

SID



بلاگ مرکز اطلاعات علمی



کارگاه‌های آموزشی



سرویس ترجمه تخصصی



فیلم‌های آموزشی

کارگاه‌ها و فیلم‌های آموزشی مرکز اطلاعات علمی

آشنایی با پایگاه‌های اطلاعات علمی بین‌المللی و ترندهای جستجو بین‌المللی و ترندهای جستجو

کاربرد نرم افزار SPSS در پژوهش

پروپوزال نویسی (علوم انسانی)

کاربرد نرم افزار End Note در استناددهی مقالات و متون علمی

صدور گواهینامه نمایه مقالات نویسندگان در SID



دوره ۶ - شماره ۱۸ - زمستان ۱۴۰۲
ویژه نامه هوش مصنوعی

جایگاه هوش مصنوعی در صحت سنجی ادله دآوری

همایون مافی، فاطمه قناد، محمدمبین اسماعیل پور

هوش مصنوعی به عنوان دلیل در محاکمه کیفری

سالار صادقی

چالش ها و موانع مسئولیت کیفری در ربات های با قابلیت هوش مصنوعی

امین امیریان فارسانی، سیدمحمد حسینی

هوش مصنوعی و تاثیر آن بر سیستم قضایی

امیررضا محمودی، مریم بحرکاظمی

تاریخچه مختصری از هوش مصنوعی: گذشته، حال و آینده هوش مصنوعی

امین حاجی وند، علی خوش منظر، صابر سیاری زهان

هوش مصنوعی در نظام عدالت کیفری: روندها و احتمالات پیشرو

سالار صادقی

هوش مصنوعی و مسئولیت قانونی

سارا صلح چی، کیان بیگلریگی

تعامل هوش مصنوعی و دیپلماسی برای پایداری محیط زیست

سبحان طیبی، نادر طیبی

جرایم هوش مصنوعی یک تحلیل بین رشته ای؛ تهدیدات و راه حل های قابل پیش بینی

زهرا وهبی

هوش مصنوعی و مردم سالاری؛ تأثیر اطلاعات غلط، ربات اجتماعی و هدف گذاری سیاسی

سارا صلح چی

کاربرد هوش مصنوعی در جرم یابی و تحقیقات جنایی؛ نمونه پژوهی: قتل های سریالی

حمیدرضا حیدرپور، محمد شهنقی، ژیللا مهرآرا

مجازانگاری استفاده اخلاقی از هوش مصنوعی با استفاده از نظریه فارابی درباره حقوق طبیعی و سعادت

محمد مهدی داور

هوش مصنوعی در نیروهای مسلح: مروری بر قابلیت ها، کاربردها و چالش ها

یاسر شاکری



Artificial Intelligence and Legal Liability

هوش مصنوعی و مسئولیت قانونی

J. K. C. Kingston
University of Brighton, BN2 4JG, UK

جی. کی. سی. کینگستون
دانشگاه برایتون، BN2 4JG، انگلستان

j.k.kingston@brighton.ac.uk

Sara Solhchi
MA. Intellectual Property Law, Faculty of Law and Political Science, University of Allameh Tabataba'i, Tehran, Iran

سارا صلح چی

کارشناس ارشد حقوق مالکیت فکری، دانشکده حقوق و علوم سیاسی، دانشگاه علامه طباطبائی، تهران، ایران

solhchis@gmail.com
<http://orcid.org/0009-0008-6224-1885>

کیان بیگلربیگی

Kian Biglarbeigi
Master student of International Law, Faculty of Law, Shahid Beheshti University, Tehran, Iran (Responsible Translator)

دانشجوی کارشناسی ارشد حقوق بین‌الملل، دانشکده حقوق، دانشگاه شهید بهشتی، تهران، ایران (مترجم مسئول)

kian.biglarbeigi@gmail.com
<http://orcid.org/0000-0001-8659-0299>

Abstract

A recent issue of a popular computing journal asked which laws would apply if a self-driving car killed a pedestrian. This paper considers the question of legal liability for artificially intelligent computer systems. This article discusses whether criminal liability could ever apply; to whom it might apply; and, under civil law, whether an AI program is a product that is subject to product design legislation or a service to which the tort of negligence applies. The issue of sales warranties is also considered. A discussion of some of the practical limitations that AI systems are subject to is also included.

Keywords: Artificial Intelligence, Civil Liability, Criminal Liability, Tort.

چکیده

یکی از شماره‌های اخیر یکی از مجلات مشهور رایانه‌ای، این پرسش را طرح کرد که اگر یک اتومبیل خودران (رانندگی خودکار بدون راننده) عابر پیاده‌ای را بکشد، چه قوانینی قابل اعمال خواهد بود. این مقاله به بررسی پرسش، پیرامون مسئولیت قانونی در سیستم‌های رایانه‌ای هوش مصنوعی می‌پردازد. در این مقاله بحث می‌شود که آیا مسئولیت کیفری اساساً قابل اعمال است یا خیر؛ بر چه کسی ممکن است اعمال شود؛ و مطابق قانون مدنی، آیا هوش مصنوعی محصولی است که بتواند موضوع قانون طراحی محصول یا خدمت قرار بگیرد به نحوی که شامل شبه جرم ناشی از قصور و غفلت شود یا خیر. مسئله گارانتی‌های فروش نیز مورد توجه قرار گرفته است. همچنین به برخی محدودیت‌های عملی که سیستم‌های هوش مصنوعی دارای آن می‌باشند، پرداخته شده است.

واژگان کلیدی: هوش مصنوعی، مسئولیت مدنی، مسئولیت کیفری، شبه جرم، بدافزار.

Received: 2023/05/28 - Review: 2023/11/10 - Accepted: 2023/12/19

دریافت مقاله: ۱۴۰۲/۰۵/۲۸ - بررسی مقاله: ۱۴۰۲/۱۱/۱۰ - پذیرش مقاله: ۱۴۰۲/۱۲/۱۹

ارجاع:

کینگستون، جی. کی. سی؛ (۱۴۰۲)، هوش مصنوعی و مسئولیت قانونی، ترجمه سارا صلح چی و کیان بیگلریگی؛ تمدن حقوقی، شماره ۱۸، ویژه‌نامه هوش مصنوعی.

Copyrights:

Copyright for this article is retained by the author (s) , with publication rights granted to Legal Civilization. This is an open-access article distributed under the terms of the Creative Commons Attribution License (<http://creativecommons.org/licenses/by/4.0>) , which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.



مقدمه

یکی از شماره‌های اخیر یکی از مجلات مشهور رایانه‌ای [۱] پرسش ذیل را طرح کرد: «سال ۲۰۲۳ است و برای نخستین بار، یک اتومیل خودران که در خیابان‌های شهر در حال حرکت است، با یک عابرپیاده تصادف کرده و موجب مرگ وی می‌شود. قطعاً این مسئله باید از طریق طرح دعوی در دادگاه تعقیب شود. اما دقیقاً کدام قوانین اعمال می‌شود؟ هیچ‌کس نمی‌داند.» این مقاله در ادامه پیشنهاد می‌کند که قوانینی که احتمالاً اعمال می‌شوند، قوانینی هستند که با تولیدات با طراحی معیوب سروکار دارند. با وجود این، معتقد است که پیروی از مسیر قانونی، مانع توسعه و پیشرفت اتومیل‌های خودران می‌شود؛ زیرا فرایند حل و فصل اختلافات مربوط به پرونده‌های طراحی محصول (در ایالات متحده آمریکا) معمولاً ده برابر بیشتر از پرونده‌های مربوط به قصور انسانی هزینه دارند و هزینه اضافی مرتبط با محصول جهت برطرف کردن مسئله نیز داخل در این هزینه‌ها نیست. این مقاله این گونه ادامه می‌دهد که در عوض، باید با چنین پرونده‌هایی مانند پرونده‌های قصور و غفلت انسانی، همان‌طور که با انسان راننده برخورد می‌شود، رفتار گردد؛ نویسنده اشاره می‌کند که یک کتاب راهنما^۱ استاندارد از حقوق شبه جرم ایالات متحده آمریکا [۲] بیان می‌کند: «وضعیت ذهنی بد، نه لازم و نه کافی جهت نمایان شدن قصور و غفلت است، رفتار همه چیز است.»

