

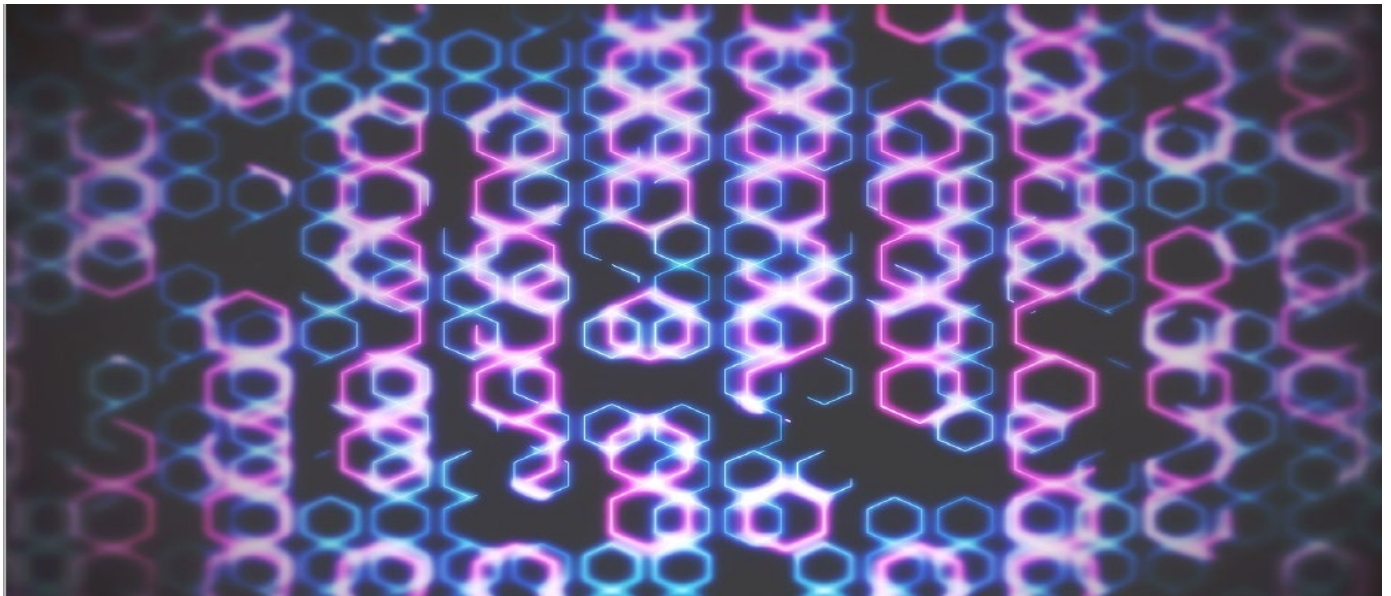


با پیشرفت هوش مصنوعی مولد، نهادهای نظارتی و واحدهای مدیریت ریسک در تلاشند تا با این پیشرفت‌ها همگام شوند

ترجمه از: مهدی حسن زاده ، کارشناسی ارشد MBA، دانشکده مدیریت و اقتصاد، دانشگاه صنعتی شریف

زیر نظر: دکتر آرش خلیلی نصر، عضو هیات علمی دانشکده مدیریت و اقتصاد، دانشگاه صنعتی شریف

هوش مصنوعی و به‌ویژه دستاورد خیره‌کننده آن، یعنی هوش مصنوعی مولد، با سرعتی چشمگیر در حال توسعه است و
قانون‌گذاران تحت فشارند تا با این پیشرفت‌ها همگام شوند



پیشرفت سریع هوش مصنوعی مولد (Gen AI) نهادهای نظارتی سراسر جهان را بر آن داشته تا برای درک، کنترل و تضمین ایمنی این فناوری تلاش کرده و در عین حال مزایای بالقوه آن را نیز حفظ کنند. در صنایع مختلف، پذیرش این فناوری چالشی تازه برای واحدهای ریسک و انطباق ایجاد کرده است: چگونه می‌توان از هوش مصنوعی مولد در چارچوب‌های نظارتی در حال تحول (و ناهمگون) استفاده کرد؟

همزمان با تلاش دولت‌ها و قانون‌گذاران برای تعریف اینکه این محیط کنترلی چگونه باید باشد، رویکردهای در حال شکل‌گیری در بسیاری موارد پراکنده و ناسازگارند؛ مسئله‌ای که مسیر حرکت سازمان‌ها را دشوارتر کرده و عدم اطمینان قابل توجهی به وجود آورده است.

