



گزارش موردی

۶۵

مرداد و شهریور ۱۴۰۰

هوش مصنوعی و کاربردهای آن
در صنعت بیمه

گزارش موردی

دوماهنامه آموزشی، پژوهشی، تحلیلی

پژوهشکده بیمه، سال یازدهم، شماره ۳، مرداد و شهریور ۱۴۰۰، دوره جدید، شماره مسلسل ۶۵

صاحب امتیاز: پژوهشکده بیمه

مدیر مسئول: دکتر لیلی نیاکان

سر دبیر: دکتر حمید کردبچه

این نشریه به موجب نامه شماره ۸۹/۲۲۴۸۲ مورخ ۸۹/۹/۲۸ وزارت فرهنگ و ارشاد اسلامی دارای مجوز نشر به

عنوان دوماهنامه آموزشی، پژوهشی، تحلیلی در زمینه علوم انسانی (مدیریت بیمه) می باشد.

مدیر داخلی: دکتر شبیم رفوا

صفحه آرا و طراح جلد: علی حسین صفری

نشانی نشریه: تهران - سعادت آباد - میدان شهید تهرانی مقدم (کاج) - خیابان سرو غربی - پلاک ۴۳

صندوق پستی: ۴۴۹۹ - ۱۹۳۹۵

تلفن: ۲۲۰۸۴۰۸۴

دورنگار: ۲۲۰۹۲۲۶۵

Workingpaper@irc.ac.ir

نشانی چاپخانه: تهران - خیابان فاطمی - میدان گلها - خیابان کاج شمالی - پلاک ۱۹۵ - چاپ کارنگ

تلفن: ۸۸۰۲۴۰۱۰

کلیه حقوق برای ناشر، محفوظ و استفاده از مطالب با ذکر مأخذ، مجاز است. مسئولیت مقالات چاپ شده بر عهده نویسنده و مترجم است و مطالب آن لزوماً بیانگر دیدگاه‌های پژوهشکده بیمه نیست. گزارش موردی در پذیرش و ویرایش مقالات آزاد است.



شناسنامه عمومی گزارش موردی

هوش مصنوعی و کاربردهای آن در صنعت بیمه	عنوان گزارش
Artificial intelligence and its applications in the insurance industry	عنوان لاتین
وحیده نورانی	مترجم
دکتر محمدجواد حیدری	ویراستار علمی
مرداد و شهریور ۱۴۰۰ - شماره ۶۵	شمارگان

پیشگفتار



“ هوش مصنوعی، این فناوری می تواند مثل انسان بیندیشد، یاد گیرد و رفتار کند. هوش مصنوعی قادر است داده ها را تفسیر کرده و از آموخته ها برای انجام انواع مختلفی از کارها استفاده کند.

”

فناوری و نیز پتانسیل کاربرد هوش مصنوعی و فرصت های آن در کسب و کارهای صنعت بیمه بیان می شوند. این گزارش نقطه شروعی برای درک بهتر ابعاد کلیدی هوش مصنوعی و نیز نحوه ارتباط آن ها با بیمه است. همچنین در این گزارش، چهار حوزه ریسک مرتبط با هوش مصنوعی مورد بررسی قرار می گیرند که شامل شفافیت و اعتماد، اخلاقیات، تعهدات و مسئولیت ها و امنیت می باشد. در اینجا لازم است از سرکار خانم وحیده نورانی به عنوان مترجم و جناب آقای دکتر محمدجواد حیدری به عنوان ویراستار علمی این شماره از نشریه و همچنین همکاران اداره کتابخانه، اسناد علمی و نشریات سپاسگزاری نمایم.

حمید کردبچه
رئیس پژوهشگاه بیمه

آنچه از بررسی مستندات و تحقیقات مشخص است این است که هوش مصنوعی هم اکنون در ارزیابی ریسک، قیمت گذاری و شناسایی تقلب بیمه ای و پیشگیری از آن به کار گرفته شده و برخی شرکت های بیمه در سراسر جهان از منافع و مزایای آن بهره می برند. اما کاربردهای هوش مصنوعی محدود به این موارد نیست و استفاده های دیگری هم در صنعت بیمه دارد که از آن جمله می توان به امکان ارائه محصولات جدید بیمه ای به کمک هوش مصنوعی اشاره کرد. با توجه به اهمیت موضوع، این شماره از گزارش موردی به هوش مصنوعی و کاربردهای آن در صنعت بیمه اختصاص یافته است. هدف این گزارش، کمک به بیمه گران برای آگاهی از این حوزه است که به سرعت در حال پیشرفت و توسعه می باشد. در این گزارش ریسک های

هوش مصنوعی به یکی از اصطلاحات پر کاربرد عصر دیجیتال تبدیل شده است. بر اساس مفهوم هوش مصنوعی، این فناوری می تواند مثل انسان بیندیشد، یاد گیرد و رفتار کند. هوش مصنوعی قادر است داده ها را تفسیر کرده و از آموخته ها برای انجام انواع مختلفی از کارها استفاده کند. این روزها وقتی اخبار و اسناد مربوط به تحولات و پیشرفت های فناوری در حوزه های مختلف علوم و صنایع مختلف را مرور می کنیم، وجه مشترک همه آن ها این است که هوش مصنوعی در همه صنایع انقلاب عظیمی ایجاد کرده و در آینده نزدیک، هر خدمت و کالایی به نوعی مرتبط با تحلیل ها و الگوریتم های هوش مصنوعی خواهد شد. در طی سال های اخیر، صنعت بیمه نیز شروع به استفاده از ابزارهای هوش مصنوعی کرده تا فعالیت های خود را سروسامان بخشد.

فهرست

مقدمه.....	۶
هوش مصنوعی چیست؟.....	۱۳
تاریخچه‌ای مختصر از پیشرفت هوش مصنوعی و روندهای اصلی آن.....	۱۸
خلاصه پیشرفت‌های اخیر.....	۲۰
نمونه‌هایی از کاربردهای فعلی هوش مصنوعی.....	۲۰
ریسک‌های هوش مصنوعی.....	۲۱
شفافیت و اعتماد.....	۲۱
اخلاقیات.....	۲۲
مسئولیت.....	۲۳
امنیت.....	۲۴
هوش مصنوعی و بیمه.....	۲۶
تأثیر هوش مصنوعی بر رشته‌های بیمه‌ای.....	۲۶
فرصت‌های توسعه کسب و کار.....	۳۱
عملیات.....	۳۲
فناوری بیمه (اینشورتک).....	۳۴
نتیجه‌گیری.....	۳۷
منابع.....	۳۹

مقدمه^۱

هوش مصنوعی - به معنی قابلیت یک ماشین در تقلید رفتار هوشمند انسان - به نحو فزاینده‌ای در جامعه به همه انواع روش‌ها، از تشخیص وضعیت پزشکی و اسکن اسناد حقوقی گرفته تا پاشیدن دقیق سم در کشاورزی برای جلوگیری از مقاومت در برابر علف‌کش‌ها و نیز در اتومبیل‌های خودران، به کار گرفته می‌شود. هوش مصنوعی پدیده جدیدی نیست و در طول ۶۰ سال، حضور قابل توجهی در مطالعات دانشگاهی داشته است. با این حال، یکی از جدیدترین پیاده‌سازی‌های برهم‌زننده هوش مصنوعی در دنیای واقعی، که رشد بسیار سریعی هم دارد، آگاهی‌هایی را درباره چالش‌های پیچیده و عمیق اخلاقی، حقوقی و اجتماعی آن ایجاد کرده است که فراتر از پیچیدگی‌های فنی روش‌ها و فرایندهای زیربنایی آن است.

ترکیب هوش مصنوعی با فناوری‌های دیگری مانند اینترنت اشیا^۲، نسل پنجم تلفن همراه^۳ و رایانش لبه‌ای^۴ یا رایانش مه^۵، امکان جمع‌آوری و خلق مجموعه‌های بزرگی از داده‌ها را فراهم می‌کند که آماده تغذیه سامانه‌های هوش مصنوعی هستند. با افزایش اثرگذاری و قابلیت‌های فنی هوش مصنوعی، ریسک‌ها و فرصت‌های استفاده از آن هم تکامل می‌یابند. هدف این گزارش کمک به بیمه‌گران برای آگاهی از این حوزه است که به سرعت در حال پیشرفت و توسعه می‌باشد. در این گزارش ریسک‌های فناوری و نیز پتانسیل کاربرد هوش مصنوعی و فرصت‌های آن در کسب‌وکارهای صنعت بیمه بیان می‌شوند.

ریسک‌های فناوری

- **عدم اطمینان:** ریسک‌های زیادی در پذیرش و استفاده از هوش مصنوعی وجود دارند؛ زیرا این فناوری هنوز ناشناخته است و نگرانی‌هایی در میان ذینفعان آن، از جمله قانون‌گذاران و نهادهای نظارتی ایجاد می‌کند.
- **اخلاقیات:** یادگیری ماشین به مجموعه‌های بزرگی از داده‌ها وابسته است که دانش از آن‌ها استخراج می‌شود. این مجموعه داده‌ها به نحوی اثربخش تبدیل به کدگذاری تاریخی می‌شوند. اما در این تاریخ‌ها، پیش‌داوری‌ها، سوگیری‌ها و

۱. این گزارش ترجمه‌ای است از: Lloyd's (2019). Taking control: artificial intelligence and insurance

۲. Internet of Things یا به اختصار IoT یک فناوری کم‌هزینه و بسیار در دسترس است و شامل عمل‌گرها و حس‌گرهایی است که در سطح وسیعی نصب شده‌اند.

3. 5G

۴. Edge Computing: در این فناوری هر دستگاهی دارای هوش مصنوعی و قدرت پردازش بسیار وسیعی است. به عبارت دیگر هر دستگاه تبدیل به مرکز داده‌های خودش می‌شود. تا کنون به این صورت بوده است که تمام دستگاه‌های هوشمند که از فضای ابری جهت ذخیره‌سازی و پردازش اطلاعاتشان استفاده می‌کردند، یعنی تمام اطلاعات را به ابرها ارسال می‌کرده‌اند و پس از پردازش توسط ابرها نتیجه به دستگاه‌ها ارسال می‌شد. با توجه به حجم گسترده دستگاه‌های هوشمند و حجم زیاد اطلاعاتی که بین این دستگاه‌ها رد و بدل می‌شود از رایانش مرزی یا لبه‌ای استفاده می‌کنند. در حال حاضر در رایانش مرزی هر دستگاه وظیفه پردازش اطلاعات خودش را دارد و تنها موارد مهم و اصلی به ابرها ارسال می‌شوند. [م]

۵. Fog Computing: همان رایانش مرزی است. [م]



بی‌عدالتی‌های اجتماعی موجود هستند که برای سالیان زیادی وجود داشته‌اند. پیامد این موضوع آن است که هوش مصنوعی به اطلاعاتی وابسته می‌شود که بر اساس آن آموزش می‌بیند و ممکن است علیه منافع انسان عمل کند.

- **اعتماد و شفافیت:** محل برخورد سوگیری و جعبه سیاه شبکه‌های عصبی عمیق^۱ (که در آن‌ها نمی‌توان نحوه رسیدن به یک نتیجه را توضیح داد)، دارای مفاهیم ضمنی متعددی برای جامعه و صنعت بیمه است. شفافیت فرایند تصمیم‌گیری و توانایی درک نحوه یک نتیجه‌گیری خاص توسط هوش مصنوعی، کلیدی برای توسعه اعتماد در فناوری خواهد بود. جهت اطمینان از اعتماد و شفافیت برای عموم و هم‌ترازی گسترده‌تر با تنوع کامل در جهان، لازم است برای فرایندهای تصمیم‌گیری، پاسخ‌گویی روشن‌تری وجود داشته باشد، بدون این که عملکرد هوش مصنوعی قربانی شود.

- **مسئولیت:** چالش قابل توجه دیگر، وضعیت حقوقی سامانه‌ها و خدمات هوش مصنوعی است. با این که برای تشویق و انگیزش مهندسان پشت سیستم، کارهای زیادی می‌توان انجام داد تا عواقب و تبعات تصمیمات خود در طراحی شبکه‌های عصبی و انتخاب مجموعه‌های داده‌ها برای یادگیری را مد نظر داشته باشند، اما باید چارچوب مستحکمی برای تعریف مسئولیت وجود داشته باشد. در کل، تولیدکننده هر محصولی، مسئولیت نقایصی را که موجب زیان به کاربران می‌شود، بر عهده دارد. با این حال، در مورد هوش مصنوعی (به ویژه روش‌های هوش مصنوعی قدرتمند)، تصمیمات اخذشده، پیامد طراحی نیستند، بلکه نتیجه‌ی تفسیر واقعیت توسط یک ماشین می‌باشند.

- **امنیت:** شایان توجه است که هوش مصنوعی، بُعد جدیدی را هم در امنیت سایبری شکل می‌دهد؛ زیرا می‌تواند یاد بگیرد و بسیار سریع‌تر از روش‌های سنتی مورد استفاده توسط محصولات قدیمی مرتبط با امنیت، به تهدیدات پاسخ دهد. با این حال، عجیب نیست که مجرمان سایبری هم شروع به استفاده از هوش مصنوعی برای یادگیری نحوه انجام همه مراحل یک حمله معمولی کرده‌اند، یعنی از شناسایی تا انجام یک حمله خاص. با پیچیده‌تر شدن و افزایش پیوندهای داخلی این سامانه‌ها، رخنه‌های امنیت سایبری برای روش‌های هوش مصنوعی به‌هم پیوسته، اهمیت بیشتری می‌یابند و احتمالاً تاثیر اقتصادی نظام‌مند عمیقی خواهند داشت.

تاثیر هوش مصنوعی بر رشته فعالیت‌های بیمه‌ای

هوش مصنوعی می‌تواند بر رشته فعالیت‌های بیمه‌ای متعددی تاثیر بگذارد. این گزارش، مفاهیم ضمنی زیادی را برای فراخوان^۲ و مسئولیت محصول، مسئولیت شخص ثالث و سایل نقلیه، جبران غرامت حرفه‌ای، مسئولیت پزشکی و ریسک‌های سایبری، تضمین درستکاری و ریسک‌های سیاسی یافته است.

مسئولیت محصول و فراخوان محصول

- راه‌های زیادی وجود دارند که یک محصول متکی بر هوش مصنوعی، بر مسئولیت محصول و فراخوان محصول تاثیر بگذارد:

1. Deep Neural Networks

۲. منظور فراخوان برای برگشت محصولاتی است که دارای ایرادات و مشکلاتی هستند. [م]



- موقعیت جدیدی که هوش مصنوعی با آن مواجه می‌شود، ولی برای آن برنامه‌ریزی نشده و می‌تواند موجب از کار افتادن یا نقص عملکرد شود.
 - نقایص محصول که می‌تواند ناشی از خطاهای ارتباطی میان دو ماشین یا میان ماشین و زیرساخت باشد.
- در کل، اگر بخشی که تحت تاثیر قرار گرفته از هوش مصنوعی به‌طور وسیع استفاده کند، فراخوان‌ها می‌توانند بزرگ‌تر و پیچیده‌تر شوند.
- از لحاظ مسئولیت، ماشین‌های هوش مصنوعی خودشان نمی‌توانند مسئول بی‌دقتی و از قلم افتادگی‌هایی باشند که موجب خسارت به اشخاص ثالث می‌شوند، بلکه با بهتر شدن آن‌ها در یادگیری و تصمیم‌گیری، این سوال پیش می‌آید که «وقتی یک ماشین مرتکب خطا می‌شود، چه کسی مسئول است؟».
- تولیدکنندگان/فروشنده‌گان و طراحان/تامین‌کنندگان هوش مصنوعی و خریداران/کاربران هوش مصنوعی، می‌توانند توسط دادگاه مسئول و مقصر شناخته شوند. آن‌چه گفته شد می‌تواند به آن معنا باشد که شرکت‌ها در آینده باید تضمین‌ها، غرامت‌ها و محدودیت‌های بیشتری را در قراردادهایشان لحاظ کنند تا بتوانند ریسک مسئولیت را کنترل نمایند.

مسئولیت شخص ثالث وسایل نقلیه

- در آینده، به سبب جابجایی احتمالی مسئولیت از راننده به وسایل نقلیه خودکار و در نتیجه به تولیدکننده، تخصیص و پوشش مسئولیت چالشی‌تر خواهد شد.
- پتانسیل قابل توجهی هم برای واکنش‌های حقوقی و قانونی در کشورهای مختلف وجود دارد که یک چشم‌انداز پیچیده بین‌المللی برای ریسک ایجاد می‌کند.

خطاهای پزشکی

- راه‌های زیادی برای بروز خطای پزشکی وجود دارند. هوش مصنوعی برای کمک به تشخیص و عارضه‌یابی مورد استفاده قرار می‌گیرد و در حوزه‌هایی مانند رادیوگرافی کاربرد دارد (Lee et al., 2017). اگر خطایی منجر به تشخیص اشتباه یا نتایج مثبت کاذب شود، خطای پزشکی تلقی شود.
- حتی اگر هوش مصنوعی به عنوان کمکی برای نظام ارجاع استفاده شود، چنان‌چه این ارجاعات موجب بررسی‌ها و انجام رویه‌هایی شوند که غیرضروری بوده‌اند، منجر به پیامدهای ضعیف‌تری برای بیمار شده و آنگاه مسئولیت‌هایی ایجاد می‌کنند.

سایبر

- با پیشرفت فناوری روبات‌های سخنگو، جدا دانستن انسان‌ها و هوش مصنوعی دشوارتر می‌شود. این وضعیت، راه را برای انجام فیشینگ^۱ در مقیاس وسیع‌تر بازتر کرده و تشخیص آن را دشوارتر می‌کند.

- چنین شرایطی، درباره این که چه نوع بیمه‌ای برای حمایت در برابر زیان‌های حاصل از این نوع از حملات سایبری موجود هستند، سؤالاتی را ایجاد می‌کند.

- مسئولیت شخص ثالث هم می‌تواند پیچیده‌تر شود. اگر کنترل یک روبات سخن‌گو برای انجام یک حمله به دست گرفته شود، چنان‌چه مالک آن نتواند ثابت کند که شاخص‌های امنیتی معقولی را برای جلوگیری از چنین حملاتی در نظر گرفته است، آن‌گاه مسئول خواهد بود.

تضمین درستکاری

- فعالیت متقلبانه از سوی کارکنان هم می‌تواند با هر یک از روش‌هایی که در بخش سایبری گفته شد، تشدید شود. بیمه‌گران تضمین درستکاری باید توجه داشته باشند که تلقب، بیشتر توسط کارکنانی انجام می‌شود که به سامانه‌های فناوری اطلاعات دسترسی دارند و نه کارکنان بخش مالی.

- وقتی پای تقلب در میان باشد، ظهور جعل‌های عمیق^۱ هم نگرانی‌هایی را ایجاد می‌کند. جعل‌های عمیق، سامانه‌های هوش مصنوعی هستند که قابلیت تولید فیلم‌ها و صوت‌های واقعی را دارند. این موضوع برای موارد جبران غرامت تقلبی، نگرانی‌هایی را ایجاد می‌کند و می‌تواند هم‌زمان با روبات‌های سخنگو مورد استفاده قرار گیرد و اعتبار و اعتماد را افزایش دهد.

ریسک‌های سیاسی

- مسلح‌سازی^۲ هوش مصنوعی می‌تواند شکل‌های متعددی به خود بگیرد، از جمله نادرستی داده‌ها، انتخاب تورش‌دار داده‌ها و استفاده غیرقانونی از داده‌ها که منجر به پیامدهای مختلفی از جمله پروپاگاندا، تغییرات رفتاری و فریبکاری می‌شود.

