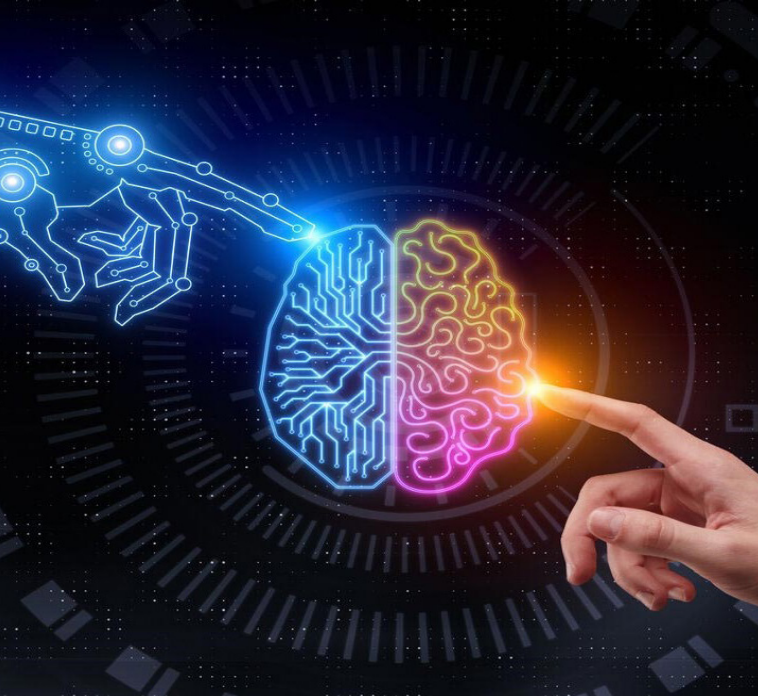




مرکز ملی فضای مجازی
پروژه نگاه فضای مجازی

عصر فضای مجازی صدوسی



هوش مصنوعی تبیین پذیر

Explainable artificial intelligence

بدر

عصر
فضای
مجازی

گزارش شماره ۱۳۰

فروردین ۱۴۰۲

هوش مصنوعی تبیین پذیر

محتوای انتشار یافته در این اثر
الزاماً بیانگر دیدگاه مرکز ملی فضای مجازی نیست

تهیه شده در: پژوهشگاه فضای مجازی - گروه مطالعات بنیادین

تهیه کننده: دکتر امیرحسین زادیوسفی
(استادیار گروه فلسفه دانشگاه تربیت مدرس)

ناظر علمی: دکتر حسین مطلبی کربکندی
(مدیر گروه مطالعات بنیادین)
دکتر مهدی خلیلی

حقوق مادی و معنوی این اثر متعلق به مرکز ملی فضای مجازی است
و استفاده از مطالب آن صرفاً با ذکر مأخذ بلامانع است.

نشانی: تهران، میدان آرژانتین، خیابان بهیقي، نبش خیابان ۱۶ غربی، پلاک ۲۰

کدپستی ۱۵۱۵۶۷۴۳۱۱

شماره تماس: ۸۶۱۲۱۰۶۱

<http://www.majazi.ir>

فهرست

۵ سخن نخست

۹ چکیده

۱۳ مقدمه

بخش اول

۱- یادگیری عمیق و مسئله جعبه سیاه.....۱۹

بخش دوم

۲- هوش مصنوعی تبیین پذیر واکنشی به مسئله جعبه سیاه.....۳۱

بخش سوم

۳- ارزشمندی هوش مصنوعی تبیین پذیر.....۳۹

رویکرد اول (ارزش ابزاری).....۴۰

رویکرد دوم (ارزش ذاتی).....۴۲

نسخه اول (استدلال از طریق مشارکت).....۴۳

نسخه دوم (استدلال از طریق علم).....۴۶

نسخه سوم (استدلال از طریق فعلیت انسان).....۴۷

بخش چهارم

۴- ملاحظات پیرامون هوش مصنوعی تبیین پذیر.....۵۳

بخش پنجم

۵- پیشنهاد راهبردی.....۶۳

۶۵..... جمع بندی

۷۱..... منابع

سخن نخست



فضای مجازی با شتاب شگرف و رو به تزایدی که در حال بسط و گسترش است تمام ساحات اجتماعی، اقتصادی، سیاسی و فرهنگی زندگی بشر را درنوردیده و هر روز بخش بزرگی از زندگی واقعی را در خود فرو برده و حیات متفاوت و جدیدی به آن می‌دهد. لذا به نظر می‌رسد دو نگاه کلان به فضای مجازی وجود دارد: نگاه اول که بالاخص در ابتدای رشد و تکوین فضای مجازی مسلط شده بود، آن را همچون ابزاری کنار سایر ابزارهای بشری تصویر می‌کرد که تنها طریقت داشت. اما نگاه دوم، در نتیجه رشد تحولات خیره‌کننده فضای مجازی و سایه گسترتری آن در حوزه‌ها و شئون بشر در یک دهه اخیر آن را چون سکویی می‌داند که بسیار فراتر از شأن ابزاری حیات انسان‌ها را سامان جدیدی داده و ادعای تمدن‌نویسی را دارد. رویکردی که از قضا از چشمان بصیر رهبر انقلاب نیز دور نمانده و انتظاری تمدنی از فضای مجازی در ایران را مطالبه داشته‌اند.

در همین راستا گزارش‌های عصر فضای مجازی تلاش می‌کند تا فهم سازمان‌ها و دستگاه‌های مرتبط با حوزه فضای مجازی را ارتقاء بخشیده و آن‌ها را برای مواجهه فعال و خردمندانه با تحولات این عرصه مهیا سازد.

سید ابوالحسن فیروزآبادی

دیرشناسی عالی‌ترین مرکز ملی فضای مجازی

چکیده



سیستم‌های هوشمند را می‌توان از یک نقطه نظر به دو بخش تقسیم کرد: سیستم‌های هوشمند خبره و سیستم‌های هوشمند مبتنی بر یادگیری عمیق. ویژگی سیستم‌های خبره این است که می‌توان فرایند اینکه سیستم چگونه به نتیجه خاصی رسیده است را بازبینی نمود. یعنی، این فرایند پوشیده نبوده و تبیین‌پذیر است. اما در مقابل، در سیستم‌های مبتنی بر یادگیری عمیق، فرایند تصمیم‌سازی توسط سیستم هوشمند حتی برای طراحان نیز معلوم نبوده و به یک بیان، فرایند تصمیم‌سازی در این سیستم‌ها مبهم است. این ابهام و تبیین‌ناپذیری به مسئله جعبه سیاه موسوم است. «هوش مصنوعی تبیین‌پذیر» واکنشی است به مسئله جعبه سیاه. هدف از این گزارش، معرفی هوش مصنوعی تبیین‌پذیر است. در این راستا، در ابتدا سیستم‌های خبره و سیستم‌های مبتنی بر یادگیری عمیق را معرفی می‌کنیم. پس از آن به مسئله جعبه سیاه اشاره خواهیم نمود. در بخش بعد، هوش مصنوعی تبیین‌پذیر را به عنوان واکنشی به مسئله جعبه سیاه بررسی خواهیم کرد. با تمایز میان این سؤال که «چگونه می‌توان از نقطه نظر تکنولوژیکی، فنی و مهندسی به هوش مصنوعی تبیین‌پذیر دست یافت؟» و این سؤال که «چرا اساساً هوش مصنوعی تبیین‌پذیر امری ارزشمند است؟»،

سؤال اول را فرو گذاشته و به بیان پاسخ‌های ارائه شده برای سؤال دوم می‌پردازیم. در ادبیات علمی پیرامون هوش مصنوعی تبیین‌پذیر چهار پاسخ به این سؤال وجود دارد که به آنها اشاره خواهیم نمود. پس از اشاره به این چهار پاسخ، ملاحظات را پیرامون این چهار دلیل به پیش می‌کشیم. در بخش بعد، پیشنهادهایی راهبردی را برای استفاده مدیران و تصمیم‌سازان عرصه هوش مصنوعی جمهوری اسلامی ایران ارائه کرده و در نهایت گزارش را با یک جمع‌بندی به پایان می‌بریم.

واژگان کلیدی

یادگیری عمیق، مسئله جعبه سیاه، هوش مصنوعی تبیین‌پذیر

مقدمه



سیستم‌های هوشمند را در یک نگاه کلی می‌توان به دو بخش تقسیم نمود:

سیستم‌های هوشمند خبره^۱ و سیستم‌های هوشمند مبتنی بر یادگیری عمیق^۲. از ویژگی‌های سیستم‌های خبره این است که می‌توان دریافت که سیستم هوشمند چگونه به یک خروجی دست یافته است. در مقابل، سیستم‌های مبتنی بر یادگیری عمیق به گونه‌ای هستند که حتی طراحان آنها نیز نمی‌توانند دریابند که سیستم هوشمند چگونه به خروجی مشخصی رسیده است. از آنجا که فرایند رسیدن به یک خروجی مشخص در سیستم‌های هوشمند مبتنی بر یادگیری عمیق، فرایندی نامعلوم است، به این سیستم‌های «جعبه سیاه»^۳ نیز اطلاق می‌شود. «هوش مصنوعی تبیین‌پذیر» رویکردی است که قصد دارد سیستم هوشمندی را طراحی کند که برخلاف سیستم‌های هوشمند مبتنی بر یادگیری عمیق بتوان این را که سیستم هوشمند چگونه به خروجی مشخصی رسیده است را تبیین کرد. در این گزارش قصد داریم «هوش مصنوعی تبیین‌پذیر» را معرفی کنیم.

-
1. Expert systems
 2. deep learning
 3. Black box

ساختار این گزارش به این شکل است: در بخش دوم به معرفی سیستم‌های هوشمند خبره و سیستم‌های مبتنی بر یادگیری عمیق می‌پردازیم. در بخش سوم، ضمن معرفی مسئله جعبه سیاه، به معرفی رویکرد هوش مصنوعی تبیین‌پذیر برای حل این مسئله خواهیم پرداخت. در بخش چهارم، به پاسخ‌های مختلفی که به سؤال «چرا اساساً هوش مصنوعی تبیین‌پذیر امری ارزشمند است؟» اشاره خواهیم کرد. در بخش پنجم، ملاحظات خود را پیرامون مسئله هوش مصنوعی تبیین‌پذیر ارائه می‌کنیم. در نهایت و در بخش ششم، پیشنهاد‌های راهبردی‌ای را برای استفاده مدیران و تصمیم‌سازان عرصه هوش مصنوعی جمهوری اسلامی بیان می‌کنیم. در نهایت، گزارش را با یک جمع بندی به پایان می‌بریم.