در این مقاله، ضمن بررسی ریسک‌های مرتبط با هوش مصنوعی و هوش مصنوعی مولد و دلایل جلب توجه نهادهای نظارتی، نقشه راهی استراتژیک ارائه می‌دهیم تا به واحدهای مدیریت ریسک در عبور از محیط نظارتی ناهمگون و در حال تغییری که نه تنها بر هوش مصنوعی مولد بلکه بر کل حوزه هوش مصنوعی متمرکز است، کمک کنیم.

چرا هوش مصنوعی مولد به تنظیم مقررات نیاز دارد؟

پیشرفت چشمگیر هوش مصنوعی، به‌ویژه در حوزه هوش مصنوعی مولد، به سرعت توجه عموم مردم را به خود جلب کرده است. برای مثال، چت‌جی‌پی‌تی¹ به یکی از سریع‌ترین پلتفرم‌های در حال رشد تبدیل شد و تنها ظرف پنج روز پس از عرضه به یک میلیون کاربر رسید. با توجه به گستره وسیع کاربردهای هوش مصنوعی مولد (از افزایش بهره‌وری و دسترسی سریع‌تر به دانش گرفته تا ایجاد تأثیر اقتصادی سالانه‌ای بین 2.6 تا 4.4 تریلیون دلار²) این استقبال گسترده چندان جای تعجب ندارد.

با این حال، انگیزه اقتصادی برای استفاده صحیح از هوش مصنوعی و هوش مصنوعی مولد وجود دارد. شرکت‌های توسعه‌دهنده هوش مصنوعی با این ریسک مواجه هستند که اگر پلتفرم‌هایشان به اندازه کافی کامل و بی‌نقص نباشند، با عواقب جدی روبه‌رو شوند؛ چراکه حتی یک خطا می‌تواند بسیار پرهزینه باشد. برای مثال، شرکت‌های بزرگ فعال در حوزه

¹ ChatGPT

² «پتانسیل اقتصادی هوش مصنوعی مولد: مرز بعدی بهره‌وری»، مکینزی، 14 ژوئن 2023.

هوش مصنوعی مولد در پی بروز پدیده هذیان‌گویی³ (تولید اطلاعات نادرست یا غیرمنطقی توسط هوش مصنوعی) شاهد افت چشمگیر ارزش بازار خود بوده‌اند.

گسترش سریع هوش مصنوعی مولد باعث برجسته‌تر شدن انواع ریسک‌ها شده است. نگرانی‌های اصلی پیرامون این فناوری شامل چگونگی توسعه مدل‌ها و سیستم‌ها، شیوه استفاده از آن‌ها و چالش‌های مرتبط است.

به‌طور کلی، موضوعاتی چون کمبود شفافیت در عملکرد سیستم‌های هوش مصنوعی مولد، داده‌های مورد استفاده برای آموزش این سیستم‌ها، مسائل مربوط به سوگیری و انصاف، نقض احتمالی حقوق مالکیت فکری، نقض حریم خصوصی، ریسک‌های مرتبط با اشخاص ثالث و نگرانی‌های امنیتی از جمله مهم‌ترین نگرانی‌هایی مطرح‌شده به شمار می‌روند.

به فهرست این نگرانی‌ها باید مسئله انتشار اطلاعات نادرست را نیز اضافه کرد؛ مواردی مانند تولید خروجی‌های اشتباه یا دستکاری‌شده و همچنین محتوای مضر یا مخرب. جای تعجب نیست که نهادهای نظارتی در تلاشند تا آسیب‌های احتمالی ناشی از این فناوری را کاهش دهند. آن‌ها می‌کوشند برای شرکت‌هایی که در زمینه توسعه یا استفاده از هوش مصنوعی مولد فعالیت می‌کنند، قطعیت قانونی⁴ ایجاد کنند. در عین حال، قانون‌گذاران می‌خواهند بدون ایجاد ترس از پیامدهای ناشناخته، از نوآوری حمایت کنند.

هدف، ایجاد استانداردهای نظارتی بین‌المللی هماهنگی است که به تقویت تجارت بین‌المللی و تسهیل انتقال داده‌ها کمک می‌کنند. در همین راستا، نوعی اجماع شکل گرفته است: جامعه توسعه‌دهندگان هوش مصنوعی مولد خود از پیشگامان حمایت از اعمال کنترل‌های نظارتی بر توسعه این فناوری هستند و خواهان اجرایی شدن این کنترل‌ها در سریع‌ترین زمان ممکن شده‌اند. اکنون دیگر بحث بر سر این مسئله نیست که آیا باید قوانین را وضع کرد یا خیر، بلکه بحث بر سر این مسئله است که این کار چگونه باید انجام شود.

