

شبه تصادفی بودن*

اودت گلدرایش*

ترجمه محمد قاسم وحیدی

نظریه دوم (رک. [۱۱]، [۱۲])، منسوب به سولومونوف^۱، کواموگوروف، و شاپیتین^۲، ریشه در نظریه محاسبه پذیری و به ویژه در مفهوم یک زبان جهانی (معادل با ماشین یا ابزار محاسبه جهانی) دارد. در این نظریه، پیچیدگی شی برحسب کوتاهترین برنامه (برای ماشین جهانی ثابتی) اندازه گرفته می شود که این شی را تولید می کند^۳. مانند نظریه شائین، پیچیدگی کواموگوروفی، کمی است، و شیهای تصادفی کامل به عنوان حالت های کرانگیزی مطرح می شوند. جالب توجه اینکه در این رهیافت، می توان گفت که یک شی واحد، و نه یک توزیع بر روی شیها، کاملاً تصادفی است. با این حال، رهیافت کواموگوروف ذاتاً احرا نشدنی است (یعنی پیچیدگی کواموگوروفی محاسبه ناپذیر است)، و بنابر تعریف، نمی توان رشته هایی با پیچیدگی کواموگوروفی بیشتر را از روی رشته های تصادفی کوتاه تولید کرد.

نظریه سوم، که آغازگران آن بلوم، گولدواسر، میکالی^۴، و یائو [۸، ۲، ۱۳] بودند، ریشه در نظریه پیچیدگی دارد و کانون توجه این مقاله است. این رهیافت صراحتاً به دنبال تدارک نظریه ای در باره تصادفی بودن کامل است که در عین حال امکان تولید کارآی رشته های کاملاً تصادفی از رشته های تصادفی کوتاه تر را فراهم آورد. آب این رهیافت، عبارت از این پیشنهاد است که شیها را برابر تلقی کنیم در صورتی که نتوان آنها را با هیچ شیوه کارآیی از هم تمیز داد. در نتیجه، توزیعی که نشود آن را به صورت کارآی از توزیع یکنواخت تمیز داد، تصادفی (یا بهتر است بگوییم «تصادفی برای کلیه مقاصد عملی») که آن را «شبه تصادفی» می نامیم) تلقی می شود. بنابراین، تصادفی بودن، خاصیت «ذاتی» اشیا (یا توزیعها) نیست بلکه وابسته به

این مقاله به شیهای متناهی می پردازد که به وسیله دنباله های دودویی متناهی به نام رشته کدگذاری می شوند. وقتی از توزیعها سخن به میان می آوریم، منظور ما توزیعهای احتمال گسسته اند با تکیه گاهی متناهی که مجموعه ای از رشته هاست. توجه خاصی به توزیع یکنواخت داریم که به ازای پارامتر طولی مانند n (که در طول بحث به صراحت یا به طور ضمنی مطرح می شود)، به هر رشته n بیتی مانند $x \in \{0, 1\}^n$: احتمالی برابر (یعنی، احتمال 2^{-n}) را تخصیص می دهد. وقتی به صورت محاوره ای صحبت از «رشته های کاملاً تصادفی» می کنیم، منظور ما رشته هایی هستند که مطابق با چنین توزیع یکنواختی برگزیده شده اند.

نیمه دوم این سده شاهد پیدایش و بسط سه نظریه درباره تصادفی بودن بوده است، مفهومی که اندیشمندان را در طول اعصار سردرگم کرده است. نخستین نظریه (رک. [۳]) که آغازگر آن شائین^۱ بود، ریشه در نظریه احتمال دارد و توجه آن معطوف به توزیعهایی است که کاملاً تصادفی نیستند. نظریه اطلاع شائین تصادفی بودن کامل را به عنوان حالتی کرانگیزی^۲ مشخص می کند که در آن گنجایش اطلاع ماکسیمم می شود (و هیچ زیادی^۳ ای در بین نیست)^۴. بنابراین، تصادفی بودن کامل وابسته به توزیعی بکتاست — همان توزیع یکنواخت. به ویژه، بنابر تعریف، نمی توان چنان رشته های تصادفی کامل را از رشته های تصادفی کوتاه تر تولید کرد.

1. Shannon 2. extreme 3. redundancy

۴. در حالت کلی، میزان اطلاع در توزیعی مانند D به صورت $-\sum_x D(x) \log_2 D(x)$ تعریف می شود. بنابراین، توزیع یکنواخت بر رشته هایی به طول n دارای اندازه اطلاع n است، و هر توزیع دیگر روی رشته های n بیتی اندازه اطلاع کمتری دارد. همچنین، برای تابعی مانند

$$f: \{0, 1\}^n \rightarrow \{0, 1\}^m$$

با $m < n$ ، توزیع حاصل از اعمال f بر رشته های n بیتی با یک توزیع یکنواخت، اندازه اطلاعی حداکثر n دارد که اکیداً کمتر از طول خروجی است.

1. Solomonov 2. Chaitin

۳. مثلاً رشته 1^n دارای پیچیدگی کواموگوروفی $O(1) + \log_2 n$ است (بر مبنای برنامه « n تا چاپ کن» که طولی نایبتر از n (مثلاً، در دستگاه دودویی)) دارد. در مقابل، استدلال شمارشی ساده ای نشان می دهد که اغلب رشته های n بیتی دارای پیچیدگی کواموگوروفی حداقل n اند.

مشاهده‌گر (و قابلیت‌های محاسباتی وی) است. برای توضیح این رهیافت، تجربه ذهنی زیر را در نظر می‌گیریم.

آلیس و باب به یکی از چهار روش زیر مشروط به خط بازی می‌کنند. در هر چهار روش، آلیس سکه‌ای را به هوا پرتاب می‌کند به طوری که تا ارتفاع زیادی بالا رود، و از باب خواسته می‌شود که برآمد آن را پیش از برخورد سکه با زمین حدس بزند. این روش‌ها از لحاظ اطلاعاتی که باب پیش از حدس زدن در اختیار دارد، با هم متفاوت‌اند. در شیق اول، باب باید حدس خود را پیش از پرتاب سکه اعلام کند. روشن است که در این حالت باب با احتمال $1/2$ برنده می‌شود. در شیق دوم، باب باید حدس خود را زمانی که سکه در هوا چرخ می‌خورد، اعلام کند. گرچه این برآمد در ساسی با حرکت سکه دقیق‌تر می‌شود، باب اطلاع دقیقی در مورد حرکت آن ندارد و بنابراین معتقدیم که در این حالت نیز باب با احتمال $1/2$ برنده می‌شود. شیق سوم، مشابه دومی است، بجز اینکه باب ابزار پیچیده‌ای در اختیار دارد که قادر است اطلاع دقیقی در مورد حرکت سکه و نیز شرایط محیط که بر برآمد تأثیر می‌گذارد، تهیه کند. با این حال، باب نمی‌تواند این اطلاع را به موقع پردازش کند که حدس خود را بهبود بخشد. در شیق چهارم، دستگاه ثابت باب مستقیماً به کامپیوتر قدرتمندی وصل است که برای حل معادلات حرکت برنامه‌ریزی شده و خروجی آن یک پیشگویی است. قابل درک است که در چنین حالتی باب می‌تواند حدس خود از برآمد را به میزان قابل ملاحظه‌ای بهبود بخشد.

نتیجه می‌گیریم که تصادفی بودن یک پیشامد به اطلاع و منابع محاسباتی که در اختیار داریم، بستگی دارد. بدین ترتیب، یک مفهوم طبیعی شبه‌تصادفی بودن، مطرح می‌شود: توزیع را شبه‌تصادفی گوئیم هرگاه هیچ شیوه کارایی نتواند آن را از توزیع یکنواخت تمیز دهد که در آن شیوه‌های کارا به الگوریتم‌های زمان چندجمله‌ای (احتمالاً) وابسته‌اند.

