

دستاوردهای پژوهشهای اخیر در ریاضیات*

(۲)

۱۶. سیستمهای ذرات برهمکنش کننده

این مبحث احتمال، با پیکربندیهای ذراتی که در طول زمان به طور تصادفی تحول می یابند سروکار دارد. ذره معمولاً طبق قانون خاص خودش حرکت می کند، نابود می شود، و یا ذرات جدیدی پدید می آورد. این قانون فقط به حالت سیستم در یک همسایگی ذره بستگی دارد.

مفهوم سیستمهای ذرات برهمکنش کننده از مطالعه مدل ایزینگ که قبلاً شرح آن آمد، نشأت گرفته است ولی امروز کاربردهای بسیار متنوعی از بررسی نظامهای زیستی گرفته تا تصویرپردازی برای پزشکی و مقاصد نظامی، پیدا کرده است. فرایند تماسی^۱، که مدلی بنیادی برای پراکنش جمعیتی از موجودات زنده است، در سال ۱۹۷۴ ارائه شد. برخلاف مدلهای مبتنی بر فرایند شاخه ای، این سیستم امکان می دهد که در هر واحد سطح فقط تعدادی محدود از عناصر وجود داشته باشند. این فید که از نظر فیزیکی موجه است، بررسی تحلیلی سیستم را خیلی مشکل می کند. ویژگیهای بنیادی حالت یک بعدی (رشد خطی سیستم در حالت بقا، و کاهش نمایی احتمال بقا در حالت زوال) در اوایل دهه اخیر محرز شد ولی ویژگیهای متناظر برای حالت مهم دوبعدی، تازه در همین اواخر اثبات شده است.

اکنون آشکالی از فرایند تماسی در حالتی که ذرات از انواع گوناگون هستند، برای مطالعه رقابت گونه ها و نظام میزبان-انگل یا شکار-شکارگر به کار می روند. مدلهای دیگری که ارتباط نزدیکی با نشأت^۲ دارند، در مطالعه آتش سوزیهای اخیر در جنگلهای یلوستون به کار گرفته شده اند. مثال سوم، بررسی توزیع و

الگوی حرکتی کریل^۱های اقیانوس منجمد جنوبی است. سیستم اخیر بسیار جالب توجه است زیرا الگوهایی را روی مقیاسهای مکانی و زمانی چندگانه نمایش می دهد. این مثالها فقط سه نمونه از مثالهای روزافزونی هستند که نشان می دهند سیستمهای ذرات برهمکنش کننده چارچوبی مناسب برای فهم مکانیسمی است که در چند پدیده اکولوژیک در کار است.

۱۷. آمار مکانی

انگیزه نیرومندی که علاقه پژوهشگران را به بسط نظریه و روشهای آمار مکانی^۱ [فضایی] جلب می کند، دسته ای از کاربردها از جمله سنجش از دور، برآورد منابع، کشاورزی و جنگل داری، اقیانوس نگاری، اکولوژی، و دیدبانی^۲ محیط زیست است. زنجیره ای که این کاربردها را به هم می پیوندد، عبارت است از مشخص سازی و بهره گیری از مجاورت^۳. در زیر، برخی از زمینه ها و امکانات مهم برای پیشرفتهای آتی را خلاصه وار ذکر می کنیم.

در سنجش از دور ژئوفیزیکی، داده ها معمولاً به صورت شبکه بندی شده به دست می آیند. و اثرات جوی ناخواسته و خطاهای مربوط به مکان و اندازه گیری، آنها را مخدوش می کند؛ در این مورد اغلب از باندهای چند طول موجی استفاده می شود که باعث می شود داده ها چندمتغیری، و غالباً به صورت کمیت های بزرگ، باشد. سؤالی که در این زمینه مطرح می شود، از این قبیل است: چگونه می توان از خطا اجتناب کرد، چگونه می توان اطلاعاتی

۱. نوعی جانور کوچک، دریایی، از سخت پوستان، که خوراک وال بدون دندان است.

2. spatial statistics 3. monitoring 4. proximity

1. contact process 2. percolation

داده می‌شود، روشهای کاهش بُعد برای داده‌های مکانی چندمتغیری مبتنی بر محکهای ساختار مکانی، ابداع مدل‌های میدانی مکانی غیرگوسی که از لحاظ ریاضی قابل تحلیل باشد برای میدانهای پارامتری پیوسته، روشهای درونبایی مکانی نابارامتری غیرخطی برای میدانهای مکانی ناپیوسته، اثبات ارتباطات نظری بین روشهای اسپلاین، و پالیتهای وینر، و مجموعه روشهای مقایسه برای درجه‌بندی تکنیکهای برآورد مکانی و تعیین دقت آنها.

۱۸. روشهای آماری برای کیفیت و بهره‌وری

بعد از جنگ دوم جهانی، کارخانه‌هایی که بهره‌وری خود و کیفیت محصولاتشان را بالا برده‌اند، به حیات خود ادامه داده‌اند و کارشان رونق یافته است ولی کارخانه‌هایی که چنین نکرده‌اند، عملکرد خوبی نداشته‌اند و یا به کلی از صحنه خارج شده‌اند. مهندسان برای بالا بردن کیفیت و بهره‌وری ناچارند از روشهای آماری در تحلیل فرایندهای تولید استفاده کنند.

در این زمینه، چهار مقوله از روشهای آماری بسیار زیاد به کار می‌روند: (۱) کنترل فرایند آماری، (۲) طرح آزمایشهای آماری، (۳) اعتمادپذیری، و (۴) نمونه‌گیری برای پذیرش. کنترل فرایند آماری مشتمل است بر روشهایی برای تبیین عملکرد یک فرایند مهندسی در طی زمان به خصوص برای تشخیص تغییرات؛ مهم‌ترین روشهای این مقوله، روشهای مبتنی بر نمودار کنترل هستند و این مبحثی است که نیاز به بازنگری دارد. بیشتر روشهایی که امروز در این زمینه مطرح‌اند، چند دهه قبل و در دوره‌ای ابداع شده‌اند که محاسبه با دست انجام می‌شد و فرض می‌کردند داده‌ها به صورت نرمال توزیع شده باشند، زیرا در غیر این صورت محاسبات فوق‌العاده دشوار می‌شد. به عنوان مثال، میزان تغییرات در روشهای مبتنی بر نمودار کنترل اغلب بر اساس دامنه داده‌ها اندازه‌گیری می‌شود که محاسبه آن آسان است. این روشها باید از اساس مورد بازنگری قرار گیرند.

طرح آزمایش، مجموعه بسیار مهمی از تکنیکهاست برای مجزا کردن عواملی که می‌توان آنها را تغییر داد تا فرایند بهتر شود؛ بنابراین، برای اینکه یک فرایند مهندسی به طور پیوسته بهتر شود، طرحهای آزمایش باید به منظور بررسی و آزمودن فرایند به طور پیوسته اجرا شوند. بیشتر پژوهشهایی که در این زمینه انجام می‌شود معطوف به فهم این موضوع است که سطح متوسط یک پاسخ چه وابستگی به عاملها دارد؛ در این مورد مدل‌هایی به کار می‌روند که در آنها واریانس یا ثابت است و یا، اگر تغییر کند، به عنوان یک پارامتر مزاحم در نظر گرفته می‌شود. ولی در بسیاری از فرایندهای مهندسی، واریانس تغییر می‌کند و همانقدر مهم است که تغییر سطح میانگین؛ و مثلاً در طرح «نیرومند» این امر صادق است. در طرح نیرومند این موضوع پذیرفته می‌شود که کیفیت بیشتر محصولات به تعداد زیادی عامل بستگی دارد. کنترل برخی از این عوامل که «پارامترهای کنترل» نامیده می‌شوند ارزان است وای کنترل عوامل دیگر، موسوم به «پارامترهای نوفه»^۱ گران یا حتی غیرممکن است. هدف، طراحی محصولی است که عملکردش در برابر پارامترهای نوفه، غیرحساس یا نیرومند باشد.