³ hallucinating

⁴ Legal certainty: به معنای اطمینان از وضوح قوانین و مقررات برای افراد و سازمان‌ها است تا بتوانند با آگاهی از حقوق و تکالیف خود فعالیت کنند.

چشم‌انداز کنونی نظارت بین‌المللی بر حوزه هوش مصنوعی

تا به امروز، هیچ کشوری قانون جامعی برای هوش مصنوعی یا هوش مصنوعی مولد تصویب نکرده است، اما چند کشور و منطقه از جمله برزیل، چین، اتحادیه اروپا، سنگاپور، کره جنوبی و ایالات متحده پیشگام تلاش‌های قانون‌گذاری در این زمینه بوده‌اند. رویکرد هر کدام از آن‌ها نسبت به قانون‌گذاری متفاوت است: از تنظیم قوانین گسترده برای هوش مصنوعی که توسط مقررات موجود در زمینه حفاظت از داده‌ها و امنیت سایبری پشتیبانی می‌شوند (اتحادیه اروپا و کره جنوبی) گرفته تا تصویب قوانین ویژه برای صنایع مختلف (ایالات متحده) و رویکردهای مبتنی بر اصول یا رهنمودهای کلی (برزیل، سنگاپور و ایالات متحده). هر یک از این رویکردها مزایا و معایب خود را دارند و انتظار می‌رود برخی بازارها در گذر زمان از اصول کلی به سمت مقررات سخت‌گیرانه‌تر حرکت کنند (شکل 1).

با وجود تفاوت در رویکردها، مضامین مشترکی در چشم‌انداز جهانی نظارت بر هوش مصنوعی پدیدار شده‌اند:

- **شفافیت:** نهادهای نظارتی به دنبال این هستند که خروجی‌های تولیدشده توسط هوش مصنوعی قابل‌پیگیری و شفاف باشند. هدف آن‌ها این است که کاربران هنگام کار با هر سیستم هوش مصنوعی بدانند که در حال تعامل با یک ماشین هستند و اطلاعاتی درباره حقوق خودشان و همچنین قابلیت‌ها و محدودیت‌های آن سیستم در اختیار داشته باشند.
- **عاملیت انسان و نظارت:** در حالت ایده‌آل، سیستم‌های هوش مصنوعی باید به‌عنوان ابزارهایی در خدمت انسان توسعه یابند. این سیستم‌ها باید کرامت انسانی و استقلال فردی را حفظ کنند و به گونه‌ای طراحی شوند که امکان کنترل و نظارت مناسب توسط انسان‌ها وجود داشته باشد.
- **پاسخگویی:** نهادهای نظارتی به دنبال سازوکارهایی هستند که مسئولیت‌پذیری، پاسخگویی و امکان جبران خسارت در ارتباط با سیستم‌های هوش مصنوعی را تضمین کنند. در عمل، این به معنای جلب حمایت مدیران ارشد، ارائه آموزش در سطح سازمان و ایجاد آگاهی از مسئولیت‌های فردی در میان کارکنان است.
- **ایمنی و استواری فنی:** قانون‌گذاران در تلاشند تا آسیب‌های غیرمنتظره و ناخواسته را کاهش دهند و این کار را از طریق اطمینان از استواری سیستم‌های هوش مصنوعی انجام می‌دهند؛ به این معنا که این سیستم‌ها باید مطابق انتظار عمل کنند، پایدار باقی بمانند، بتوانند خطاهای کاربران را اصلاح کنند، در برابر تلاش‌های مخرب اشخاص ثالث برای دستکاری سیستم مقاوم باشند و برای رفع هرگونه ناکامی در برآورده کردن این معیارها، راهکارهای جایگزین و اقدامات جبرانی مناسبی داشته باشند.

شکل 1

مقررات مرتبط با حکمرانی هوش مصنوعی در کشورهای مختلف متفاوت است.

