

<http://doi.org/10.22133/mtlj.2025.485015.1386>

Basis of Civil Liability Arising from the Artificial Intelligence Performance

Ebrahim Shoarian Sattari^{1*} , Fatemeh Pourahmad² 

¹ Ph.D. in Private Law, Faculty of Law and Social Sciences, Tabriz University, Tabriz, Iran.

² MA. Student in Private Law, Faculty of Law and Social Sciences, Tabriz University, Tabriz, Iran.

Article Info	Abstract
Original Article	With the increasing progress of science and technology, especially in the field of artificial intelligence, it is necessary to respond to one of its most important challenges, i.e. the civil liability resulting from the operation of these systems, and provide appropriate solutions to achieve the two goals of "damage compensation" and "not creating obstacles in the way of the progress of science and technology" in parallel. In this article, with an interdisciplinary study and descriptive-analytical method, the necessity of granting legal personality to artificial intelligence has been proven, so that the damages caused to people by intelligent systems can be compensated, with a proper basis and considering the types of artificial intelligence systems and their level of risk. So that for systems with less risk, such as Google Translate, the basis of fault is considered, and for high-risk systems, such as self-driving cars, Metaverse, and humanoid robots, the basis of risk is considered. In addition to these two traditional bases, the basis of economic analysis has also been examined .After proving the legal personality of artificial intelligence systems, considering the existing laws and procedures regarding the damages caused by these systems, it is determined that each of the factors involved in this technology, from the designer and manufacturer to the consumer, has the responsibility and the extent of this responsibility can be evaluated with better quality regarding the evaluation of the supervisory group mentioned in the "artificial intelligence law" which was recently approved in the European Union. Based on the above-mentioned issues, proposals for legislation have also been presented.
Received: 22-10-2024	
Accepted: 14-04-2025	
Keywords: Artificial Intelligence, Basis of Civil Liability, Elements of Civil Liability, Self-Driving, Metaverse	
*Corresponding author e-mail: e_shoarian@tabrizu.ac.ir	
How to Cite: Shoarian Sattari, S., & Pourahmad, F. (2026). Basis of Civil Liability Arising from the Artificial Intelligence performance. <i>Modern Technologies Law</i> , 7(13), 299-327.	
Published by University of Science and Culture https://www.usc.ac.ir Online ISSN: 2783-3836	

1. Introduction

Humanity has gone through four industrial and technological revolutions and is currently in the fourth industrial revolution, which is referred to as artificial intelligence, which is based on the digital revolution. The origin of the science of artificial intelligence can be attributed to Alan Turing, a computer science genius and the discoverer of the decoding algorithm for the German cipher machine (called Enigma) during World War II, in 1950. Turing, the first person to express the concept of artificial intelligence, explained that machines can think like humans and designed a test called the "Turing test" This test is a method for measuring the level of intelligence of a machine, in which a person, as a judge, sits in conversation with a machine and a human and tries to distinguish the machine from the human. In this process, if the machine can make the judge make a mistake in his judgment, it has successfully passed the test. So far, there has been no single, precise definition of artificial intelligence that is agreed upon by scientists in this field, but most definitions are based on four beliefs: 1) systems that think rationally; 2) systems that act rationally; 3) systems that think like humans; and 4) systems that act like humans. There are four approaches to defining artificial intelligence: 1) human-like performance; 2) human-like thinking; 3) rational thinking; and 4) rational action. According to the performance of intelligent systems, artificial intelligence is classified into three categories: 1) Artificial Narrow Intelligence; 2) Artificial General Intelligence; and 3) Artificial Super Intelligence. Narrow AI systems focus only on limited tasks or a single problem. Narrow or weak AI does not support simultaneous tasks but performs its task with the highest level of accuracy. Also, this system cannot identify the reasons, but facilitates uniform work. Examples of this type of system include Google search and Google ranking, chatbots, Siri by Apple, and Alexa by Amazon. General or strong AI can mimic the cognitive abilities of the human brain and is actually a great solution for doing an unfamiliar task. Whenever there are problems and differences in decision-making, strong AI can perform well there and is capable of performing a variety of tasks. General or strong AI is like humans in thinking and reasoning and is flexible. Super AI doesn't exist yet, but it will be the result of the advancement and development of general AI. This type of AI is a hypothetical system that goes beyond the human brain. As artificial intelligence enters all areas, this type of system is manifested in any physical or non-physical object; for example, a robot is used as a servant, a cook, or in specialized fields that each have specific rules, such as medicine and judgment. In this research, while proving the necessity of granting legal personality in order to attribute liability to artificial intelligence and facilitate compensation, we seek to determine a suitable basis for this type of liability by studying the existing laws and procedures in this field. Therefore, the existing laws and judicial procedures in this area have been examined in two separate sections, under the headings of 1) the necessity of granting legal personality to artificial intelligence from the perspective of facilitating compensation for damages and 2) the basis of liability has been examined based on the theories of fault and risk and the theory of economic analysis.

2. Methodology

In this research, we tried to collect and gather information from primary sources (including laws, regulations, directives, bylaws, judicial decisions, etc.) and secondary sources (such as books and articles, theses, research reports, etc.) in scientific research data centers, official and reputable websites, libraries, and similar resources. From the authors' perspective, the title of the study suggests that first, the necessity of granting legal personality to artificial intelligence should be addressed, and second, the appropriate basis for attributing civil liability to artificial intelligence should be examined. The logical order of the study requires examining various laws and guidelines in this area, followed by suggestions for legislation to the legislator for the positive impact of the article and its applicability, which will both comply with legal principles and improve the speed and quality of the development of new technologies. Therefore, the studies are organized under two main headings. The first main heading includes "The necessity of granting legal personality to artificial intelligence from the perspective of

facilitating compensation for damage,” and the second main heading includes “The basis of civil liability of artificial intelligence”.

3. Results and Discussion

With the advancement of science and technology in the field of artificial intelligence, which today encompasses various aspects of people's lives, many legal challenges arise. Law, as a dynamic science, must have an appropriate response to these challenges based on the requirements of time and place. To attribute damages caused by intelligent systems to them, these types of systems must be granted legal personality. Globally, phenomena such as Indian temples and the Ecuadorian ecosystem have legal personality; therefore, artificial intelligence with these dimensions and influence in all aspects of social and personal life must have legal personality in the first place. In addition, an example given for the need to grant legal personality to artificial intelligence is the example of the legal personality of company, with the explanation that the personality of companies is not simply due to the system and the inanimate office, but rather due to its consequences in the legal system and society. In granting legal personality to these types of systems, it was suggested to the legislator that, considering the level of the system and the level of risk it poses to society, it should accept two types of joint and several liability for artificial intelligence systems, that is, like existing companies, they should have legal personality and different forms of liability (joint and several). To provide financial means and assets for the registered system, the involved parties who register open an account for the artificial intelligence so that a percentage of the system's revenues is deposited into this account, and these amounts are spent solely to compensate for possible losses and improve the quality of the system. It is also possible to take advantage of the distributive justice approach, and in cases where the cause of the loss cannot be identified and the damage compensated in any way, insurance companies and government support funds can play their role effectively. The main goal of granting legal personality to artificial intelligence is to facilitate proving damages and seeking compensation for them, and one of the important reasons is to increase precautions and monitoring and the quality of intelligent systems. As a result of examining the three theories of fault, risk theory, and economic analysis, and combining goals such as compensation for damage, increasing the quality of artificial intelligence systems, promoting social welfare, and advancing science and technology, each system must be examined separately from a technical and economic perspective to determine a basis that is appropriate for achieving all of the aforementioned goals and depending on the level of risk that the system manufacturers present in the instructions or product advertisements, it can be included as a fault basis or a risk basis. A relative basis is applied with technical and economic considerations. In recent years, the European Union has issued various guidelines to organize AI activities in various fields, such as self-driving cars, surgical robots, the metaverse, and other areas, but finally, in 2024, the European Union passed the “Artificial Intelligence Law” to develop a framework for the activities of AI systems. By establishing the supervisory group foreseen in this law, the duties of the factors involved in AI are assessed and the fault of each actor, from the designer and manufacturer to the user, is determined with greater precision. By studying the existing cases in this field, especially self-driving cars and surgical robots, it becomes clear that the roles and duties of each factor are examined in detail in the courts, and ultimately, it has been proven that so far, the advertising aspect has been very effective in determining the liability of artificial intelligence and goes beyond the actual performance of these products. The criticism leveled at the “Artificial Intelligence Law” and all existing laws and guidelines are that when drafting a law, matters such as defining and explaining existing terms and the persons involved in detail and determining their type of responsibility must be observed, and the drafting of laws must be carried out in a specialized manner in a specific field with the cooperation of experts.

4. Conclusions and Future Research

Overall, it seems that research on the various levels of AI and responsible persons needs interdisciplinary research in different fields to increase the transparency of the parties involved and the manner and extent of their responsibility.

Given that the legal personality of AI is a necessity for the flow of this technology in the present era and was also proven in the present study, the civil liability of the AI system itself and the persons involved in it was also proven.

Given the rate and speed of progress of AI in various subjects, it is predicted that legal studies in various scientific fields in which AI is integrated will increase, and national and international legislation will also increase in this area.





مبنای مسئولیت مدنی ناشی از عملکرد هوش مصنوعی

ابراهیم شعاریان ستاری^{*۱} (ID)، فاطمه پوراحمد^۲ (ID)

^۱ استاد، دانشکده حقوق و علوم اجتماعی، دانشگاه تبریز، تبریز، ایران.

^۲ دانش‌آموخته کارشناسی ارشد، دانشکده حقوق و علوم اجتماعی، دانشگاه تبریز، تبریز، ایران.

اطلاعات مقاله	چکیده
مقاله پژوهشی	
تاریخ دریافت: ۱۴۰۳/۰۸/۰۱	
تاریخ پذیرش: ۱۴۰۴/۰۱/۲۵	
واژگان کلیدی: هوش مصنوعی مبنای مسئولیت مدنی ارکان مسئولیت مدنی خودران متاورس	با پیشرفت روزافزون علم و فناوری، به‌ویژه در حوزه هوش مصنوعی، باید به یکی از مهم‌ترین چالش‌های آن، یعنی مسئولیت مدنی ناشی از عملکرد این سیستم‌ها، پاسخ داد و راهکارهای مناسبی ارائه کرد تا بتوان دو هدف «جبران خسارت» و «عدم ایجاد مانع در راه پیشرفت علم و فناوری» را به موازات پیش برد. در این مقاله، با مطالعه میان‌رشته‌ای و روش توصیفی - تحلیلی، لزوم اعطای شخصیت حقوقی به هوش مصنوعی اثبات شده است تا خسارات وارد شده به اشخاص به‌واسطه سیستم‌های هوشمند، با مبنایی مناسب و با در نظر گرفتن انواع سیستم‌های هوش مصنوعی و میزان خطر آن‌ها، جبران‌پذیر باشد؛ به‌طوری‌که برای سیستم‌های با خطر کمتر، مانند گوگل ترجمه، مبنای تقصیر و برای سیستم‌های پرخطر، مثل اتومبیل خودران، متاورس و ربات‌های انسان‌نما، مبنای خطر در نظر گرفته شده است. علاوه بر این دو مبنای سنتی، مبنای تحلیل اقتصادی نیز بررسی شده است. پس از اثبات شخصیت حقوقی سیستم‌های هوش مصنوعی، با مطالعه قوانین و رویه‌های موجود درباره خسارت‌های ایجاد شده از طریق این سیستم‌ها، مشخص می‌شود که هریک از عوامل دخیل در این فناوری، از طراح و تولیدکننده تا مصرف‌کننده، مسئولیت دارند و میزان این مسئولیت، با توجه به ارزیابی گروه ناظر ذکر شده در «قانون هوش مصنوعی»، که اخیراً در اتحادیه اروپا تصویب شده است، با کیفیت بهتری ارزیابی می‌شود. برحسب مباحث ذکر شده، پیشنهاداتی نیز در راستای قانون‌گذاری ارائه شده است.
*نویسنده مسئول رایانامه: e_shoarian@tabrizu.ac.ir	
نحوه استناددهی: شعاریان ستاری، ابراهیم، و پوراحمد، فاطمه (۱۴۰۵). مبنای مسئولیت مدنی ناشی از عملکرد هوش مصنوعی. حقوق فناوری‌های نوین، ۷(۱۳)، ۲۹۹-۳۲۷.	
ناشر: دانشگاه علم و فرهنگ https://www.usc.ac.ir شاپای الکترونیکی: ۲۷۸۳-۳۸۳۶	

مقدمه

بشر تاکنون به چهار دوره انقلاب صنعتی و فناوری دست یافته است و هم‌اکنون در انقلاب صنعتی چهارم قرار دارد که از آن با عنوان «هوش مصنوعی»^۱ یاد می‌شود و مبتنی بر انقلاب دیجیتال است. مبدأ علم هوش مصنوعی را می‌توان در سال ۱۹۵۰ به آلن تورینگ، نابغه علم رایانه و کاشف الگوریتم رمزگشایی از ماشین رمزگذار آلمانی‌ها (به نام اینگما) در جنگ جهانی دوم، منسوب کرد. تورینگ، اولین شخص بیان‌کننده مفهوم هوش مصنوعی، تشریح کرد که ماشین‌ها می‌توانند مانند انسان فکر کنند و آزمونی را تحت عنوان «آزمون تورینگ» طراحی کرد. این آزمون روشی برای سنجیدن میزان هوشمندی ماشین است که در جریان آن، یک شخص در مقام قاضی با یک ماشین و یک انسان به گفت‌وگو می‌نشیند و سعی می‌کند ماشین را از انسان تشخیص دهد. در این فرایند اگر ماشین بتواند قاضی را در داوری خود دچار اشتباه کند، توانسته است آزمون را با موفقیت پشت‌سر بگذارد (Abouzari, 2022, p. 26).

تاکنون تعریف واحد و دقیقی برای هوش مصنوعی ارائه نشده است که دانشمندان این علم بر آن توافق داشته باشد؛ اما بیشتر تعریف‌ها بر چهار باور استوارند: (۱) سیستم‌هایی که به‌طور منطقی فکر می‌کنند؛ (۲) سیستم‌هایی که به‌طور منطقی عمل می‌کنند؛ (۳) سیستم‌هایی که مانند انسان فکر می‌کنند؛ (۴) سیستم‌هایی که مانند انسان عمل می‌کنند (Giovanni & Branting, 2020, p. 7). درواقع، چهار رویکرد در تعریف هوش مصنوعی وجود دارد که عبارت‌اند از: (۱) عملکرد انسان‌گونه؛ (۲) تفکر انسان‌گونه؛ (۳) تفکر عقلایی؛ (۴) کنش عقلایی (Russell & Norving, 2020, pp. 14-18). با توجه به عملکرد سیستم‌های هوشمند، هوش مصنوعی به سه دسته طبقه‌بندی می‌شود: (۱) هوش مصنوعی محدود؛ (۲) هوش مصنوعی عمومی؛ (۳) ابر هوش مصنوعی.^۴

سیستم‌های هوش مصنوعی محدود فقط روی کارهای محدود با مشکل تمرکز می‌کند. هوش مصنوعی محدود با ضعیف از وظایف هم‌زمان پشتیبانی نمی‌کند؛ اما وظیفه خود را با بالاترین سطح دقت انجام می‌دهد. همچنین، این سیستم نمی‌تواند دلایل را شناسایی کند؛ اما موجب تسهیل کار یکنواخت می‌شود. ازجمله نمونه‌های موجود از این نوع سیستم، می‌توان به جست‌وجوی گوگل و رتبه‌بندی گوگل، چت‌بات، سیری توسط اپل و الکسا توسط آمازون اشاره کرد.

