

Submit Date: 21 January 2026
Revise Date: 17 May 2026
Accept Date: 26 May 2026
Initial Publish: 28 July 2026
Final Publish: 23 October 2027

The Encyclopedia of Comparative Jurisprudence and Law

Analysis and Examination of the Government's Civil Liability for Damages Arising from Artificial Intelligence with an Approach Based on Islamic Jurisprudence

Fahime Taherloo^{1*}

1. MA, Department of Private Law, CT.C., Islamic Azad University, Tehran, Iran

* Corresponding Author's Email: fahimehtaherloo@gmail.com

ABSTRACT

The rapid expansion of artificial intelligence and its integration into governmental structures have raised new issues concerning the government's liability for the potential damages caused by this technology. This study aims to analyze and examine the foundations and manifestations of the government's civil liability for damages arising from artificial intelligence from the perspective of Islamic jurisprudence, seeking to provide a framework for governmental accountability in this emerging field. This research is theoretical in nature and was conducted using a descriptive–analytical method. In the present study, the principles of trustworthiness, honesty, impartiality, and originality of the work were observed. The government's civil liability for damages caused by artificial intelligence can be identified in areas such as environmental damages, damages to cultural heritage, and damages related to judicial systems. The principles of justice, expediency, no-harm (La Zarar), and destruction (Itlaf) in Islamic jurisprudence, together with the theory of risk and the principle of full compensation for damages in law, provide strong foundations for the government's civil liability regarding damages caused by artificial intelligence. These foundations oblige the government to prevent, supervise, and compensate for damages resulting from emerging technologies with the aim of guaranteeing citizens' rights.

Keywords: Government civil liability, damages, artificial intelligence.

How to cite: Taherloo, F. (2027). Analysis and Examination of the Government's Civil Liability for Damages Arising from Artificial Intelligence with an Approach Based on Islamic Jurisprudence. *The Encyclopedia of Comparative Jurisprudence and Law*, 5(4), 1-17.



تاریخ ارسال: ۱ بهمن ۱۴۰۴

تاریخ بازنگری: ۲۷ اردیبهشت ۱۴۰۵

تاریخ پذیرش: ۵ خرداد ۱۴۰۵

تاریخ چاپ اولیه: ۶ مرداد ۱۴۰۵

تاریخ چاپ نهایی: ۱ آبان ۱۴۰۶

دانشنامه فقه و حقوق تطبیقی

تحلیل و بررسی مسئولیت مدنی دولت در قبال خسارات ناشی از هوش مصنوعی با رویکردی بر فقه

فهیمة طاهرلو^۱*

۱. کارشناسی ارشد، گروه حقوق خصوصی، واحد تهران مرکزی، دانشگاه آزاد اسلامی، تهران، ایران

* پست الکترونیک نویسنده مسئول: fahimehtaherloo@gmail.com

چکیده

گسترش روزافزون هوش مصنوعی و ورود آن به ارکان حکومتی، مسائل جدیدی را در زمینه مسئولیت دولت در قبال خسارات احتمالی این فناوری مطرح کرده است. این پژوهش با هدف تحلیل و بررسی مبانی و مصادیق مسئولیت مدنی دولت در قبال خسارات ناشی از هوش مصنوعی با رویکردی بر فقه، به دنبال ارائه چارچوبی برای پاسخگویی دولت در این حوزه نوظهور است. این تحقیق از نوع نظری بوده و به روش توصیفی - تحلیلی صورت گرفته است. در تحقیق حاضر، اصل امانتداری، صداقت، بی طرفی و اصالت اثر رعایت شده است. مسئولیت مدنی دولت در قبال خسارات ناشی از هوش مصنوعی را می توان در مصادیقی چون خسارات زیست محیطی، خسارات میراث فرهنگی و دستگاه های قضایی جستجو کرد. اصول عدالت، مصلحت، لاضرر و اتلاف در فقه، همراه با نظریه خطر و اصل جبران کامل خسارت در حقوق، مبانی قدرتمندی برای مسئولیت مدنی دولت در قبال خسارات هوش مصنوعی فراهم می کنند. این مبانی، دولت را به پیشگیری، نظارت و جبران خسارات ناشی از فناوری های نوین، با هدف تضمین حقوق شهروندان، مکلف می سازند.

کلیدواژگان: مسئولیت مدنی دولت، خسارات، هوش مصنوعی.

نحوه استناددهی: طاهرلو، فهیمة. (۱۴۰۶). تحلیل و بررسی مسئولیت مدنی دولت در قبال خسارات ناشی از هوش مصنوعی با رویکردی بر فقه. دانشنامه فقه و حقوق تطبیقی، ۵(۴)، ۱۷-۱۹.

حقوقی استوار است و در چه مصادیقی قابل تحقق می‌باشد؟ و فرضیه اصلی چنین بیان می‌شود که «به نظر می‌رسد مبنای مسئولیت دولت در این حوزه ترکیبی از مسئولیت مبتنی بر تقصیر و مسئولیت مبتنی بر خطر است و در مصادیق مرتبط با تصمیم‌گیری‌های عمومی و خدمات هوشمند، دولت مکلف به جبران خسارات ناشی از نقص یا خطای سامانه‌های هوش مصنوعی خواهد بود.» نظر به مطالب یادشده، تحقیق حاضر با عنوان «مسئولیت مدنی دولت در قبال خسارات ناشی از هوش مصنوعی با رویکردی بر فقه» به مطالعه و بررسی می‌پردازد.