داده‌های گردآوری شده تا نوامبر 2023، غیرجامع

تعدادی از کشورهای بدون قوانین کلی هوش مصنوعی	پیشنهاد یا تکمیل قوانین کلی هوش مصنوعی	نوع سیاست: اصول غیرالزام‌آور (مثال،
• استرالیا • هند • نیوزیلند • عربستان سعودی	• برزیل • کانادا • چین • کره جنوبی • اتحادیه اروپا	• ژاپن • سنگاپور • امارات متحده عربی • بریتانیا • ایالات متحده • سایر کشورهای عضو OECD

منبع: سازمان همکاری و توسعه اقتصادی (OECD)، تحلیل مکی‌نری

- **تنوع، عدم تبعیض و رعایت انصاف:** از دیگر اهداف نهادهای نظارتی این است که سیستم‌های هوش مصنوعی عاری

از هرگونه سوگیری عمل کنند و خروجی‌های آن‌ها به تبعیض یا رفتار ناعادلانه با افراد منجر نشود.

- **حریم خصوصی و حکمرانی داده‌ها:** نهادهای نظارتی می‌خواهند توسعه و استفاده از سیستم‌های هوش مصنوعی

طبق قوانین موجود حریم خصوصی و حفاظت از داده‌ها باشد و این سیستم‌ها داده‌هایی با کیفیت و یکپارچگی بالا را پردازش کنند.

- **رفاه اجتماعی و زیست‌محیطی:** نهادهای نظارتی تمایل شدیدی دارند که سیستم‌های هوش مصنوعی پایدار باشند،

با محیط‌زیست (برای مثال، در مصرف انرژی) سازگار باشند و برای همه افراد مفید باشند. همچنین خواستار نظارت و ارزیابی مداوم اثرات بلندمدت این فناوری بر افراد، جامعه و دموکراسی هستند.

با وجود برخی وجوه مشترک در اصول راهنمای مرتبط با هوش مصنوعی، شیوه اجرا و عبارات به کاررفته توسط یک نهاد

نظارتی یا منطقه با نهادهای نظارتی و مناطق دیگر متفاوت است. بسیاری از قوانین هنوز جدید هستند و بنابراین ممکن

است به‌طور مکرر به‌روزرسانی شوند (شکل 2). این موضوع برنامه‌ریزی استراتژیک بلندمدت در زمینه هوش مصنوعی را

برای سازمان‌ها دشوار می‌کند.

این مسائل برای سازمان‌ها چه معنایی دارد؟

ممکن است برخی سازمان‌ها وسوسه شوند که صبر کنند و ببینند چه مقرراتی در زمینه هوش مصنوعی تصویب می‌شوند. اما اکنون زمان اقدام است. سازمان‌هایی که در این مرحله اقدام نکنند، ممکن است در آینده با ریسک‌های بزرگ حقوقی، اعتباری، سازمانی و مالی روبه‌رو شوند. برخی بازارها از جمله ایتالیا، به دلیل نگرانی‌های مربوط به حریم خصوصی، شکایات مربوط به نقض حق نشر از سوی چندین فرد و نهاد و دعاوی مربوط به تهمت و افترا، دسترسی به چت‌جی‌پی‌تی را ممنوع کرده‌اند.

احتمالاً در آینده موانع بیشتری وجود خواهد داشت. هرچه تأثیرات منفی هوش مصنوعی بیشتر شناخته و اطلاع‌رسانی شوند، نگرانی‌های عمومی نیز افزایش خواهد یافت. این موضوع می‌تواند موجب بی‌اعتمادی عمومی نسبت به شرکت‌های توسعه‌دهنده یا استفاده‌کننده از هوش مصنوعی شود.

سازمان‌ها در حین برنامه‌ریزی استراتژی‌های بلندمدت هوش مصنوعی خود با چالش رعایت

مقررات متنوع روبرو هستند.

شکل 2

تلاش‌های سیاستی و نظارتی مرتبط با حکمرانی هوش مصنوعی در سراسر جهان همچنان در حال انجام است.