ممکن است این مسئله حتی پیش از سال ۲۰۲۳ میلادی نیز رخ دهد. نویسنده این مقاله به تازگی اتومبیلی کرایه کرده است که امکانات ایمنی بسیاری دارد. یکی از این ویژگی‌ها این است که اگر رادارهای خودرو یک برخورد قریب‌الوقوع را درحالی‌که اتومبیل با سرعت بین چهار تا نوزده مایل بر ساعت حرکت می‌کند تشخیص دهد، موتور اتومبیل جهت جلوگیری از برخورد از کار می‌افتد. نویسنده در حال عقب آمدن جهت خروج از یک راه فرعی اختصاصی، بسیار نزدیک به یک پرچین حرکت کرد؛ اتومبیل هشدار صدای نزدیک شدن را به صدا در آورد و موتور از کار افتاد. با این حال، حتی زمانی که فرمان می‌چرخید تا خودرو از پرچین عبور کند و درحالی‌که خودرو در حالت دنده عقب بود، موتور روشن نمی‌شد. نویسنده مجبور شد که اتومبیل را در حالت دنده معمولی قرار دهد، مقدار کمی رو به جلو حرکت کند تا بتواند عقب رفتن خود را ادامه دهد. همه این اتفاقات در یک جاده فرعی اختصاصی رخ داد. با این حال، اگر زمانی رخ می‌داد که یک کامیون سنگین از عقب با سرعت یکسان در حال نزدیک شدن به اتومبیل می‌بود و نزدیک بود از پشت با خودرو برخورد کند و آن را از مسیر به بیرون پرتاب کند؛ بسیاری از رانندگان خطر خراشیده شدن رنگ خودرو با یک پرچین را به نشستن در اتومبیلی که از روشن شدن امتناع می‌کند و مانور موردنظر را کامل می‌کند؛ ترجیح می‌دهند. به زودی بدیهی به نظر می‌آید که برخی رانندگان، امکانات امنیتی را به جهت حضور و نقش داشتن در تصادف‌های شدید مؤثر بدانند.

هدف این پژوهش این است که ظرفیت‌های موجود یا کمبود وجود آن‌ها در هوش مصنوعی را بررسی کند و سپس دوباره به پرسش این که مسئولیت قانونی در پرونده‌های فوق چه جایگاهی دارد، بپردازد. نخست، مهم است که بیان شود این مقاله از اصطلاح «هوش مصنوعی» چه مقصودی دارد. تحقیقاتی در این حوزه وجود دارد که هوش مصنوعی را هر چیزی که با هر روشی از هوش انسانی تقلید می‌کند، هوش مصنوعی می‌دانند؛ افراد دیگری نیز وجود دارند که فکر می‌کنند تنها برنامه‌های هوش مصنوعی آن‌هایی هستند که طریقه تفکر انسان را تقلید می‌کنند. عده دیگری نیز در عرصه سیستم‌های اطلاعاتی، بسیاری از برنامه‌های هوش مصنوعی را به عنوان سیستم‌های پیچیده اطلاعاتی با هوش مصنوعی حقیقی که مختص تصمیم‌گیری در سطح پیشرفته که گاهی به عنوان عقل و خرد مشخص شده، دسته‌بندی می‌کنند. در این مقاله، هر سیستم رایانه‌ای که قادر به تشخیص یک موقعیت یا رویداد باشد و به شکل «اگر این موقعیت پیش آمد پس این اقدام را پیشنهاد کرده یا اخذ می‌کند»؛ به عنوان یک سیستم هوش مصنوعی در نظر گرفته می‌شود.

۱- مسئولیت قانونی

۱-۱- مسئولیت کیفری

ارجاعاتی که ذیل این بحث اشاره می‌شوند، اساساً به قوانین ایالات متحده آمریکا بازمی‌گردند؛ هرچند، بسیاری از نظام‌های قضایی دیگر نیز چنین قوانین مشابهی در حیطه‌های مربوطه دارند. در ارجاع [۳]، گابریل هالوی به بیان این که چطور، و این که آیا، نهادهای هوش مصنوعی ممکن است از لحاظ کیفری مسئول شناخته شوند، می‌پردازد. قوانین کیفری به طور عادی عنصر مادی جرم^۲ (یک رفتار) و عنصر معنوی^۳ (قصد ذهنی) را ضروری می‌دانند و هالوی به شکل یاری‌دهنده‌ای قوانین را به این شکل طبقه‌بندی می‌کند: آن‌هایی که در آن، عنصر مادی شامل یک فعل می‌شود و آن‌هایی که عنصر مادی در آن دربرگیرنده یک ترک فعل است؛

آن‌هایی که عنصر معنوی در آن نیازمند اطلاع یا آگاهی است؛ آن‌هایی که عنصر معنوی در آن صرفاً قصور و غفلت است (یک شخص متعارف باید آن‌ها را بداند)؛ و جرایمی که در آن‌ها مسئولیت محض است که لازم نیست هیچ عنصر معنوی در آن ثابت شود.

هالوی با طرح سه مدل قانونی که به وسیله آن‌ها جرایم ارتکاب یافته توسط سیستم‌های هوش مصنوعی شناخته می‌شوند؛ ادامه می‌دهد:

اول، ارتکاب از طریق شخصی دیگر.^۴ اگر یک جرم به وسیله یک شخص معلول ذهنی، کودک یا حیوان صورت پذیرد، مرتکب به جهت نبود اهلیت ذهنی برای تشکیل عنصر معنوی جرم، به عنوان یک فاعل بی‌گناه در نظر گرفته می‌شود (این امر حتی در مورد جرایم دارای مسئولیت محض نیز صادق است). هرچند اگر فاعل بی‌گناه توسط شخص دیگر آموزش داده شده باشد (برای نمونه اگر صاحب سگی، سگ خود را جهت حمله به یک شخص آموزش دهد) سپس شخص آموزش‌دهنده از لحاظ کیفری مسئول است.^۵ بر طبق این مدل، برنامه‌های هوش مصنوعی نیز می‌توانند تحت عنوان فاعل بی‌گناه محسوب شوند و برنامه‌نویس نرم‌افزار یا کاربر فاعل معنوی جرم در نظر گرفته شوند.

2- Actus Reus

3- Mens Rea

۴- فاعل معنوی جرم.

۵- به ارجاع شماره [۴] جهت رویه قضایی ایالات متحده آمریکا متوجه رجوع شود.