الگوریتمی را زمان چندجمله‌ای می‌نامند هرگاه یک چندجمله‌ای مانند p موجود باشد به طوری که برای هر ورودی x ، الگوریتم در مدت زمانی با کران $p(|x|)$ اجرا شود که در آن $|x|$ طول رشته x را نشان می‌دهد. بنابراین، زمان اجرای چنان الگوریتمی میانه روانه به صورت تابعی از ورودی آن رشد می‌کند. الگوریتم احتمالاتی، الگوریتمی است که می‌تواند گام‌هایی تصادفی را اختیار کند که در آن، بی‌آنکه از کلیت کاسته شود، می‌توان گفت هرگام تصادفی عبارت است از اتخاذ تصمیم در باره اینکه کدام یک از دو گام از پیش تعیین شده باید در مرحله بعدی انتخاب شود به طوری که هرگام ممکن با احتمال $1/2$ اختیار شود. این گزینه‌ها، «پرتاب‌های سکه درونی» الگوریتم نامیده می‌شوند.

تعریف مولدهای شبه‌تصادفی

به بیانی نادقیق، هر مولد شبه‌تصادفی برنامه‌ای (یا الگوریتمی) است که رشته‌های تصادفی کوتاه را کش می‌آورد و آنها را به دنباله‌های شبه‌تصادفی بلند تبدیل می‌کند. ما بر سه جنبه بنیادین در مفهوم مولد شبه‌تصادفی تأکید می‌کنیم: ۱. کارایی. مولد باید کارا باشد. از آنجا که کارایی محاسبه را به زمان چندجمله‌ای بودن آن وابسته می‌دانیم، فرض می‌کنیم که این مولد به کمک یک الگوریتم زمان چندجمله‌ای تعینی قابل اجرا باشد.

این الگوریتم رشته‌ای به نام D را به عنوان ورودی دریافت می‌کند. این بذر، مقداری کراندار از تصادفی بودن را در اختیار خود می‌گیرد که به وسیله ایزاری که «دنباله‌های شبه‌تصادفی را تولید می‌کند» به کار گرفته می‌شود. این فرموله‌بندی، چنان ایزاری را شیوه‌ای تعینی تلقی می‌کند که بر بذری تصادفی اعمال می‌شود.

۲. کش‌آوری. از مولد انتظار می‌رود که بذر ورودی را کش بیاورد و به دنباله خروجی طولانی‌تری تبدیل کند. به طور مشخص، مولد بذره‌ای به طول n بیت را به خروجی‌های طولانی‌تر $l(n)$ بیتی کش می‌آورد که در آن $l(n) > n$. تابع l اندازه کش‌آوری (یا تابع کش‌آوری) مولد نامیده می‌شود.

۳. شبه‌تصادفی بودن. خروجی مولد از دید هر مشاهده‌گر کارایی باید تصادفی به نظر آید. یعنی، هیچ شیوه کارایی نباید قادر به تمیز دادن خروجی یک مولد بذری تصادفی از یک دنباله واقعاً تصادفی با همان طول باشد. جمله اخیر به مفهومی عام از «تمایزناپذیری [از نظر محاسباتی]» اشاره دارد که قالب کل این رهیافت است.

برای روشن شدن مطالب بالا، طرح زیر را برای یک مولد شبه‌تصادفی در نظر بگیرید. بذر آن متشکل از جفتی از اعداد صحیح 32 بیتی مانند x و N است، و خروجی 10000 بیتی آن با مجذور کردن مکرر عدد جاری x به پیمانه N و بیرون آوردن بیت با کمترین معنی در هر نتیجه در مراحل بینابینی به دست می‌آید. (یعنی، فرض کنید (به پیمانه N) $x_{i+1} = x_i^2$ ، برای 10000 ، $i = 1, \dots, 10000$ ، و خروجیها $b_1, b_2, \dots, b_{10000}$ هستند، که در آن $x_i \equiv x_i \pmod{N}$ و b_i بند با کمترین معنای x_i است). این فرایند را می‌توان به بذرهایی به طول n (در اینجا ما از $n = 64$ استفاده کرده‌ایم) و خروجی‌هایی به طول $l(n)$ (در اینجا $l(n) = 10000$) تعمیم داد. چنان فرایندی مطمئناً فقره‌های (۱) و (۲)ی بالا را برآورده می‌کند، در حالی که این پرسش که آیا فقره (۳) برقرار است یا خیر، قابل بحث است (پس از آنکه تعریفی دقیق برای آن ارائه شد). با اشاره پیشاپیش به بخشی از بحث‌های آتی، متذکر می‌شویم که، تحت این فرض که تجزیه اعداد صحیح بزرگ، دشوار است، صورت اندک تغییر یافته‌ای از فرایند بالا درواقع یک مولد شبه‌تصادفی است.

تمایزناپذیری [از نظر محاسباتی]

از لحاظ شهودی، دو شیء از نظر محاسباتی تمایزناپذیرند هرگاه هیچ شیوه کارایی نتواند آنها را از هم تمیز دهد. مطابق آنچه در نظریه پیچیدگی معمول است، صورتبندی ظریف این موضوع مستلزم تحلیلی مجانبی است (یا بهتر است بگوئیم بررسی زمان اجرای الگوریتم‌ها به عنوان تابعی از طول ورودی آنها). بنابراین، شیئهای مورد بحث، دنباله‌هایی نامتناهی از توزیعها هستند که در آنها هر توزیع تکیه‌گاهی متناهی دارد. چنان دنباله‌ای یک جرگه^۲ توزیعی نامیده می‌شود. ما جرگه‌های توزیعی به شکل $\{D_n\}_{n \in \mathbb{N}}$ را در نظر می‌گیریم که در آن برای تابعی مانند $l: \mathbb{N} \rightarrow \mathbb{N}$ ، تکیه‌گاه هر D_n زیرمجموعه‌ای از $\{0, 1\}^{l(n)}$ است. به علاوه نوعاً l یک چندجمله‌ای مثبت خواهد بود. برای

۱. بررسی مجانبی (یا تابعی) در این رهیافت جنبه اساسی ندارد. می‌توان کل رهیافت را برحسب ورودی‌های با طولهای ثابت و مفهومی مناسب از پیچیدگی الگوریتم‌ها به پیش برد.

ولی چنان شیوه‌ای پرزحمت‌تر است.

بنابراین، مولدهای شبیه‌تصادفی برنامه‌های تعیینی کاراً (یعنی زمان چندجمله‌ای) هستند که بذره‌های کوتاه به تصادف انتخاب شده را منبسط کرده به صورت دنباله‌های بی‌نتی شبیه‌تصادفی بلندتر درمی‌آورند به‌طوری که دنباله‌های اخیر به صورت تمایزناپذیر محاسباتی از دنباله‌های واقعاً تصادفی، به وسیله الگوریتم‌های کاراً (یعنی، زمان چندجمله‌ای) تعریف می‌شوند. نتیجه می‌گیریم که هر الگوریتم تصادفی شده کاراً نحوه اجرای خود را، در صورتی که به‌جای پرتابهای سکه درونی دنباله‌هایی قرار دهیم که به‌وسیله یک مولد شبیه‌تصادفی تولید شده‌اند، حفظ می‌کند. یعنی

ساختمان ۳. (کاربرد نوعی مولدهای شبیه‌تصادفی). فرض کنید که A یک الگوریتم زمان چندجمله‌ای احتمالاتی باشد، و فرض کنید که $\rho(n)$ معرف کرانی بالا برای پیچیدگی تصادفی بودن آن (یعنی تعداد سکه‌هایی که A روی ورودی‌های n بیتی پرتاب می‌کند) باشد. همچنین تصور کنید که $A(x, r)$ معرف خروجی A روی ورودی x و دنباله پرتاب سکه $r \in \{0, 1\}^{\rho(|x|)}$ باشد. حال فرض کنید که G یک مولد شبیه‌تصادفی با تابع کش‌آورنده $l: \mathbb{N} \rightarrow \mathbb{N}$ باشد. در این صورت AG یک الگوریتم تصادفی شده است که روی ورودی x به صورت زیر عمل می‌کند. عدد $k = k(|x|)$ را به‌عنوان کوچکترین عدد صحیح به‌طوری که $l(k) \geq \rho(|x|)$ قرار می‌دهد، به‌طور یکنواخت $s \in \{0, 1\}^k$ را انتخاب می‌کند، و خروجی $A(x, r)$ را می‌دهد که در آن r پیشوند به طول $\rho(|x|)$ بیت $G(s)$ است.

