره بررسی می‌شود. روش علمی و کاربرد آن در چند مقوله حقوقی در چارچوب احتمالاتی توضیح داده خواهد شد. نتیجه این است که به این ترتیب، مجموعه‌ای از روشهای عمومی و

رضایت‌بخش آماری داریم که اجرای آنها باید موجبات پیشرفت کار آمار را فراهم کند.

دیدگاه ارائه‌شده در این مقاله بیشتر به ساخت الگو توجه دارد تا به استنباط صوری، و از این لحاظ با جدیدترین آرا توافق دارد. یک دلیل برای این تغییرکانون توجه، این است که استنباط صوری یک روال یا کار نظام‌مند در قالب حساب احتمال است. ولی ساخت الگو نمی‌تواند خیلی نظام‌مند باشد. این مقاله حاصل تجربیات من از ششمین کنفرانس آمار بیزی است که در شهر والنسیا در ژوئن سال ۱۹۹۸ برگزار شد [۵]. اگر چه من تحت‌تأثیر کیفیت کلی مقالات عرضه‌شده و پیشرفتهای به‌دست آمده قرار گرفتم ولی به نظرم رسید که بسیاری از شرکت‌کنندگان فلسفه بیزی را آنچنان که شایسته است قدر نمی‌نهند. این مقاله حاصل تلاشی است برای توضیح دیدگاه من درباره این فلسفه، و بازتاب پنجاه سال تجربه آماری و تحول دیدگاه شخصی من است که از فراوانی‌گرایی به مکتب عینی‌گرایی بیز و سپس به نگرش ذهنی‌گرا که در این مقاله درباره آن بحث می‌شود، روی آورده‌ام. هیچ تلاشی نکرده‌ام که فلسفه‌های دیگر را به تفصیل تحلیل کنم، فقط نشان داده‌ام که نتایج آنها در چه مواردی با نتایج ما مغایر هستند و روشهای عملی حاصل را با هم مقایسه کرده‌ام.

۲. آمار

برای بحث راجع به فلسفه آمار لازم است تا حدی روشن کنیم که منظور از آمار چیست؛ البته نه آنکه تعریف دقیقی بیاوریم، بلکه در حدی که بتوانیم تصویری کلی از موضوع داشته باشیم. نظر این مقاله این است که آمار مطالعه عدم حتمیت است [۲۵b] و آماردانان متخصص رسیدگی به مسأله عدم حتمیت‌اند. آنان ابزارهایی نظیر خطاهای استاندارد و سطوح معنی‌دار بودن را برای اندازه‌گیری عدم حتمیت ارائه داده‌اند. برای تحقیق در این مطلب که این تعریف چقدر با کاری که ما عملاً انجام می‌دهیم مطابقت دارد،

۳. عدم حتميت

با پذيرش اين نظر كه آمار مطالعه عدم حتميت است، تحقيق درباره اين پديده لازم مي آيد. رهيافت علمي به عدم حتميت به معني اندازه گيري آن است زيرا همچنان كه كلوين گفته است، فقط با انتساب عدد به مفهوم علمي است كه آن مفهوم را مي توان به درستي درك كرد. دليل اندازه گيري فقط اين نيست كه مي خواهيم برداشت دقيق تري مثلاً از اين موضوع داشته باشيم كه حتميت ما در پيش بيني وضع بازار بورس كمتر از پيش بيني طلوع خورشيد در روز بعد است. بلكه تلفيق عدم حتميت ها هم هست. فقط در حالات استثنائي ممكن است يك مورد عدم حتميت در يك مسأله وجود داشته باشد و در حالت كلي موارد متعددي از آن وجود دارد. هنگام جمع آوري داده ها، در واحد نمونه گيري عدم حتميت وجود دارد و در عدد گزارش شده در نمونه گيري نيز عدم حتميت هست. در يك مسأله آماري نوعي هم در داده ها و هم در پارامتر عدم حتميت وجود دارد. بنابراين مسأله اساسي، تركيب عدم حتميت ها است. قوانين تركيب اعداد بسيار آسان است. اعداد از دو طريق جمع و ضرب تركيب مي شوند و اين تركيب به غناي مفاهيم رياضي منجر مي شود. ما مي خواهيم عدم حتميت ها را اندازه بگيريم تا بتوانيم آنها را تركيب كنيم. سياستمداري گفته است كه قيود را به اعداد ترجيح مي دهد. ولي متأسفانه تركيب قيود مشكل است.

اين اندازه گيري چگونه ميسر مي شود؟ همه اندازه گيري ها بر مبناي مقايسه با يك استاندارد صورت مي پذيرد. ما براي اندازه گيري طول از خط قرمز-نارنجي مربوط به ايزوتوپ كريتون ۸۶ به عنوان استاندارد استفاده مي كنيم. نقش اساسي مقايسه دال بر اين است كه هيچ مطلقي در دنياي اندازه گيري وجود ندارد. اين نکته اي است كه در بخش ۱۱ به آن برمي گرديم. بنابراين لازم است استاندارد براي عدم حتميت بيابيم. استانداردهاي گوناگوني پيشنهاده شده است اما ساده ترين آنها يعني بازياي شانس از نظر تاريخي اولين پيشنهاده است. اين بازياها متضمن نخستين مواردی از عدم حتميت هستند كه به طور نظام مند بررسي شده اند. بنابراين بگذاريد يك بازی ساده را به عنوان استاندارد به كار ببريم.

فرض كنيد كيسه اي شامل N توپ كوچك است كه آنقدر كه مهندسي مدرن توانسته، آنها را شبیه به هم ساخته است. تصور كنيد يك توپ به طور تصادفي از كيسه بيرون آورده مي شود. براي اينكه مفهوم اين جمله روشن باشد بايد معني «تصادفي» مشخص شود. فرض كنيد كه توپها از ۱ تا N شماره گذاري شده باشند و همچنين فرض كنيد، بدون آنكه هزينه اي براي شما منظور شود، چنانچه توپ شماره ۵۷ را بيرون بياوريد جايزه اي به شما تعلق گيرد. اکنون اين حالت را در نظر بگيريد كه همان جايزه به شما تعلق مي گيرد اگر شماره ۱۲ را بيرون بياوريد. اگر اين دو حالت يا به طور كلي تر هر شماره بين ۱ تا N براي شما تفاوت نداشته باشد، آنگاه توپ به طور تصادفي بيرون آورده شده است. توجه كنيد كه تعريف تصادفي بودن در اينجا ذهني و وابسته به شماست. آنچه براي يك نفر تصادفي است ممكن است براي شخص ديگري تصادفي نباشد. ما به اين مقوله در بخش ۸ برخورد مي گشت.

حال كه مفهوم بيرون كشيدن تصادفي توپ از كيسه مشخص شد، اعداد روي توپها را فراموش كنيد و فرض كنيد R تا از اين N توپ قرمز و بقيه

مي توان به چهار سري مجله منتشر شده توسط انجمن سلطنتي آمار^۱ در سال ۱۹۹۷ مراجعه كرد. در اين انتشارات مي توان مقوله هاي گوناگون حتي راجع به حسابداری اجتماعی و قوانین حقوقی ثابت يافت. بجز چند استثنا نظير بخش الگوريتمها (كه بعداً حذف شده است) و مقاله اي درباره آموزش، تمام مقالات يا مستقيماً با عدم حتميت سروكار دارند يا درباره مشخصه هاي مطرح شده در مسائلي هستند كه به نحوی با عدم حتميت مرتبط اند. شاهدهي بر ديده گاه ما، اين است كه آمار نقش بسيار مهم تري در مقوله هاي شامل تغييرات، كه به عدم حتميت منجر مي شوند، دارد تا در مقوله هايي كه با موضوعات دقيق سروكار دارند. براي مثال، كشاورزی ارتباط نزديكي با آمار دارد، در صورتي كه فيزيك داراي چنين ارتباطي نيست. توجه كنيد كه فقط بررسي عدم حتميت است كه براي ما اهميت دارد نه خود موضوعي كه غير حتمتي است. مثلاً ما سازوكار بارش باران را مطالعه نمي كنيم بلكه علاقه منديم بدانيم باران مي بارد يا خير. اين امر، باعث شده كه علم آمار در موقعيتي عجيب قرار بگيرد كه ما آماردانان را به ديگران وابسته مي كند. پيش بيني بارش باران هم به هواشناسان و هم به آماردانان مربوط مي شود. ما تنها در مقام نظريه پرداز است كه مي توانيم بدون وابستگي به ديگران عمل كنيم. حتي در آن حالت هم چنانچه خيلي از علوم فاصله بگيريم آسيب خواهيمديد. در اين مقاله لفظ «مشتري» به فردي مانند دانشمند يا حقوقدان اطلاق مي شود كه در رشته كاري خودش با عدم حتميت مواجه است و براي حل مسائليش به آمار مراجعه مي كند.