بخش اول



یادگیری عمیق و مسئله جعبه سیاه

همان‌طور که گفتیم سیستم‌های هوشمند، از یک نقطه نظر، به دو دسته تقسیم می‌شوند: سیستم‌های هوشمند خبره و سیستم‌های هوشمند مبتنی بر یادگیری عمیق. اجازه دهید درباره این دو گونه سیستم هوشمند توضیح دهیم. سیستم‌های خبره که می‌توان آنها را سیستم‌های «اگر - آنگاه» نامید سیستم‌های هستند که از چهار بخش اصلی تشکیل می‌شوند. بخش اصلی سیستم‌های خبره چیزی است که به «پایگاه دانش» موسوم است. پایگاه دانش در یک سیستم خبره مجموعه همه اطلاعاتی است که طراحان به سیستم خبره داده‌اند و اغلب به صورت گزاره‌های شرطی بیان می‌شوند. به بیان دقیق‌تر، در پایگاه دانش قوانین مختلفی به صورت جملات شرطی اگر - آنگاه تعبیه شده است. بخش دوم سیستم‌های خبره چیزی است که به «موتور استنتاج» معروف است. موتور استنتاج که شاید بتوان گفت مهم‌ترین بخش یک سیستم خبره است با ترکیب ورودی‌هایی که کاربر به سیستم می‌دهد با آنچه در پایگاه دانش در آن ذخیره شده است، خروجی‌هایی را تحویل می‌دهد. موتور استنتاج از دو طریق می‌تواند به نتایج دسترسی یابد. روش اول به روش مبتنی بر داده و روش دیگر روش مبتنی بر هدف نام دارد. در

روش اول سیستم با استفاده از داده‌ها و نیز گزاره‌های شرطی اگر - آنگاه، از مقدم این گزاره‌ها شروع کرده، به نتایج مناسب می‌رسد. در روش دوم، سیستم برعکس عمل می‌کند و برای نتیجه‌ای خاص به دنبال مقدماتی که می‌توانند به آن نتیجه منجر شوند می‌گردد. برای مثال، فرض کنید که سیستم خبره‌ای برای تشخیص و عیب‌یابی خوردنها طراحی کرده‌ایم. هنگامی که ما به سیستم خبره این ورودی را می‌دهیم که «از گزوز این ماشین دود سیاه رنگی خارج می‌شود» آنگاه این سیستم خبره بر اساس این اطلاعات که «اگر از گزوز ماشینی دود سیاه رنگی خارج شد، آنگاه واشرهای سیلندر و روغن دان بازبینی شود» که در پایگاه دانش آن وجود دارد نتیجه خواهد کرد که «واشرهای سیلندر و روغن دان بازبینی شود». بخش سوم یک سیستم خبره، «امکانات توضیح»^۱ نام دارد که کاربر یا طراح را قادر می‌سازد که بتواند تمام مراحل را که سیستم خبره از ورودی به سمت خروجی طی نموده است دنبال نماید. درواقع، این قابلیت این امکان را به کاربر یا طراح می‌دهد تا دریابد که سیستم خبره چگونه به نتیجه‌ای خاص رسیده است. بخش چهارم یک سیستم خبره، چیزی است که به «رابط کاربر» موسوم است. رابط کاربر این امکان را به کاربر می‌دهد تا بتواند داده‌های را به عنوان ورودی به سیستم خبره بدهد. معمولاً رابط کاربر شامل یک پردازنده زبانی است که می‌تواند زبان طبیعی را که قابل درک برای کاربر انسانی است به داده‌های قابل درک برای سیستم خبره تبدیل کند (Nath, 2015). همین اندازه برای معرفی سیستم‌های خبره کافی است. آنچه درباره سیستم‌های خبره برای ما مهم است ویژگی و شاخصه دوم این سیستم‌ها، یعنی «امکانات توضیح» است. این ویژگی ما را قادر خواهد ساخت از جزئیات اینکه

چگونه سیستم خبره به نتیجه‌ای خاص رسیده است آگاه شویم. حال، که دانستیم سیستم‌های هوشمند خبره چه هستند، اجازه دهید به بررسی سیستم‌های هوشمند مبتنی بر یادگیری عمیق بپردازیم.



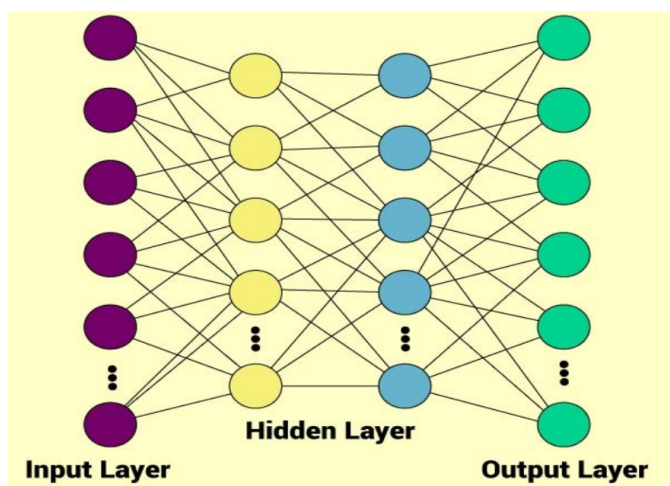
یادگیری عمیق زیرمجموعه‌ای از یادگیری ماشین^۱ است که نوعی فناوری هوش مصنوعی است. به طور خاص، یادگیری عمیق نوعی از یادگیری ماشین است که هدف آن این است که به کامپیوتر بیاموزد که چگونه از طریق مثال‌ها و تجارب بیاموزد. یادگیری ماشین مستلزم این است که داده‌های زیادی به سیستم داده شود تا سیستم بتواند از طریق آنها بیاموزد. باید به این نکته توجه کرد که یادگیری عمیق از نقطه نظر نحوه

کار، با سایر انواع یادگیری ماشین متفاوت است. در یادگیری عمیق سیستم خود، به تنهایی، از داده‌ها یاد می‌گیرد؛ بنابراین ایده یادگیری عمیق این است که کامپیوتر به طور مستقل یاد بگیرد. (بدون اینکه انسان به او بگوید دنبال چه چیزی باشد).

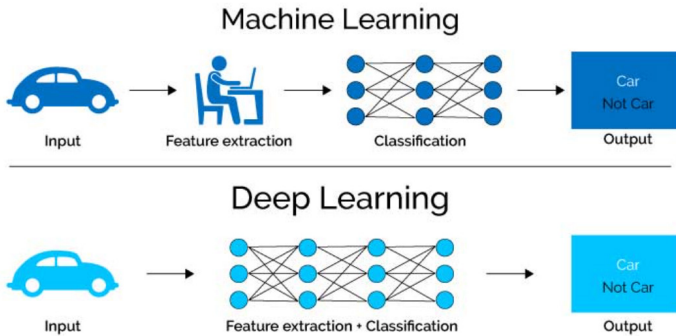
باید به این نکته توجه کرد که یادگیری عمیق از نحوه کارکرد مغز انسان الهام گرفته شده است. در یادگیری عمیق از چیزی به نام شبکه‌های عصبی مصنوعی^۱ استفاده می‌شود. بنابراین، برای پاسخ به سؤال «یادگیری عمیق چیست؟»، در ابتدا باید بدانیم که شبکه‌های مصنوعی عصبی چه هستند.

شبکه عصبی مصنوعی یک سیستم رایانه‌ای است که از نحوه کار مغز انسان الهام گرفته شده است. این شبکه عصبی مصنوعی کمابیش مانند یک فیلتر عمل می‌کند. شبکه‌های عصبی مصنوعی دارای لایه‌هایی از «گره‌ها»^۲ یا «نورون‌ها»^۳ هستند. این گره‌ها هر کدام ورودی را تجزیه و تحلیل می‌کنند و سپس نتایج را با گره‌های موجود در لایه بعدی به اشتراک می‌گذارند. بنابراین، ما به شبکه عصبی مصنوعی داده‌هایی را برای پردازش می‌دهیم. این داده‌ها ورودی سیستم هستند. سپس، تمام گره‌های موجود در لایه اول ورودی را پردازش می‌کنند. سپس خروجی آنها به عنوان ورودی برای لایه بعدی از گره‌ها مورد استفاده قرار گرفته و دوباره پردازش می‌شوند. این فرایند برای هر لایه اتفاق می‌افتد. این لایه‌های میانی «لایه‌های پنهان»^۴ نام دارند. آنچه در نهایت سیستم به عنوان خروجی به ما تحویل می‌دهد پاسخ شبکه عصبی مصنوعی به ورودی‌ای است که به آن داده بودیم (Grover, 2019).

1. artificial neural networks (ANN)
2. nodes
3. Neurons
4. hidden layers



بنابراین، در پاسخ به این سؤال که «یادگیری عمیق چیست؟» باید بگوییم که «یادگیری عمیق یک شبکه عصبی مصنوعی بسیار عظیم است». به طور معمول، یک شبکه عصبی مصنوعی فقط چند لایه پنهان (حداکثر ۲ تا ۳ لایه) دارد که ورودی از طریق آن فیلتر می‌شود. اما در یادگیری عمیق تعداد این لایه‌های میانی بسیار زیاد است (مثلاً میلیون‌ها لایه میانی). واژه «عمیق» در عنوان «یادگیری عمیق» دقیقاً به‌خاطر وجود همین لایه‌های میانی بسیار زیاد است. به طور خلاصه، یادگیری عمیق یک شبکه عصبی مصنوعی عظیم است که به ماشین‌ها اجازه می‌دهد از طریق مثال یاد بگیرند. هرچه نمونه‌ها و مثال‌های بیشتری به سیستم‌های مبتنی بر یادگیری عمیق ارائه شود، سیستم عملکرد بهتری خواهد داشت.



اجازه دهید برای درک بهتر عملکرد سیستم‌های مبتنی بر یادگیری عمیق مثالی بزنیم. فرض کنید کودک نوپایی اولین کلمه‌ای فراگرفته است «گربه» است و می‌تواند به خوبی مصادیق گربه را از سایر حیوانات تمایز دهد. اجازه دهید ببینیم این کودک چگونه توانسته است به این توانایی دست یابد. این کودک در هر بار به شیئی اشاره کرده و واژه گربه را برای نامیدن آن بکار می‌برد. والدین او نیز در هر بار به او این بازخورد را می‌دهند که آیا شیئی که کودک به آن اشاره کرده است یک گربه است یا نه. کودک به مرور زمان و از طریق تجارب مختلف و بازخوردهای والدینش در نهایت به این توانایی خواهد رسید که بتواند گربه را از غیر آن متمایز کند. سیستم‌های مبتنی بر یادگیری عمیق نیز به همین شکل هستند. آن‌ها نیز از طریق تجربه و از طریق بازخوردی که به آن داده می‌شود به مرور زمان می‌توانند بیاموزند که خروجی خواسته شده از آن‌ها را به کاربر تحویل دهند.