بحث

مباحث متفرع در راستای موضوع این تحقیق، به شرح ذیل مطرح می‌گردد.

مفهوم هوش مصنوعی

در گستره کاربرد واژگان، «هوش» که بازتاب‌دهنده مفهوم *Intelligence* در زبان انگلیسی است، در دانشنامه معتبر بریتانیکا، به مثابه توانایی بارز انسان در انطباق مؤثر با محیط از طریق دگرگونی رفتار، اصلاح شرایط زیست‌محیطی، یا ابداع راه‌حل‌های نوین تعریف شده است (Ahmadi et al., 2013).

در معنای اصطلاحی نیز، هوش به مفاهیمی همچون ادراک، فهم، بصیرت و قدرت تحلیل اشاره دارد. لغت‌نامه دهخدا، واژه «هوشمند» را به فردی ذکاوت‌مند و هوشیار اطلاق کرده و آن را هم‌معنی با اصطلاحات «ذو ذكاء» و «ذکی» به معنای سریع‌الفهم معرفی نموده است (Dehkhoda, 1998). این مجموعه معانی، بیانگر ظرفیت ذهنی انسان در دریافت، پردازش و پاسخگویی عقلانی به محرک‌های پیرامونی است.

در کنار واژه «هوش»، اصطلاح «مصنوعی» که ریشه در «صنع» دارد، در لغت به معنای ساخته دست بشر و در تقابل با امور طبیعی به کار می‌رود (Dehkhoda, 1998). در پرتو این معنا، «هوش مصنوعی» به مجموعه‌ای از سیستم‌ها و روش‌های محاسباتی اطلاق

در سال‌های اخیر، گسترش هوش مصنوعی و ورود آن به حوزه‌های مختلف زندگی بشر، از خدمات عمومی تا تصمیم‌گیری‌های حکومتی، پرسش‌های بنیادینی را درباره حدود و ثغور مسئولیت دولت در قبال زیان‌های احتمالی این فناوری برانگیخته است. ماهیت هوش مصنوعی که بر پردازش خودکار داده‌ها، یادگیری مستمر و تولید تصمیم‌های شبه‌مستقل استوار است، موجب شده تا نتایج عملکرد آن همیشه قابل پیش‌بینی نبوده و در برخی موارد از کنترل انسان نیز خارج شود؛ وضعیتی که پیچیدگی‌های تازه‌ای را در انتساب خطا یا زیان ایجاد می‌کند. این تحول فناورانه، در کنار بهره‌گیری روزافزون دولت‌ها از سامانه‌های هوش مصنوعی در امور امنیتی، قضایی، رفاهی، شهری و نظارتی، سبب شده تا مصادیق مسئولیت مدنی دولت در قبال خسارات ناشی از هوش مصنوعی در قلمروهایی چون تصمیم‌یارهای قضایی، سامانه‌های خودکار شناسایی هویت، الگوریتم‌های تخصیص منابع عمومی، خودروهای خودران تحت نظارت دولت و سامانه‌های پیش‌بینی و هشداردهی هوشمند، با چالش‌های حقوقی تازه‌ای مواجه شود.

در چنین بستری، این پرسش مطرح می‌شود که با توجه به ویژگی‌های خاص این فناوری، دولت در برابر خسارات ناشی از استفاده، نظارت ناکافی، خطای الگوریتمی یا نقص در تنظیم‌گری مسئولیت چگونه و بر چه مبنایی باید پاسخگو باشد. بررسی این موضوع بدون توجه به مبنای مسئولیت مدنی دولت در قبال خسارات ناشی از هوش مصنوعی کامل نخواهد بود؛ مبنایی که می‌تواند بر تقصیر، خطر، نقض الزامات تنظیم‌گری یا تعهدات ایجابی دولت در تضمین ایمنی فناوری‌های نو استوار شود و هر یک پیامدهای متفاوتی برای احراز مسئولیت ایجاد می‌کند.

در این راستا، سؤال اصلی پژوهش این است که «مسئولیت مدنی دولت در قبال خسارات ناشی از هوش مصنوعی بر چه مبنایی

کاراکترهای چاپی به صورت خودکار و بر پایه نمونه‌های پیشین شناسایی می‌شوند و این خود نمونه‌ای رایج از کاربرد یادگیری ماشین است (Aghajani, 2021; Sarabadi & Amir, 2018).