- از منظر پوشش ریسک سیاسی، هوش مصنوعی می‌تواند در ایجاد یا تشدید رویدادهای سیاسی، جنگ‌ها، اقدامات تروریستی، آشوب‌های مدنی و سایر شکل‌های جرایم سیاسی، هم در بازارهای توسعه‌یافته و هم در حال توسعه، سهم دارد.

- بی‌ثباتی‌های ناشی از انتقال هوش انسانی به ماشین^۳ و برهم‌زنندگی‌های اقتصادی حاصله، پتانسیل وارد کردن نیروهایی را برای برگزاری تظاهرات و ایجاد تلاطم دارد که این‌ها می‌توانند منجر به دخالت حکومت، تبعیض‌گزینی و وقفه در کسب‌وکارها شوند.

فرصت‌های کسب‌وکار برای بیمه‌گران

- در کل، هر شرکتی که سامانه‌های الگوریتم‌محور را به شرکت‌های غنی از داده‌ها ارائه می‌کند (مانند تشخیص تقلب در فروش آنلاین)، ممکن است به دنبال پوشش بیمه در برابر ریسک الگوریتم‌هایی باشد که موجب تصمیمات اشتباه

1. Deep Fakes
2. Weaponisation
3. Autonomation

می‌شوند و نیز تاثیری که بر مشتریان شرکت‌های هوش مصنوعی می‌گذارند.

- شرکت‌های جدیدی در حوزه دفاع در برابر اطلاعات دروغ، در حال ظهور هستند تا فناوری‌هایی را برای فیلتر کردن اخبار جعلی، تشخیص و حذف روبات‌های اوباش مجازی^۱ و تعیین صحت و سقم و اعتبار تصاویر و ویدئوها تهیه کنند. بیمه‌گران می‌توانند بررسی کنند که چه نوع محصولاتی (مثلا جبران غرامت حرفه‌ای و محصولات سایبری) و چه به شکلی می‌توانند برای این کسب‌وکارهای جدید مفید باشند.

- با توسعه این حوزه و افزایش کاربردهای آن، فرصت‌هایی برای ارائه خدمات مدیریت ریسک فراهم می‌شوند. اگرچه پیشرفت هوش مصنوعی تقریباً در مراحل اولیه خود به سر می‌برد، اما متخصصان دانش هوش مصنوعی از قبل با تقاضای زیادی از سوی شرکت‌هایی مواجه هستند که در حال تلاش برای مدیریت ریسک خود می‌باشند. شرکت‌های خدمات تخصصی در حال ظهور هستند که هدفشان جلوگیری از خسارت است. با افزایش پیچیدگی‌ها، تقاضا برای این خدمات افزایش خواهد یافت.

- بیمه‌گران با اتخاذ نقش رهبری در این فضا، دانش مورد نیاز برای راهنمایی درباره بهترین شیوه‌های هوش مصنوعی را در اختیار بیمه‌شدگان قرار خواهند داد و بدین وسیله زیست بوم هوش مصنوعی را که در آن کار می‌کنند، شکل و توسعه می‌دهند.

استفاده از هوش مصنوعی برای بهبود فرایندها

از منظر عملیاتی، شرکت‌های بیمه از قبل در حال بهره‌گیری از پتانسیل هوش مصنوعی برای ارائه ارزش به روشی کارا تر بوده‌اند که نمونه‌های آن شامل موارد زیر است:

- خدمات مشتری: روبات‌های سخن‌گو که می‌توانند پیشنهاد محصول بدهند یا محصولات را شخصی‌سازی کنند، به شکایات رسیدگی نمایند، ارتباطات با مشتریان را بهبود دهند و تراکنش‌های ساده را پردازش کنند.

- صدور بیمه: این کار می‌تواند به وسیله هوش مصنوعی، تقویت شده و سرعت آن افزایش یابد. مدل‌های مبتنی بر هوش مصنوعی می‌توانند نرخ‌های حق بیمه را بر اساس ارزیابی‌های ریسک تاریخی تعیین کرده و مبتنی بر پروفایل ریسک مشتری، مظنه‌های سفارشی بدهند.

- تشخیص تقلب: امروزه تشخیص تقلب، عملیاتی است که بیشتر توسط انسان انجام می‌شود، اما می‌تواند با استفاده از هوش مصنوعی، خودکارسازی شود. تشخیص تقلب را می‌توان با افزودن داده‌های خارجی از سوابق تاریخی، حسگرها و تصاویر، خودکارتر کرد تا هزینه‌های تعمیر و استهلاک بهتر تخمین زده شوند.

- دعاوی خسارت: هوش مصنوعی می‌تواند با خودکارسازی تشخیص تصویر، جستجو در پایگاه‌های بزرگ داده برای دعاوی مشابه، تعداد دعاوی خسارتی را که نیازمند تحلیل انسان هستند کاهش دهد و به این طریق از تقلب جلوگیری کند، اعتبار آن‌ها را بسنجد و مشتریان را از پیشرفت پردازش خساراتشان آگاه سازد.

- مدل سازی، مدیریت ریسک و قیمت گذاری: با هوش مصنوعی و ترکیب آن با فناوری های دیگری نظیر اینترنت اشیا، بیمه گران و ارائه دهندگان مدل، به داده هایی دسترسی خواهند داشت که می توانند با آن ها مدل هایی را تغذیه کنند که یاد می گیرند و با سرعتی بیشتر از همتایان دستی خود، انطباق می یابند. این وضعیت مثلاً می تواند نرخ های بهره را بر اساس ارتباطات بانک مرکزی پیش بینی کند یا به پیش بینی داده های جافتاده در یک مجموعه داده کمک کند.
- توسعه کسب و کار: سامانه های هوش مصنوعی می توانند فروش جانبی و نرخ های تغییر محصول را بهبود دهند. می توانند محصولاتی متناسب سازی شده را پیشنهاد داده و به شناسایی مشتریان جدید کمک کنند.

نتایج

- با افزایش به کارگیری هوش مصنوعی در سامانه ها و فرایندهای سازمان ها، آن ها به بیمه ای نیاز خواهند یافت که در برابر گستره وسیعی از ریسک ها از آن ها محافظت کند.
- بیمه گران باید با توسعه مدل ها و محصولات مناسب، به این فناوری واکنش نشان دهند.
- بیمه گران می توانند روش های کاری حال حاضر خود را با استفاده از هوش مصنوعی در همه بخش های زنجیره ارزش، از اولین اعلام تا تسویه خسارات و جلوگیری از ریسک، بهبود بخشند.
- با این که هوش مصنوعی دارای پتانسیل برای توسعه کسب و کار و افزایش کارایی های عملیاتی است، اما بیمه گران باید از ریسک های مرتبط با استفاده از این فناوری جدید که به سرعت در حال رشد است، آگاه باشند.

افق زمانی

- در کوتاه مدت، می توان انتظار داشت که شاهد پذیرش روزافزون هوش مصنوعی با سرعتی نسبتاً محدود باشیم و مشکلاتی در زمینه پیشنهادات و مدل سازی رفتار و بهینه سازی آن به وجود بیاید.
- برای میان مدت، مسائل پیچیده تری بروز می کنند که از آن جمله می توان به خودروهای بدون راننده و خدمات کارکنان به مشتری اشاره کرد.
- در بلندمدت تر، شاهد جابجایی به سمت روش های هوش مصنوعی قدرتمند با قابلیت های سامانه های هوشمندانه ای خواهیم بود که از هوش انسان قابل تشخیص نیستند.

بینش های قابل اجرا

- بر اساس تحقیقات، ۱۰ بینش قابل اجرا برای بیمه گران ارائه می شوند:
- (۱) توجه به سوگیری: سامانه های هوش مصنوعی، به ویژه راهکارهای یادگیری ماشین، نیازمند مجموعه داده های جامع و معتبری هستند که بر مبنای آن ها یاد بگیرند. بینش های حیاتی تر برای منابع داده های مورد استفاده جهت توسعه سامانه های هوش مصنوعی، با توجه به تاثیر اجتماعی و سیاسی که ناشی از فرایندهای انتخاب هستند، به کار گرفته می شود.
- یکی از سوالاتی که می توان در این زمینه پرسید این است که «آیا با انتخاب، خطا یا به صورت تصادفی، سوگیری های اجتماعی هم وجود دارند؟» اگر وجود دارند، پس در الگوریتم ها و انتخاب های آن ها هم تورش هایی وجود خواهند داشت.

۲) فکر کردن به کارگزاری یا دلالتی داده: مناطق و فرهنگ‌های مختلف جهان، دارای رویکردهای متفاوتی در کاربردهای داده‌ها هستند؛ اما رویکردشان در به دست آوردن، به اشتراک گذاری، تجارت و استفاده از داده‌ها هر چه که باشد، این داده‌ها افق‌های هوش مصنوعی را توسعه خواهند داد.

مثلا در انگلستان، برنامه‌هایی برای توسعه چارچوبی جهت ذخیره‌سازی و به اشتراک گذاری دسترسی به داده‌ها در میان سازمان‌ها و دامنه‌ها وجود دارند. هر چه مجموعه داده‌ها بزرگ‌تر باشند، بیمه‌گران بهتر می‌توانند ریسک‌ها را فهمیده و قیمت گذاری کنند.

۳) ارائه شفافیت فنی و توضیح‌پذیری: آگاهی رو به رشدی از این موضوع وجود دارد که سامانه‌های هوش مصنوعی، به ویژه در موقعیت‌های بحرانی باید از نظر فنی، شفاف و قابل توضیح باشند. اگر هوش مصنوعی بخواهد مورد اعتماد قرار گرفته و باور شود، شهروندان و مصرف‌کنندگان، این ویژگی‌ها را طلب خواهند کرد.

۴) طراحی از طریق همکاری: ریشه‌های هوش مصنوعی به عنوان زیررشته‌ای از علم رایانه، منجر به فناوری‌هایی شده است که در یک جامعه کوچک خلق شده‌اند. آینده سامانه‌های عملی هوش مصنوعی از مشارکت اجتماع‌هایی وسیع‌تر و متنوع‌تر با دیدگاه‌های علمی و افق‌های اجتماعی گسترده‌تر بهره خواهد برد.

۵) تحلیل و پایش: چه نگاهمان به سامانه‌های موروثی باشد و چه به طراحی‌های جدید، راهکارهای هوش مصنوعی باید در طول زمان پایش، تحلیل و بازبینی شوند.

سوالات

برخی از سوالاتی که می‌توان در زمینه هوش مصنوعی پرسید این است که «آیا دنیا آنقدری تغییر کرده است که سوگیری‌های درون آن را تشدید یا جبران کند؟ جوامع و فرهنگ‌های مختلف مزایای سامانه هوش مصنوعی را چگونه تجربه می‌کنند؟ آیا خطای مجموعه آزمون^۱ رو به افزایش است؟».

۶) پشتیبانی از مقررات و چارچوب‌های اخلاقی: راهکارهای منصفانه‌تر و نمونه‌تر هوش مصنوعی، در نتیجه‌ی مقررات و چارچوب‌های اخلاقی دقیق و قابل اجرا، ارائه خواهند شد که توسط جامعه، هم‌تراز با تنوع کامل دنیای واقعی توسعه می‌یابند.

۷) افزایش آگاهی: هوش مصنوعی باید بیشتر باز شود. نه این که هر کسی بتواند در هر طرحی درگیر شود، اما اگر هوش مصنوعی قصد دارد که جهان را تغییر داده و بهتر کند، این هدف تنها در صورتی قابل دستیابی است که صداهای بیشتری در مباحثات شنیده شوند.

۸) سرمایه‌گذاری در مهارت‌ها و شایستگی‌ها: دانش تخصصی در کنار همکاری، رهبری و ارتباطات، همگی برای متخصصان و رهبران لازم هستند تا سایر مسائل و فرصت‌ها را مورد توجه قرار دهند و در جذب و بازآموزی مهارت‌های نیروی کار فعلی، سرمایه‌گذاری کنند.

۱. Test set error: خطایی است که هنگام راه‌اندازی یک مدل آموزش داده‌شده روی یک مجموعه داده رخ می‌دهد که قبلاً هرگز در معرض آن قرار نگرفته است. افزایش خطای مجموعه آزمون نشان می‌دهد که مدل و واقعیت در حال تغییر هستند.



بخش بزرگی از موفقیت‌ها یا شکست‌های پروژه‌های هوش مصنوعی، ناشی از نحوه مدیریت خوب یا بد فرایندهای تغییر در سازمان، واحد یا تیم است. فناوری هوش مصنوعی حتی اگر عملکرد بسیار خوبی داشته باشد، اما اگر کارکنانی که از آن استفاده می‌کنند، مفهوم آن را از ابتدا به خوبی متوجه نشوند، پشتیبانی از این پروژه‌ها به احتمال بسیار زیادی با شکست مواجه می‌شود.

(۹) مشارکت با اینشورتک: بسیاری از پیشرفت‌های فناوریانه در هوش مصنوعی، تحت هدایت شرکت‌های نوآفرین^۱ کوچک قرار دارند که اغلب دارای پیوندهای قوی با دانشگاه‌ها هستند. این شرکت‌ها در حال خروج از آزمایشگاه‌های خود هستند تا پیشرفت‌های هوش مصنوعی را عملیاتی کنند و ابزارهایی عالی را برای پرورش متخصصان هوش مصنوعی و پیوندهایی در شرکت‌های بزرگ‌تر ارائه دهند.

(۱۰) بررسی تاثیر هوش مصنوعی بر محصولات و فرایندها: بیمه‌گران باید بررسی کنند که مشتریان‌شان چگونه هوش مصنوعی را خواهند پذیرفت تا راهکارهای جدید را ارائه کنند و چگونه ریسک‌های حاصل از هوش مصنوعی بر رشته فعالیت‌های فعلی تاثیر خواهد گذاشت.

هوش مصنوعی چیست؟

چیستی هوش مصنوعی

هوش مصنوعی امروزه در همه جا حاضر است. از قلب‌های هنری گرفته تا بهداشت و سلامت و از حمل‌ونقل تا کشاورزی. بخشی نیست که به دنبال بهره‌گیری از منافع حاصل از روش‌های و رویکردهایی نباشد. در واقع از کنفرانس دارتموث در سال ۱۹۵۶^۲ ظاهر شدند و آغاز این حوزه تحقیقاتی از آن‌جا بود.

در طول مدت زمان مداخله، هوش مصنوعی مجموعه‌هایی از فراز و فرودها را تجربه کرده است؛ فراز در زمانی که جریانات پولی برای انجام تحقیقات سرازیر شدند و هوش مصنوعی به عنوان پیشنهاددهنده راهکارهای هیجان‌انگیز برای برخی چالش‌ها شناخته شد و فرود برای زمانی که راهکارهای کاربردی آن، دور و بعید به نظر آمدند.

این گزارش نقطه شروعی برای درک بهتر ابعاد کلیدی هوش مصنوعی و نیز ریسک‌ها و فرصت‌های آن است. سوالات غالبی که در این گزارش به آن‌ها پرداخته می‌شود شامل ریسک‌ها و فرصت‌هایی که هوش مصنوعی ایجاد می‌کند و نحوه ارتباط آن‌ها با بیمه است.

این گزارش از طریق یک فرایند تحقیق نظام‌مند در دو مرحله انجام شده است: مرور ادبیات و سپس استخراج نظرات متخصصان دانشگاهی، فناوری، اقتصاد و بیمه‌گران.

اما دنیا در آینده چگونه به تاثیر هوش مصنوعی نگاه می‌کند؟ هوش مصنوعی هم‌زمان هیجان‌آمیز و همراه با نگرانی دیده

1. Startups
2. Dartmouth Conference of 1956

می‌شود؛ زیرا از قبل به بسیاری از ابعاد زندگی ما وارد شده یا دست‌درازی کرده است. بخش بیمه هم‌چنان شاهد نوآوری‌هایی در حوزه‌های مختلفی مانند قیمت‌گذاری، رسیدگی به خسارات و تشخیص تقلب خواهد بود.

در کل، دنیای بیمه در حال مواجه شدن با دامنه‌ای از ریسک‌های مرتبط با هوش مصنوعی است که به همراه منابع حاصل از فرصت‌های نوآوری، بهره‌وری را افزایش می‌دهند که دامنه آن از محصولات آماده مصرف تا اثبات مفهوم^۱ از طریق آزمایشگاه‌های نوآوری و شراکت با شرکت‌های نوآفرین را در بر می‌گیرد.

در هر حال، هوش مصنوعی جدید نیست و در طول ۶۰ سال گذشته، حضور قابل توجهی در دیدگاه‌های آکادمیک داشته است. با این وجود، رشد سریع و برهم‌زننده پیاده‌سازی هوش مصنوعی در سال‌های اخیر، آگاهی‌های زیادی را درباره پیچیدگی‌های آن ایجاد کرده و چالش‌های اخلاقی، حقوقی و اجتماعی متعددی را برانگیخته است که فراتر از پیچیدگی‌های فناوریانه روش‌ها و فرایندهای اساسی آن هستند.

از ترکیب هوش مصنوعی با سایر فناوری‌ها مانند اینترنت اشیا (حسگرها و به کاراندازنده‌های کم‌هزینه و بسیار در دسترس که به وفور نصب شده‌اند)، ۵G و رایانش مرزی یا مه، امکان گردآوری و خلق مجموعه‌های بزرگی از داده‌ها فراهم می‌شود که آماده تغذیه سامانه‌های هوش مصنوعی هستند.

با این که همه این‌ها امکان انجام کارهای هوشمندتر و سریع‌تری را فراهم می‌کنند، اما نگرانی‌هایی در میان ذینفعانی چون قانون‌گذاران وجود دارد که ناشی از ماهیت جعبه سیاه این راهکارها است (Brown et al., 2017). درک و توضیح نحوه عملکرد این جعبه‌های سیاه، چالشی کلیدی برای آینده محسوب می‌شود.

مجموعه اصطلاحات فنی رایج

برای کمک به فهم هوش مصنوعی، اصطلاحات کلیدی زیر می‌توانند برای آشنایی با موضوع استفاده شوند. تعریف مفاهیمی که به شرح تغییرات حاصل از هوش مصنوعی کمک می‌کنند، در گذار از سامانه‌های خودکار به سمت سامانه‌های خودمختار یا مستقل مفید هستند. سامانه خودکار^۲ به معنی سامانه‌هایی است که از دستورالعمل‌های ازپیش‌برنامه‌ریزی‌شده پیروی می‌کنند و خودمختار یا مستقل به معنی سامانه‌هایی است که مستقلاً عمل می‌کنند. این تمایز در قلب چالش‌های اجتماعی و ریسک‌های نوظهور قرار دارد.

قاب ۱- اصطلاحات کلیدی

هوش مصنوعی

هوش مصنوعی یک ابرواژه برای مجموعه‌هایی رو به تکامل از فناوری‌هایی است که در چند دهه گذشته ظهور و افول کرده‌اند. به همین دلیل، تعریف ساده‌ای از هوش مصنوعی وجود ندارد که در همه زمینه‌ها جواب دهد.

لغت‌نامه مریام-وبستر^۱، هوش مصنوعی را به صورت زیر تعریف می‌کند:

۱. شاخه‌ای از علم رایانه که مرتبط با شبیه‌سازی رفتار هوشمند در رایانه‌ها است.

۲. قابلیت یک ماشین در تقلید رفتار هوشمند انسانی.