اجازه دهید به برخی مثال‌ها و کاربردهای یادگیری عمیق اشاره کنیم. یکی از کاربردهای الگوریتم‌های یادگیری عمیق برای تشخیص و دسته

بندی اشیاء است. برای مثال، می‌توان سیستمی طراحی کرد که چهره افراد را تشخیص دهد و یا بتواند اشیاء خاصی را که مطلوب ماست، مثلاً ماشین پراید را، از سایر اشیاء دیگر تشخیص دهد. یکی دیگر از کاربردهای موفق یادگیری عمیق در پردازش زبان طبیعی است. برای مثال، مترجم گوگل که می‌تواند گفتارها و متون‌هایی را از زبان مبدأ به زبان مقصد ترجمه کند نمونه‌ای است از کاربردهای یادگیری عمیق. یکی دیگر از کاربردهای یادگیری عمیق در خودروهای خودران است. این خودروها از طریق یادگیری عمیق می‌آموزند که شرایط محیطی را، همانند یک راننده انسانی، تشخیص دهند. برای مثال، یاد خواهند گرفت هنگامی که کودکی به وسط خیابان پدید آن را به عنوان یک انسان تشخیص دهد و سپس ترمز نماید. به نظر می‌رسد همین اندازه برای آشنایی با یادگیری عمیق کافی است (Grover, 2019). حال اجازه دهید «مسئله جعبه سیاه»^۱ را معرفی کنیم.

هر سیستم محاسباتی یک ورودی را دریافت کرده، بر روی آن پردازشی می‌کند و در نهایت خروجی‌ای را تحویل می‌دهد. اما یک سیستم محاسباتی یا یک برنامه هوشمند «جعبه سیاه» نامیده می‌شود به شرطی که به ما این امکان را بدهد که ورودی و خروجی را مشاهده کنیم، ولی هیچ فرایندی و عملکرد بین ورودی و خروجی را مشاهده نکنیم. بنابراین جعبه سیاه هوش مصنوعی به این واقعیت اشاره دارد که با بیشتر ابزارهای مبتنی بر هوش مصنوعی، ما نمی‌دانیم که آنها چگونه کارهایی را که انجام می‌دهند را انجام می‌دهند. به عبارت دیگر، ما از داده‌هایی که سیستم هوشمند با آنها شروع می‌کند (یعنی ورودی) آگاهییم (به عنوان مثال، عکس‌های حیوانات). همچنین، ما از پاسخی که

سیستم هوشمند تولید می‌کند (یعنی خروجی) نیز آگاهیم (مثلاً، جدا کردن تصاویر گربه از تصاویر سایر حیوانات). اما به‌خاطر مسئله جعبه سیاه، ما نمی‌دانیم که چگونه توانسته است از ورودی‌هایی که به آن داده‌ایم آن خروجی‌های خاص را به ما تحویل دهد. اما این عدم آگاهی هنگامی که سیستم هوشمند خروجی‌های نادرست، ناخواسته و مشکل‌دار را به ما تحویل بدهد، خود را به صورت یک مسئله نمایان می‌کند.

اما می‌توان پرسید «چه چیزی سبب بروز مسئله جعبه سیاه می‌شود؟» در پاسخ به این سؤال باید گفت که رایج‌ترین ابزارهایی که از مسئله جعبه سیاه رنج می‌برند، ابزارهایی هستند که از شبکه‌های عصبی مصنوعی یا همان یادگیری عمیق که در بالا درباره آن سخن گفتیم استفاده می‌کنند. همان‌طور که در بالا اشاره کردیم، شبکه‌های عصبی مصنوعی دارای بی‌شمار لایه‌های مخفی از گره‌ها (نورون‌ها) هستند. این گره‌ها هر یک ورودی داده شده را پردازش کرده و خروجی خود را به لایه بعدی گره‌ها منتقل می‌کنند. همان‌طور که بیان کردیم، یادگیری عمیق، یک شبکه عصبی مصنوعی عظیم، با بی‌شمار لایه میانی پنهان است که خود می‌تواند از طریق مثال‌ها و تجربه‌ها بیاموزد. اما نکته اینجاست که ما نمی‌توانیم دریابیم که گره‌ها چه چیزهایی را یاد گرفته‌اند. در واقع، ما از خروجی بین لایه‌ها اطلاعی نداشته و فقط نتیجه را دریافت می‌کنیم. از این‌رو، ما نمی‌توانیم نحوه تجزیه و تحلیل داده‌ها توسط گره‌ها را دریابیم و در واقع (بر اساس تعریف بالا) با یک جعبه سیاه روبرو خواهیم بود (Zednik, 2019).

از آنجا که سیستم‌های هوشمند مبتنی بر یادگیری عمیق برگرفته و الهام گرفته از کارکرد مغز انسان هستند، مسئله جعبه سیاه درباره مغز

انسان نیز مطرح می‌شود. چرا که ما دقیقاً نمی‌دانیم که مغز ما چگونه کار می‌کند و چگونه به نتیجه یا یک خروجی خاص می‌رسد. در واقع، به هیچ وجه مشخص نیست که در سطحی پایه‌ای، چگونه تصمیمات خود را (مثلاً تصمیم به نوشیدن یک لیوان آب)، اخذ می‌کنیم. در واقع، برای اینکه بتوانیم چنین تصمیم ساده‌ای را اتخاذ کنیم تعداد زیادی نورون و سیناپس در مغزمان فعالیت می‌کنند.

حال که دانستیم مسئله جعبه سیاه چگونه مسئله‌ای است اجازه دهید توضیح دهیم چرا جعبه سیاه یک مسئله و مشکل قلمداد می‌شود و اینکه چه رویکردی برای حل این مسئله وجود دارد. در بخش بعد به این مطلب می‌پردازیم.

بخش دوم



هوش مصنوعی تبیین پذیر و انکشف به مسئله جعبه سیاه

در بخش قبل بیان کردیم که مسئله جعبه سیاه این مسئله است که حتی طراحان سیستم هوشمند نیز نمی دانند که چگونه سیستم هوشمند به خروجی خاصی رسیده است. اما چرا این مسئله باید یک مشکل قلمداد شود. در ادامه به این سؤال پاسخ می دهیم.

گاهی اوقات سیستم های هوشمند مبتنی بر یادگیری عمیق به نتایجی دست پیدا می کنند که مخالف با اهدافی است که طراحان آن برای آن در نظر گرفته اند. در واقع، گاهی اوقات این سیستم ها از داده ها و تجربیات طوری می آموزند که نتایجی برخلاف آنچه برای آن طراحی شده اند را ارائه می دهند. برای مثال، یک سیستم مبتنی بر یادگیری عمیق که برای تشخیص تصاویر اسب طراحی شده بود، در نهایت یاد گرفته بود که تصاویری را که بر روی آن برچسب کپی رایت وجود دارد را از تصاویری که این برچسب را ندارد متمایز کند (چرا که بالاتفاق بر روی تصاویر اسبی که به آن نمایش داده می شود برچسب کپی رایت نیز وجود داشت).

یکی از اشتباهات و خطاهای سیستم های هوشمند، سیستم هوش مصنوعی ساخت شرکت IBM است که به پزشکان در حوزه تشخیص

بیماری یاری می‌رساند. در سال ۲۰۱۳ شرکت IBM اولین محصول تجاری خود را برای درمان سرطان تولید کرد. اما این در حالی بود که برخی از پزشکان بیان کردند که این سیستم هوشمند توصیه‌های اشتباهی را برای درمان سرطان ارائه می‌دهد که می‌تواند سبب مشکلات جدی و حتی نتایج مرگبار شود (Peng, 2018).

مثال دیگر، کشته‌شدن یک عابر پیاده در اثر تصادف با یک اتومبیل خودران اوبر^۱ بود که در ۲۸ مارس سال ۲۰۱۸ در آریزونا، آمریکا اتفاق افتاد. بعد از بررسی شرکت اوبر اعلام کرد که ماشین خودران بعد از اینکه تشخیص داده بود که یک عابر پیاده در مقابلش هست ولی تصمیم گرفته بود که هیچ اقدامی را انجام ندهد (Peng, 2018).

می‌توان حالت‌های فرضی مختلف دیگری را نیز برای اشتباهات اتومبیل‌های خودران در نظر گرفت. برای مثال، سیستم هوشمند اتومبیل‌های خودران می‌تواند کودکان کوتاه قامت را به‌عنوان یک عابر انسانی شناسایی نکند یا در صورتی که امر دائر شود میان تصادف با یک ماشین دیگر و تصادف با یک عابر پیاده این‌طور تشخیص دهد که تصادف با عابر پیاده بهتر است از تصادف با یک ماشین دیگر. حتی سیستم هوشمند به دلایلی نامعلوم ممکن است این را آموخته باشد که در برابر عابرانی که با خود یک جعبه حمل می‌کنند توقف نکند.

1. Uber self-driving



یکی دیگر از خطاهای سیستم‌های هوشمند مبتنی بر یادگیری عمیق، سیستمی است که شرکت آمازون برای استخدام افراد استفاده کرده است. شرکت آمازون در خلال سال‌های ۲۰۱۴ تا ۲۰۱۷ از سامانه هوشمندی برای بررسی بی‌شمار رزومه ارسالی و ارائه پیشنهاداتی استفاده می‌کرده است. اما این شرکت بعدها متوجه شد که سیستم هوشمند به دلایل نامشخصی متقاضیان مرد را نسبت به متقاضیان زن در اولویت قرار می‌دهد (Peng, 2018). همین مقدار برای مثال‌ها کافی است.

مثال‌های بالا، همگی حاکی از یک مشکل اخلاقی است که ناشی از نادیده گرفتن مسئله جعبه سیاه است. به نظر می‌رسد در صورتی که سیستم‌های هوشمند بتوانند نحو نادرستی بیاموزند و در نتیجه خروجی‌هایی را ارائه کنند که مطلوب ما نبوده است آنگاه دیگر نمی‌توان به این سیستم‌ها اعتماد نمود. برای مثال، در حوزه پزشکی اگر سیستم هوشمند طوری آموخته باشد که توصیه‌های

اشتباهی برای درمان ارائه دهد، آنگاه این مسئله جان بیماران را به خطر خواهد انداخت. همچنین، اگر سیستم هوشمند در بررسی رزومه‌های متقاضیان برای کار طوری آموخته باشد که نسبت به مردان سوگیری داشته باشد، آنگاه فرصت‌های شغلی برای زنان لایق از دست خواهد رفت و این نتیجه‌ای جز تبعیض جنسیت و نابرابری نخواهد داشت. نیز در صورتی که ماشین‌های خودران به اشتباه بیاموزند که در برابر عابروانی خاص (مثلاً عابروانی که با خود یک جعبه حمل می‌کنند) توقف نکنند، آنگاه جان افراد به خطر خواهد افتاد.