«بینایی ماشین»، که در واقع تجسمی از توانایی ادراک بصری برای ماشین‌ها و رایانه‌هاست، به ماشین‌ها این قابلیت را می‌بخشد تا تصاویر را تحلیل کرده و مفاهیم نهفته در آن‌ها را درک کنند؛ فرایندی که به طور شگفت‌انگیزی شبیه به عملکرد سیستم بینایی انسان عمل می‌کند. در دنیای ماشین‌ها و ربات‌ها، اطلاعات بصری از طریق دوربین‌های دیجیتال دریافت می‌شود و در نهایت، آنچه توسط این سیستم‌ها پردازش می‌گردد، در قالب ماتریسی از اعداد نمایان می‌شود. مأموریت اصلی «بینایی ماشین»، نه تنها پردازش این داده‌های عددی، بلکه تبدیل آن‌ها به درک و مفاهیمی است که برای انسان قابل فهم باشد. یکی از موانع قابل توجه در این عرصه، مواجهه با حجم عظیم اطلاعات بصری است که ماشین‌ها باید با آن دست‌وپنجه نرم کنند. بنابراین، می‌توان گفت هدف «بینایی ماشین»، شبیه‌سازی دقیق عملکرد سیستم بینایی انسان در محیط‌های ماشینی است. این حوزه، دامنه کاربردهای گسترده‌ای را دربرمی‌گیرد و به نظر می‌رسد بیش از هر چیز در بازرسی‌های خودکار و هدایت ربات‌ها در محیط‌های صنعتی مورد استفاده قرار می‌گیرد. برخی دیگر از کاربردهای چشمگیر «بینایی ماشین» شامل تضمین کیفیت محصولات، دسته‌بندی خودکار کالاها، جابه‌جایی هوشمندانه مواد، هدایت دقیق ربات‌ها و انجام اندازه‌گیری‌های نوری با دقت بالا است. اولین گام حیاتی در اجرای فرایند «بینایی ماشین»، ثبت تصویر است که معمولاً با استفاده از ترکیبی از دوربین، لنز و نورپردازی دقیق صورت می‌گیرد و طراحی این مجموعه باید با نهایت دقت انجام شود تا پردازش‌های بعدی با موفقیت همراه باشد. سیستم‌های نرم‌افزاری «بینایی ماشین» با بهره‌گیری از تکنیک‌های پیشرفته پردازش تصاویر دیجیتال،

می‌گردد که به منظور بازآفرینی یا شبیه‌سازی رفتارهای هوشمندانه انسانی طراحی شده‌اند و قابلیت انجام اموری چون یادگیری، استنتاج و حل مسئله را دارا هستند.

با این وجود، در قلمرو هوش مصنوعی، برداشت صاحب‌نظران از مفهوم «هوش»، تفاوتی بنیادین با رویکردهای سنتی دارد. روان‌شناسان عمدتاً بر سنجش ضریب هوشی و فرایندهای روانی متمرکز هستند، در حالی که پژوهشگران هوش مصنوعی، هدف خود را بر تعریف عملی و قابل اندازه‌گیری از هوش بنا نهاده‌اند. از دیدگاه آنان، ورود به مباحث انتزاعی و فلسفی پیرامون ماهیت عقل و ذهن، غالباً به نتایج قطعی نمی‌انجامد، چرا که این مفاهیم از منظر نظری، ذهنی و از منظر علمی، نیازمند معیارهای کارکردی هستند. بنابراین، هوش در هوش مصنوعی، بیشتر به قابلیت سامانه‌های ماشینی در تحلیل داده‌ها، یادگیری، تصمیم‌گیری و اجرای کنش‌های هدفمند اطلاق می‌شود (Ghaeminia, 2006).

ماهیت هوش مصنوعی

«یادگیری ماشین» به عنوان بخشی کلیدی از هوش مصنوعی، به مطالعه و طراحی سیستم‌هایی می‌پردازد که قابلیت یادگیری از داده‌ها را دارا هستند. در این رویکرد، ماشین از طریق تعدیل داده‌های خود بر اساس ورودی‌ها، رفتار خود را به نتایج مطلوب نزدیک‌تر می‌سازد؛ به عبارت دیگر، ماشین باید توانایی تحلیل داده‌ها را کسب کند. به عنوان مثال، در حوزه تشخیص پزشکی، می‌توان با استفاده از یادگیری ماشین، سیستمی را آموزش داد که با تحلیل تصاویر پزشکی مانند رادیوگرافی یا *MRI*، قادر به شناسایی ناهنجاری‌ها و بیماری‌های احتمالی باشد. این سیستم پس از یادگیری از مجموعه وسیعی از تصاویر سالم و بیمار، می‌تواند تصاویر جدید را تحلیل کرده و با دقت بالا، احتمال وجود بیماری را در بیماران تشخیص دهد. کاربردهای متنوعی برای یادگیری ماشین وجود دارد، از جمله تشخیص نوری کاراکتر که در آن،

مردن آن، از قرن بیستم آغاز شد. امروزه، واژه «ربات» به طور گسترده به هرگونه ماشین ساخته دست انسان اطلاق می شود که قادر به انجام وظایفی است که معمولاً توسط انسانها به انجام می رسد. در حال حاضر، بخش عظیمی از رباتها در محیطهای صنعتی، به ویژه کارخانهها، برای تولید محصولاتی همچون خودرو و موارد مشابه به کار گرفته می شوند. علاوه بر این، رباتها در مأموریت های اکتشافی در اعماق اقیانوسها یا در سیارات دیگر نیز نقشی حیاتی ایفا می کنند؛ به نظر می رسد، این کاربردها نشان دهنده توانایی رباتها در فعالیت در محیطهای غیرقابل دسترس برای انسان است. رباتها معمولاً از سه مؤلفه اصلی تشکیل شده اند: «مغز» که عموماً یک کامپیوتر مرکزی است، بخش مکانیکی که وظیفه حرکت و انجام عمل را بر عهده دارد، یعنی محرکها، و مجموعه ای از سنسورها که اطلاعات محیطی را جمع آوری می کنند. این سنسورها می توانند انواع مختلفی داشته باشند، از جمله سنسورهای بینایی برای درک تصاویر، سنسورهای صوتی برای شنیدن، سنسورهای تشخیص دما، نور، تماس و حرکت. با توجه به تنوع کاربردها و ویژگی های رباتها، آنها به دسته های گوناگونی طبقه بندی می شوند. با این حال، وجه مشترک در تمامی این انواع رباتها این است که «این شاخه از هوش مصنوعی در تلاش است تا رباتها را به گونه ای برنامه ریزی کند که بتوانند اعمال هوشمندانه، نظیر دیدن، شنیدن و واکنش نشان دادن به محرکهای محیطی را به انجام رسانند» (Sarabadi & Amir, 2018).