هوش مصنوعی عمومی یا قوی توانایی‌های شناختی مغز انسان را تقلید می‌کند و در واقع، راه‌حلی عالی برای انجام دادن کارهای ناآشناست. هر زمان که در تصمیم‌گیری مشکلات و تفاوت‌هایی وجود داشته باشد، هوش مصنوعی قوی به‌خوبی در آنجا عمل می‌کند و قادر به انجام انواع وظایف است. هوش مصنوعی عمومی یا قوی در تفکر و استدلال، مانند انسان و انعطاف‌پذیر است.

ابر هوش مصنوعی هنوز وجود ندارد؛ اما نتیجه پیشرفت و توسعه هوش مصنوعی عمومی خواهد بود. این نوع هوش مصنوعی یک سیستم فرضی فراتر از مغز انسان است (Khan, 2021, pp. 1-4). با توجه به ورود هوش مصنوعی به همه عرصه‌ها، این نوع سیستم در هر شیء فیزیکی یا غیر آن متجلی می‌شود. برای مثال، ربانی به‌عنوان خدمتکار، آشپز یا در حوزه‌های تخصصی، که هرکدام احکام خاصی دارد، نظیر پزشکی و قضاوت، به کار گرفته می‌شود (Rajabi, 2019, p. 450).

1. Artificial Intelligence
2. Artificial Narrow Intelligence (ANI)
3. Artificial General Intelligence (AGI)
4. Artificial Super Intelligence (ASI)

در این پژوهش، ضمن اثبات لزوم اعطای شخصیت حقوقی در راستای انتساب مسئولیت به هوش مصنوعی و تسهیل جبران خسارت، درصدد هستیم با مطالعه قوانین و رویه‌های موجود در این حوزه، مبنایی مناسب برای این نوع مسئولیت تعیین کنیم؛ از همین رو، قوانین و رویه‌های قضایی موجود در این حوزه در دو بخش مجزا، ذیل عناوین «ضرورت اعطای شخصیت حقوقی به هوش مصنوعی از منظر تسهیل جبران خسارت» و «مبنای مسئولیت براساس نظریات تقصیر و خطر و تحلیل اقتصادی» بررسی شده‌اند.

۱. ضرورت اعطای شخصیت حقوقی به هوش مصنوعی از منظر تسهیل جبران خسارت

انتساب مسئولیت به یک موجود در جامعه نیازمند اعطای شخصیت حقوقی به آن پدیده است. شخصیت حقوقی هنگامی پدید می‌آید که دسته‌ای از افراد که منافع و فعالیت مشترک دارند یا پاره‌ای از اموال، که به اهداف خاصی اختصاص یافته‌اند، در کنار هم قرار گیرند و قانون آن‌ها را طرف حق و تکلیف بشناسد و برایشان شخصیت مستقلی قائل شود (Safai & Qasemzadeh, 2017, p. 164). حال سؤال این است که آیا هوش مصنوعی دارای شخصیت حقوقی است که حقوق و تکالیف داشته باشد؟

اولین فردی که اعطای شخصیت حقوقی برای هوش مصنوعی را مطرح کرد لارنس بی سولوم در سال ۱۹۹۲ بود (Solum, 1992, p. 1238). در سال ۲۰۱۷ عربستان سعودی برای اولین بار به ربات انسان‌نمایی به اسم «سوفیا» تابعیت اعطا کرد (arabnews, 2017). در همان سال، هوش مصنوعی با شخصیت پسری با نام «شیبویا میرای» اولین ربات هوش مصنوعی جهان بود که اقامت توکیو را دریافت کرد (Tokyo: "Artificial intelligence 'boy'...", 2017). همچنین در همان سال، پارلمان اروپا به موجب قطعنامه‌ای در مورد قوانین مدنی رباتیک، اصطلاح «شخصیت الکترونیکی» را مطرح کرد (European Parliament, 2017)؛ اما این ایده را تعدادی از اندیشمندان، از نظر حقوقی و اخلاقی، به شدت نقد کردند (Open Letter to the European Commission Artificial Intelligence and Robotics, n.d). در سال ۲۰۱۹ گروه تخصصی مسئولیت و فناوری‌های جدید منصوب کمیسیون اروپا نیز موضع مشابهی در قبال این قطعنامه اتخاذ کرد. در سال ۲۰۲۰ دپارتمان سیاست‌گذاری درباره حقوق شهروندی و امور قانون اساسی اداره کل سیاست‌های داخلی^۱ اتحادیه اروپا به درخواست کمیته امور حقوقی مطالعه‌ای انجام داد که براساس آن، استدلال‌های ارائه شده در تقویت انتقادات به قطعنامه مذکور ناکافی قلمداد و اذعان شد که این انتقادات بی‌توجه به فناوری‌های موجود و در حال توسعه مطرح شده است؛ یکی از استدلال‌هایی که در این مطالعه برای لزوم اعطای شخصیت حقوقی به هوش مصنوعی مطرح شده، مقایسه شخصیت حقوقی شرکت‌ها با فناوری‌های نوین است (European Parliament, 2020). در این راستا، پیشنهادی برای قانون‌گذار بدین شرح ارائه می‌شود: برای اعطای یک شخصیت حقوقی متناسب با سطح عملکردی سیستم هوش مصنوعی لازم است قانون‌گذار، همانند شرکت‌ها، مسئولیت هوش مصنوعی را برحسب میزان مسئولیت خود و سازندگان و طراحان مشخص کند. برای مثال، دو نوع مسئولیت تضامنی و نسبی را برای سیستم‌های هوش مصنوعی تعریف کند؛ با این تفاوت که در شرکت‌ها مؤسسان در ابتدا و اعضا در طول فعالیت شرکت طی فرایند قانونی نوع شرکت را مشخص می‌کنند؛ اما راجع به هوش مصنوعی، قانون‌گذار چهارچوبی را به این منظور تعیین می‌کند. برای مثال هنگامی که یک سیستم هوشمند خطرات متعددی برای جامعه دارد یا با جان انسان رابطه مستقیم دارد، سازندگان و عوامل دخیل در این نوع سیستم را به صورت تضامنی در جبران خسارات وارده مسئول می‌داند و این نوع سیستم را به عنوان «سیستم هوش مصنوعی با مسئولیت تضامنی» ثبت می‌کند؛ یعنی هنگامی که این سیستم موجب ورود خساراتی شود، سیستم و سازندگان سیستم به طور تضامنی در جبران خسارت مسئولیت دارند.

در صورتی که اشخاص حقیقی به طراحی و ساخت سیستم هوش مصنوعی اقدام کنند، با چالش کمتری روبه‌رو می‌شویم؛ زیرا محصول تولید شده برحسب میزان خطر در دو گروه مسئولیت تضامنی و مسئولیت نسبی قابلیت ثبت را دارد؛ ولی هنگامی که یک شخص حقوقی قصد طراحی و تولید یک سیستم هوش مصنوعی را داشته باشد، ثبت سیستم هوشمند برحسب میزان خطر سیستم خواهد بود، نه میزان مسئولیت نوع

1. Policy Department for Citizens' Rights and Constitutional Affairs Directorate-General for Internal Policies

خاصی از شرکت؛ با این توضیح که وقتی شرکتی تحت عنوان شرکت با مسئولیت محدود ثبت می‌شود، تمامی اقدامات آن شرکت تحت قواعد راجع به مقررات شرکت با مسئولیت محدود خواهد بود؛ اما سیستم‌های هوش مصنوعی این‌گونه نیست. فرض کنید شرکتی تحت عنوان شرکت با مسئولیت محدود A تأسیس شود که در زمینه تولید خودران‌ها فعالیت دارد. حال فرض کنید این شرکت دارای سه محصول به اسم‌های B و C و D است. اگر محصول B سیستم کم‌خطر داشته باشد و دو محصول دیگر سیستم‌های پرخطر داشته باشند، سیستم کم‌خطر تحت عنوان سیستم هوشمند B با مسئولیت نسبی ثبت می‌شود و سیستم C و D تحت عنوان با مسئولیت تضامنی ثبت می‌شود و تمام خسارات و مسئولیت‌های مربوط به این محصول، با قواعد مسئولیت تضامنی جبران می‌شود.

سؤالی که پیش می‌آید این است که یک سیستم هوش مصنوعی چگونه می‌تواند درصدد جبران خسارت برآید؟ تمکن مالی این نوع سیستم‌ها چگونه تأمین می‌شود؟

در این راستا نیز برای قانون‌گذار پیشنهاد می‌شود که هنگام ثبت سیستم‌های هوشمند و اعطای شخصیت حقوقی به این نوع سیستم‌ها همانند شرکت، حساب بانکی مخصوصی برای سیستم‌ها افتتاح شود؛ اما حداقل مبلغی که می‌توان با آن اجازه اعطای شخصیت حقوقی به سیستم هوشمند را داد باید برحسب میزان ایجاد خطر نوع فعالیت سیستم تعیین شود؛ به نحوی که در ابتدای امر سیستم یک پشتوانه مالی متناسبی برای جبران خسارات احتمالی داشته باشد. علاوه بر این، برای اطمینان از توانایی جبران خسارت توسط سیستم‌های هوش مصنوعی، به‌ویژه در سیستم‌هایی با مسئولیت محدود، باید سالانه درصد مشخصی از درآمد از طریق سیستم را به حساب افتتاح شده واریز شود که این مبالغ جمع شده فقط برای جبران خسارت احتمالی سیستم نگهداری و در نهایت در همین راستا هزینه می‌شود. یکی از موارد دیگری که می‌توان از این مبلغ جمع شده استفاده کرد که در راستای افزایش ایمنی سیستم و کاهش خطرات احتمالی باشد. چنانچه عوامل مؤثر به بهانه افزایش کیفیت و کارایی و ایمنی سیستم از سرمایه مذکور بهره‌مند شوند؛ ولی اقدامی در این راستا انجام ندهند یا اقداماتی که صورت گرفته نیاز به مبلغ کمتری نسبت به هزینه صرف شده داشته باشد و مبلغ اضافی را که برداشت کرده‌اند در راستای منافع شخصی خود صرف کرده باشند؛ در این صورت، برای جبران هزینه صرف شده، علاوه بر ضمانت اجرایی که قانون‌گذار مقرر می‌دارد، باید مسئولیت سازندگان و عوامل دخیل در برابر شرکت و ثالث از مسئولیت محدود به مسئولیت تضامنی تغییر یابد.

برای اجرای صحیح پیشنهادات قانونی فوق باید یک کارگروه تخصصی برای نظارت سالیانه بر عملکرد سیستم‌های هوش مصنوعی و سازندگان و طراحان آن در قانون پیش‌بینی شود و علاوه بر آن، سیستمی جامع برای گزارش‌های ماهانه از نحوه عملکرد هوش مصنوعی و خسارات به وقوع پیوسته، اختلالات به وجود آمده در سیستم هوشمند، میزان سود حاصل از سیستم هوشمند و سایر اطلاعاتی که در روند بهبود حوزه فنی بر هوش مصنوعی و نظارت قانونی لازم است، طراحی شود.

علاوه بر پشتوانه مالی در نظر گرفته شده فوق، استفاده از مفهوم عدالت توزیعی نیز در تسهیل جبران خسارت و افزایش کیفیت سیستم‌های هوشمند، کارآمد است.

در حقوق ایران، تأثیر رویکرد عدالت توزیعی در مواردی همچون اجتماع اسباب در وقوع خسارت، مشهود است؛ چراکه بسیاری از حوادث زیان‌بار، معلول چند عامل یا سبب است که تعیین سبب مسئول از میان اسباب مختلف کاری دشوار است. از طرفی هم تقسیم مسئولیت به تساوی و مطالبه خسارت از تمام اسباب مجمل، که به‌طور قطع چند سبب در وقوع زیان نقش داشته‌اند، به نحوی که ضرر از مجموع اسباب ایجاد شده باشد و با توجه به دخالت همه اسباب، ترجیح یکی از آن‌ها بر دیگری ترجیح بلامرجه بوده و فاقد توجیه است (Kazemi, 2012, p. 85). پس باید مسئولیت بین همه اسباب توزیع شده و جبران خسارت بر عهده همگی باشد. علاوه بر این مورد، در قانون بیمه اجباری،

صندوق تأمین خسارت بدنی براساس تأثیرپذیری از رویکرد عدالت توزیعی، در مواردی که قانون‌گذار مشخص کرده، پرداخت خسارت را به عهده گرفته است (Nasrarkan et al., 2019, p. 312).

پذیرش عدالت توزیعی در مسئولیت مدنی ناشی از عملکرد هوش مصنوعی به صورت جزئی سودمند است؛ بدین معنا که پذیرش عدالت توزیعی به عنوان روشی غالب برای جبران خسارت در موضوع سیستم‌های هوشمند تأثیرات منفی به همراه دارد که عبارت است از: (۱) ایجاد مانع در راه پیشرفت فناوری‌های نوین: اگر طراحان، مخترعان، تولیدکنندگان، مصرف‌کنندگان و تمامی عوامل دخیل در طراحی و فروش و مصرف سیستم‌های هوشمند بر این باور باشند که هرگونه خسارت ناشی از این سیستم‌ها موجب تقسیم مسئولیت به طور مساوی می‌شود؛ بدون اینکه بررسی‌های فنی دقیقی به عمل آید، انگیزه‌ای در تحقیق و توسعه و تولید این فناوری نوین باقی نمی‌ماند. همچنین شرکت‌های بیمه نیز از پذیرش جبران خسارت همه‌جانبه سیستم‌های هوشمند، به‌ویژه در سیستم‌های پرریسک، اجتناب می‌کنند. پس بهترین راه حل همانی است که به عنوان پیشنهاد در فوق برای قانون‌گذار ذکر شد و فقط در مواردی که به هیچ طریقی نتوان مقدار مسئولیت و اسباب را به طور دقیق تعیین کرد، از عدالت توضیعی بهره‌مند می‌شویم؛ (۲) نبود رغبت به افزایش کیفیت و امنیت سیستم‌های هوش مصنوعی: چنانچه طراحان و تولیدکنندگان خود را در سایه حمایت بی‌حدوحصر شرکت‌های بیمه و صندوق‌های حمایتی ببینند؛ به‌طوری که مطمئن باشند هرگونه نقص فنی که موجب خسارت شود تحت پوشش و حمایت ارگان‌های مختلف است، تلاشی برای بهبود کیفیت سیستم نخواهند کرد؛ اما اگر طبق پیشنهادی که به قانون‌گذار ارائه شد، هر سیستمی از ابتدا نوع مسئولیتش مشخص باشد و همچنین درصدی از درآمد سیستم را در حسابی اندوخته و صرف کیفیت و جبران خسارت ناشی از آن کنند، وضعیت مساعدتر و کیفیت بیشتری را در این حوزه خواهیم داشت.