پس موضع فلسفي اي كه در اينجا مي پذيريم اين است كه آمار اساساً مطالعه عدم حتميت است و نقش آماردان كمك كردن به مشثريهائي از حوزه هاي ديگر است كه در حوزه كار خودشان با عدم حتميت مواجه مي شوند. در عمل اين محدوديت وجود دارد كه آمار غالباً با داده ها سروكار دارد و رابطه بين عدم حتميت، يا تغيير پذيري داده ها و عدم حتميت در خود مقوله، آن چيزي است كه آماردان به بررسي اش مي پردازد. برخي از نويسندگان حتي داده ها را به داده هاي داراي فراواني كه قابل تكرار به صورت تقريباً يکسان باشند محدود مي كنند. عدم حتميت در خارج از داده ها، جاذبه زيادي براي آماردان ندارد. مطالعه عدم حتميت منحصر به آماردانان نيست. احتمال دانان در اين باره بحث مي كنند كه تصادفي بودن بخشي از يك سيستم چگونه بر ديگر بخشهاي آن تأثير مي گذارد. بنابراين الگوي [مدلي] از يك فرايند تصادفي پيش بيني هايي در مورد داده هايي كه آن فرايند توليد مي كند به دست مي دهد. گذر از فرايند به داده ها روشن است؛ مشكل زماني پيش مي آيد كه مي خواهيم از داده ها به سوي فرايند برويم. اين مقاله بيشتر به اين مرحله، كه آن را معمولاً استنباط مي نامند، و عملي كه توليد مي كند اختصاص دارد.

توجه كنيد كه عدم حتميت همه جا هست و منحصر به علوم يا حتي داده ها نيست. از جمله، بعضي از جنبه هاي الهيات از آن نشأت مي گيرد [۲]. بنابراين شناسايي آمار به عنوان مطالعه عدم حتميت نقش گسترده اي را به آماردانان محول مي كند. اگر بتوان مکتبي فلسفي پديد آورد كه همه عدم حتميت ها را در بر بگيرد، پيشرفت مهمي در شناخت ما از جهان به دست خواهد آمد. ولي در حال حاضر دستيابي به چنين هدف بزرگي حتمي به نظر نمي رسد.

۴. عدم حتمیت و احتمال

گفتیم که یک دلیل اصلی اندازه‌گیری عدم حتمیت‌ها، این است که بتوانیم آنها را با هم ترکیب کنیم. بنابراین بگذارید ببینیم که روش پیشنهاد شده چگونه ما را به این هدف می‌رساند. فرض کنید از N توپ داخل کیسه، R تا قرمز، B تا آبی و بقیه سفید هستند. در این صورت عدم حتمیت مربوط به استخراج یک توپ رنگی از کیسه عبارت است از $R/N + B/N = (R+B)/N$ یعنی برابر با مجموع عدم حتمیت‌های مربوط به قرمز و مربوط به آبی است. همین نتیجه برای هر دو پیشامد ناسازگار که عدم حتمیت‌هایشان به ترتیب R/N و B/N باشد به دست می‌آید و ما قانون جمع را برای عدم حتمیت‌های پیشامدهای ناسازگار داریم.

حال فرض کنید R توپ قرمز و $N - R$ توپ سفید در کیسه است و S تا N توپ خالدار و $N - S$ تا ساده‌اند. کیسه حالا دارای چهار نوع توپ است که تعدادی از آنها، مثلاً T تا، قرمز و خالدارند. بنابراین عدم حتمیت مربوط به استخراج توپ قرمز خالدار T/N می‌باشد که برابر است با $R/N \times T/R$ یعنی حاصلضرب عدم حتمیت یک توپ قرمز در عدم حتمیت یک توپ خالدار در میان قرمزها. باز همین نتیجه برای هر دو پیشامد که با توپهای رنگی و توپهای خالدار مقایسه شوند، به دست می‌آید و ما قانون ضرب را برای عدم حتمیت‌ها داریم.

این قوانین ضرب و جمع به همراه قانون تحدب که می‌گوید اندازه R/N همیشه داخل بازه (محدب) به طول واحد قرار می‌گیرد، قوانینی هستند که احتمال را، لاقلاً برای تعداد متناهی از پیشامدها، تعریف می‌کنند (بخش ۵ را ببینید). بنابراین نتیجه می‌گیریم که اندازه‌های عدم حتمیت را می‌توان با استفاده از حساب احتمال بیان کرد. در جهت عکس، قوانین احتمال به سادگی به قوانین مربوط به نسبتها تحویل می‌یابد. با توجه به این مطلب می‌توان دریافت که چرا روشهای فراوانی‌گرا غالباً بسیار مفید هستند: تلفیق عدم حتمیت‌ها را می‌توان بر اساس نسبتها یا فراوانیها در یک گروه، در اینجا گروه توپها، مطالعه کرد. پایه ریاضی احتمال بسیار ساده است و شاید تعجب‌آور باشد که این پایه ساده منتج به نتایج مفید و پیچیده می‌شود.

نتیجه می‌گیریم که آماردانان با عدم حتمیت سروکار دارند و عدم حتمیت را می‌توان به وسیله احتمال اندازه‌گیری کرد. ممکن است تصور شود که توضیح مجملی که در اینجا داده شد برای چنین نتیجه‌گیری مهمی کفایت نمی‌کند، لذا بیاید به دیدگاههای دیگر نیز نظری بیندازیم.

از لحاظ تاریخی، عدم حتمیت مرتبط با بازیهای شانسی و قمار بوده است. بنابراین یک راه اندازه‌گیری عدم حتمیت از طریق قمارهایی است که به عدم حتمیت وابسته‌اند. شرط‌بندی به مبلغ s بر روی وقوع پیشامدی خاص برای بردی به مبلغ w در واقع شاخصی از عدم حتمیت وقوع آن پیشامد را برحسب بخت [نسبت] w/s به ۱ به دست می‌دهد. حال به جای ترکیبی از پیشامدها خانواده‌ای از قمارها داریم. برای مشخص‌تر شدن این مفهوم، مسابقات اسبدوانی را در نظر بگیرید. در مورد دو اسب که در یک مسابقه شرکت دارند، هم شرط‌بندی جداگانه روی آنها و هم شرط‌بندی روی این پیشامد که حداقل یکی از آنها مسابقه را ببرد، قابل بررسی است؛ و اگر دو اسب در مسابقه‌های مختلف شرکت کنند می‌توان روی این پیشامد شرط‌بندی کرد که هر دو اسب مسابقه‌هایشان را ببرند. با استفاده از این گونه ترکیبات می‌توان چیزی را که در انگلستان «دفترچه هلندی»^۱ خوانده می‌شود

سفیدند و رنگ توپها تأثیری در انتخاب شما ندارد. این پیشامد غیرحتمی را در نظر بگیرید که توپی که به تصادف بیرون کشیده می‌شود قرمز باشد. عقیده بر این است که به این ترتیب استاندارد برای عدم حتمیت به دست می‌آید و اندازه آن R/N یعنی نسبت تعداد توپهای قرمز به کل توپهای داخل کیسه است. در اینجا هیچ نکته عمیقی وجود ندارد و این فقط شکلی از یک فرض است که بازیهای شانسی مبتنی بر آنها هستند. حال این را به پیشامد یا گزاره‌ای که می‌تواند رخ بدهد یا ندهد و یا درست باشد یا نباشد، تعمیم می‌دهیم. پیشنهاد می‌شود که عدم حتمیت وقوع یک پیشامد را از طریق مقایسه با استاندارد اندازه‌گیری کنید. مثلاً اگر فکر می‌کنید عدم حتمیت یک پیشامد معادل است با درآوردن تصادفی یک توپ قرمز از کیسه‌ای شامل N توپ که R تای آن قرمزند، آنگاه پیشامد مورد نظر دارای عدم قطعیتی برابر با R/N برای شماس است. انتخاب R و N به شما بستگی دارد. به ازای N داده شده، به آسانی می‌توان نشان داد که بیش از یک چنین R ای وجود ندارد. حال شاخصی از عدم حتمیت برای هر پیشامد در دست است. قبل از آنکه جلوتر برویم بگذارید فرایند اندازه‌گیری را به دقت بررسی کنیم.

در اینجا یک فرض جدی اتخاذ شده مبنی بر اینکه مفهوم عدم حتمیت را می‌توان از جنبه‌های دیگر جدا کرد. در بحث راجع به تصادفی بودن، مفید بود که یک جایزه را تحت شرایط مختلف، مثلاً بیرون آوردن توپ ۵۷ یا توپ ۱۲، مقایسه کنیم. اما لزوماً میزان جوایز نمی‌تواند در مقایسه یک پیشامد با یک استاندارد مؤثر باشد. به عنوان مثال فرض کنید پیشامدی که عدم قطعیتش مورد مطالعه است این باشد که سال بعد قرار است یک بمب اتمی در فاصله ۵۰ مایلی شما منفجر شود. در این صورت جایزه ۱۰۰۰۰ دلاری برای درآوردن یک توپ قرمز و برای گریز از تشعشعات اتمی دارای ارزش کاملاً متفاوتی خواهد بود. فرایند اندازه‌گیری که توضیح داده شد مبتنی بر این فرض است که شما می‌توانید عدم حتمیت انفجار هسته‌ای را از پیامدهای ناخوشایند آن جدا کنید. رمزی [۲۴] که کار او در بخش ۴ مورد بحث قرار خواهد گرفت، مفهومی به نام «پیشامد اخلاقاً خنثی» را معرفی کرد که در مورد آن، مقایسه با مسئله کیسه توپها مشکلات کمتری ایجاد می‌کند. بمب اتمی از نظر اخلاقی خنثی نیست.