روشهایی که در مقوله «اعتمادپذیری» قرار می‌گیرند، برای مطالعه نیم‌عمر مؤلفه‌ها و سیستمها و مربوط ساختن این نیم‌عمرها به عواملی که عملکرد را مشخص می‌کنند، به کار می‌روند. در این زمینه، لازم است جامعه پژوهشگران

را که از طول موجهای گوناگون به دست می‌آید با هم ترکیب کرد، چگونه می‌شود الگوهایی استخراج کرد، کار تا چه حدی خوب پیش می‌رود، داده‌ها را تا چه اندازه می‌توان افزایش داد، و فایده داده‌های اضافی چیست. ملاحظات مربوط به مجاورت، تأثیر عمیقی بر رهیافتهای آماری برای پاسخگویی به همه این پرسشها خواهد گذاشت.

در روشهای حذف اثرات ناخواسته نمی‌توان این واقعیت را از نظر دور داشت که آن اثرات غالباً به طور خیالی دقیق مشخص نشده‌اند. شیوه‌های استخراج الگوهای مربوطه از ترکیب اطلاعات چندباندی باید قویاً متکی باشند به مدل‌های احتمالی که به اندازه کافی برمایه‌اند. از سوی دیگر، برآوردهای دقت آماری باید تا حد امکان مستقل از مدل باشند. این الزامات، مسائل مهمی در پیش آماردانان پژوهشگر می‌نهد، که دشواریهای فنی و محاسباتی آنها به اندازه قابل ملاحظه‌ای بیشتر از دشواریهای مسائل مربوطه در تحلیل سریهای زمانی یا فرایندهای تصادفی است.

داده‌های مکانی مربوط به برآورد منابع و دیدبانی محیط زیست نوعاً به شکل شبکه‌های منظم به دست نمی‌آیند و نسبتاً پراکنده‌اند. اگرچه داده‌های حاصل از حفاری مغزه‌ای ممکن است تفکیک عمودی خوبی داشته باشد، ولی احتمال دارد حفاری بیشتر در مناطق «پرعیار» انجام شده باشد. خواست ما معمولاً برآورد کل منابع، توزیع فراوانی منابع، و برنامه معقولی برای بهره‌برداری است که مبتنی بر برآوردهای موضعی از منابع، همراه با شاخصهای عدم قطعیت باشد تا معلوم شود در آینده به چه میزان کارهای استخراجی دیگر نیاز هست. مجموعه روشهای آماری اولیه که هنوز هم کاربرد گسترده‌ای دارد، و مبتنی بر مدلسازی فرایند گاوسی است، با توجه به خصلت داده‌های مربوط به منابع که بسیار خط‌آمیزند، چندان مناسب نیست و پژوهش بسیار بیشتری لازم است تا روشهایی کشف شوند که با توجه به کیفیات خاص داده‌های منابع، مناسب باشند.

در دیدبانی محیط زیست، نمونه‌گیری ممکن است بیشتر در نواحی بسیار آلوده انجام شده باشد. انتخابی بودن و تصادفی بودن داده‌های محیط‌زیستی، مشکلات مهمی در مدلسازی و استنباط آماری پیش می‌آورد. به علاوه، در هر محل دیدبانی، معمولاً یک سری زمانی از داده‌ها وجود دارد که تحت تأثیر تغییرات فصلی است. چیزی که معمولاً می‌خواهیم بدانیم این است که چگونه کیفیت محیط زیست تغییر می‌کند. چون داده‌های ثبت شده معمولاً نمی‌تواند مقدار تغییر [در محیط زیست] را که در میان تغییرات شدید طبیعی قابل اندازه‌گیری باشد بر ملا سازد، به روشهای مناسبی برای ترکیب اطلاعات گرفته شده از چندین محل دیدبانی نیاز داریم. همچنین افزایش و تجدید آرایش منابع دیدبانی مستلزم پژوهشی در مسائل طرح آماری است که محدوده آن از مسائل ایده‌آلی شده و ساده‌ای که جوابهایی برای آنها داریم، بسیار فراتر می‌رود.

در طی دهه گذشته، گامهای بلندی در جهت ایجاد نظریه و روشهایی برای حل مسائل مکانی با استفاده از ابزارهای آماری برداشته شده است. پس، پژوهش موردنیازی که شرح آن در بالا آمد از مبنای محکمی برخوردار است. پیشرفت مهمی که اخیراً از پژوهشهای مربوط به آمار مکانی به دست آمده، ابداع و کاربرد مدل‌های شبکه‌ای مارکوفی مکانی انعطاف‌پذیر برای تصویربرداری و کاربرد تکنیکهای نیرومند از قبیل الگوریتم متروپولیس برای اجرای این مدل‌هاست. از جمله سایر پیشرفتهای می‌توان از موارد زیر نام برد: تکنیکهای هموارسازی مکانی که با پیچیدگی موضعی الگوهای مکانی تطبیق

1. noise

روشهایی ابداع کنند که مهندسان به وسیله آنها بتوانند به جای تحلیل شاخصهای «زمان تا خرابی»، شاخصهای تباهی^۲ را تحلیل کنند که حاوی اطلاع بیشتری است. همچنین در این مبحث لازم است تحقیقات بیشتری در زمینه مدل‌های داده‌هایی که از آزمایشهای تسریع شده زمان تا خرابی به دست می‌آیند، انجام شود.

نمونه‌گیری برای پذیرش، عبارت است از روشهایی برای امتحان کردن محموله‌ای از یک کالا، برای اینکه معلوم شود آن محموله قابل پذیرش هست یا نه. در بعضی از موارد، هر عدد از کالاهای محموله امتحان می‌شود ولی در غالب موارد، نمونه‌ای از کالاها انتخاب می‌شود و با امتحان کردن آن، درباره کل محموله نتیجه‌گیری می‌شود. روشهای تازه‌ای برای پرداختن به این موضوع لازم به نظر می‌رسد. برخی از تحقیقات گذشته نشان داده است که روشهای بیزی راه مناسبی هستند که می‌توان در پیش گرفت.

نمونه‌گیری برای پذیرش، عبارت است از روشهایی برای امتحان کردن محموله‌ای از یک کالا، برای اینکه معلوم شود آن محموله قابل پذیرش هست یا نه. در بعضی از موارد، هر عدد از کالاهای محموله امتحان می‌شود ولی در غالب موارد، نمونه‌ای از کالاها انتخاب می‌شود و با امتحان کردن آن، درباره کل محموله نتیجه‌گیری می‌شود. روشهای تازه‌ای برای پرداختن به این موضوع لازم به نظر می‌رسد. برخی از تحقیقات گذشته نشان داده است که روشهای بیزی راه مناسبی هستند که می‌توان در پیش گرفت.

۱۹. کهداهای گراف

کهدای از یک گراف، گرافی است که از یک زیرگراف گراف اولیه با انقباض یاها به دست می‌آید. چند نتیجه مرتبط با هم در مورد کهداهای گراف، که برخی به نظریه «محض» گرافها تعلق دارند و برخی به طرح الگوریتمها، در زمره دستاوردهای جدید به شمار می‌روند. با به دست آمدن این نتایج، راههای تازه‌ای برای تحقیق گشوده شده و مسائل و فرصتهای تازه‌ای پدیدار گشته است.