حال می‌توان گفت که مسئله جعبه سیاه مستلزم یک مسئله مهم‌تر، یعنی مسئله «عدم اعتماد به هوش مصنوعی» است. در واقع، در صورتی که ما ندانیم که آیا سیستم هوشمندی که از آن استفاده می‌کنیم قابل اعتماد است یا خیر، تمایلی برای استفاده از آن نخواهیم داشت.

در واکنش به مسئله جعبه سیاه و به دنبال آن، مسئله عدم اعتماد، در چند سال اخیر جریانی به راه افتاده است که به «هوش مصنوعی تبیین‌پذیر»^۱ موسوم است. هوش مصنوعی تبیین‌پذیر هوش مصنوعی‌ای است که می‌توان فرایند تصمیم‌سازی‌اش را درک نمود. در واقع، در هوش مصنوعی تبیین‌پذیر انسان‌ها می‌توانند دریابند که یک سیستم هوشمند چگونه به تصمیم خاصی رسیده است. بنابراین، با هوش مصنوعی تبیین‌پذیر می‌توان اشکالات و خطاهای سیستم‌های هوشمند را دریافت و از این‌ها از بروز مشکلات اخلاقی، ناعدالتی و حتی فاجعه‌بار جلوگیری نمود.

1. Explainable Artificial Intelligence (XAI)

در ادبیات علمی در باب هوش مصنوعی تبیین پذیر معمولاً دو سؤال از یکدیگر متمایز می‌شوند (Colaner, 2021). پرسش اول این است که :

«چگونه می‌توان از نقطه نظر تکنولوژیکی، فنی و مهندسی به هوش مصنوعی تبیین پذیر دست یافت؟»^۱

پرسش دومی که می‌توان در حوزه هوش مصنوعی تبیین پذیر پرسید این پرسش است که:

«چرا اساساً هوش مصنوعی تبیین پذیر امری ارزشمند است؟»

ما در این گزارش با پرسش اول کاری نخواهیم داشت. پرسش اول پرسشی است مربوط به مسائل فنی در حوزه هوش مصنوعی. در عوض، در بخش‌های بعد به پرسش دوم خواهیم پرداخت.

۱. برای نمونه ببینید: (Vellido et al; 2012. Kim et al; 2016. Tamagnini et al; 2017. Guidotti et al; 2018. Ayesha et al; 2020)

بخش سوم



ارزشمندی هوش مصنوعی تبیین پذیر

در ادبیات علمی شکل گرفته پیرامون هوش مصنوعی تبیین پذیر چهار دلیل عمده برای پاسخ به این پرسش که «چرا اساساً هوش مصنوعی تبیین پذیر امری ارزشمند است؟» ارائه شده است. در ادامه این چهار پاسخ را بررسی خواهیم کرد. این چهار پاسخ در دو دسته تقسیم بندی می شوند:

پاسخ هایی که به ارزش ابزاری هوش مصنوعی تبیین پذیر اشاره دارند و پاسخ هایی که به ارزش ذاتی هوش مصنوعی تبیین پذیر اشاره دارند. منظور از ارزش ابزاری این است که هوش مصنوعی تبیین پذیر به این دلیل خوب است که وسیله و ابزاری است برای دست یافتن به چیزهای دیگری که برای ما مطلوب و ارزشمند هستند. اما در مقابل، پاسخ هایی که سعی دارند ارزش هوش مصنوعی تبیین پذیر را بر اساس یک ارزش ذاتی توضیح دهند تلاش می کنند بیان کنند هوش مصنوعی تبیین پذیر به خاطر ویژگی های خوبی که دارد ارزشمند است. در این رویکرد بیان نمی شود که هوش مصنوعی تبیین پذیر خوب است چون نتایج خوبی برای ما به دنبال دارد. بلکه در عوض، خود سیستم تبیین پذیر به خودی خود

ویژگی‌های ارزشمندی دارد، فارغ از اینکه آیا نتایج حاصل از آن برای ما مطلوب باشد یا نه. در ادامه این پاسخ‌ها را بررسی خواهیم کرد.

رویکرد اول (ارزش‌ابزاری)

پاسخ اول به این پرسش که «چرا اساساً هوش مصنوعی تبیین‌پذیر امری ارزشمند است؟» (که در بخش‌های قبل مختصراً به آن اشاره شد) این است که امر مطلوبی را ذکر کنیم که بتوان از طریق هوش مصنوعی تبیین‌پذیر، به آن دست‌یافت. در واقع، در این پاسخ سعی می‌شود به نحو غیرمستقیم توضیح داده شود که چرا ارائه تبیین امر ارزشمندی است. به بیان دیگر، بر اساس این رویکرد هوش مصنوعی تبیین‌پذیر ابزاری است برای دستیابی به آن فایده خوب (Colaner, 2021). یک مثال خوب برای این پاسخ، صورت‌بندی است که آژانس پروژه‌های تحقیقاتی پیشرفته دفاعی ایالات متحده آمریکا^۱ که هدف آن تلاش برای فهم تکنولوژی‌های نو ظهور در حوزه امنیت ملی است، ارائه داده است. یکی از اهداف این آژانس تولید سیستم‌های هوشمندی است که برای انسان‌ها توضیح‌پذیر باشند. از نظر آن‌ها اگر انسان‌ها موفق شوند هوش مصنوعی تبیین‌پذیر را تولید کنند، آنگاه شرایطی برای کاربران محقق خواهد شد که بتوانند به نحو درستی به هوش مصنوعی اعتماد کرده و به نحو مؤثری آن را مدیریت کنند^۲ (DARPA, 2016). همان‌طور که پیداست «اعتماد» و «مدیریت» اهداف مطلوبی هستند که می‌توان از طریق تبیین به آن‌ها دست یافت. به بیان دیگر، می‌توان گفت دو امر ارزشمند به نام‌های «اعتماد» و «مدیریت»

1. Defense Advanced Research Projects Agency (DARPA)
2. appropriately trust, and effectively manage

وجود دارد که می‌توان از طریق ایجاد هوش مصنوعی تبیین‌پذیر به آن‌ها دست یافت؛ اموری که نمی‌توان از طریق سیستم‌های هوشمند فعلی (که تبیین‌پذیر نیستند) به آن‌ها دست پیدا کرد. بنابراین، می‌توان گفت که تبیین‌پذیری در این رویکرد ارزشی ابزاری دارد. یعنی، تبیین‌پذیری ارزشمند است زیرا سبب می‌شود به دو ارزش دیگر، یعنی اعتماد به و مدیریت هوش مصنوعی دستیابیم.

این رویکرد، یعنی ارزش ابزاری تبیین‌پذیری (برای دستیابی به دو ارزش اعتماد و مدیریت) همچنین در «آیین‌نامه عمومی حفاظت از داده‌ها»^۱ نیز مورد توجه قرار داده شده است. یکی از پایه‌های این چارچوب مقرراتی، تأکید بر داشتن حق تبیین و توضیح^۲ است. برای مثال، اگر فردی که برای استخدام در شرکت آمازون اقدام کرده و رزومه وی توسط بررسی اولیه سیستم هوشمند رد شده است، این حق را دارد که بپرسد «چرا درخواست من رد شده است؟». در واقع، از نظر این چارچوب مقرراتی، هنگامی که یک الگوریتم هوشمند تصمیمی را اتخاذ می‌کند، این حق برای افراد (مثلاً کسانی که از آن تصمیم متأثر می‌شوند) وجود دارد که بدانند چرا سیستم هوشمند به آن تصمیم رسیده است (Wachter et al, 2017).

همان‌طور که پیداست در رویکرد نخست تبیین‌پذیری امری ارزشمند است زیرا وسیله و ابزاری است برای دست‌یافتن به ارزش‌های دیگری همچون، «اعتماد»، «قابلیت مدیریت» و «حق تبیین». بنابراین، می‌توان استدلالی را مانند زیر برای رویکرد نخست ترتیب داد:

(۱) اگر چیزی مانند الف ابزار و راهی برای رسیدن به ارزش دیگری

مانند ب باشد، آنگاه الف نیز خود دارای ارزش خواهد بود.
 ۲) هوش مصنوعی تبیین‌پذیر ابزاری است برای رسیدن به اموری همچون «اعتماد»، «قابلیت مدیریت» و «حق تبیین».
 ۳) «اعتماد»، «قابلیت مدیریت» و «حق تبیین» اموری ارزشمند هستند.
 ۴) بنابراین، هوش مصنوعی تبیین‌پذیر ارزشمند است.

رویکرد دوم (ارزش ذاتی)

برخلاف رویکرد اول در پاسخ به این سؤال که «چرا اساساً هوش مصنوعی تبیین‌پذیر امری ارزشمند است؟»، رویکرد دیگری وجود دارد که بیان می‌دارد که اساساً هوش مصنوعی تبیین‌پذیر ذاتاً امری ارزشمند است. رویکرد دوم، با این سؤال آغاز می‌کند که آیا اساساً می‌توان تبیین‌ها را اموری دانست که ذاتاً دارای ارزش هستند؟ این امکان که تبیین و توضیح امری ذاتاً ارزشمند باشد امری معنادار خواهد بود در صورتی که بپذیریم اگر انسان بخشی از یک فرایند غیر تبیین‌پذیر باشد، آنگاه این امر مستقیماً بر روی شرایط انسانی وی تأثیر خواهد گذاشت (با صرف‌نظر از اینکه آیا خروجی فرایند تبیین‌پذیر مطلوب وی باشد یا نه). به بیان دیگر، از نقطه‌نظر این رویکرد تبیین‌پذیری ذاتاً ارزشمند است فارغ از اینکه تبیین‌پذیری چه نتایجی را به دنبال دارد. این استدلال گاه با نام «استدلال فردیت»^۱ شناخته می‌شود زیرا مبتنی بر سؤالاتی در باب فردیت از قبیل خودمختاری فردی^۲ و کرامت انسانی^۳ است.

در ادبیات بحث سه نسخه از استدلال فردیت ارائه شده است که در

1. personhood argument
 2. individual autonomy
 3. human dignity

ادامه هر یک را بیان می‌کنیم.