باتوجه به انواع ربات‌ها، که به سه دسته خودکار، نیمه‌مستقل و کاملاً مستقل یا خودمختار تقسیم می‌شوند، ربات‌های خودمختار و هوشمند دارای علم و اختیار بوده و با توجه به الگوریتم‌های خاص، قادر به ادراک محیط پیرامون و واکنش مناسب در شرایط خاص هستند (Borges, 2018, p. 173). راجع به حقوق ربات‌ها، برخی معتقدند که در طبیعت همه پدیده‌های جاندار و بی‌جان دارای حقوق هستند. در مقابل، عده‌ای نیز به دلیل برخوردار نبودن ربات‌ها از قوه احساس، بر شایسته نبودن آن‌ها برای داشتن حقوق اذعان دارند؛ بالاین حال، معتقدند که اگر ربات‌ها و سیستم‌های هوشمند به حدی توسعه یابند که دارای قدرت انتخاب و تصمیم‌گیری عاقلانه باشند، در آن مرحله باید ربات‌ها را واجد حقوق و تکالیف تلقی کرد (Vaseghi, 2020, p. 329). این نظر قابل انتقاد است؛ زیرا باتوجه به درگیری همه‌جانبه اشخاص با فناوری هوشمند در زندگی روزمره، امکان اعطای شخصیت حقوقی به هوش مصنوعی، در صورتی که فقط خودمختار بوده و قابلیت تفکر داشته باشد، امری نادرست است و آثار زیان‌باری را در پی خواهد داشت. هوش مصنوعی در هر طبقه از اختیار و استقلال که باشد، باید مسئول تقصیر خود باشد و این مسئولیت با اعطای شخصیت محقق می‌شود. نباید از هیچ ربات و سیستم هوشمندی انتظار داشت که کاملاً همانند اختیار و اراده انسان عمل کند تا به مرحله وجود شخصیت برسد؛ زیرا در این حالت، باید سالیان طولانی منتظر اختراع چنین سیستمی باشیم و طی این مدت، آثار زیان‌بار و خسارات زیادی از جانب هوش مصنوعی‌های موجود، بدون مسئولیت باقی می‌ماند. بدین ترتیب، این فناوری خودآموز باید همانند سایر اشخاص در قبال اعمال و تصمیمات خود مسئول باشد (Beranger, 2021, p. 109). برخی نیز اظهار داشته‌اند که تفاوت ربات‌های هوشمند با اشخاص حقیقی که دارای اراده مستقل و اشخاص حقوقی که شامل مجموعه‌ای از افراد حقیقی هستند، موجب عدم پذیرش صلاحیت موضوعات حقوقی می‌شود؛ زیرا ربات‌ها همانند انسان‌ها رفتار هدفمند و آگاهانه ندارند (Lixin & Fanaei, 2023, p. 235). ایرادی که بر این رویکرد وارد می‌شود این است که با توجه به تفاوت شرکت و سیستم هوش مصنوعی، اعطای شخصیت به هوش مصنوعی ارجحیت دارد؛ زیرا شرکت متشکل از انسان‌هاست و از طریق آن‌ها فعالیت می‌کند؛ در حالی که یک سیستم هوشمند توسط انسان‌ها ساخته می‌شود؛ اما با ابتکار عمل خود فعالیت می‌کند و افراد فقط در مواردی به‌طور غیر مستقیم در فعالیت‌های آن دخالت دارند (Chesterman, 2021, p. 126). همچنین وقتی در صحنه بین‌المللی موجودیت‌های غیرانسانی که تبعاتشان کمتر یا مساوی هوش مصنوعی است، شخصیت حقوقی دارند؛ مثل معابد در

هند^۱، رودخانه‌ای در نیوزلند^۲ و کل اکوسیستم اکوادور^۳. پس می‌توان نتیجه گرفت که دولت می‌تواند برای موجودیت‌های جدید، مثل هوش مصنوعی، شخصیت قائل شود (Chesterman, 2020, p. 820). توبر، پژوهشگر حقوقی و جامعه‌شناس، این استدلال را مطرح کرده است که «هیچ دلیل قانع‌کننده‌ای برای محدود کردن انتساب کنش منحصرأ به انسان‌ها و سیستم‌های اجتماعی وجود ندارد» (Teubner, 2007, p. 5). همچنین باید توجه داشت که اعطای شخصیت حقوقی به هوش مصنوعی باعث تشویق نوآوری و توسعه فناوری شود (Turner, 2021, p. 187).

عدم شناسایی شخصیت برای سیستم‌های هوش مصنوعی، دست‌کم برای سیستم‌های سطح بالا، موجب غافل ماندن از خطرات این فناوری می‌شود. پیشرفت هوش مصنوعی به‌اندازه‌ای خارج از کنترل و نگران‌کننده شده است که در سال ۲۰۲۳ پیشنهاد شد همه آزمایشگاه‌های هوش مصنوعی برای حداقل شش ماه آزمایش سیستم‌های قدرتمندتر از GPT-4 را تعلیق کنند که این توقف باید عمدی و قابل تأیید باشد و بازیگران کلیدی این عرصه را شامل شود. اگر چنین توقفی به‌سرعت اعمال نشود، دولت‌ها باید وارد عمل شده و یک تعلیق را وضع کنند (Future of Life Institute, 2023).

بسیاری از محققان درمورد پدیده‌های جدید تحت عنوان «دام انسان‌سازی»^۴ را نگرانی اساسی دارند. این نگرانی درمورد شرکت‌ها نیز، که به مرور دارای شخصیت حقوقی شدند، مطرح بود. حال این سؤال مطرح می‌شود که قانون‌گذار چگونه باید با این نهادها در طول فرایند شکل‌گیری برخورد کند؟

محققان آلمانی راه‌حلی تحت عنوان «اهلیت قانونی جزئی»^۵ برای این چالش ارائه داده‌اند. در واقع مراد از «اهلیت قانونی جزئی» اهلیتی است که قانون اهلیت نهادی را به رسمیت می‌شناسد و در حدود به‌رسمیت شناخته‌شده، برای آن شخصیت و اعتبار قائل می‌شود و بعد با توجه به عملکرد نهاد مذکور، حقوق و تکالیف جدیدی را به آن بار می‌کند (Schirmer, 2020, as cited in Wischmeyer & Rademacher, 2020). ضرورتی ندارد اشخاصی که دارای حقوق هستند و شخصیتشان به رسمیت شناخته شده است، حتی در یک سیستم، حقوق و تعهدات یکسانی داشته باشند (Bryson et al., 2017, p. 280).

در واقع باید در نظر داشت که لزوم اعطای شخصیت حقوقی برای هوش مصنوعی در راستای تسهیل جبران خسارت است؛ زیرا اثبات رابطه سببی در هوش مصنوعی برای عموم اشخاص مسئله‌ای دشوار بوده و نیازمند دانش آکادمیک فنی و حقوقی است. با در نظر گرفتن دلایل و نحوه اعطای شخصیت حقوقی به شرکت‌ها، برای هوش مصنوعی نیز قائل به شخصیت حقوقی می‌شویم؛ اما اختیارات سیستم‌های هوشمند به تبع نوع عملکرد و هدفشان متفاوت خواهد بود. برای مثال همه شرکت‌ها توانایی انعقاد قرارداد را دارند؛ اما آیا یک خودران توانایی انعقاد قرارداد را دارد؟

در این موارد باید به هدف سیستم توجه کرد. خودران در مقام انعقاد یک قرارداد نیست؛ بلکه به‌طور مشخص اهداف خودران تسهیل و افزایش کیفیت رانندگی و کاهش تخلفات رانندگی است و اگر به دنبال اعطای شخصیت به این گونه سیستم‌ها هستیم، برای تسهیل اثبات و جبران خسارت است و همچنین با اعطای شخصیت حقوقی و مسئول دانستن سیستم هوشمند موجب تلاش عوامل دخیل به‌منظور افزایش کیفیت و

1. Shriomani Gurudwara Prabandhak ... vs Shri Som Nath Dass & Ors on 29 March, 2000, LAWS (SC)-2000-3-67, Supreme Court of India, available at: <https://indiankanon.org/doc/1478973/> (Last Visited: 2024.02.10)

2. Te Awa Tupua (Whanganui River Claims Settlement) Act 2017, available at: <https://www.legislation.govt.nz/act/public/2017/0007/latest/whole.html>

3. Constitution of the Republic of Ecuador 2008, As Amended to 2021, available at: <https://constitutions.unwomen.org/en/countries/americas/ecuador?%5B0%5D=provisioncategory%3A694>

4. Humanization trap

5. Teilrechtsfähigkeit

دقت در طراحی سیستم می‌شود. شخصیت حقوقی دارای آثار حقوقی مختلفی است؛ اما این مقاله به‌صورت خاص بر انتساب مسئولیت مدنی و تسهیل جبران خسارت تمرکز دارد.

۲. مبنای مسئولیت مدنی هوش مصنوعی

یکی از استانداردهای مهم ارائه شده برای سازمان‌دهی هوش مصنوعی توسط گروه اروپایی اخلاق در علم و فناوری‌های نوین کمیسیون اتحادیه اروپا، داشتن مسئولیت است (European Commission & European Group on Ethics, 2018).

مسئولیت ناشی از عملکرد هوش مصنوعی در دو نوع مدنی و کیفری قابل بررسی است که در این پژوهش تمرکز بر مسئولیت مدنی است. مسئولیت مدنی هنگامی مطرح می‌شود که شخصی ملزم به جبران خسارتی است که بدون مجوز قانونی به حق دیگری لطمه بزند و بر اثر آن، وی متضرر شود (Emami, 2016, p. 663). زمانی که ضرری بر متضرر حادث می‌شود، باید رابطه میان خسارت و فعل زیان‌بار مشخص شود. برای نیل به این هدف، مبنای مسئولیت مدنی راه‌حل منطقی و مناسبی است (Badini, 2005, p. 76). مبنای مسئولیت یک توجیه علمی و عملی برای انتساب ضرر به زیان‌زننده و دلیلی برای جبران خسارت از سوی او است (Jourdain, 2006, p. 76).

انتساب مسئولیت مدنی به هوش مصنوعی مستلزم تعریف مبنایی مناسب و لازمه تعیین این مبنای، بررسی قوانین در این عرصه است. مطابق گزارش گروه تخصصی مسئولیت و فناوری‌های نوین کمیسیون اروپا، مقررات ایمنی محصول اتحادیه اروپا نمی‌تواند به طور کامل احتمال ورود خسارت ناشی از عملکرد این فناوری‌ها را منتفی سازد (Expert Group, 2019).

همان‌طور که در ۸ می ۲۰۲۳ معاون ارشد کمیسیون اتحادیه اروپا بر لزوم تسریع قانون‌گذاری در هوش مصنوعی تأکید کرد (dig.watch, 2023)، عدم کفایت قوانین موجود، جوامع را بر آن داشته که در این عرصه، به قانون‌گذاری مناسب اقدام کنند.

در آگوست ۲۰۲۳ اداره تحقیقات حقوقی جهانی ایالات متحده آمریکا مطالعاتی پیرامون قوانین مصوب در عرصه هوش مصنوعی انجام داد و نتایج این مطالعات را منتشر کرد. طبق این بررسی‌ها مشخص شد که دامنه قوانین در کشورهای مختلف متفاوت است. برخی از قوانین، مثل قوانین کشورهای آرژانتین، بلاروس، مصر، پرتغال، قطر، صربستان و نیز در قوانین پیشنهادی کشورهای همچون چین، شیلی، برزیل و کلمبیا چهارچوب‌های کلی، مثل توسعه و تجاری‌سازی سیستم‌های هوش مصنوعی، اصول اخلاقی و حقوق اساسی برای توسعه و همچنین تأسیس نهادهایی برای نظارت و مشورت دولت در این عرصه تدوین شده است. در برخی از قوانین دیگر، در کشورهای مثل استونی، یونان، مجارستان، لهستان و چین، کاربردهای خاص سیستم‌های هوش مصنوعی، مانند سیستم‌های پرداخت و تصمیم‌گیری در مشاغل مختلف توسط عوامل دولتی و خصوصی، وسایل نقلیه خودران، محتوای آنلاین، تولید متن یا تصویر، پردازش سیگنال گفتار، به تصویب رسیده است. در برخی از کشورها، قوانین درباره شفافیت عملکرد سیستم‌های هوشمند و جلوگیری از سوءاستفاده تقنینی اتخاذ شده است. در سال ۲۰۲۱ در اتحادیه اروپا پیشنهادهایی برای تدوین قوانینی در راستای شفافیت برای سیستم‌های هوش مصنوعی کم‌خطر و ممنوعیت‌ها و الزامات اجباری برای برخی از سیستم‌های پرخطر ارائه شد. در سال ۲۰۲۲ دستورالعملی برای وضع قوانین یکسان برای برخی از جنبه‌های مسئولیت مدنی در قبال خسارات ناشی از سیستم‌های هوش مصنوعی معرفی شد. علاوه بر کشورهای مختلفی که در حال تلاش برای تدوین مقررات لازم در راستای سازمان‌دهی سیستم‌های هوش مصنوعی هستند، سازمان‌های بین‌المللی نیز در این مسیر گام‌هایی برداشته‌اند (Cantekin, 2023).

آخرین قانون در حوزه هوش مصنوعی در ۱۳ ژوئن ۲۰۲۴ تصویب و منتشر شد. اتحادیه اروپا بعد از چند سال مطالعه و بررسی و ارائه انواع مقررات و دستورالعمل‌ها، قانونی تحت عنوان «قانون هوش مصنوعی»^۱ تصویب کرد. در این قانون، به مسئولیت سیستم‌های هوشمند اشاره شده و ارائه‌دهندگان این سیستم‌ها در مرکزیت مسئولیت قرار گرفته‌اند. طبق بند ۱ ماده ۳۱ قانون مذکور، نهادی براساس قوانین ملی تعیین می‌شود که دارای شخصیت حقوقی خواهد بود. بند ۲ و ۳ ماده فوق بیانگر تعیین مسئولیت‌ها، خط مشی و تمامی اقدامات ضروری از سوی ارگان‌های

1. European Union Artificial Intelligence Act (EU AI Act)

لازم برای سازمان‌هایی است که سیستم‌های هوش مصنوعی را ارائه می‌دهند. نظارت شفاف و صحیح بر سازمان‌های ارائه‌دهنده سیستم‌های هوش مصنوعی مستلزم وجود نهادهای مستقل از این سازمان‌هاست که بند ۴ به این امر مهم پرداخته است و مقرر می‌دارد: «نهادهای ناظر باید از ارائه‌دهنده سیستم هوش مصنوعی پرخطری که راجع به آن [سیستم] فعالیت‌های ارزیابی انطباق انجام می‌دهند، مستقل باشند. به علاوه، نهادهای ناظر باید از هر اپراتور دیگری که در سیستم‌های هوش مصنوعی ارزیابی شده نفع اقتصادی دارد و رقابتی ارائه‌دهنده نیز مستقل باشد. این امر استفاده از سیستم‌های پرخطر ارزیابی شده هوش مصنوعی که برای عملکرد نهاد ارزیابی انطباق ضروری هستند یا استفاده از چنین سیستم هوش مصنوعی پرخطر برای اهداف شخصی را منع نمی‌کند». بند ۹ ماده ۵۷ این قانون نظارت‌های خاصی را در راستای تقویت و توسعه نوآوری و رقابت در فناوری هوش مصنوعی پیش‌بینی کرده است. طبق بند ۷ ماده ۶۰: «هر حادثه جدی شناسایی شده در جریان آزمایش در شرایط دنیای واقعی باید طبق ماده ۷۳ [این قانون] به مرجع نظارت بر بازار ملی گزارش شود. ارائه‌دهنده یا ارائه‌دهنده آتی باید اقدامات مدیریت ریسک فوری را انجام دهد و در صورت عدم موفقیت [در این امر]، تا زمانی که چنین اقداماتی انجام نگرفته باشد، باید آزمایش در شرایط دنیای واقعی را تعلیق نماید یا به نحو دیگری بدان خاتمه دهد. ارائه‌دهنده یا ارائه‌دهنده آتی باید برای احیای سریع سیستم هوش مصنوعی به محض خاتمه آزمایش در شرایط دنیای واقعی، آیینی مقرر نماید.»^۱

قانون مذکور جمع‌آوری شده از همه بررسی و مطالعات و دستورالعمل‌های پارلمان اتحادیه اروپا در چند سال اخیر است. یکی از مهم‌ترین نقدهایی که می‌توان بر این قانون و قوانین مشابه در این حوزه وارد کرد، عدم پیش‌بینی عوامل دخیل در هوش مصنوعی به طور جزئی است.