حوزه رباتیک امروزه شاهد رشدی خیره کننده است و رباتهای نوین با اهداف متنوع در حوزه های عمومی، تجاری و نظامی طراحی و به خدمت گرفته می شوند. بسیاری از این رباتها وظایفی را بر عهده دارند که به دلیل ماهیت خطرناکشان برای انسان، توسط انسان قابل انجام نیست؛ از جمله این وظایف می توان به خنثی سازی بمبها و مینها، بازرسی دقیق لاشه های کشتی ها و

اطلاعات کلیدی را استخراج کرده و براساس آن، تصمیم گیری های لازم را انجام می دهند. چنین سیستمی، به نظر می رسد، می تواند در سناریوهای مختلفی از جمله شناسایی افراد مشکوک در اماکن عمومی، ارزیابی میزان هوشیاری رانندگان با تحلیل حالات چشم آنها، فراهم کردن امکان تعامل افراد دارای معلولیت جسمی با کامپیوتر از طریق فرمانهای چشمی، ایجاد نظم و انضباط در ترافیک شهری و همچنین نظارت نامحسوس در محیطهای مختلف، به کار گرفته شود (Aghajani, 2021; Jafari et al., 2016).

«پردازش زبان طبیعی»، به مثابه پلی میان علوم کامپیوتر، زبان شناسی و هوش مصنوعی، به بررسی تعامل پیچیده میان رایانه ها و زبان انسانی می پردازد. به تعبیری دیگر، این حوزه بر کاربرد رایانه ها در تحلیل و درک عمیق زبان، چه به صورت گفتاری و چه نوشتاری، تمرکز دارد. کاربردهای «پردازش زبان طبیعی» به طور کلی به دو دسته عمده، یعنی کاربردهای نوشتاری و گفتاری، تقسیم می شوند. در قلمرو کاربردهای نوشتاری، می توان به استخراج اطلاعات خاص و کلیدی از متون، ترجمه روان متون به زبانهای مختلف، و یافتن مستندات مورد نظر در پایگاههای داده نوشتاری، مانند جست و جوی یافتن کتابهای مرتبط در یک کتابخانه، اشاره کرد. در بخش کاربردهای گفتاری، این حوزه شامل سیستمهای پیشرفته پرسش و پاسخ میان انسان و رایانه، خدمات خودکار و هوشمند ارتباط با مشتری از طریق تلفن، سیستمهای آموزشی تعاملی برای دانش آموزان، و سیستمهای کنترلی است که با فرمان صوتی عمل می کنند؛ به نظر می رسد، نمونه بارز این کاربردها، دستیارهای صوتی هستند که امروزه در بسیاری از دستگاهها یافت می شوند (Razavi Salestani et al., 2015).

قدمت تلاش برای ساخت ماشینهایی که بتوانند وظایف را به صورت خودکار انجام دهند، به دورانهای گذشته بازمی گردد؛ اما پژوهشهای عملیاتی و بررسی کاربردهای رباتها به معنای

دیگر مأموریت‌های مشابه اشاره کرد. «سیستم خبره»، به سامانه‌های کامپیوتری اطلاق می‌شود که توانایی تصمیم‌گیری یک فرد متخصص در یک حوزه خاص را شبیه‌سازی می‌کنند. اساس کار این سیستم‌ها بر گردآوری، سازماندهی و تحلیل تجربیات و دانش موجود برای اتخاذ یا پیشنهاد تصمیمات استوار است. نتایج حاصل از سیستم‌های خبره، به نظر می‌رسد، می‌تواند شامل تصمیمات بسیار مؤثر و قابل اتکایی در حوزه‌های مختلف باشد. تحقیقات حاکی از آن است که سیستم‌های خبره به‌عنوان کارآمدترین نوع سیستم‌ها در برنامه‌ریزی تولید و کنترل موجودی در شرکت‌های صنعتی شناخته شده‌اند و علاوه بر افزایش کارایی، سرعت اجرای وظایف را نیز به‌طور چشمگیری بهبود می‌بخشند (Katouz et al., 2017). در این راستا، کامپیوترها به‌طور ویژه برای اخذ تصمیم در شرایط واقعی برنامه‌ریزی می‌شوند. سیستم‌های خبره، برخلاف سیستم‌هایی که صرفاً از دستورالعمل‌های برنامه‌نویسی پیروی می‌کنند، با بهره‌گیری از دانش تخصصی خود و از طریق استنتاج منطقی، به حل مسائل پیچیده می‌پردازند، مشابه عملکرد یک فرد متخصص. یکی از مثال‌های بارز این سیستم‌ها، سامانه‌های تشخیص بیماری است. در این سامانه‌ها، تجربیات چندین متخصص در کنار داده‌های دریافتی از مراجعان، وارد کامپیوتر می‌شود؛ سپس کامپیوتر با طرح سؤالاتی از مراجعه‌کننده و مقایسه پاسخ‌ها با پایگاه دانش خود، به تشخیص بیماری فرد می‌پردازد (Sarabadi & Amir, 2018). شایان ذکر است که اگرچه این سیستم‌ها نقش «خبره» را ایفا می‌کنند، اما توانایی آن‌ها محدود به اطلاعاتی است که به آن‌ها ارائه شده است و به‌خودی‌خود فاقد قابلیت یادگیری مستقل هستند.