نسخه اول (استدلال از طریق مشارکت)

هنگامی که تصمیمات و اقدامات یک سیستم هوشمند غیر قابل تبیین و توضیح باشد، آنگاه انسان‌ها به عنوان موجوداتی که قرار است از تصمیمات و اقدامات هوش مصنوعی متأثر شوند دیگر نمی‌توانند در فرایند تصمیم‌سازی مشارکتی داشته باشند. برای مثال، همان‌طور که در بخش‌های قبل بیان کردیم هوش مصنوعی در حوزه پزشکی تصمیماتی را برای بیمار می‌گیرد که نه بیمار و نه پزشک معالج وی هیچ مشارکتی در اتخاذ آن تصمیم نداشته‌اند. بنابراین، می‌توان ادعا کرد هرچقدر که یک سیستم هوشمند بیشتر غیر قابل تبیین و توضیح باشد، ما انسان‌ها کمتر می‌توانیم در اتخاذ یک تصمیم مشارکت داشته باشیم. برای تقریب به ذهن، اجازه دهید مثالی را از حوزه عدالت اجتماعی ارائه کنیم. محققین حوزه عدالت اجتماعی بیان می‌کنند هنگامی که گروهی از انسان‌ها از فرایند تصمیم‌سازی برای همه انسان‌ها کنار گذاشته می‌شوند، آنگاه ناعدالتی اجتماعی رخ داده است. برای مثال، یانگ^۱ این مثال را بیان می‌کند که کارگران یک کارخانه از این اطلاعیه که رئیس یک کارخانه بزرگ در شهر، قصد دارد کارخانه‌اش را تعطیل کند بسیار خشمگین هستند. حرف آن‌ها این است که چرا تصمیم‌سازان این کارخانه بدون اطلاع و در جریان گذاشتن، مشورت و هشدار قبلی به کارگران کارخانه، چنین تصمیمی را اتخاذ کرده‌اند (Young, ۱۹۹۰). باید به این نکته دقت کرد که نکته این مثال نتایج ناعادلانه

تعطیلی کارخانه برای کارگران نیست. نکته این مثال این است که کارگران از این که چرا آن‌ها در فرایند تصمیم‌سازی برای مسئله به این مهمی دخالت داده نشده‌اند عصبانی هستند.

در واقع، می‌توان فرض کرد که مدیران کارخانه به کارگران این وعده را بدهند که بعد از تعطیلی کارخانه به آن‌ها مبلغ زیادی پول بلاعوض خواهند داد و حق بیمه آن‌ها را نیز پرداخت خواهند کرد. اگرچه این شرایط و این نتیجه به لحاظ اقتصادی عادلانه به نظر می‌رسد، ولی نکته این است که کارگران کماکان می‌توانند از تصمیم اتخاذ شده ناراضی باشند، تنها به این دلیل که آن‌ها در جریان این تصمیم مهم قرار داده نشده و مشارکتی برای اتخاذ این تصمیم نداشته‌اند.

در واقع، ممکن است افراد بتوانند با نتایج یک تصمیم کنار بیایند ولی صرفاً به این دلیل که آن‌ها در جریان اتخاذ تصمیم قرار نگرفته‌اند احساس ناعدالتی نمایند. همچنین، این اصل که عدالت مستلزم گوش فرادادن به سخن ذی‌نفعان در خلال یک فرایند تصمیم‌سازی است، از نقطه‌نظر اصول مدیریت حقوق سیاسی نیز بحث شده است (Zhu et al, 2019). با توجه به این توضیحات می‌توان نسخه اول از استدلال فردیت را مانند زیر بیان کرد (Colaner, 2019):

(۱) در مواقعی که نتایج و خروجی‌های یک فرایند تصمیم‌سازی بر روی فرد یا جامعه تأثیر می‌گذارد، کرامت انسانی مستلزم فرصت مشارکت معنادار در آن فرایند است.

(۲) الگوریتم‌های مبتنی بر یادگیری عمیق اساساً غیرقابل‌درک و

توضیح هستند.

۳) بنابراین، نمی‌توان در فرایند تصمیم‌سازی الگوریتم‌های مبتنی بر یادگیری عمیق مشارکت معنادار داشت.

۴) بنابراین، استفاده از الگوریتم‌های مبتنی بر یادگیری عمیق در تصمیم‌سازی برای فرد یا جامعه، مستلزم نقض کرامت انسانی خواهد شد.

همان‌طور که در استدلال بالا پیداست، از آنجاکه سیستم‌های هوشمند مبتنی بر یادگیری عمیق تبیین‌ناپذیر هستند، بنابراین، انسان‌ها نمی‌توانند در فرایند تصمیم‌سازی آن‌ها مشارکتی داشته باشند و از آنجاکه مشارکت در یک فرایند تصمیم‌سازی که نتایج آن متوجه فرد و جامعه است، شرط ضروری برای کرامت انسانی است، پس می‌توان نتیجه گرفت استفاده از یادگیری عمیق مستلزم نقض کرامت انسانی است. در استدلال بالا باید به این نکته دقت شود که برای ما مهم نیست که نتایج حاصل از تصمیمات مبتنی بر یادگیری عمیق درست است یا خیر، اخلاقی است یا خیر، قانونی است یا خیر، برایمان سودمند است یا خیر. حتی اگر نتایج حاصل از یادگیری عمیق درست، اخلاقی، قانونی و سودمند نیز باشد، صرفاً به این دلیل که ما انسان‌ها در فرایند تصمیم‌سازی مشارکتی نداشته‌ایم استفاده از آن منافی امری است که ارزش ذاتی دارد (یعنی کرامت انسانی).

می‌توان استدلال بالا را به نحو ایجابی نیز بیان کرد:

۱) استفاده از هوش مصنوعی تبیین‌پذیر به این معنی است که بتوان در فرایند تصمیم‌سازی مشارکت داشت.

(۲) مشارکت در فرایند تصمیم‌سازی ذاتاً ارزشمند است
 (۳) بنابراین، استفاده از هوش مصنوعی تبیین‌پذیر ذاتاً ارزشمند است.
 همان‌طور که مشخص است این استدلال قصد ندارد ارزشمندی
 استفاده از هوش مصنوعی تبیین‌پذیر را به‌خاطر ارزش دیگری همچون
 اعتماد و مدیریت‌پذیری بیان کند. این استدلال بیان می‌کند صرف
 هوش مصنوعی تبیین‌پذیر، فارغ از اینکه نتیجه تصمیمات آن چه
 باشد، امری ارزشمند است، تنها به این دلیل که حق مشارکت در
 تصمیمات برای انسان‌ها محفوظ است.

نسخه دوم (استدلال از طریق علم)

نسخه دوم استدلال فردیت، برخلاف نسخه اول که بر «مشارکت»
 در امر تصمیم‌سازی تأکید می‌کرد، بر «علم» تأکید دارد
 (Colaner, 2021). این استدلال می‌توان مانند زیر بیان کرد:
 (۱) کرامت انسانی شامل این است که من بدانم چرا آنچه برای
 من اتفاق افتاده است، اتفاق افتاده است

(۲) اما در صورت استفاده از سیستم‌های مبتنی بر یادگیری عمیق
 (تبیین‌ناپذیر) علم به چرایی اینکه چرا چیزی برای من اتفاق
 افتاده است غیرممکن است

(۳) بنابراین، در صورت استفاده از سیستم‌های هوشمند مبتنی بر
 یادگیری عمیق، کرامت انسانی نقض شده است.

همان‌طور که پیداست نسخه دوم، بر این تأکید دارد که اصل اینکه
 فرد بداند برای او چه اتفاقی افتاده است و چرا سیستم هوشمند
 تصمیمی خاصی را برای وی اتخاذ کرده است امری ارزشمند و مهم

است. زیرا در این صورت است که کرامت انسانی انسان‌ها حفظ شده است. همان‌طور که پیداست، در استدلال بالا برای ما مهم نیست که نتایج حاصل از تصمیمات مبتنی بر یادگیری عمیق درست است یا خیر، اخلاقی است یا خیر، قانونی است یا خیر، برایمان سودمند است یا خیر. حتی اگر نتایج حاصل از یادگیری عمیق درست، اخلاقی، قانونی و سودمند نیز باشد، صرفاً به این دلیل که ما انسان‌ها به اینکه چرا سیستم هوشمند تصمیم خاصی را گرفته است علم نداریم، استفاده از آن دارای اشکال می‌شود. در مقابل، از آنجاکه استفاده از هوش مصنوعی تبیین‌پذیر امکان علم را به ما می‌دهد همین امر فی ذاته، سبب ارزشمندی استفاده از هوش مصنوعی تبیین‌پذیر می‌شود.

نسخه سوم (استدلال از طریق فعلیت انسان)^۱

نسخه سوم از استدلال فردیت که نیشن کولانر^۲ آن را ارائه می‌کند، بر این امر تأکید دارد که در صورتی انسان‌ها می‌توانند به خود و جامعه خود را فعلیت بخشند که برای خود و جامعه تصمیم اتخاذ نمایند. به بیان دیگر، ایده اصلی این نسخه از استدلال فردیت این است که اتخاذ تصمیم برای تحقق و فعلیت انسانیت انسان امری ضروری است (Colaner, 2021). نکته مهمی که باید به آن توجه کرد است که این نسخه از استدلال با نسخه اول ارتباط دارد اما دقیقاً همان نیست.

نکته این است که نسخه اول بر روی مسئله «قدرت در مشارکت» تأکید داشت اما این نسخه بر مسئله «فعلیت انسان» تأکید دارد.

این نسخه از استدلال ریشه در ایده‌های ارسطو در باب ماهیت و طبیعت انسان دارد. از نظر ارسطو ما انسان‌ها بیش از موجوداتی همچون زنبورها، اجتماعی هستیم.

در واقع از نظر ارسطو انسانیت ما انسان‌ها در گرو زندگی اجتماعی با یکدیگر است. بنابراین، این خصیصه سبب خواهد شد که ما انسان‌ها به این سمت متمایل شویم که چگونه جامعه خود را سامان دهیم تا بتوانیم در سایه عدالت در کنار یکدیگر زندگی کنیم. ارسطو در این باب بیان می‌دارد:

فردی که قادر به زندگی در جامعه نیست، یا فردی که به دلیل خودکفا بودن هیچ نیازی ندارد، یا باید یک جانور باشد یا یک خدا. چنین فردی هیچ بخشی از یک جامعه نیست [...] عدالت پیوند انسان‌ها در یک جامعه است (Aristotle, 1984).

با توجه به این نکته که لازمه اجتماعی بودن انسان تصمیم‌سازی برای خود و جامعه است و از طرف دیگر فعلیت گوهر انسانی نیز مبتنی بر تصمیم‌سازی وی برای خود و جامعه است، می‌توان نسخه سوم از استدلال فردیت را، به‌صورت تقریبی، مانند زیر بیان کرد:

- (۱) فعلیت کامل انسانیت انسان مستلزم این است که وی برای زندگی خود و جامعه‌اش تصمیم‌سازی کند (نه سیستم هوشمند).
- (۲) اما هوش مصنوعی تبیین‌ناپذیر این فرصت را به انسان نمی‌دهد تا بتواند برای خود و جامعه تصمیم بگیرد.
- (۳) بنابراین، استفاده از هوش مصنوعی تبیین‌ناپذیر سبب خواهد

شد انسانیت انسان فعلیت نیابد.