برای تدوین یک قانون جامع و کامل و ابهام‌زدایی در عرصه هوش مصنوعی، نیاز به قانونی است که موارد ذیل را پیش‌بینی کند:

- ۱) ارائه تعاریف علمی و صحیح از انواع هوش مصنوعی و تمامی اشخاصی که در روند عملکرد این سیستم دخیل هستند؛
- ۲) تبیین نحوه ثبت هوش مصنوعی به عنوان یک شخصیت حقوقی و تعیین نحوه نظارت و اداره سیستم‌های هوشمند؛
- ۳) تدوین قوانین خاص برای رشته‌های تخصصی. برای مثال در زمینه استفاده از هوش مصنوعی در حوزه پزشکی، لازم است تمامی موازین قانونی درحدامکان پیش‌بینی و تدوین شود یا به عنوان مثالی دیگر، اگر در زمینه نگارش قراردادها سیستم هوشمندی طراحی شود که ادعای توانایی بر نگارش یک قرارداد حرفه‌ای را دارد، بایستی با رعایت چهارچوب حقوق قراردادها و توجه به اصول فنی، یک قانون تخصصی در این عرصه تصویب شود.

نکته حائز اهمیت دیگر این است که در تدوین چنین قوانینی، باید متخصصان مرتبط رشته‌های علمی مطالعات مشترکی داشته باشند تا نتیجه مطلوبی حاصل شود. مثلاً برای تدوین قوانین و دستورالعمل‌هایی در حوزه تجهیزات دارای سیستم هوش مصنوعی عمران و سازه‌های ساختمانی، نیازمند مشارکت مهندسان عمران و وکلای تخصصی است؛ چراکه چنین متخصصانی به پیش‌بینی عوامل دخیل، اشکالات فنی، اصول قانونی حوزه عمران و ساخت انواع سازه‌ها اشراف کامل دارند و در نتیجه قوانین مطلوبی را تدوین می‌کنند.

در همه مقررات ملی اروپایی، جز برخی موارد استثنایی، اصل مسئولیت ناشی از تقصیر پیش‌بینی شده است (Zakerinia, 2023, 150). در راستای تعیین مبنای مناسب، پارلمان اروپا در اکتبر ۲۰۲۰ قطعنامه‌هایی همراه با توصیه‌هایی به کمیسیون درمورد نظام حقوقی مسئولیت مدنی برای هوش مصنوعی تصویب کرد. پارلمان توصیه کرد که برای سیستم‌های پرخطر، مبنای خطر در نظر گرفته شود. پارلمان در قطعنامه بعدی خود در ۳ مه ۲۰۲۲، درمورد هوش مصنوعی در عصر دیجیتال، تأکید کرد ضمن آنکه سیستم‌های هوش مصنوعی پرخطر باید تحت قوانین سخت‌گیرانه مسئولیت (همراه با پوشش بیمه اجباری) قرار گیرند، مسئولیت هر فعالیت، دستگاه یا فرایند دیگری که با سیستم‌های هوش مصنوعی هدایت شده و به صدمه یا خسارت منجر می‌شود، مسئولیت مبتنی بر تقصیر باشد (Madiega, 2023).

1. eur-lex.europa.eu

براساس مواد ۴۸ و ۵۲ قانون هوش مصنوعی اتحادیه اروپا، سیستم‌های هوش مصنوعی پرخطر به سیستم‌هایی اطلاق می‌گردد که با حقوق اساسی اشخاص در ارتباط است. علاوه بر این مورد، نوع دیگر این سیستم‌ها، هوش مصنوعی مستقل است.^۱

۱-۲. نظریه تقصیر و انطباق آن با هوش مصنوعی

مسئولیت مدنی مبتنی بر تقصیر شخصی را مسئول می‌داند که مرتکب تقصیر شده است و بار اثباتی آن اصولاً برعهده متضرر است (Safai & Rahimi, 2019, pp. 61-62). مسئولیت مدنی در حقوق ایران مبتنی بر تقصیر است و به‌طور اخص در مورد اتلاف از نظریه خطر و علت پیروی شده است (Emami, 2016, pp. 675-676). در الزامات خارج از قرارداد تقصیر همیشه خلاف اصل است (Katouzian, 2021, p. 22). در صورت بروز یک خسارت ممکن است عوامل مختلفی دخالت داشته باشند. تشخیص این عوامل و همچنین نحوه و میزان تأثیر آن‌ها در خسارت وارده اهمیت بسزایی دارد. در مباحث حقوقی، از این عوامل با عنوان تعدد اسباب یاد می‌شود و نظریات مختلفی پیرامون این بحث مطرح شده است.

از میان عوامل متعدد دخیل در وقوع خسارت، صرفاً عللی مبنای مسئولیت مدنی هستند که عرف آن‌ها را مسبب بروز خسارت بدانند. برای تفکیک چنین عللی، دو اصطلاح «شرط» و «سبب» مطرح شده است. مقصود از «شرط» تمام عواملی است که با پیوستن آن‌ها به یکدیگر خسارت ایجاد می‌شود، ولی «سبب» فقط به عواملی اشاره دارد که از نظر عرفی و حقوقی خسارت مستند به آن‌هاست. برای تشخیص سبب بودن یا نبودن باید به این سؤال پاسخ داد که آیا در صورت عدم بروز عامل مدنظر، همچنان خسارت به وقوع پیوسته به همان شکل ایجاد می‌شود یا خیر؟ در صورت مثبت بودن جواب، مشخص می‌شود که عامل مذکور سبب مؤثری نبوده است. در تعیین اسباب مؤثر به ملاک‌های مختلفی توجه می‌شود که از جمله آن‌ها ایراد خسارت مستقیم است. در این نوع خسارت، که به «اتلاف» مشهور است، شخصیت یا اعضا و جوارح و ابزاری که در اختیار خود دارد موجب ورود خسارت به دیگری می‌شود. این نوع سبب واضح‌ترین علتی است که عرف آن را به راحتی سبب عرفی و حقوقی خسارت می‌پذیرد. داشتن عمد و وقوع یک امر غیر معمول در تشخیص سبب مؤثر، از دیگر ملاک‌های تعیین کننده است. منظور از امر غیر معمول، امری خلاف روال عادی و معمول مدنظر عرف و قوانین است. در تعیین تقصیر، خطا، رابطه سببیت و همچنین دامنه مسئولیت مدنی، قابلیت پیش‌بینی از اهمیت بسزایی برخوردار است و نظریه پردازان بسیاری قابلیت استناد خسارت را به قابلیت پیش‌بینی موکول کرده‌اند؛ البته این به معنای محدودیت موارد انتساب خسارت به موارد قابل پیش‌بینی نیست؛ زیرا اگر عرف رابطه سببیت را احراز کند، همچنان خسارت به فعل مذکور منتسب می‌شود (Babaei, 2022, pp. 87-100).

به‌طورکلی، تعیین تقصیر و میزان آن برای هرکدام از اسباب منوط به احراز رابطه سببیت است؛ حال امکان دارد سبب مستقیم یا غیر مستقیم باشد و فقط تفاوت آن در بار اثباتی باشد؛ به این صورت که اگر سببی به‌طور غیر مستقیم موجب ورود خسارت شود، مدعی ورود خسارت باید رابطه سببیت را اثبات کند؛ مثلاً برای اثبات بی‌احتیاطی تولیدکننده ربات‌های خودمختار، زیان‌دیده باید نقض وظیفه قانونی مراقبت تولیدکننده و در نتیجه آن، ورود آسیب به خود را ثابت کند (Hekmatnia et al., 2019, p. 244)؛ چراکه مشاور کمیسیون عمومی اروپا در مورد انطباق قوانین مسئولیت مدنی با عصر دیجیتال و هوش مصنوعی بیان داشته است که مسئولیت نقض مفروض یک وظیفه عینی مراقبت همچنان از ایده تقصیر ناشی می‌شود (Wendehorst et al., 2022). در نظریه تقصیر، مادامی که دادگاه خواننده را قاصر تشخیص ندهد، وی مسئول نیست؛ بدین معنا که رفتار خواننده باید کمتر از استاندارد تعیین شده در قانون باشد (Cooke, 2010, p. 32).

در هوش مصنوعی ارزیابی تقصیر و تشخیص این موضوع که کدام فعل یا ترک فعل تعدی یا تفریط است، نیازمند تعیین و ارائه استانداردهایی است که این استانداردها به تناسب هر علمی به طور مجزا از سوی متخصصان آن علم تعریف شده‌اند و کارشناسانی مثل گروه

1. eur-lex.europa.eu

ناظر ذکر شده در قانون هوش مصنوعی، مصوب اتحادیه اروپا، می‌تواند با بررسی هریک از عوامل طراحی تا تولید و استفاده از هوش مصنوعی و تطبیق استانداردهای مدنظر برای عملکرد هرکدام، تقصیر و میزان آن را نسبت به هر کدام از عوامل ارزیابی کنند.

با توجه به ارزیابی استانداردهای موجود در یک فناوری و سطح توقعات مصرف‌کننده، می‌توان با سهولت مسبب عیب را تشخیص داد. تمام عوامل مؤثر در ساخت سیستم هوشمند، از جمله طراحان، تولیدکنندگان سخت‌افزار و نرم‌افزار، توسعه‌دهندگان و پشتیبانان، همگی می‌توانند مسئول شناخته شوند (Rahbar & Dehghanpour Farashah, 2021, p. 536).

هنگامی که یک هوش مصنوعی پیشرفته، که براساس الگوریتم‌ها و توانایی یادگیری خود برای رسیدن به اهدافش فعالیت می‌کند، به شخص ثالثی ضرر وارد کند، می‌توان با اعمال محدودیت‌هایی، خود آن را مسئول قلمداد کرد. عوامل دخیل در طراحی هوش مصنوعی به میزانی که ورود خسارت به کنش آن‌ها بستگی دارد، مسئول هستند؛ اما در صورتی که نتوان ضرر را به آن‌ها منتسب کرد، فقط هوش مصنوعی با هویتی که مانند شخصیت حقوقی برای او احراز شده است مسئول شناخته می‌شود (Valipour & Esmaeili, 2021, p. 7).

براساس گزارش کمیته سیاست اقتصاد دیجیتال اتحادیه اروپا، بازیگران هوش مصنوعی به اشخاصی اطلاق می‌شود که نقش فعالی در چرخه حیات سیستم هوش مصنوعی ایفا می‌کنند؛ از جمله سازمان‌ها و افرادی که هوش مصنوعی را مستقر یا اجرا می‌کنند. چرخه حیات سیستم هوش مصنوعی معمولاً شامل چندین مرحله است که عبارت‌اند از: برنامه‌ریزی و طراحی، جمع‌آوری و پردازش داده‌ها، ساخت مدل(ها) و/یا تطبیق مدل(های) موجود برای وظایف خاص، آزمایش، ارزیابی، تأیید و اعتبارسنجی، در دسترس قرار دادن برای استفاده/استقرار، اجرا و نظارت و بازنشتگی/از کار افتادن. این مراحل اغلب به‌صورت تکراری انجام می‌شوند و لزوماً متوالی نیستند (OECD, 2023).

نقض یک وظیفه مراقبتی که ناشی از سهل‌انگاری (بی‌احتیاطی یا بی‌مبالاتی) باشد، می‌تواند سیستمی هوشمند را در مسیری هدایت کند که برای آن طراحی نشده است یا عدم نظارت مؤثر بر سیستم، تقصیر را متوجه فردی می‌کند که این وظیفه را برعهده دارد (Wendehorst, 2022, p. 209).

نفوذ هوش مصنوعی در همه عرصه‌ها با سرعت چشمگیری توسعه می‌یابد که یکی از مهم‌ترین حوزه‌های پیشرفت آن، پزشکی و حوزه سلامت است. گفتنی است که برای استفاده از سیستم هوشمند به‌منظور درمان و جراحی باید به بیمار اطلاع داده شود تا بیمار رضایت آگاهانه داشته باشد (Neri et al., 2020, p. 1). طبق بررسی سرویس تحقیقات پارلمان اروپا، هوش مصنوعی در حوزه پزشکی علاوه بر تأثیرات مثبت فراوان، خطراتی نیز به دنبال دارد. در واقع با توجه به این خطرات، بازیگران دخیل در عرصه پزشکی و بهداشتی هوش مصنوعی عبارت‌اند از: تولیدکنندگان، طراحان، پزشکان، متخصصان مراقبت‌های بهداشتی، اپراتورها و بیماران (European Parliamentary Research Service, 2022). طبق گزارش مشترکی که دفتر پاسخ‌گویی دولت آمریکا و آکادمی ملی پزشکی این دولت منتشر کردند، درمان مؤثر بستگی به تشخیص دقیق و به‌موقع دارد. خطاهای تشخیصی رایج‌ترین، فاجعه‌بارترین و پرهزینه‌ترین خطاهای پزشکی هستند. یادگیری ماشین¹، که زیرشاخه‌ای از هوش مصنوعی است، می‌تواند به تشخیص بالینی دقیق‌تر و زود هنگام‌تر منجر شود. سه گروه ذی‌نفع که در توسعه و استقرار فناوری‌های تشخیص پزشکی مشارکت دارند، نهادهای تحقیق و توسعه، طراحان و تنظیم‌کننده‌ها و کاربران نهایی هستند. علاوه بر این، هنگامی که براساس نتایج نادرست یک سیستم صدمه‌ای به بیمار وارد شود، خود سیستم هم می‌تواند دارای مسئولیت باشد، به‌خصوص اینکه پدیده «جعبه سیاه» در سیستم‌ها کنترل ارائه‌دهندگان را از بین می‌برد (In-text Citation). با وجود چنین چالش‌هایی در بررسی ضرر منتسب به هوش مصنوعی، پیش‌بینی و ارزیابی تمامی اختلالات در سیستم هوشمند نیز با چالش‌هایی روبه‌رو خواهد بود (Takhshid, 2021, p. 238). مشکل جعبه سیاه در واقع ناتوانی در درک کامل فرایند تصمیم‌گیری هوش مصنوعی و ناتوانی در پیش‌بینی تصمیمات یا خروجی‌های سیستم هوشمند است؛

1. Machine Learning (ML)

بعضی از جعبه‌سیاه‌ها قوی هستند؛ بدین توضیح که این نوع فرایند تصمیم‌گیری کاملاً برای انسان مبهم است و نمی‌توان آن را تحلیل کرد؛ اما بعضی دیگر از جعبه‌سیاه‌ها ضعیف هستند؛ یعنی با اینکه فرایند تصمیم‌گیری آن برای انسان مبهم است، اما می‌توان مهندسی معکوس یا کاوش کرد تا نوعی رمزگشایی جزئی در آن صورت گیرد (Bathae, 2018, pp. 905-906). هوش مصنوعی این پتانسیل را دارد که روابط مراقبتی را تغییر دهد و در واقع موجب جابه‌جایی و تغییر میزان مسئولیت پزشک شود؛ اما این بدان معنا نیست که رابطه سنتی پزشک و بیمار کاملاً حذف و نادیده گرفته می‌شود (Mittelstadt, 2021). سنجش پزشکان بخشی از تخصص آن‌هاست که اطلاعات را به‌طور مستقل ارزیابی کنند و هرچه یک خطا تهدیدآمیزتر باشد باید مبنای تصمیم خود را با جدیت بیشتری زیر سؤال ببرند (Molnar Gabor, 2022, pp. 350-351). میزان تحمل تقصیر از یک سو و شدت آسیب‌های ناشی از آن از سوی دیگر، می‌تواند معیار اصلی برای تعیین و محاسبه زیان‌زننده و میزان خسارت از جانب عاملان باشد (Giuffrida, 2019, p. 452).

در کارکرد یک سیستم، باید آگاهانه عمل کرد و در نظارت آن دقت داشت. برای مثال، تا قبل از استفاده از پلنفرم یادگیری ماشینی به اسم «روتِر»^۱ بیمارستان‌ها و مراکز درمانی نیویورک در فصل واگیری آنفلوآنزا با ازدحام مراجع روبه‌رو بودند؛ ولی تماس‌گیرندگان با روتِر در مسیر مناسبی به مراکز درمانی هدایت می‌شدند. این سیستم در شروع سال نو، دچار اختلال شد و در نتیجه این اختلال، بیماران آمبولانس‌های خارج از ورودی بیمارستان، در ترافیک ماندند و حتی برخی از آن‌ها نیز فوت شدند. طبق تحقیقاتی که به دستور دولت انجام شد، نظارت‌کنندگان بر سیستم روتِر از الگوی مسیریابی غیرمعمول در شب سال نو آگاه بودند؛ اما در حل این مشکل اقدامی نکردند (Arnold & Toner, 2021, pp. 14 - 15).