دوم، عواقب طبیعی و محتمل. در این مدل بخشی از برنامه هوش مصنوعی که جهت اهداف خوب در نظر گرفته شده بود، به طور نامطلوبی فعالیت کرده و اقدام مجرمانه‌ای را انجام می‌دهد. هالوی نمونه‌ای ارائه می‌کند (نقل شده از ارجاع شماره [۵]) که در آن یک کارمند ژاپنی کارخانه موتورسیکلت به وسیله یک ربات با هوش مصنوعی که نزدیک او کار می‌کرد، کشته شد. ربات به اشتباه کارمند را تهدید می‌کرد. انجام مسئولیتش تشخیص داده، و محاسبه می‌کند که مفیدترین راه جهت حذف این تهدید، هل دادن وی به سوی یک ماشین در حال کار در نزدیکی آن‌هاست. ربات با استفاده از بازوهای بسیار قوی هیدرولیک خود، کارمند غافلگیر شده را به ماشین کوبیده، بلافاصله او را می‌کشد و بعد به ادامه انجام وظایفش باز می‌گردد. استفاده قانونی از مسئولیت «نتیجه طبیعی و محتمل» جهت تعقیب شرکای جرم است. اگر هیچ نقشه یا تبانی قابل اثبات نباشد باز هم ممکن است (در حقوق ایالات متحده امریکا) که شریک جرم از لحاظ قانونی مسئول شناخته شود؛ با این فرض که اعمال مجرمانه مرتکب نتیجه محتمل یا طبیعی طرحی بوده (این عبارت ریشه در ارجاع شماره [۶] دارد) که شریک آن را یاری یا ترغیب کرده [۷]، به این شرط که شریک جرم از ابعاد مجرمانه طرح آگاه بوده باشد. بنابراین کاربران یا (با احتمال بیشتر) برنامه‌نویسان ممکن است دارای مسئولیت قانونی شناخته شوند اگر بدانند که جرمی کیفری، نتیجه طبیعی یا محتمل برنامه ایشان یا استفاده از یک نرم‌افزار است. به هر حال، قبل از اعمال این اصل باید میان برنامه‌های هوش مصنوعی که می‌دانند طرح مجرمانه در راه است (برای مثال برنامه‌ریزی شده‌اند تا یک طرح مجرمانه را اعمال کنند) و آن‌هایی که نمی‌دانند (برای هدف دیگری برنامه‌ریزی شده‌اند) تمایز قائل شد. جرایمی که در آن‌ها عنصر معنوی دانستن و آگاهی لازم است قابل تعقیب در دسته اخیر نیستند (اما آن‌هایی که لازمه وقوع آن‌ها، عنصر معنوی انسان معقول و یا مسئولیت محض (حادثه) است، قابل تعقیب‌اند).

سوم، مسئولیت مستقیم (بلاواسطه) (این مدل هم عنصر مادی و هم عنصر معنوی را به یک سیستم هوش مصنوعی انتساب می‌دهد). بحث انتساب عنصر مادی جرم به یک سیستم هوش مصنوعی به طور نسبی آسان است. اگر سیستمی اقدامی انجام دهد که منجر به فعل مجرمانه‌ای شود، یا در انجام اقدامی که وظیفه اجرای آن را داشته شکست بخورد، در اینجا عنصر مادی یک جرم حاصل شده است. تعیین عنصر معنوی به مراتب دشوارتر است و به همین جهت است که در اینجا سه سطح از عنصر معنوی مهم می‌شوند. برای جرایم با مسئولیت محض، در جایی که هیچ قصدی برای انجام یک جرم لازم نیست، درواقع ممکن است بتوان برنامه‌های هوش مصنوعی را دارای مسئولیت کیفری به حساب آورد. با توجه به مثال اتومبیل خودران، سرعت غیرمجاز جرمی با مسئولیت محض است؛ بنابراین مطابق با نظر هالوی، اگر کاشف به

عمل آمد که اتومبیلی خودران از محدودیت سرعت برای جاده‌ای که در آن است، گذر کرده؛ قانون ممکن است مسئولیت کیفری را به برنامه هوش مصنوعی که در آن زمان در حال راندن اتومبیل بوده نسبت دهد. این امکان، مسائلی دیگری را ایجاد می‌کند که هالوی آن‌ها را به اختصار ذکر کرده است. از جمله، دفاعیات (آیا برنامه‌ای که دچار نقص در عملکرد بوده دفاعیه را شبیه به دفاعیات انسانی مبتنی بر جنون طرح کند؟ یا اگر به وسیله یک ویروس الکترونیکی تحت تأثیر قرار گرفته باشد می‌تواند دفاعیه شبیه به تهدید و اجبار یا مستی طرح کند؟)؛ و مجازات (چه کسی یا چه چیزی برای یک جرمی که سیستم هوش مصنوعی مستقیماً مسئول آن بوده مجازات می‌شود؟)

۱-۲- دفاع بدافزار تروجان^۶

در زمینه دفاعیاتی که در مقابل مسئولیت سیستم‌های هوش مصنوعی صورت می‌پذیرد، ذکر بسیاری از پرونده‌ها که در آن خواننده متهم به جرایم سایبری شده بود و موفق شده دفاعی ارائه دهد مبتنی بر این امر

۶- تروجان اساساً یک برنامه خرابکارانه است که خود را بی‌ضرر جلوه می‌دهد تا مردم را در مورد دانلود کردن خود فریب دهد. تروجان‌ها یکی از نخستین انواع بدافزارهایی هستند که پدید آمده‌اند. دفاع بر این اساس برداشتی مبتنی بر فناوری از دفاع کلاسیک SODDI است که گمان می‌رود در سال ۲۰۰۳ میلادی در بریتانیا ظاهر شد. یا (ii) ارتکاب یک جنایت سایبری از طریق رایانه متهم، بر این اساس که یک بدافزار (مانند اسب تروجان، ویروس، کرم، ربات اینترنتی یا سایر برنامه‌ها) یا مرتکب دیگری که از چنین بدافزاری استفاده می‌کند، مسئول این ارتکاب بوده است. در اینجا یک شخصی که متهم به یک جرم غیرسایبری اعتراف کند که درحالی‌که از نظر فنی ممکن است متهم مسئول ارتکاب جرم باشد، اما او به دلیل دخالت بدافزار فاقد قصد مجرمانه یا دانش لازم است. این عبارت به خودی خود یک اصطلاح حقوقی ثابت نیست و از متون اولیه متخصصان شواهد دیجیتالی که به طور خاص به تروجان‌ها اشاره می‌کنند نشأت می‌گیرد، با توجه به استفاده روزافزون از برنامه‌های تروجان توسط هکرها و افزایش تبلیغات در مورد دفاع، استفاده از آن احتمالاً گسترده‌تر می‌شود.

See: 1- Bowles, S., Hernandez-Castro, J., "The first 10 years of the Trojan Horse defence", Computer Fraud & Security, January 2015, Vol. 2015 (1), pp. 5-13, page 5. 2- Steel, C. M. S., "Technical SODDI Defences: the Trojan Horse Defence Revisited", DFSL V9N4 (<http://ojs.jdfsl.org/index.php/jdfsl/article/viewFile/258/236>). 3- Šepec, M., "The Trojan Horse Defence -- a Modern Problem of Digital Evidence", Digital Evidence and Electronic Signature Law Review, 9, (2012), p. 1. 4- Brenner, S., Carrier, B., Henninger, J., 'The Trojan Horse Defense in Cybercrime Cases' (2004) 21 Santa Clara Computer and High Technology Law Journal 1, page 18. See the case of Eugene Pitts (2003). 5- Šepec, M., "The Trojan Horse Defence -- a Modern Problem of Digital Evidence", Digital Evidence and Electronic Signature Law Review, 9, (2012), page 2. 6- Brenner, S., Carrier, B., Henninger, J., 'The Trojan Horse Defense in Cybercrime Cases' (2004) 21 Santa Clara Computer and High Technology Law Journal 1, p. 11. 7- Kao, D.Y., Wang, S.J., Huang, F., 'SoTE: Strategy of Triple-E on solving Trojan defense in Cyber-crime cases' (2010) Computer Law and Security Review 26, p. 55.