(۴) اما هوش مصنوعی تبیین‌پذیر چنین مشکلی را پدید نمی‌آورد.

(۵) بنابراین، هوش مصنوعی تبیین‌پذیر ارزشمند است.

بنا بر استدلال فوق، هوش مصنوعی تبیین‌پذیر، برخلاف هوش مصنوعی تبیین‌ناپذیر، به این دلیل که سبب می‌شود انسان بتواند از طریق دخالت در فرایند تصمیم‌سازی برای خود و جامعه حقیقت خود را فعلیت بخشد، امری ارزشمند است. همان‌طور که پیداست، این استدلال بیان نمی‌کند هوش مصنوعی تبیین‌پذیر به دلیل اینکه سبب می‌شود نتایج مطلوبی به دست آید امری ارزشمند است. بلکه این استدلال صرفاً نظر از نتایج عملی هوش مصنوعی تبیین‌پذیر (از قبیل کاهش خطاها، اعتماد و ...) خود تبیین‌پذیری را امری دارای ارزش می‌داند.

به بیان دیگر، در استدلال بالا برای ما مهم نیست که نتایج حاصل از تصمیمات مبتنی بر یادگیری عمیق درست است یا خیر، اخلاقی است یا خیر، قانونی است یا خیر، برایمان سودمند است یا خیر. حتی اگر نتایج حاصل از یادگیری عمیق درست، اخلاقی، قانونی و سودمند نیز باشد، صرفاً به این دلیل که ما انسان‌ها از طریق اعمال تصمیم نتوانسته‌ایم گروه خود را فعلیت دهیم، استفاده از آن دارای اشکال می‌شود.

ممکن است تصور شود توصیه استدلال بالا این است که اساساً باید از استفاده از سیستم‌های هوشمند برای تصمیم‌سازی دست کشید. در پاسخ، نیثن کولانر بیان می‌کند که این امر، لازمه

استدلال بالا نیست. ما می‌توانیم آنچه ارسطو بیان کرده است را جدی تلقی کنیم ولی درعین حال بر این عقیده باشیم که جایی برای تصمیم‌سازی هوش مصنوعی نیز وجود دارد. در این صورت می‌توان انواع مختلف تصمیم‌سازی هوش مصنوعی را از یکدیگر متمایز کرد:

- چه تصمیماتی را می‌توان به صورت کامل به هوش مصنوعی سپرد؟ (برای مثال، توصیه‌هایی برای خرید)
 - چه تصمیماتی می‌باید حاصل تعامل میان انسان و هوش مصنوعی باشد؟ (برای مثال، تصمیم برای استخدام افراد)
 - چه تصمیماتی باید توسط هوش مصنوعی اتخاذ شود اما با این امکان که انسان بتواند آن را نپذیرد؟ (برای مثال، درخواست دادن برای کارت اعتباری یا وام)
 - چه تصمیماتی باید منحصراً در حیطه اختیارات انسان باشد نه هوش مصنوعی؟ (برای مثال، انتخاب یک رئیس‌جمهور).
- نیشن کولانر بیان می‌کند که در این حالت، بحث بر سر پاسخ به سؤالات فوق خواهد بود. از نظر وی تصمیماتی از نوع دوم و سوم، با فرض هوش مصنوعی غیر تبیین‌پذیر نمی‌تواند محقق شود. از این رو، استدلال بالا می‌تواند کماکان در مورد تصمیمات از نوع دوم و سوم جاری باشد و بتوان نتیجه گرفت که هوش مصنوعی تبیین‌ناپذیر مانعی خواهد بود برای تصمیمات از نوع دوم و سوم.

بخش چهارم



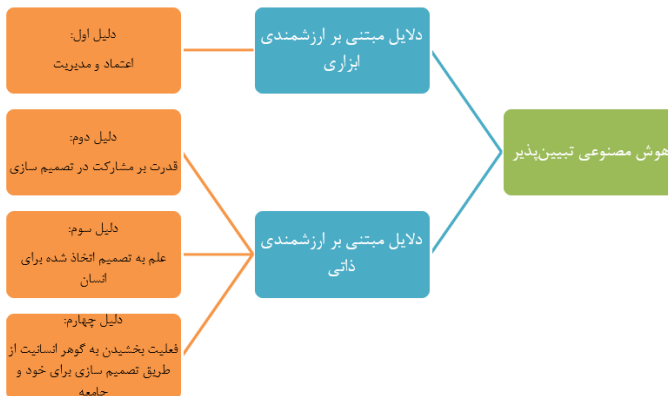
■ ملاحظات پیرامون هوش مصنوعی تبیین پذیر

همان‌طور که در بخش‌های قبل ملاحظه کردیم، اساساً چهار دلیل عمده برای ارزشمندی هوش مصنوعی تبیین‌پذیر ارائه شده است که خود به دو بخش دلایل ابزاری و دلایل ذاتی تقسیم می‌شوند. بر اساس دلایل ابزاری، هوش مصنوعی تبیین‌پذیر ارزشمند است زیرا وسیله‌ای است برای دستیابی به ارزش‌های دیگری همچون اعتماد. در واقع، هوش مصنوعی تبیین‌پذیر به ما کمک می‌کند که بتوانیم به نحو معقولی به هوش مصنوعی اعتماد کرده و از آن استفاده کنیم و در صورت بروز اشکال و خطر، بتوانیم دریابیم که اشکال کار کجاست. بنابراین، بر اساس این نگاه استفاده از هوش مصنوعی تبیین‌پذیر ارزشمند است اما نه به‌خودی‌خود، بلکه به این سبب که منافعی را برای ما به دنبال می‌آورد. بر اساس دلایل ذاتی، هوش مصنوعی تبیین‌پذیر، صرف‌نظر از اینکه چه منافع و فوایدی را برای ما به دنبال می‌آورد امری ارزشمند است. این دلایل خود از سه طریق این مطلب را اثبات می‌کنند.

بنابراین، دلایل دوم تا چهارم به‌قرار زیرند: دلیل دوم اینکه، هوش مصنوعی تبیین‌پذیر ارزشمند است زیرا قابلیت مشارکت

در تصمیم‌سازی را میسر می‌کند (حتی در مواردی که نتایج به‌دست‌آمده از یک سیستم هوشمند همگی اخلاقی، قانونی و منطبق بر منافع و تمایلات ما باشد)؛ دلیل سوم اینکه، هوش مصنوعی تبیین‌پذیر ارزشمند است زیرا سبب خواهد شد که ما به اینکه چرا فلان تصمیم درباره‌مان گرفته شده است علم پیدا کنیم (حتی در مواردی که نتایج به‌دست‌آمده از یک سیستم هوشمند همگی اخلاقی، قانونی و منطبق بر منافع و تمایلات ما باشد) و دلیل چهارم اینکه، هوش مصنوعی تبیین‌پذیر ارزشمند است زیرا سبب خواهد شد که ما انسان‌ها، با اتخاذ تصمیم برای خود و جامعه، خوشتن واقعی خود را فعلیت بخشیم. اما به نظر می‌رسد در میان دلایل ذکر شده برای ارزشمندی هوش مصنوعی تنها دلیل اول قابل قبول بوده و سه دلیل دیگر قابل دفاع نباشد.

در ادامه در این باره توضیح می‌دهیم.



همان‌طور که دیدیم، دلایل دوم تا چهارم، برخلاف دلیل اول، دلایل ذاتی هستند و نه دلایل ابزاری. اجازه دهید دلیل دوم را بررسی کنیم. بر اساس دلیل دوم هوش مصنوعی تبیین‌پذیر ارزشمند است زیرا قابلیت مشارکت در تصمیم‌سازی را میسر می‌کند.

همان‌طور که بیان کردیم این دلیل با نتایجی که یک هوش مصنوعی برای ما به بار می‌آورد کاری ندارد، نتایج هر چه می‌خواهد باشد (به سود یا به ضرر ما)، آنچه ارزشمند است این است که ما انسان‌ها بتوانیم در فرایند تصمیم‌سازی مشارکت کنیم. در واقع، نفس مشارکت امری مهم و ارزشمند است. سپس این دلیل بیان می‌کرد از آنجاکه در یادگیری عمیق چنین چیزی میسر نبوده ولی در هوش مصنوعی تبیین‌پذیر میسر است، بنابراین هوش مصنوعی تبیین‌پذیر امری ارزشمند است. اما به نظر می‌رسد در اینجا اشکالی وجود داشته باشد. اشکال این است که همین عدم مشارکت در تصمیم‌سازی، در حوزه‌های انسانی نیز به چشم می‌خورد.

در واقع، در بسیاری از مواقع، کارکنان یک شرکت، کارگران یک کارخانه، شهروندان یک شهر و مردمان یک کشور هیچ مشارکتی در تصمیم‌سازی‌های خرد و کلان مسئولین ندارند. حال می‌توان گفت، اگر استفاده از هوش مصنوعی غیر تبیین‌پذیر صرفاً به این خاطر که نمی‌توان در امر تصمیم‌سازی مشارکت نمود مذموم و دارای اشکال است، آنگاه باید این را نیز پذیرفت که اصل تصمیم‌سازی مسئولین برای زیر دستان خود (در سطوح مختلف) بدون مشارکت آن‌ها نیز امری مذموم و دارای اشکال است.

اما به نظر می‌رسد که پذیرفتن مورد اخیر این لازمه را خواهد

داشت که همه افراد در همه تصمیمات مشارکت کنند که به لحاظ عملی امری است غیرممکن.

البته ممکن است کسی بیان کند در نظام‌های دموکراتیک این مردم هستند که مسئولین را انتخاب کرده و تصمیم‌سازی در امور کشور را به او تفویض می‌کنند. از این‌رو، به نحو غیرمستقیم مردم در تصمیم‌سازی‌ها مشارکت کرده‌اند.

اما به نظر می‌رسد پاسخ بالا چندان موفقیت‌آمیز نباشد. زیرا فرض کنید که درباره مسئله‌ای برگزیدگان و نمایندگان مردم تصمیم الف را اتخاذ می‌کنند اما اکثریت مردمی که به آن نمایندگان رأی داده‌اند با اتخاذ تصمیم الف مخالف‌اند. به نظر می‌رسد در این موارد می‌توان به نحو معقولی مدعی شد که مردم (اگرچه خود نمایندگان را انتخاب کرده‌اند) در تصمیم‌سازی مشارکت نکرده‌اند؛ چراکه اکثریت آن‌ها نظری خلاف مسئولین خود دارند.

حال اجازه دهید به دلیل سوم بپردازیم. بر اساس دلیل سوم، اصل علم به اینکه چگونه هوش مصنوعی تصمیمی را برای ما اتخاذ کرده است امری ارزشمند است؛ امری که هوش مصنوعی جعبه سیاه به دلیل نامعلوم بودن فرایند تصمیم‌سازی، تأمین نمی‌کند. اما به نظر می‌رسد در اینجا نیز مشکلی وجود داشته باشد.