سیستم «روتِر» جزو سیستم‌های کم‌خطر بوده و به‌طور مستقیم با جان انسان‌ها ارتباط ندارد و صرفاً در راستای انتظام به امور طراحی شده است. به همین سبب، سیستم کم‌خطری به‌شمار می‌رود؛ بنابراین از این لحاظ مبنای تقصیر بر آن اعمال می‌شود و به دلیل این که عوامل نظارت‌کننده بر سیستم از اختلال سیستم آگاه بودند و نحوه اختلال و مشکلات پیش‌آمده کاملاً قابل پیش‌بینی بوده است، همه مسئولیت ناشی از این اختلال، متوجه عوامل ناظر بر سیستم است؛ البته اگر به غیر از ناظران اشخاص دیگری نیز مثل طراح سیستم و مسئول فنی از این اختلال آگاه باشند، مسئولیت متوجه آن‌ها نیز هست. در این مثال، کاربران دخالتی در خسارت به وجود آمده نداشتند.

برای اعمال مبنای تقصیر در هوش مصنوعی، باید هر عنصر را با دقت بررسی کرد تا میزان تقصیر هر سبب مشخص شود؛ زیرا تقصیر جزئی هر مسبب پیامدهای زیادی را در پی دارد. در پرونده «سینگ علیه ادواردز»^۲ در اکتبر ۲۰۰۹، سینگ برای یک عمل جراحی قلب به مرکز پزشکی «پراویدنس» مراجعه کرد. در حین جراحی، مانیتور شرکت ادواردز دچار مشکل شد و دستگاه‌های ایمن را خاموش کرد و باعث شد کاتتر، که به قلب سینگ متصل بود، گرم شود و قلبش دچار سوختگی شد. بر اثر درمان این عارضه و مصرف داروها، بیمار به سرطان خون مبتلا شد. سینگ دعوای مسئولیت محصولات علیه ادواردز را مطرح کرد و ادواردز مسئولیت خود را پذیرفت؛ اما مدعی مسئولیت پراویدنس نیز شد. شواهد ارائه شده در محکمه نشان داد که ادواردز از اوایل سال ۱۹۹۸ می‌دانست نقصی در دستگاه مانیتور قلب وجود دارد و این مانیتور حاوی نرم‌افزار مخربی بود که می‌توانست محرک‌های ایمن را، که برای اطمینان از گرم شدن بیش از حد کاتتر بود، از بین ببرد. در یادداشتی در ۱۷ ژوئیه ۱۹۹۸ از طرف طراح اصلی نرم‌افزار ادواردز، به کشف اختلال در نرم‌افزار اشاره شده بود؛ با این حال این مانیتورها در سال ۲۰۰۰، بدون حذف این نرم‌افزار، عرضه شدند. دستگاه چندین بار اختلال گرم شدن و آتش گرفتن را بروز داد. این حوادث حکایت از آن داشت که استفاده از این ابزارها مسئولانه نبود و درنهایت، دادگاه ۹۹/۹ درصد تقصیر را به ادواردز و ۱/۱ درصد تقصیر را به پراویدنس منسوب و به تناسب آن، غرامت را تعیین کرد.

1. Routr

2. SINGH v. EDWARDS LIFESCIENCES CORPORATION, Court of Appeals of Washington, Division 1, No. 61823-7-I., July 06, 2009, available at: <https://caselaw.findlaw.com/court/wa-court-of-appeals/1100881.html>

در پرونده یاد شده، به دلیل اینکه مانیتور و دستگاه‌های ایمن ارتباط مستقیمی با عمل جراحی نداشتند و صرفاً یکی از تجهیزاتی بودند که باید تحت نظارت مرکز درمانی و شرکت سازندگان استفاده می‌شدند، جزو سیستم‌های کم‌خطر به‌شمار می‌روند؛ در نتیجه مبنای تقصیر بر آن جاری است و عوامل دخیل در خسارت، که مرکز درمانی پراویدنس و شرکت ادواردز بود، به تناسب سببیت خود مسئول جبران خسارت شناخته شدند.

احتمال طرح سؤالی در استدلال فوق وجود دارد و آن هم این است که چرا برای تمام سیستم‌هایی که در حوزه پزشکی وجود دارد، مبنای خطر جاری نشود؛ زیرا تجهیزات پزشکی در حوزه سلامت و جان انسان قرار دارد؟

برای پاسخ به این سؤال، چنین عنوان می‌شود که برای قانون‌گذاری در عرصه فناوری و به‌ویژه هوش مصنوعی باید به نحوی اقدام کرد که مانعی برای مسیر پیشرفت علم و فناوری به وجود نیاید و تدوین قوانین با لحاظ تمامی جوانب صورت گیرد. اگر تمامی تجهیزاتی را که با سلامت و ایمنی انسان به هر نحوی ارتباط دارند، در زمره سیستم‌های پرخطر و مشمول نظریه خطر قرار گیرند، هیچ مخترع و طراحی انگیزه و جرتی برای ساخت تجهیزات منطبق با علم روز دنیا نخواهد داشت. این مسئله یکی از مهم‌ترین عوامل عدم پیشرفت فناوری در کشور خواهد بود. استفاده مسئولانه از هوش مصنوعی خطرات آن را کاهش می‌دهد (Dignum, 2019, p. 66). بررسی انواع عملکرد سیستم هوشمند، از جمله طراحی و راه‌اندازی مکانیسم‌های تضمین ایمنی همچون نظارت و پاسخ به سیگنال‌های اخطار در هنگام بروز مشکل، نصب امکاناتی که در هنگام بروز مشکل حاد به کاربران اجازه قطع کامل سیستم را می‌دهند، مثل ترمز اضطراری و ارائه خدمات پس از فروش به‌صورت برخط و بی‌وقفه، نشان‌دهنده اقدامات احتیاطی طراحان سیستم بوده که درصد خسارت انتسابی به آن‌ها را کاهش می‌دهد (Zakerinia, 2024, pp. 163-164). البته باید کیفیت کارکرد این خدمات و امکانات دقیقاً مطابق تبلیغات و دستورالعمل‌های ارائه شده از سوی تولیدکنندگان باشد و مصرف‌کنندگان نیز طبق دستورالعمل‌ها، سیستم‌های هوشمند را به کار گیرند. چنین رویکردی در آرای دادگاه‌ها نیز مشاهده می‌شود. برای مثال، در سال ۲۰۲۰ شکایتی طرح شد که در آن، رأی به عدم تقصیر تسلا صادر شد. در این پرونده، مالک تسلا اتوپایلویت و سیستم کروز کنترل آگاه از ترافیک را فعال کرد. خودرو حدود ۲۰ تا ۳۰ مایل در ساعت رانندگی می‌کرد که ناگهان از جاده منحرف شد. در این هنگام، کیسه هوای خودرو باز شد و فک و دندان‌های راننده شکست. در این پرونده، تسلا براساس دستورالعمل‌هایی که برای رانندگی ارائه داده بود استدلال کرد که راننده باید دستش روی فرمان باشد، ولی راننده چنین اقدامی را انجام نداده بود. قاضی نیز در این پرونده با استناد به سهل‌انگاری راننده و هشدارهایی که تسلا درباره استفاده از محصول داده بود، حکم به جبران خسارت نداد^۱.

در پرونده فوق، خودران طراحی شده از نوع سیستم کم‌خطر است؛ بدین دلیل که تسلا هیچ ادعایی مبنی بر تمام‌هوشمندبودن خودران نکرده و دستورالعمل‌ها و هشدارهای لازم را برای احتیاط و عملکرد نظارتی راننده داده بود. به همین منظور، حکم دادگاه مبنی بر عدم تقصیر تسلا صادر شد؛ البته در مورد چنین سیستم‌هایی نباید بررسی‌های صورت گرفته به رعایت و عدم رعایت دستورالعمل‌ها محدود شود. در این پرونده، خسارتی که به وجود آمد به دلیل باز شدن کیسه هوای خودرو است. روند بررسی باید چنین باشد که کدام دلیل و به چه میزان سبب ورود خسارت است؟ اگر باز شدن کیسه هوا استاندارد بوده و ناشی از اختلال در سیستم نبوده و صرفاً سهل‌انگاری راننده که موجب خروج خودرو از جاده شده، تصمیم دادگاه کاملاً درست است؛ اما اگر با بررسی‌های فنی مشخص شود که حتی با خروج خودرو از جاده با سرعت و استانداردهای تعریف شده برای سیستم، نباید کیسه هوا باز می‌شد؛ اما نقص فنی موجب باز شدن کیسه هوا شده است. در این صورت، علاوه بر راننده، باید تسلا نیز مسئول شناخته می‌شد. در واقع صرف سهل‌انگاری راننده نباید باعث شود که بررسی‌های فنی همه‌جانبه روی محصول صورت نگیرد.

1. JUSTINE HSU v. TESLA, Los Angeles County Superior Court of California, May 14, 2020, available at: <https://unicourt.com/case/ca-la23-justine-hsu-vs-tesla-inc-541298>

با توجه به پرونده‌های مطرح شده در می‌یابیم که وجود نهادی نظارتی برای عملکرد و رعایت استانداردها و همچنین انطباق تبلیغات صورت گرفته محصولات دارای هوش مصنوعی با واقعیت عملکردی آن‌ها بسیار حائز اهمیت است. این نهاد نظارتی، در آخرین قانون این عرصه، یعنی «قانون هوش مصنوعی» مصوب اتحادیه اروپا در سال ۲۰۲۴، پیش‌بینی شده است. این نهاد بسیاری از چالش‌ها را حل کرده و تقصیر و عدم تقصیر هریک از عوامل دخیل در سیستم‌های هوشمند را، از طراح و تولیدکننده تا مصرف‌کننده، مشخص کند؛ بدین طریق که با سنجش معیارهای مشخص شده توسط نهاد نظارتی، می‌توان ارزیابی کرد که کدام یک از اشخاص دخیل وظایف خود را انجام داده یا نداده‌اند و یا به‌طور ناقص رعایت کرده‌اند. این نکته‌ای مثبت در قانون مذکور است؛ اما نقد وارد بر این قانون این است که به‌طور کلی راه درمان را پیش‌بینی کرده است؛ درحالی که باید در این موضوع مهم و خطیر پیشگیری را مقدم بدانیم؛ با این توضیح که باید سازوکاری را طراحی کرد که نظارت‌های فنی و قانونی در بدو امر صورت گیرد و نوع مسئولیت هریک مشخص باشد.

۲-۲. نظریه مسئولیت مبتنی بر خطر و انطباق آن با هوش مصنوعی

طبق این نظریه، هنگامی که شخصی از فعالیت خود نفعی را کسب کرد؛ هرچند تقصیری در میان نباشد، ولی در نتیجه فعالیت وی خسارتی به شخصی دیگر وارد شود، باید درصدد جبران خسارت آن برآید و زیان‌دیده فقط باید ثابت کند که ضرر وارد شده ناشی از فعالیت سودآور خوانده است (Safai & Rahimi, 2019, pp. 62-63).

وسایل نقلیه خودران یکی از صنایع‌های مهمی است که هوش مصنوعی در آن توسعه یافته است. این خودران‌ها به پنج سطح دسته‌بندی می‌شوند: از سطح صفر که هیچ‌گونه اتوماسیونی ندارد تا سطح پنج که کاملاً خودران است (Barzegar & Elham, 2020, pp. 201-207). طبق تعریف ارائه شده از کمیسیون حقوقی انگلستان، ولز و کمیسیون حقوقی اسکاتلند، وسایل نقلیه خودران به وسایل نقلیه‌ای اطلاق می‌شود که می‌توانند خود را، بدون کنترل یا نظارت یک فرد، دست‌کم برای بخشی از سفر، برانند. با وجود مزایای گسترده، مانند بهبود وضعیت ایمنی جاده‌ها، کاهش تصادفات و ایجاد زمینه جدید برای رشد اقتصادی، این فناوری خطرات بالقوه زیادی را در پی دارد. بسیاری از وسایل نقلیه خودران در برخی شرایط به راننده انسان نیاز دارند تا کنترل آن را در دست گیرد. به این فرد «کاربر مسئول»^۱ گفته می‌شود. این کاربر باید تمام شرایط یک راننده را داشته باشد. کاربر مسئول در رانندگی پویا باید در برابر تخلفات رانندگی یا مجازات مدنی مصون باشد؛ ولی مسئولیت‌هایی همچون بیمه در آن به قوت خود باقی است. برای اینکه یک وسیله نقلیه خودران تلقی شود باید کل وظیفه رانندگی پویا را انجام دهد که شامل کنترل فرمان، ترمز، تشخیص و واکنش در برابر اشیاء و رویدادهاست. این نوع وسیله نقلیه باید بتواند محیط رانندگی را کنترل کند و به سایر کاربران در جاده واکنش نشان دهد. استدلال کمیسیون برای این مصونیت، عدم تقصیر و عدم اشتباه در کاربر مسئول است؛ بدین شرح که هنگامی که کنترل همه‌جانبه با سیستم رانندگی خودکار است، اشتباهی متوجه کاربر نیست. در مقررات فرانسه نیز رویکرد مشابهی وجود دارد. حال اگر کاربر عمداً در عملکرد سیستم رانندگی خودران دخالت کند، این مصونیت از بین می‌رود (Law Commission of England and Wales & Scottish Law Commission, 2022). برای جلوگیری از تلفات جانی و مالی، لازم است ایمنی وسایل نقلیه خودران قبل از استقرار گسترده آن‌ها، تأیید و اعتبارسنجی شود (Moss et al., 2021, p. 1). در سال ۲۰۱۹ در فلوریدا یک خودروی تسلا و کامیون دچار سانحه شدند که به فوت راننده خودروی تسلا منجر شد. طبق بررسی‌های انجام شده، هر دو راننده از لحاظ جسمی سلامت کامل داشتند و هیچ نوع الکل و مواد مخدري مصرف نکرده بودند و گواهینامه‌های لازم را نیز داشتند. در این سانحه، سیستم‌های هشدار خودرو و ترمز اضطراری خودکار قبل از تصادف فعال نشده بود و راننده تسلا بیش از حد استاندارد به خودکار بودن خودرو تکیه کرده بود. این درحالی است که مالک تسلا در دفترچه راهنما بارها در مورد عدم تکیه کامل به اتوماتیک بودن هشدار داده بود. شرکت تسلا گزارشی داد که این خودرو برای عبور از ترافیک یا جلوگیری از تصادف در سرعت‌های بالا طراحی نشده است. هیئت ملی ایمنی حمل‌ونقل ایالات متحده علت تصادف را چنین تشخیص داد: عبور راننده

1. User in Charge

کامیون و ندادن حق تقدم به خودرو، بی توجهی راننده خودرو به دلیل اتکای بیش از حد به سیستم خودکار رانندگی، محدودیت در طراحی سیستم خودکار توسط شرکت تسلا. یکی دیگر از عواملی که در این تصادف نقش داشت، ناکامی اداره ملی ایمنی ترافیک در توسعه روشی برای تأیید ادغام سیستم حفاظتی قابل قبول توسط سازندگان وسیله نقلیه خودکار بود که استفاده از سیستم‌های کنترل خودکار خودرو را به شرایطی خاص محدود می‌کرد (National Transportation Safety Board, 2019).