که یک تروجان یا برنامهٔ بدافزار شبیه به آن، رایانهٔ او را در اختیار گرفته بوده و از طریق رایانهٔ خواننده به ارتکاب جرم می‌پرداخته بی‌آن که خواننده از آن خبر داشته باشد، مهم است. بررسی چنین مواردی نشان می‌دهد که [۸]؛ این موارد شامل پرونده‌ای در انگلستان می‌شود که در آن رایانه‌ای با تصاویر ناپسند از کودکان، درگیر یازده برنامهٔ تروجان شده بود و پروندهٔ انگلیسی دیگری، جایی که دفاعیات یک نوجوان هکر رایانه در برابر اتهام حملهٔ سایبری و از دسترس خارج کردن خدمات این بود که حمله از طریق رایانهٔ خواننده به وسیلهٔ یک برنامهٔ تروجان انجام شده که متعاقباً رایانه خود را پیش از آن که مورد ارزیابی دادگاهی قرار گیرد، پاک کرده بود. وکیل خواننده موفق شد تا هیات منصفه را قانع کند که چنین سناریویی ورای تردید منطقی قرار نمی‌گیرد.

۱-۳- حقوق مدنی: شبه جرم و نقض گارانتی

۱-۱-۳- قصور

زمانی که یک نرم‌افزار معیوب است، یا زمانی که یک طرف در نتیجهٔ استفاده از نرم‌افزار آسیب می‌بیند، رویه‌های قضایی ناشی از آن، معمولاً به جای مسئولیت کیفری، اقامهٔ دعوی شبه جرم قصور می‌کنند. [۹] گستر [۱۰] سه عنصر را که به طور معمول باید ایجاد شوند تا ادعای قصور ثابت شود، ذکر می‌کند: اول، خواننده وظیفهٔ مراقبت داشته؛ دوم، خواننده آن وظیفه را نقض کرده؛ سوم، آن نقض، سبب ورود آسیب به خواهان شده است.

با توجه به نکتهٔ اول، گستر بیان می‌کند که تردید اندکی پیرامون این که آیا فروشندهٔ نرم‌افزار وظیفهٔ مراقبت و حمایت در برابر خریدار دارد وجود دارد؛ اما تصمیم در این که چه استانداردی از حمایت باید رعایت شود، دشوار است. اگر سیستم مدنظر، «سیستم متخصص» شناخته شود، گستر پیشنهاد می‌کند که استاندارد متناسب، حمایتی از سوی یک متخصص یا دست کم حرفه‌ای است. در مورد نکتهٔ دوم، گستر راه‌های زیادی را بیان می‌کند که ممکن است یک سیستم هوش مصنوعی وظیفهٔ مراقبت را نقض کرده باشد که شامل: خطاهای عملکرد برنامه که می‌توانستند توسط توسعه‌دهنده تشخیص داده شوند؛ پایگاه دانشی نادرست یا ناکافی، مستندات یا هشدارهای نادرست یا ناکافی؛ عدم بروزرسانی مدام دانش، ورودی‌های معیوب تأمین‌کنندهٔ کاربری، کاربری با تکیهٔ بیش از حد بر خروجی؛ یا استفاده از برنامه جهت نیل به هدفی نادرست می‌باشد. در مورد نکتهٔ سوم، پرسش این که آیا یک سیستم هوش مصنوعی می‌تواند

طوری تلقی شود که باعث ایجاد آسیب شود، هنوز مورد تردید است. پرسش کلیدی این است که شاید سیستم هوش مصنوعی یک اقدام را در یک موقعیت به خصوص پیشنهاد کند (همان‌گونه که بسیاری از نظام‌های تخصصی این کار را انجام می‌دهند)، یا اقدامی انجام دهد (مانند خودران بودن و تجهیزات ایمنی خودکار اتومبیل). در پرونده اول، می‌بایست حداقل یک فاعل دیگر نیز حضور داشته باشد، بنابراین اثبات علیت دشوارتر از پرونده اخیر است که به مراتب آسان‌تر است.

گستره همچنین به بیان یک استثناء مطابق با حقوق ایالات متحده آمریکا در مورد «مسئولیت محض قصور» می‌پردازد. این امر بر محصولاتی اعمال می‌شود که مشخصاً یا به‌طور غیرمعقولی خطرناک هستند، هنگامی که با رفتاری عادی، با اراده، آگاهی و از نظر عقلی قابل پیش‌بینی، از آن‌ها استفاده می‌شود و باعث ایجاد آسیب می‌شوند (برخلاف ضرر اقتصادی و مالی). وی بحث می‌کند که آیا نرم‌افزار حقیقتاً یک «محصول تولیدی» است یا به‌طور محض یک «خدمت» است؛ وی به پرونده‌ای اشاره می‌کند که برقر در آن به‌عنوان یک محصول در نظر گرفته شده بود [11]، و بنابراین به سمت تعریف نرم‌افزار به‌عنوان یک محصول به جای یک خدمت تمایل دارد.

با این فرض که نرم‌افزار در واقع یک محصول است، برعهده توسعه‌دهندگان سیستم‌های هوش مصنوعی است که اطمینان حاصل کنند که سیستم‌های شان عاری از هرگونه عیوب طراحی، عیوب در تولید و یا هشدارها و دستورالعمل‌های ناکافی است. کول [12] بحث مفصل‌تری راجع به این سوال که آیا نرم‌افزار یک محصول است یا یک خدمت ارائه می‌دهد. نتیجه او این است که به حساب آوردن تمام سیستم‌های هوش مصنوعی به‌عنوان خدمت، در بهترین حالت تا حدی قابل اجراست و ترجیح می‌دهد که هوش مصنوعی را به جای محصول، یک خدمت در نظر بگیرد، اما بیان می‌کند که قانون در این حوزه بد تعریف شده است. کول به برخی پرونده‌های قضایی در زمینه «وظیفه مراقبت» اشاره می‌کند که سیستم‌های هوش مصنوعی باید آن را رعایت کنند:

(ارجاع شماره [۱۴])، در یک مدرسه دعوای قصور را در برابر یک دفتر آماری که (گویا) با محاسبات اشتباه از ارزش مدرسه که دچار سوختگی شده بود، منجر به این امر شد که مدرسه از ضرر غیرمشمول شدن بیمه آسیب ببیند، اقامه دعوا نمود. وظیفه‌ای که اینجا شناسایی شد، وظیفه تأمین اطلاعات با مراقبت معقول بود. دادگاه عواملی را ازجمله: وجود، در صورت وجود، ضمانت صحت را در نظر

گرفت. آگاهی متهم مبنی بر این که شاکی به اطلاعات اتکاء می‌کند. محدودیت مسئولیت بالقوه برای یک گروه کوچک؛ عدم وجود مدرکی مبنی بر هرگونه تصحیح پس از کشف؛ نامطلوب بودن الزام یک طرف بی‌گناه برای تحمل بار اشتباهات حرفه‌ای دیگری؛ و ترویج تکنیک‌های احتیاطی در میان ارائه‌دهندگان اطلاعاتی (ابزار).