فارغ از اینکه هوش مصنوعی تبیین‌پذیر به لحاظ تکنولوژیکی و فنی و مهندسی قابل دستیابی باشد یا نه، مسئله این است که اگر تصمیم‌سازی هوش مصنوعی غیر تبیین‌پذیر امری است مشکل‌دار (آن‌گونه که دلیل سوم بیان می‌دارد)، آنگاه تصمیم‌سازی انسان‌ها نیز امری است مشکل‌دار. اجازه دهید توضیح دهیم.

ما در زندگی خود تصمیمات مختلفی را از تصمیمات ساده تا تصمیمات پیچیده اتخاذ می‌کنیم. برای مثال، فرض کنید تصمیم می‌گیریم تا به سمت یخچال رفته و لیوان آبی را بنوشیم. باتوجه به اینکه مغز انسان یک مجموعه عظیم نورونی است (همان‌گونه که سیستم‌های یادگیری عمیق ملهم از مغز انسانی هستند)، این سؤال مطرح است که آیا ما واقعاً تبیینی از این مسئله داریم که چه فرایند نورونی در مغزمان رخ داده است که در نتیجه آن ما تصمیم به نوشیدن آب گرفته‌ایم؟ به نظر می‌رسد جواب این سؤال خیر باشد. ما هیچ اطلاعی از اینکه چه فرایندهایی در مغزمان رخ می‌دهد نداریم. در واقع می‌توان گفت مغز ما همانند یک جعبه سیاه است که فرایند تصمیم‌سازی در آن برای ما نامعلوم است.

بنابراین، به نظر می‌رسد در تصمیمات انسانی نیز تبیین‌ناپذیری وجود دارد. حال می‌توان از مدافعان دلیل سوم پرسید چرا استفاده از تصمیمات سیستم هوشمند تبیین‌ناپذیر امری مشکل‌دار بوده ولی استفاده از تصمیمات انسان‌ها (که در آن نیز تبیین‌ناپذیری وجود دارد) بلا اشکال است. در واقع، اگر بپذیریم تصمیم‌سازی توسط سیستم هوشمند جعبه سیاه مشکل‌دار است، آنگاه باید این را نیز بپذیریم که تصمیم‌سازی مغز انسان‌ها نیز مشکل‌دار است. اما به نظر می‌رسد مطلب اخیر نپذیرفتنی نباشد.

حال اجازه دهید به بررسی دلیل چهارم بپردازیم. بر اساس دلیل چهارم هوش مصنوعی تبیین‌پذیر ارزشمند است زیرا سبب خواهد شد که ما انسان‌ها، با اتخاذ تصمیم برای خود و جامعه، خویشتن واقعی خود را فعلیت بخشیم. اما به نظر می‌رسد در اینجا

نیز مشکلی وجود داشته باشد. این دلیل بر این تأکید دارد که اصل تصمیم‌سازی سبب بالفعل شدن حقیقت و گوهر انسان خواهد شد. سپس بیان می‌کند که این امر در هوش مصنوعی تبیین‌ناپذیر محقق نیست. اما به نظر می‌رسد به خلاف ادعای دلیل چهارم، این امر در هوش مصنوعی تبیین‌ناپذیر محقق است. درست است که ما نمی‌توانیم در فرایند تصمیم‌سازی یک سیستم هوشمند تبیین‌ناپذیر دخالتی کنیم ولی در اصل استفاده یا عدم استفاده از سیستم هوشمند تصمیمی را اتخاذ کرده‌ایم. در واقع، ما پیش از آنکه سیستم هوشمند تصمیمی را برای ما یا جامعه ما اتخاذ کند، تصمیم گرفته‌ایم که از سیستم هوشمند تبیین‌ناپذیر استفاده کنیم. بنابراین، به نظر می‌رسد که حتی در سیستم هوشمند تبیین‌ناپذیر ملاحظه مدنظر دلیل چهارم (یعنی اعمال تصمیم‌سازی توسط انسان به‌منظور فعلیت بخشیدن به خود و جامعه) تأمین خواهد شد.

از نظر نگارنده، شاید بهترین دلیل برای استفاده از سیستم‌های هوشمند تبیین‌پذیر، همان دلیل اول باشد. در واقع، شاید بهترین دلیل این باشد که استفاده از هوش مصنوعی تبیین‌پذیر ارزشمند است زیرا سبب خواهد شد بتوانیم اشکالات موجود در سیستم‌های هوشمند را کاهش داده و بتوانیم اساساً برای استفاده از سیستم‌های هوشمند به آن‌ها اعتماد کنیم.

همچنین، حتی ممکن است بتوان دلایل را ارائه کرد که نشان دهند استفاده از سیستم‌های هوشمند تبیین‌ناپذیر، علی‌رغم وجود مسئله جعبه سیاه، امری بی‌مشکل است.

اجازه دهید یکی از این دلایل را ارائه کنیم. ایده اصلی این دلیل

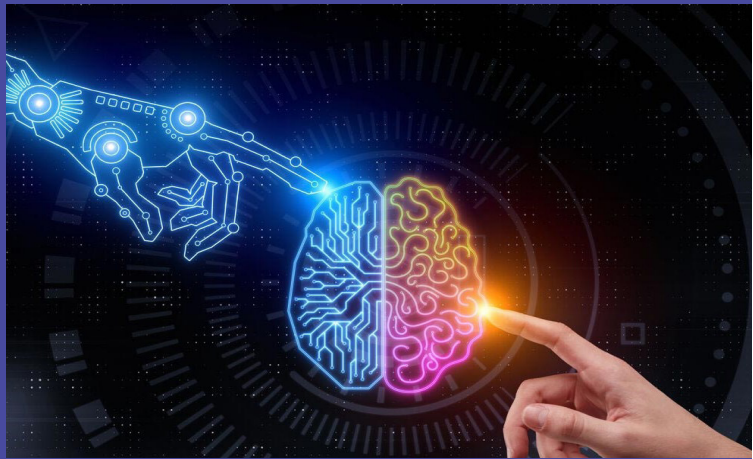
این است که همان‌طور که تصمیمات اشتباه و خطا در میان انسان‌ها امری پذیرفته شده است، اشتباهات و خطاهای سیستم‌های هوشمند نیز می‌باید امری پذیرفته شده باشد.

توضیح آنکه، انسان‌ها به‌صورت فردی و یا به‌صورت جمعی بعضاً تصمیمات اشتباهی اتخاذ می‌کنند. اما مسئله این است که این تصمیمات اشتباه و خطا بعضاً اتفاق می‌افتد. حال اگر یک سیستم هوشمند گاهی اوقات دچار تصمیمات اشتباه و خطا شود، به نحو معقولی می‌توان آن را همانند خطاهای انسانی در نظر گرفت. به بیان دیگر، در صورتی که تصمیمات اشتباه یک سیستم هوشمند بسیار کمتر از تصمیمات درست آن باشد، آنگاه می‌توان به نحو معقولی بر روی آن اعتماد کرده و از آن استفاده کرد. این ایده برخاسته از نظریه اتکاگرایی^۱ در معرفت‌شناسی است.

بر اساس نظریه اتکاگرایی در باب توجیه، اگر باورهای تولید شده توسط یک قوای شناختی در اکثر مواقع درست باشند، آنگاه آن قوه شناختی و به‌تبع باورهای تولید شده توسط آن معتبر بوده و می‌توان به آن اتکا کرد (Lemos, 2007, p.85). به نظر می‌رسد می‌توان همین ایده را نیز در مورد سیستم‌های هوشمند تبیین‌ناپذیر نیز به کار برد. اگر یک سیستم هوشمند در اکثر مواقع (مثلاً ۹۰ درصد از مواقع) تصمیمات درستی اتخاذ کند و تنها در مواقع اندکی (مثلاً ۱۰ درصد مواقع) تصمیمات اشتباهی اتخاذ کند، آنگاه این سیستم قابل‌اتکا بوده و می‌توان به آن اعتماد کرد. ولی اگر برائز آزمایشات مشخص شد که یک سیستم هوشمند در اکثر مواقع تصمیمات اشتباهی اتخاذ می‌کند سیستم مذکور قابل‌اتکا نبوده و نمی‌توان از

آن استفاده کرد. بنابراین، به نظر می‌رسد حتی بتوان استدلال کرد که استفاده از سیستم‌های غیر تبیین‌پذیر با وجود مسئله جعبه سیاه امری بی‌مشکل است.

بخش پنجم



بر اساس آنچه در بالا بیان شد، نگارنده پیشنهادهای راهبردی زیر را برای استفاده مدیران و تصمیم‌سازان عرصه هوش مصنوعی جمهوری اسلامی ایران ارائه می‌کند:

همان‌طور که دیدیم، مسئله هوش مصنوعی تبیین‌پذیر مسئله نوپایی در عرصه هوش مصنوعی است که هنوز ابهامات زیادی در آن وجود دارد. پیشنهاد می‌شود بودجه‌های تحقیقاتی‌ای برای کاوش ابعاد مختلف مسئله اختصاص داده شود تا بتوان برای سؤالات زیر پاسخ‌های قابل قبولی یافت:

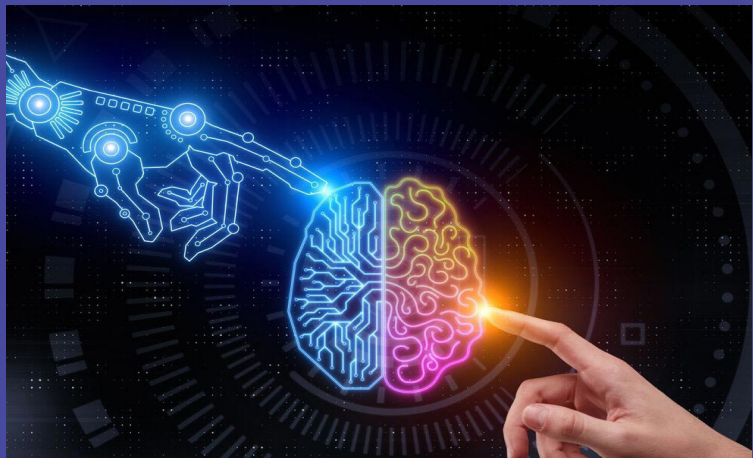
«آیا دستیابی به هوش مصنوعی تبیین‌پذیر ضرورت دارد؟»، «آیا دلایل معقولی برای حرکت به سمت هوش مصنوعی تبیین‌پذیر وجود دارد؟» و «آیا دستیابی به هوش مصنوعی تبیین‌پذیر به لحاظ تکنولوژیکی و فنی و مهندسی اساساً ممکن است؟».