در یک مسأله قدیمی و کاملاً حل شده در نظریه جریان شبکه‌ای، گرافی مفروض است که برخی از رأسهایش را «مبدأ» و رأسهای دیگرش را «مقصد» می‌نامند. سؤال این است که آیا مثلاً ده مسیر در گراف وجود دارد که هر یک از مبدأ شروع و به مقصدی ختم شود و این مسیرها هم را قطع نکنند؟ چگونه می‌توان برنامه‌ای برای کامپیوتر نوشت که جواب این سؤال را بدهد؟ یک راه، تهیه فهرست همه مسیرها و آزمودن همه ترکیبات ده تا از آنهاست؛ ولی چنین راهی حتی برای یک گراف کوچک خیلی طولانی است. خوشبختانه (چون این مسأله اهمیت زیاد و کاربردهای فراوانی دارد)، الگوریتم دیگری در دست است که غیر مستقیم تر ولی بسیار کاراست. بنابراین، مسأله را می‌توان کاملاً حل شده دانست.

اگر مسأله را قدری تغییر دهیم، و فرض کنیم ده مبدأ و ده مقصد وجود دارد و بخواهیم مسیر اول از مبدأ اول تا مقصد اول، مسیر دوم از مبدأ دوم تا مقصد دوم، و ... الی آخر، باشد، آن وقت چه می‌شود؟ این مسأله خیلی مشکلتر است. حتی در مورد دو مسیر (دو مبدأ و دو مقصد به جای ده تا)، مسأله مشکل است، و مسأله مربوط به سه مسیر تا همین اواخر حل نشده بود. یکی از احکام مهم به دست آمده در باره کهداهای گراف این است که به ازای هر تعداد (مثلاً ده) مسیر، الگوریتم کارایی برای حل مسأله ده مسیر وجود دارد (کلمه «کارا» در اینجا یک اصطلاح فنی است و به این معنی است که زمان اجرای الگوریتم محدود به یک چندجمله‌ای بر حسب تعداد رؤس گراف است).

دستاوردهای دیگر، اثباتی است از یک حدس قدیمی متعلق به واگنر، حاکی از اینکه در هر گردانی نامتناهی از گرافها، گرافنی وجود دارد که شامل گراف دیگری است. (یک گراف «شامل» گرافی دیگر است اگر دومی را بتوان

2. degradation

می‌توان حکم مشابهی ارائه کرد که از کلیت بسیار برخوردار است و کاربردهای زیادی در عاوم کامپیوتر نظری دارد. فرض کنید می‌خواهیم الگوریتم کارایی برای آزمودن گرافها طرح کنیم که معلوم کند یک گراف، ویژگی خاصی را دارد یا نه. در مورد برخی از ویژگیها، یافتن چنین الگوریتمی غیر ممکن است؛ ولی فرض کنید گرافی که خاصیت مورد نظر را دارد، شامل هیچ گرافی بدون آن خاصیت نیست. (مثلاً این ویژگی که گراف قابل ساختن باشد بدون اینکه گره نخورد، در این شرط صدق می‌کند). پس الگوریتم کارایی برای آزمودن یک گراف از احاط داشتن آن خاصیت وجود دارد، اگرچه ممکن است تاکنون کسی آن را نیافته باشد. باید تأکید کرد که دانستن اینکه الگوریتم کارایی وجود دارد، حتی اگر کسی هنوز آن را پیدا نکرده باشد، خود به خود اطلاع مهم و با معنایی است.

۲۰. اقتصاد ریاضی

گرچه بحث ریاضی درباره عملکرد بازار در قرن گذشته آغاز شد، نخستین توصیف ریاضی دقیق از مبانی اقتصادی عملکرد بازار در اواخر دهه ۱۹۴۰ و اوایل دهه ۱۹۵۰ به عمل آمد. این توصیف به مدل معروف تعادل عمومی (GE) تحت فرض کامل بودن بازارها انجامید یعنی تحت این فرض که همه کالاها را می‌توان در بازارهایی که در یک زمان کالا عرضه می‌کنند و خریدار و فروشنده در این بازارها از اعتبار کامل برخوردارند، خرید و فروش کرد. در حالت تعادل، عرضه با تقاضا برابر است و هر خانوار طبق علائق خود و متناسب با بودجه‌اش عمل می‌کند.

ولی مدل بازارهای کامل، اشکال عمده‌ای دارد. برنامه‌ریزی برای پیشامدها و احتمالات آتی یک قسمت اساسی مسأله تخصیص آماری است و مدل تعادل عمومی را می‌شود به این صورت تعبیر کرد که در این مدل، زمان و عدم قطعیت به وسیله طبقه‌بندی کالاها بر حسب پیشامد [خرید یا فروش] و زمان، منظور می‌شود. ولی تنها یک قید در مورد بودجه، این دیدگاه غیرواقعیتانه را بر این تعبیر تحمیل می‌کند که همه معامله‌ها یکبار به انجام می‌شوند. تحقیقات اخیر در زمینه مدل‌های بازار ناکامل که این اشکال را رفع می‌کند، به پیشرفت مهمی نائل آمده و مسائل و چشم‌اندازهای تازه‌ای را مطرح کرده است.

در مدل تعادل عمومی با فرض ناکامل بودن بازارها (GEI)، معامله‌گران نمی‌توانند همه کالاهای ممکن را یکبار معامله کنند. برای سادگی، فرض کنید کالاهایی فاسدشدنی داشته باشیم که بتوان آنها را در زمان صفر مصرف کرد و یا تحت هر یک از S حالت ممکن، در زمان یک، به علاوه، سرمایه معامله‌گران ممکن است در بعضی حالات زیاد و در بعضی حالات کم باشد. در زمان صفر، معامله‌گران مجازند تعداد محدودی از داراییها را که تحویل آنها به صورت ترکیبات مختلفی از کالاها، بسته به حالت، میسر است معامله کنند. مثلاً ذخیره انبار یک کارخانه، یک نوع دارایی است که در حالاتی که

جمع دو این عدد در N مرحله پیمایی مسلماً آسان است ولی آیا می‌توان این کار را در یک مرحله به وسیله N پردازنده موازی انجام داد؟ واقعیت این است که نمی‌توان چنین کاری کرد و حتی بدتر از این، در اولین نگاه به نظر می‌رسد که مسئله جمع را اصلاً نمی‌توان با استفاده از N پردازنده سریعتر از یک پردازنده حل کرد. بنابراین، مثال جمع در نگاه اول خیلی ناامید کننده به نظر می‌آید زیرا اگر عمل جمع را نتوان به طور موازی انجام داد، چه امیدی به حل موازی مسائل دیگر می‌توان داشت؟

خوشبختانه، امروز به دست آوردن الگوریتم‌های موازی سریع برای بسیاری از مسائل محاسباتی در علوم و مهندسی امکان پذیر است. مثلاً در مورد عمل ساده جمع، الگوریتمی وجود دارد که در صورت استفاده از $N/\log(N)$ پردازنده، فقط $\log(N)$ مرحله دارد؛ اگرچه همان طور که دیدیم، وجود چنین الگوریتمی به هیچ وجه واضح و بدیهی نیست. طی پنج سال گذشته، پیشرفت چشمگیری در کشف روشهای جالب و بسیار کارا برای حل برخی مسائل مهم که ذاتاً ترتیبی به نظر می‌رسند، به دست آمده است. از جمله این روشها، الگوریتم‌هایی برای حساب، محاسبات ماتریسی، عملیات روی چندجمله‌ایها، معادله‌های دیفرانسیل، همبندی گرافها، انقباض و محاسبه درخت، جورسازی^۱ گرافها، مسائل مجموعه مستقل، برنامه‌ریزی خطی، هندسه محاسباتی، جورسازی رشته‌ها، و برنامه‌ریزی پویا را می‌توان نام برد. همچنین پیشرفت زیادی در زمینه مسائل نهفته در طراحی ماشینهای موازی حاصل شده است. مثلاً انواع متعددی از شبکه مخابراتی برای وصل کردن پردازنده‌های یک ماشین موازی به یکدیگر ابداع شده و الگوریتم‌های سریعی برای هدایت داده‌های صحیح به مکان صحیح در زمان صحیح به دست آمده است. این کار تا حد زیادی جنبه ریاضی دارد و مبتنی است بر استفاده فراوان از ترکیبات، تحلیل احتمالاتی، و جبر. درواقع، معمولاً در کامپیوترهای موازی از شبکه‌های ارتباطی و الگوریتم‌های مسیره‌ای که مبتنی بر ترکیبات هستند استفاده می‌شود. پس در این زمینه هم، این دستاوردهای اساسی، امکانات زیادی برای پیشرفتهای بعدی فراهم کرده است.