شرکت سازنده خودران‌ها علاوه بر آن که مسئولیت اخذ مجوزهای مربوطه را دارد، برای شناسایی مسبب اصلی زیان نیز اولین گزینه‌ای است که باید بررسی شود. طبق ادعای شرکت تسلا، خودروهای این شرکت قابلیت لازم یک خودران را دارند؛ اما شایان توجه است که باید در مورد تمامی تبلیغات راجع به فناوری‌های هوش مصنوعی با آگاهی نکته‌سنجی کرد؛ چراکه برخلاف ادعاهای تسلا، از این شرکت به دلیل تبلیغات گمراه‌کننده خود شکایت شده است.^۲ براساس یک پژوهش، وسایل خودران برای ارائه اطلاعات موقعیت‌یابی و ناوبری به «سیستم جهانی ناوبری ماهواره‌ای»^۳ متکی هستند و به همین علت، در مناطق پرترافیک با تراکم داده‌ها و در نتیجه تداخل یا انسداد داده‌ها روبه‌رو می‌شود. این آثار باعث اختلال سیستم و قطعی سرویس خودران‌ها می‌شود. استفاده از حسگرهای ارزان‌قیمت در خودران‌ها نیز در این اختلال عملکردی مزید بر علت می‌شود و در این مورد نمی‌توان طراح یا سازنده نرم‌افزار را مسئول دانست (Parsa, 2024, pp. 120-121). به نظر می‌رسد در این استدلال، باید قائل به تفکیک شد؛ فرض کنید طراح یا سازنده خودران مدعی پیشرفته بودن آن در همه شرایط ترافیک باشد، در حالی که به ضعف عملکردی سیستم یا ناتوانی حسگرها هیچ اشاره‌ای نکرده است. در این صورت، چه بسا علاوه بر مسئولیت مدنی، مسئولیت کیفری نیز متوجه او خواهد بود. در فرض دیگر، اگر طراح و سازنده تذکرات لازم را در قالب دستورالعمل، ضعف عملکردی سیستم را به کاربران ارائه دهد، اما کاربر به طور کامل به خودران بودن سیستم متکی باشد و حادثه‌ای رخ دهد که ناشی از بی احتیاطی او نسبت به دستورالعمل‌های ارائه شده باشد، در این حالت طراح و سازنده مسئولیتی نخواهند داشت. در سال ۲۰۱۶ راننده‌ای در فلوریدا با تکیه بر اطمینانی که نماینده فروش تسلا برای عملکرد رانندگی خودکار خودرو داده بود، در حال استراحت در خودروی خودران بود که با خودرویی که به طور معمول در راه متوقف شده بود، تصادف کرد و راننده تسلا دچار جراحت دائمی شد. دلیل خریداری این خودرو توسط راننده تسلا این بود که تسلا در تبلیغات خود مدعی شده بود سیستم رانندگی خودکار این خودرو قابلیت تشخیص خطر و توقف را دارد.^۴ با توجه به این که استفاده از خودروهای خودران خطرناک‌تر و تهدیدآمیزتر از خودروهای دیگر هستند (Freitas et al., 2023, p. 17) در نتیجه استفاده از وسیله نقلیه خودران موقعیتی خطرناکی را برای کاربر و دیگران ایجاد می‌کند. در یکی از مطالعات انجام شده بر مبنای الزامات اساسی، همچون جامعه، فرد، فن و دانش، پیش‌نیازهایی را برای خودروهای خودران طبقه‌بندی کرده‌اند که عبارت‌اند از: ۱. بسترسازی جاده‌ها؛ ۲. زیرساخت فنی و تجهیزات؛ ۳. دانش؛ ۴. کاربران؛ ۵. نقش مدیران و ۶. فرهنگ قوانین (Mahmoudi et al., 2024, p. 103).

اخیراً پرونده‌های زیادی علیه سیستم‌های هوش مصنوعی به جریان افتاده است که در این میان، شکایات علیه خودروهای خودران سهم درخور توجهی دارد.

سال ۲۰۲۳ پرونده‌ای طرح شد که طبق آن، در سال ۲۰۱۹ در میامی یک کامیون با خودروی تسلا مدل ۳ تصادف کرد و به فوت راننده تسلا منجر شد. در این پرونده، همسر متوفی علیه تسلا شکایتی را طرح کرد و قاضی دادگاه شواهد منطقی پیدا کرد که نشان می‌داد ایلان ماسک

1. <https://tesla.com> (2019)

2. <https://electrek.co/>

3. Global Navigation Satellite System (GNSS).

4. Hudson v. Tesla, Inc., Circuit Court of the Ninth Judicial Circuit, 2018-CA-011812-O, available at: <https://blogs.gwu.edu/law-eti/ai-litigation-database-search/case-detail-page/?pid=26>

blogs.gwu.edu (2019), Banner, Kim v. Tesl, the Fourth District Court of Appeal State of Florida, Case No. 4D2023-303, 2019,(2024,7,19): <https://blogs.gwu.edu/law-eti/ai-litigation-database-search/case-detail-page/?pid=150>

و سایر مدیران تسلا با علم به معیوب و ناقص بودن فناوری خودران، خودروهای خود را تبلیغ کردند و مجوز دادند که خودروها به شیوه‌ای ناایمن رانندگی شوند. قاضی عنوان کرد: «تسلا درگیر یک راهبرد بازاریابی است که محصولات را خودکار نشان دهد و اظهارات عمومی ایلان ماسک در مورد اعتماد به این فناوری، تأثیر بسزایی را در باور قابلیت‌های محصول داشته است»^۱.

خودران مذکور در این پرونده از نوع سیستم پرخطر است؛ چراکه شرکت سازنده با ترفندهای تبلیغاتی سعی در تمام هوشمند نشان دادن خودران داشته است.

در سال ۲۰۲۴، شکایتی علیه تسلا در ۱۱ ایالت توسط ۲۳ شاکی مطرح شد. آن‌ها ادعا می‌کنند نقص فنی سخت‌افزاری و نرم‌افزاری در مدل‌های (۲۰۱۳S) تا (۲۰۲۰X) باعث شتاب‌گیری ناخواسته می‌شود. همچنین گزارشی وجود دارد مبنی بر اینکه یک تسلا مدل ۳ به‌طور ناگهانی شتاب گرفته و پس از خروج از بزرگراه، با یک ساختمان اداری برخورد کرده و به فوت یک کارمند منجر شده است^۲.

در زمان نگارش این مقاله، رأی برای این شکایات صادر نشده است؛ اما شکایات مشابه با مشکل فنی مذکور به دفعات زیادی علیه تسلا مطرح شده است. در بررسی سیستم‌های هوشمند علاوه بر بررسی‌های فنی و دستورالعمل‌های ارائه شده توسط تولیدکنندگان، باید به نحوه تبلیغات محصول نیز دقت لازم را داشت. در حالی که بروشور یک محصول بر نظارت و کنترل راننده بر فرمان تأکید دارد، فیلم تبلیغاتی آن راننده را در حال خواب یا خوردن و آشامیدن نشان می‌دهد. این پیام متناقض به مشتری القا می‌کند که خودران کاملاً هوشمند است و هیچ نیازی به نظارت راننده ندارد. در چنین حالتی، باید مسئولیت تولیدکنندگان را نیز در نظر داشت و صرفاً به تذکرات دستورالعمل‌ها اکتفا نکرد تا صرف ارائه یک دستورالعمل دستاویزی برای فرار از مسئولیت نباشد و تولیدکننده از تبلیغات اغواگرانه اجتناب کند؛ زیرا تبلیغات اغواگرانه در عرصه سیستم‌های هوش مصنوعی خسارات جبران‌ناپذیری را بر جای می‌گذارد.

یکی دیگر از فناوری‌های هوش مصنوعی، که به تازگی توجه بسیاری از شرکت‌ها و اشخاص را به خود جلب کرده است، «متاورس»^۳ است. متاورس به معنای پلتفرمی جدید اما سه بعدی، جنبشی و تا حدودی قابل لمس است و گام بعدی در توسعه فناوری دیجیتال است (Florid, 2022, p. 2). در متاورس، قادر خواهیم بود دنیای فیزیکی را، که شامل افراد مختلف و مکان‌ها و اشیاء است، از طریق پلتفرم‌های واقعی توسعه یافته (واقعیت گسترده)^۴ ملاحظه کنیم و تعامل داشته باشیم. واقعیت گسترده یا واقعیت متقابل طیفی را نشان می‌دهد که هریک از فناوری‌ها در محیطی کاملاً مجازی و تولیدشده توسط یارانه در سطح بالا و در طرف دیگر، واقعیت فیزیکی (بدون مجازی سازی) قرار دارد (Pimentel et al., 2022, pp. 3-4). طبق گزارش ارائه شده از سوی سرویس تحقیقات پارلمان اروپا، متاورس فرصت‌ها و خطراتی را به همراه دارد که گستره آن هنوز مشخص نیست. در برخی از مسائل، خط‌مشی و پیامدهای احتمالی آن در حوزه‌هایی از جمله رقابت تجاری، حفاظت از داده‌ها، تراکنش‌های مالی، امنیت سایبری، سلامت، دسترسی و دامنه فراگیری آن شناسایی شده است (Madiega et al., 2022). پیامدهای بالقوه استفاده از فناوری متاورس در سایر حوزه‌ها نیز مشخص شده است؛ برای مثال، تأثیر متاورس در حوزه سلامت و بهداشت بدین گونه است که پزشکان می‌توانند در محیطی مجازی آموزش ببینند و اطلاعات بیومتریک بیماران را در دستگاه‌های عینک‌مانند، که به‌نحو واقعیت گسترده یا واقعیت افزوده عمل می‌کنند، مشاهده کنند (Zhu, 2022).

مفهوم متاورس حداقل از سال ۱۹۹۲ استفاده شده است و در حال حاضر به طور کلی به مفهوم یک دنیای مجازی فراگیر و پایدار اشاره دارد که در آن کاربران می‌توانند با سایر کاربران و محیط اطراف ارتباط و تعامل داشته باشند و در فعالیت‌های اجتماعی مشابه تعاملات دنیای فیزیکی شرکت کنند. مفهوم متاورس پس از اقدامات مارک زاکربرگ (مدیر عامل فیسبوک)، به ویژه با تغییر نام فیسبوک به «متا» در اکتبر ۲۰۲۱،

1. blogs.gwu.edu (2019), Banner, Kim v. Tesl, the Fourth District Court of Appeal State of Florida, Case No. 4D2023-303, 2019,(2024,7,19): <https://blogs.gwu.edu/law-eti/ai-litigation-database-search/case-detail-page/?pid=150>

2.blogs.gwu.edu (2020),Inkie Lee v. Tesla, Inc., (2:20-c-v-00570), District Court, C.D. California, (2024,7,19): <https://blogs.gwu.edu/law-eti/ai-litigation-database-search/case-detail-page/?pid=126>

3. Metaverse

4. Extended Reality (XR)

بیشترین توجه را به خود جلب کرد. یکی از ویژگی‌های ضروری در فناوری دنیای مجازی، واقع‌گرایی روان‌شناختی است که این ویژگی بر توانایی فناوری متاورس برای انتقال کاربر از دنیای فیزیکی به دنیای مجازی تمرکز دارد. در طول توسعه متاورس سه نکته کلیدی مطرح می‌شود که عبارت‌اند از: مقبولیت اجتماعی، امنیت و حریم خصوصی و مسئولیت‌پذیری (Rawal et al., 2022, pp. 4-14). هر شخصی در متاورس دارای «آواتار» خواهد بود. هنگامی که کاربران از طریق آواتارهای خود تعامل دارند، ممکن است حوادثی رخ دهد که در دنیای فیزیکی این نوع حوادث به مسئولیت و جبران خسارت منجر می‌شود. در این شرایط، کاربران انتظار دارند که حقوق آواتارهایشان در متاورس محافظت شود. یکی از راه‌های حل این مشکل، مسئول دانستن آواتارها در قبال اعمالشان است (Cheong, 2022, p. 470).

استدلال موجود برای اعطای شخصیت حقوقی به آواتارها در فصل مربوط به اعطای شخصیت حقوقی به هوش مصنوعی بررسی شد. با توجه به توضیحات ارائه شده، استفاده از فناوری متاورس موقعیت‌های خطرناکی را برای افراد ایجاد می‌کند که تبعات جانی و مالی فراوانی دارد. با توجه به سطح پیشرفته هوش مصنوعی، که در این نوع فناوری به کار رفته، اثبات علت ورود زیان از طرف زیان‌دیده بسیار دشوار و گاه ناممکن است؛ در نتیجه در این نوع از فناوری هوشمند، مبنای قرار دادن نظریه خطر بهترین راه حل خواهد بود.

برای یافتن مبنای مناسب در حوزه هوش مصنوعی، باید سطح و میزان خطر در سیستم‌های هوشمند بررسی شود. در سیستم‌هایی که سطح و خطرات آن‌ها پایین است، مانند سیستم ترجمه گوگل، سیری و امثال این‌ها، مبنای تقصیر مناسب است؛ ولی در هوش مصنوعی با سطح پیشرفته و تهدیدآمیزتر، مانند ربات انسان‌نما، خودروهای خودران و فناوری متاورس، مبنای خطر راه حل مناسبی است؛ زیرا اگر فقط مبنای خطر را معیار قرار دهیم، در نوآوری و پیشرفت فناوری تأثیر منفی می‌گذارد و طراحان به دلیل فرض مسئولیت، انگیزه‌ای برای ارائه محصولات خود نخواهند داشت. در مقابل، در صورت مبنای قرار دادن نظریه تقصیر در همه شرایط، حق زیان‌دیده احقاق نمی‌شود. علاوه بر این، مبنای خطر موجب افزایش دقت و کیفیت در طراحی و تولید سیستم در تولیدکنندگان و طراح آن می‌شود؛ بنابراین لازم است مبنایی نسبی برای تعیین مسئولیت در نظر گرفته شود. در واقع، در حال حاضر رژیم‌های مسئولیت کنونی موجود تا جایی از زیان‌دیدگان حمایت بنیادین می‌کنند که ویژگی‌های خاص فناوری‌های نوظهور، مانند هوش مصنوعی، در نظر گرفته شود؛ بنابراین به جای در نظر گرفتن اصول جدید برای مسئولیت که اعمال آن‌ها مستلزم اصلاح رژیم فعلی مسئولیت است، می‌توان به انطباق دو مبنای موجود مسئولیت پرداخت (Yaniv & Justin, 2020, p. 196).

برای اعمال هر کدام از نظریه تقصیر و خطر باید به ادعای سازندگان آن سیستم و نحوه عملکرد آن دقت کرد که این مورد را قانون‌گذار باید مدنظر داشت باشد. برای مثال خودران‌ها از سطح یکسانی از هوشمند بودن برخوردار نیستند؛ بنابراین نمی‌توان به طور مطلق عنوان کرد که خودران‌ها بر مبنای نظریه تقصیر یا خطر دارای مسئولیت هستند. اگر یک سازنده خودرانی ادعا کند که این خودرو کاملاً هوشمند است و نیازی به دخالت و نظارت راننده ندارد، در این صورت این مدل از خودران جزو سیستم‌های پرخطر بوده و مبنای خطر بر آن جاری می‌شود؛ اما اگر خودران به عنوان سیستمی معرفی شود که نیازمند نظارت راننده است؛ در این صورت مبنای تقصیر بر این نوع خودران مناسب است. در حوزه پزشکی اگر سیستمی طراحی شود که به طور مستقیم با جان انسان مرتبط باشد، یعنی دستگاهی طراحی شود که به طور هوشمند به عنوان جراحی اقدام به جراحی کند، در این صورت برای این سیستم مبنای خطر جاری است؛ ولی اگر سیستمی باشد که به عنوان یک سیستم تسهیل امور باشد و به طور مستقیم اقدامی انجام ندهد، در زمره مواردی خواهد بود که مبنای تقصیر بر آن اعمال می‌شود.

در نتیجه دو مبنای سنتی تقصیر و خطر به طور کامل و مطلق پاسخ‌گوی سیستم‌های هوش مصنوعی نیستند و براساس بررسی هر سیستم به طور مجزا، مبنای متناسب تعیین می‌شود. همچنین برای جبران خسارات وارده از طریق سیستم‌های هوشمند پیشنهادی ارائه شد تا برای سیستم حسابی افتتاح شود تا درصدی از عواید آن به طور سالانه در حساب مذکور واریز شود که این مبلغ را صرفاً می‌توان برای جبران خسارت و ارتقای کیفیت سیستم به کار گرفت. در چنین تفسیری نظریه نوینی تحت عنوان «نظریه تحلیل اقتصادی» مطرح می‌شود.