در گام بعد و پس از مشخص‌شدن پاسخ سؤالات فوق نیاز است تا بر اساس پاسخ‌های به‌دست‌آمده در گام قبلی، تصمیمات درستی اتخاذ شود.

برای مثال، اگر به این نتیجه رسیدیم که دستیابی به هوش

مصنوعی تبیین‌پذیر ضرورت داشته و امری ممکن است، آنگاه باید سیاست‌های راهبردی برای دستیابی به آن اتخاذ شود.

جمع بندی



در این گزارش در ابتدا توضیحی درباره انواع مختلف سیستم‌های هوشمند ارائه شد: سیستم‌های هوشمند خبره و سیستم‌های هوشمند مبتنی بر یادگیری عمیق. ویژگی سیستم‌های خبره این است که می‌توان فرایند اینکه سیستم چگونه به نتیجه خاصی رسیده است را بازبینی نمود. یعنی، این فرایند پوشیده نبوده و تبیین‌پذیر است. اما در مقابل در سیستم‌های مبتنی بر یادگیری عمیق، فرایند تصمیم‌سازی توسط سیستم هوشمند حتی برای طراحان نیز معلوم نبوده و به یک بیان، فرایند تصمیم‌سازی در این سیستم‌ها مبهم است.

پس از بیان این مطلب، مسئله جعبه سیاه را شرح دادیم. مسئله جعبه سیاه به همین مسئله عدم تبیین‌پذیری و نامشخص بودن فرایند رسیدن به یک خروجی در سیستم‌های مبتنی بر یادگیری عمیق اشاره دارد. پس از توضیحاتی در باب مسئله جعبه سیاه، بیان کردیم که یکی از رویکردها برای در پاسخ به مسئله جعبه سیاه، رویکردی است که به «هوش مصنوعی تبیین‌پذیر» معروف است. پس از تمایز میان این سؤال که «چگونه می‌توان از نقطه‌نظر تکنولوژیکی، فنی و مهندسی به هوش مصنوعی تبیین‌پذیر دست‌یافت؟» و این سؤال که «چرا اساساً هوش مصنوعی تبیین‌پذیر امری ارزشمند

است؟»، سؤال اول را فرو گذاشته و به بیان پاسخ‌های ارائه شده برای سؤال دوم پرداختیم. در ادبیات بحث چهار پاسخ به این مسئله وجود دارد. این چهار دلیل خود به دو بخش دلایل ابزاری و دلایل ذاتی تقسیم می‌شوند.

بر اساس دلایل ابزاری، هوش مصنوعی تبیین‌پذیر ارزشمند است زیرا وسیله و ابزاری است برای دستیابی به ارزش‌های دیگری همچون اعتماد. در واقع، هوش مصنوعی تبیین‌پذیر به ما کمک می‌کند که بتوانیم به نحو معقولی به هوش مصنوعی اعتماد کرده و از آن استفاده کنیم و در صورت بروز اشکال و خطر، بتوانیم دریابیم که اشکال کار کجاست. بنابراین، بر اساس این نگاه استفاده از هوش مصنوعی تبیین‌پذیر ارزشمند است اما نه به‌خودی‌خود، بلکه به این سبب که منفعی را (مثلاً «اعتماد» به آن را)، برای ما به دنبال می‌آورد. بر اساس دلایل ذاتی، هوش مصنوعی تبیین‌پذیر، صرف‌نظر از اینکه چه منافع و فوایدی را برای ما به دنبال می‌آورد امری ارزشمند است.

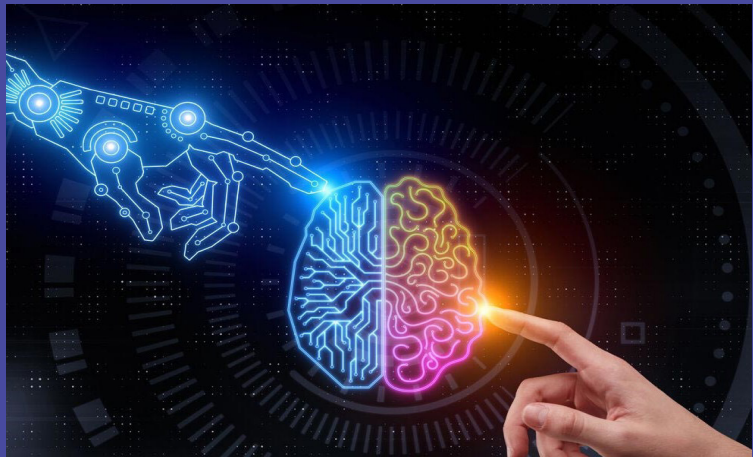
این دلایل خود از سه طریق این مطلب را اثبات می‌کنند. بنابراین، دلایل دوم تا چهارم به‌قرار زیرند: دلیل دوم اینکه، هوش مصنوعی تبیین‌پذیر ارزشمند است زیرا قابلیت مشارکت در تصمیم‌سازی را میسر می‌کند (حتی در مواردی که نتایج به‌دست‌آمده از یک سیستم هوشمند همگی اخلاقی، قانونی و منطبق بر منافع و تمایلات ما باشد)؛ دلیل سوم اینکه، هوش مصنوعی تبیین‌پذیر ارزشمند است زیرا سبب خواهد شد که ما به اینکه چرا فلان تصمیم درباره‌مان گرفته شده است علم پیدا کنیم (حتی در مواردی که نتایج به‌دست‌آمده از یک سیستم هوشمند همگی اخلاقی، قانونی و منطبق بر منافع و تمایلات ما باشد) و دلیل چهارم اینکه، هوش مصنوعی تبیین‌پذیر ارزشمند است زیرا سبب خواهد شد که ما

انسان‌ها، با اتخاذ تصمیم برای خود و جامعه، خویشتن واقعی خود را فعلیت بخشیم. در بخش بعدی به برخی ملاحظات در باب هوش مصنوعی تبیین‌پذیر اشاره کردیم.

یکی از ملاحظات بر این نکته تأکید داشت که دلایل دوم تا چهارم قابل چالش بوده و شاید بهترین دلیل برای دستیابی و استفاده از هوش مصنوعی تبیین‌پذیر، همان دلیل اول باشد. ملاحظه دیگر، این بود که حتی می‌توان استدلال کرد که استفاده از هوش مصنوعی غیر تبیین‌پذیر امری بلا مشکل است. زیرا، همان‌طور که این مسئله که انسان‌ها در تصمیم‌سازی دچار خطا می‌شوند مسئله‌ای پذیرفته شده است، خطا و اشتباهات هوش مصنوعی قابل اعتماد نیز می‌توان پذیرفت.

در پایان پیشنهادهاتی را برای استفاده مدیران و تصمیم‌سازان عرصه هوش مصنوعی جمهوری اسلامی ایران ارائه کردیم. یکی از این پیشنهادها راهبردی این بود که باتوجه به نوظهور بودن هوش مصنوعی تبیین‌پذیر و مشخص نبودن ابعاد مختلف آن، راه‌اندازی پروژه‌های تحقیق برای شناسایی ابعاد مختلف مسئله امری ضروری است.

منابع



[1] Aristotle J (1984) The complete works of Aristotle: the revised oxford translation. Edited by Jonathan Barnes. Princeton University Press, Princeton (Volume Two. Bollington Series LXXI.2)D

[2] Ayesha S, Hanif MK, Talib R (2020) Overview and comparative study of dimensionality reduction techniques for high dimensional data. Inf Fusion 58-59:44

[3] Buckner, Cameron and James Garson. (2019). «Connectionism», The Stanford Encyclopedia of Philosophy, Edward N. Zalta (ed.), URL = <<https://plato.stanford.edu/archives/fall2019/entries/connectionism/>>.

[4] Colaner, N. (2021). Is explainable artificial intelligence intrinsically valuable?. AI & SOCIETY, 8-1.

[5] DARPA (2016) Explainable Artificial Intelligence (XAI) Program, Broad Agency Announcement Explainable Artificial Intelligence (XAI) DARPA-BAA16-D

[6] Grover, Richa. (2019). Deep Learning - Overview, Practical Examples, Popular Algorithms. Retrived from:

<https://www.analyticssteps.com/blogs/deep-learning-overview-practical-examples-popular-algorithms>

[7] Guidotti R, Monreale A, Ruggieri S, Turini F, Giannotti F, Pedreschi D (2018) A survey of methods for explaining black box models. *ACM Comput Surv* 1:(5)51

[8] Kim B, Khanna R, Kovejo O (2016) Examples are not Enough, Learn to Criticize! Criticism for Interpretability. *Proceedings of Advances in Neural Information Processing Systems*, Barcelona, Spain, 29.

[9] Lemos, N. (2007). *An introduction to the theory of knowledge*. Cambridge University Press.

[10] Nath, Prasenjit. (2015). *AI & Expert System in Medical Field: A study by survey method*. AITHUN. Volume-I. 100.

[11] Peng, Tony. (10 .(2018 AI Failures. Retrived from: <https://medium.com/syncedreview/-2018in-review-10-ai-failures-c18faadf5983>

[12] Tamagnini P, Krause J, Dasgupta A, Bertini E (2017) Interpreting Black-Box Classifiers Using Instance-Level Visual Explanations. In *Proceedings of HILDA'17*, Chicago, IL, USA.

[13] Vellido A, Jos'e D Martin-Guerrero, J Lisboa (2012) Making machine learning models interpretable. *Proceedings of European Symposium on Artificial Neural Networks, Computational Intelligence and Machine Learning*. Bruges, Belgium.

[14] Wachter, S., Mittelstadt, B., & Floridi, L. (2017). Why a right to explanation of automated decision-making does not exist in the general data protection regulation. International Data Privacy Law, 99-76 ,(2)7.

[15] Young IM (1990) Justice and the politics of difference. Princeton University Press, PrincetonD

[16] Zednik, C. (2019). Solving the black box problem: a normative framework for explainable artificial intelligence. Philosophy & Technology, 24-1.

[17] Zhu YN, Sun L, Li C (2019) Inclusive Leadership, Proactive Personality and Employee Voice: A Voice Role Identity Perspective. Proceedings of the Academy of Management Annual Meeting.



مرکز ملی فضایی مجازی
پژوهشگاه فضایی مجازی

csri.majazi.ir

حوزه فضای مجازی به اندازه انقلاب اسلامی اهمیت دارد. این فضا مثل یک رودخانه پر از آب خروشان است که می آید و دائماً هم بر آب آن افزوده و خروشان تر می شود. اگر ما بر این رودخانه تدبیر کنیم و برنامه داشته باشیم، زه کشی کنیم و هدایت کنیم این رودخانه را تا به سد بریزد، می شود فرصت. اگر رهايش کنیم و برنامه ای برای آن نداشته باشیم می شود یک تهدید.



csri.ac.ir