۲۲. الگوریتم‌های تصادفی (شده)

دانشمندان کامپیوتر طی ۱۵ سال گذشته به مزایای فراوان الگوریتم‌هایی که در حین اجرا، به اصطلاح «ناس می‌ریزند» [یعنی مبتنی بر انتخاب تصادفی عدد، داده، یا رویداد هستند] پی برده‌اند. برای کارهای بسیار متنوعی، از آزمایش اول بودن یک عدد گرفته تا تخصیص منابع در سیستم‌های کامپیوتر پیچیده، ساده‌ترین و کاراترین الگوریتم‌هایی که در حال حاضر می‌شناسیم، الگوریتم‌های تصادفی شده هستند. بنابراین، افزایش آگاهی در زمینه این گونه الگوریتم‌ها، مسأله‌ای است که پژوهشگران را به مبارزه می‌طلبد و زاینده امکاناتی در آینده است.

تقریباً از همان آغاز عصر کامپیوتر، مولدهای اعداد تصادفی، به کمک تکنیک‌های نمونه‌گیری پیچیده‌ای که مجموعاً روش مونت کارلو نامیده می‌شوند، در شبیه‌سازی سیستم‌های پیچیده‌ای که شامل صف بندی و سایر پدیده‌های تصادفی هستند، و نیز در برآورد انتگرالهای چند بعدی و سایر کمیت‌های ریاضی به کار می‌رفته‌اند.

عامل مهمی که توجه متخصصان کامپیوتر را به کاربردهای گسترده‌تر

برنامه‌های تولید خوب پیش می‌رود، مقدار زیادی کالا و در غیراین حالات، مقدار کمی کالا تحویل می‌دهد.

به کمک ابزارهایی از تئوری دیفرانسیل می‌توان نشان داد که تعادل GEI به ویژگی‌های بسیار شگفت‌انگیز متعددی دارد. نخست اینکه، درست برخلاف بازارهای کامل، اگر دارایی‌ها کمتر از حالتها باشند، آنگاه تعادلهای GEI «نوعاً» ناکارا هستند. ناکارایی این تعادلهای نه فقط به دلیل نبود برخی بازارهای دارایی است، بلکه از بازارهایی هم که وجود دارند، به درستی استفاده نمی‌شود. خصوصاً تعجب آور دیگری که مدل GEI دارد، ویژگی‌های خاصی است که تعادل پولی می‌تواند در این مدل داشته باشد. اگر کالایی وجود داشته باشد که هیچ اثری بر مطلوبیت نداشته باشد، و جزو مایه اولیه هیچ معامله‌گری نباشد، دوتا از ویژگی‌های پول را دارد. اگر دارایی‌ها قابل تحویل به این پول باشند، آنگاه تحت شرایط قضیه ناکارایی، معمولاً $S - 1$ بعد متمایز برای تخصیص کالا در حالت تعادل وجود دارد. ولی اگر بازار دارایی «کامل» باشد، معمولاً فقط تعادلهای منفرد وجود دارد.

کنار گذاشتن فقط یک دارایی، موجب برش عدم تعین از حالت صفر بعدی به $S - 1$ بعدی می‌شود. این بعد، هر قدر تعداد دارایی‌ها کم باشد، ثابت است. بدون استفاده از تئوری دیفرانسیل، دلیل قابل قبولی برای پیش‌بینی چنین نتیجه‌ای وجود نداشت و مسلماً اثبات آن نیز به صورتی قانع‌کننده میسر نبود. توجه کنید که بدون یک کالای پولی معمولاً تعدادی متناهی حالت تعادل وجود دارد. در مدل GEI پول نقش چشمگیری دارد.

این ایده‌های ریاضی کاربردی دیگری هم در مدل GEI دارند. مدتها تصور می‌شد که شاید یک مدل GEI وجود نداشته باشد که به مفهومی «نیرومند» باشد. ترفندی که برای تحلیل صحیح مسأله به کار می‌رفت، پذیرفتن این موضوع بود که محدودیت بودجه GEI را می‌توان بر حسب حیطه بازده‌های مالی در تمام S حالت - و بنابراین بر حسب یک خمینه گراسمان که با استفاده از فضای حالات ساخته شود - بیان کرد. استدلال‌هایی مبتنی بر نظریه درجه نشان می‌دهند که دستگاه معادلاتی که به وسیله تعادل GEI بر این خمینه گراسمان تعریف می‌شود، نوعاً جواب دارد.

۲۱. الگوریتم‌ها و معماریهای موازی

پیشرفتهای چشمگیر در تکنولوژی تولید کامپیوتر، تأثیر چشمگیری هم بر نحوه استفاده ما از کامپیوتر داشته است. مثلاً امروز قویترین کامپیوترها در واقع مرکب از تعداد زیادی پردازنده کوچک (یعنی تراشه) هستند که با هم جمع شده و یک ماشین موازی پدید آورده‌اند. در چنین ماشین موازی، پردازنده‌ها معمولاً ماشینهای ترتیبی سنتی هستند که با هم برای حل یک مسأله بزرگ، کار می‌کنند. قاعدتاً انتظار داریم که با گردهم آوردن N پردازنده برای انجام دادن یک کار خاص، آن کار N بار سریعتر از وقتی که فقط یک پردازنده کار می‌کند، انجام شود. اگرچه از لحاظ شهودی به نظر می‌رسد که N پردازنده باید بتواند یک مسأله را N بار سریعتر از یک پردازنده تنها حل کنند، این امر همیشه امکان پذیر نیست. در واقع، ممکن است طرح الگوریتم‌هایی موازی که در ماشینهای N پردازنده، N بار سریعتر از ماشینهای تک پردازنده اجرا شوند بسیار مشکل باشد.

برای اینکه ببینیم وقتی می‌خواهیم یک الگوریتم ترتیبی را به صورت موازی درآوریم چه مشکلاتی ممکن است پیش آید، مثالی بسیار مقدماتی می‌آوریم که عبارت است از کار «پیش‌بافتاده» جمع کردن دو عدد N رقمی.

1. matching

از خاصیت‌های دنباله‌های تصادفی را دارند هر چند که با فرایندی کاملاً تعینی تولید می‌شوند. در حال حاضر، پژوهش عمیقی در ویژگی‌های مولدهای اعداد شبه‌تصادفی در جریان است. هدف این پژوهش، نشان دادن این نکته است که مادامی که بذر تصادفی باشد، محصول آن را نمی‌توان به وسیله هیچ آزمون محاسباتی با زمان چندجمله‌ای، از یک دنباله کاملاً تصادفی تمیز داد.

و بالاخره، بیشتر تحقیقات نظری جدید معطوف به رابطه بین مولدهای شبه‌تصادفی و مفهوم تابع یکطرفه است که مفهوم اساسی در رمزنگاری است. تابع یکطرفه تابعی است که محاسبه آن آسان ولی وارون کردنش مشکل است. نشان داده‌اند که هر تابع یکطرفه را می‌توان برای ساختن یک مولد اعداد شبه‌تصادفی که توجیه دقیقی داشته باشد به کار برد. متأسفانه، پژوهشگران نظریه پیچیدگی محاسبه هنوز حتی وجود توابع یکطرفه را معلوم نکرده‌اند. این یکی از مسائل مهم متعددی است که باقی مانده و باید حل شود.