«شبکه‌های عصبی مصنوعی»، سیستم‌هایی هستند که قادر به تقلید از عملکردهای سیستم‌های عصبی طبیعی و ویژگی‌های مغز انسان می‌باشند. این شبکه‌ها در واقع مدل‌هایی هستند که با الهام از ساختار سیستم عصبی مرکزی حیوانات، به‌ویژه مغز، طراحی

شده‌اند و توانایی یادگیری و شناسایی الگوها را دارا هستند. اگرچه شبکه‌های عصبی سابقه دیرینه‌ای دارند، اما پیشرفت‌های چشمگیری که در اواخر دهه ۷۰ و اوایل دهه ۸۰ میلادی رخ داد، فرصت‌های تحقیقاتی و بحث‌های گسترده‌تری را در این زمینه فراهم آورد. از ویژگی‌های برجسته شبکه‌های عصبی مصنوعی، توانایی یادگیری آن‌هاست؛ این به معنای قابلیت شبکه در دریافت و ذخیره‌سازی اطلاعات برای استفاده آتی است. قابلیت تعمیم‌پذیری و پردازش موازی نیز از دیگر مزایای این شبکه‌هاست که به‌طور قابل توجهی سرعت پردازش اطلاعات را افزایش می‌دهد. ساختار شبکه‌های عصبی از ترکیب نوروها، گره‌ها یا واحدهای پردازش اطلاعات تشکیل شده است که به‌صورت سلسله‌مراتبی در لایه‌های مختلفی سازماندهی می‌شوند؛ این لایه‌ها شامل لایه ورودی، لایه پنهان، یعنی میانی، و لایه خروجی می‌باشند. همان‌طور که اشاره شد، این شبکه‌ها که بر پایه مدل شبکه عصبی ذهن انسان طراحی شده‌اند، کاربردهای گسترده و متنوعی دارند. به‌عنوان مثال، شبکه‌های عصبی مصنوعی در زمینه‌هایی چون سیستم‌های کنترلی، رباتیک، تشخیص متون و پردازش تصویر به کار می‌روند. علاوه بر این، به نظر می‌رسد، این شبکه‌ها ابزاری قدرتمند برای شناسایی الگوهای ناشناخته در داده‌ها محسوب می‌شوند و یکی از کاربردهای کلیدی آن‌ها نیز در حوزه پیش‌بینی قرار دارد (Biranvand & Sahraeyan, 2018; Ebrahimi & Pasha, 2018).

«الگوریتم ژنتیک»، به‌عنوان یکی از شاخه‌های خلاقانه هوش مصنوعی، ریشه‌های خود را در نظریه تکامل طبیعی چارلز داروین می‌دواند. این الگوریتم، شگفتی‌های فرایند انتخاب طبیعی را بازآفرینی می‌کند و بر پایه انتخاب برترین افراد برای بقا و انتقال ویژگی‌های ژنتیکی به نسل‌های بعدی عمل می‌نماید. به‌طور مشابه، «الگوریتم ژنتیک» با الهام از همین مفهوم، به جست‌وجوی بهینه‌ترین راه‌حل ممکن برای مسائل پیچیده از میان مجموعه‌ای از

حوزه، با در نظر گرفتن طیف گسترده کاربردهای هوش مصنوعی، به پیگیری آن در زیرشاخه‌های متعددی مانند شبکه‌های عصبی، پردازش زبان طبیعی، رباتیک، مسابقات هوش مصنوعی، سیستم‌های خبره، یادگیری ماشین، استراتژی‌های تکاملی و بینایی ماشین پرداخته‌اند. در ادامه به دو مورد از مهم‌ترین مصادیق مسئولیت مدنی دولت در قبال خسارات ناشی از هوش مصنوعی پرداخته می‌شود.

مسئولیت مدنی دولت در قبال خسارات زیست‌محیطی هوش مصنوعی

در تحلیل مسئولیت مدنی دولت در قبال خسارات زیست‌محیطی ناشی از هوش مصنوعی، می‌توانیم از چارچوب حقوقی جرایم زیست‌محیطی به‌عنوان یک الگو بهره ببریم. تخریب محیط زیست، به‌طور کلی، هرگونه عملی است که به‌طور جدی به اکوسیستم آسیب رسانده و سلامت و بقای بشر را به مخاطره اندازد. به نظر می‌رسد، این نوع جرایم، مانند بسیاری دیگر، دارای سه رکن اساسی هستند: رکن قانونی، رکن مادی و رکن معنوی، یعنی عمدی. رکن قانونی: این رکن به اصل قانونی بودن جرایم و مجازات‌ها اشاره دارد که در آن، هیچ عملی جرم تلقی نمی‌شود مگر آنکه قانون صراحتاً آن را جرم‌انگاری کرده باشد. در چارچوب هوش مصنوعی، می‌توان قوانین مربوط به مسئولیت در قبال سیستم‌های خودکار و الگوریتم‌ها را مد نظر قرار داد. به‌عنوان مثال، اگر یک سیستم هوش مصنوعی که توسط دولت توسعه یافته یا مدیریت می‌شود، در فرایند پردازش داده‌های زیست‌محیطی خطایی مرتکب شود که منجر به تخریب گسترده یک منطقه شود، این سؤال مطرح می‌شود که کدام قانون مسئولیت دولت را در این زمینه مشخص می‌کند. در قوانین موجود، موادی مانند مواد ۶۷۵ تا ۶۹۶ قانون مجازات اسلامی و همچنین برخی مواد قوانین خاص مانند قانون هوای پاک، مانند مواد ۲۹ و ۳۱، به جرایم زیست‌محیطی اشاره دارند. در بحث هوش مصنوعی، قانون‌گذار