تفاوت‌های تعاریف هوش مصنوعی، بر مبنای چیزی است که هوش مصنوعی برای به دست آوردن آن کار می‌کند. مثلاً، شرکت‌های فناوری و نوآفرین، هوش مصنوعی را به عنوان علمی تعریف می‌کنند که ماشین‌های هوشمندی را خلق می‌کنند که قابلیت انجام وظایف را به صورت بلادرنگ در سطح تخصص انسانی دارند.

برخی وظایفی که روش‌های هوش مصنوعی با کیفیت به خوبی می‌توانند انجام دهند، به شرح زیر هستند:

- پردازش مقادیر بزرگی از داده‌ها با سرعت بالاتر و بنابراین ظرفیتی بیش از انسان‌ها؛
- یادگیری از مجموعه داده‌های نمونه و انطباق با آن برای حل مشکلات جدید بر مبنای آنچه که از مجموعه داده‌ها استخراج می‌شوند؛

- شناسایی اشیاء و ارتباط آن‌ها با یک موقعیت خاص؛

- استخراج وضعیت آینده یک شی یا موقعیت؛

- شناسایی تصمیم بهینه بر اساس گذشته، حال و آینده‌ی استنباط شده.

یادگیری ماشین

یادگیری ماشین زیرمجموعه‌ای از هوش مصنوعی است که اغلب از تکنیک‌های آماری برای دادن توانایی یادگیری به رایانه (مثلاً بهبود تدریجی عملکرد در یک وظیفه خاص) به وسیله داده‌ها استفاده می‌کند، بدون این که صریحاً برای انجام این کار برنامه‌ریزی شده باشد.

یادگیری عمیق

یادگیری عمیق شاخه‌ای از یادگیری ماشین است که از شبکه‌های عصبی پیچیده برای تصمیم‌گیری استفاده می‌کند و اغلب متشکل از چندین میلیون عصب است که در دهه‌ها یا صدها ساختار لایه‌بندی شده قرار گرفته‌اند. شکل رایج شبکه عصبی عمیق، شبکه عصبی پیچشی^۲ است که می‌تواند سیگنال‌های تک‌بعدی (مانند داده‌های سری زمانی یا صوت) یا سیگنال‌های دو بعدی (مانند عکس و ویدئو) را پردازش کند تا وجود یک شی یا ویژگی را در سیگنال تشخیص دهند. این شبکه‌های عصبی مصنوعی مانند مغز انسان مدل می‌شوند و در آن‌ها گره‌های عصبی مانند یک تار عنکبوت به هم متصل هستند. برای اطلاعات بیشتر، به قاب ۲ مراجعه شود.

در این گزارش، هوش مصنوعی به عنوان ابرواژه‌ای استفاده می‌شود که همه موارد بالا را پوشش می‌دهد (مگر این که چیز دیگری گفته شود).

یادگیری بانظارت و بدون نظارت

- دو رویکرد اصلی وجود دارند که در یادگیری ماشین استفاده می‌شوند: بانظارت و بدون نظارت.^۱
- رویکرد بانظارت زمینه‌ای را تعریف می‌کند که دارای اهداف قبلی در رابطه با مقادیری است که باید برای داده‌های مورد نظر به کار گرفته شوند. بنابراین، هدف یادگیری بانظارت ایجاد سامانه‌ای است که می‌تواند به‌نحو قابل اعتمادی الگوها و طبقه‌بندی‌های از پیش تعیین شده را تشخیص دهند. در زمینه بیمه، این یعنی نرخ‌های حق بیمه که می‌توانند بر اساس ویژگی‌های دقیق موجودیتی که به دنبال خرید بیمه است، تعیین شوند.
 - رویکرد بدون نظارت، به جستجوها برای خوشه‌بندی اشاره دارد که طی آن در داده‌ها جستجو انجام می‌شود تا بدون برچسب‌ها و طبقه‌بندی‌های از پیش تعریف شده، الگوها و ویژگی‌های ساختاری پیدا شوند. در حوزه بیمه، این رویکرد می‌تواند به معنی جستجو برای یافتن الگوهایی در میان مشتریان باشد که ممکن است منجر به ارجاع محصولات^۲ شود و نیز مشتریانی که ممکن است در کنار یک محصول بیمه‌ای به سایر محصولات هم علاقمند باشند.

جدول ۱- از هوش مصنوعی با نظارت تا بدون نظارت

بدون نظارت	نیمه نظارت شده	بانظارت
- داده‌های ورودی (x) را در اختیار دارید و هیچ متغیر خروجی مرتبطی ندارید. - هدف، مدل کردن ساختار اساسی برای یادگیری بیشتر درباره داده‌ها است. - الگوریتم‌ها به حال خودشان رها می‌شوند تا ساختار مورد نظر در داده‌ها را کشف و ارائه کنند.	- این رویکرد معمولاً از مقدار کمی از داده‌های برچسب‌گذاری شده^۳ با مقادیر زیادی از داده‌های بدون برچسب استفاده می‌کند.	- متغیرهای ورودی (x) و متغیرهای خروجی Y را در اختیار دارید و از الگوریتمی برای یادگیری تابع نگاشت از ورودی به خروجی استفاده می‌کنید $Y=f(x)$ - هدف این است که تابع نگاشت را تخمین بزنید، به نحوی که وقتی داده‌های ورودی جدید (x) را دارید، بتوانید متغیرهای خروجی (Y) را برای آن داده‌ها پیش‌بینی کنید.

منبع: Lloyd's, 2018

هوش مصنوعی ضعیف (محدود) و قوی (گسترده)

هوش مصنوعی اغلب به دو نوع ضعیف و قوی تقسیم می‌شود. تفاوت در این است که هوش مصنوعی ضعیف یا محدود بر حل یک مساله خاص در حوزه‌های محدودی تمرکز دارد (مانند روبات‌های سخن‌گو برای خدمات مشتری). با هوشمندتر شدن هوش مصنوعی، به واسطه مقدار داده‌های یادگیری شده و توانایی در حل مسائلی که قبلاً هرگز نمی‌توانسته حل کند، می‌توان آن را قوی نامید.

1. Supervised and Unsupervised

۲. به بازاریابی ارجاعی اشاره دارد که به زبان ساده یعنی همان خرید یک مشتری به خاطر نظر یا پیشنهاد دیگران. [مترجم]

3. Labelled Data



هوش مصنوعی ضعیف یا محدود	هوش مصنوعی عمومی
- سامانه‌هایی که فقط می‌توانند مانند انسان‌ها رفتار کنند و کارهای درست انجام دهند. - اما نحوه فکر کردن انسان را توضیح نمی‌دهند. - بر یک وظیفه محدود تمرکز می‌کنند. - مانند «دیپ بلو» آی‌بی‌ام^۱ در سال ۱۹۹۷	- شبیه‌سازی اصیل استدلال انسانی - ساخت سامانه‌هایی که نه تنها فکر می‌کنند، بلکه نحوه فکر کردن انسان را هم توضیح می‌دهند. - آموزش ندیده و بدون نظارت

منبع: Lloyd's, 2018

هوش مصنوعی قوی که قابلیت استدلال فراتر از ظرفیت انسان دارد، در حال حاضر قطعا در قلمروی افسانه‌های علمی قرار می‌گیرد.، اگر چه برای بسیاری از افراد و نهادها، انگیزه‌ای قوی جهت تحقیقات مستمر می‌باشند، اما مسئله این است که آیا چنین سامانه‌هایی امکان‌پذیر یا حتی هدفی مطلوب هستند. سامانه‌های هوشمند امروزی که قابلیت تبدیل شدن به یک شاهکار را دارند، تا همین چند سال پیش چالش بزرگ غیرقابل حلی محسوب می‌شدند. مثلا هوش مصنوعی‌هایی که می‌توانند صحنه‌های بصری را توصیف کنند یا با زبان طبیعی، به سؤالاتی درباره تصاویر پاسخ دهند (پاسخ سوال بصری). این توانایی‌ها نباید با خودآگاهی یا توانایی استدلال، اشتباه گرفته شوند؛ بلکه نشان‌دهنده قابلیت شبکه عصبی برای یادگیری از هزاران کلمه و عبارت محرک هستند تا واسط مناسبی بین انسان و رایانه ایجاد کنند.

در مواجهه با این دستاوردها و اصطلاحات مشترک، وسوسه‌انگیز است که خطوط موازی عمیقی را میان هوش مصنوعی امروزی که حاصل شبکه‌های عصبی عمیق هستند و شبکه‌های عصبی مغز انسان بکشیم. با این که مشابهت‌های ظاهری وجود دارند (یعنی هر دو از عصب‌های به هم متصل تشکیل شده‌اند) اما عملیاتی که توسط اعصاب انجام می‌شود با عملیات معماری‌ها تفاوت‌های قابل توجهی دارند. اعصاب زیستی مانند شبکه‌های عصبی عمیقی که توسط انسان ساخته می‌شوند، در ساختارهای لایه‌ای متوالی قرار ندارند، بلکه در الگوهای نامنظمی به هم متصلند که با گذشت زمان و با شلیک همزمان سلول‌های عصبی به صورت ناهمگن تغییر می‌کنند و زمان‌بندی آن‌ها، ویژگی‌ها و رفتارهای جدیدی را ایجاد می‌کند.

اعصاب زیستی می‌توانند در طول زمان فراموش کنند و در کمیت زیادی رخ دهند. با این که می‌توان به دنبال تقلید از برخی از این ویژگی‌ها در شبکه‌های عصبی عمیق بود، اما دلیلی وجود ندارد که نتیجه بگیریم فناوری شبکه عصبی عمیق در شکل فعلی خود می‌تواند آگاهی مصنوعی^۲ به وجود آورد. با این وجود، دلیلی هم وجود ندارد که شک کنیم پیچیدگی سامانه عصبی زیستی نمی‌تواند روزی در یک شبیه‌سازی رایانه‌ای، مدل شود.

برای درک بهتر اثری که هوش مصنوعی می‌تواند در جهان (از جمله در دنیای بیمه)، در سطحی وسیع داشته باشد و خواهد داشت، باید به تاریخ شصت ساله آن نگاهی بیندازیم. علی‌رغم دوران طولانی آبتستی آن، کاربرد هوش مصنوعی

1. IBM's Deep Blue
2. Artificial Consciousness

اخیرا تاثیر گسترده‌ای در جهان داشته است. قبلا، هوش مصنوعی به عنوان علمی با شاخه‌های نظری و کاربردی دیده می‌شد. به جز کاربردهای نظامی، سایر کاربردهای آن شامل نمایش‌های^۱ اثبات مفهوم یا نمونه‌هایی با تمرکز محدود، مانند بازیکن شطرنج، وسایل نقلیه نیازمند کمک انسان، روبات‌های کارخانه و سایر سامانه‌های خبره^۲ بودند.

تاریخچه‌ای مختصر از پیشرفت هوش مصنوعی و روندهای اصلی آن

هوش مصنوعی به مدت بیش از شصت سال وجود داشته است و صرفا از همین اواخر که تاثیر بالقوه آن بر آینده کار و بسیاری از دیگر ابعاد زندگی ما آشکار شده، تبدیل به بحثی غالب شده است.

تاریخچه هوش مصنوعی را می‌توان در سه موج مورد بررسی قرار داد که هر یک دارای دوره‌های کوتاه اما قابل توجه کمبود تامین مالی بوده‌اند که اغلب از این دوره‌ها با عنوان زمستان‌های هوش مصنوعی یاد می‌شود. این مراحل را می‌توان با رویکردهای فناورانه غالب هم مورد بررسی قرار داد.

- موج اول: اولین موج پیشرفت هوش مصنوعی متمرکز بر رویکردهای نمادینی بود که در آن، مسائل به عنوان مسیری مارپیچ مطرح می‌شدند که باید مورد بررسی و کاوش قرار می‌گرفتند. پس از آن که اولین زمستان هوش مصنوعی تمام شد، عصر جدیدی در دهه ۱۹۷۰ و ۱۹۸۰ آغاز شد و سامانه‌های خبره که به آن‌ها سامانه‌های دانش‌محور^۳ هم می‌گویند، رشد کردند.

این سامانه‌ها مجموعه‌هایی از دانش سطح بالایی را تشکیل دادند که از سامانه‌های خبره گردآوری شده بودند. با شروع رشد کار در شبکه‌های عصبی، دومین زمستان هوش مصنوعی رخ داد و جوهی که به تحقیق اختصاص داده می‌شد، دوباره خشک شدند.

- موج دوم: بعدا در دهه ۱۹۹۰، تحقیقات توسعه یافتند و از سامانه‌های متمرکز سلسه مراتبی دور شدند. تحقیقاتی جایگزین آن‌ها شدند که به عامل‌های همکار^۴ (که مستقلا تعامل می‌کنند، پیام‌ها را به اشتراک می‌گذارند و از تجربه‌شان یاد می‌گیرند)، می‌پرداختند.

به لطف تلاقی عوامل فناورانه‌ای مانند رشد ذخیره ارزان‌قیمت داده‌ها و قدرت پردازش و نیز در دسترس بودن داده‌های زیاد و اتصالات حاصل از اینترنت و اطلاعات توزیع شده، این رویکرد پژوهشی توانست شکوفا شود.

- موج سوم: هوش مصنوعی، از رویکرد نمادین مبتنی بر اطلاعات ساختاریافته، به رویکردهای زیرنمادین^۵ با محتوای

1. Demo

2. Expert Systems

3. Knowledge-based Systems

۴. Collaborating agents: در دانش رایانه، عامل‌های نرم‌افزاری (software agent)، بخشی از یک نرم‌افزارند که جهت کمک به کاربر یا

نرم‌افزاری دیگر در چارچوب روابط واسطه‌ای کار می‌کنند. در واقع کاربران، قدرت تصمیم‌گیری در مورد این که در هر زمانی چه اقدامی باید صورت بگیرد را به عامل‌ها همانند یک واسطه (گماشته) و می‌گذارند. به جز عامل‌های همکار، عامل‌هایی نظیر کاربر، خریدار، مراقبتی و نظارتی و ... هم وجود دارند. [م]

5. Sub symbolic



غیرساختارمند و ساختارهای کنترلی غیرمتمرکز تغییر کرد.

این کاربردهای تخصصی می‌توانند در وظایف فردی خود که از داده‌ها برای آموزش الگوریتم‌های مرتبط استفاده می‌کنند و معمولاً به عنوان هوش مصنوعی محدود شناخته می‌شوند، به برتری برسند؛ هر چند که هدف بزرگ‌تر آن‌ها ارائه هوش مصنوعی عمومی یا هوش عمومی مصنوعی است (که با عنوان هوش مصنوعی کامل / قوی / وسیع هم شناخته می‌شود).

آزمون‌های مختلفی برای رسیدن به این هدف دشوار پیشنهاد شده‌اند که مهم‌ترین آن‌ها آزمون کلاسیک تورینگ^۱ است که توسط آلن تورینگ^۲ در سال ۱۹۵۰ توسعه داده شد تا امکان تقلید ظرفیت شناختی انسان را ارزیابی کند (Turing, 1950).

نکته: نحوه کارکرد یادگیری عمیق

پذیرش سریع هوش مصنوعی در چند سال گذشته در سطح تجاری، با پیشرفت‌های فناورانه در یادگیری ماشین، به ویژه یادگیری عمیق سرعت گرفته که دستاوردهای عملکردی تحول‌آمیزی به همراه دارند و در عین حال به دامنه وسیع‌تری از کاربردها تعمیم می‌یابند. یادگیری عمیق به معنای استفاده از شبکه‌های عصبی پیچیده (یا عمیق) برای حل مسائل یادگیری ماشین مانند طبقه‌بندی (مثلاً درک تصویر یا صحنه) یا رگرسیون (مثلاً برازش مدل) است.

عصب‌ها سنگ بنای شبکه‌های عصبی هستند؛ یعنی شبیه‌سازی دیجیتال سلول‌های عصبی زیستی که اطلاعات سایر اعصاب شبکه را جمع کرده و به آن‌ها اطلاعات می‌دهد. رفتار شبکه عصبی با الگوی اتصال میان عصب‌های آن (معماری شبکه) و نیز نحوه دستکاری اطلاعات توسط هر عصب تعیین می‌شود. معماری شبکه‌های عصبی هنوز نیازمند طراحی دستی توسط انسان است. با این حال، وقتی طراحی صورت بگیرد، رفتار شبکه از طریق اهمیت یا وزنی تعریف می‌شود که هر عصب یاد می‌گیرد تا به اطلاعات دریافتی از مجاوران خود نسبت دهد. این وزن‌ها پارامترهایی هستند که در طول یادگیری آن، توسط شبکه یاد گرفته می‌شوند.

شبکه‌های عصبی معمولاً طراحی می‌شوند تا عصب‌ها را در ساختاری لایه‌بندی شده به هم متصل کنند. شبکه‌های عصبی کلاسیک دهه‌های گذشته، فقط حاوی یک یا دو لایه بودند. برعکس، شبکه‌های عصبی عمیق می‌توانند بیش از صد لایه داشته باشند که به ترتیب ده‌ها میلیون پارامتر وزنی قرار می‌گیرند. نتیجه این است که این شبکه‌ها می‌توانند وظایف پیچیده‌تری را انجام دهند، اما نیازمند حجم بیشتری از داده‌های آموزشی هستند و تلاش رایانه‌ای بیشتری برای یاد دادن به آن‌ها مورد نیاز است. این موضوع بسیار مهم است، چرا که تعداد الگوهای یک شبکه عصبی می‌تواند تشخیص دهد، متناسب با (البته نه به‌طور خطی) تعداد عصب‌ها و بنابراین وزن‌ها در شبکه است.

یادگیری عمیق حوزه‌ای است که سرعت توسعه بالایی دارد و چالش‌های قابل توجهی در آن باقی مانده‌اند که کمترین آن زمان طولانی برای آموزش (که مانع یادگیری بلادرنگ از داده‌ها می‌شود) و نیز نیاز به حجم بسیار بالایی از داده‌ها است. این مشکلات در تکنیک‌های کلاسیک‌تر برای یادگیری ماشین، مانند ماشین‌های بردار پشتیبان^۳ و مدل‌های گرافی که پیش از ظهور یادگیری

1. Turing test

2. Alan Turing

۳. Support Vector Machines (SVMs): یکی از روش‌های یادگیری با نظارت است که از آن برای طبقه‌بندی و رگرسیون استفاده

می‌شود. [م]

ماشین این حوزه را رهبری می‌کردند، رخ نمی‌داد. آن‌ها هنوز هم در چشم‌انداز هوش مصنوعی جایگاهی دارند؛ یعنی در جایی که داده‌های آموزشی تفسیر شده، پراکنده هستند یا وقتی داده‌ها باید از طریق آموزش تدریجی در طول زمان، یاد گرفته شوند. شاید قابل توجه‌ترین مسئله در یادگیری عمیق، توانایی محدود ما برای توضیح یا تفسیر رفتار یک شبکه‌ی آموزش دیده است. به راحتی می‌توان عملکرد شبکه عمیق یا در واقع هر سامانه یادگیری ماشین را کمی سازی کرد؛ حتی سامانه‌ای که صرفاً مقداری از داده‌های آموزشی را نگهداری نموده (و از آن با عنوان بخش داده‌های آزمون^۱ یاد می‌شود) و شبکه آموزش دیده را از لحاظ نگهداری داده‌ها برای سنجش میزان دقت تصمیم اخذ شده، ارزیابی می‌کند.