می‌توان ثابت کرد که پیدا کردن x ‌های بلندی که روی آنها رفتار ورودی-خروجی AG به‌طور محسوس متفاوت از آن A باشد، نشدنی است، گرچه AG ممکن است از تعداد پرتابهای به مراتب کمتری در مقایسه با A استفاده کند. این موضوع در قضیه زیر مدون شده است که در آن F معرف الگوریتمی است در پی یافتن x به گونه‌ای که $A(x)$ و $AG(x)$ (به کمک الگوریتمی مانند D) قابل تمیز باشند.

قضیه [فرعی] ۴. فرض کنید A و G به‌صورتی باشند که در بالا گفته شد، برای هر الگوریتم D ، فرض کنید $\Delta_{A,D}(x)$ معرف اختلاف در رفتار A و AG ، به قضاوت D ، روی ورودی x باشد یعنی $\Delta_{A,D}(x)$ به صورت زیر تعریف شود

$$\left| \Pr_{r \sim U_{\rho(n)}}[D(x, A(x, r)) = 1] - \Pr_{s \sim U_{l(n)}}[D(x, AG(x, s)) = 1] \right|$$

که در آن احتمالات روی همه U_m ‌ها و نیز روی کلمه پرتابهای سکه D گرفته شده‌اند، در این صورت برای هر جفت از الگوریتم‌های زمان چندجمله‌ای احتمالاتی F و D ، هر چندجمله‌ای مثبت p ، و همه n ‌های به قدر کافی بزرگ داریم

$$\Pr \left[\Delta_{A,D}(F(1^n)) > \frac{1}{p(n)} \right] < \frac{1}{p(n)}$$

که در آن احتمال روی همه پرتابهای سکه F گرفته می‌شود.

این قضیه با نشان دادن این موضوع ثابت می‌شود که یک سمتایی (A, F, D) را که حکم را نقض می‌کند می‌توان به الگوریتمی مانند D'

چنان D_n ‌ای، فرایند انتخاب e مطابق توزیع D_n را با $D_n \sim e$ نشان می‌دهیم بنابراین، برای محمولی مانند P ، احتمال آن را که $P(e)$ برقرار باشد وقتی که e مطابق با D_n توزیع (یا انتخاب) شده باشد با $\Pr_{e \sim D_n}[P(e)]$ نشان می‌دهیم.

تعریف ۱. تمایزناپذیری محاسباتی ([۱۳، ۸]). دو جرگه احتمالات $\{X_n\}_{n \in \mathbb{N}}$ و $\{Y_n\}_{n \in \mathbb{N}}$ نامزد محاسباتی نامیده می‌شوند هرگاه برای هر الگوریتم زمان چندجمله‌ای احتمالاتی مانند A ، هر چندجمله‌ای مثبت مانند p ، و همه n ‌های به قدر کافی بزرگ،

$$\left| \Pr_{x \sim X_n}[A(x) = 1] - \Pr_{y \sim Y_n}[A(y) = 1] \right| < \frac{1}{p(n)}$$

احتمال روی X_n (به ترتیب، Y_n) و نیز روی پرتابهای سکه الگوریتم A اختیار می‌شود.

در اینجا چند تذکر لازم است. نخست اینکه اجازه داده‌ایم الگوریتم A ، که متمایزکننده نام دارد، احتمالاتی باشد. این کار، شرط را صرفاً قویتر می‌کند، و برای چندین جنبه مهم رهیافت ما اساسی به نظر می‌رسد. دوم، ما پیشامدهایی را که کران بالایی احتمال رخ دادن آنها عکس یک چندجمله‌ای است، صرف‌نظرکردنی تلقی می‌کنیم. این امر با مفهوم کارایی (یعنی محاسبات زمان چندجمله‌ای) به‌خوبی جفت‌وجور می‌شود: پیشامدی که با احتمالی صرف‌نظرکردنی (به‌عنوان تابعی از پارامتری مانند n) رخ می‌دهد در صورتی هم که آزمایش به دفعاتی برابر یک چندجمله‌ای برحسب n تکرار شود، با احتمالی صرف‌نظرکردنی رخ خواهد داد. سوم، وقتی در تعریف بالا به A اجازه دهیم که تابعی دلخواه (به جای الگوریتم زمان چندجمله‌ای احتمالاتی) باشد، یک مفهوم تمایزناپذیری آماری را به‌دست خواهیم آورد. مفهوم اخیر معادل با این شرط است که فاصله تغییرات بین X_n و Y_n (یعنی $\sum_i |X_n(z) - Y_n(z)|$) نسبت به n ناچیز است.

با‌آر می‌شویم که تمایزناپذیری محاسباتی مفهومی بسیار آزادتر از تمایزناپذیری آماری است (رگ. [۱۳]). حالتی مهم، آن است که توزیعها به توسط یک مولد شبیه‌تصادفی (به صورتی که ذیلاً تعریف می‌شود) تولید شوند: چنان توزیعهایی از نظر محاسباتی از توزیع یکنواخت تمایزناپذیرند، اما از نظر آماری از توزیع یکنواخت غیرقابل تمایز نیستند.

تعریف ۲. (مولدهای شبیه‌تصادفی ([۱۳، ۲])). یک الگوریتم زمان چندجمله‌ای تعیینی مانند G مولد شبیه‌تصادفی نامیده می‌شود هرگاه تابعی کش‌آورنده مانند $l: \mathbb{N} \rightarrow \mathbb{N}$ موجود باشد به‌طوری که دو جرگه احتمال زیر، که با $\{G_n\}_{n \in \mathbb{N}}$ و $\{R_n\}_{n \in \mathbb{N}}$ نشان داده می‌شوند، تمایزناپذیر محاسباتی باشند: ۱. توزیع G_n به عنوان خروجی G روی بذری که به‌طور یکنواخت انتخاب شده، در $\{0, 1\}^n$ تعریف می‌شود.