حال به فرضی که اتخاذ نشده است توجه کنید. برای هیچ پیشامدی، از جمله انفجار بمب اتمی، فرض نشده است که شما توانایی اندازه‌گیری و تعیین R را دارید (و همچنین تعیین N را که دقت شما را در اندازه‌گیری عدم حتمیت منعکس می‌کند). بلکه فرض می‌کنیم که دوست دارید این کار را انجام دهید اگر بدانید چگونه انجام می‌شود. چیزی که در مورد هر فرایند اندازه‌گیری فرض می‌شود این است که معقول باشد و نه آنکه به سادگی انجام شود. اگر نمی‌دانیم چطور فاصله خود را تا ماه اندازه بگیریم به این معنی نیست که باور نداریم فاصله‌ای تا آن وجود دارد. دانشمندان دقت فراوان صرف اندازه‌گیری دقیق طول کرده‌اند زیرا مفهوم فاصله را با توجه به نور کریپتون باور دارند. به طور مشابه، معقول به نظر می‌رسد که سعی کنیم عدم حتمیت را اندازه بگیریم.

1. Dutch book

احتمال برمی گردیم. مقاله رمزی چندان مورد توجه قرار نگرفت تا آنکه ساویج [۲۵] همان سؤال را مطرح کرد و با فرضیهایی متفاوت با فرضهای رمزی ولی مرتبط با آنها به همان نتیجه یعنی به احتمال رسید. ساویج همچنین ارتباطی بین این موضوع و ایده های دوفیتی برقرار کرد. از آن زمان تا کنون کسان دیگری نیز در این حوزه کندوکاو کرده و نتایج مشابهی به دست آورده اند. متن محبوب من در این زمینه، در میان منتهایی که از دقت کامل برخوردارند، فصل ششم [۱۱] است. اثر برناردو و اسمیت [۷] هم شرح عالی جدیدی از این موضوع است.

۵. احتمال

نتیجه بحث بالا این است که شاخصهای عدم حتمیت باید از قوانین حساب احتمال تبعیت کند. قواعد دیگر نظیر قوانین منطق فازی یا نظریه امکان^۱ که به مینیم و ماکسیم وابسته اند نه به مجموع و حاصلضرب، و همچنین چند قانون دیگر که آماردانان از آنها استفاده می کنند، کنار گذاشته می شوند. بخش ۶ را ببینید. تمام این استنتاجات، چه بر پایه توپهای داخل کیسه ها، چه بر پایه قمار و چه بر پایه قوانین امتیازبندی و یا تصمیم گیری، مبتنی بر فرضهاست. چون این فرضیات منجر به نتیجه گیریهای بسیار مهم می شوند، لازم است به دقت بررسی شوند. متأسفانه اکثریت بزرگ آماردانان این کار را نمی کنند. بعضی از آنها نقش محوری احتمال را نفی می کنند بدون آنکه دلالی برای آن ارائه کنند. هدف این مقاله این نیست که بررسی عمیقی راجع به این مطلب انجام دهد ولی اجازه دهد به یک فرض به خاطر نتایج بعدی اش نگاهی بیندازیم. همانند همه فرضهای دیگر تصور می شود که این فرض هم خود به خود بدیهی است و خلاف آن را احتماقه می دانیم. این فرض مبتنی است بر یک مفهوم اولیه یعنی اینکه یک پیشامد از پیشامد دیگر محتمل تر است، کلمه «محتمل» در اینجا به مفهوم فنی آن مورد نظر نیست و فقط به معنای معمولی آن در زبان طبیعی به کار رفته است. می نویسیم « $A > B$ » هرگاه « A محتمل تر از B » باشد. فرض مورد نظر این است که اگر A_1 ، A_2 با هم و B_1 ، B_2 با هم ناسازگار باشند آنگاه $A_i > B_i$ ، $i = 1, 2$ ، نتیجه می دهد $A_1 \cup A_2 > B_1 \cup B_2$ (در اینجا $A_1 \cup A_2$ به معنی وقوع پیشامدی است که رخ می دهد اگر و تنها اگر A_1 یا A_2 رخ بدهد). اگر بدانیم فرد بعدی که از در وارد می شود محتمل تر است زنی غیر سفیدپوست باشد (A_1) تا مردی غیر سفیدپوست (B_1) یا زنی سفیدپوست باشد (A_2) تا مردی سفیدپوست (B_2)، آنگاه ممکن است احساس ناخرسندی کنیم که زن بودن واردشونده ($A_1 \cup A_2$) کمتر از مرد بودن او ($B_1 \cup B_2$) محتمل باشد. مباحثی که در بالا به آنها اشاره شد با فرضها یا اصول موضوعی از این نوع شروع می شوند و نکته مهم این است که همه آنها به این واقعیت می انجامند که احتمال تنها وسیله قانع کننده برای بیان عدم حتمیت است.

جمله اخیر کاملاً دقیق نیست. برخی از نویسندگان با تعمق در اصول موضوع، مخالفتهایی با آن ابراز کرده اند. والی در نقد زیبایی [۲۷] سعی کرده با یک جفت عدد به نام احتمالات بالایی و پایینی به جای یک عدد به نام احتمال، سیستمی را بنا کند. نتیجه کار، سیستمی پیچیده تر است، که به نظر من این پیچیدگی ضرورتی ندارد. با این حال می پذیرم که گاهی اوقات

به کار گرفت. می گویند دنباله ای از شرط بندیها با بختها [نسبتها] مشخص تشکیل یک دفترچه هلندی می دهد هرگاه بتوان مجموعه ای از سرمایه گذاریها را چنان انجام داد که در کل، سرمایه گذاشته شده افزایش یابد. مسئول شرط بندی هرگز بختهایی را که بر اساس آنها می توان دفترچه هلندی ساخت فاش نمی کند. به آسانی می شود نشان داد [۱۶] که اجتناب از دفترچه هلندی معادل است با احتمالات مستخرج از بختها، که در قوانین تحدب، جمع، و ضرب احتمالات که در بالا گفته شد تبعیت می کنند. برای بخت p ، احتمال متناظر برابر است با $p = (1 + \circ)^{-1}$. به طور خلاصه، استفاده از بخت به همراه اجتناب از دفترچه هلندی ما را مانند قبل به احتمال رهنمون می سازد. دوفیتی روش دیگری را نیز معرفی کرده است. فرض کنید از شما خواسته شده است که میزان عددی عدم حتمیت یک پیشامد را از نظر خودتان مشخص کنید و به شما گفته شده که امتیازی به اندازه $S(x, E)$ می گیرید هرگاه این میزان را x بگیرید؛ $E = 1$ و $E = 0$ به ترتیب متناظر با وقوع و عدم وقوع پیشامدند. برای پیشامدهای متفاوت امتیازها جمع می شوند. جریمه ای که دوفیتی به کار گرفت عبارت است از $S(x, E) = (x - E)^2$ ولی بجز چند تابع خاص که منجر به نتایج مضحکی می شوند هر تابع دیگری را نیز می توان تابع جریمه در نظر گرفت. بنابراین اگر شما x را به عنوان تابعی از احتمال به کار بگیرید یعنی تابعی از چیزی که از سه قانون احتمال تبعیت می کند، مطمئن خواهید بود که از شخص دیگری که روش دیگری به کار می گیرد بهتر عمل می کنید به این معنی که جریمه کمتری می گیرد. روش دوفیتی روش جذابی است زیرا عملاً می توان آن را به کار برد و برای تعلیم کسانی که وضع هوا را پیش بینی می کنند مورد استفاده قرار گرفته است. این روش معمولاً محکی تجربی برای کیفیت برآوردهای احتمالاتی شما، و در نتیجه آزمونی از توانایی شما به عنوان آماردان، به دست می دهد (بخش ۹ را ببینید). نکته مهم این است که این روش هم به احتمال ختم می شود. بعداً گروه سومی از روشها را مطالعه می کنیم که در آنها هم استفاده از احتمال برای توصیف عدم حتمیت الزامی است.

انجمن سلطنتی آمار انگلستان و بسیاری دیگر از گروههای آماری در آغاز به منظور جمع آوری و انتشار داده ها تأسیس شدند. این هدف با بحث ما در مورد عدم حتمیت توافق دارد زیرا یکی از انگیزه های جمع آوری داده ها، کاهش عدم حتمیت است. امروز هم بخش مهمی از فعالیتهای آماری به این منظور انجام می شود. اغلب دولتها دارای اداراتی هستند که وظیفه اصلی آنها تهیه و ارائه آمار است. خیلی طول نکشید که آماردانان به این فکر افتادند که چگونه از داده ها می توان به بهترین وجه استفاده کرد، و استنباط آماری نوین پا به عرصه وجود نهاد. آنها شاهد بودند که استنباط آماری به عنوان مبنای تصمیمها و عملها به کار می رود و همراه دیگران شروع به بررسی مسأله تصمیم گیری کردند. حال توضیح می دهیم که چگونه این موضوع هم به احتمال منتهی می شود.