با این‌که این وضعیت در محیط آزمایشگاهی کاملاً رضایت‌بخش است، اما موفقیت یادگیری ماشین، محرک کاربردهای تجاری وسیعی است و در سامانه‌های مبتنی بر یادگیری ماشین در حجم انبوه، موقعیت‌های زیادی وجود خواهند داشت که برای آن غیرممکن است یک مورد آزمون را پیش‌بینی کند (ناشناخته‌های ناشناخته). مثلاً، چگونه می‌توان یک مجموعه داده آزمون را برای همه شرایط ممکن که یک خودروی بدون راننده با آن مواجه می‌شود، توسعه داد؟ در برخی مراحل، مانند همه نرم‌افزارها، باید بر مبنای پوشش آزمون، با قضاوت شخصی تشخیص دهیم که آیا سامانه برای انتشار آماده است یا خیر.

خلاصه پیشرفت‌های اخیر

پس از انفجاری که در زمان‌های قدیم‌تر در زمینه توسعه هوش مصنوعی رخ داد، این فناوری از دوره‌های طولانی رشد تاب‌آور بهره برده است. سطح پوشش این پیشرفت اکنون از کاربردهای نظامی و تحقیقات دانشگاهی عبور کرده و به بخش خصوصی رسیده است.

رشد و گسترش هوش مصنوعی، پدیده‌ای جهانی است و نقاط کانونی زیادی در دنیای توسعه یافته برای آن وجود دارد؛ به ویژه در چین که تا حدی به دلیل دسترسی به مقادیر بسیار زیادی داده‌های گردآوری شده از جمعیت عظیم خود و نیز تعهد دولت به تامین مالی ۱۵۰ میلیارد دلاری برای چند سال آینده، حضور قابل توجه و رو به رشدی در این حوزه دارد (Reuters, 2018).

تاثیر گسترده هوش مصنوعی در چند سال گذشته ناشی از چهار عامل مکمل بوده است:

- رشد فزاینده و کاهش هزینه توسعه قدرت پردازش رایانه‌ها؛
- دسترسی به منابع سرشار و ارزان‌تر داده‌ها؛
- رشد گسترده اتصال شبکه‌ای در اینترنت؛
- رشد نمایی داده‌های در دسترس (کلان داده‌ها) در سراسر جهان.

نمونه‌هایی از کاربردهای فعلی هوش مصنوعی

هوش مصنوعی پس از سال‌ها ارائه در پروژه‌های نمایشی در آزمایشگاه‌های دانشگاه‌ها، در دامنه وسیعی از فعالیت‌های بخش‌های مختلفی (از جمله بیمه) کاربرد پیدا کرده است.

هوش مصنوعی در بسیاری از سایر ابعاد محیط کار نیز تحول ایجاد می‌کند، به این معنا که خدمات تقویت‌شده با هوش مصنوعی، کار انسان‌ها را ردگیری و پایش می‌کنند.

مثال‌های زیر، نمایی کلی از کاربرد هوش مصنوعی در بخش‌های مختلف را ارائه می‌دهند، اما این فهرست کامل نیست. مثال‌ها به شرح زیر هستند:

- در پزشکی، اسکن‌های دیجیتال می‌توانند با سامانه‌های آموزش‌دیده، پردازش و تحلیل شوند تا تومورها و سایر ناهنجاری‌ها را شناسایی کنند.
- وکلا می‌توانند مقادیر عظیمی از اسناد حقوقی را برای شناسایی الگوهای نقض قانون، پردازش کنند.
- ارائه‌دهندگان خدمات سرگرمی می‌توانند بر مبنای چند انتخاب و بررسی تاریخچه ترجیحات افراد، پیشنهادات شخصی‌سازی‌شده به مشتریان ارائه دهند.
- دوربین‌ها و سایر وسایلی که به‌طور سنتی پیچیده هستند، اکنون می‌توانند به‌طور خودکار خودشان را با نیازهای اپراتورهایی که صلاحیت کمتری دارند، تنظیم کنند.

ریسک‌های هوش مصنوعی

با این‌که هوش مصنوعی در بسیاری از عرصه‌های زندگی مانند بهداشت و درمان، حمل‌ونقل و تولید، منافع را ایجاد می‌کند، اما ریسک‌هایی هم وجود دارند که موجب نگرانی می‌شوند. استیون هاو کینگ^۱ و ایلان ماسک^۲، ریسک‌های مهمی را که برای بشریت وجود دارد، یا در واقع ریسک‌هایی را که در زمان تبدیل شدن هوش مصنوعی محدود به سامانه‌های قدرتمند به وجود می‌آیند، برشمرده‌اند.

این احتمالات فناورانه چالش‌هایی را برای نوع بشر و سیاره زمین ایجاد می‌کنند؛ زیرا سازوکارهای موجود، ما را برای مدیریت مسئولیت‌ها، اخلاقیات، مقررات و تعهدات به چالش می‌کشند. در این گزارش، چهار حوزه ریسک مرتبط با هوش مصنوعی مورد بررسی قرار می‌گیرند که شامل شفافیت و اعتماد، اخلاقیات، تعهدات و مسئولیت‌ها و امنیت می‌باشد.

شفافیت و اعتماد

یکی از اصلی‌ترین نگرانی‌هایی که در مورد هوش مصنوعی وجود دارد، در جایی است که یک جعبه سیاه اثربخش با مجموعه قواعد و شرایط ناشناخته‌ای ایجاد می‌شود که ممکن است آزمون‌های نامناسب یا حتی غیرقانونی داشته باشد تا به نتایج مطلوب برسد.

جعبه سیاه‌های هوش مصنوعی به شکل شبکه‌های عصبی عمیقی ساخته می‌شوند که به وسیله آن‌ها، فرایند یادگیری ماشین، لایه‌هایی از شبکه‌های عصبی را ایجاد می‌کند. بین لایه ورودی و لایه خروجی، لایه‌های متعددی در شبکه عصبی وجود دارند که توسط فرایند یادگیری ماشین، آموزش می‌بینند و بازبینی می‌شوند.

1. Stephen Hawking
2. Elon Musk

شفافیت فرایند تصمیم‌گیری و توانایی همه در فهم نحوه رسیدن هوش مصنوعی به نتایجی خاص، کلیدی برای توسعه اعتماد در فناوری است. توجه رو به رشدی به نیاز برای سامانه‌های قابل توضیح وجود دارد و رفع این نیاز در پس توسعه هوش مصنوعی توضیح‌پذیر قرار گرفته که جعبه‌ای شیشه‌ای است و متضاد جعبه سیاه می‌باشد.

جهت اطمینان از اعتماد و شفافیت برای عموم، مسئولیت‌پذیری و پاسخ‌گویی روشن‌تری برای فرایندهای تصمیم‌گیری، لازم است که بتواند بدون قربانی کردن عملکرد هوش مصنوعی انجام شود. راهکارهای منصفانه‌تر و نمایا‌تر، تنها به این علت ارائه می‌شوند که مقررات و چارچوب‌های اخلاقی مستحکم و قابل اجرایی توسط اجتماعی وسیع‌تر توسعه داده شوند که هم‌تراز با تنوع جهان واقعی باشند.

اخلاقیات

گزارش «هوش مصنوعی در زمان حال» (Solon et al., 2017) که در سال ۲۰۱۷ منتشر شد، سوالات انتقادی اجتماعی درباره هوش مصنوعی را مورد بررسی قرار داده و بر چهار حوزه تمرکز داشته است: نیروی کار و خودکارسازی، سوگیری و فراگیر بودن، حقوق و مسئولیت‌ها، اخلاقیات و حاکمیت.

یکی از حیاتی‌ترین حوزه‌هایی که توجه بسیار زیادی را به خود جلب کرده است، مسئله سوگیری^۱ یا تورش است. سوگیری می‌تواند به صورت عمدی یا سهوی و از طریق انتخاب داده‌ها و پیش‌زمینه فرهنگی خود توسعه‌دهندگان ایجاد شود.

قبلاً مشخص شده است که سامانه‌های هوش مصنوعی که در این چند سال اخیر توسعه داده شده‌اند، با داده‌هایی آموزش دیده‌اند که حاوی سوگیری‌های زیان‌بار زیادی هستند که امروزه قابل قبول نمی‌باشند (مانند سوگیری‌های جنسیتی و نژادی).

سوگیری‌های ذاتی موجود در گروه‌های اجتماعی نسبتاً کوچکی که عمدتاً درگیر توسعه سامانه‌های هوش مصنوعی بوده‌اند، می‌توانند با استفاده از گروه‌های وسیع‌تر و متنوع‌تر در توسعه و آزمون سامانه‌های آینده و نیز افراد حرفه‌ای که به طور مداوم وظایف یاد گرفته شده توسط سامانه هوش مصنوعی را انجام می‌دهند تا با باورهای معاصر سازگار شود، تقابل داشته باشند.

یادگیری ماشین به مجموعه‌های بزرگ داده‌هایی بستگی دارد که دانش از آن‌ها استخراج می‌شود. این مجموعه‌های داده‌ها، کدگذاری تاریخی می‌شوند. با این حال، در این تاریخ‌ها پیش‌قضاوت‌ها، سوگیری‌ها و بی‌عدالتی‌های اجتماعی دیده می‌شوند که برای سال‌ها وجود داشته‌اند. پیامد این‌ها آن است که هوش مصنوعی به اطلاعاتی وابسته می‌شود که بر اساس آن آموزش دیده و ممکن است علیه منافع انسان عمل کند. به عنوان مثال، مشخص شد الگوریتمی که آمازون به عنوان ابزار جذب نیرو آزمایش می‌کرد، جنسیت‌گرا و بنابراین غیر قابل استفاده است (BBC, 2018). این سامانه هوش مصنوعی بر مبنای داده‌هایی آموزش دیده بود که توسط متقاضیان در طول یک بازه زمانی ده ساله ارسال شده بود و بسیاری از این متقاضیان مرد بودند. بنابراین باعث شد سامانه این گونه آموزش ببیند که کاندیداهای مرد ارجحیت دارند.

در سال ۲۰۱۶، گزارشی که توسط یک سازمان روزنامه‌نگاری تحقیقاتی به نام پروپابلیکا^۱ انجام شد، ادعا کرد نرم‌افزاری که از هوش مصنوعی استفاده می‌کرد و در دادگاه‌های ایالات متحده برای ارزیابی ریسک استفاده می‌شد، علیه زندانیان سیاه‌پوست سوگیری دارد. الگوریتم آن در پیش‌بینی این که چه کسی دوباره خلاف می‌کند، مجرمان سیاه‌پوست را دو برابر بیشتر از مجرمان سفیدپوست به عنوان مجرم آینده تشخیص می‌داد؛ چرا که ریسکشان کمتر از همتایان سیاه‌پوست‌شان برچسب‌گذاری شده بود (ProPublica, 2016).

در جابه‌جایی از تصمیمات مبتنی بر انسان به سمت تصمیمات مبتنی بر هوش مصنوعی، هوش مصنوعی لزوماً مبتنی بر کدهای اخلاقی نیست و نه فیلتر خردمندانه‌ای دارد و نه ترسی از تصمیمات غلط، تنبیه یا حس مسئولیت‌پذیری که همگی مسائل مهمی هستند.

توجه به این مسائل، صرفاً به معنی تلاش برای ایجاد سامانه‌های قابل توضیحی نیست که از طریق آن‌ها بتوان تصمیمات را تا حدی موشکافی کرد. درجه‌ای از نظارت مورد نیاز است تا اطمینان حاصل شود که سامانه‌های هوش مصنوعی طوری آموزش می‌بینند که بتوانند تصمیماتی را اتخاذ کنند که برای زمانه زندگی ما و زمان‌های آینده مناسب باشند. اخلاقیات و مسئولیت‌ها، همیشه مضمونی مهم و حیاتی باقی خواهند ماند و در همه ابعاد هوش مصنوعی رسوخ خواهند کرد و با رشد این حوزه، پیچیدگی و تعامل‌پذیری آن‌ها می‌تواند افزایش یابد.

سوگیری می‌تواند به شکل‌های زیادی وارد هوش مصنوعی شود که یا غیرعمدی از طریق طراحان سامانه‌ها و مهندسان نرم‌افزار است، یا نظام‌مند و از طریق مجموعه‌های داده‌ها یا روش‌های تحلیل داده‌ها رخ می‌دهد. همراهی سوگیری‌های الگوریتمی (عمدی یا سهوی) با جعبه سیاهی که نمی‌توان فهمید چگونه به نتیجه خاصی رسیده است، می‌تواند دارای مفاهیم زیادی برای محصولات، خدمات و ریسک شهرت شرکت‌ها باشد.

مسئولیت

یکی از مهم‌ترین چالش‌ها، وضعیت حقوقی سامانه‌ها و خدمات هوش مصنوعی است. با این که کارهای زیادی می‌توان انجام داد تا مهندسان پشت سیستم‌ها انگیزه پیدا کرده و تشویق شوند که پیامدهای تصمیماتشان در طراحی شبکه‌های عصبی و انتخاب مجموعه‌های داده‌ها را برای یادگیری مورد توجه قرار دهند، اما باید چارچوب‌های قانونی مستحکمی هم برای تعریف مسئولیت وجود داشته باشد.

سامانه‌های پزشکی، حمل‌ونقل و مالی، سامانه‌هایی هستند که تصمیمات مرتبط با مرگ و زندگی یا حداقل تصمیماتی را اتخاذ می‌کنند که زندگی را به چالش می‌کشند. چه این مسئولیت را کسی بپذیرد یا نپذیرد، سامانه‌های هوش مصنوعی یا روبات‌ها باید هویتی قانونی داشته باشند؛ اصطلاحی که مشخص می‌کند کدام موجودیت دارای حقوق و تعهدات قانونی است و می‌تواند وارد قرارداد شود یا تحت پیگرد قضایی قرار گیرد.

چارچوب‌های قانونی و رهنمودهای اخلاقی، در سطوح ملی و بین‌المللی با سرعت‌های متفاوتی توسعه می‌یابند و ذینفعان

بیشتری نسبت به کسانی که این فناوری‌ها را خلق کرده‌اند، ایجاد می‌کنند. دولت‌ها و سایر سیاست‌گذاران متوجه این واقعیت می‌شوند که باید مطمئن شوند جامعه آن‌ها نه تنها از مزایای هوش مصنوعی بهره می‌برد، بلکه از پیامدهای منفی بالقوه آن نیز محافظت می‌شود و آن‌ها در این باره مسئول هستند. سوالاتی که در این زمینه بروز می‌کنند شامل مواردی از این دست هستند: «آیا نظام آموزشی ما نیروی کار پویا، همکار و آگاه به علم داده را پرورش می‌دهد که آماده یادگیری مادام‌العمر باشد؟» یا «آیا صنایع ما می‌توانند از قدرت تصمیم‌گیری مبتنی بر هوش مصنوعی در همه مراحل کار خود استفاده کنند؟». در کل، تولیدکننده یک محصول، مسئول معایبی است که موجب خسارت به کاربران می‌شود. با این حال، در مورد تصمیمات هوش مصنوعی (به ویژه هوش مصنوعی قدرتمند)، زیان‌های احتمالی حاصله، پیامد طراحی نیستند، بلکه حاصل تفسیر واقعیت توسط یک ماشین می‌باشند.

امنیت

آسیب واقعی که دسترسی غیرمجاز می‌تواند به هوش مصنوعی وارد کند، به زمینه و بستر کار بسیار حساس است. مثلاً اگر هوش مصنوعی برای یک سامانه مهم ایمنی استفاده شود، می‌تواند منشأ تغییر کند و سامانه وابسته به آن در بهترین حالت غیرقابل استفاده و در بدترین حالت ناایمن می‌شود. این وضعیت بیش از همه منجر به تنزل امنیت زنجیره تامین خواهد شد. نگرانی‌ها و ریسک‌های خاصی (و البته فرصت‌هایی) مرتبط با دسترسی منبع‌باز به هوش مصنوعی وجود دارند. بستن نرم‌افزار، از سوءاستفاده از فناوری جلوگیری می‌کند و نیز مانع توسعه‌دهندگان هوش مصنوعی می‌شود که ممکن است نیت‌های زیان‌باری داشته باشند و شاید برای نفرات بعدی دشوار باشد که شاخص‌های ضدفعالیت^۱ را توسعه دهند.

تاریخچه حوزه امنیت سایبری نشان داده است که نرم‌افزار منبع‌باز از هر نوعی که باشد، ویژگی‌های خوب و بد را با هم دارد؛ یعنی تا حدی خوب است و تا حدی بد. بحثی وجود دارد در این رابطه که اگر نرم‌افزاری منبع‌باز باشد، افراد زیادی آن را امتحان و در آن مداخله می‌کنند. هم‌چنین این خطر وجود دارد که کدی بتواند کپی شود و به منظور دیگری مورد استفاده قرار گیرد. در میان توسعه‌دهندگان بدافزارها، رایج است که نرم‌افزاری را تکامل دهند که توسط دیگران نوشته شده است. کد بدخواهانه می‌تواند حاوی کدهای کوچکی باشد که به عنوان اثبات مفهوم برای نشان دادن آسیب‌پذیری‌ها منتشر می‌شوند و کدهای بدخواهانه پیشرفته‌تر می‌توانند از آن بهره ببرند.

سرعت شگفت‌انگیز دستاوردهای فناوری‌ها در هوش مصنوعی، مسائل جدیدی را در حیطه آثار معکوس بالقوه آن و تأثیرش بر اجتماع، که ناشی از حضور آن در همه ابعاد زندگی است، به وجود می‌آورد. فناوری‌های فعلی هوش مصنوعی امکان تصمیم‌گیری خودکار با استفاده از داده‌ها و توانایی انجام اقدامات یا شناسایی الگوها در سیلاب داده‌های موجود در دنیای مدرن را فراهم می‌کند. پتانسیل سوگیری الگوریتمی یا سوگیری در داده‌هایی که سامانه بر اساس آن آموزش می‌بیند، ریسک‌های جدیدی را ایجاد می‌کند؛ زیرا اتکای ما به فناوری هوش مصنوعی برای شناسایی داده‌های مربوطه و اقدام بر اساس آن‌ها رو به افزایش است. این موضوع اخیراً در پلتفرم‌های شبکه‌های اجتماعی آشکار شده است که در

آن‌ها، نمایه‌سازی جمعیت‌شناختی و محبوبیت فرد، بر پیشنهاد استوری‌های جدید برای او اثر می‌گذارند.

این وضعیت منجر به نگرانی‌هایی در خصوص نقش هوش مصنوعی در دوقطبی‌سازی عقاید در جامعه مدرن و در زمان‌های بحرانی شده است. رویدادهایی مانند انتخابات و همه‌پرسی، فرصت‌هایی را برای بهره‌گیری از رفتارهای شناخته‌شده در چنین الگوریتم‌هایی، جهت ایجاد سوگیری در عقاید عموم ایجاد می‌کند.

به صورت کلی‌تر، فصل مشترک امنیت سایبری و هوش مصنوعی در حال حاضر چندان مورد بررسی و تحقیق قرار نگرفته است. با این که سامانه‌های متعدد امنیت سایبری از هوش مصنوعی، مثلاً برای کمک به تشخیص هک کردن و دستورالعمل شبکه استفاده می‌کنند، اما تحلیل‌های کمی از امنیت سایبری هوش مصنوعی و سطح حمله بالقوه این سامانه‌ها صورت گرفته است. این شاید تا حدی به سبب اتکای سامانه‌های هوش مصنوعی بر شبکه‌های عصبی عمیق باشد که تفسیر یا عیب‌زدایی آن‌ها دشوار است. برای توجه به این دغدغه، تحقیقاتی در زمینه هوش مصنوعی توضیح‌پذیر^۱ در حال انجام است که به جایی فراتر از ارزیابی عملکرد شبکه‌های عصبی عمیق می‌روند تا بفهمند که چرا یک شبکه عصبی عمیق آن‌گونه عمل می‌کند.