۲۳. الگوریتم چندقطبی سریع

در زمینه الگوریتم‌های مربوط به مسائلی از قبیل باریکه‌های ذرات در فیزیک پلاسما، آکوستیک زیرآبی، مدلسازی هواکولی، و حتی جنبه‌های بسیار مهمی از آنرویدینامیک، امکانات زیادی برای پیشرفت وجود دارد. یک نوع محاسبه بنیادی که در این کاربردها اهمیت اساسی دارد، محاسبه برهمکنشها در سیستمی بس ذره‌ای است. این برهمکنشها اغلب بلند برد هستند، بنابراین همه جفتهای ذرات را باید در نظر گرفت. به خاطر این قید، محاسبه مستقیم همه برهمکنشها در سیستمی متشکل از N جسم مستلزم عملیاتی است که تعداد آنها از مرتبه N^2 است. این نوع محاسبه را محاسبه پتانسیل N جسمی می‌نامیم.

تلاشهایی به منظور کاهش پیچیدگی محاسباتی مسأله پتانسیل N جسمی به عمل آمده است. قدیمیترین روشی که در این زمینه به کار می‌رود، روش ذره در سلول (PIC) است که مستلزم عملیاتی است که تعداد آنها از مرتبه $N \log(N)$ است. متأسفانه، این روشها نیز مستلزم شبکه‌های است که تفکیک محدودی را به دست می‌دهد و وقتی توزیع ذره ناپیکنواخت باشد، ناکاراست. یکی از روشهای جدیدتر، روش «حل‌کننده‌های سلسله مراتبی»^۱ است که روشهایی بی‌شبکه برای شبیه‌سازیهای بس ذره‌ای هستند و پیچیدگی محاسباتی آنها نیز از مرتبه $N \log(N)$ برآورد شده است.

در روش چندقطبی سریع (FMM) که اخیراً پیدا شده است، مقدار کار لازم برای محاسبه همه برهمکنشهای دوه‌دو با خطایی در حدود خطای گرد کردن، بدون توجه به توزیع ذره، فقط متناسب با N است. این روش همچون روش حل‌کننده‌های سلسله مراتبی، یک الگوریتم «تجزیه و حل» است که مبتنی بر ایده تجمع ذرات روی مقیاسهای طولی مکانی گوناگون است؛ و در حقیقت مبتنی است بر ترکیبی از فیزیک (انبساطهای چندقطبی)، ریاضیات (نظریه تقریب)، و علم کامپیوتر، و کاربرد آن رو به فزونی است. تکنیکهایی که در این زمینه ابداع می‌شوند، کاربردهای صنعتی بلاواسطه‌ای دارند. نتیجه استفاده سرراست از شبیه‌سازیهای ذرات، قابل ملاحظه است و تقریباً بی‌واسطه به دست می‌آید این نوع شبیه‌سازیها در مدلسازیهای ادوات ریز موجی با باریکه الکترونی پرتوان (مانند زیروترونها و لیزرهای آزاد الکترون)، باریکه‌های ذرات، ادوات گداخت کنترل‌شده، و نظام اینها، به کار می‌رود.

«تصادفی‌سازی» جلب کرد، کشف دو الگوریتم تصادفی کارا برای آزمایش اول بودن یک عدد بود، که در سال ۱۹۷۵ روی داد. هر دو الگوریتم مبتنی بر نشانه‌های [شواهد] مرکب بودن عددند. مثال ساده‌ای از مفهوم «نشانه» از قضیه‌ای به دست می‌آید که منسوب به فرماست. این قضیه می‌گوید که اگر n یک عدد اول باشد، به ازای هر عدد صحیحی چون m که مضربی از n نباشد، $m^{(n-1)} - 1$ مضربی از n است. اگر این محاسبه برای عددی مثل m انجام شود و نتیجه‌ای که قضیه فرما پیش‌بینی می‌کند به دست نیاید، آنگاه n مرکب است (یعنی اول نیست)؛ در این صورت، m را نشانه‌ای از مرکب بودن n می‌خوانند. آزمونهای ذکر شده مبتنی بر انواعی از نشانه‌ها هستند که قدری پیچیده‌تر از این نوع‌اند. کارایی این آزمونها از قضایایی ناشی می‌شود که نشان می‌دهند اگر n مرکب باشد، آنگاه بیشتر اعداد صحیح بین 1 و $n-1$ ، نشانه این امر خواهند بود. یک جنبه جالب توجه این آزمونها این است که نشانه‌ای از اول بودن به دست نمی‌دهند، ولی این نقطه ضعف با تعریف نشانه اول بودن، به جای مرکب بودن، رفع شد، یعنی با نشان دادن اینکه اگر n عدد اولی باشد، بیشتر اعدادی که به تصادف انتخاب می‌شوند، نشانه‌ای از این واقعیت خواهند بود. الگوریتمهای تصادفی متعدد دیگری هم بر اساس فراوانی نشانه‌ها وجود دارد.

معلوم شده است که تکنیکهای تصادفی در ساختن الگوریتم برای کارهایی از قبیل مرتب‌کردن، جستجو، و هندسه محاسباتی نیز بسیار کارا هستند. مسأله تهیه فهرست تمام تلاقیهای تعدادی پاره‌خط واقع در صفحه، مثال ساده‌ای در این زمینه است. الگوریتم نتوی نسبتاً ساده‌ای وجود دارد که در آن، پاره‌خطها یکی پس از دیگری در نظر گرفته می‌شوند و تلاقیهای هر پاره‌خط جدید با همه خطهای قبلی گزارش می‌شود. اگر پاره‌خطها اتفاقاً به ترتیبی بسیار نامناسب در نظر گرفته شوند، زمان اجرای این الگوریتم فوق‌العاده زیاد خواهد بود. ولی می‌توان نشان داد که اگر پاره‌خطها به ترتیبی تصادفی پردازش شوند، به احتمال بسیار زیاد، الگوریتم خیلی سریع خواهد بود.

به علاوه، تصادفی‌سازی نقش فوق‌العاده مهمی در طراحی سیستمهای کامپیوتری توزیع شده دارند، یعنی سیستمهایی که در آنها تعداد زیادی کامپیوتر پراکنده و دور از هم که با رابطهای مناسبی بهم مربوط شده‌اند، به صورت اجزاء یک سیستم واحد با هم کار می‌کنند. در این سیستمها باید برای این مسأله که هر یک از کامپیوترها یا رابطها ممکن است از کار بیفتند و یا در یک زمان، سرعتهای متفاوتی داشته باشند، تدبیری اندیشیده شود، و اطمینان حاصل شود که زمان صرف‌شده برای ارتباط بین پردازنده‌ها نامعقول نیست. روشهای تصادفی به خصوص در محول کردن کارهای محاسباتی به هر یک از پردازنده‌ها و انتخاب مسیرهایی برای ارتباط که داده‌ها در طول آنها جریان یابند بسیار کارا هستند. می‌توان نشان داد که در وضعیتها و حالات بسیار متنوعی، احاطه تصادفی کارها به پردازنده‌ها و تخصیص تصادفی داده‌ها به پیمانه‌های حافظه، همراه با مسیریابی تصادفی پیامها، به احتمال زیاد نتیجه‌ای تقریباً بهینه به همراه خواهد داشت.