۲. توزیع R_n به‌عنوان توزیع یکنواخت بر $\{0, 1\}^{l(n)}$ تعریف می‌شود. یعنی، با نشان دادن توزیع یکنواخت روی U_m با $\{0, 1\}^m$ قید می‌کنیم که برای هر الگوریتم زمان چندجمله‌ای احتمالاتی مانند A ، هر چندجمله‌ای مثبت مانند p ، و همه n ‌های به‌قدر کافی بزرگ،

$$\left| \Pr_{s \sim U_n}[A(G(s)) = 1] - \Pr_{r \sim U_{l(n)}}[A(r) = 1] \right| < \frac{1}{p(n)}$$

احتمالاتی مانند A' ، هر چند جمله‌ای مثبت $p(\cdot)$ ، و همه n های به قدر کافی بزرگ،

$$\Pr_{x \sim U_n} [A'(f(x)) \in f^{-1}(f(x))] < \frac{1}{p(n)}$$

که در آن U_n توزیع یکنواخت بر $\{0, 1\}^n$ است

در مورد تابعهای یکطرفه، کاندیداهای رایج بر مهارناپذیری^۱ حدسی تجزیه اعداد صحیح به عوامل اول، مسألهٔ الگاریتمی گسسته، و کدگذاری کدهای تصادفی خطی مبتنی هستند. ناشدنی بودن وارونسازی f ، مفهومی ضعیف از پیشگویی ناپذیری را عاید می‌کند: فرض کنید $b_i(x)$ معرف بیت i ام x باشد. در این صورت، برای هر الگوریتم زمان چندجمله‌ای مانند

$$\Pr_{i,x} [A(i, f(x)) \neq b_i(x)] > 1/2^n$$

که در آن احتمال به‌طور یکنواخت روی $\{1, 2, \dots, n\}$ و $i \in \{0, 1\}^n$ گرفته شده است. مفهومی قویتر (و در واقع، قویترین مفهوم ممکن) از پیشگویی ناپذیری همانا محمول «هسته‌ای» است. به بیان غیر دقیق، یک محمول قابل محاسبهٔ زمان چندجمله‌ای چون b ، هستهٔ [اصلی] تابعی مانند f نامیده می‌شود هرگاه هر الگوریتم کارا، با معلوم بودن $f(x)$ ، بتواند $b(x)$ را با احتمال موفقی که تنها به‌طور صرف نظرکردنی بیشتر از نیم باشد، حدس بزند.

تعریف ۷. (محمول هسته‌ای [۲]). یک محمول قابل محاسبهٔ زمان چندجمله‌ای مانند $b: \{0, 1\}^* \rightarrow \{0, 1\}$ هستهٔ [اصلی] تابعی مانند f نامیده می‌شود هرگاه برای هر الگوریتم زمان چندجمله‌ای احتمالاتی مانند A' ، هر چند جمله‌ای مثبت $p(\cdot)$ ، و همه n های به قدر کافی بزرگ،

$$\Pr_{x \sim U_n} [A'(f(x)) = b(x)] < \frac{1}{2} + \frac{1}{p(n)}$$

آشکار است که هرگاه b هستهٔ تابع یک‌به‌یک محاسبه‌پذیر زمان چندجمله‌ای f باشد، آنگاه f باید یکطرفه باشد.^۲ نتیجه آنکه هر تابع یکطرفه را می‌توان اندکی بهبود بخشید به‌طوری که محمولی هسته‌ای داشته باشد.

قضیهٔ ۸. (یک هستهٔ عام [۷]). فرض کنید که f تابع یکطرفهٔ دلخواهی باشد و تصور کنید g به صورت $g(x, r) \stackrel{\text{تعریف}}{=} (f(x), r)$ تعریف شود که در آن $|x| = |r|$. حال فرض کنید $b(x, r)$ معرف ضرب داخلی در دایره‌های دودویی x و r ده دهمانهٔ ۲ باشد، در این صورت محمول b دلتا هستهٔ تابع g است.

برهانی از این قضیه را می‌توان در [۵، پیوست C ۲۰] یافت. سرانجام به ساختمان مولدهای شبه‌تصادفی می‌رسیم.

قضیهٔ [فرعی] ۹. (شیوهٔ ساده‌ای برای ساخت مولدهای شبه‌تصادفی).

1. intractability

۲. تابعی که یک‌به‌یک نیستند، ممکن است محمولهای هسته‌ای با ماهیت نظریهٔ طلاع داشته باشند، اما این تابعها در اینجا به درد ما نمی‌خورند. مثلاً تابعی به شکل

$$f(\sigma, x) = {}^\circ f'(x)$$

(برای $\sigma \in \{0, 1\}$) دارای یک محمول هسته‌ای «نظریهٔ اطلاع» مانند $b(\sigma, x) = \sigma$ هستند.

برگرداند که خروجی G را از توزیع یکنواخت متمایز می‌کند، و این برخلاف فرض است. با استدلالهای مشابهی می‌توان ثابت کرد که فرایند تصادفی شدهٔ کارا (خواه مانند بالا الگوریتمی باشد و خواه محاسبه‌ای چندبخشی)، اگر شبه‌تصادفی بودن (با تعریف بالا) جانشین تصادفی بودن واقعی شود، رفتارش را حفظ می‌کند. بنابراین، با در دست داشتن مولدهای شبه‌تصادفی با تابع کش‌آوری بزرگ، می‌توان به‌طور قابل ملاحظه‌ای پیچیدگی تصادفی را در هر کاربرد کارا کاهش داد.

تقویت تابع کش‌آوری

از مولدهای شبه‌تصادفی، به صورتی که در بالا تعریف شده‌اند، تنها انتظار می‌رود که ورودی خود را یک بیت کش بیاورند؛ مثلاً کش‌آوردن ورودیهایی به طول n بیت به خروجیهایی به طول $(n+1)$ بیت کفایت می‌کند. روشن است که مولدهایی با چنین تابعهای کش‌آوری میانه‌روی در عمل کمتر به درد می‌خورند. در مقابل، می‌خواهیم مولدهای شبه‌تصادفی با تابعهای کش‌آوری به دلخواه طویل داشته باشیم. نابار شرط کارایی، تابع کش‌آوری می‌تواند حداکثر چندجمله‌ای باشد. نتیجه آن می‌شود که مولدهای شبه‌تصادفی با کوچکترین تابعهای کش‌آوری ممکن را می‌توان برای ساختن مولدهای شبه‌تصادفی با هر تابع کش‌آوری چندجمله‌ای موردنظر به کار برد. (بنابراین، وقتی بحث از وجود مولدهای شبه‌تصادفی در میان است، می‌توانیم تابعهای کش‌آوری مشخص را نادیده بگیریم.)

قضیهٔ ۵. (۵، بخش ۲.۳.۳). فرض کنید G دلتا مولد شبه‌تصادفی با تابع کش‌آوری $l(n) = n+1$ باشد و تصور کنید که l' هر تابع کش‌آوری دلخواه باشد به‌طوری که $l'(n)$ در زمان چندجمله‌ای بر حسب n قابل محاسبه باشد، همچنین فرض کنید $G_1(x)$ معرف پیشوند به طول $|x|$ بیت $G(x)$ باشد، و $G_2(x)$ معرف آخرین بیت $G(x)$ باشد (یعنی $G(x) = G_1(x)G_2(x)$). در این صورت

$$G'(S) \stackrel{\text{تعریف}}{=} \sigma_1 \sigma_2 \dots \sigma_{l'(|S|)}$$

که در آن $x_i = G_1(x_{i-1})$ ، $\sigma_i = G_2(x_{i-1})$ ، $x_0 = s$ به ازای

$$i = 1, \dots, l'(|s|)$$

دلتا مولد شبه‌تصادفی با تابع کش‌آوری l' است.

چگونگی ساختن مولدهای شبه‌تصادفی

روشهای شناخته شده برای ساختن این مولدها، مشکل محاسباتی را در قالب تابعهای یکطرفه (که در زیر تعریف می‌شود) به مولدهای شبه‌تصادفی بودن تبدیل می‌کنند. به بیان غیر دقیق، یک تابع محاسبه‌پذیر زمان چندجمله‌ای، یکطرفه نامیده می‌شود هرگاه هر الگوریتم کارا بتواند آن را تنها با احتمال موفقی صرف نظرکردنی وارون کند. برای سهولت، تنها تابعهای یکطرفهٔ طول نگهدار را در نظر می‌گیریم.