تکنیک‌هایی در حال توسعه هستند تا سوگیری در آموزش یک شبکه (مثلاً شناسایی این که آیا تصمیم‌دارای سوگیری است یا تبعیض غیرقانونی اعمال می‌کند) یا شناسایی راه‌هایی که یک شبکه می‌تواند با دروندادهای خصمانه فریب بخورد، آشکار کنند. مثلاً، اخیراً نشان داده شده است که سیستم‌های شناسایی بصری^۲ می‌توانند با طبقه‌بندی بد اشیا یا حتی حذف برخی از اشیا با وارد کردن برچسب‌های خصمانه در صحنه، فریب بخورند (BBC, 2017; Open AI, 2017).

توانایی نابینا کردن یک شبکه عصبی عمیق و ندیدن حضور اشیا^۳ که در تصویری به وضوح قابل مشاهده هستند، نشان‌دهنده تفاوت‌ها میان ادراکات انسانی و فناوری‌های هوش مصنوعی است. هم‌چنین این موضوع بر اهمیت بحث‌های اخلاقی در حیطه تاب‌آوری و توضیح‌پذیری هوش مصنوعی و نیز حریم خصوصی داده‌ها و سوگیری الگوریتمی در فناوری‌های کنونی یادگیری ماشین تاکید دارد.

در نهایت، شایان توجه است که هوش مصنوعی، بعد جدیدی را در امنیت سایبری شکل داده است؛ زیرا می‌تواند یاد بگیرد و بسیار سریع‌تر از روش‌های سنتی مورد استفاده محصولات امنیتی قدیمی، به تهدیدات واکنش نشان دهد. با این حال، عجیب نیست که مجرمان هم شروع به استفاده از هوش مصنوعی کرده‌اند تا یاد بگیرند که همه مراحل یک حمله سایبری را چگونه هدایت کنند؛ یعنی از شناسایی تا انجام حمله. این مسابقه تسلیحاتی اجتناب‌ناپذیر است. با پیچیده‌تر شدن و اتصال داخلی این سامانه‌ها به یکدیگر، رخنه‌های امنیت سایبری برای هوش مصنوعی به هم پیوسته و در هم تنیده، اهمیت بیشتری می‌یابد و احتمالاً آثار اقتصادی نظام‌مندی خواهد داشت.

هوش مصنوعی و بیمه

تأثیر هوش مصنوعی بر رشته‌های بیمه‌ای

مسئولیت محصول و فراخوان محصول

محصول متکی بر هوش مصنوعی، به طرق متعددی می‌تواند معیوب باشد. حتی اگر آزمون محصول یک ماشین متکی بر هوش مصنوعی بسیار دقیق باشد، ممکن است در مواقعی ماشین خودمختار با موقعیتی مواجه شود که در طول دوره آزمون با آن مواجه نشده بود.

این وضعیت می‌تواند منجر به توقف، فروپاشی یا سوء کارکرد شود و نقص در طراحی هم می‌تواند خطرناک باشد. برای تعاملات انسان و ماشین، ایمنی اهمیت زیادی دارد. عیوب محصول حتی می‌تواند حاصل خطاهای ارتباطی میان دو ماشین یا میان ماشین و زیرساخت باشد.

نکته

فراخوان‌ها، به ویژه اگر بخش تأثیر گرفته، از هوش مصنوعی در سطح وسیعی استفاده کند، می‌توانند بزرگ‌تر و پیچیده‌تر شوند. مثلاً، اگر مجموعه تصادف‌هایی رخ دهند و این تصادفات مربوط به هوش مصنوعی مورد استفاده در اتومبیل‌های بدون راننده باشند، این وضعیت می‌تواند منجر به فراخوانی بزرگ در میان تولیدکنندگان و کشورهای مختلف شود و در کنار آن دادگاه‌های پرهزینه و بلندمدتی تشکیل شود و زمان زیادی صرف اثبات این شود که طراحی هوش مصنوعی مشکلی نداشته است.

ماشین‌های هوش مصنوعی نمی‌توانند خودشان مسئول اقدامات قصورآمیز یا خطاهای خود باشند که موجب خسارت به اشخاص ثالث می‌شوند. اما با بهتر شدن این ماشین‌ها در تفکر، یادگیری و تصمیم‌گیری، سوال این است که وقتی ماشین خطایی می‌کند، چه کسی مسئول است؟

فروشنده‌گان یا تولیدکنندگان محصولات، طراحان و تامین‌کنندگان هوش مصنوعی و خریداران و کاربران آن، ممکن است توسط دادگاه مقصر شناخته شوند. در آینده، بیمه‌گران احتمالاً شاهد افزایش شرکت‌هایی باشند که از تضامین قراردادی، جبران غرامت و محدودیت‌ها برای کنترل ریسک مسئولیت هوش مصنوعی استفاده می‌کنند.

مسئولیت شخص ثالث وسایل نقلیه

در آینده به سبب جابه‌جایی احتمالی مسئولیت از رانندگان به سمت خودروهای بدون راننده و بنابراین به تولیدکنندگان، تخصیص و پوشش مسئولیت چالشی‌تر خواهد شد. پتانسیل قابل توجهی هم برای واکنش‌های حقوقی و قانونی در کشورهای مختلف وجود دارد و چشم‌انداز پیچیده‌ای برای ریسک دیده می‌شود. در حوزه‌های قانونی، بیمه‌نامه‌های فعلی خودرو به‌طور کلی بر این اصل استوار هستند که کاربر خودرو مسئول خسارات ناشی از اشتباهات رانندگی فردی و نقص خودرو به دلیل عدم نگهداری صحیح است.

نکته

با خودروی کاملاً خودران، این مسئولیت ممکن است از بیمه‌شدگان به تولیدکننده خودرو یا تامین‌کنندگان آن، مانند ارائه‌دهندگان نرم‌افزار هوش مصنوعی منتقل شود. در برخی مناطق، ممکن است مدل‌های جدید بسط مسئولیت محصول پذیرفته شوند و تولیدکنندگان، مسئولیت شخص ثالث را برای عیوب محصول بر عهده بگیرند. در نتیجه، بازار برای پوشش مسئولیت شخص ثالث ممکن است در برخی مناطق به شدت کاهش پیدا کند، که شاید همراه با افزایش تقاضای بیمه مسئولیت محصول برای تولیدکننده باشد.

در انگلستان، در زمینه وسایل نقلیه قابل خودکارسازی سطح ۴ یا ۵^۱، مقررات سال ۲۰۱۸ تأیید کرده است که بیمه‌شده قادر خواهد بود در وهله اول تحت بیمه‌نامه شخص ثالث وسیله نقلیه ادعای خسارت داشته باشد. در این صورت بیمه‌گر قادر خواهد بود بازایی مناسب را در صورت قصور تولیدکننده انجام دهد. این کار موجب جلوگیری از موقعیت‌هایی می‌شود که افراد از جبران غرامت (به عنوان مدعی شخص اول بدون پوشش بیمه‌نامه) یا پیگیری خسارت از تولیدکنندگان بر مبنای مسئولیت محصول، که می‌تواند فرایندی پیچیده و طولانی باشد، مستثنی می‌شوند.

نکته: جبران غرامت حرفه‌ای و روبات‌های مشاور

مشاوره روباتی، مشاوره‌ای است که توسط یک الگوریتم انجام می‌شود. ممکن است از مشاوره غلطی که توسط شرکتی داده می‌شود، مسئولیتی ایجاد شود اما نه توسط انسان‌های متخصصی که معمولاً تحت پوشش بیمه‌نامه مسئولیت حرفه‌ای قرار دارد. آنچه گفته شد، با مسئولیت محصول متفاوت است و ممکن است خسارتی ایجاد نشود یا طراحی سامانه معیوب نباشد. با این حال، توصیه‌ای که هوش مصنوعی ارائه می‌دهد می‌تواند دارای خطا باشد.

اما وقتی هوش مصنوعی برای ارائه مشاوره حرفه‌ای استفاده می‌شود، چه اتفاقی می‌افتد؟ پیشرفت‌هایی در قابلیت‌های روبات‌های مشاور به وجود آمده است. این قابلیت‌ها از مشاوره روباتیک برای سرمایه‌گذاری‌ها تا کاربردهای بالقوه در فناوری قانون^۲ و حتی روبات‌های پاسخگو^۳ که درمان رفتار شناختی^۴ را ارائه می‌کنند، در بر می‌گیرد که در آن‌ها از الگوریتم‌ها برای ایجاد پیام استفاده می‌شود (Techemergence, 2018; Woebot, 2018; Hudsonmckenzie, 2018).

مشاوره روباتی برای سرمایه‌گذاری، از پلتفرمی استفاده می‌کند که برخی پارامترهای داده‌شده توسط مصرف‌کننده را دریافت کرده و با الگوریتم‌ها، پورتفویی از دارایی‌ها را می‌سازد. چنین روبات‌های مشاوره، هوش مصنوعی را به میزان‌های مختلفی به کار می‌گیرند و برخی از یادگیری ماشین برای شبیه‌سازی عملکرد گذشته استفاده می‌کنند (Techemergence, 2018). تا امروز، درباره این که میزان اشتباه ریسک چه قدر باید باشد یا چه مبلغی باید سرمایه‌گذاری شود، مشاوره‌های چندانی به مصرف‌کنندگان داده نشده است. اگر این وضعیت تغییر کند و مشاوره روباتیک ابعاد حرفه‌ای‌تری به خود بگیرد، آن‌گاه برای مشاوره غلط، مسئولیت ایجاد می‌شود.

۱. در خودکارسازی سطح ۴، ماشین اغلب اوقات می‌تواند بدون دخالت راننده حرکت کند، اما در پاره‌ای از مواقع باید راننده دخالت داشته باشد و در خودکارسازی سطح ۵، اتومبیل در همه شرایط خودران است. [م]

2. LawTech
3. Chatbots
4. Cognitive Behaviour Therapy (CBT)



این وضعیت در مورد هر روبات مشاوره که مشاوره حقوقی و یا در هر حوزه دیگری که خطا می‌تواند موجب نقض قرارداد یا زیان‌آور باشد، نیز صدق می‌کند. آنچه گفته شد چالش‌ها و فرصت‌هایی را برای بیمه‌گران جبران غرامت حرفه‌ای ایجاد می‌کند. در مرحله‌ی گذار که متخصصان از هوش مصنوعی برای کمک استفاده می‌کنند و کل تصمیم‌گیری توسط هوش مصنوعی انجام نمی‌شود، این وضعیت پیچیده است. در این جاست که مسئولیت هم‌چنان به بیمه‌نامه جبران غرامت حرفه‌ای تخصیص می‌یابد. برخی روبات‌های سخنگو با سایر پلتفرم‌ها مانند فیسبوک، در حال یکپارچه شدن هستند. یک روبات سخنگوی درمان‌گر مانند ووبات^۱، پذیرش و راهنمایی گام‌به‌گام را انجام می‌دهد. تصور می‌شود که تعاملات مشابه انسان، به مشارکت افراد کمک کرده و خدمات مشفقانه روانشناسی رایگان می‌دهد.

بیشتر روبات‌های سخنگوی موجود، توصیه نمی‌کنند که کسی صرفاً بر اساس حرف آن‌ها عمل کند که این می‌تواند با بلوغ هوش مصنوعی در آینده تغییر کند. اما در این صورت، چه کسی مسئول مشاوره‌ی ارائه‌شده است؟ ارائه‌دهندگان اپلیکیشن یا طراحان؟ باید در قرارداد، به روشنی مشخص شود که چه کسی مسئول است. اگر داده‌های شخصی محدود به جزییات ساده نباشند، بلکه داده‌های کافی برای ایجاد یک نمایه روانشناختی وجود داشته باشند، آن‌گاه هر گونه رخنه داده‌ها می‌تواند بسیار جدی باشد و به‌طور بالقوه افراد را در معرض آسیب‌پذیری قرار دهد. نحوه ثبت تاریخچه این چت‌ها نیز مشکلاتی را ایجاد می‌کند. اگر آن‌ها نگهداری شوند، شانس رخنه داده‌ها افزایش می‌یابد. اما اگر پاک شوند، مشکلاتی در دعاوی قضایی به وجود خواهد آمد، به ویژه اگر یک دوره نهفتگی میان مشاوره و ادعای خسارت وجود داشته باشد.

قصور پزشکی

قصور پزشکی به طرق مختلفی می‌تواند بروز کند. هوش مصنوعی برای کمک به تشخیص استفاده و در حوزه‌هایی مانند رادیوگرافی به کار گرفته می‌شود (Lee et al., 2017).

اگر خطایی منجر به تشخیص اشتباه یا نتیجه مثبت اشتباهی شود، می‌تواند مشمول قصور پزشکی باشد. حتی اگر هوش مصنوعی به عنوان کمکی برای ارجاع استفاده شود و اگر این ارجاع‌ها موجب آزمایشات یا رویه‌های غیرضروری شوند، یا پیامدهای بد یا ضعیفی برای بیمار داشته باشند، آن‌گاه موضوع مسئولیت بروز می‌کند.

نکته

کمیسیون قانون موضوعه می‌تواند از لحاظ هزینه دفاع، منجر به دعاوی خسارت بزرگ‌تری هم بشود. زیرا ممکن است بیمه‌شده شواهد مدعی را به چالش بکشد و با وجود سوابق قانونی محدود، حل مسئله می‌تواند با استفاده از راهکارهای جایگزین، دشوار باشد. هوش مصنوعی‌هایی که ضعیف طراحی شده یا الگوریتم‌هایی که برای کاری غلط بهینه شده باشند، می‌توانند خطای نظام‌مند ایجاد کنند و منجر به دعاوی حقوقی متعددی شوند. وقتی مسئولیت را بتوان به وضوح به هوش مصنوعی نسبت داد، به راحتی قابل مستثنی شدن از بیمه‌نامه قصور پزشکی است. اما اگر هوش مصنوعی به عنوان کمک استفاده شده باشد یا در استفاده از آن‌ها غفلت شده باشد، مثلاً درست راه‌اندازی نشده باشد، تعیین مسئولیت و علت نزدیک می‌تواند پیچیده شود.

سایبری

روبات‌های سخن‌گو به‌طور روزافزونی به عبور موفق از آزمون تورینگ^۱ نزدیک می‌شوند که اولین بار الن تورینگ در سال ۱۹۵۰ مدعی آن شد. برای گذراندن آزمون، رایانه باید رفتاری هوشمندانه از خود نشان دهد تا انسان ارزیاب به دشواری بتواند پیش‌بینی کند که محتوای ارائه‌شده در گفتگو، توسط انسان تولید شده یا ماشین. دعاوی خسارت متعددی وجود دارند که این اتفاق در آن‌ها رخ داده است، مانند ویژگی نسخه نمایشی دستیار هوش مصنوعی گوگل که قادر است به تلفن شما پاسخ دهد (Extremetech, 2018). این که آزمون تورینگ گذرانده شود یا نه، لزوماً نقطه عطفی نیست که بیمه‌گران منتظر آن باشند. اگر روبات بتواند دارای برخی ویژگی‌های انسانی شود، پس پتانسیل سوءاستفاده نیز افزایش می‌یابد.

اگر یک سامانه هوش مصنوعی برای ارسال ایمیل‌ها، تماس‌ها یا پیام‌های فیشینگ و پاسخ به دریافت‌کننده استفاده شود، آن‌گاه این اتفاق می‌تواند روش انجام کلاهبرداری‌های مهندسی اجتماعی را تغییر دهد. این مشکل عوامل متعددی دارد که می‌توانند موجب کلاهبرداری‌های موفق‌تری شوند. این مسئله می‌تواند موجب بروز موارد زیر شود:

(۱) سفسطه‌ها می‌توانند به سرعت افزایش یابند. به جای ارسال یک ایمیل پیش‌فرض برای همه، روبات‌های هوش مصنوعی می‌توانند پیام‌های اولیه خود را شخصی‌سازی کنند. اگر سامانه‌های هوش مصنوعی از طریق رخنه‌های قبلی یا داده‌هایی که خریداری می‌شوند، به هر گونه اطلاعات شخصی دسترسی داشته باشند، این شخصی‌سازی می‌تواند وسیع باشد.

(۲) با یادگیری ماشین، روبات‌های هوش مصنوعی می‌توانند یاد بگیرند که کدام نوع ایمیل‌ها یا حملات، بهتر کار می‌کنند. این ایمیل‌ها ممکن است بیشترین شباهت را به رفتار انسان نداشته باشند، اما به راحتی قادر باشند آسیب‌پذیرترین کاربران را که مستعد این نوع حمله هستند، هدف قرار دهند.

(۳) سامانه‌های هوش مصنوعی می‌توانند به صورت بالقوه در وب‌سایت‌های جعلی گنجانده شوند، که با عنوان سایت‌های اصلی نمایش داده می‌شوند. این سایت می‌تواند شبکه اجتماعی، بانک یا سایتی جعلی باشد که قادر است پذیرندگان را به خود جلب کند.

(۴) فیشینگ هوش مصنوعی می‌تواند خودکارسازی شده و در حجم‌های بالا ارسال شود.

(۵) هوش مصنوعی ممکن است حتی طراحی خود کلاهبردار هم نباشد و از طریق هک کردن یا منبع‌بازهای موجود به دست آمده باشد.

(۶) یک حمله هیبریدی که در آن از انسان‌ها در بخش‌هایی از یک زنجیره کلاهبرداری پیچیده استفاده می‌شود، می‌تواند برای کمبودها از قابلیت‌های هوش مصنوعی استفاده کند.

(۷) با این که گسیل کردن هوش مصنوعی در امنیت سایبری می‌تواند توانایی مدافعان را برای شناسایی و جلوگیری از حملات افزایش دهد، اما برخی نقص‌های دفاعی امنیت سایبری در سامانه‌های هدف را نیز می‌تواند آشکار کند.

نکته

این مسائل سوالاتی را درباره این که چه نوع پوشش‌های محافظی در برابر انواع مختلف زیان‌ها، می‌توانند در دسترس باشند، ایجاد می‌کند.

این که یک حمله فیشینگ مبتنی بر هوش مصنوعی به قدر کافی مختل‌کننده هست که به عنوان یک حمله سازشکارانه حساب شود، ممکن است سوالاتی را درباره پوشش ایجاد کند. آن‌چه به عنوان یک رویداد بیمه‌پذیر شناخته می‌شود، باید به خوبی تعریف گردد. مسئولیت برای خسارات شخص ثالث هم می‌تواند پیچیده باشد. اگر یک روبات بیمه‌شده ربوده شود و برای اهداف مجرمانه مورد استفاده قرار گیرد، آن‌گاه ممکن است تحت پوشش یک بیمه‌نامه سایبری قرار گیرد. اگر به سبب استفاده از روبات برای اهداف مجرمانه، خسارات پایدار باشند، این موضوع بروز می‌کند.

اگر بیمه‌شده به علت قصور و غفلت، اقدامات مناسب برای محافظت از ربوده شدن روبات خود را انجام ندهد، آن‌گاه مسئول شناخته می‌شود. پوشش برای زیان‌های شخص اول مانند عدم‌النفع ناشی از ربایش روبات، ممکن است سراسرتر باشد.

وظیفه‌شناسی

فعالیت متقابلانه از طرف کارکنان هم می‌تواند توسط هر یک از روش‌های ذکرشده در بخش سایبری تشدید شود. بیمه‌گران باید توجه داشته باشند که تقلب به‌طور فزاینده از سوی کارکنانی انجام خواهد شد که به سامانه‌های فناوری اطلاعات دسترسی دارند، نه کارکنانی که دارای اختیارات مالی هستند.