همه کاربردهایی که برای تصادفی‌سازی برشمرديم، وابسته به این فرض است که الگوریتمها، یا برنامه‌های کامپیوتری، به رشته‌ای از بیت‌های تصادفی مستقل دسترسی دارند. معمولاً کامپیوتر از اعداد شبه‌تصادفی استفاده می‌کند که به وسیله فرایندی تکرری و کاملاً تعینی از یک عدد اولیه، به نام بذریه نقطه، تولید می‌شوند. این مولدها معمولاً در معرض آزمونهای آماری خاصی قرار می‌گیرند تا مسلم شود که رشته‌هایی از اعداد که از آنها به دست می‌آیند، برخی

1. hierarchical solvers

به حساب آمده و کنار گذاشته شده بود. این ایده عبارت بود از اینکه در جستجو برای یک رأس بهینه، از درون چندوجهی عبور کنیم حال آنکه در روش سادگی حرکت فقط در روی چندوجهی انجام می شد. دشواری کار در مورد چنین رهیافت «نقطه درونی»، پیدا کردن جهت و مسافت صحیحی است که در هر مرحله باید طی شود تا رهیافت نمر بخش باشد. راه رفع این مشکل، استفاده از تبدیلهای تصویری و یک تابع پتانسیل الگوریتمی برای هدایت جستجو است؛ با این راه حل، زمان اجرای $O(n^{3/5}L^2)$ به دست می آید. ولی موضوع اصلی، بهبود نظری زمان اجرای روش بیضیواری نبود. نکته مهمتر این است که به نظر می رسد این الگوریتم (و چندین صورت مختلف آن)، وقتی با تکنیکهای مناسب مبتنی بر ماتریسهای تنگ اجرا شود، قادر به رقابت با روش سادگی است. علاوه بر آن، الگوریتم گفته شده در موارد بزرگ یاتبهگون، مزایای قابل ملاحظه ای از لحاظ زمان اجرا دارد. در واقع، این الگوریتم کاربردهای عملی مهمی در طرح شبکه های مخابراتی بزرگ و در حل مسائل بزرگ مقیاس مربوط به برنامه ریزی لجستیک و زمانبندی یافته است.

از زمانی که نخستین گزارشها از کاربردپذیری بالقوه این رهیافت انتشار یافته، پژوهشهای بسیار در زمینه روشهای نقطه درونی انجام شده است. کندوکاو گسترده ای در روابط بین این الگوریتم و الگوریتمهای پیشین صورت گرفته است. مثلاً این الگوریتم را می توان نوعی الگوریتم «تابع مانع» الگاریتمی» یا حتی کاربردی از روش نیوتن (در فضایی که به طرز مقتضی تبدیل شده باشد) دانست. این الگوریتم در حد، وقتی طول مرحله بینهایت کوچک باشد، میدانی برداری در درون چندوجهی پدید می آورد که همه مسیرهای حدیش به رأس بهینه می رسند. در این زمینه، الگوریتم را می توان به مثابه دنبال کردن تقریبی چنین مسیری، با برداشتن گامهای کوتاه در روی خطهای مماس، در نظر گرفت. این خود انواعی از الگوریتم را مطرح می سازد که در آنها، در طول خمهایی که معرف تقریبهای سری توانی بالامرتبه تر میدان برداری هستند، گام بر می داریم. انواع دیگر این الگوریتم، مبتنی بر طی کردن تقریبی مسیری خاص هستند که «مسیر مرکزی» نامیده می شود. انواع اخیر به زمانهای اجرای بهتر و بهتری برای بدترین حالت انجامیده است، و بهترین مورد در حال حاضر، با زمان اجرای $O(n^{1/5}L^2)$ ، مبتنی بر استفاده هوشمندانه از پیشرفتهای اخیر در زمینه ضرب ماتریسی سریع است.

الگوریتمها و بصیرتهای جدید همچنان با آهنگی سریع پدید می آیند و گرچه به نظر نمی رسد که این جریان به جوابهایی با زمان چندجمله ای برای مسائل NP تمام دشوارتر بینجامد، امید زیادی هست که رهیافت [مبتنی بر] نقطه درونی، حوزه مسائلی را که می شود پاسخهای مفیدی برایشان یافت، گسترش دهد.

۲۵. برنامه ریزی خطی تصادفی

مدلهای تعینی برای مسائل برنامه ریزی خطی و راه-دهای آنها، از ۴۰ سال پیش که برای نخستین بار مطرح شدند تاکنون پایه های پیشرفتهای خارق العاده کامپیوتر به پیش رفته اند. در زمان حاضر، حکومتها و مؤسسات صنعتی هر روز با حل مدلهایی سروکار دارند که در آنها عدم قطعیتی در کار نیست، یعنی به حل مسائلی برای طرح ریزی و زمانبندی هدفها می پردازند که تعینی خطی

یک کاربرد صنعتی بلاواسطه دیگر، در دینامیک مولکولی است که تکنیکی است برای مطالعه خواص مایعات (و سایر حالت های ماده) به وسیله شبیه سازی کامپیوتری. وقتی مواضع و سرعت های اولیه برای تعدادی از ذرات نماینده مشخص شده باشد، مسیرهای آنها با انتگرال گیری از قانون دوم حرکت نیوتن معلوم می شود. در شبیه سازیهای اولیه، فقط سیالات غیرقطبی مورد نظر بود. در این مورد، میزان محاسبه در هر مرحله زمانی متناسب با تعداد ذرات N است. در سیالات قطبی وضع کاملاً فرق می کند؛ یک جمله کوانتی به تابع پتانسیل افزوده می شود و همه برهمکنشهای دوه دو باید به حساب آیند، یعنی با یک محاسبه متعارف N جسمی سروکار داریم. روش FMM، شبیه سازی سیستمهای شیمیایی بسیار بزرگتری، نسبت به سیستمهای قبلی، را امکان پذیر می سازد. مثلاً بررسی یک مولکول پروتئین در آب، مستلزم در نظر گرفتن حرکت دهها هزار اتم در دهها هزار مرحله زمانی است. ادامه مطالعات مشروح در زمینه سینماتیک شیمیایی، احتمالاً در درازمدت فواید و دستاوردهای واقعی خواهد داشت.

۲۴. روشهای [مبتنی بر] نقطه درونی برای برنامه ریزی خطی

برای بسیاری از مسائل تخصیص منبع می توان مدلهایی از نوع مسئله «برنامه ریزی خطی» ساخت. در این نوع مسئله، تابعی خطی را روی رئوس یک چندوجهی n بعدی بهینه می سازند. الگوریتم سنتی سادگی برای حل این مسئله، که بر اساس حرکت کردن روی مرز چندوجهی است، از زمان کشف آن در چهل سال پیش تا کنون ارزش و تأثیر زیادی داشته است، ولی یک اشکال نظری مهم دارد و آن اینکه، مدت زمان اجرای آن در حالات نامساعد ممکن است به صورت تابعی نمایی از تعداد متغیرها افزایش یابد. الگوریتم بسیار مطلوبتر، دست کم از لحاظ نظری، الگوریتمی است که زمان اجرایش «در بدترین حالت»، محدود به چندجمله ای باشد.

اولین نمونه از چنین الگوریتم مطلوبی، روش بیضیواری بود که در ۱۹۷۶ کشف شد. زمان اجرای این الگوریتم، $O(n^4 L^2)$ بود که n تعداد متغیرها و L شاخصی از تعداد بیت های لازم برای توصیف ورودی است. خاصیت مطلوب دیگری که این الگوریتم دارد، این است که می توان آن را در مورد «برنامه ریزی محدب» کلیتر به کار برد؛ به علاوه، در این الگوریتم، توصیف کامل جسم محدبی که بهینه سازی روی آن انجام می گیرد لازم نیست و فقط یک زیرروال «جعبه سیاه» لازم است که در مورد هر نقطه مفروض، معین می کند که آن نقطه در چندوجهی هست یا نه، و اگر نبود، آبرصفحه ای مشخص می کند که نقطه را از چندوجهی جدا می سازد. پس روش بیضیواری، نسبت به روش سادگی، برای دسته بسیار بزرگتری از مسائل قابل کاربرد است و وجودش به پیدایش بسیاری الگوریتمهای بازمان چندجمله ای برای مسائلی که قبلاً حل نشده بودند انجامیده است. ولی پژوهشگران به زودی دریافتند که برای برنامه ریزی خطی، بهتر شدن کران زمان اجرای الگوریتم در بدترین حالت، ربطی به بهتر اجرا شدن آن در عمل ندارد.