پاسخ‌های صحیح و محتمل می‌پردازد. به نظر می‌رسد، حوزه کاربرد «الگوریتم ژنتیک» بسیار وسیع و رو به گسترش است و با توجه به پیشرفت‌های مستمر در عرصه‌های علوم و فناوری، استفاده از این روش در بهینه‌سازی و حل مسائل، به‌طور فزاینده‌ای رایج شده است. «الگوریتم‌های ژنتیک» در زمینه‌های متعددی کاربرد دارند، از جمله در حوزه‌هایی چون بیوانفورماتیک، فیلوژنتیک، علوم محاسباتی، مهندسی شیمی، فیزیک، ریاضیات و داروسازی. علاوه بر این، به نظر می‌رسد، این الگوریتم‌ها در بسیاری از حوزه‌های پژوهشی دیگر نیز مورد استفاده قرار می‌گیرند، مانند حل مسائل عددی، حل مسئله مدار هامیلتونی، شناسایی ساختار مولکول‌های پروتئین، مسائل NP طراحی شبکه‌های عصبی، تولید توابع برای خلق تصاویر، مدل‌سازی شناختی و تئوری بازی‌ها (Sivandian & Rastgari, 2014). به‌عنوان مثال، در طراحی یک شبکه عصبی، «الگوریتم ژنتیک» می‌تواند برای یافتن بهترین ساختار و پارامترهای بهینه مورد استفاده قرار گیرد تا عملکرد شبکه به حداکثر برسد.

نظر به آنچه گذشت، می‌توان گفت که هوش مصنوعی، قلمروی وسیع و چندوجهی است که مجموعه‌ای از شاخه‌ها و فنون تخصصی را دربرمی‌گیرد و به سامانه‌های رایانه‌ای امکان می‌پردازد تا با چالش‌های پیچیده‌ای همچون سیستم‌های خبره، بینایی ماشین و سایر فناوری‌های پیشرفته روبه‌رو شوند (Salehabadi et al., 2019).

مصادیق مسئولیت مدنی دولت در قبال خسارات ناشی از هوش مصنوعی

امروزه، ردپای هوش مصنوعی در عرصه‌های گوناگونی از جمله پزشکی، هوافضا، اکتشافات علمی، صنایع نظامی، پیش‌بینی آب‌وهوا، نقشه‌برداری و شناسایی عوارض جغرافیایی، تشخیص صدا، گفتار و دستخط، و همچنین در بازی‌ها و نرم‌افزارهای رایانه‌ای به‌وضوح مشهود است. به همین دلیل، متخصصان این

باید به تناسب جرم و مجازات توجه ویژه داشته باشد و چارچوب‌های روشنی برای مسئولیت‌پذیری دولت در قبال خطاهای احتمالی سیستم‌های هوش مصنوعی تعریف کند (Hayati, 2010).

رکن مادی: این رکن به وقوع فعل یا ترک فعلی اشاره دارد که منجر به تحقق جرم می‌شود. در مورد خسارات زیست‌محیطی ناشی از هوش مصنوعی، این رکن می‌تواند شامل اقدامات یا عدم اقدامات سیستم‌های هوش مصنوعی باشد. به‌عنوان مثال، اگر یک الگوریتم هوش مصنوعی که برای مدیریت منابع آب طراحی شده است، به دلیل نقص فنی یا الگوریتم نادرست، منجر به رهاسازی فاضلاب صنعتی در یک رودخانه شود، این ترک فعل، یعنی عدم مدیریت صحیح، یا فعل، یعنی رهاسازی نادرست، رکن مادی جرم را تشکیل می‌دهد. در قوانین فعلی، مواردی مانند عدم نصب فیلترهای کاهنده آلودگی بر روی خروجی دودکش‌ها به‌عنوان مصادیق تحقق مادی جرم، یعنی فعل یا ترک فعل، ذکر شده است (Taghizadeh Ansari, 1995). در مورد هوش مصنوعی، می‌توان این را به نقص در کدهای نرم‌افزاری، خطاهای پردازشی یا حتی عدم به‌روزرسانی سیستم‌های نظارتی تعمیم داد.