۲-۳. نظریه تحلیل اقتصادی

نظریه تحلیل اقتصادی رویکردی جدید در مسئولیت مدنی بود که از دهه ۱۹۶۰ مورد توجه قرار گرفت. در این نظریه، به اثر قواعد مسئولیت مدنی و نتیجه اعمال آن بر رفاه اجتماعی و کارآمدی توجه ویژه شده است. در واقع هدف از مسئولیت مدنی ارتقای رفاه اجتماعی از سه طریق ایجاد انگیزه برای کاهش حوادث و خطرها، تحمیل و استناد درست و مناسب خسارات ناشی از حوادث و در آخر، کاهش هزینه‌های اداری نظام مطالبه خسارت است. نظریه تحلیل اقتصادی از دو نظریه سنتی تقصیر و خطر استفاده می‌شود. اگر نظریه تقصیر مبنا قرار گیرد، شخص در صورتی مسئول است که احتیاط لازم را اعمال نکرده باشد و چنانچه نظریه خطر مبنا قرار گیرد، شخص مسئول همه خسارات ناشی از فعالیت خویش است؛ در نتیجه تمام تلاش خود را در راستای رعایت احتیاط‌های لازم به کار می‌گیرد تا از بروز حادثه جلوگیری کند و چنانچه هزینه انجام دادن اقدامات احتیاطی بیش از منفعت حاصل از فعالیت باشد، سطح فعالیت کاهش خواهد یافت و به تناسب آن، احتمال بروز حادثه نیز کم خواهد شد (Safai & Rahimi, 2019, pp. 69-70).

مسئولیت مدنی یکی از مهم‌ترین مسائل حقوقی است که می‌توان آن را با سنجش کارایی اقتصادی بررسی کرد. با بررسی نظرات مختلف در این موضوع، می‌توان به الگویی مناسب و کارآمد دست یافت. مسئولیت مدنی با کارآمدی اقتصادی اهدافی را از قبیل بازدارندگی، درونی کردن هزینه‌های خارجی و توزیع ضرر دنبال می‌کند و می‌توان این سه شیوه را برای دستیابی به مسئولیت مدنی با تحلیل کارآمدی اقتصادی بیان کرد. نظریه‌های مختلفی از جمله وضعیت برتر پارتو، وضعیت بهینه پارتو، حداکثرسازی مطلوبیت و حداکثرسازی ثروت در این زمینه وجود دارد. نظریه اخیر کارایی بیشتری به نسبت به سایر نظریه‌ها دارد. در این نظریه، براساس حوادث یک‌جانبه و دوجانبه مبنای تقصیر و خطر سنجیده شده و مبنای مناسبی تعیین می‌شود. به‌طورکلی در حوادثی که به‌طور یک‌جانبه بوده و نیازمند احتیاط از یک طرف است انتخاب مبنای تقصیر و مبنای خطر به‌نحو مطلوبی بهینه اجتماعی را تأمین می‌کند؛ اما در حوادثی که نیازمند به‌کارگیری و رعایت اقدامات احتیاطی است و همچنین در شرایطی که زیان‌زندگان متعدد باشند، مبنای تقصیر مناسب‌تر است (Razazzadeh et al., 2024, p. 15).

همان‌طور که در مباحث فوق بررسی شد، مبنا قرار دادن هریک از نظریه‌های تقصیر و خطر به‌طور مطلق در سیستم هوش مصنوعی امکان‌پذیر نیست و باید میزان خطر و تأثیر هر سیستمی به‌طور مجزا بررسی شده تا مشمول هریک از نظریه‌های تقصیر و خطر شود. همچنین عنوان شد که یکی از دلایل چنین بررسی‌هایی نبود موانع در مسیر پیشرفت فناوری‌های نوین است. از دلایل دیگر، می‌توان به ارتقای رفاه اجتماعی اشاره کرد. در واقع چنانچه تولیدکنندگان برحسب میزان خطر و تبعاتی که بر جامعه و دیگران تحمیل می‌کنند، دارای مسئولیت بوده و باید کیفیت سیستم‌هایی را که طراحی و تولید می‌کنند افزایش دهند و از طرفی مصرف‌کنندگان نیز باید از علم و اطلاعات قابل قبولی برخوردار باشند تا بتوانند از سیستم‌های هوش مصنوعی به‌نحو صحیح و مطلوبی استفاده کنند. افزایش کیفیت در سیستم‌ها و کاهش خطرات آن‌ها از یک‌سو و افزایش آگاهی و علم عمومی در جامعه از سوی دیگر موجب ارتقای رفاه اجتماعی می‌شود. در نظریه تحلیل اقتصادی، در صورتی که زیان‌زندگان متعدد باشند، مبنای تقصیر اعمال می‌شود؛ اما برای سیستم هوش مصنوعی چنین دیدگاهی مناسب نیست؛ زیرا به‌طورکلی در طراحی سیستم‌های هوشمند، تعداد عوامل دخیل بیش از دو نفر است که هریک از آن‌ها دارای علم تخصصی هستند. به همین دلیل اثبات تقصیر این عوامل برای عموم مردم دشوار است؛ در نتیجه، برای تعیین مبنای مناسب مسئولیت مدنی ناشی از هوش مصنوعی، باید به تناسب سطح سیستم و میزان هوشمندی آن توجه کرد. همچنین، لازم است تحلیل‌های اقتصادی و هموار کردن مسیر پیشرفت را نیز در نظر گرفت تا در کنار جبران خسارت، رفاه اجتماعی و فناوری ارتقا یابد.

نتیجه‌گیری

با پیشرفت علم و فناوری در حوزه هوش مصنوعی، که امروزه جوانب مختلفی از زندگی اشخاص را در بر گرفته است، چالش‌های حقوقی زیادی به وجود می‌آید. حقوق به‌منزله علمی پویا باید براساس مقتضیات زمانی و مکانی پاسخ مناسبی برای این چالش‌ها داشته باشد. برای انتساب

خسارات وارده از سیستم‌های هوشمند به آن‌ها، باید به این نوع سیستم‌ها شخصیت حقوقی اعطا شود. در سطح جهان، پدیده‌هایی همچون معابد هند و اکوسیستم آکوادور دارای شخصیت حقوقی هستند؛ بنابراین هوش مصنوعی با این ابعاد و نفوذ در همه جنبه‌های زندگی اجتماعی و شخصی به طریق اولی باید دارای شخصیت حقوقی باشد. علاوه بر آن، مثالی که برای لزوم اعطای شخصیت حقوقی به هوش مصنوعی بیان می‌شود، مثال شخصیت حقوقی شرکت‌هاست؛ با این توضیح که شخصیت شرکت‌ها صرفاً به دلیل سیستم و دفتر اداری بی‌جان نیست؛ بلکه به دلیل تبعات آن در نظام حقوقی و جامعه است. در اعطای شخصیت حقوقی به این نوع سیستم‌ها، به قانون‌گذار چنین پیشنهاد شد که با توجه به سطح سیستم و میزان خطری که برای جامعه ایجاد می‌کند، دو نوع مسئولیت تضامنی و نسبی را برای سیستم‌های هوش مصنوعی بپذیرد؛ یعنی همانند شرکت‌های موجود، دارای شخصیت حقوقی و قالب‌های مختلفی (تضامنی و نسبی) از مسئولیت باشند به منظور فراهم کردن تمکن مالی و دارایی برای سیستم ثبت شده عوامل دخیل که اقدام به ثبت می‌کنند حسابی را برای هوش مصنوعی افتتاح کرده تا درصدی از عواید سیستم در این حساب واریز شود. این مبالغ صرفاً در راستای جبران خسارات احتمالی و ارتقای کیفیت سیستم صرف می‌شود. همچنین می‌توان از رویکرد عدالت توزیعی نیز تأثیر پذیرفت و در مواردی که نتوان به هیچ طریقی سبب زیان را شناسایی و خسارت را جبران کردن، شرکت‌های بیمه و صندوق‌های حمایتی از طرف دولت، نقش خود را به طور مؤثر ایفا کنند. هدف اصلی اعطای شخصیت حقوقی به هوش مصنوعی در راستای تسهیل اثبات خسارات و درصدد برآمدن جبران آن است و یکی از دلایل مهم، افزایش اقدامات احتیاطی و نظارت و کیفیت بر سیستم‌های هوشمند است.

در نتیجه بررسی سه نظریه تقصیر، نظریه خطر و نظریه تحلیل اقتصادی و در راستای جمع اهدافی همچون جبران خسارت، افزایش کیفیت سیستم‌های هوش مصنوعی، ارتقای رفاه اجتماعی و پیشرفت علم و فناوری، با هدف تعیین مبنایی که برای نیل به همه اهداف مذکور مناسب باشد، باید هر سیستمی به لحاظ فنی و اقتصادی به طور مجزا بررسی شود و برحسب میزان خطری که سازندگان سیستم در دستورالعمل‌ها یا تبلیغات محصول ارائه می‌دهند شامل مبنای تقصیر یا مبنای خطر شود. در واقع یک مبنایی نسبی با ملاحظات فنی و اقتصادی اعمال می‌شود. در سال‌های اخیر، دستورالعمل‌های گوناگونی برای سازمان‌دهی فعالیت‌های هوش مصنوعی در عرصه‌های مختلف، مثل خودران‌ها، ربات‌های جراح، متاورس و سایر حوزه‌ها، توسط اتحادیه اروپا منتشر شده است؛ اما در نهایت در سال ۲۰۲۴ «قانون هوش مصنوعی» به تصویب اتحادیه اروپا رسید تا چهارچوبی برای فعالیت سیستم‌های هوش مصنوعی تدوین شود. با تعیین گروه ناظر پیش‌بینی شده در این قانون، وظایف عوامل دخیل در هوش مصنوعی ارزیابی شده و تشخیص تقصیر هر کدام از بازیگران از طراح و تولیدکننده تا کاربر با دقت بیشتری صورت می‌گیرد. با مطالعه پرونده‌های موجود در این حوزه، به ویژه خودروهای خودران و ربات‌های جراحی، مشخص می‌شود که در دادگاه‌ها نیز نقش و وظایف هرکدام از عوامل با جزئیات بررسی قرار می‌شود و در نهایت اثبات شده است که تاکنون جنبه تبلیغاتی در تعیین مسئولیت هوش مصنوعی بسیار مؤثر بوده و فراتر از عملکرد واقعی این محصولات است.

نقدی که بر «قانون هوش مصنوعی» و تمامی قوانین و دستورالعمل‌های موجود وارد است این است که باید در تدوین قانون مواردی همچون تعریف و تبیین اصطلاحات موجود و اشخاص دخیل به طور دقیق و کامل و تعیین نوع مسئولیت آن‌ها رعایت شود و همچنین تدوین قوانین باید به صورت تخصصی در رشته خاص با همکاری متخصصان صورت گیرد.

منابع

- امامی، حسن (۱۳۹۵). حقوق مدنی. ج ۳۷، ج ۱، تهران: انتشارات اسلامی.
- ابوذی، مهنوش (۱۴۰۱). حقوق و هوش مصنوعی. ج ۲، تهران: نشر میزان.
- بابانی، ایرج (۱۴۰۱). حقوق مسئولیت مدنی و الزامات خارج از قرارداد، ج ۳. تهران: میزان.
- بادینی، حسن (۱۳۸۴). فلسفه مسئولیت مدنی. ج ۱. تهران: شرکت سهامی انتشار.
- برزگر، محمدرضا و الهام، غلامحسین (۱۳۹۹). مسئولیت کیفری کاربر خودرو خودران در قبال صدمات وارده توسط آن، پژوهش حقوق کیفری، ۲۰۱-۲۰۲، ۲۲۰۹-۲۲۰۹. <https://doi.org/10.22054/jclr.2020.39009.1838>
- پارسا، ناهید (۱۴۰۳). بررسی مسئولیت مدنی ناشی از نقض نرم‌افزار در خودروهای خودران با نگاهی به حقوقی آمریکا. مطالعات حقوق تطبیقی معاصر، ۳۴(۸۵)، ۱۱۵-۱۴۷. <https://doi.org/10.22034/law.2024.55488.3245>
- تخشید، زهرا (۱۴۰۰). مقدمه‌ای بر چالش‌های هوش مصنوعی در حوزه مسئولیت مدنی. حقوق خصوصی، ۱۸(۱)، ۲۲۷-۲۵۰. <https://doi.org/10.22059/jolt.2021.319529.1006965>
- حکمت‌نیا، محمود، محمدی، مرتضی و واثقی، محسن (۱۳۹۸). مسئولیت مدنی ناشی از تولید ربات‌های مبتنی بر هوش مصنوعی خودمختار. حقوق اسلامی، ۶(۶۰)، ۲۳۱-۲۵۸. [لینک](#)
- ذاکری‌نیا، حانیه (۱۴۰۲). ماهیت و مبنای مسئولیت مدنی ناشی از هوش مصنوعی در حقوق ایران و کشورهای اتحادیه اروپا. حقوق خصوصی، ۲۰(۱)، ۱۳۵-۱۵۲. [لینک](#)
- ذاکری‌نیا، حانیه و غلامپور، زهرا (۱۴۰۳). الگوریتم‌های معقول و متعارف و تقویت نظریه قابلیت انتساب مسئولیت مدنی هوش مصنوعی. حقوق فناوری‌های نوین، ۵(۹)، ۱۶۸-۱۵۵. <https://doi.org/10.22133/mtlj.2023.406159.1224>
- راسل، استوارت و نوروگ، پیتر (۱۳۹۹). هوش مصنوعی (رهیافتی نوین). ترجمه عین‌الله جعفرنژاد قمی، بابل: علوم بارانه‌ای.
- رزازاده، مهنوش، نصیران نجف‌آبادی، داود و سلطانی، رضا (۱۴۰۳). کارآمدی اقتصادی به عنوان مبنای مسئولیت مدنی. مطالعات فقه اقتصادی، ۳(۳)، ۱-۱۶. <https://www.doi.org/10.22034/ejs.2024.430978.1635>
- رهبر، نوید و دهقانپور فراشاه، سبحان (۱۴۰۰). بررسی تطبیقی مبنا مسئولیت مدنی در تصادفات وسایل نقلیه خودران. مطالعات حقوق تطبیقی، ۲۲(۲)، ۵۲۳-۵۴۳. <https://doi.org/10.22059/jcl.2021.320449.634169>
- رجبی، عبدالله (۱۳۹۸). ضمان در هوش مصنوعی. مطالعات حقوق تطبیقی، ۲۰(۲)، ۴۴۹-۴۶۶. <https://doi.org/10.22059/jcl.2019.274782.633787>
- ژودن، پاتریس (۱۳۹۴). اصول مسئولیت مدنی. ترجمه مجید ادیب. ج ۴. تهران: میزان.
- سارتور، گیواندی و برنتینگ، کارل (۱۳۹۹). کاربردهای قضائی هوش مصنوعی. ترجمه مهنوش ابوذر و محمد سعید شفیعی، ج ۱. تهران: میزان.
- صفایی، حسین و رحیمی، حبیب‌الله (۱۳۹۸). مسئولیت مدنی خارج از قرارداد. ج ۱۲. تهران: سمت.
- صفایی، حسین و قاسم‌زاده، مرتضی (۱۳۹۶). حقوق مدنی؛ اشخاص و محجورین. ج ۲۳. تهران: سمت.
- کاتوزیان، ناصر (۱۴۰۰). دوره مقدماتی حقوق مدنی؛ وقایع حقوقی - مسئولیت مدنی. تهران: انتشارات گنج دانش.
- کاظمی، محمود (۱۳۹۱). اثر فعل زیان‌دیده بر مسئولیت مدنی عامل زیان. دیدگاه‌های حقوق قضایی، ۵۷(۱۷)، ۷۹-۱۰۴. [لینک](#)