تعریف ۶. (تابع یکطرفه). هر تابع یکطرفه مانند f ، یک تابع محاسبه‌پذیر زمان چندجمله‌ای است به‌طوری که برای هر الگوریتم زمان چندجمله‌ای

فرض کنید b یک محدول هسته‌ای تابع محاسبه پذیر زمان چندجمله‌ای f باشد، در این صورت

$$G(s) \stackrel{\text{تعریف}}{=} f(s)b(s)$$

(یعنی $f(s)$ و $b(s)$ یک مولد شبه تصادفی است).

به یک معنی، نکته اصلی در برهان قضیه بالا عبارت از اثبات آن است که پیشگویی ناپذیری خروجی G (که بنا به تعریف واضح است) مستلزم شبه تصادفی بودن آن است. این حقیقت که پیشگویی ناپذیری (بیت بعدی) و شبه تصادفی بودن در حالت کلی معادل اند، در توصیف دیگری از موضوع که در زیر می‌آید، به صراحت ثابت می‌شود.

توصیف دیگر

شیوه‌ای که برای ساخت مولدهای شبه تصادفی، از طریق قضیه‌های ۵ و ۹، عرضه کردیم، متفاوت اما مشابه با طرز ساخت مولدهای شبه تصادفی است که به وسیله بلوم و میکالی [۲] مطرح شده است: با مفروض بودن تابع کش‌آوری $l: \mathbb{N} \rightarrow \mathbb{N}$ و یک تابع یکطرفه یک‌به‌یک با هسته b ، تعریف می‌کنیم

$$G(s) \stackrel{\text{تعریف}}{=} b(x_0)b(x_1) \cdots b(x_{l(|s|)-1})$$

که در آن $x_0 = s$ و $x_i = f(x_{i-1})$ به ازای $i = 1, \dots, l(|s|) - 1$. یک مصداق ملموس، بر مبنای این فرض که تجزیه اعداد صحیح بزرگ کاری دشوار است، در شکل ۱ شرح داده شده است. شبه تصادفی بودن G در دو گام با استفاده از مفهوم پیشگویی ناپذیری (بیت بعدی) برقرار شده است. جرگه‌ای چون $\{Z_n\}_{n \in \mathbb{N}}$ پیشگویی ناپذیر نامیده می‌شود هرگاه هیچ ماشین زمان چندجمله‌ای احتمالاتی که پیشوندی از Z_n را به دست می‌آورد، قادر به پیشگویی بیت بعدی آن با احتمالی که به طور صرف نظرناکردنی بزرگتر از $1/2$ است، نباشد.

گام ۱. ابتدا ثابت می‌کنیم که جرگه $\{G(U_n)\}_{n \in \mathbb{N}}$ که در آن U_n توزیع یکنواخت روی $\{0, 1\}^n$ است، پیشگویی ناپذیر (از لحاظ بیت بعدی) است (از راست به چپ [۲]). به بیانی غیردقیق، اگر بتوان $b(x_i)$ را از روی $b(x_{i+1}) \cdots b(x_{l(|s|)-1})$ پیشگویی کرد آنگاه می‌توان $b(x_i)$ را با مفروض بودن $f(x_i)$ پیشگویی کرد (یعنی، با محاسبه $x_{i+1}, \dots, x_{l(|s|)-1}$ و به این ترتیب، به دست آوردن $b(x_{i+1}) \cdots b(x_{l(|s|)-1})$). اما این نتیجه با فرض هسته در تناقض است.

گام ۲. سپس از نتیجه‌گیری یائو استفاده می‌کنیم که بنابر آن یک جرگه شبه تصادفی است اگر و تنها اگر پیشگویی ناپذیر (از لحاظ بیت بعدی) باشد (رک. ۴، بخش ۳.۳.۴).

روشن است که اگر بتوان بیت بعدی را در یک جرگه پیشگویی کرد، آنگاه می‌توان این جرگه را از جرگه یکنواخت تمیز داد (که صرف نظر از توان محاسبه، پیشگویی ناپذیر است). ولی ما طرف دیگر را که کمتر بدیهی است لازم داریم. می‌توان نشان داد که پیشگویی ناپذیری (بیت بعدی) مستلزم تمایزناپذیری از جرگه یکنواخت است. به طور مشخص، توزیعهای «دورگه» ای را در نظر بگیرید که در آن دورگه نام، 2 بیت نخست را از جرگه مورد بحث و بقیه

فرض می‌کنیم که تجزیه اعداد صحیحی که حاصلضرب دو عدد اول بزرگ اند (هر یک هم‌نشت با ۳ به پیمانه ۴) ناشدنی باشد. تحت این فرض، مجذور کردن به پیمانه چنان اعداد صحیحی یک تابع یکطرفه خواهد بود. به علاوه، مجذور کردن به پیمانه چنان N ای روی مانده‌های مربعی به پیمانه N یک‌به‌یک است، و کم‌معناترین بیت شناسه، یک هسته متناظر است [۱].

مولد شبه تصادفی زیر از یک الگوریتم زمان چندجمله‌ای استفاده می‌کند که وقتی $4m$ بیت به آن خورانده شود، عددی m بیتی تولید می‌کند به طوری که وقتی ورودیهای $4m$ بیتی به صورت یکنواخت توزیع شده باشند، خروجی اساساً یک عدد اول تصادفی m بیتی (هم‌نشت با ۳ به پیمانه ۴) است.

ورودی: یک بذر n بیتی $s = abc$ که در آن $|a| = |b| = 4n/10$. گامهای آغازی:

۱. با استفاده از a ، عدد اول $n/10$ بیتی (به پیمانه ۴) $p \equiv 3$ را تولید کنید.

۲. به همین نحو، از b استفاده کرده، (به پیمانه ۴) $q \equiv 3$ تولید کنید.

۳. p را در q ضرب کنید و N را به دست آورید.

۴. فرض کنید (به پیمانه N) $x_0 \leftarrow c^2$.

تکرارها: برای $i = 0, \dots, l(n) - 1$

۵. فرض کنید h_i کم‌معناترین بیت x_i باشد

۶. فرض کنید (به پیمانه N) $x_{i+1} \leftarrow x_i^2$.

خروجی: $b_0, b_1, \dots, b_{l(n)-1}$.

شکل ۱ یک مولد شبه تصادفی بر مبنای اجرا نشدنی بودن تجزیه به عوامل.

را از توزیع یکنواخت اختیار می‌کند. بنابراین، تمیز دادن دورگه‌های کرانگینی مستلزم تمیز دادن برخی دورگه‌های هم‌سایه است که به نوبه خود مستلزم پیشگویی ناپذیری بیت بعدی (در جرگه مورد بحث) است.

شرطی عام برای وجود مولدهای شبه تصادفی

به خاطر آورید که با مفروض بودن هر تابع یک‌به‌یک یکطرفه، می‌توان به آسانی یک مولد شبه تصادفی ساخت. درواقع، می‌توان شرط یک‌به‌یک بودن را کنار گذاشت، اما شیوه فعلاً معلوم برای ساخت — در حالت کلی — کاملاً پیچیده است. با این حال، داریم:

قضیه ۱۰. (در باره وجود مولدهای شبه تصادفی [۹]). مولدهای شبه تصادفی موجودند اگر و تنها اگر تابعهای یکطرفه موجود باشند.

برای نشان دادن اینکه وجود مولدهای شبه تصادفی مستلزم وجود تابعهای یکطرفه است، مولدی شبه تصادفی مانند G را با تابع کش‌آوری $l(n) = 2n$ در نظر بگیرید. برای هر $x, y \in \{0, 1\}^n$ ، تعریف کنید $G(x)$ $\stackrel{\text{تعریف}}{=} f(x, y)$

روی کلیه تابعهایی است که $\{0, 1\}^{l_D(n)}$ را به $\{0, 1\}^{l_R(n)}$ می نگارند.