رکن معنوی، یعنی عمدی: این رکن به قصد و نیت مجرمانه اشاره دارد. در بسیاری از جرایم زیست‌محیطی، وجود «قصد مجرمانه» یا «عمد» ضروری است. قصد به معنای اراده آگاهانه برای انجام عملی است که قانون آن را منع کرده است. در مورد خسارات ناشی از هوش مصنوعی، بحث «عمد» پیچیده‌تر می‌شود. آیا می‌توان هوش مصنوعی را به داشتن «قصد» متهم کرد؟ در رویه‌های حقوقی فعلی، مسئولیت به فرد یا نهادی که سیستم را طراحی، پیاده‌سازی یا مدیریت کرده است، نسبت داده می‌شود. اگر مشخص شود که دولت یا مسئولین آن از خطرات احتمالی سیستم هوش مصنوعی آگاه بوده‌اند اما اقدام لازم را برای پیشگیری از وقوع خسارت انجام نداده‌اند، می‌توان این را نوعی «عمد» یا «تقصیر» تلقی کرد. به‌عنوان

مثال، اگر دولت از ضعف‌های امنیتی یک سیستم هوش مصنوعی که مسئولیت نظارت زیست‌محیطی را بر عهده دارد، آگاه باشد و برای رفع آن اقدامی نکند، و در نتیجه این ضعف منجر به تخریب یک منطقه حفاظت‌شده شود، مسئولیت مدنی دولت قابل طرح خواهد بود. برخی از جرایم زیست‌محیطی «مقید» هستند، یعنی تحقق آن‌ها منوط به ورود خسارت است، مانند بند «د» ماده ۱۲ قانون شکار و صید که به آلودگی آب و از بین رفتن آبزیان اشاره دارد. در مقابل، جرایم «مطلق» وجود دارند که تحقق آن‌ها نیازی به نتیجه مجرمانه ندارد، مانند ممنوعیت سوزاندن زباله در معابر عمومی طبق ماده ۲۴ قانون هوای پاک، که ضمانت اجرای آن در ماده ۳۱ همان قانون مشخص شده است. در مورد هوش مصنوعی، خسارات زیست‌محیطی می‌تواند هم مطلق، مانند انتشار آلودگی بیش از حد مجاز توسط یک سیستم صنعتی خودکار، و هم مقید، مانند ایجاد یک فاجعه زیست‌محیطی در اثر یک خطای محاسباتی در مدل پیش‌بینی آب‌وهوا، باشد.

معاونت در جرم: در قوانین جرایم زیست‌محیطی، علاوه بر مباشر جرم، به معاونت نیز پرداخته شده است. در صورتی که فردی به‌طور عمدی به مجرم در تحقق جرم کمک کند، معاون جرم محسوب می‌شود. در مورد هوش مصنوعی، اگرچه خود سیستم قادر به معاونت نیست، اما افرادی که با ارائه اطلاعات نادرست، عدم نظارت کافی یا سهل‌انگاری در پیاده‌سازی، به وقوع خسارت زیست‌محیطی توسط سیستم کمک کنند، می‌توانند مشمول عنوان معاونت در جرم شناخته شوند (Ghasemi, 2012).

در نهایت، تحلیل این مفاهیم در بستر هوش مصنوعی، نیازمند بازنگری و احتمالاً تدوین قوانین جدید برای پوشش دادن مسئولیت‌های ناشی از تصمیمات و اقدامات سیستم‌های خودکار است.

منصوب باشند. مواد ۵۵۸، ۵۶۳، ۵۶۵، ۵۶۶ و ۵۶۹ قانون مجازات اسلامی به این جرایم پرداخته‌اند.

فرض کنید یک سیستم هوش مصنوعی پیشرفته مسئولیت مدیریت و نظارت بر اماکن تاریخی را بر عهده دارد. این سیستم وظیفه دارد از طریق دوربین‌ها و سنسورهای خود، هرگونه فعالیت مشکوک یا تخریب احتمالی را شناسایی و به مقامات مربوطه گزارش دهد. حال، اگر به دلیل نقص در الگوریتم‌های یادگیری ماشین یا خطای نرم‌افزاری، این سیستم قادر به تشخیص حفاری‌های غیرمجاز در یک محوطه باستانی نباشد یا اطلاعات نادرستی در مورد وضعیت یک بنای تاریخی ارائه دهد، خسارات جبران‌ناپذیری به میراث فرهنگی وارد خواهد شد. در چنین سناریویی، مسئولیت مدنی دولت در قبال این خسارات، از طریق عدم نظارت کافی یا استفاده از فناوری‌های ناکارآمد، مطرح می‌شود.

برای مثال، اگر یک الگوریتم هوش مصنوعی که برای تحلیل تصاویر ماهواره‌ای جهت شناسایی تجاوز به حریم آثار تاریخی طراحی شده است، به دلیل عدم دقت کافی یا عدم آموزش بر روی داده‌های متنوع، قادر به تفکیک سازه‌های جدید از بقایای تاریخی نباشد و در نتیجه، مجوز ساخت‌وساز در نزدیکی یک اثر ثبت ملی صادر شود، این می‌تواند منجر به تخریب بخش‌هایی از آن اثر گردد. در این حالت، دولت به‌عنوان متولی حفاظت از میراث فرهنگی، مسئول جبران خسارات وارده خواهد بود، زیرا فناوری مورد استفاده، نتوانسته است وظیفه نظارتی خود را به‌درستی انجام دهد. این نشان می‌دهد که حتی در حوزه صیانت از گذشته، هوش مصنوعی می‌تواند چالش‌های جدیدی در مسئولیت‌پذیری دولت ایجاد کند.