به نظر می آید کسی که اولین قدم را به سمت تصمیم گیری برداشت رمزی بود. وی سؤال ساده ای را مطرح کرد: «وقتی با عدم حتمیت مواجه ایم چگونه تصمیم بگیریم؟» و چند فرض منطقی اتخاذ کرد و از آنها قضایی را استنتاج نمود. قضیه ای که الان منظور ماست آن است که می گوید شاخصهای عدم حتمیت باید از قوانین حساب احتمال تبعیت کنند. بنابراین باز هم به

اختلاط از نوع دسته‌ای از قوانین است که به نام «چیزهای مطمئن» معروف‌اند. به تسامح می‌توان گفت، چنانچه هر اتفاقی بیفتد (هر چه باشد B_n) باور شما در مورد یک پدیده برابر k باشد آنگاه این باور به‌طور غیرشرطی k است. برخی از آماردانان به نظر می‌رسد که از قابلیت اختلاط فقط زمانی که به نفع آنهاست استفاده می‌کنند: ما جنبه کاربردی آن را در بخش ۸ بررسی می‌کنیم. توجه کنید که قوانین احتمالی که اینجا به آنها اشاره کردیم مانند اصول موضوعی که در کتابهای درسی احتمال آورده می‌شود نیستند. اینها، بجز قانون ۴، نتایجی از فرضهای ساده‌تر دیگری هستند.

۶. معنی دار بودن و اطمینان

عکس‌العمل بسیاری از آماردانان نسبت به این حکم که باید از احتمال استفاده کنند این است که آنها این کار را می‌کنند و دیدگاهی که در اینجا توضیح داده شده فقط مهر تأییدی است بر کاری که از قبل می‌کرده‌اند. مجلات پر است از احتمالات: توزیعهای نرمال و دوجمله‌ای به‌وفور دیده می‌شود؛ خانواده‌ی نمایی همه جا هست. حتی ممکن است ادعا شود از هیچ شاخص دیگری برای اندازه‌گیری عدم حتمیت استفاده نمی‌شود؛ عده کمی از آماردانان ممکن است به منطق فازی اشاره کنند. ولی این عکس‌العمل درست نیست؛ آماردانان شاخصهایی از عدم حتمیت را به‌کار می‌برند که با استفاده از قوانین حساب احتمال با یکدیگر ترکیب نمی‌شوند.

گیریم فرض H این باشد که یک دارو بی‌تأثیر است و یا یک عامل اجتماعی تأثیری بر میزان جنایت ندارد. پزشک یا جامعه‌شناس درباره H نظر حتمی ندارد و داده‌هایی جمع‌آوری می‌شوند به این امید که این عدم حتمیت برطرف شود و یا دست‌کم کاهش یابد. برای راهنمایی در مورد عدم حتمیت از آماردانی کمک خواسته می‌شود و این آماردان ممکن است پیشنهاد کنند که مشتری شاخصی از عدم حتمیت مثل سطح زیر منحنی یا سطح معنی‌دار بودن را در مورد H به‌عنوان فرض صفر به‌کار گیرد. یعنی فرض کند H درست است و احتمال داده‌های مشاهده شده را تحت این فرض محاسبه کند. در واقع این شاخص اعتقادی است که می‌توان به H داشت و هر چه این احتمال کوچکتر باشد اعتقاد به H کمتر است.

این روش در واقع مغایر با بحث بالاست زیرا در بالا تأکید بر این بود که عدم حتمیت در مورد H باید با احتمال H اندازه‌گیری شود. سطح معنی‌دار بودن چنین احتمالی را به‌دست نمی‌دهد. این تفاوت را می‌توان به‌صورت زیر مطرح کرد.

سطح معنی‌دار بودن — احتمال مشاهده جنبه خاصی از داده‌ها، با فرض اینکه H درست است؛

احتمال — احتمالی که شما برای H قائل‌اید، با فرض مشاهده داده‌ها. «مغلطه دادستان» که در محافل حقوقی معروف است، در واقع در اثر خلط کردن $p(A|B)$ با $p(B|A)$ ، دو کمیتی که به‌ندرت با هم برابرند، رخ می‌دهد. تفاوت بین سطح معنی‌دار بودن و احتمال تقریباً مانند مغلطه دادستان است: «تقریباً» به این دلیل که گرچه در ساختار مثال دادستان، B را می‌توان برابر با H در نظر گرفت ولی داده‌ها به نحو دیگری بررسی می‌شوند. A به‌عنوان داده‌ها در احتمال مورد استفاده قرار می‌گیرد. طرفداران سطح معنی‌دار بودن خیلی زود دریافته‌اند که کافی نیست فقط از داده‌ها استفاده کنند بلکه لازم

روش احتمالاتی کاملاً دقیق نیست و این عدم دقت می‌تواند با به‌کارگیری احتمالاتی بالایی و پایینی ترمیم شود. هدف این است که این جفت عدد، گزاره‌های احتمالاتی را دقیق‌تر کنند، ولی احتمال به تنهایی شامل شاخصی از دقت خودش هست. عقیده من این است که ساده بر پیچیده ترجیح دارد، به شرط آنکه کارساز باشد و من اساساً با تیغ اکام^۱ موافقم.

با این نتیجه‌گیری که عدم حتمیت تنها توسط احتمال به طرز قانع‌کننده‌ای بیان می‌شود، سه قانون یا سه اصل موضوع حساب احتمال را به‌طور رسمی بیان می‌کنیم. احتمال به دو عنصر وابسته است: پیشامد غیرحتمی و شرایطی که تحت آنها به آن پیشامد توجه می‌کنید. مثلاً در بیرون کشیدن گوی از کیسه، احتمالی که برای قرمز بودن گوی قائل می‌شوید به این شرط وابسته است که گوی به تصادف از کیسه بیرون کشیده شود. برای احتمال وقوع A وقتی که می‌دانیم یا فرض می‌کنیم که B رخ داده است می‌نویسیم $p(A|B)$ و می‌خوانیم احتمال A به شرط (یا به فرض) B ، این قوانین عبارت‌اند از: (الف) قانون ۱ (تحدب): برای همه A و B ها، $0 \leq p(A|B) \leq 1$ و

$$p(A|A) = 1$$

(ب) قانون ۲ (جمع): اگر A و B ناسازگار باشند و C مفروض باشد،

$$p(A \cup B|C) = p(A|C) + p(B|C)$$

(ج) قانون ۳ (ضرب): برای همه A و B و C ها

$$p(AB|C) = p(A|BC)p(B|C)$$

(اینجا AB به معنی پیشامدی است که رخ می‌دهد اگر A و B هر دو رخ دهند). قانون تحدب گاهی اوقات توسط رابطه زیر تقویت می‌شود: $p(A|B) = 1$ تنها اگر A نتیجه منطقی B باشد. قانون جمع به قانون کرامول معروف است.

قانون جمع ممکن است فقط یک نوع ظریف‌کاری ریاضی به نظر برسد ولی در حقیقت نتایج عملی مهمی دارد که در بخش ۸ نشان داده خواهد شد. درباره این قانون نکته‌ای شایان ذکر است. با سه قانونی که در بالا ذکر شد، قانون جمع دو پیشامد را می‌توان به هر تعداد متناهی از پیشامدها تعمیم داد. هیچ‌یک از رویکردهای بسیاری که قبلاً درباره آنها بحث شد به برقراری این قانون برای بینهایت پیشامد منجر نمی‌شود. ولی معمول است که فرض شود قانون جمع برای تعداد نامتناهی از پیشامدها برقرار است زیرا خلاف آن منتهی به نتایج ناخوشایند می‌شود. برای صراحت بخشیدن به این فرض می‌توان قانون جمع را اصلاح کرد و یا قانون چهارمی اضافه کرد که همراه سه قانون بالا تعمیم قانون جمع را به بینهایت پیشامد دوبه‌دو ناسازگار به‌دست دهد. من شخصاً ترجیح می‌دهم قانون چهارمی به شکل زیر اضافه کنیم.

(د) قانون ۴ (قابلیت اختلاط): اگر $\{B_n\}$ افزاری، متناهی یا نامتناهی، از C باشد و به‌ازای همه مقادیر m ، $p(A|B_n C) = k$ و $p(A|C) = k$ (به آسانی می‌توان نشان داد که قانون ۴ وقتی که افزاز متناهی است از قوانین ۱ تا ۳ به‌دست می‌آید. این تعریف از آن دوفینیتی است). قابلیت

1. Occam's razor

اشاره‌ای است به اصلی منسوب به فیلسوف انگلیسی ویلیام اکام (در اصل، آکم) مبنی بر اینکه در توضیح یک پدیده باید حداقل مفروضات آورده شود.

چگونه می‌توان چند مجموعه داده درباره یک فرض را که هر یک سطح معنی‌دار بودن برای خود دارد با هم ترکیب کرد؟ نتایج دو آزمون t استودنت برای میانگینهای μ_1 و μ_2 با آزمون T دو متغیره برای (μ_1, μ_2) سازگار نیست [۱۸].

پس این نتیجه‌گیری که احتمال تنها شاخص عدم حتمیت است صرفاً یک تعریف و تمجید خشک و خالی نیست بلکه به بسیاری از فعالیتهای پایه‌ای آماری مربوط می‌شود. ساویج ایده‌های خودش را در تلاشی برای توجیه عمل آماری ارائه داد. او با تشکیک در جنبه‌هایی از آن نتایجی به دست آورد که برای خودش شگفت‌انگیز بود. حال از اختلاف نظرها بگذریم و به سمت ایده‌های سازنده‌ای برویم که از شناسایی نقش اساسی احتمال در آمار نتیجه شده‌اند.