برای سهولت فرض کنید که $l_D(n) = n$ و $l_R(n) = 1$. در این صورت تابعی که به طور یکنواخت از بین 2^n تابع (از یک جرگه شبه تصادفی) انتخاب شده است، در زمان چندجمله‌ای برحسب n ، یک رفتار ورودی-خروجی تمایزناپذیر از رفتار تابعی که به تصادف از بین کلیه 2^n تابع بولی انتخاب شده است، نشان می دهد. این را با توزیعی روی 2^n دنباله که با یک مولد شبه تصادفی به کار رفته در مورد یک بذر تصادفی n بیتی تولید شده است، مقایسه کنید. بذر اخیر غیر قابل تمایز از دنباله‌ای است که به طور یکنواخت از بین همه چندجمله‌ای برحسب 2^n دنباله انتخاب شده است. در عین حال، تابعهای شبه تصادفی را می توان از روی هر مولد شبه تصادفی ساخت.

قضیه ۱۲. (چگونگی ساختن تابعهای شبه تصادفی [۶]). فرض کنید که G یک مولد شبه تصادفی با تابع کش آدوی $l(n) = 2n$ باشد، همچنین فرض کنید $G_0(s)$ (به ترتیب، $G_1(s)$) معرف نخستین (به ترتیب، آخرین) $|s|$ بیت در $G(s)$ باشد، و

$$G_{\pi[s] \dots \sigma_1 \sigma_0}(s) \stackrel{\text{تعریف}}{=} G_{\sigma_1 \dots \sigma_0}(G_{\sigma_1}(G_{\sigma_0}(s)) \dots)$$

در این صورت جرگه تابعی $\{f_s : \{0, 1\}^{|s|} \rightarrow \{0, 1\}^{|s|} \mid s \in \{0, 1\}^*\}$ که در آن $G_x(s) \stackrel{\text{تعریف}}{=} f_s(x)$ ، شبه تصادفی با پارامترهای طول، $l_D(n) = l_R(n) = n$ است.

شیوه ساخت بالا را می توان به آسانی در مورد هر نوع پارامترهای طول (که کرانی به صورت چندجمله‌ای داشته باشند):

$$l_D, l_R : \mathbb{N} \rightarrow \mathbb{N}$$

به کار برد.

متذکر می شویم که تابعهای شبه تصادفی برای استخراج نتایج منفی در نظریه یادگیری محاسباتی و نظریه پیچیدگی به کار رفته اند.

کاربردپذیری مولدهای شبه تصادفی

شبه تصادفی بودن نقشی با اهمیت روزافزون در محاسبات ایفا می کند: اغلب در طرح الگوریتمهای دنباله‌ای، موازی، و توزیع شده به کار گرفته می شود و البته نقشی مرکزی در رمزنگاری دارد. هر چند طرح چنین الگوریتمهایی با استفاده آزادانه از تصادفی بودن، آسان است اما در عین حال مطلوب آن است که استفاده از تصادفی بودن در اجراهای واقعی به حداقل برسد زیرا تولید کردن بیتهای کاملاً تصادفی از طریق سخت افزار ویژه بسیار گران است. بنابراین، مولدهای شبه تصادفی (به صورتی که در بالا تعریف شدند) جزء اصلی در یک «جعبه ابزار الگوریتمی» است: این مولدها برنامه های مترجم خودکاری برای برگرداندن برنامه هایی که آزادانه از تصادفی بودن استفاده می کنند به برنامه هایی که استفاده صرفه جویانه ای از تصادفی بودن به عمل می آورند، در اختیار می گذارند.

درواقع، «مولدهای اعداد شبه تصادفی» با نخستین کامپیوترها به عرصه درآمدند. با این حال، در اجراهای نوعی از مولدهایی استفاده می کنند که

به طوری که f محاسبه پذیر زمان چندجمله‌ای (و طول نگهدار) باشد. باید f یکطرفه باشد چه در غیر این صورت می توان $G(U_n)$ را از U_{2n} با تلاش برای وارون کردن و امتحان کردن نتیجه، تمیز داد: وارون کردن f روی توزیع دامنه آن به توزیع $G(U_n)$ باز می گردد، درحالی که احتمال اینکه U_{2n} تحت f دارای وارون باشد، صرف نظر کردنی است.

مسیر جالب توجه، ساختن مولدهای شبه تصادفی براساس هر تابع یکطرفه داخواه است. در حالت کلی (وقتی که f ممکن است یک به یک نباشد)، جرگه $f(U_n)$ ممکن است شبه تصادفی نباشد، و بنابراین ساختمان ۹، (یعنی $G(s) = f(s)b(s)$)، که در آن b یک هسته f است) نمی تواند مستقیماً مورد استفاده قرار گیرد. در عوض این ساختمان همراه با چند ایده دیگر به کار می رود (رگ. [۹]). مثلاً، این ایده ها و بیشتر از آن، جزئیات به اجرا درآوردن آنها به مراتب پیچیده تر از آن هستند که بتوان در اینجا توصیفشان کرد. درواقع شیوه دیگری برای ساخت مولدهای شبه تصادفی بر مبنای هر تابع یکطرفه داخواه ارزش بیشتری دارد.

تابعهای شبه تصادفی

مولدهای شبه تصادفی امکان تولید کارآی دنباله های شبه تصادفی بلند را از بذرهای تصادفی کوتاه فراهم می آورند. تابعهای شبه تصادفی (که در زیر تعریف می شوند) از این هم نیرومندترند: آنها امکان دسترسی کارآی مستقیم به دنباله شبه تصادفی عظیمی را (که بررسی بیت به بیت آن ناشدنی است) فراهم می آورند. به عبارت دیگر، تابعهای شبه تصادفی را می توان در هر کاربرد کارآ (مثلاً قابل ذکرتر از همه در رمزنگاری) به جای تابعهای واقعاً تصادفی قرار داد. یعنی، تابعهای شبه تصادفی از تابعهای تصادفی به وسیله ماشینهای کارآیی که مقادیر تابعها را در شناسه هایی به انتخاب خود به دست می آورند، قابل تمایز نیستند. چنان ماشینهایی مکاشفه نامیده می شوند؛ و اگر M چنان ماشیننی باشد و f یک تابع باشد، آنگاه $M^f(x)$ معرف محاسبه M روی ورودی x است وقتی یرسشهای M به وسیله تابع f پاسخ داده می شوند.

تعریف ۱۱. (تابعهای شبه تصادفی [۶]). یک تابع شبه تصادفی (جرگه)، با پارامترهای طول $l_D, l_R : \mathbb{N} \rightarrow \mathbb{N}$ ، مجموعه ای از تابعهای

$$F \stackrel{\text{تعریف}}{=} \{f_s : \{0, 1\}^{l_D(|s|)} \rightarrow \{0, 1\}^{l_R(|s|)} \mid s \in \{0, 1\}^*\}$$

است که در شرطهای زیر صدق می کند:

• (محاسبه کارآ). یک الگوریتم کارآی (یعنی) موجود است به طوری که وقتی یک بذر s و یک شناسه $l_D(|s|)$ بیتی، x داده شده باشد، مقدار $f_s(x)$ به طول $l_R(|s|)$ را باز می گرداند.

(بنابراین، بذر s «توصیف کارآیی» از تابع f_s است.)