نکته

وقتی صحبت از تقلب به میان می‌آید، ظهور جعل عمیق هم می‌تواند نگرانی‌هایی را ایجاد کند. جعل‌های عمیق، سامانه‌های هوش مصنوعی هستند که قابلیت تولید صوت، ویدئو و تصاویر ظاهراً واقعی را دارند. این روش مربوط به موارد تقلب هویت نیز هست و می‌تواند در کنار روبات‌های سخن‌گو برای افزایش اعتبار و ایجاد اعتماد هم مورد استفاده قرار گیرد.

حملات سایبری که از جعل عمیق بهره می‌گیرند، می‌توانند تهدیدی را به روش‌های جدید و غیرمنتظره ایجاد کنند. مثلاً متخصصان امنیت سایبری، آسیب‌پذیری‌های شبکه بیمارستانی را آشکار کرده‌اند. محققان بدافزاری را معرفی کردند که از یادگیری عمیق برای حذف یا ایجاد لخته‌های خون یا تومورها در تصاویر سه بعدی پزشکی مانند سی‌تی‌اسکن و ام‌آر‌آی استفاده می‌کند. تشخیص غلط می‌تواند منجر به پیامدهای بالینی زیان‌بار، کارشکنی در تحقیقات یا تقلب در بیمه‌های پزشکی و درمان شود. این فناوری برای انجام حملات مخفیانه به رهبران سیاسی یا اقدام تروریستی هم می‌تواند استفاده شود (BBC, 2019).

ریسک‌های سیاسی

همان‌گونه که در بخش ریسک‌های هوش مصنوعی گفته شد، مسلح‌سازی‌های بالقوه می‌تواند شکل‌های مختلفی به خود بگیرد که از آن جمله می‌توان به نادرستی داده‌ها، انتخاب نامتوازن داده‌ها و استفاده غیرقانونی از داده‌ها اشاره کرد که به پیامدهای بسیار زیادی از جمله پروپاگاندا، تغییرات رفتاری و فریبکاری منجر می‌شود. سامانه‌های هوش مصنوعی می‌توانند از مزیت ظرفیت بهبودیافته برای تحلیل رفتارها، خلق و خوها و باورهای انسان بر اساس داده‌های موجود بهره ببرند (Brundage, et al., 2018).

نکته

از منظر پوشش ریسک سیاسی، هوش مصنوعی می‌تواند در ایجاد یا تشدید رویدادهای قوه قهریه سیاسی، مانند مصادره، جنگ، اقدام تروریستی، آشوب‌های شهری و سایر شکل‌های جرم سیاسی، نه تنها در بازارهای در حال توسعه، بلکه حتی در بازارهای توسعه‌یافته هم سهمیم باشد. در چند سال گذشته، شاهد (سوء)استفاده از هوش مصنوعی در کمپین‌های سیاسی بوده‌ایم (The Conversation, 2017). به‌طور مشابه، پروپاگاندا‌های هدفمند رایانشی و جعل‌های عمیق هم می‌توانند موجب توزیع کاراتر، شدید و وسیع اطلاعات غلط و تشدید ماجراهای گمراه‌کننده شوند. در صدر ریسک‌های یادشده، ناپایداری ناشی از سامانه قطع خودکار و اختلال اقتصادی، می‌تواند نیروی پیشران بالقوه‌ای برای تظاهرات اعتراض‌آمیز و تحریک باشند که منجر به مداخله دولت، تبعیض گزینشی و عدم‌النفع می‌شوند.

فرصت‌های توسعه کسب‌وکار

در کل، هر شرکتی که سامانه‌های مبتنی بر الگوریتم را برای شرکت‌های غنی از داده ارائه کند (مثلاً تشخیص تقلب در فروش آنلاین)، ممکن است در برابر ریسک الگوریتم‌هایی که تصمیمات غلط می‌گیرند و نیز بابت تاثیر آن‌ها بر مشتریان شرکت‌های هوش مصنوعی، به دنبال دریافت پوشش بیمه‌ای باشد. قبل از بیمه کردن خروجی‌های الگوریتم (خودیادگیرنده بالقوه)، بیمه‌گران باید نظریه آماری^۱، زیرساخت و اعتمادپذیری آن را آزمون کنند. به علاوه، شرکت‌های جدیدی در حوزه دفاع در برابر اطلاعات غلط در حال ظهور هستند که:

- فناوری‌هایی را برای فیلتر کردن اخبار جعلی ارائه می‌کنند؛

- روبات‌های اوباش مجازی را تشخیص داده و حذف می‌کنند؛

- اطلاعات و اعتبار تصاویر و ویدئوها را تایید می‌کنند.

آن‌چه گفته شد، ممکن است فرصتی باشد برای بیمه‌گران تا بررسی کنند که چه نوع محصولات (مثلاً جبران غرامت حرفه‌ای و محصولات سایبری) و به چه شکلی می‌توانند برای این کسب‌وکارهای جدید مفید باشند.

با پیشرفت این حوزه و افزایش کاربردها، فرصتی‌هایی هم برای ارائه خدمات مدیریت ریسک ظهور می‌کنند. اگر چه ما در مراحل نسبتاً ابتدایی توسعه هوش مصنوعی به سر می‌بریم، اما متخصصان دانش هوش مصنوعی از قبل تقاضای زیادی از طرف شرکت‌هایشان برای مدیریت ریسک داشته‌اند. شرکت‌های خدمات تخصصی در حال ظهور هستند که هدفشان پیش‌گیری از خسارت است. با افزایش پیچیدگی، تقاضا برای این خدمات افزایش می‌یابد.

عملیات

هوش مصنوعی، هم ریسک‌ها و هم فرصت‌هایی را برای بخش بیمه فراهم می‌کند. از منظر عملیاتی، شرکت‌های بیمه از قبل در حال بهره‌برداری از پتانسیل هوش مصنوعی برای ارائه ارزش به روشی کارا تر بوده‌اند. بسیاری از مراحل در زنجیره ارزش بیمه (اگر نگوییم همه مراحل)، توسط هوش مصنوعی متحول خواهند شد. ماهیت رقابتی مدل‌های کسب و کار خدمات محور، در حال انگیزه دادن به همه بخش‌هاست تا سامانه‌های پشتگاه^۱ خود را در حول تجربه مشتری بهبود دهند و توجه خاصی هم به کاهش هزینه‌ها و افزایش سرعت وجود دارد.

موارد زیر، رایج‌ترین کاربردهای هوش مصنوعی است:

- روبات‌های سخن‌گو و دستیاران هوش مصنوعی برای مشتری: آن‌ها می‌توانند محصولات را پیشنهاد داده و شخصی‌سازی کنند، به شکایات رسیدگی نمایند و بر مبنای تحلیل احساسات و پردازش تعاملات ساده، ارتباطات با مشتری را بهبود بخشند.

- تشخیص تقلب در خدمات مالی از جمله بیمه: در حال حاضر، این عملیات بیشتر متکی بر انسان است. اما دسترسی بیشتر به فناوری‌های رایانشی به همراه پیشرفت‌های هوش مصنوعی، به معنای این است که تشخیص تقلب می‌تواند با استفاده از یادگیری اجرایی عمیق، برای کاوش داده‌های ساختاریافته و بدون ساختار، خودکارسازی شود تا الگوهای جدید بالقوه تقلب را تشخیص دهد (Bouchti, Chakroun, Abbar, & Okar, 2017).

در انگلستان، اداره تقلبات از یک روبات پایلوت برای پردازش نیم میلیون سند در روز استفاده می‌کند تا محتوای حقوقی حرفه‌ای را با سرعت ۲۰۰۰ برابر بیشتر از یک وکیل، اسکن کند (SFO, 2018).

در بیمه، فرایندها می‌توانند با افزودن داده‌های بیرونی از رکوردهای تاریخی، حسگرها و تصاویر، خودکارتر شوند تا هزینه‌های تعمیر و استهلاک را بهتر تخمین بزنند.

- صدور بیمه: صدور بیمه می‌تواند از طریق هوش مصنوعی تقویت شده و سریع‌تر صورت گیرد، به ویژه وقتی مجموعه داده‌های چندگانه‌ای در فرایند دخیل هستند. مدل‌های مبتنی بر هوش مصنوعی می‌توانند بر مبنای ارزیابی ریسک‌های گذشته، حق بیمه را پیش‌بینی و مظنه‌هایی را برای یک محصول خاص ارائه کنند که متناسب با پروفایل ریسک مشتری است.

- دعاوی خسارت: هوش مصنوعی می‌تواند به کاهش تعداد خساراتی که نیازمند تحلیل انسانی و تعامل هستند، با خودکارسازی شناسایی تصویر و جستجوی پایگاه‌های بزرگ داده برای خسارات مشابه، هرگونه ریسک تقلب را برآورد کنند، اعتبار دعاوی خسارت را بسنجند و با مشتریان تعامل داشته باشند تا خسارات موجود را به‌روزرسانی کنند.

- مدلسازی، مدیریت ریسک و قیمت‌گذاری: با هوش مصنوعی و ترکیب آن با سایر فناوری‌ها مانند اینترنت اشیا، بیمه‌گران و مدلسازان می‌توانند به داده‌هایی دسترسی داشته باشند که مدل‌هایی را تغذیه می‌کنند که می‌توانند یاد بگیرند

و با سرعت بیشتری انطباق یابند. با این بازارها مثلاً می‌توان نرخ‌های بهره را بر مبنای ارتباطات بانک مرکزی پیش‌بینی کرد یا به پیش‌بینی قسمت‌های خالی در یک مجموعه داده کمک نمود (Panlilio, Canagaretna, Perkins, du Preez, and Lim, 2018).

- توسعه کسب و کار: سیستم‌های هوش مصنوعی می‌توانند فروش جانبی و نرخ‌های تبدیل محصول را بهبود بخشند، محصولات متناسب با نیازهای مشتری را پیشنهاد دهند و به شناسایی مشتریان جدید کمک کنند.

نکته: مهندسی ریسک اموال تجاری

هدف مشاوره ریسک این است که بتوان درباره ریسک، به بیمه‌گران و مشتریان مشاوره تخصصی داد. این مشاوره، به بیمه‌گران در تصمیم‌گیری برای پذیرش ریسک و قیمت آن و نیز به مشتریان در فهم جایی که می‌توانند اقداماتی را برای جلوگیری از ریسک انجام دهند، کمک می‌کند. با این حال، مشاوران ریسک باید داده‌های مربوط به ریسک را برای ارزیابی اخذ کنند تا به انتخاب و مدیریت ریسک کمک کنند؛ کاری که به سبب حجم زیاد عرضه‌ها و اطلاعات مورد نیاز برای ارزیابی دقیق ریسک، روزبه‌روز چالشی‌تر می‌شود. برای آمیختن این‌ها، مشاوره ریسک هم دارای منابع محدود و هم حساس به زمان است. به دلیل افزایش تقاضا برای مهندسی، رسیدن به اهداف رشد چالش برانگیز شده است. یک راه برای رفع این چالش، افزایش افراد (که به سبب ماهیت تخصصی این نقش و اهداف هزینه‌ای، همیشه عملی نیست) یا یافتن راهی برای خودکارتر کردن فرایند پردازش است.

تصور کنید که از هوش مصنوعی در دنیای مهندسی ریسک اموال استفاده شود تا قابلیت ایجاد شود که یک مشاور ریسک بتواند مانند هزار نفر فکر کند. با خودکارسازی فرایند مطالعه گزارش‌های مهندسی ریسک از طریق پردازش زبان طبیعی^۱، بهره‌وری مشاوران ریسک به میزان زیادی افزایش می‌یابد و منجر به صدور آگاهانه‌تر بیمه و رضایت بالاتر مشتریان می‌شود. خودکارسازی هم‌چنین می‌تواند فشار را از روی مهندسی ریسک که نیازمند منابع و زمان بسیار زیادی است، بردارد. تصور کنید که قادر باشید منابع داده‌ای غنی و وسیع را از گزارش‌های مهندسی به دست آورید و آن را در جایی ذخیره کنید تا بینش‌های کلیدی را در سطح حساب یا پورتفو، یا درباره نحوه کمک به بیمه‌گر برای زمان تمدید یا شناسایی ریسک‌های مشابه بر مبنای نقاط داده‌ای مشابه به دست آورید.

هوش مصنوعی می‌تواند با استفاده از پردازش زبان طبیعی، گزارش‌های مهندسی را ظرف مدت چند دقیقه بخواند و دانش و بینش‌هایی را از صفحات متعدد (گاهی بیش از ۱۰۰ صفحه) به دست آورد. با استفاده از دانش مشاوران ریسک برای غنی کردن درک پردازش زبان طبیعی از مهندسی، راهکاری می‌تواند ایجاد شود که محتوای گزارش‌های مهندسی را درک کند، مفاهیم مهندسی را از اسناد به دست آورد و برای مقاصد مانند ارزیابی ریسک یا تولید بینش داده‌ای استفاده کند. علاوه بر این، الگوریتم‌های پیچیده‌تر پردازش زبان طبیعی می‌توانند روی داده‌ها به کار گرفته شوند تا ارزیابی کنند که آیا ریسکی خوب است یا بد و در نهایت در تصمیم جهت صدور یا عدم صدور بیمه برای آن سهم شوند.

۱. Natural Language Processing: پردازش زبان طبیعی یا به اختصار NLP، یکی از شاخه‌های هوش مصنوعی است که به تعاملات بین رایانه و انسان، از طریق زبان طبیعی می‌پردازد. هدف غایی NLP، خواندن، رمزگشایی، فهم و درک زبان انسان با روشی ارزشمند است. بیش‌تر روش‌های پردازش زبان طبیعی برای استخراج و فهم معنای زبان انسانی، مبتنی بر تکنیک‌های یادگیری ماشین هستند. [م]

ارائه قابلیت زبان طبیعی نه تنها نیازمند فناوری بزرگی است، بلکه مهم‌تر از آن باید ترکیب درستی از تیم و همکاری در آن ایجاد شود. پردازش زبان طبیعی نیازمند زمان و تعهد تیمی است که هم مفاهیم آن و هم حوزه کسب و کار مربوطه را بشناسد. مثلاً، آکسا ایکس‌ال^۱ و اکسپرت سیستم^۲ با تیمی جهانی از شش مشاور ریسک و فناوران جهانی کار می‌کنند تا راهکاری ارائه دهند که همه گزارش‌های مهندسی با سرعت و دقت زیادی قابل پردازش باشند، در حالی که قبلاً فقط چند نمونه قابل خواندن بودند. سرعت زمان برای مرور گزارش‌ها یعنی این که اکنون ۵۰ گزارش می‌تواند در مدت ۵ دقیقه یا کمتر، خوانده شوند و مرور بیشتر به‌طور متوسط نیازمند ۸ ساعت زمان است، کاری که قبلاً ۱۰ تا ۲۰ ساعت برای هر حساب کاربری زمان نیاز داشت. باید توجه شود که برای آموزش بنیان‌های مشاوره ریسک به سامانه، مشارکت مستمر متخصصان مشاوره ریسک ضروری است. راهکارهای پردازش زبان طبیعی تنها به اندازه کیفیت متخصصان موضوع مناسب هستند، چرا که ماشین می‌تواند صرفاً آن‌چه را که آموخته است انجام دهد.

با فراهم کردن بینش‌های باکیفیت‌تر و سریع‌تر مهندسی، ریسک‌های درستی بیمه می‌شوند که منجر به حصول اهداف رشد شرکت می‌گردد. وقتی زمان مشاوران ریسک آزاد شود تا وقت بیشتری را با مشتریان بگذرانند، تجربه مشتریان افزایش می‌یابد و احتمالاً منجر به نتایج بهتری در مدیریت ریسک می‌شود. هم‌زمان، نیاز به توسعه نیروی کار مهندسی کاهش می‌یابد، زیرا دارای ریسک از دست رفتن دانش و این احتمال است که این ریسک‌ها بدون بینش کامل بیمه شوند یا پوشش بیمه‌ای آن‌ها رد شود. در نهایت، می‌توان نتیجه گرفت که به‌کارگیری هوش مصنوعی در فرایند مهندسی، برای قیمت‌گذاری خودکار یا پویا و نیز فراهم کردن منبع جدید داده‌های مهندسی قابل دسترس در کل سازمان، یک نیاز اساسی است که وقتی به سایر منابع غنی از داده متصل شود، بینش‌های جدید و گران‌بها و تصمیمات بهتری را حاصل می‌کند. سرانجام این که، تمرکز بر آینده و نوآوری، اثر مثبتی بر شهرت شرکت دارد.

فناوری بیمه (اینشورتک)

بیمه در حال ورود به زمانه‌ای است که برای پیشرفت‌های استثنایی در فناوری و کسب و کارهای زیاد اینشورتکی، از هوش مصنوعی برای برهم‌زنندگی بیشتر استفاده و بازار بزرگ‌تری را برای نوآوری و خدمات جدید خلق می‌کنند. نوآوری اینشورتکی در سطح جهانی در حال وقوع است، اما بسیاری از کسب و کارها همه چیز را در سطح داخلی می‌بینند.

زیست‌بوم‌های دیجیتال مرزهای ملی را نمی‌شناسند. یک سرمایه‌گذار فرشته^۳ به نام مهرداد پیروزرام گفته است شرکت‌هایی که چشم‌اندازی بین‌المللی و بدون مرز ندارند، بقا نخواهند یافت (Piroozram, 2017). در این بخش، سه مطالعه موردی درباره زاستی^۴، ژئولکت^۵ و لایر^۶ ارائه می‌شود.

1. AXA XL
2. Expert System
3. Angel Investor
4. Zasti
5. Geollect
6. Layr



قاب ۹- زاستی^۱

زاستی یک پلتفرم فناوری مبتنی بر ابر هوش مصنوعی است که با استفاده از الگوریتم‌های یادگیری عمیق اختصاصی ساخته شده و هدف آن کمک به کسب و کارها و مشتریان آن‌هاست تا بتوانند ریسک را پیش‌بینی کنند و کارایی را بهبود بخشند. این پلتفرم فناوری، راهکارهای تشخیص‌دهنده و پیش‌بینی‌کننده را برای مسائل کسب و کار ارائه می‌دهد. زاستی داده‌های موجود را تحلیل می‌کند، ناهنجاری‌ها را شناسایی می‌کند، الگوهای کاربرد را مشخص و سپس پیش‌بینی‌ها و تشخیص‌های دقیقی را از طریق الگوریتم‌های تنظیم شده به صورت عمودی^۲، ارائه می‌دهد. مفهوم تنظیم عمودی از این استراتژی مشتق شده است که مدل‌های خاصی انتخاب می‌شوند تا یک وظیفه خاص (مثلاً پیش‌بینی شکستن قطعات) را بر مبنای مجموعه داده‌های یک صنعت یا بخش خاص (بازار عمودی)، بهینه کنند. این کار، پلتفرم مذکور را قادر می‌سازد که تشخیص‌ها و پیش‌بینی‌های متنوعی را انجام دهد و از مزیت انتقال یادگیری هم بهره‌بردار. آنچه گفته شد با پلتفرم‌های از پیش ساخته که به صورت غیرسفارشی ساخته می‌شوند و فقط می‌توانند برای برخی مسائل محدود مورد استفاده قرار گیرند، متفاوت است. استفاده از مدل‌سازی اکتشافی برای انتخاب مدل خودکار، منجر به هزینه کمتر، دقت بیشتر و پردازش سریع‌تر می‌شود. مدل‌های اکتشافی از تقریب‌ها برای یافتن راهکارها، به روشی سریع‌تر از روش‌های کلاسیک بهره‌می‌برند.