۷. استنباط

فرمولبندی‌ای که در سراسر این قرن خدمات شایسته‌ای به آمار کرده بر اساس داده‌هایی است که به‌ازای هر مقدار از پارامتر مورد نظر، یک توزیع احتمال دارند. این فرمولبندی با این نظر که عدم حتمیت داده‌ها باید به زبان احتمالاتی بیان شود، سازگار است. مطلوب آن است که از داده‌ها چیزی درباره پارامتر استنباط کنیم. معمولاً همه جنبه‌های پارامتر، مورد نظر نیست، بنابراین می‌نویسیم (θ, a) به این معنی که می‌خواهیم چیزی در مورد θ استنباط کنیم و a در اصطلاح فنی پارازیت یا پارامتر مزاحم است. اگر داده‌ها را با x نمایش دهیم آنگاه $p(x|\theta, a)$ احتمال x است وقتی که θ و a داده شده باشند. یک مثال ساده می‌تواند توزیعی نرمال با میانگین θ و واریانس a باشد ولی این فرمولبندی بسیاری از حالت‌های پیچیده را نیز در بر دارد.

با این روش، عدم حتمیت درون داده‌ها مطابق میل هر شخص کنترل می‌شود. پارامتر نیز غیرحتمی است. در حقیقت این عدم حتمیت است که دغدغه اصلی آماردان است و طبق این روش، آن نیز باید با یک احتمال $p(\theta, a)$ توصیف شود. با این کار ما از روش معمول عدول کرده‌ایم. غالباً می‌گویند که پارامترها بنا به فرضی کمیتهایی تصادفی هستند. ولی چنین نیست. این اصول موضوع هستند که فرض می‌شوند و خاصیت تصادفی بودن از آنها نتیجه می‌شود. با در دست بودن هر دو احتمال و استفاده از حساب احتمال می‌توان در میزان عدم حتمیت با توجه به داده‌ها تجدیدنظر کرد:

$$p(\theta, a|x) \propto p(x|\theta, a)p(\theta, a) \quad (۱)$$

که ثابت تناسب فقط به x بستگی دارد و نه به پارامترها. از آنجا که a مورد علاقه ما نیست می‌توانیم با استفاده از حساب احتمال آن را حذف کنیم:

$$p(\theta|x) = \int p(\theta, a|x) da. \quad (۲)$$

رابطه (۱) قانون ضرب و رابطه (۲) قانون جمع است. مسأله استنباط به کمک این دو رابطه حل می‌شود و یا به عبارت بهتر چارچوبی برای حل آن فراهم می‌آید. رابطه (۱) قضیه بیز است که نام آن بنا به تصادفی تاریخی به کل این رویکرد داده شده است و به آن روش بیزی می‌گویند. این نام‌گذاری ناسف‌انگیز با چند نام‌گذاری بدتر همراه است. غالباً $p(\theta)$ را توزیع پیشین و $p(\theta|x)$ را توزیع پسین می‌نامند. اینها نامناسب‌اند زیرا پیشین و پسین اصطلاحاتی

است تعدادی داده «فرین‌تر» را به شکل دم یک توزیع وارد کنند. همان‌طور که جفریز [۱۹] اشاره می‌کند، این سطح شامل داده‌هایی است که ممکن بود مشاهده شوند ولی نشدند.

حال از آزمون فرض به برآورد نقطه‌ای می‌رویم. پارامتر θ ممکن است اثر یک دارو یا تأثیر یک عامل اجتماعی بر میزان جنایت باشد. باز θ غیرحتمی است، لذا ممکن است داده‌هایی جمع‌آوری شده از آماردانی مشورت خواسته شود (امید این است که ترتیب غیر از این باشد). آماردان معمولاً یک بازه اطمینان را پیشنهاد می‌کند. در این امر که مبتنی بر عدم حتمیت اندازه‌گیری شده است، از چگالی احتمال θ و شاید بازه‌ای از آن چگالی استفاده می‌شود. دوباره ما با تضادی شبیه به آنچه در مغلفه دادستان بود مواجه‌ایم:

اطمینان — احتمال آنکه بازه شامل θ باشد؛

احتمال — احتمال آنکه θ در این بازه قرار گیرد.

عبارت اول گزاره‌ای احتمالاتی در مورد یک بازه، مشروط به مشاهده θ است؛ حال آنکه عبارت دوم گزاره‌ای در مورد θ مشروط به داده‌هاست. غالباً این دو را با هم اشتباه می‌گیرند. مهمتر از این اشتباه این است که نه سطوح معنی‌دار بودن و نه بازه‌های اطمینان هیچ‌یک بر اساس قوانین حساب احتمال ترکیب نمی‌شوند.

آیا این اشتباه حائز اهمیت است؟ از لحاظ نظری مسلماً چنین است زیرا استفاده از شاخصی که از قوانین حساب احتمال تبعیت نکند نهایتاً ناقض بعضی از فرضهای پایه‌ای خواهد بود، فرضهایی که بدیهی در نظر گرفته شده‌اند و نقض آنها ایجاد اشکال می‌کند. ولی از لحاظ عملی این موضوع آنقدر واضح نیست و لازم است درباره پیامدهای عملی آنها توضیح داده شود. آماردانان تمایل دارند هر مسأله را جداگانه مطالعه کنند و نیازی نمی‌بینند که نتایج را با هم ترکیب کنند؛ این ترکیب نتایج است که مشکل به وجود می‌آورد، و این را در مثال رنگ-جنسیت در بخش ۵ دیدیم. به عنوان مثال به ندرت امکان دارد که بتوان دفترچه هلندی را با استفاده از گزاره‌های سطوح معنی‌داری ساخت. معلوم است که برخی از برآوردهای رایج غیر قابل قبول‌اند. روشن‌ترین مثال از یک اشکال مهم، ناشی از رابطه بین سطح معنی‌دار بودن و اندازه (n) نمونه‌ای است که این سطح بر آن بنا شده است. تعبیر «معنی‌دار در سطح ۵٪» بستگی به n دارد حال آنکه احتمال ۵٪ همیشه ثابت است. آماردانان به ارتباط بین گزاره‌ای که بیان می‌کنند و اندازه نمونه‌ای که گزاره مبتنی بر آن است توجه چندانی نکرده‌اند. دلایلی نظری وجود دارد که فکر کنیم به دست آوردن سطح معنی‌دار ۵٪ خیلی ساده است [۳]. و اگر این مطلب درست باشد بسیاری از آزمایشها امیدی واهی به تأثیری سودمند، که در واقع وجود ندارد، ایجاد می‌کنند.

گزاره‌های منفرد آماری معمولاً اشکالی ندارند. مشکل زمانی به وجود می‌آید که بخواهیم آنها را با هم تلفیق کنیم. به عنوان مثال، بازه‌های اطمینان برای یک پارامتر غالباً قابل قبول هستند ولی برای چندین پارامتر نه. حتی میانگین یک نمونه از توزیع نرمال نیز در ابعاد بالا مناسب به نظر نمی‌رسد. در آزمایشی با تیمارهای گوناگون، آزمونهای منفرد خوب‌اند ولی مقایسه‌های چندگانه دچار مشکل می‌شوند. اما حقیقت علمی با تلفیق نتایج آزمایشهای گوناگون محقق می‌شود: هنوز مبحث فراتحلیل^۱ حوزه‌ای دشوار از آمار است.

1. meta-analysis

متأسفانه این کار معمولاً منجر به تناقضی با خاصیت اختلاط می‌شود. برای مثال، فرض کنید می‌دانید که θ فقط مقادیر صحیح مثبت را می‌تواند اختیار کند و اطلاع دیگری درباره مقادیر آن ندارید. آنگاه از مفهوم معمولی چهل چنین برمی‌آید که همه مقادیر همسان هستند: $p(\theta = i|I) = c$ به‌ازای همه i ها. قانون جمع برای n پیشامد دوه‌دو ناسازگار $\{\theta = i\}$ با ضابطه $nc > 1$ نتیجه می‌دهد $c = 0$ زیرا بنابر اصل تحذب، هیچ احتمالی نباید از یک بیشتر شود. حال اعداد صحیح مثبت را به مجموعه‌هایی سه عضوی افراز کنید که در هر یک از آنها دو عدد فرد و یک عدد زوج وجود داشته باشد مثلاً به‌صورت $A_n = (4n-3, 4n-1, 2n)$. اگر E پیشامد آن باشد که θ زوج است آنگاه $p(E|A_n) = 1/3$ و با توجه به خاصیت اختلاط، $p(E) = 1/3$. افراز دیگری را مرکب از مجموعه‌هایی در نظر بگیرید که هر یک شامل دو عدد زوج و یک عدد فرد است مثلاً به‌صورت $B_n = (4n-2, 4n, 2n-1)$. در این صورت $p(E|B_n) = 2/3$ و بنابراین $p(E) = 2/3$ که با مقدار قبلی متفاوت است. در واقع با انتخاب مناسب افراز، $p(E)$ می‌تواند هر مقداری را در بازه یکه اختیار کند. چنین توزیعی را ناسره می‌نامند.