نمونه‌ها بر اساس نوع سیاست یا تلاش و زمان پیشنهاد گنجانده شده‌اند؛ غیر جامع



منبع: سازمان همکاری و توسعه اقتصادی؛ تحلیل مکینزی

یک خطا در این مرحله می‌تواند بسیار پرهزینه باشد. سازمان‌ها ممکن است به دلیل نقض قوانین با جریمه‌های سنگینی مواجه شوند (برای مثال، طبق قانون پیشنهادی اتحادیه اروپا برای هوش مصنوعی، جریمه‌ای تا سقف 7 درصد از درآمد سالانه جهانی). تهدید دیگری که وجود دارد، زیان مالی ناشی از کاهش اعتماد مشتریان یا سرمایه‌گذاران است که می‌تواند

به کاهش قیمت سهام، از دست دادن مشتریان یا کند شدن روند جذب مشتری منجر شود. انگیزه برای اقدام سریع زمانی دوچندان می‌شود که بدانیم اگر مدل‌های حکمرانی و سازمانی مناسب برای هوش مصنوعی از ابتدا ایجاد نشوند، ممکن است در آینده به علت تغییر مقررات، وقوع رخنه‌های اطلاعاتی یا حملات سایبری، این مدل‌ها به اصلاحات اساسی نیاز پیدا کنند. اصلاح سیستم پس از بروز مشکل، نه تنها بسیار پرهزینه است، بلکه اجرای یکپارچه آن در سراسر سازمان هم دشوار خواهد بود.

آینده دقیق تعهدات قانونی هنوز مبهم است و ممکن است بسته به منطقه جغرافیایی و نقش خاص هوش مصنوعی در زنجیره ارزش، متفاوت باشد. با این حال، برخی اقدامات وجود دارند که سازمان‌ها می‌توانند همین امروز انجام دهند تا در برابر تغییرات قریب‌الوقوع قوانین آمادگی داشته باشند.

این اقدامات پیش‌دستانه را می‌توان در چهار حوزه کلیدی دسته‌بندی کرد که از تلاش‌های موجود برای حفاظت از داده‌ها، حریم خصوصی و امنیت سایبری نشأت می‌گیرند، چراکه این حوزه‌ها اشتراکات زیادی با یکدیگر دارند:

شفافیت: ایجاد یک فهرست و طبقه‌بندی از مدل‌ها، دسته‌بندی آن‌ها بر اساس مقررات و ثبت تمام کاربردهای آن‌ها در یک مخزن مرکزی که برای افراد داخل و خارج از سازمان قابل مشاهده باشد. همچنین باید مستندات دقیق و روشنی از کاربردهای داخلی و خارجی هوش مصنوعی و هوش مصنوعی مولد، عملکرد این سیستم‌ها، ریسک‌ها و کنترل‌هایشان تهیه شود. علاوه بر این، لازم است توضیح داده شود که هر مدل چگونه توسعه یافته است، چه خطراتی دارد و قرار است به چه شکلی استفاده شود.

حکمرانی: پیاده‌سازی یک ساختار حکمرانی برای هوش مصنوعی و هوش مصنوعی مولد که نظارت، اختیار و پاسخگویی کافی را هم در داخل سازمان و هم در تعامل با اشخاص ثالث و نهادهای نظارتی تضمین کند. در این ساختار، باید تمام نقش‌ها و مسئولیت‌های مربوط به مدیریت هوش مصنوعی و هوش مصنوعی مولد مشخص شوند و برنامه‌ای برای مدیریت حوادث احتمالی ناشی از این فناوری‌ها تدوین گردد. این ساختار حکمرانی باید آن قدر قوی باشد که تغییرات نیروی انسانی و گذر زمان آن را تضعیف نکند، ولی در عین حال آن قدر منعطف و چابک باشد که بتواند با تغییرات فناوری، نیازهای جدید کسب‌وکار و الزامات جدید نظارتی سازگار شود.

مدیریت داده‌ها، مدل و فناوری: هوش مصنوعی و هوش مصنوعی مولد به مدیریت قوی داده‌ها، مدل و فناوری نیاز دارند:

- مدیریت داده‌ها: داده‌ها پایه و اساس هر مدل هوش مصنوعی و هوش مصنوعی مولد هستند. کیفیت داده‌های ورودی مستقیماً روی کیفیت خروجی مدل تأثیر می‌گذارد. مدیریت صحیح و قابل‌اعتماد داده‌ها به معنای آگاهی کامل از منابع داده، دسته‌بندی آن‌ها، اطمینان از کیفیت و اصل و نسب داده‌ها، رعایت مالکیت فکری و محافظت از حریم خصوصی است.