پلتفرم هوش مصنوعی، کلان‌داده‌های نقاط داده‌ای متعددی مانند پارامترهای بیمه‌نامه، داده‌های خسارت، پارامترهای هواشناسی، داده‌های جرایم، اینترنت اشیا و داده‌های حسگرها را پردازش می‌کند. پلتفرم می‌تواند نسبت‌های ریسک و پروفایل‌های ریسکی را شناسایی کند که بیمه‌گران را قادر می‌سازد تا بیمه‌نامه‌ها را برای هر مشتری به صورت بلادرنگ شخصی‌سازی کند.

در بخش بیمه، زاستی راهکارهایی را در رشته‌های مختلف مانند بیمه‌های املاک تجاری، هوانوردی، درمان و اتومبیل ارائه می‌کند که کاربردهای آن به شرح زیر هستند:

- زاستی از تحلیل پیش‌بینی‌گر برای پیش‌بینی شکستن قطعات حیاتی یا حوادث زنجیره تامین و ارائه هوش پیشگیرانه^۳ استفاده می‌کند که بیمه‌گران را قادر می‌سازد از عدم‌النفع ناشی از تاخیر پرواز، آتش‌سوزی، حوادث زنجیره تامین، آب‌وهوا یا شکستن ماشین‌آلات اجتناب کنند. این کار به شرکت‌های بیمه اجازه می‌دهد پرداخت‌های خسارات پرهزینه را کم کرده یا به صورت بلادرنگ بیمه‌نامه‌های جدید شخصی‌سازی شده ارائه کنند.

- در هوانوردی، زاستی با پردازش داده‌های تاخیر پرواز، داده‌های خسارات و آب‌وهوا در طول ۵ سال، مدلی را ساخته است. این مدل می‌تواند تاخیرهای پرواز را برای یک پرواز، ایرلاین یا فرودگاه خاص و برای بازه زمانی یک هفته، یک ماه، یک فصل، نیم سال یا یک سال پیش‌بینی کند. این اطلاعات به بیمه‌گران کمک خواهند کرد که خسارات را کاهش داده و بیش‌تری درباره خسارات برآورد شده برای یک بازه زمانی خاص را به دست آورند. این امکان بیمه‌گران را قادر می‌سازد با به اشتراک‌گذاری هوش اجتناب از ریسک، به مشتریان خود ارزش ارائه دهند.

- برای بیمه اتومبیل، پلتفرم هوش مصنوعی از یک شبکه عصبی پیچشی^۴ برای شناسایی خسارات یک اتومبیل بر مبنای ویدئوها و تصاویر استفاده می‌کند. وقتی مدل برای یک کشور، وسیله نقلیه، مدل‌ها و ویژگی‌های خاص آموزش دید، می‌تواند میزان خسارت را شناسایی و تعیین کند. بر اساس تصویر، هوش مصنوعی می‌تواند مشخص کند که آیا تصویر متعلق به خودروی بیمه‌شده است یا خیر و بدین ترتیب امکان تشخیص تقلب را فراهم نماید.

1. ZASTI
2. Vertically-tuned Algorithms
3. Provide Preventive
4. Convolutional Neural Network



قاب ۱۰- لایر

لایر یک پلتفرم مبتنی بر ابر است که کسب و کارهای کوچک می‌توانند روی آن بیمه مسئولیت بخرند یا آن را مدیریت کنند. لایر از یک مدل توزیع بیمه‌ای استفاده می‌کند که به مالکان کسب و کارها اجازه می‌دهد، استقلال بیشتری روی بیمه‌ای که خریداری می‌کنند داشته باشند. این حالت از طریق یک مولد^۱ پیشنهاد بیمه‌ای انجام می‌شود که می‌تواند به صورت بلادرنگ مظنه بدهد، بهترین شرکت را برای ارائه پوشش شناسایی نموده و ظرف مدت کوتاه ۱۲ دقیقه، پوشش را تامین کند. مالک کسب و کار می‌تواند مبلغ و وسعت پوشش در میان رشته‌های مختلف مانند مسئولیت عمومی، مسئولیت حرفه‌ای، مسئولیت کارفرما، سائیری، اموال، و .. را تنها با چند کلیک متناسب‌سازی کند.

لایر در حال استفاده از هوش مصنوعی برای این است که با اثربخشی هزینه‌ای و از طریق پیش‌بینی قیمت‌گذاری و بیمه‌گر مناسب، وارد بازار تجارت‌های کوچکی شود که فروش کمی در آن اتفاق می‌افتد و حاشیه سود آن نیز پایین است. فناوری لایر یاد می‌گیرد که کدام بیمه‌گر می‌خواهد کدام نمایه‌ها را بیمه کند و در چه قیمتی، سپس فرصت‌های دقیق مورد نظر خود را با اختیارات لازم به بیمه‌گران ارائه می‌دهد. به علاوه، با رشد مجموعه داده‌های لایر، آن‌ها از هوش مصنوعی برای شناسایی ریسک به صورت بلادرنگ و پیش‌بینی مشتریان کسب و کارهای خرد استفاده خواهند کرد که این کار به آن‌ها امکان خواهد داد فروش جانبی و بیش‌فروشی را هم خودکارسازی کنند.

در نهایت، لایر در فرایند یکپارچگی با خدمات مالی کسب و کارهای کوچک، از سامانه‌های آنلاین دفترداری^۲ یا پردازشگرهای پرداخت^۳، قرار دارد تا به صورت خودکار داده‌های تایید اعتبار شده اشخاص ثالث را بازیابی کند و بتواند از آن‌ها جهت صدور بیمه برای مشتریان تجاری کوچک استفاده کرده و بیمه‌گری آسان و کارایی را تسهیل نماید.

منبع: withlayr.com

قاب ۱۱- ژئولکت^۴

ژئولکت یک هوش جغرافیایی و یک شرکت تحلیل‌گر است که بینش‌هایی را از منابع متعدد داده‌ای مانند تصاویر ماهواره، داده‌های موقعیتی، رسانه‌های اجتماعی و سایر جریانات منبع‌باز گردآوری می‌کند. هوش جغرافیایی، دربرگیرنده کاربرد الگوریتم‌های پویا برای مکان‌یابی، کمی‌سازی، توصیف و درک داده‌های فضایی است.

داده‌های تحلیل‌شده ممکن است بلادرنگ، تاریخی، فضایی یا حتی ماهیتاً غیرفضایی باشند. تحلیل داده‌ها فراتر از روش‌های آماری سنتی است و از الگوریتم‌های یادگیری ماشین برای ارائه بینش‌های جدید از داده‌ها استفاده می‌کند. با رشد مستمر عمق و وسعت مجموعه داده‌ها، الگوریتم‌های ژئولکت هم به صورت مستمر بهبود یافته و نتایجی را ارائه می‌کنند که به‌طور فزاینده مستحکم و مفصل‌تر می‌شوند. روابط، الگوها، ریسک‌ها و فرصت‌هایی که از قبل شناخته‌شده هستند، از طریق یک داشبورد تعاملی، بصری‌سازی^۵ می‌شوند. این خروجی‌ها می‌توانند برای کمک به ارزیابی ریسک و تصمیم‌گیری، مورد پرس‌جو و دستکاری قرار بگیرند.

مثال خوبی از مزایای روش‌های ژئولکت، در تایید اعتبار دعاوی خسارت است. هوش جغرافیایی ژئولکت می‌تواند نتایج سریع،

1. Generator
2. Bookkeeping
3. Payment Processors
4. Geollect
5. Visualization



کمّی شده و بسیار دقیقی را از مکان‌های دور فراهم کند و فرآیند را تسریع نماید که این کار به کاهش هزینه‌ها کمک می‌کند. یادگیری ماشین، با ایجاد امکان خودکارسازی و معتبر کردن این فرآیند، نقش مهمی را در این میان ایفا می‌کند. تا امروز، ژئولکت ارزش این رویکرد را در بیمه و دفاع و تامین امنیت صنایع مختلف نشان داده است. مثلاً، ژئولکت در حال حاضر توسط کلوب P&I انگلستان استفاده می‌شود و اطلاعاتی را درباره فعالیت، زیرساخت و ریسک‌های بندر ارائه می‌کند. منبع: geollect.com

نتیجه‌گیری

آینده

سامانه‌های امروزی هوش مصنوعی، حول فناوری‌های یادگیری ماشین (مانند یادگیری عمیق) ساخته می‌شوند که می‌توانند تصمیمات خودکاری را بر مبنای مقادیر بسیار زیادی از داده‌ها اتخاذ کنند. این تصمیمات، تعمیم‌های مبتنی بر نمونه‌هایی هستند که در طول آموزش سامانه توسط انسان‌ها تهیه می‌شوند. عملکرد این سامانه‌ها نه تنها به موجود بودن مثال‌های آموزشی وابسته است، بلکه پیوند محکمی هم با حوزه‌ای دارد که سامانه در آن عمل می‌کند (مثلاً صدور بیمه، پیشنهاد محصول و ...). هر چه دامنه محدودتر و منسجم‌تر تعریف شود، هوش مصنوعی بهتر عمل خواهد کرد؛ زیرا پوشش داده‌های آموزشی بهتر خواهد بود و بنابراین بیشتر نمایای موقعیت‌هایی است که هوش مصنوعی با آن مواجه می‌شود.

در آینده، هوش مصنوعی باید بر این وابستگی بنیادین به نظارت انسانی در طول آموزش غلبه کند، تا در حوزه‌ها و دامنه‌های وسیع‌تری (که حاوی موقعیت‌های کمتر قابل پیش‌بینی هستند) قابل استفاده باشد. آموزش بدون نظارت سامانه‌های هوش مصنوعی، حوزه تحقیقاتی فعالی است که احتمالاً منجر به پیشرفت‌هایی در پنج سال آینده خواهد شد. این شیوه آموزش هنوز هم نیازمند حجم‌های بزرگی از داده‌ها است که باید در طول زمان آموزش به آن ارائه شوند؛ با این حال، آن داده‌ها نیازمند تفسیر انسان نیستند. آن‌چه گفته شد مثلاً می‌تواند سامانه هوش مصنوعی را قادر سازد رفتارهای ناهنجار یا غیرمعمول را که ممکن است ریسک بیشتری را در طول صدور بیمه ایجاد کنند، شناسایی نماید. آموزش بدون نظارت، تناسب خوبی هم با سناریوهایی دارد که در آن حجم‌های بزرگی از داده‌ها قابل گردآوری هستند (مثلاً از اینترنت یا دوربین‌های وسایل نقلیه)، اما تفسیرشان غیرممکن است.

حوزه تحقیقاتی فعال دیگر در زمینه هوش مصنوعی، یادگیری با نظارت ضعیف است که در آن، فقط داده‌هایی که به میزان کم یا به‌طور غیردقیق تفسیر شده‌اند، می‌توانند مورد استفاده قرار گیرند تا هوش مصنوعی را آموزش دهند. با این وجود، چالش‌های بلندمدت‌تر در جایی برای هوش مصنوعی باقی خواهند ماند که داده‌های آموزشی بسیار کمی وجود داشته باشند، چه در تشخیص موارد نادر (مثلاً در حوزه پزشکی) چه در جایی که تفسیر داده‌ها در مقیاس زیاد، به‌طور بازدارنده‌ای گران یا به دلایل مربوط به حریم خصوصی، غیرعملی است.

جریان اصلی رسانه‌ای به طور مکرر پیشرفت‌های دانشگاه‌ها و شرکت‌هایی مانند گوگل دیپ‌ماینده^۱ را در استفاده از هوش مصنوعی برای بازهای ویدئویی یا بازی‌های کلاسیکی مانند شطرنج، گزارش می‌کنند. انگیزه استفاده از چنین بازی‌هایی برای مطالعه تکنیک‌های نوظهور هوش مصنوعی این است که آن‌ها نماینده محیط‌های کنترل‌شده‌ای هستند که در آن‌ها تنها گستره محدودی از اقدامات می‌توانند انجام شوند. فضای اقدامات ممکن است که می‌توانند در دنیای واقعی انجام شوند، بسیار بزرگ‌تر است و مثلاً ناوبری رباتیک یا تدارکات عملیاتی و پیشرفت‌های جدید، باید برای مقابله با طیف وسیع‌تری از اقدامات و رویدادهایی که یک سامانه هوش مصنوعی ممکن است در دنیای واقعی با آن روبرو شود، صورت گیرند. تعمیم تکنیک‌هایی مانند یادگیری کیوی عمیق^۲ برای گسیل شدن در دنیای واقعی، حوزه‌ای فعال است که احتمالاً در پنج سال آینده به ثمر می‌نشیند. در بلندمدت‌تر، شاهد جابه‌جایی به سمت هوش مصنوعی قدرتمندی خواهیم بود که در آن تشخیص خردمندی سامانه از هوش انسانی غیر قابل تمایز است.

نتایج

به عنوان نتیجه باید گفت که هوش مصنوعی در حال اثرگذاری بر همه بخش‌ها و همه ابعاد دنیای ماست و در آینده نزدیک هم به این تأثیرات ادامه خواهد داد. در بیمه، هوش مصنوعی از طریق پیشگیری از ریسک، بر همه مراحل زنجیره ارزش بیمه، یعنی از اعلام‌های اولیه تا تسویه خسارات تأثیر خواهد گذاشت. واردشوندگان جدید، برهم‌زنندگان و بازیگران قدیمی بازار، همگی در مورد نقش‌ها و وابستگی متقابلشان بازاندیشی خواهند کرد. از آنجایی که کسب‌وکارها به‌طور روزافزونی هوش مصنوعی را وارد سامانه‌ها و فرایندهایشان می‌کنند، به بیمه نیاز دارند تا آن‌ها را از دامنه وسیعی از ریسک‌های بالقوه محافظت کند. بنابراین، با این که هوش مصنوعی پتانسیل عظیمی دارد، اما بیمه‌گران باید از ریسک‌های مرتبط با استفاده از این فناوری جدید که سرعت پیشرفت زیادی هم دارد، آگاه باشند و با توسعه مدل‌ها و محصولات مناسب، به تقاضاهای جدید پاسخ دهند.

1. Google DeepMind

۲. Deep Q-learning: تکنیک یادگیری تقویتی است که با یادگیری یک تابع اقدام/مقدار، سیاست مشخصی را برای انجام حرکات مختلف در وضعیت‌های مختلف دنبال می‌کند. یکی از نقاط قوت این روش، توانایی یادگیری تابع مذکور بدون داشتن مدل معینی از محیط می‌باشد. [م]



منابع

- Abbott, R. (2018). The Reasonable Computer: Disrupting the Paradigm of Tort Liability, 86 Geo. Wash. L. Rev.1.
- Abbott, R and Bogenschneider, BN (2018) Should Robots Pay Taxes? Tax Policy in the Age of Automation Harvard Law & Policy Review, 12 (1). pp. 145-175. Retrieved from <http://epubs.surrey.ac.uk/821099>.
- Abbott, R. (2019) Everything is Obvious, 66 UCLA. L. Rev. 2.
- Ackerman, E (2071), Slight Street Sign Modifications Can Completely Fool Machine Learning Algorithms, IEE Spectrum Journal, Retrieved 6 Feb 2019 from <https://spectrum.ieee.org/cars-that-think/transportation/sensors/slight-street-signmodifications-can-fool-machine-learning-algorithms>.
- AI HLEG, (2018), Draft Ethics Guidelines for Trustworthy AI, Retrieved 6 Feb 2019 from https://ec.europa.eu/futurium/en/system/files/ged/ai_hleg_draft_ethics_guidelines_18_december.pdf.
- AI powered 'Robo-Lawyer' helps step up the SFO's fight against economic crime. (2018). Retrieved June 24, 2018, from <https://www.sfo.gov.uk/2018/4/10/ai-powered-robo-lawyer-helps-step-up-the-sfos-fight-against-economic-crime>.
- Artificial Intelligence Market in Insurtech: By Type; By Application - Forecast (2017-2022). (2018). Retrieved April 1, 2018, from <https://www.researchandmarkets.com/research/23dqwr/artificial>.
- AXA XL, (2019), XL Catlin Signs Landmark Agreement With Oxbotica, Retrieved 16 Aug 2018 from <https://axaxl.com/media/xl-catlin-signs-landmark-agreement-with-oxbotica>, Published: January 13, 2016, Author: Sinead Finlay AXA XL.
- Azadian, B. J. S., & Fahy, G. M. (2018). Artificial Intelligence Artificial Intelligence and the Law : Navigating "Known Unknowns," 35 (2).
- Battista, B, and Roli, F, (2018) "Wild patterns: Ten years after the rise of adversarial machine learning." Pattern Recognition 84: 317-331.
- BBC, 2017. AI image recognition fooled by single pixel change. Available at: <https://www.bbc.co.uk/news/technology-41845878>.
- BBC, 2018. Amazon scrapped 'sexist AI' tool, Retrieved from <https://www.bbc.co.uk/news/technology-45809919>, 10 October 2018.
- BBC, 2019. Computer virus alters cancer scan images, Retrieved from <https://www.bbc.co.uk/news/technology-47812475>.
- Bouchti, A. El, Chakroun, A., Abbar, H., & Okar, C. (2017). Fraud detection in banking using deep reinforcement learning. 2017 Seventh International Conference on Innovative Computing Technology (INTECH), (Intech), 58-63. <https://doi.org/10/1109/INTECH.2017/80102446>.
- Brookings Institution, (2014), Products Liability and Driverless Cars: Issues and Guiding Principles for Legislation, John Villasenor, April 2014, Retrieved 16 August 2018 from https://www.brookings.edu/wpcontent/uploads/2016/06/Products_Liability_and_Driverless_Cars.pdf.

فراخوان ارسال مقاله

پژوهشکده بیمه وابسته به بیمه مرکزی جمهوری اسلامی ایران با هدف ارتقاء، بسط، گسترش و نهادینه کردن علم بیمه با رویکرد مطالعه موردی یک موضوع خاص بیمه‌ای و تحلیل مباحث آن، نشریه «گزارش موردی» را منتشر می‌کند. این نشریه قابل استفاده برای کسانی است که به دنبال مباحث خاص بیمه‌ای به صورت تئوریک هستند که از آن میان می‌توان به دانشجویان بیمه و اقتصاد، مدیران عالی رتبه صنعت بیمه کشور، اساتید دانشگاه‌ها و دست‌اندرکاران صنعت بیمه اشاره کرد؛ لذا از کلیه استادان، پژوهشگران، صاحب‌نظران و کارشناسان محترم برای ارائه مقالات دعوت به عمل می‌آید.

الف. شرایط پذیرش مقاله

۱. مقالات می‌توانند به صورت تألیفی یا ترجمه باشند. باید همراه با مقالات ترجمه شده، نسخه اصلی آنها نیز ارسال شود.
۲. مقالات باید به مطالعه موردی یک موضوع خاص بیمه‌ای پردازند.
۳. حجم مقالات باید با توجه به شرایط مندرج در بند «ب» حداقل ۴۰ صفحه باشد.
۴. مقالات ارسالی نباید قبلاً در نشریه‌های داخلی و خارجی یا مجموعه مقالات سمینارها و مجامع علمی چاپ شده باشند و نباید همچنین برای انتشار به جای دیگر واگذار شده باشند.
۵. مقالات باید دارای فهرست منابع و مأخذ مستند و اطلاعات کتاب‌شناختی معتبر باشند.
۶. مقالات ترجمه‌ای حداکثر در سال ۲۰۱۵ چاپ شده باشند (مگر در موارد خاص و با تأیید داور).
۷. مقالات ارسال شده پس از طی فرایند داوری و تأیید سردبیر نشریه قابل پذیرش و چاپ می‌باشند.
۸. حق ویرایش مقالات برای نشریه محفوظ است.
۹. مسئولیت مطالب، نظریات و اطلاعات ارائه شده در مقاله‌ها و صحت و سقم آنها بر عهده مؤلف(ان)/ مترجم(ان) است.
۱۰. دریافت مقاله به صورت الکترونیکی امکان‌پذیر است.
۱۱. مقالات دریافت شده به مؤلف(ان)/ مترجم(ان) بازگردانده نمی‌شوند.