متأسفانه بیشتر تلاشها در جهت تولید $p(\theta|I)$ با استفاده از دیدگاه «لازم» به یک ناسرگی منجر می‌شود که نه تنها با خاصیت اختلاط سازگار نیست بلکه پیامدهای غیر قابل قبول دیگری هم دارد. در این زمینه، مثلاً [۱۰] را ببینید. دیدگاه لازم را اولین بار جفریز [۱۹] مورد بررسی قرار داد. برناردو [۶] و دیگران پیشرفتهای واقعاً چشمگیری را موجب شده‌اند ولی مسأله هنوز حل نشده است. در این مقاله این نظر اتخاذ می‌شود که احتمال گزاره‌ای است که یک شخص با آگاهی مشخص در مورد یک کمیت غیر حتمی می‌سازد. از این شخص با عنوان «شما» یاد می‌کنیم. $p(A|B)$ میزان باور شما در مورد A است وقتی B را می‌دانید. با این دید، درست نیست از احتمال به‌صورتی یاد کنیم که گویی منحصر به فرد است؛ فقط «احتمال از نظر شما» معنی دارد. بیان شرطها به همان اندازه اهمیت دارد که پیشامد غیر حتمی مورد نظر.

در مثال بالا که θ فقط مقادیر صحیح مثبت را اختیار می‌کند به محض آنکه مفهوم چیزی را که این حرف یونانی به آن اشاره دارد فهمیدید، دیگر جاهل نیستید. در آن صورت $p(\theta = i) = a_i$ با شرط $\sum a_i = 1$. برخی از تحقیقات آماری به جهت سردرگمی بین حرف یونانی و واقعیتی که حرف نماینده آن است مخدوش می‌شوند. من همه را به مبارزه می‌طلبم که کمیتی واقعی ارائه کنند که واقعاً نسبت به آن جاهل باشند. نکته دیگر این است که رایانه نمی‌تواند مثالی از θ تولید کند. از آنجا که $p(\theta \leq n|I) = nc = 0$ و $p(\theta > n) = 1$ بنابراین θ باید قطعاً بزرگتر از هر عددی باشد که شما به فکرتان می‌رسد و یا رایانه می‌تواند محاسبه کند.

توجه به مطلب دیگری نیز ضروری است. معمولاً هنگامی که بخشی از پیشامد شرطی‌کننده مجهول است اما فرض می‌شود که رخ داده است، باز از همین مفهوم و نمادها استفاده می‌شود. به‌عنوان مثال، همه آماردانان نماد $p(x|\theta, a, K)$ یا به عبارت دقیق‌تر $p(x|\theta, a, K)$ را به‌کار می‌برند. توجه کنید که آنها در اینجا مقدار پارامتر را نمی‌دانند. چیزی که با این نماد بیان می‌شود عدم حتمیت درباره x است اگر پارامترها مقادیر θ و a را اخذ کنند.

نسبی هستند که به داده‌ها مربوط می‌شوند. پسین امروز، پیشین فرداست. ولی این دو اصطلاح چنان رواج یافته‌اند که اجتناب کامل از آنها ناممکن است.

- حال مواضعی را که به آنها دست یافته‌ایم خلاصه‌وار ذکر می‌کنیم.
- (الف) آمار مطالعه عدم حتمیت است.
 - (ب) عدم حتمیت باید با احتمال اندازه‌گیری شود.
 - (ج) عدم حتمیت مربوط به داده‌ها به این طریق، مشروط به پارامترها، اندازه‌گیری می‌شود.
 - (د) عدم حتمیت پارامتر نیز با احتمال اندازه‌گیری می‌شود.
 - (ه) استنباط در قالب حساب احتمال، عمدتاً به کمک روابط (۱) و (۲)، انجام می‌شود.

درباره موارد (الف) و (ب) قبلاً بحث کردیم و (ج) را نیز عموماً می‌پذیرند. مورد (ه) از موارد (الف) تا (د) نتیجه می‌شود. اعتراض عمده به روش بیزی، به مورد (د) مربوط می‌شود. بنابراین اجازه دهید به آن بپردازیم.

۸. ذهنیت

در سطح پایه‌ای که بحث را از آنجا شروع کردیم، روشن است که عدم حتمیت‌ها در نظر اشخاص مختلف معمولاً متفاوت است. ممکن است درباره بازیهایی شانس‌ی که خوب اجرا شوند توافقی وجود داشته باشد به‌طوری که بتوان آنها را به‌عنوان استاندارد به‌کار برد. ولی در بسیاری از موضوعات، حتی در علوم، ممکن است توافق وجود نداشته باشد. در زمان نوشتن این سطور دانشمندان درباره برخی از مسائل مربوط به GMF (مواد غذایی‌ای که به‌طور ژنتیکی دستکاری شده‌اند) اختلاف نظر دارند. بنابراین لازم به نظر می‌رسد که ذهنیت را به‌نحوی مناسب در نمادهایمان منعکس کنیم. راه مرجع این است که مفهوم آگاهی شخص را، در زمانی که قضاوت احتمالاتی می‌کند، وارد کار کنیم. فرض کنید K نمایانگر آن باشد، لذا یک نماد بهتر از $p(\theta)$ ، نماد $p(\theta|K)$ ، احتمال θ به شرط K ، است که بر اساس داده‌ها به $p(\theta|x, K)$ تغییر می‌یابد. این نماد ارزشمند است زیرا بر شرطی بودن احتمال تأکید می‌کند (بخش ۵ را ببینید). این نماد به دو مؤلفه وابسته است: چیزی که عدم حتمیت آن توصیف می‌شود و آگاهی‌ای که این عدم حتمیت بر آن بنا شده است. حذف مؤلفه شرطی‌کننده غالباً به سردرگمی می‌انجامد. تمایز بین پیشین و پسین با تکیه بر شرایط متفاوتی که احتمال پارامتر مورد نظر در آن شرایط تعیین می‌شود، بهتر توصیف می‌شود.

این نظر مطرح شده است که اگر دو نفر آگاهی یکسانی راجع به یک موضوع داشته باشند، عدم حتمیت و در نتیجه احتمال از نظر آنها باید یکسان باشد. این را دیدگاه لازم می‌نامند. این دیدگاه دو اشکال دارد. اولاً توضیح اینکه منظور از داشتن آگاهی یکسان چیست و همچنین تشخیص آن در عمل، دشوار است. نکته دوم اینکه اگر احتمال لزوماً به‌دست می‌آید پس باید بتوانیم آن را بدون ارجاع به شخص محاسبه کنیم. تاکنون هیچ‌کس نتوانسته این محاسبه را به‌صورت کاملاً قانع‌کننده‌ای انجام دهد. یک راه، استفاده از مفهوم چهل است. اگر سطح آگاهی I مشخص شود، آنگاه فقدان آگاهی در مورد θ مشخص می‌شود و $p(\theta|I)$ تعریف می‌گردد و $p(\theta|K)$ را به‌ازای هر K توسط فرمول بیزی رابطه (۱) با تبدیل I به K می‌توان محاسبه کرد.

هر چند معنای آن، چنانکه توضیح داده خواهد شد، با معنای معمول این اصطلاح در آمار فرق دارد و با مدلهای رایج در علوم هم متفاوت است. کار آماردان این است که برای دنیای غیرحتمی تحت مطالعه الگو بسازد. هنگامی که این کار انجام شد، عدم حتمیت جنبه‌های خاص مورد نظر را می‌توان به وسیله حساب احتمال و بر اساس اطلاعات موجود محاسبه کرد. بنابراین مطالعه ما دو جنبه دارد: ساختن الگو و تحلیل آن. دومی اساساً خودبه‌خود انجام می‌شود؛ علی‌الاصول می‌توان آن را توسط ماشین انجام داد. ولی اولی نیازمند به تماس نزدیک با واقعیت است. با دستکاری در جمله‌ای از دوفینتی و قدری غلو می‌توان گفت: «وقتی که الگو می‌سازی، فکر کن، و وقتی الگو را داری فکر نکن و آن را به رایانه بسپار». بار دیگر این مطلب را تکرار می‌کنیم که شخصی که احتمالات مورد نظر او بررسی می‌شود، یا به زبان این مقاله، «شما»، آماردان نیست بلکه یک مشتری است؛ غالباً محقق است که خواستار اظهارنظر آماری شده است. وظیفه آماردان ترجمه عدم حتمیت‌های محقق به زبان احتمال و محاسبه آنها با استفاده از اعدادی است که به دست می‌آورد. هر الگو صرفاً انعکاس واقعیت در فکر شماست و مانند احتمال، نه شما را و نه دنیای شما را توصیف نمی‌کند، بلکه فقط رابطه بین شما و آن دنیا را وصف می‌کند. لذا نام بردن از آن به عنوان الگوی واقعی درست نیست. زمانی می‌توان این کلمه را به کار برد که غالب مردم با آن الگو موافق باشند. مثلاً توزیع نرمال دو تغییری را ممکن است بتوان الگوی واقعی قد پدران و پسران نامید.