• (شبه تصادفی بودن). برای هر ماشین مکاشفه زمان چندجمله‌ای احتمالاتی مانند M ، هر چندجمله‌ای مثبت مانند p ، و همه n های به قدر کافی بزرگ،

$$\left| \Pr_{f \sim F_n}[M^f(1^n) = 1] - \Pr_{p \sim R_n}[M^p(1^n) = 1] \right| < \frac{1}{p(n)}$$

که در آن F_n معرف توزیع روی $f_s \in F$ است که با انتخاب s به طور یکنواخت در $\{0, 1\}^n$ به دست آمده است و R_n معرف توزیع یکنواخت

شبیه‌تصادفی بودن، یک رهیافت رفتارشناختی است. به علاوه، توزیعهای احتمالی موجودند که یکنواخت نیستند (و حتی از نظر آماری به یک توزیع یکنواخت نزدیک نیستند) اما به وسیله هیچ شیوه کارایی از یک توزیع یکنواخت قابل تمایز نیستند. بنابراین، توزیعهایی که از نظر هستی‌شناختی بسیار متفاوت باشند، بنابر دیدگاه رفتارشناختی که در تعریف بالا اتخاذ شد، معادل تلقی می‌شوند.

دیدگاه نسبی‌گرایانه‌ای در باره تصادفی بودن

شبیه‌تصادفی بودن در بالا نسبت به مشاهده‌گر آن تعریف شد. منظور از آن، توزیعی است که نمی‌توان آن را به وسیله هیچ مشاهده‌گر کاراً (یعنی، زمان چندجمله‌ای) از یک توزیع یکنواخت تمیز داد. با این حال، دنباله‌های شبیه‌تصادفی را می‌توان از دنباله‌های تصادفی به وسیله کامپیوترهای بینهایت قدرتمند (که در اختیار ما نیستند!) تمیز داد. به طور مشخص، یک ماشین زمان‌نمایی می‌تواند به آسانی، خروجی یک مولد شبه‌تصادفی را از رشته‌ای با همان طول که به طور یکنواخت انتخاب شده است، تمیز دهد (مثلاً، فقط با امتحان کردن کلیه بذرهای ممکن). بنابراین، شبه‌تصادفی بودن، امری ذهنی، وابسته به تواناییهای مشاهده‌گر است.

تصادفی بودن و دشواری محاسباتی

شبیه‌تصادفی بودن و دشواری محاسباتی نقشهایی دوگان یکدیگر دارند: تعریف شبه‌تصادفی بودن بر این حقیقت متکی است که قراردادن محدودیتهای محاسباتی بر مشاهده‌گر به توزیعهایی منجر می‌شود که یکنواخت نیستند ولی با این حال نمی‌توان آنها را از توزیع یکنواخت تمیز داد. به علاوه، ساختمان مولدهای شبه‌تصادفی متکی بر حدسهایی در باره دشواری محاسباتی (یعنی، وجود تابع‌های یکطرفه) است، و این امر اجتناب‌ناپذیر است: با مفروض بودن یک مولد شبه‌تصادفی، می‌توانیم تابع‌های یکطرفه بسازیم. بنابراین، شبه‌تصادفی بودن نابیهی و دشواری محاسباتی را می‌توان با هم جابه‌جا کرد.

تعمیم

شبیه‌تصادفی بودن را به صورتی که در بالا مرور شد، می‌توان حالت خاص مهمی از یک پارادایم [الگوی] کلی تلقی کرد. یک صورزنبدی عام مولدهای شبه‌تصادفی عبارت است از مشخص کردن سه جنبه بنیادی — اندازه کش‌آوری مولدها، رده متمایزکنندگانی که انتظار می‌رود مولدها فریبشان بدهند (یعنی الگوریتمهایی که نسبت به آنها تمایزناپذیری محاسباتی باید برقرار باشد)، و منابعی که مولدها مجاز به استفاده از آنها هستند (یعنی، پیچیدگی محاسباتی خود آنها). در توصیف بالا، ما بر مولدهای زمان چندجمله‌ای عطف توجه کردیم (بنابراین اندازه کش‌آوری چندجمله‌ای را داشتیم) که هر مشاهده‌گر زمان چندجمله‌ای احتمالاتی را فریب می‌دهد. حالت‌های گوناگون دیگری نیز مورد توجه‌اند، و به اختصار بعضی از آنها را مورد بحث قرار می‌دهیم. برای تفصیل بیشتر، رک. [۵].

مفاهیم ضعیف‌تر تمایزناپذیری محاسباتی

هر گاه منظور ما این باشد که به جای دنباله‌های تصادفی، که به وسیله الگوریتمی به‌کار گرفته شده‌اند، دنباله‌های شبه‌تصادفی را قرار دهیم، می‌توان

بنابر تعریف بالا شبه‌تصادفی نیستند. در عوض، نشان داده می‌شود که این مولدها، در بهترین حالت، درخی آزمونه‌ای آماری خاص را با موفقیت از سر می‌گذرانند. (رک. [۱۰]). با این حال، این حقیقت که یک «مولد اعداد شبه‌تصادفی» برخی آزمونه‌ای آماری را با موفقیت می‌گذراند، به این معنی نیست که آزمون جدیدی را خواهد گذراند یا برای کاربردی (آزمون نشده) در آینده مناسب خواهد بود. به علاوه، رهیافت قراردادن مولد در معرض برخی آزمونه‌ای خاص، قادر به تأمین نتایج کافی از نوعی که در بالا بیان شد، نیست (یعنی، به شکل «برای کلیه مقاصد عملی، استفاده از خروجی مولد، همان حسن را دارد که استفاده از پرتابهای سکه ناریب واقعی»). در مقابل، رهیافتی که در تعریف ۲ مندرج است، چنان کلیتی را هدف خود قرار می‌دهد و درواقع برای به‌دست آوردن آن طرح‌ریزی شده است: مفهوم تمایزناپذیری محاسباتی، که زمینه تعریف ۲ است، کلیه کاربردهای ممکن کاراً را در برمی‌گیرد و مبتنی بر این فرض است که برای همه آنها، دنباله‌های شبه‌تصادفی به‌خوبی دنباله‌های واقعاً تصادفی‌اند.

مولدها و تابع‌های شبه‌تصادفی در رمزنگاری اهمیت اساسی دارند. آنها نوعاً برای ایجاد طرح‌های رمزگذاری کلید شخصی^۱ و تعیین صحت به‌کار می‌روند (رک. [۵]، بخش‌های ۲.۵.۱ و ۲.۶.۱). به عنوان مثال، فرض کنید که دو گروه در یک رشته تصادفی m بیتی مانند s ، که یک تابع شبه‌تصادفی را (مطابق تعریف ۱۱ $l_D(n) = l_R(n) = n$) مشخص می‌کنند، شریک‌اند، و s برای رقیب آنها نامعلوم است. در این صورت دو گروه مزبور می‌توانند پیام‌هایی رمزی را با XOR کردن پیام با مقدار s در نقطه‌ای تصادفی به یکدیگر ارسال کنند، یعنی برای رمزگذاری $m \in \{0, 1\}^n$ ، فرستنده به‌طور یکنواخت $r \in \{0, 1\}^n$ را انتخاب و $(r, m \oplus f_s(r))$ را به گیرنده می‌فرستد. امنیت این طرح رمزگذاری بر این حقیقت متکی است که برای هر رقیب از نظر محاسباتی ممکن (نه تنها برای استراتژیهای رقیب که به تصور درآمده و مورد امتحان قرار گرفته‌اند)، مقادیر تابع f_s روی چنان r هایی تصادفی به نظر می‌رسد.

درونمایه‌های نظری مولدهای شبه‌تصادفی

اینک به اختصار برخی جنبه‌های نظری مولدهای شبه‌تصادفی را، به صورتی که در بالا تعریف شدند، مورد بحث قرار می‌دهیم.