براساس (ارجاع شماره [۱۵])، کول وظیفه تأمین نتایج منطقی از ورودی‌های غیرمعقول را بررسی می‌کند. او از (ارجاع شماره [۱۶]) جهت بیان این که توسعه‌دهندگان هوش مصنوعی احتمالاً وظیفه مثبت جهت فراهم کردن تکنیک‌های بررسی خطای ورودی ارزان، بی‌خطر و ساده را برعهده دارند، اما خاطرنشان می‌کند که این قواعد ممکن است در جایی که برنامه هوش مصنوعی در حال انجام یک عمل است که خطا در ورودی آن، ممکن است مستقیماً جان مردم را به خطر بیندازد (برای نمونه تجویز دارو برای یک بیمار) اعمال نشود؛ در چنین پرونده‌هایی او به طور جایگزین اعمال قاعده «فعالیت‌های بسیار خطرناک و دستورالعمل‌ها» را توصیه می‌کند [۱۷].

کول بیان می‌کند که سیستم‌های هوش مصنوعی باید از موانع و محدودیت‌های شان آگاه باشند و این اطلاعات باید به خریدار اطلاع داده شود. این امر به خوبی تثبیت شده است که فروشنده وظیفه دارد تا به خریدار هر عیبی که می‌داند را اطلاع دهد؛ اما چطور عیوب یا ضعف‌های ناشناس ممکن است مشخص شوند و سپس موارد موجود را به آن‌ها اطلاع‌رسانی کرد؟

۲-۱-۳- نقض گارانتی

اگر یک سیستم هوش مصنوعی واقعاً یک محصول است، باید با گارانتی فروخته شود. حتی اگر هیچ ضمانت‌نامه صریحی توسط فروشنده (یا خریداری شده توسط کاربر) ارائه نشده باشد، یک ضمانت ضمنی وجود دارد که (برای استفاده از عبارت قانون فروش کالاهای بریتانیا در سال ۱۹۷۹ میلادی)، «رضایت بخش است و برای یک زمان معقول.» برخی از حوزه‌های قضایی اجازه می‌دهند که ضمانت‌های ضمنی توسط بندهایی در قرارداد باطل شوند. با این حال، زمانی که یک سیستم هوش مصنوعی داخلی در سایر کالاها (مانند خودرو) خریداری می‌شود، بعید به نظر می‌رسد که چنین استثنائات قراردادی (مثلاً بین سازنده خودرو و تأمین‌کننده نرم‌افزار هوش مصنوعی) با موفقیت به خریدار آن منتقل شود.

۱-۴- مسئولیت حقوقی: خلاصه

به نظر می‌رسد، پرسشی که آیا سیستم‌های هوش مصنوعی می‌توانند از لحاظ حقوقی مسئول شناخته شوند؛ دست کم به سه عامل بستگی دارد: اول، محدودیت‌ها و موانع سیستم‌های هوش مصنوعی و این که آیا این محدودیت‌ها شناخته شده و به خریدار اطلاع داده شده‌اند یا خیر؛ دوم، این که آیا سیستم هوش مصنوعی محصول است یا خدمت؛ سوم، آیا جرم به عنصر معنوی نیاز دارد یا جرمی با مسئولیت محض است. اگر یک سیستم هوش مصنوعی مسئول شناخته شود، پرسشی که پیش می‌آید این است که باید تحت کدام عنوان، فاعلی بی‌گناه، شریک جرم یا مرتکب شناسایی گردد. بخش پایانی این مقاله به مورد اول از این سه عنصر می‌پردازد.

۲- محدودیت‌های سیستم‌های هوش مصنوعی

محدودیت‌های مختلفی که نظام‌های هوش مصنوعی مشمول آن هستند را می‌توان به دو دسته تقسیم کرد: محدودیت‌هایی که انسان‌های متخصص با دانش مشابه نیز مشمول آن هستند؛ محدودیت‌های فناوری هوش مصنوعی در مقایسه با انسان‌ها.

۲-۱- محدودیت‌هایی که بر سیستم‌های هوش مصنوعی و متخصصان انسانی تأثیر می‌گذارد

محدودیت‌هایی که هم سیستم‌های هوش مصنوعی و هم انسان‌های متخصص را متأثر می‌کند؛ به دانشی که مختص مسئله و مشکل است؛ مرتبط هستند. اول، دانش ممکن است به سرعت تغییر یابد. این امر موجب می‌شود که هم انسان‌ها و هم سیستم‌های هوش مصنوعی، هر دو، بدانند که به‌روزترین دانش چیست و همچنین تشخیص دهند کدام بخش از دانش قبلی به روز نیست. این که این یک مسئله است، تقریباً به شکل کامل به حوزه و حیطه بستگی دارد: در مثال ما از یک اتومبیل با رانندگی خودکار، علم و دانشی که جهت رانندگی یک اتومبیل لازم است، حقیقتاً به آرامی تغییر می‌کند. با این حال، در جهان امنیت سایبری، دانش مربوط به سوءاستفاده‌ها و رفع و رجوع آن‌ها به شکل روزانه تغییر می‌کند. دوم، دانش ممکن است جهت شناخت تمام احتمالات بسیارگسترده باشد. سیستم‌های هوش مصنوعی در حقیقت می‌توانند در انجام اموری مانند (امکان جست‌وجوی هزاران یا حتی صدها هزار راه‌حل) بهتر از انسان‌های متخصص عمل کنند، اما هنوز برخی از امور وجود دارند که مقیاس آن از چیزی که گفته شد

گسترده‌تر باشد. این موضوع معمولاً در کارهای برنامه‌ریزی و طراحی صدق می‌کند، جایی که تعداد برنامه‌ها یا طرح‌هایی ممکن است نزدیک به بی‌نهایت باشد. (در مقابل، جدول‌بندی و شکل‌دهی، که نیازمند برنامه‌ریزی و طراحی درون یک چهارچوب معین است، کمتر پیچیده هستند، هرچند که ممکن است گزینه‌های محتمل به هزاران گزینه نیز برسد.) در چنین مواردی، سیستم‌های هوش مصنوعی می‌توانند تعهد بدهند که در اغلب قضایا پاسخ‌های خوبی ارائه می‌کنند، اما نمی‌توانند ضمانت دهند که در تمام قضایا بهترین راهکار را ارائه می‌کنند.

از منظر حقوقی، می‌تواند این طور استدلال شود که راه‌حل برای چنین مشکلاتی به عهده فروشنده است تا به خریدار یک سیستم هوش مصنوعی درباره وجود این محدودیت‌ها هشدار دهد. در عرصه‌هایی که سریعاً تغییر می‌کنند، اگر فروشنده روشی برای به‌روزرسانی مکرر دانش سیستم ارائه ندهد، ممکن است از نظر قانونی غیرمنطقی تلقی شود. این امر سبب طرح این پرسش می‌شود که زمینه مرزهای «تغییر-سریع» کجاست. مانند همیشه، آزمون حقوقی راجع به قابلیت معقول بودن است، که معمولاً در برابر عمر مورد انتظار یک سیستم هوش مصنوعی قیاس می‌گردد؛ بنابراین اگر پیش‌بینی شده بود که دانش آن به شکل سالانه تغییر داده شود (برای مثال سیستم هوش مصنوعی که مشغول محاسبه مسئولیت مالیات شخصی است)، پس احتمالاً برای فروشنده، هشدار پیرامون این که دانش استفاده شده مورد تغییر قرار می‌گیرد؛ معقول به نظر می‌رسد. هرچند ممکن است برای فروشنده‌ای که به‌روزرسانی‌های خودکار دانش را تأمین می‌کند عاقلانه نباشد، زیرا پیچیدگی قانون مالیات‌طوری است که هر به‌روزرسانی صرفاً نیازمند دانلود فایل‌های داده و اطلاعات نیست؛ بلکه ممکن است نیازمند یک سیستم تازه تأسیس و امتحان شده باشد.