بنابراین، به نظر می رسد که طرح الگوریتمهای برنامه ریزی با زمان چندجمله ای، از لحاظ نظری کار ظریفی است ولی فایده عملی ندارد. این طرز فکر در سال ۱۹۸۴ با ارائه نوع جدیدی از الگوریتم برنامه ریزی خطی با زمان چندجمله ای، تغییر کرد. این نوع الگوریتم مبتنی بر استفاده هوشمندانه ای از یک ایده قدیمی است، ایده ای که مدتها قبل از آن غیرعملی

ویژه مطابقت دارند. مقدار اطلاعاتی که متعاقباً به دست آمد، مسائل آماری جدیدی مطرح کرده است.

نقشه‌های فیزیکی، جایگاه‌های نسبی قطعاتی از DNA را که هویت مشخص و قابلیت دودمان‌سازی دارند، به دست می‌دهند. در دسترس بودن یک نقشه فیزیکی، تعیین کامل توالی یک زنجیره DNA را آسانتر می‌کند. اگر نقشه جایگاه‌های مجموعه‌ای از دودمانها^۱ - که هر یک طولی در حدود دهها هزار نوکلئوتید دارد - در اختیار باشد، چند آزمایشگاه می‌توانند با همکاری هم و به طور همزمان به تعیین توالی تک تک دودمانها بپردازند. می‌توان انتظار داشت که تحلیل آماری پارامترهای طرح، از قبیل تعداد آزمونهای برنده و اینکه چه آزمونهایی انتخاب شوند، فوایدی در فرایند نقشه‌برداری داشته باشد. فرایند تعیین فیزیکی محل دودمانها روی ژنوم، با شناخت پارامترهای طرح و منابع تغییریه که ذاتی فرایند است، بسیار آسان می‌شود.

به دست آوردن اطلاعات مربوط به توالی DNA فقط نخستین گام در زیست‌شناسی مولکولی جدید است. این توالی بعداً مورد بررسی گسترده قرار می‌گیرد تا ارتباط آن با اطلاعاتی که قبلاً درباره توالیهای DNA به دست آمده روشن شود. چون در جریان تکامل، ویژگیهای مهم زیست‌شناختی، از جمله خصوصیات توالی، محفوظ می‌ماند، این بررسیها ممکن است در فهم کارکرد بخشهای ویژه‌ای از توالی بسیار مهم باشد. اغلب استفاده از کامپیوتر برای اجرای الگوریتمهای پیچیده لازم می‌آید؛ تحلیل ریاضی الگوریتمها، هم برای تغییر نامیه و آموزنده‌ای از نتایج و هم برای کسب اطمینان از اینکه الگوریتمی که به صورت برنامه درآمده عملیاتش را با سرعت کافی انجام می‌دهد، لازم است.

توانایی ما در شناخت DNA و توالیهای پروتئین باعث پیشرفت در مطالعه تکامل مولکولی شده است. اکنون نمونه‌گیری از توالی مربوط به یک زن در جامعه‌ای مشخص، دارد امکان‌پذیر می‌شود. این امر، مسائل جدید زیادی را مطرح می‌کند. تکامل مولکولی در درازمدت و کوتاه‌مدت چگونه انجام می‌شود؟ با استفاده از مدل‌های فرایند تصادفی برای تکامل مولکولی، هم ترسیم درختهای تکامل و هم تعیین آهنگهای تکامل انجام‌پذیر است. برخی از مهم‌ترین کارهایی که در این زمینه انجام می‌شود، با تعیین ناحیه‌های کدگذاری پروتئین یا ژنهای موجود در DNA آغاز می‌گردد. تعیین جایگاه یک زن ۶۰۰ حرفی که به صورت قطعاتی کوچک در طول ۱۰۰۰۰ یا ۲۰۰۰۰ حرف DNA پراکنده است، کاری طاقت‌فرسا ولی ضروری است و به تحلیل ترکیبانی و آماری پیچیده‌ای نیاز دارد.

علم زیست‌شناسی مولکولی به سرعت در حال پیشرفت است. میزان داده‌های مربوط به توالی DNA و پروتئین در آینده، زمینه را برای رشد این مبحث از احاطه ریاضی مهیا می‌کند. ماهیت این علم چنان است که ارتباط نزدیک ریاضیدانان و زیست‌شناسان را ضروری می‌سازد. هر دو رشته مسلماً از چنین همکاری سود خواهند برد.

۲۷. آمار زیستی و اپیدمیولوژی

اپیدمیولوژی با توزیع بیماری در جوامع انسانی و عوامل مؤثر در آن سروکار دارد و موضوعات متنوعی را از قبیل تفاوت میزان بروز بیماری در نواحی جغرافیایی گوناگون، استانداردهای قرار گرفتن در معرض تابش در محل کار، و

و ریاضی هستند و گاه شامل هزارها متغیر همراه با یک تابع هدف خطی یا غیرخطی و هزاران قید به صورت نابرابری‌اند. پیش‌فرض این کار مثلاً این است که اطلاعات ما درباره انواع تکنولوژی که در آینده در دسترس خواهد بود، قطعی و دقیق است. در نتیجه، جوابهای حاصل از مدل‌های تعیینی ناقص است زیرا در این جوابها، پیشامدها و احتمالات آتی به درستی منظور نشده است. گرچه به آسانی می‌توان مدل‌های ریاضی را مجدداً صورت‌بندی کرد تا عدم قطعیت در آنها ملحوظ شود، ولی در این صورت اندازه دستگاههای ریاضی که باید حل شود در اکثر کاربردهای عملی فوق‌العاده زیاد خواهد بود. در حل مدل‌های تصادفی، قسمت مشکل کار عبارت بوده (و هست) از محاسبه امکانات.

بنابراین، علی‌رغم پیشرفتی که حاصل شده، هنوز دسته‌ای از مسائل «تصمیم» که اهمیت بسیار دارند، حل نشده‌اند، یعنی مسائل مربوط به یافتن جواب «بهینه» برای برنامه‌های ریاضی و خطی تصادفی. در اینجا «تصادفی» به معنی نبود قطعیت و یقین در مورد مسائلی از قبیل در دسترس بودن منابع، رقابت خارجی، یا تأثیرات اغتشاشات سیاسی و احتمالی است. چون موضوع این‌گونه مسائل، تخصیص بهینه منابع کمیاب بدون داشتن اطلاعات قطعی است، ملحوظ کردن عدم قطعیت در صورت‌بندی مسئله از اهمیت اساسی برخوردار است. اگر چنین مسائلی را بتوان در حالت کلی حل کرد، توانایی ما در طرح‌ریزی و زمان‌بندی برنامه‌های مربوط به شرایط پیچیده بسیار افزایش خواهد یافت.

در حال حاضر در ایالات متحده و کشورهای دیگر تلاش شدیدی برای حل دسته‌های خاصی از برنامه‌های تصادفی خطی و غیرخطی در جریان است. پیشرفتهای مهم اخیر در زمینه سخت‌افزار کامپیوتر، به خصوص پدید آمدن کامپیوترهای بزرگ چندپردازنده‌ای و پردازنده‌های موازی، محرک و انگیزه‌ای برای این تلاش بوده است. امید می‌رود که ترکیبی از سه نوع تکنیک - روشهای اصلاح شده برای تجزیه ریاضی مسائل بزرگ مقیاس به مسائل کوچکتر، تکنیکهای مهم نمونه‌گیری مونت کارلو، و استفاده از پردازنده‌های موازی - پیشرفتهای مهمی در آینده نسبتاً نزدیک به بار آورد.