در یک سناریوی معمولی چه نوع عدم حتمیت‌هایی وجود دارد؟ مسأله اساسی استنباط و استقرا، پیش‌بینی داده‌های آتی از روی داده‌های گذشته است. رصدهای مکرر حرکت اجرام آسمانی محاسبه مواضع آینده آنها را امکان‌پذیر می‌کند. مطالعات کلینیکی روی یک دارو به پزشک امکان می‌دهد که وضع آینده بیماری را که آن دارو برایش تجویز شده، پیش‌بینی کند. گاهی اوقات داده‌های غیرحتمی در گذشته‌اند، نه در آینده. هرگاه گزارشی از یک واقعه تاریخی به جا نمانده باشد، مورخ از تمام شواهدی که در دست دارد استفاده می‌کند تا شاید مشخص کند چه اتفاقی به وقوع پیوسته است. دادگاه جنایی می‌کوشد بر اساس شواهدی که پس از جنایت به دست آمده، معلوم کند چه اتفاقی افتاده است. ولی ما در اینجا فرض می‌کنیم که داده‌های گذشته، x ، برای استنباط راجع به داده‌های آینده، y ، به کار می‌روند (همان‌طور که در الفبای انگلیسی x قبل از y است). با این ترتیب، مسأله ما محاسبه $p(y|x, K)$ است. برای سادگی، اطلاعات قبلی را که در طول کار ثابت است در نمادگذاری منظور نمی‌کنیم و می‌نویسیم $p(y|x)$.

یک راه ممکن این است که سعی کنیم $p(y|x)$ را مستقیماً محاسبه کنیم. این کار معمولاً دشوار است، هر چند می‌توان آن را نمونه‌ای از عمل در نظام استاد و شاگردی دانست. اینجا شاگرد باید پای درس معلم بنشیند و اطلاعات (داده‌های x) را جذب کند. پس از سالها کسب تجربه در این زمینه، شاگرد ما می‌تواند خودش به تنهایی استنباط کند که چه ممکن است رخ دهد. مشاهده مکرر استفاده از فلان مواد برای ساختن تایلر این توانایی را به او می‌دهد که از آن مواد در ساختن تایلر برای خودش استفاده کند. ولی راه بهتری وجود دارد و آن مطالعه ارتباط‌های بین x و y و سازوکار عمل‌کننده است. قوانین نیوتن امکان محاسبه جذر و مد را می‌دهند. علم مواد به طراحی و ساخت تایلر

(و آگاهی ما راجع به آنها به اندازه K باشد). لازم نیست که در این چارچوب بین فرض و واقعیت فرق گذاشته شود و نماد $p(A|B)$ کفایت می‌کند. موضع فلسفی من این است که عدم حتمیت شخصی شما از طریق احتمال شخصی شما در مورد یک کمیت غیرحتمی بر اساس سطح آگاهی شما، واقعی یا فرضی، بیان می‌شود. این را نگرش شخصی یا ذهنی به احتمال می‌نامیم.

بسیاری از افراد، مخصوصاً وقتی با مفاهیم علمی سروکار دارند، فکر می‌کنند گزاره‌هایشان عینی‌اند و از طریق یک احتمال یکتا بیا می‌شوند، و از مزاحمت ذهنیات هراس دارند. این هراس را می‌توان با توجه کردن به واقعیت و نحوه بازتاب واقعیت در حساب احتمال کاهش داد. ما در چارچوب علوم بحث می‌کنیم ولی این رویکرد به طور کلی درست است. علم حقوق مثال دیگری به دست می‌دهد. فرض کنید θ کمیت علمی مورد نظر باشد. (در حقوق جنایی می‌توان θ را برحسب اینکه متهم مرتکب جنایت شده یا نه، برابر ۱ یا ۰ گرفت). در ابتدا محقق در مورد θ کم می‌داند چون اطلاع پایه‌ای K اندک است، و دو محقق عقاید مختلفی خواهند داشت که با احتمالهای $p(\theta|K)$ بیان می‌شود. سپس آزمایشهایی ترتیب داده می‌شود، داده‌های x ای به دست می‌آیند، و این احتمالات به صورتی که قبلاً گفته شد به $p(\theta|x, K)$ تبدیل می‌شوند. می‌توان نشان داد [۱۴] که تحت شرایطی که معمولاً حاصل می‌شود، چنانچه تعداد داده‌ها افزایش یابد، عقاید مختلف راجع به θ همگرا می‌شود و معمولاً به معلوم شدن θ و یا حداقل، تعیین آن با دقت زیاد، می‌انجامد. چیزی که ما در عمل مشاهده می‌کنیم این است: با اینکه دانشمندان ممکن است در ابتدا نظراتی متفاوت و در برخی موارد مغایر هم داشته باشند ولی نهایتاً به توافق می‌رسند. همان‌طور که شخصی گفته است: عینیت در حقیقت همان توافق است. بنابراین توافق خوبی بین تجربه علمی و پارادایم بیزی وجود دارد. مواردی هست که تقریباً همه درباره یک احتمال توافق دارند. این را در مبحث تعویض‌پذیری در بخش ۱۴ بررسی خواهیم کرد.

حال می‌توان فهمید که چرا نوعی اکراه در پذیرش نکته (د) که همان استفاده از توزیع احتمال برای توصیف عدم حتمیت پارامتر است، وجود دارد. علت این است که جنبه اساسی ذهنی از نظر دور مانده است. وقتی تعداد داده‌ها کم است مقدار $p(\theta, a)$ در اذهان مختلف متفاوت است. وقتی داده‌ها افزایش می‌یابند، توافق حاصل می‌شود. توجه کنید که $p(x|\theta, a)$ نیز ذهنی است. وقتی دو آماردان در تحلیل یک مجموعه داده‌ها دو الگوی مختلف به کار می‌برند، این موضوع آشکارا تصدیق می‌شود. ما دوباره به این مطلب در بخشهای ۹ و ۱۲ برمی‌گردیم تا نقش الگوها را بررسی کنیم.

۹. الگوها

موضوع الگوها [مدلها] را دراپر [۱۲] به دقت از دیدگاه بیزی تحلیل کرده است و برای شرحی مفصل‌تر از آنچه در اینجا می‌آید باید به مقاله او مراجعه کرد. موضع فلسفی ارائه‌شده در این مقاله این است که عدم حتمیت فقط باید برحسب احتمال مورد نظر شما توصیف شود. اجرای این نظر نیاز به ساختن توزیعهای احتمال برای تمام عناصر غیرحتمی در پدیده واقعی مورد مطالعه دارد. تعیین کامل احتمال، الگو یا مدل (احتمال) نامیده می‌شود،

احتمالاتی و با گذار از توزیع توأم (θ, a) به چگالی حاشیه‌ای θ ، رابطه (۲)، برطرف کرد.

معرفی پارامترها ساخت الگو را به محاسبه $p(x|\theta, a)$ و $p(\theta, a)$ تقلیل می‌دهد. $p(y|\theta, a)$ نیز مطرح می‌شود ولی تعیین آن شبیه به $p(x|\theta, a)$ است و نیازی نیست به‌طور جداگانه مورد بحث قرار گیرد. در اینجا تفاوت بین استفاده ما از «الگو» و استفاده معمولی از آن که فقط مبتنی بر توزیع داده‌ها به شرط در دست داشتن پارامترهاست، دیده می‌شود. تعریف ما شامل توزیع پارامترهاست زیرا آنها بخش مهمی از عدم حتمیت موجود را شکل می‌دهند. غالب متون فعلی مربوط به الگوها معطوف به داده‌هاست و ما در بخش بعد به این موضوع می‌پردازیم. در حال حاضر دوباره به این نکته اشاره می‌کنیم که حتی $p(x|\theta, a)$ نیز ذهنی است. معمولاً به این دلیل آن را، به غلط، عینی می‌پندارند که در بسیاری موارد، اطلاعات مردم درباره داده‌ها بیشتر است تا درباره پارامترها، و ما در بخش ۸ دیدیم که با افزایش اطلاعات، مردم به توافق نزدیک می‌شوند.

چرا سلسله مراحل تعیین $p(x|\theta, a)$ و $p(\theta, a)$ و بعد از آن $p(\theta|x)$ را انجام می‌دهیم؟ اگر $p(\theta, a)$ قابل محاسبه است چرا مستقیماً $p(\theta|x)$ را محاسبه نمی‌کنیم تا از بعضی محاسبات پیچیده پرهیز کرده باشیم؟ این سؤال را، با استفاده از اصطلاحاتی که من به آنها علاقه‌ای ندارم، می‌توان این طور نیز بیان کرد: اگر توزیع پیشین شما مستقیماً قابل تعیین است چرا توزیع پسین شما نباشد؟ بخشی از جواب در اطلاعاتی نهفته است که معمولاً درباره چگالی داده‌ها در دست است ولی دلیل اساسی، تمایل ما به سازگاری مفاهیم است. مجموعه‌ای از گزاره‌های راجع به عدم حتمیت سازگار نامیده می‌شوند هرگاه از قوانین حساب احتمال تبعیت کنند. مثلاً دو گزاره $p(A|B) = 0.7$ و $p(B|A) = 0.4$ سازگار نیستند. $(B \sim A)$ نمایانگر متمم مجموعه B است. فرض کنید A گزاره‌ای راجع به داده‌های x است و B گزاره‌ای درباره پارامتر θ . دو گزاره اول که راجع به عدم حتمیت داخل داده‌ها هستند با اولین گزاره راجع به پارامتر یعنی $p(B|A) = 0.5$ سازگارند. [فرض کنید $0.36 = 0.4/1.1$]. ولی هر سه گزاره با دومین گزاره راجع به پارامتر یعنی $p(B| \sim A) = 0.3$ سازگار نیستند. با فرض $p(B) = 0.36$ مقدار سازگاری ۰.۲۲ است. روش استاندارد این اطمینان را به شما می‌دهد که برای هر مقدار داده‌ها، A یا $\sim A$ ، زمینه را آماده کرده‌اید، و استنباط‌های نهایی راجع به B در مجموع معنی دارند.