در مقابل، سیستم‌های هوش مصنوعی که به شوراها محلی کمک می‌کنند تا مزایای خانوار را محاسبه کنند، ممکن است براساس این فرض (به ظاهر غیرقابل تزلزل) ساخته شده باشند که ازدواج بین یک مرد و یک زن بوده است. با این حال، این دانش اکنون تغییر کرده است تا ازدواج بین هر دو انسان بالغ مجاز باشد. آیا منطقی است که از فروشنده بخواهیم به خریدار هشدار دهد که این قوانین نیز ممکن است تغییر کنند؟ چنین تغییری در حال حاضر بسیار بعید به نظر می‌رسد. اما در ایالات متحده آمریکا، تلاش‌هایی توسط یک زن برای ازدواج با سگ خود و توسط یک مرد برای ازدواج با رایانه شخصی‌اش

صورت گرفته است و همچنین تلاش‌های طولانی مدتی از سوی برخی گروه‌های مذهبی جهت قانونی کردن چندهمسری انجام شده است. در پرونده آمریکایی Kociemba v Searle [۱۸]، یک تولیدکننده دارو به جهت عدم هشدار به خریداران پیرامون آن که استفاده از یک داروی خاص که در ارتباط با بیماری‌های التهابی لگن بود؛ مسئول شناخته شد، با وجود این که محصول به وسیله سازمان غذا و دارو «امن و مؤثر» شناخته شده بود. بنابراین به نظر می‌رسد، در محدوده‌ای که ممکن است به طور منطقی به هشدار نیاز باشد، درواقع به دانش نیاز است تا تأیید نظارتی.

مایکایتین و دیگران [۱۹] راجع به مسائل مربوط به مسئولیت حقوقی یک سیستم هوش مصنوعی که به شناسایی و گزینش متخصصان پیوند خورده، بحث می‌کنند. آن‌ها دو پرونده‌ای را نقل قول می‌کنند (ارجاعات [۲۰] و [۲۱]) که در آن بیمارستان‌ها به جهت شکست در گزینش پزشکان با صلاحیت کافی و مناسب جهت تأمین خدمات درمانی که باید فراهم می‌کردند، مسئول شناخته شدند؛ با این قیاس، توسعه‌دهندگان هوش مصنوعی نیز می‌توانند مسئول شناخته شوند مگر این که متخصصانی با صلاحیت کافی را در حوزه‌های مدنظر انتخاب کنند یا این که به خریداران هشدار دهند که صلاحیت متخصصان قابل تسری به سایر حوزه‌هایی که ممکن است از این سیستم استفاده شود، نیست. راه‌حل مطرح شده به وسیله مایکایتین و دیگران، استفاده از متخصصان دارای مجوز و گواهینامه است. آن‌ها بر این امر اشاره دارند که استانداردهایی که نهادهای صدور گواهی نیاز دارند، برخی مواقع جهت تعیین عملکرد یک متخصص در حد سطح موردنظر، استفاده می‌شوند [۲۲]. آن‌ها حتی پیشنهاد می‌دهند که ممکن است اخذ گواهی برای خود سیستم هوش مصنوعی هم مطلوب باشد. کمیسیون بورس و امنیت ایالات متحده آمریکا به ویژه مشتاق به شیوه مذکور بوده است؛ این کمیسیون الزام می‌دارد که سیستم پیشنهاد دهنده بازار سهام به عنوان مشاور اقتصادی ثبت گردد [۲۳] و توسعه‌دهندگان برنامه‌های توصیه سرمایه‌گذاری نیز به عنوان مشاوران سرمایه‌گذاری دسته‌بندی شوند [۲۴].

۲-۲- محدودیت‌های سیستم‌های هوش مصنوعی که متخصصان انسانی را تحت تأثیر قرار نمی‌دهد

محدودیت اصلی این است که سیستم‌های هوش مصنوعی دچار کمبود دانش عمومی هستند. انسان‌ها همواره دارای مقادیر عظیمی از دانش هستند که ممکن است به صورت مستقیم مرتبط با یک وظیفه

خاص نباشد، اما امکان ارتباط را نیز دارد. برای مثال، زمان رانندگی یک اتومبیل، توصیه می‌شود که هنگام عبور از یک مدرسه، خصوصاً اگر در خارج از مدرسه ردیفی از ماشین‌های پارک شده وجود داشته باشد، یا شما می‌دانید که زمان کاری مدرسه موقعی است که شما در حال گذر از آنجا هستید پایان می‌پذیرد، به آرامی رانندگی کنید. دلیل آن جلوگیری از خطر پریدن کودکان از پشت ماشین‌های پارک شده است، زیرا دانش عمومی یک انسان راننده شامل این واقعیت می‌شود که بعضی از کودکان مهارت‌های ضعیفی در ایمنی راه‌ها دارند. یک اتومبیل خودکار این امر را نمی‌داند، مگر این که با یک قاعده خاص یا دسته‌ای از قواعد عمومی راجع به اماکن غیرمعمول خطرناک برنامه‌نویسی شده باشد.

البته باید اعتراف کرد، مواقعی هست که انسان‌ها در استفاده از دانش عمومی خود جهت تشخیص شرایط خطرناک شکست می‌خورند: همان‌طور که یکی از مفسرین یک بار بیان کرد، «تفاوت میان یک جاروبرقی تسمه‌ای^۷ و یک ژنراتور وان دی گراف^۸ چیست؟ بسیار جزئی است، اما هرگز رایانه شخصی خود را با چنین جاروبرقی تمیز نکنید.» هرچند، بدون دانش عمومی، سیستم‌های هوش مصنوعی شانس برای تشخیص چنین شرایطی ندارند. یک مسئله مرتبط این است که سیستم‌های هوش مصنوعی به شکل مشهوری در عملکرد سطوح پایین و با ظرافت ضعیف هستند. این امر زمانی مشاهده می‌شود که موارد لبه‌ای^۹ (مواردی که یک متغیر در آن ارزش زیادی دارد) یا موارد گوشه‌ای^{۱۰} (مواردی که متغیرهای متعددی در آن ارزش زیادی دارند) در دستور کار قرار می‌گیرد. زمانی که انسان‌ها با شرایطی مواجه می‌شوند که پیشتر معتقد بودند غیرممکن یا با احتمال بسیار اندک است، معمولاً می‌توانند اقداماتی را برگزینند که اثر مثبت بر روی شرایط دارد. زمانی که سیستم‌های هوش مصنوعی با چنین موقعیتی روبرو می‌شوند که برای آن برنامه‌ریزی نشده‌اند، به طور کلی اصلاً نمی‌توانند فعالیت کنند.