- لی شین، یانگ و فنایی، مهنوش (۱۴۰۲). وضعیت حقوقی ربات‌های هوشمند در حقوق چین. حقوق فناوری‌های نوین، ۴۵(۸)، ۲۴۳-۲۲۷.
<https://doi.org/10.22133/mtlj.2023.387130.1175>
- محمودی، محمد، جعفری، محمد و پیشدار، مهسا (۱۴۰۳). شناسایی کاربردها و الزامات به‌کارگیری هوش مصنوعی در محصولات نوین خودرویی. مطالعات مدیریت کسب و کار هوشمند، ۱۲(۴۷)، ۷۹-۱۰۹. <https://doi.org/10.22054/ims.2023.72410.2292>
- نصرارکن، حجت، محمدرضاپور اوزان، بابک و مولایی، یوسف (۱۳۹۸). بررسی جایگاه عدالت توزیعی در مسئولیت مدنی. تحقیقات حقوقی تطبیقی ایران و بین الملل، ۱۲(۴۵)، ۳۱۱-۳۳۷. [لینک](#)
- واثقی، محسن (۱۳۹۹). امکان‌سنجی اعطای شخصیت حقوقی به ربات‌های هوشمند با تکیه بر مصوبه اتحادیه اروپا: شخص الکترونیک، مجلس و راهبرد، ۲۷(۱۰۳)، ۳۰۷-۳۳۳. [لینک](#)
- ولی‌پور، علی و اسماعیلی، محسن (۱۴۰۰). امکان‌سنجی مسئولیت مدنی هوش مصنوعی عمومی ناشی از ایجاد ضرر در حقوق مدنی، اندیشه معاصر حقوقی، ۲(۶)، ۸-۱. <https://doi.org/10.22034/lth.2021.248596>
- Arnold, Z., & Toner, H. (2021). *AI Accidents: An Emerging Threat; What Could Happen and What to Do*. Center for Security and Emerging Technology (CSET). [Link](#)
- European Parliament. (2020). *Artificial Intelligence and Civil Liability*. [Link](#)
- European Parliamentary Research Service. (2022). *Artificial intelligence in Healthcare: Applications, Risks, and Ethical and Societal Impacts*. European Union. [Link](#)
- Law Commission of England and Wales & Scottish Law Commission. (2022). *Automated Vehicles: Joint Report*. [Link](#)
- Barfield, W. (2018). Liability for autonomous and artificial intelligence robots. *Paladyn, Journal of Behavioral Robotics*, 9, 193-203. <https://doi.org/10.1515/pjbr-2018-0018>
- Barfield, W. (2018). Liability for autonomous and artificial intelligence robots. *Paladyn, Journal of Behavioral Robotics*, 9, 193-203. <https://doi.org/10.1515/pjbr-2018-0018>
- Beranger, J. (2021). *Societal Responsibility of Artificial Intelligence*, Vol. 4, London: Wiley.
- Borges, G. (2018). New liability concepts: The potential of insurance and compensation funds. In *Liability for artificial intelligence and the internet of things* (pp. 230-250). Nomos Verlag. https://doi.org/10.1007/978-3-031-28497-7_34
- Bryson, J. J., Diamantis, M. E., & Grant, T. D. (2017). Of, for, and by the people: The legal lacuna of synthetic persons. *Artificial Intelligence and Law*, 25, 273-291. <https://doi.org/10.1007/s10506-017-9214-9>
- Rawal, B. S., Mentges, A., & Ahmad, S. (2022, August). The rise of metaverse and interoperability with split-protocol. In *2022 IEEE 23rd International Conference on Information Reuse and Integration for Data Science (IRI)* (pp. 192-199). IEEE. <https://dl.acm.org/doi/10.1109/IRI54793.2022.00051>

- Cantekin, K. (2023). *Regulation of artificial intelligence around the world*. The Law Library of Congress, Global Legal Research Directorate. [Link](#)
- Cheong, B. C. (2022). Avatars in the metaverse: Potential legal issues and remedies. *International Cybersecurity Law Review*, 3, 467-494. <https://doi.org/10.1365/s43439-022-00056-9>
- Chesterman, S. (2020). Artificial intelligence and the limits of legal personality. *International & Comparative Law Quarterly*, 69(4), 819-844. <https://doi.org/10.1017/S0020589320000366>
- Chesterman, S. (2021). *We, the robots? Regulationg artificial intelligence and the limits of the law*. Cambridge University Press
- National Transportation Safety Board. (2019). *Collision Between Car Operating with Partial Driving Automation and Truck-Tractor Semitrailer Delray Beach, Florida, March 1, 2019*. [Link](#)
- Cooke, J. (2010). *Law of Tort*, 9th Edition, London: Longman.
- Dignum, V. (2019). *Responsible Artificial Intelligence: How to Develop and Use AI in a Responsible Way*, Berlin: Springer Nature.
- European Commission: Directorate-General for Research and Innovation & European Group on Ethics in Science and New Technologies. (2018). *Statement on artificial intelligence, robotics and 'autonomous' systems*. Publications Office. <https://doi.org/10.2777/531856>
- European Parliament. (2017, February 16). *Resolution with recommendations to the Commission on Civil Law Rules on Robotics*. EUR-Lex. [Link](#)
- Digital Watch. (2023, May 8). *EU urged to speed up AI regulation, says Margrethe Vestager*. [Link](#)
- Expert Group on Liability and New Technologies. (2019). *New Technologies Formation, Liability for Artificial Intelligence and Other Emerging Digital Technologies*. European Union. [Link](#)
- Floridi, L. (2022). Metaverse: A matter of experience. *Philosophy and Technology*, 35(73), 1-7. <https://doi.org/10.1007/s13347-022-00568-6>
- De Freitas, J., Zhou, X., Atzei, M., Boardman, S., & Di Lillo, L. (2023). Public Perception and Autonomous Vehicle Liability. *Harvard Business School Research Paper Series Working*. 23-036. [Link](#)
- Giuffrida, I. (2019). Liability for AI decision-making: Some legal and ethical considerations. *Fordham Law Review*, 88(2), 439-456. [Link](#)
- Khan, Hanif. (2021). Different Types of Artificial Intelligence Systems, *University of Botswana*, available at: [Link](#)
- Wendehorst, C., Koch, B. A., van den Berg, M. D. F. B., & Vogenauer, S. B. (2022). *Response of the European Law Institute to the European Commission's Public Consultation on Civil Liability Adapting Liability Rules to the Digital Age and Artificial Intelligence*. European Law Institute. [Link](#)
- Madiega, T. (2023). *Artificial Intelligence Liability Directive*. European Parliamentary Research Service. [Link](#)

- Madiega, T., Baczynska, M., & O'Brien, J. (2022). *Metaverse: Opportunities, risks and policy implications*. European Parliamentary Research Service. [Link](#)
- Mittelstadt, B. (2021). *The impact of artificial intelligence on the doctor patient relationship*. Council of Europe. [Link](#)
- Molnár-Gábor, F. (2020). Artificial Intelligence in Healthcare: Doctors, Patients and Liabilities, *Regulating Artificial Intelligence*, Wischmeyer, T. & Rademacher, T., Switzerland: Springer Nature Switzerland AG. https://doi.org/10.1007/978-3-030-32361-5_15
- Moss, R. J., Tawfiq, J., & Takahashi, A. (2021). *Autonomous vehicle risk assessment*. Stanford Center for AI Safety. [Link](#)
- Eri, E., Neri, G., Sartori, S., Faggioni, L., & Della Gala, R. (2020). Artificial intelligence: Who is responsible for the diagnosis? *La Radiologia Medica*, 125, 517-521. <https://doi.org/10.1007/s11547-020-01135-9>
- OECD. (2023). *AI, Data Governance and Privacy: Synergies and Areas of International Co-Operation* (OECD Artificial Intelligence Papers, No. 22). OECD Publishing. [Link](#)
- Open Letter to the European Commission Artificial Intelligence and Robotics. (n.d.). [Link](#)
- Future of Life Institute. (2023, March 22). *Pause giant AI experiments: An open letter*. [Link](#)
- Pimentel, D., Fagan, J., Leone, V., & McGivney, M. (2022). *An introduction to learning in the metaverse*. Meridian Treehouse. [Link](#)
- Saudi Arabia becomes first country to grant citizenship to a robot. (2017, October 26). *Arab News*. [Link](#)
- Solum, L. B. (1992). Legal personhood for artificial intelligences. *North Carolina Law Review*, 70(4), 1231-1287. [Link](#)
- Te Awa Tupua (Whanganui River Claims Settlement) Act 2017*. (2017). [Link](#)
- Tokyo: Artificial Intelligence 'Boy' Shibuya Mirai becomes world's first AI bot to be granted residency. (2017, November 6). *Newsweek*. [Link](#)
- Tubner, G. (2006). Rights of non-humans? Electronic agents and animals as new actors in politics and law. *Journal of Law and Society*, 33(4), 497-521. [Link](#)
- Turner, J. (2021). *Robot Rules: Regulating Artificial Intelligence*, Berlin: Springer Nature.
- United States Government Accountability Office & National Academy of Medicine. (2022). *Artificial Intelligence in Health Care: Benefits and Challenges of Machine Learning Technologies for Medical Diagnostics*. [Link](#)

- Wendehorst, C. (2022). Liability for artificial intelligence: The need to address both safety risks and fundamental rights risks. In S. Voeneke, P. Kellmeyer, O. Mueller, & W. Burgard (Eds.), *The Cambridge handbook of responsible artificial intelligence: Interdisciplinary perspectives* (pp. 187-209). Cambridge University Press.
- Wischmeyer, T., & Rademacher, T. (Eds.). (2020). *Regulating artificial intelligence*. Springer Nature Switzerland AG.
- Yaniv, B., & Ferland, J. (2020). Artificial intelligence and damages: Assessing liability and calculating damages. In G. D'Agostino, A. Gaon, & C. Piovesan (Eds.), *Leading legal disruption: Artificial intelligence and a toolkit for lawyers and the law* (pp. 165-197). Thomson Reuters. [Link](#)
- Zhu, L. (2022). *The metaverse: Concepts and issues for Congress*. Congressional Research Service. [Link](#)

[In Persisn]

- Abdollah, R. (2019). Liability of Artificial Intelligence; the Reflection of Developments in the Liability Rule, *Comparative Law Studies*, 10(2), 449-466. <https://doi.org/10.22059/jcl.2019.274782.633787>
- Abouzari, M. (2022). *Law and Artificial Intelligence*. Edition:2, Tehran: Mizan Publishing.
- Emami, H. (2016). *Civil law* (37th ed.). Islamic Publications.
- Babaei, I. (2022). *Tort law* (3rd ed.). Mizan Legal Foundation.
- Badini, H. (2005). *Philosophy of civil responsibility* (1st ed.). Publishing Company.
- Barzegar, M., & Elham, G. (2020). Criminal liability of the self-driving cars for the injury caused by it. *Journal of Criminal Law Research*, 8(30), 201-229. <https://doi.org/10.22054/jclr.2020.39009.1838>
- Parsa, N. (2024). Civil liability of software defects in autonomous vehicles and the inefficiency of existing rules, with a review of US law. *Contemporary Comparative Legal Studies*, 15(34). 115-147. <https://doi.org/10.22034/law.2024.55488.3245>
- Takhshid, Z. (2021). An introductory study on the challenges of artificial intelligence in tort law. *Private Law*, 18(1), 227-250. <https://doi.org/10.22059/jolt.2021.319529.1006965>
- Hekmatnia, M., Mohammadi, M., & Vaseghi, M. (2019). Civil liability for damages caused by robots based on autonomous artificial intelligence. *Islamic Law*, 16(60), 231-258. [Link](#)
- Zakerinia, H. (2023). The nature and basis of civil liability arising from artificial intelligence in Iranian and EU members' laws. *Private Law*, 20(1), 135-152. [Link](#)
- Zakerinia, H., & Gholampour, Z. (2024). Reasonable algorithms and strengthening the "opposability" theory on the civil liability of artificial intelligence. *Modern Technologies Law*, 5(9), 155-168. <https://doi.org/10.22133/mtlj.2023.406159.1224>

- Rajabi, A. (2019). Guarantee in Artificial Intelligence. *Comparative Law Studies*, 10(2), 446-466. <https://doi.org/10.22059/jcl.2019.274782.633787>
- Russell, S. J., & Norving, P. (2020). *Artificial intelligence: A modern approach* (4th ed.). Computer Science Publications.
- Razazzadeh, M., Nassiran Najafabadi, D., & Soltani, R. (2024). Economic efficiency as the basis of civil responsibility. *Economic Jurisprudence Studies*, 6(3), 1-16. <https://www.doi.org/10.22034/ejs.2024.430978.1635>
- Rahbar, N., & Dehghanpour Farashah, S. (2021). The foundation of civil liability for accidents including autonomous vehicles: A comparative study. *Comparative Law Review*, 12(2), 523-543. <https://doi.org/10.22059/jcl.2021.320449.634169>
- Jourdain, P. (2006). *Les principes de la responsabilité civile* (4th ed.). Mizan Legal Foundation.
- Giovanni, S., & Branting, K. (2020). *Judicial application of artificial intelligence* (1st ed.). Mizan Legal Foundation.
- Safai, H., & Rahimi, H. (2019). *Civil liability (non-contractual obligations)* (4th ed.). Samt Publication.
- Safai, H., & Qasemzadeh, M. (2017). *Civil law: Persons and persons under legal incapacity* (5th ed.). Samt Publication.
- Katouzian, N. (2021). *Elementary courses of Iranian civil law: Juridical facts tort liability* (Rev. and expanded ed.). Ganjedanesh Publication.
- Kazemi, M. (2012). The effect of the injured person's act in civil responsibility. *Judicial Law Views Quarterly (Law Views)*, 17(57), 79-104. [Link](#)
- Lixin, Y., & Fanaei, M. (2023). Legal status of intelligent robots in Chinese law. *Modern Technologies Law*, 4(8), 243-227. <https://doi.org/10.22133/mtlj.2023.387130.1175>
- Mahmoudi, M., Jafari, M., & Pishdar, M. Identifying the applications and requirements of using artificial intelligence in new automotive products. *Business Intelligence Management Studies*, 12(47), 79-109. <https://doi.org/10.22054/ims.2023.72410.2292>
- Nasrarkan, H., Mohammadrezapour Ozan, B., & Molaei, Y. (2019). Investigating the position of distributive justice in civil liability. *Quarterly Journal of Iranian and International Comparative Legal Research*, 45(12), 311-337. [Link](#)
- Vaseghi, M. (n.d.). Feasibility study of granting legal personality to smart robots based on European Union resolution (electronic person -2017). *Majlis and Rahbord*, 27(103), 337-307. [Link](#)

Valipour, A., & Esmaeili, M. (2021). Feasibility study of civil liability of artificial general intelligence due to damage in civil law. *The Journal of Contemporary Legal Thought*, 3(2), 1-8.

<https://doi.org/10.22034/lth.2021.248596>

<https://indiankanoon.org/doc/1478973/> (Last Visited: 2024.02.10)

<https://www.legislation.govt.nz/act/public/2017/0007/latest/whole.html>

<https://constitutions.unwomen.org/en/countries/americas/ecuador?f%5B0%5D=provisioncategory%3A694>

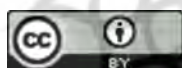
<https://eur-lex.europa.eu/>

<https://tesla.com>

<https://electrek.co/>

COPYRIGHTS

©2026 by the authors. Published by University of Science and Culture. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution 4.0 International (CC BY 4.0) <https://creativecommons.org/licenses/by/4.0/>



پژوهشگاه علوم انسانی و مطالعات فرهنگی
پرتال جامع علوم انسانی