مراجع

1. Aitken, C. G. G. and Stoney, D. A. (1991) *The Use of Statistics in Forensic Science*. Chichester: Horwood.
2. Bartholomew, D. J. (1988) Probability, statistics and theology (with discussion). *J. R. Statist. Soc. A*, **151**, 137-178.
3. Berger, J. O. and Delampady, M. (1987) Testing precise hypotheses (with discussion). *Statist. Sci.*, **2**, 317-352.
4. Berger, J. O. and Wolpert, R. L. (1988) *The Likelihood Principle*. Hayward: Institute of Mathematical Statistics.

کمک می‌کند. اغلب روشهای مدرن استنباطی را می‌توان از طریق پارامتری مانند θ بیان کرد که نشان‌دهنده ارتباط بین دو مجموعه از داده‌هاست. این بحث را با استفاده از حساب احتمال بسط می‌دهیم تا θ را نیز شامل شود:

$$p(y|x) = \int p(y|\theta, x)p(\theta|x) d\theta$$

معمولاً فرض می‌کنند که وقتی θ معلوم است داده‌های گذشته مطرح نیست. به زبان احتمال، با فرض در دست داشتن θ ، داده‌های x و y مستقل‌اند. با توجه به این واقعیت و با تعیین $p(\theta|x)$ به وسیله قضیه بیز خواهیم داشت

$$p(y|x) = \int p(y|\theta)p(x|\theta)p(\theta) d\theta / \int p(x|\theta)p(\theta) d\theta$$

حال باید عدم حتمیت مربوط به پارامتر یا $p(\theta)$ و عدم حتمیت‌های دو دسته داده، $p(x|\theta)$ و $p(y|\theta)$ را به شرط در دست داشتن پارامتر محاسبه کرد. در بسیاری موارد، استنباط را می‌توان در $p(\theta|x)$ متوقف کرد و $p(y|\theta)$ را به عهده دیگران گذاشت. در مثال بالا در مورد دارو، دکتر نیاز دارد توزیع کارایی θ را بداند تا $p(y|\theta)$ را برای مریض مورد نظرش محاسبه کند. گرچه موضوع اصلی استنباط در بسیاری از موارد تعیین $p(\theta|x)$ است و استنباط بحق برحسب آن بیان می‌شود ولی بررسی $p(y|x)$ هم فایده مهمی دارد. این فایده از این حقیقت نشأت می‌گیرد که y در نهایت مشاهده خواهد شد؛ دکتر نهایتاً خواهد دید چه بر سر مریضش خواهد آمد. پارامتر معمولاً مشاهده نمی‌شود. عدم حتمیت θ غالباً باقی می‌ماند ولی عدم حتمیت y از میان می‌رود. این خصیصه امکان می‌دهد که کارایی استنباط را با استفاده از تعمیم قانون امتیازدهی بخش ۴ نمایش دهیم. اگر استنباط مورد نظر، $p(y|x)$ باشد و بعداً مشاهده کنیم که y برابر y_0 است، آنگاه تابع امتیاز $\{p(\cdot|x), y_0, S\}$ نشان می‌دهد که استنباط تا چه حدی خوب بوده و به این وسیله توانایی مشتری و آماردان محک زده می‌شود. این روش در علم هواشناسی مثلاً برای پیش‌بینی بارندگی فردا، مورد استفاده قرار گرفته است. چنین روشهایی به آسانی برای $p(\theta|x)$ در دسترس نیستند. یکی از انتقادهایی که علیه سطوح معنی‌دار بودن مطرح می‌شود این است که بررسی چندانی در این باره به عمل نیامده که چه تعداد از فرضیهایی که در سطح ۰.۵ رد شده‌اند، بعداً درست از آب درآمده‌اند. دلیلی ندارد که فکر کنیم این سطح ۰.۵ است. از نظریه چنین برمی‌آید که باید خیلی بالاتر باشد و در نتیجه معنی‌دار بودن خیلی سریع حاصل می‌شود، اگر هواشناسی پیش‌بینی کند فقط ۰.۵٪ روزها بارندگی خواهد بود و بعداً دیده شود ۰.۲۰٪ روزها باران باریده، آن هواشناس زیاد هم قابل قدردانی نیست. بهترین راه ارزیابی کیفیت استنباطها بررسی $p(\theta|x)$ از طریق احتمالهای داده‌ها یا $p(y|x)$ است که آنها تولید می‌کنند.

همان‌طور که قبلاً اشاره شد، معمولاً لازم است پارامتر مزاحم a ، علاوه بر θ ، وارد کار شود، تا ارتباط بین x و y را به قدر کفایت توصیف کند، و استقلال بین آنها را به شرط (θ, a) برقرار سازد. در مثال دارو، a ممکن است خصوصیات شخصی بیمار باشد. پارامترهای مزاحم مشکلات بسیار بزرگی را برای برخی از اشکال استنباط ایجاد می‌کنند، مثلاً استنباطهایی که فقط مبتنی بر درست‌نمایی‌اند. اما این مسائل را می‌توان به راحتی در چارچوب

- 20a. Lehmann, E. L. (1983) *Theory of Point Estimation*. New York: Wiley.
- 20b. — (1986) *Testing Statistical Hypotheses*. New York: Wiley.
21. O'Hagan, A. (1995) Fractional Bayes factors for model comparison (with discussion). *J. R. Statist. Soc. B*, **57**, 99-138.
22. Onions, C. T. (ed.) (1956) *The Shorter English Dictionary*. Oxford: Clarendon.
23. Pearson, K. (1892) *The Grammar of Science*. London: Black.
24. Ramsey, F. P. (1926) Truth and probability. In *The Foundations of Mathematics and Other Logical Essays* (ed. R. B. Braithwaite), pp. 156-198. London: Kegan Paul.
- 25a. Savage, L. J. (1954) *The Foundations of Statistics*. New York: Wiley.
- 25b. — (1977) The shifting foundations of statistics. In *Logic, Laws and Life: Some Philosophical Complications* (ed. R. G. Colodny), pp. 3-18. Pittsburgh: Pittsburgh University Press.
26. Stein, C. (1956) Inadmissibility of the usual estimation of the mean of a multivariate normal distribution. In *Proc. 3rd Berkeley Symp. Mathematical Statistics and Probability* (eds J. Neyman and E. L. Scott), vol. 1, pp. 197-206. Berkeley: University of California Press.
27. Walley, P. (1991) *Statistical Reasoning with Imprecise Probabilities*. London: Chapman and Hall.
5. Bernardo, J. M. (1999) Nested hypothesis testing: the Bayesian reference criterion. In *Bayesian Statistics 6* (eds J. M. Bernardo, J. O. Berger, A. P. Dawid and A. F. M. Smith). Oxford: Clarendon.
6. Bernardo, J. M., Berger, J. O., Dawid, A. P. and Smith, A. F. M. (eds) (1999) *Bayesian Statistics 6*. Oxford: Clarendon.
7. Bernardo, J. M. and Smith, A. F. M. (1994) *Bayesian Theory*. Chichester: Wiley.
8. Box, G. E. P. (1980) Sampling and Bayes inference in scientific modelling (with discussion). *J. R. Statist. Soc. A*, **143**, 383-430.
9. Box, G. E. P. and Cox, D. R. (1964) An analysis of transformations (with discussion). *J. R. Statist. Soc. B*, **26**, 211-252.
10. Dawid, A. P., Stone, M. and Zidek, J. V. (1973) Marginalization paradoxes in Bayesian and structural inference (with discussion). *J. R. Statist. Soc. B*, **35**, 189-233.
11. DeGroot, M. H. (1970) *Optimal Statistical Decisions*. New York: McGraw-Hill.
12. Draper, D. (1995) Assessment and propagation of model uncertainty (with discussion). *J. R. Statist. Soc. B*, **57**, 45-98.
13. Duckworth, F. (1998) The quantification of risk. *RSS News*, **26**, no. 2, 10-12.
14. Edwards, W. L., Lindman, H. and Savage, L. J. (1963) Bayesian statistical inference for psychological research. *Psychol. Rev.*, **70**, 193-242.
15. Eggleston, R. (1983) *Evidence, Proof and Probability*. London: Weidenfeld and Nicolson.
- 16a. de Finetti, B. (1974) *Theory of Probability*, vol. 1. Chichester: Wiley.
- 16b. — (1975) *Theory of Probability*, vol. 2. Chichester: Wiley.
17. Fisher, R. A. (1935) *The Design of Experiments*. Edinburgh: Oliver and Boyd.
18. Healy, M. J. R. (1969) Rao's paradox concerning multivariate tests of significance. *Biometrics*, **25**, 411-413.
19. Jeffreys, H. (1961) *Theory of Probability*. Oxford: Clarendon.

- Dennis V. Lindley, "The philosophy of statistics", *The Statistician* (3) **49** (2000) 293-337.

به علت طولانی بودن مقاله، ترجمه آن به دو قسمت تقطیک شده است. قسمت دوم در شماره آینده خواهد آمد.

* دنيس ليندلى، انگلستان

thombayes@aol.com