- مدیریت مدل: سازمان‌ها می‌توانند اصول و چارچوب‌هایی محکمی برای توسعه و استفاده از هوش مصنوعی و هوش مصنوعی مولد ایجاد کنند تا ریسک‌های سازمان را به حداقل برسانند و اطمینان حاصل کنند که تمام مدل‌های هوش مصنوعی و هوش مصنوعی مولد، انصاف و کنترل‌های سوگیری، عملکرد صحیح، شفافیت، وضوح و امکان نظارت انسانی را رعایت می‌کنند. آموزش همه کارکنان سازمان درباره توسعه صحیح و استفاده درست از این فناوری‌ها، برای کاهش ریسک‌ها کاملاً ضروری است. همچنین لازم است طبقه‌بندی ریسک و چارچوب ریسک سازمان طوری تدوین شود که ریسک‌های مربوط به هوش مصنوعی مولد را هم در بر بگیرد، نقش‌ها و مسئولیت‌های مرتبط با مدیریت این ریسک‌ها تعیین شوند و ارزیابی‌ها و کنترل‌های ریسک با مکانیزم‌های مناسب آزمایش و نظارت برای کنترل و کاهش ریسک‌ها برقرار گردد. مدیریت داده و مدل باید فرایندی چابک و مستمر داشته باشد و نباید فقط به‌عنوان کاری نمادین در ابتدای پروژه‌های توسعه تلقی شود.

- مدیریت فناوری و امنیت سایبری: باید ساختار امنیت سایبری و فناوری قدرتمندی ایجاد شود که شامل کنترل دسترسی، فایروال‌ها، لاگ‌ها، نظارت و موارد مشابه باشد تا از دسترسی غیرمجاز یا سوءاستفاده جلوگیری کرده و بتواند حوادث احتمالی را خیلی زود شناسایی کند.

حقوق فردی: کاربران را آموزش دهید. آن‌ها باید بدانند که با سیستم هوش مصنوعی تعامل دارند و همچنین دستورالعمل‌های روشنی برای نحوه استفاده از این سیستم‌ها در اختیارشان قرار دهید. لازم است یک نقطه تماس مشخص ایجاد شود تا شفافیت لازم را فراهم سازد و کاربران بتوانند حقوق خود را اعمال کنند، از جمله چگونگی دسترسی به داده‌ها، نحوه عملکرد مدل‌ها و شیوه انصراف از به‌کارگیری پلتفرم. در نهایت باید رویکردی مشتری‌محور در طراحی و کاربرد هوش مصنوعی اتخاذ شود که هم پیامدهای اخلاقی استفاده از داده‌ها و هم تأثیر بالقوه آن‌ها بر مشتریان را در نظر بگیرد. از آنجاکه همواره آنچه قانونی است الزاماً اخلاقی نیست، در اولویت قرار دادن ملاحظات اخلاقی در استفاده از هوش مصنوعی از اهمیت بالایی برخوردار است.

هوش مصنوعی و هوش مصنوعی مولد همچنان تأثیر قابل توجهی بر بسیاری از سازمان‌ها خواهند گذاشت، چه آن‌هایی که ارائه‌دهنده مدل‌های هوش مصنوعی هستند و چه آن‌هایی که از این سیستم‌ها استفاده می‌کنند. با وجود اینکه چشم‌انداز نظارتی به سرعت در حال تغییر است و همچنان در مناطق جغرافیایی و صنایع مختلف هماهنگ نشده و ممکن است غیرقابل پیش‌بینی به نظر برسد، سازمان‌هایی که از هم‌اکنون در جهت بهبود نحوه ارائه و به کارگیری هوش مصنوعی گام برمی‌دارند، از مزایای ملموسی برخوردار خواهند شد.

ناکامی در مدیریت محتاطانه هوش مصنوعی و هوش مصنوعی مولد می‌تواند به خسارات حقوقی، اعتباری، سازمانی و مالی منجر شود. با این حال، سازمان‌ها می‌توانند با تمرکز بر شفافیت، حکمرانی، حقوق فردی و مدیریت فناوری و داده‌ها، خود را برای آینده آماده کنند. رسیدگی به این حوزه‌ها، بنیانی محکم برای حکمرانی داده‌ها و کاهش ریسک آینده فراهم می‌کند و در عین حال عملیات مرتبط با امنیت سایبری، مدیریت و حفاظت از داده‌ها و استفاده مسئولانه از هوش مصنوعی را تسهیل می‌کند. شاید این موضوع مهم‌تر از هر چیز دیگری باشد که اجرای اقدامات حفاظتی به سازمان‌ها کمک می‌کند تا جایگاه خود را به عنوان ارائه‌دهندگان مورد اعتماد تثبیت کنند.

این مقاله ترجمه مقاله ذیل از شرکت مکنزی است:

As gen AI advances, regulators—and risk functions—rush to keep pace



The Institute for Management Insight Change