۲۶. کاربردهای آمار در [مطالعه] ساختار DNA

هر زنجیره DNA را می‌توان به صورت رشته‌ای از بازها، A، C، G، T، نشان داد که اطلاعات مربوط به رشد و کارکرد یک موجود زنده را در بردارد. بنابراین، تلاش زیادی برای تعیین توالیهای DNA به عمل آمده است. در سال ۱۹۷۶ روشهایی برای تعیین سریع توالی پیدا شد و نتیجه آن، افزایش سرسام‌آور اطلاعات کمی ژنتیکی بود. امروز بیش از ۲۵ میلیون نوکلئوتید، متعلق به انواع بسیار گوناگونی از موجودات زنده، در قطعاتی با میانگین طول ۱۰۰۰، مشخص شده است. بهسازی روشهای تعیین توالی همچنان ادامه دارد، و اکتشافات وابسته به آن در زیست‌شناسی پایه، مبهوت‌کننده است.

دو نوع نقشه برای DNA رسم می‌شود: نقشه‌های ژنتیکی و نقشه‌های فیزیکی. هر دو نوع نقشه با استفاده از آزمونهای محدودکننده‌ای که توالی DNA را با الگوهای خاصی می‌بُرند و آن را به قطعاتی تجزیه می‌کنند، تهیه می‌شوند. طولهای این قطعه‌ها را می‌توان با قدری خطای آزمایشی ذاتی اندازه گرفت. در سال ۱۹۸۰ این فکر مطرح شد که تغییرات اندکی در توالی DNA ممکن است قطعاتی با طولهای متفاوت پدید آورد که می‌شود آنها را به عنوان «نشانهگر» به کار گرفت و تقریباً با جایگاههای ویژه روی کروموزومهای

1. clones

مثلاً در مطالعات «آینده‌نگر» که اندازه‌گیریهای مکرری روی یک شخص واحد در طی زمان انجام می‌شود و یا در مطالعات ژنتیکی الگوهای بیماری در خانواده‌ها، پیش می‌آید. روشهای بهتری نیز برای تعیین مقدار خطاهای اندازه‌گیری و برای تصحیح اثر این خطاها در ضعیف جلوه‌دادن ارتباط بین «در معرض قرار گرفتن» و بیماری، لازم است. پیشرفتهای جدید در محاسبات آماری و نظریه آماری تحلیلی شبه‌درست‌نمایی بر اساس معادلات برابردکننده تعمیم‌یافته، نویدبخش پیشرفتهای سریع در این زمینه است.

و بالاخره، ایدز که مهمترین بیماری همه‌گیر عصر ماست، مسائل مبرمی در برابر متخصصان آمارزیستی و اپیدمیولوژی قرار داده است. یکی از این مسأله‌ها، برابردکردن میزان آلودگی به ویروس ایدز در گذشته، حال، و آینده، به منظور تعیین تعداد مبتلایان به ایدز در آینده است. مسأله دیگر، به دست آوردن اطلاعات بیشتر درباره الگوهای انتقال ویروس ایدز در داخل گروههایی از افراد مستعد به ابتلا و نیز در بین این گروههاست تا برنامه‌ریزی برای سازماندهی بهینه منابع جامعه به منظور جلوگیری از آلودگی بیشتر امکان‌پذیر شود. جمع‌آوری داده‌های مربوط به ایدز به‌ندرت سرراست و قاعده‌مند است؛ در اغلب موارد بعضی از اطلاعات اساسی در دست نیست و بنابراین، نتایج بررسیها به‌ندرت تن به تحلیلی سرراست می‌دهند. از دیدگاه ریاضی، مسأله برابردکردن میزانهای آلودگی به ویروس ایدز با استفاده از داده‌های مربوط به میزانهای بروز ایدز و توزیع مدت زمان از آلودگی تا تشخیص می‌تواند نوع نام‌تعارفی از مسأله بازپیش^۱ به حساب آید. این مسأله با مسائلی ارتباط دارد که در مبحث تصویربردازی مطرح می‌شوند، مبحثی که پیشرفت سریعی در آن صورت گرفته است. ولی فقدان یا کمبود داده‌های کلیدی باعث می‌شود که این مسأله فوق‌العاده دشوارتر باشد.

• این مقاله، بخش دوم از ترجمه یکی از یوستهای گزارش معروف به گزارش دیوید II است که بخش اول آن را در شماره گذشته (شماره ۱ و ۲، سال ۵) خواندید. در اینجا از متخصصانی در زمینه‌های مختلف، به‌خصوص خانم دکتر الهی و آقایان دکتر صادقی، دکتر عمیدی، مهندس کیانی، و دکتر مرنزی، که برای ترجمه برخی از عبارتها و اصطلاحات این مقاله با آنها مشورت شده است، سپاسگزاری می‌شود.

ارزیابی تأثیر واکسیناسیون با استفاده از آزمایشهای میدانی تصادفی، در بر می‌گیرد.

در بیشتر تحقیقاتی که در اپیدمیولوژی بیماریهای مزمن انجام می‌گیرد، دو طرح مطالعاتی متمایز یعنی «گروهان»^۱ و «موردشاهد»^۲ی به‌کار می‌رود. در بررسیهای گروهانی، میزان قرارگرفتن در معرض [عاملی که مظنون به بیماری‌زایی است] و متغیرهای کمکی مربوطه برای گروه مشخصی از افراد سالم اندازه‌گیری می‌شود و افراد گروه در طی زمان تحت نظر قرار می‌گیرند تا معلوم شود که کدامیک از آنها به بیماری موردنظر مبتلا می‌شوند. یا بر اثر آن می‌میرند. روشها و مفاهیم تحلیل بقا، به‌خصوص مدل رگرسیون خطرهای متناسب، تأثیر زیادی در بررسی آماری داده‌های گروهانی داشته‌اند؛ و چارچوب ریاضی دقیقی برای فرمول‌بندی مفاهیم اساسی اپیدمیولوژی یعنی میزان بروز و مخاطره نسبی فراهم کرده‌اند. تغییر جدیدی از تکنیکهای قدیمی اپیدمیولوژی به زبان نظریه کلاسیک آمار ارائه شده، و راه برای ابداع روشهای انعطاف‌پذیرتر و نیرومندتر در تحلیل داده‌ها هموار شده است.

در بررسیهای «موردشاهد»ی با نمونه‌هایی از افراد بیمار و تندرست سروکار داریم که سابقه در معرض قرارگرفتن آنها معلوم است. ثابت شدن این موضوع که نسبت مربوط به در معرض قرارگرفتنی که از روی داده‌های بیمار-شاهدی محاسبه می‌شود، تقریباً برابر با نسبت میزانهای بروز بیماری بین در معرض قرارگرفتگان و قرارنگرفتگان است، اهمیت بسیار زیادی در تثبیت اعتبار عملی این طرح داشته است. اخیراً متخصصان آمارزیستی و اقتصادسنجی، مستقل از هم، روشهایی ابداع کرده‌اند برای تحلیل داده‌های بیمار-شاهدی و سایر داده‌ها در حالتی که نمونه بر اساس پیامدی که بیش از همه موردنظر است، انتخاب شود. لازم است این روشها گسترش بیشتری یابند تا در طرحهای طبقه‌بندی‌شده کلاسیک به‌کار آیند و نیز در مواردی که فقط اطلاعاتی جزئی در مورد تعداد زیادی از آزمودنیهای نمونه‌گیری‌شده داریم، مفید واقع شوند.

همچنین تحقیقات بیشتری در زمینه روشهای تحلیل داده‌های اپیدمیولوژیک در حالتی که برآمدها وابسته‌اند، لازم است؛ چنین حالتی

1. deconvolution

1. cohort 2. case - control