ب. نحوه نگارش مقاله

۱. مقاله حداقل در ۵۰ صفحه A4 با فاصله خطوط multiple 1.1 و حاشیه‌های ۲ سانتی‌متر از هر طرف در نرم‌افزار Word تایپ شود.
۲. نوع قلم و اندازه آن مطابق با شرایط مندرج در جدول (۱) باشد.
۳. اصول نگارش زبان فارسی به طور کامل رعایت شود و از به کار بردن اصطلاحات انگلیسی که معادل فارسی آنها در فرهنگستان زبان فارسی تعریف شده است، حتی‌الامکان خودداری شود.

فراخوان ارسال مقاله

جدول (۱). نوع قلم و اندازه

موقعیت استفاده	نام قلم	اندازه قلم
عنوان مقاله	B lotus پررنگ	۱۶
متن مقاله	B lotus	۱۴
تیترهای اصلی	B lotus پررنگ	۱۵
تیترهای فرعی	B lotus پررنگ	۱۴
عناوین جدول‌ها و شکل‌ها	B lotus	۱۴
متن جدول‌ها، شکل‌ها و منابع	B lotus	۱۴
پاورقی فارسی	B lotus	۱۱
پاورقی انگلیسی	Times New Roman	۱۰

ج. شیوه تنظیم منابع

در ذکر منابع، سبک Harvard (ویرایش سال ۲۰۱۱) رعایت شود؛ به عنوان مثال:

• منابع انتهایی متن

- کتاب

نام خانوادگی، حرف ابتدای نام، سال انتشار. عنوان کتاب (ایتالیک). نام مترجم. ویرایش (اگر کتاب ویرایش اول باشد، نیازی به ذکر نوبت ویرایش نیست). محل نشر (باید نام شهر ذکر شود نه کشور): ناشر. شماره صفحات (در صورت موجود بودن).

- Redman, P., 2008. *Business and the organization*. 3rd ed. Chester: Pearson.

- جهانخانی، ع. و پارسائیان، ع.، ۱۳۷۶. مدیریت سرمایه‌گذاری و ارزیابی اوراق بهادار. تهران: دانشکده مدیریت. صص ۵-۲۰.

توجه: اگر منبعی بیش از یک نویسنده داشته باشد، نام نویسندگان را به ترتیب زیر ذکر می‌کنیم:

- Barker, R., Kirk, J. and Munday, R.J., 1988. *Narrative analysis*. 3rd ed. Bloomington: Indiana University Press.

در صورتی که منبع بیش از سه نویسنده داشته باشد، لزومی به ذکر نام تمام نویسندگان نیست و فقط نام نویسنده اول را ذکر می‌کنیم و به جای نام نویسندگان دیگر از عبارت "et al." همکاران استفاده می‌کنیم.

- Grace, B. et al., 1988. *A history of the world*. Princeton, NJ: Princeton University Press.

- اگر در میان منابع مقاله، دو منبع با تاریخ انتشار و نام نویسنده مشابه داشته باشیم برای تفکیک منابع از حروف

فراخوان ارسال مقاله

"a,b" / الف، ب" در کنار تاریخ انتشار استفاده می‌کنیم و در منابع درون‌متنی نیز این حروف را در کنار تاریخ انتشار ذکر می‌کنیم.

- Soros, G., 1966a. *The road to serfdom*. Chicago: University of Chicago Press.

- Soros, B., 1966b. *Beyond the road to serfdom*. Chicago: University of Chicago Press.

و در منابع درون‌متنی، این دو منبع را به این شکل ذکر می‌کنیم:

-(Soros, 1966a)

-(Soros, 1966b)

- منابع الکترونیکی (E-book; pdf; journal online) ...

نام خانوادگی، حرف ابتدای نام، سال. عنوان (ایتالیک). [نوع منبع الکترونیکی] محل نشر: ناشر. Available through (قابل دسترسی از طریق): نام منبع الکترونیکی (شامل وب‌سایت یا آدرس اینترنتی) [تاریخ دسترسی].
- Fishman, R., 2005. *The rise and Fall of suburbia*. [e-book] Chester: Castle Press. Available through: Anglia Ruskin University Library website <<http://libweb.anglia.ac.uk>> [Accessed 5 June 2005].

- مقاله

نام خانوادگی و حرف ابتدای نام، سال. عنوان مقاله. عنوان نشریه (ایتالیک)، شماره جلد (شماره بخش یا فصل انتشار)، صفحه یا صفحات.

- Boughton, J.M., 2002. The Bretton Woods proposal: a brief look. *Political Science Quarterly*, 42(6), p.564.

• منابع درون‌متنی

(نام خانوادگی نویسنده، سال انتشار اثر)

- (کریمی، ۱۳۸۷)؛ (Wang, 2008)

توجه: اگر اثری بیش از چهار نویسنده داشته باشد، فقط نام نویسنده اول را ذکر می‌کنیم و به جای نام دیگر نویسندگان از عبارت "et al." / همکاران" استفاده می‌کنیم:

- (نفیسی و همکاران، ۱۳۸۹)؛ (Soha, et al., 1995)

برای اطلاعات بیشتر در زمینه سبک Harvard به آدرس الکترونیکی زیر مراجعه کنید:

<http://libweb.anglia.ac>.

علاقه‌مندان برای دریافت اطلاعات تکمیلی می‌توانند به نشانی زیر مراجعه فرمایند:

آدرس: تهران - سعادت آباد - میدان شهید تهرانی مقدم (کاج) - خیابان سرو غربی - شماره ۴۳ - دفتر نشریه گزارش موردی - شماره تماس ۲۲۰۸۴۰۸۴. جهت مکاتبه با نشریه به آدرس الکترونیکی workingpaper@irc.ac.ir مراجعه نمایید.

فهرست گزارش‌های منتشرشده در پژوهشکده بیمه

- گزارش موردی ۱ (دی ۱۳۸۹): کلیات اقتصاد برنامه‌های بیمه اجتماعی
- گزارش موردی ۲ (اسفند ۱۳۸۹): آمارهای حوادث جاده‌ای در کشورهای منتخب و تحلیل خسارت‌های پرداختی بیمه شخص ثالث در ایران
- گزارش موردی ۳ (فروردین و اردیبهشت ۱۳۹۰): اوراق بهادار بیمه‌ای
- گزارش موردی ۴ (خرداد و تیر ۱۳۹۰): نقش شاخص‌ها در انتقال ریسک در صنعت بیمه
- گزارش موردی ۵ (مرداد و شهریور ۱۳۹۰): شاخص‌های پایه‌ای نرخ بیمه زلزله ساختمان‌های ایران
- گزارش موردی ۶ (مهر و آبان ۱۳۹۰): اصلاح سیستم خدمات درمانی در ژاپن: کنترل هزینه‌ها، ارتقای کیفیت و تضمین برابری
- گزارش موردی ۷ (آذر و دی ۱۳۹۰): بیمه در کشورهای در حال توسعه: بهره‌گیری از فرصت‌های موجود در بیمه‌های خرد
- گزارش موردی ۸ (بهمن و اسفند ۱۳۹۰): پولشویی و روش‌های جلوگیری از آن در صنعت بیمه
- گزارش موردی ۹ (فروردین و اردیبهشت ۱۳۹۱): کاربرد ملی مقررات ساختمان در مدیریت ریسک و نرخ‌گذاری بیمه آتش‌سوزی
- گزارش موردی ۱۰ (خرداد و تیر ۱۳۹۱): پیشگیری شناسایی و مقابله با کلاهبرداری در بیمه
- گزارش موردی ۱۱ (مرداد و شهریور ۱۳۹۱): تجربه کشور هندوستان در حذف تعرفه‌های بیمه‌های غیرزندگی
- گزارش موردی ۱۲ (مهر و آبان ۱۳۹۱): تدوین بیمه‌نامه زلزله در بخش مسکن و ارائه مدلی کاربردی جهت بررسی نقش بیمه در بهبود کیفیت ساختمان در ایران
- گزارش موردی ۱۳ (آذر و دی ۱۳۹۱): رابطه بیمه و رشد اقتصادی- تحلیل نظری و تجربی
- گزارش موردی ۱۴ (بهمن و اسفند ۱۳۹۱): ارزیابی و تحلیل ریسک قراردادهای بیمه زندگی: ترکیب رویکردهای اکچوئرال و مالی
- گزارش موردی ۱۵ (فروردین و اردیبهشت ۱۳۹۲): دوگزارش بیمه‌ای: ارزیابی عملکرد صنعت بیمه کشور و تبیین چشم‌انداز آینده (مقاله اول)- بررسی و سنجش سطح رضایت‌مندی بیمه‌گذاران (مشتریان) شرکت‌های فعال در صنعت بیمه کشور (مقاله دوم)
- گزارش موردی ۱۶ (خرداد و تیر ۱۳۹۲): بیمه سلامت و بیمه نوین سلامت (مطالعه موردی: کشورهای چین، ژاپن و کره جنوبی)
- گزارش موردی ۱۷ (مرداد و شهریور ۱۳۹۲): راهکارهای عملی افزایش تقاضای بیمه زندگی انفرادی و تدوین چهارچوبی برای ارائه بیمه‌های زندگی جدید
- گزارش موردی ۱۸ (مهر و آبان ۱۳۹۲): چهارچوب نظارتی انجمن بین‌المللی ناظران بیمه
- گزارش موردی ۱۹ (آذر و دی ۱۳۹۲): کاربرد منطق فازی، شبکه عصبی و الگوریتم ژنتیک در صنعت بیمه
- گزارش موردی ۲۰ (بهمن و اسفند ۱۳۹۲): بررسی عملکرد ساختمان‌ها و تأسیسات زیربنایی و عملکرد صنعت بیمه در زلزله
- مرداد ۱۳۹۱ آذربایجان شرقی

گزارش موردی ۲۱ (فروردین و اردیبهشت ۱۳۹۳): مطالعه تطبیقی مقررات‌گذاری و نظارت صنعت بیمه در کشورهای منتخب توسعه یافته

گزارش موردی ۲۲ (خرداد و تیر ۱۳۹۳): مطالعه عوامل ریسک و فاکتورهای مؤثر بر محاسبه حق‌بیمه در بیمه‌های اتومبیل در ایران و جهان

گزارش موردی ۲۳ (مرداد و شهریور ۱۳۹۳): مطالعه روش‌های محاسباتی و ارزیابی ریسک‌های مختلف بیمه‌های زندگی

گزارش موردی ۲۴ (مهر و آبان ۱۳۹۳): مشتری‌مداری در صنعت بیمه

گزارش موردی ۲۵ (آذر و دی ۱۳۹۳): بیمه ریسک در قراردادهای تأمین مالی پروژه.

گزارش موردی ۲۶ (بهمن و اسفند ۱۳۹۳): تعیین روش بهینه محاسبه سیستم پاداش – جریمه در بیمه‌های شخص ثالث

گزارش موردی ۲۷ (فروردین و اردیبهشت ۱۳۹۴): حفظ سلامت افراد در بازارهای نوظهور با حمایت بیمه

گزارش موردی ۲۸ (خرداد و تیر ۱۳۹۴): سیری در پیشرفت‌های اخیر بیمه‌های دریایی و هوایی

گزارش موردی ۲۹ (مرداد و شهریور ۱۳۹۴): بیمه‌گری ریسک‌های تجاری دائم‌التغییر

گزارش موردی ۳۰ (مهر و آبان ۱۳۹۴): بررسی وضعیت اتکایی اجباری در ایران و کشورهای منتخب

گزارش موردی ۳۱ (آذر و دی ۱۳۹۴): بررسی ساختار و کارکرد ناظر بیمه‌ای در ایران و کشورهای منتخب

گزارش موردی ۳۲ (بهمن و اسفند ۱۳۹۴): معرفی قراردادهای بیمه زندگی متصل به سهام و روش‌های قیمت‌گذاری

گزارش موردی ۳۳ (فروردین و اردیبهشت ۱۳۹۵): بازار بیمه وسایل نقلیه موتوری در اروپا، نوامبر ۲۰۱۵ (قسمت اول)

گزارش موردی ۳۴ (خرداد و تیر ۱۳۹۵): بازار بیمه وسایل نقلیه موتوری در اروپا، نوامبر ۲۰۱۵ (قسمت دوم)

گزارش موردی ۳۵ (مرداد و شهریور ۱۳۹۵): مطالعه تطبیقی شرایط عمومی بیمه‌نامه آتش‌سوزی و ارائه پیشنهادها اصلاحی

گزارش موردی ۳۶ (مهر و آبان ۱۳۹۵): مطالعه تطبیقی طرح‌های مستمری بازنشستگی خصوصی در کشورهای منتخب

گزارش موردی ۳۷ (آذر و دی ۱۳۹۵): مطالعه عوامل ریسک و فاکتورهای مؤثر بر محاسبه حق‌بیمه در رشته بیمه‌های باربری

گزارش موردی ۳۸ (بهمن و اسفند ۱۳۹۵): حاکمیت شرکتی در مؤسسات بیمه

گزارش موردی ۳۹ (فروردین و اردیبهشت ۱۳۹۶): بررسی شرایط عمومی بیمه بدنه اتومبیل در ایران و سایر کشورها

گزارش موردی ۴۰ (خرداد و تیر ۱۳۹۶): بررسی شرایط پوشش حوادث راننده توسط بیمه‌گران خارجی و ارائه شرایط عمومی

پیشنهادی برای ایران

گزارش موردی ۴۱ (مرداد و شهریور ۱۳۹۶): مطالعه شرایط عمومی بیمه‌های باربری و ارائه پیشنهادها اصلاحی

گزارش موردی ۴۲ (مهر و آبان ۱۳۹۶): یافته‌های تجربی در خصوص قیمت‌گذاری بیمه وسایل نقلیه موتوری در آلمان، اتریش و

سوئیس

گزارش موردی ۴۳ (آذر و دی ۱۳۹۶): مدیریت ریسک‌های کلیدی

گزارش موردی ۴۴ (بهمن و اسفند ۱۳۹۶): اصول اساسی بیمه انجمن بین‌المللی ناظران بیمه (اصول ۱ الی ۱۳)



- گزارش موردی ۴۵ (فروردین و اردیبهشت ۱۳۹۷): اصول اساسی بیمه انجمن بین‌المللی ناظران بیمه (اصول ۱۴ الی ۱۷)
- گزارش موردی ۴۶ (خرداد و تیر ۱۳۹۷): اصول اساسی بیمه انجمن بین‌المللی ناظران بیمه
- گزارش موردی ۴۷ (مرداد و شهریور ۱۳۹۷): فن‌آوری و نوآوری در بخش بیمه
- گزارش موردی ۴۸ (مهر و آبان ۱۳۹۷): راهنمای نظام نظارت مبتنی بر ریسک برای شرکت‌های بیمه
- گزارش موردی ۴۹ (ویژه‌نامه سمینار توسعه بیمه‌های زندگی): بیمه زندگی: تمرکز بر مصرف‌کننده
- گزارش موردی ۵۰ (ویژه‌نامه بیست و پنجمین همایش بین‌المللی بیمه و توسعه): تحولات فناوری مالی در صنعت بیمه
- گزارش موردی ۵۱ (بهمن و اسفند ۱۳۹۷): شناسایی و مطالعه ضوابط و مقررات نحوه حضور مؤسسات بیمه خارجی در کشور چین
- گزارش موردی ۵۲ (فروردین و اردیبهشت ۱۳۹۸): کلان داده و بیمه: پیامدهایی برای نوآوری، رقابت و حریم خصوصی
- گزارش موردی ۵۳ (خرداد و تیر ۱۳۹۸): گزارش توانگری و شرایط مالی شرکت هانوفر ری
- گزارش موردی ۵۴ (مرداد و شهریور ۱۳۹۸): گزارش تحلیل مغایرت با استانداردهای حسابداری در خصوص ذخیره فنی تکمیلی و خطرات طبیعی و ارائه پیشنهاد اصلاحی
- گزارش موردی ۵۵ (مهر و آبان ۱۳۹۸): سیستم شناسایی تقلب بیمه‌ای
- گزارش موردی ۵۶ (آذر و دی ۱۳۹۸): بررسی عملکرد ماده ۳۰ قانون الحاق برخی مواد به قانون تنظیم بخشی از مقررات مالی دولت (۲) و بند الف تبصره ۱۰ قانون بودجه سالیانه (عوارض قانونی بیمه اجباری شخص ثالث)
- گزارش موردی ۵۷ (بهمن و اسفند ۱۳۹۸): بیمه زندگی در عصر دیجیتال؛ تحولات بنیادی پیش‌رو
- گزارش موردی ۵۸ (فروردین و اردیبهشت ۱۳۹۹): گزارش ریسک‌های آینده گروه آکسا و اوراسیا
- گزارش موردی ۵۹ (خرداد و تیر ۱۳۹۹): وضعیت حاکمیت شرکتی در کشورهای عضو سازمان همکاری و توسعه اقتصادی در سال ۲۰۱۷
- گزارش موردی ۶۰ (مرداد و شهریور ۱۳۹۹): درک سودآوری در بیمه زندگی
- گزارش موردی ۶۱ (مهر و آبان ۱۳۹۹): بررسی اجمالی شبکه فروش صنعت بیمه و چگونگی نظارت بر آن در کشورهای توسعه یافته و در حال توسعه
- گزارش موردی ۶۲ (آذر و دی ۱۳۹۹): ضرورت بهره‌وری در صنعت بیمه
- گزارش موردی ۶۳ (فروردین و اردیبهشت ۱۴۰۰): تحول دیجیتال صنعت بیمه در کشورهای منتخب: فناوری‌ها، قوانین و مقررات و نهادسازی
- گزارش موردی ۶۴ (خرداد و تیر ۱۴۰۰): تغییر اقلیم و صنعت بیمه: انجام اقداماتی به عنوان مدیران ریسک و سرمایه‌گذاران از دیدگاه مدیران سطح C در صنعت بیمه

Working Paper

Bimonthly Journal of Insurance Research Center (IRC)

New Edition, Vol.11, No.3, August & September, 2021, Serial No. 65

Concessioner: Insurance Research Center (IRC)

Director- in- Charge: Leili Niakan (PhD)

Editor- in- Chief: Hamid Kordbache (PhD)

Translator: Vahideh Nourani

Scientific Editor: Mohammad Javad Heidari (PhD)

Internal Manager: Shabnam Refoua (PhD)

Layout & Cover Designer: Ali hossein Safari

Publishing Technical Supervisor: Mohsen Bagherian

Publication Address: No. 43, West Sarv Ave., Saadat Abad, Tehran/ Iran

P.O. Box: 19395 -4499

Tel: (+9821) 22084084

Fax: (+9821) 22092265

Workingpaper@IRC.ac.ir

Printing House Address: Karang Printing House, no. 195, North kaj St., Golha Sq., Fatemi St., Tehran.

Tel: (+9821) 88024010

All rights reserved for IRC. Quoting the content is allowed if acknowledged. The views expressed in this journal are those of the authors and do not necessarily reflect the views of IRC. Working Paper reserves the right to edit the articles.