رفتارشناخت در برابر هستی‌شناخت

تعریف ما از مولدهای شبه‌تصادفی مبتنی بر مفهوم تمایزناپذیری محاسباتی است. ماهیت رفتارشناختی مفهوم اخیر به بهترین نحو از مقایسه آن با رهیافت کولموگوروف شایسته به تصادفی بودن، مدلل می‌شود. به بیان غیردقیق، یک رشته تصادفی-کولموگوروفی است هر گاه طول آن برابر با طول کوتاهترین برنامه‌ای باشد که آن را تولید می‌کند. این کوتاهترین برنامه را می‌توان «توضیح واقعی» پدیده‌ای که به وسیله آن رشته توصیف می‌شود، تلقی کرد. بنابراین یک رشته تصادفی-کولموگوروفی رشته‌ای است که توضیح اساساً ساده‌تری (یعنی کوتاه‌تری) در مقایسه با خودش نداشته باشد. در نظر گرفتن ساده‌ترین توضیح یک پدیده را می‌توان رهیافتی هستی‌شناختی تلقی کرد. در مقابل، در نظر گرفتن تأثیر پدیده (بر یک مشاهده‌گر) به عنوان زمینه تعریف

1. private-key 2. exclusive or

مراجع

1. W. ALEXI, B. CHOR, O. GOLDBREICH, and C. P. SCHNORR, RSA/Rabin functions: Certain parts are as hard as the whole, *SIAM J. Comput.* **17** (1988), 194-209.
2. M. BLUM and S. MICALI, How to generate cryptographically strong sequences of pseudo-random bits, *SIAM J. Comput.* **13** (1984), 850-864.
3. T. M. COVER and G. A. THOMAS, *Elements of Information Theory*, Wiley, New York, 1991.
4. O. GOLDBREICH, *Foundation of Cryptography-Fragments of a Book*, February 1995, available from <http://theory.ics.mit.edu/~oded/frag.html>.
5. ———, *Modern Cryptography, Probabilistic Proofs and Pseudorandomness*, Algorithms and Combinatorics, vol. 17, Springer-Verlag, New York, 1998.
6. O. GOLDBREICH, S. GOLDWASSER, and S. MICALI, How to construct random functions, *J. ACM* **33** (1986), 792-807.
7. O. GOLDBREICH and L. A. LEVIN, Hard-core predicates for any one-way function, *21st ACM Symposium on the Theory of Computing*, 1989, pp. 25-32.
8. S. GOLDWASSER and S. MICALI, Probabilistic encryption, *J. Comput. System Sci.* **28** (1984), 270-299.
9. J. HASTAD, R. IMPAGLIAZZO, L. A. LEVIN, and M. LUBY, A pseudorandom generator from any one-way function, *SIAM J. Comput.* **28** (1999), 1364-1396.
10. D. E. KNUTH, *Seminumerical Algorithms*, The Art of Computer Programming, Vol. 2, Addison-Wesley, Reading, MA, 1969; second edition, 1981.
11. L. A. LEVIN, Randomness conservation inequalities: Information and independence in mathematical theories, *Inform. and Control* **61** (1984), 15-37.
12. M. LI and P. VITANYI, *An Introduction to Kolmogorov Complexity and Its Applications*, Springer-Verlag, New York, 1993.
13. A. C. YAO, Theory and application of trapdoor functions, *23rd IEEE Symposium on Foundations of Computer Science*, 1982, pp. 80-91.

- Oded Goldreich, "Pseudorandomness", *Notices Amer. Math. Soc.*, (10) **46** (1999) 1209-1216.

* اودت گلدرایش، مؤسسه علوم وایزمن، اسرائیل

oded@wisdom.weizmann.ac.il

تلاش کرد که از اطلاعات مربوط به الگوریتم هدف سود برده شود. در بالا صرفاً از این حقیقت استفاده کرده‌ایم که الگوریتم هدف در زمان چندجمله‌ای اجرا می‌شود. ولی اگر مثلاً بدانیم که الگوریتم از فضای کاری بسیار کمی استفاده می‌کند، آنگاه ممکن است قادر باشیم بهتر عمل کنیم. همین‌طور ممکن است بتوانیم بهتر عمل کنیم هر گاه بدانیم که تحلیل الگوریتم تنها به برخی خاصیت‌های دنباله تصادفی که به کار می‌برد، و وابسته است (مثلاً، استقلال دوبه‌دوی عناصر آن)، در حالت کلی، مفاهیم ضعیف‌تر نمایانپذیری محاسباتی نظیر فریب دادن الگوریتم‌های فضا کراندار، مدارهای ژرفا ثابت، و حتی آزمون‌های مشخص (مثلاً، آزمون کردن استقلال دوبه‌دوی دنباله) به‌طور طبیعی پیش می‌آیند. مولدهایی که دنباله‌هایی را تولید می‌کنند که چنان آزمون‌هایی را فریب می‌دهند، در کاربردهای گوناگونی مفیدند؛ اگر در کاربرد مورد نظر، تصادفی بودن به‌طریقی محدود به کاربرده شود، در این صورت دادن دنباله‌هایی با کیفیت تصادفی پایین به خورد آن، ممکن است کافی باشد. نیازی به گفتن نیست که نویسنده از صورت‌بندی دقیق مشخصه‌هایی با چنان کاربردها و ساخت دقیق مولدهایی هواداری می‌کند که آن نوع از آزمون‌ها را که پیش می‌آیند، فریب دهد.

سایر مفاهیم کارایی مولدها

پاراگراف پیشین بر یک جنبه از مسأله شبه تصادفی بودن تمرکز داشت: منابع یا نوع مشاهده‌گر (یا متمایزکننده بالقوه) پرسش مهم دیگر این است که آیا چنان دنباله‌های شبه تصادفی را می‌توان از دنباله‌های بسیار کوتاه‌تر تولید کرد و به چه بهایی (یعنی با چه میزان پیچیدگی). در سرتاسر این مقاله شرط کرده‌ایم که فرایند تولید دست‌کم به قدر متمایزکننده، کارا باشد^۱. این شرط درواقع، «منصفانه» و طبیعی است. مجاز داشتن مولد به اینکه پیچیده‌تر از متمایزکننده باشد (یعنی، از زمان یا فضای بیشتر استفاده کند)، غیرمنصفانه به نظر می‌رسد اما باز هم پیامدهای جالب توجهی در چارچوب تلاش برای «باصصادی‌سازی»^۲ رده‌های پیچیدگی تصادفی دارد. به عنوان مثال، می‌توان مولدهایی را در نظر گرفت که نسبت به طول بذری در زمان نمایی کار می‌کنند. در برخی حالتها، چنین آسانگیریهایی (یعنی مجاز داشتن مولدهای زمان نمایی) چیزی از دست نمی‌دهیم. برای ملاحظه دایل آن، یک استدلال نوعی مبتنی بر باد تصادفی‌سازی را در نظر می‌گیریم که مرکب از دو گام است: ابتدا به جای دنباله‌های واقعاً تصادفی الگوریتم، دنباله‌های تصادفی تولید شده از بذری بسیار کوتاه‌تر را قرار می‌دهیم، و سپس به‌طور تعینی به تمام بذری ممکن رجوع می‌کنیم و در پی یافتن رفتار با بیشترین فراوانی الگوریتم اصلاح شده برمی‌آیم. در چنین حالتی، پیچیدگی تعینی در هر حال نسبت به طول بذری، نمایی است. نفع این کار آن است که ساختن مولدهای زمان نمایی ممکن است آسان‌تر از ساختن مولدهای زمان چندجمله‌ای باشد.

۱. درواقع شرط کرده‌ایم که مولد کارا تر از متمایزکننده باشد، یعنی شرط کرده‌ایم که مولد یک الگوریتم زمان چندجمله‌ای ثابت باشد، در حالی که متمایزکننده مجاز است که هر الگوریتمی با زمان اجرای چندجمله‌ای باشد.

2. derandomization