به عنوان نمونه، در مثال رانندگی اتومبیل که در ابتدای مقاله ارائه شد، موقعیت فرضی که اتومبیل از روشن شدن امتناع می‌کند درحالی که یک کامیون سنگین درحال زیر گرفتن آن است، یک مورد لبه‌ای

7- Belt-driven Vacuum Cleaner

8- Van de Graff Generator

9- Edge Cases

10- Corner Cases

است. به علاوه، به نظر نمی‌رسد که سیستم ایمنی اتومبیل با توجه به ذهنیت رانندگان شهری طراحی شده باشد؛ اتومبیل به رانندگان هشدار می‌دهد تا پیش از حرکت از ایمنی جاده مطمئن شوند، اما این واقعیت را در نظر نمی‌گیرد که در یک شهر، یک راه ممکن است برای مقطع زمانی کوتاهی ایمن باشد، بنابراین، این نوع موارد «لبه‌ای» بیش از حد انتظار رایج است. در مورد پرونده‌های گوشه‌ای، بیست و ششم سپتامبر ۱۹۸۳ میلادی روزی بود که ماهواره هشدار زودهنگام اتحاد جماهیر شوروی ابتدا یکی، سپس دو و در نهایت پرتاب پنج موشک هسته‌ای ایالات متحده آمریکا را تشخیص داد. سیاست استاندارد اتحاد جماهیر شوروی در آن زمان، اقدام تلافی‌جویانه به وسیله موشک‌های خودش بود؛ و این اتفاق در مقطعی از زمان رخ داد که تنش‌های سیاسی شدیدی میان ایالات متحده آمریکا و اتحاد جماهیر شوروی وجود داشت. افسر مسئول تنها چند دقیقه برای تصمیم راجع به این که چه اقدامی انجام دهد فرصت داشت و اطلاعات بیشتری نیز در دست نبود؛ او تصمیم گرفت تا آن پیام را به‌عنوان هشدار نادرست در نظر بگیرد با این استدلال که، «هنگامی که مردم جنگی را آغاز می‌کنند، آن را با شلیک تنها پنج موشک آغاز نمی‌کنند». بعدها معلوم شد که ماهواره، انعکاس نور خورشید از ابرها را با ویژگی گرمایی پرتاب موشک اشتباه گرفته است. گردش ماهواره برای جلوگیری از چنین خطاهایی طراحی شده بود، اما در آن روز (نزدیک به نقطه اعتدال شبانه روز) موقعیت ماهواره، وضعیت خورشید و موقعیت زمین‌های موشکی ایالات متحده آمریکا همگی باهم ترکیب شدند تا موجب پنج خوانش اشتباه شوند. اگر یک سیستم هوش مصنوعی آن روز در اتحاد جماهیر شوروی کنترل پرتاب موشک را عهده داشت، ممکن بود در تشخیص ماهواره دچار دشواری و اشکال شود و آن را شناسایی نکند؛ لذا موشک‌ها را پرتاب کند. بنابراین هوش مصنوعی به طور حقوقی مسئول خرابی‌های به بار آمده می‌شد. هرچند مبهم است که آیا وکلایی پیگیری پرونده را قبول می‌کردند یا خیر.

مسئله سوم این است که سیستم‌های هوش مصنوعی ممکن است به جهت ورودی‌های با کیفیت پایین‌تر، فاقد اطلاعاتی باشند که انسان‌ها استفاده می‌کنند. این دقیقاً پرونده اتومبیل با سیستم ایمنی است؛ روش‌های ورودی آن تنها ردیاب‌های رادار نسبتاً با برد کوتاه است؛ که نمی‌توانند میان یک پرچین یا کامیون سنگین تمایزی قائل شود یا حتی شیئی را که تقریباً دور است اما به‌سرعت نزدیک می‌شود تشخیص دهند، بنابراین ممکن است اگر پرونده‌ای در رابطه با تصادف «ناشی از این سیستم‌های ایمنی» به

دادگاه بیاید، تمرکز بر این امر باشد که سیستم‌های هوش مصنوعی چقدر برای مقابله با این ورودی‌های غیردقیق برنامه‌ریزی شده‌اند. موضوع اطلاعات غیرنمادین نیز وجود دارد. در دنیای مدیریت دانش، خواندن این ادعاها رایج است که دانش بشری هرگز نمی‌تواند به طور کامل در سیستم‌های کامپیوتری گنجانده شود، زیرا بسیار درون یافتی است [۲۵]. کینگستون [۲۶] بحث می‌کند که چنین دیدگاهی به طور کلی نادرست است، زیرا بر پایه فهم ضعیف از انواع متنوع دانش ضمنی طرح شده است؛ اما او می‌پذیرد که اطلاعات غیرنمادین (اطلاعات بر پایه اعداد، اشکال، درک چیزهایی مانند بافت‌ها، یا اطلاعات فیزیولوژیکی مثل حرکت ماهیچه‌های یک رقصنده باله) و مهارت‌ها یا دانشی که از این اطلاعات به دست می‌آید، فراتر از مقیاس هر سیستم هوش مصنوعی است.

در برخی زمینه‌ها، این اطلاعات غیرنمادین حیاتی هستند؛ پزشکان حین معاینه بیماران، برای نمونه، همزمان با این که اطلاعات بسیاری از گفتار بیمار کسب می‌کنند اطلاعاتی را نیز از زبان بدن او در می‌یابند. برخی از انتقاداتی که به خدمات تشخیصی تلفنی^{۱۱} فعلی انگلستان وارد شده است، می‌تواند به پزشکان متخصصان فاقد این نوع اطلاعات بازگردد. در مثال رانندگی اتومبیل، اطلاعات غیرنمادین ممکن است شامل نوربالا توسط رانندگان دیگر برای ارسال پیام از یک خودرو به خودروی دیگر باشد؛ چنین اطلاعاتی حیاتی نیستند اما مهم است که راننده‌ای باشیم که به دیگران احترام می‌گذارد.

نتیجه

این‌طور بیان شد که مسئولیت قانونی سیستم‌های هوش مصنوعی حداقل به سه عامل بستگی دارد: اول، این که آیا هوش مصنوعی خدمت یا محصول است. این امر در قانون بد تعریف شده است و مفسرین متفاوت، دیدگاه‌های مختلفی را در این مورد عرضه می‌کنند. دوم، اگر جرم کیفری رخ داد، چه عنصر معنوی‌ای مورد نیاز است. بعید به نظر می‌رسد که برنامه‌های هوش مصنوعی باعث نقض قوانینی شوند که نیازمند آگاهی و علم به ارتکاب فعل مجرمانه‌اند؛ اما بسیار ممکن است که قوانینی را نقض کنند که «یک انسان متعارف می‌دانست» که این دست از اقدامات ممکن است باعث جرم شوند و تقریباً قطعی است که آن‌ها می‌توانند جرایم با مسئولیت محض مرتکب شوند. سوم، این که محدودیت‌های

سیستم‌های هوش مصنوعی باید به خریداران اعلام شود. از آنجایی که سیستم‌های هوش مصنوعی هر دو محدودیت‌های عمومی و خاص را دارند، پرونده‌های حقوقی پیرامون چنین موضوعاتی ممکن است بر پایه گفتار خاص هر هشدار دربارۀ این محدودیت‌ها باشند.

همچنین سوالی مطرح شد که چه کسی باید مسئول شناخته شود؟ این موضوع بستگی به این دارد که کدام یک از مدل‌های سه‌گانه هالوی اعمال می‌شود (ارتکاب از طریق شخصی دیگر یا فاعل معنوی جرم؛ نتیجه طبیعی یا محتمل؛ و یا مسئولیت بلافصل و مستقیم): در جرایم ارتكابی (از طریق) شخصی دیگر یا فاعل معنوی جرم، فرد تعلیم دهنده سیستم هوش مصنوعی (چه کاربر و چه برنامه‌نویس)، احتمالاً مسئول تلقی شود. در جرایم در نتیجه طبیعی یا محتمل، مسئولیت ممکن است برعهده هر شخصی باشد که باید پیش‌بینی می‌کرده که محصول ممکن است بدین طریق مصرف شود؛ برنامه‌نویس، فروشنده (یک محصول) یا تأمین‌کننده خدمات. کاربر در اینجا کمتر مورد سرزنش قرار می‌گیرد مگر آن که دستورالعملی که همراه با محصول/خدمت بیان می‌شود؛ محدودیت‌های سیستم را بیان کند و نتایج ممکن از استفاده ناصحیح را با جزئیات طرح کرده باشد. سیستم‌های هوش مصنوعی همچنین ممکن است در جرایم مسئولیت محض نیز دارای مسئولیت شناخته شوند که در آن برنامه‌نویس احتمالاً مقصر دانسته می‌شود. هرچند در تمام موارد برنامه‌نویس مسئول شناخته می‌شود، ممکن است تردیدهای آتی نیز وجود داشته باشد که تقصیر با برنامه‌نویس، طراح برنامه، متخصصی که اطلاعات را فراهم کرده، یا مدیری که متخصص یا برنامه‌نویس یا طراح نالایق منصوب کرده است؛ باشد.

ملاحظات اخلاقی: موارد مربوط به اخلاق در پژوهش و نیز امانتداری در استناد به متون و ارجاعات مقاله تماماً رعایت گردیده است.

تعارض منافع: تعارض منافع در این مقاله وجود ندارد.

تأمین اعتبار پژوهش: این پژوهش بدون تأمین اعتبار مالی نگارش یافته است.

منابع

- Greenblatt N. A.: Self-Driving Cars and the Law. IEEE Spectrum, p. 42 (16 February 2016).
- Dobbs D. B.: Law of Torts. West Academic Publishing (2008).
- Hallevy G.: The Criminal Liability of Artificial Intelligence entities. <http://ssrn.com/abstract=1564096> (15 February 2010).
- Morrissey v. State, 620 A. 2d 207 (Del. 1993) ; Conyers v. State, 367 Md. 571, 790 A. 2d 15 (2002) ; State v. Fuller, 346 S. C. 477, 552 S. E. 2d 282 (2001) ; Gallimore v. Commonwealth, 246 Va. 441, 436 S. E. 2d 421 (1993).
- Weng Y-H, Chen C-H and Sun C-T: Towards the Human-Robot Co-Existence Society: On Safety Intelligence for Next Generation Robots, 1 Int. J. Soc. Robot. 267, 273 (2009).
- United States v. Powell, 929 F. 2d 724 (D. C. Cir. 1991).
- Francis Bowes Sayre: Criminal Responsibility for the Acts of Another, 43 Harv. L. Rev. 689 (1930).
- Brenner S. W. , Carrier B. , Henninger J.: The Trojan Horse Defense in Cybercrime Cases, 21 Santa Clara High Tech. L. J. 1 <http://digitalcommons.law.scu.edu/chtj/vol21/iss1/1> (2004).
- Tuthill G. S.: Legal Liabilities and Expert Systems, AI Expert (Mar. 1991).
- Gerstner M. E.: Comment, Liability Issues with Artificial Intelligence Software, 33 Santa Clara L. Rev. 239. <http://digitalcommons.law.scu.edu/lawreview/vol33/iss1/7> (1993).
- Ransome v. Wisconsin Elec. Power Co. , 275 N. W. 2d 641, 647-48. Wis. (1979).
- Cole G. S. , 1990, Tort Liability for Artificial Intelligence and Expert Systems, 10 Computer L. J. 127.
- Restatement (Second) of Torts: Section 552: Information Negligently Supplied for the Guidance of Others. (1977).
- Independent School District No. 454 v. Statistical Tabulating Corp 359 F. Supp. 1095. N. D. Ill. (1973).
- Stanley v. Schiavi Mobile Homes Inc. , 462 A. 2d 1144. Me. (1983).
- Helling v. Carey 83 Wash. 2d 514, 519 P. 2d 981 (1974).
- Restatement (Second) of Torts: Sections 520-524. *op. cit.*
- Kociemba v. GD Searle & Co. , 683 F. Supp. 1579. D. Minn. (1988).
- Mykytyn K. , Mykytyn P. P. , Lunce S.: Expert identification and selection: Legal liability concerns and directions. AI & Society, 7, 3, pp. 225-237 (1993)
- Joiner v Mitchell County Hospital Authority, 186 S. E. 2d 307. Ga. Ct. App. (1971).
- Glavin v Rhode Island Hospital, 12 R. I. 411, 435, 34 Amer. Rep. 675, 681 (1879).

- Bloombecker R.: Malpractice in IS? Datamation, 35, pp. 85-86 (1989).
- Warner E.: Expert Systems and the Law. In Boynton and Zmud (eds) Management Information Systems, Scott Foresman/Little Brown Higher Education, Glenview IL, pp. 144-149 (1990).
- Hagendorf W. , 1990, Bulls and Bears and Bugs: Computer Investment Advisory programs That Go Awry. Computer Law Journal, X.
- Jarce, H. , 2010, Sharing Tacit Knowledge. <http://www.jarce.com/2010/01/sharing-tacit-knowledge/>.
- Kingston J. , 2012, Tacit Knowledge: Capture, Sharing, And Unwritten Assumptions. Journal of Knowledge Management Practice, Vol. 13, No. 3.

Legal Civilization

No.18- Winter 2024

ISSN: 2873-1841
ISSN: 2873-1922

The Place of Artificial Intelligence in the Validation of Arbitration Evidence

Homayoun Mafi, Fatemeh Ghanad, Mohammad Amin Esmacilpour

Artificial Intelligence in the Criminal Justice System: Leading Trends and Possibilities

Salar Sadeghi

Challenges and Obstacles of Criminal Liability in Robots with Artificial Intelligence Capabilities

Amin Amirian Farsani, Sayyed Mohammad Hosseini

Artificial Intelligence and its Effect on the Judicial System

Amirreza Mahmoudi, Maryam Bahrekazemi

A Brief History of Artificial Intelligence: On the Past, Present, and Future of Artificial Intelligence

Amin Hajivand, Ali Khosh Manzar, Saber Sayari Zuhan

Artificial Intelligence in the Criminal Justice System: Leading Trends and Possibilities

Salar Sadeghi

Artificial Intelligence and Legal Liability

Sara Solhchi, Kian Biglarbeigi

Artificial Intelligence and Diplomacy Interaction for Environmental Sustainability

Sobhan Tayebi, Nader Tayebi

Artificial Intelligence Crime an Interdisciplinary Analysis of Foreseeable Threats and Solutions

Zahra Vahabi

Artificial Intelligence and Democracy: The Impact of Disinformation, Social Bots and Political Targeting

Sara Solhchi

The Use of Artificial Intelligence in Crime Detection and Criminal Investigations; Case Study: Serial Murders

Hamidreza Heydarpour, Mohammad Shahanaghi, Zhila Mehrara

Ethical Permissibility of Using Artificial Intelligence through the Lens of Al-Farabi's Theory on Natural Rights and Prosperity

Mohamad Mahdi Davar

Artificial Intelligence in the Military: An Overview of the Capabilities, Applications, and Challenges

Yasser Shakeri

SID



بلاگ مرکز اطلاعات علمی



کارگاه‌های آموزشی



سرویس ترجمه تخصصی



فیلم‌های آموزشی

کارگاه‌ها و فیلم‌های آموزشی مرکز اطلاعات علمی

آشنایی با پایگاه‌های اطلاعات علمی بین‌المللی و ترندهای جستجو بین‌المللی و ترندهای جستجو

کاربرد نرم افزار SPSS در پژوهش

پروپوزال نویسی (علوم انسانی)

کاربرد نرم افزار End Note در استناددهی مقالات و متون علمی

صدور گواهینامه نمایه مقالات نویسندگان در SID