مسئولیت مدنی دولت در قبال خسارات ناشی از دستگاه‌های قضائی

مسئولیت مدنی دولت در قبال خسارات ناشی از دستگاه‌های قضائی، موضوعی پیچیده و چندوجهی است که ریشه در اصول

مسئولیت مدنی دولت در قبال خسارات میراث فرهنگی

در بررسی مسئولیت مدنی دولت نسبت به خسارات ناشی از هوش مصنوعی، یکی از مصادیق قابل توجه، حفاظت از میراث فرهنگی و تاریخی کشور است. آثار تاریخی و فرهنگی، صرفاً اموال مادی نیستند، بلکه ریشه‌های هویتی و فرهنگی یک ملت را در خود دارند و متعلق به نسل‌های گذشته، حال و آینده‌اند. بنابراین، هرگونه تعرض به این آثار، نه‌تنها به بنیان‌های فرهنگی جامعه آسیب می‌زند، بلکه آرامش و آسایش عمومی را نیز مختل می‌سازد. این دیدگاه، جرایم علیه میراث فرهنگی را در زمره جرایم علیه آسایش عمومی قرار می‌دهد؛ اعمالی که نظم عمومی جامعه را بر هم زده و مصالح همگانی را نقض می‌کند (Sarikhan, 2005).

تخریب، حفاری‌های غیرمجاز، تغییر کاربری، مرمت‌های غیرقانونی، و هرگونه دستکاری در تزیینات، تأسیسات، اشیاء و نقوش اماکن تاریخی، باعث تزلزل پایه‌های فرهنگی یک ملت شده و اعتبار فرهنگی آن را در سطح بین‌المللی خدشه‌دار می‌کند. این آثار، که گنجینه دانش، تجربه و تاریخ یک کشور محسوب می‌شوند، نقشی حیاتی در تعیین هویت تاریخی و فرهنگی جامعه ایفا می‌کنند و قانون‌گذار برای حراست از آن‌ها، مجازات‌هایی را برای مرتکبان جرایم مرتبط پیش‌بینی کرده است.

این میراث گران‌بها، که هم جنبه مادی و هم جنبه معنوی دارند، باید با دقت مورد مراقبت قرار گیرند تا نسل حاضر بتواند از آن بهره‌مند شده و به آیندگان منتقل گردد. در این راستا، دولت وظیفه دارد با افرادی که به فعالیت‌های غیرقانونی در حوزه میراث فرهنگی مشغول هستند، قاطعانه برخورد کند (Arzhang et al., 2019).

به‌طور خلاصه، موضوع تخریب آثار تاریخی شامل اشیاء منقول و غیرمنقول تاریخی، فرهنگی یا مذهبی می‌شود. در خصوص آثار غیرمنقول، شرط ثبت در فهرست آثار ملی ضروری است و برای آثار منقول، باید در اماکن ثبت‌شده در فهرست آثار ملی موجود یا

حقوقی بنیادین دارد. در نظام حقوقی ایران، اصل بر عدم مسئولیت دولت در قبال اعمال حاکمیتی و قضایی است، زیرا قضاوت، امری حاکمیتی تلقی شده و استقلال دستگاه قضا از اصول مسلم آن است. با این حال، این اصل مطلق نبوده و در مواردی استثنائاتی بر آن وارد شده است. یکی از مهم‌ترین این استثنائات، مربوط به «اشتباهات قضایی» است که می‌تواند منجر به ورود خسارت به شهروندان شود. این اشتباهات ممکن است ناشی از خطای قاضی در تشخیص موضوع، تفسیر نادرست قانون یا حتی نقص در تحقیقات قضایی باشد.

در قانون اساسی جمهوری اسلامی ایران، اصل ۱۷۱ به‌صراحت به این موضوع پرداخته و مقرر می‌دارد: «هرگاه در اثر تقصیر یا اشتباه قاضی در موضوع یا در حکم یا در تطبیق حکم بر مورد خاص، ضرر مادی یا معنوی متوجه کسی گردد، در صورت تقصیر، مقصر طبق موازین اسلامی ضامن است و در غیر این صورت خسارت به وسیله دولت جبران می‌شود، و در هر حال از متهم اعاده حیثیت می‌گردد»؛ این اصل، مبنای اصلی پذیرش مسئولیت مدنی دولت در قبال اشتباهات قضایی است. همچنین، قانون مسئولیت مدنی مصوب ۱۳۳۹ و آیین‌نامه اجرایی اصل ۱۷۱ قانون اساسی، سازوکارهای اجرایی و تعیین میزان خسارت را مشخص کرده‌اند. با این حال، تمایز قائل شدن میان «خطای قضایی»، یعنی تقصیر قاضی، و «اشتباه در استنباط قضایی»، که لزوماً تقصیر محسوب نمی‌شود، از اهمیت بالایی برخوردار است. در صورتی که ثابت شود قاضی با عمد یا تقصیر، موجب ورود زیان به فردی شده است، دولت مکلف به جبران خسارت خواهد بود. این جبران خسارت می‌تواند شامل ضررهای مالی، از دست دادن فرصت‌های شغلی، یا حتی لطمات روحی و روانی باشد. فرایند احراز تقصیر قاضی و تعیین میزان خسارت، معمولاً از طریق هیئت‌های رسیدگی به تخلفات اداری قضات و یا دیوان عالی کشور صورت می‌گیرد. این سازوکارها تلاش می‌کنند تا ضمن حفظ استقلال دستگاه قضا،

حقوق شهروندان در برابر خطاهای احتمالی قضایی را نیز تضمین نمایند.

در دوران معاصر و با ظهور فناوری‌های نوین، بحث مسئولیت مدنی دولت در قبال خسارات ناشی از «هوش مصنوعی» در دستگاه قضایی نیز مطرح شده است. به‌عنوان مثال، اگر سیستمی مبتنی بر هوش مصنوعی در فرایند دادرسی، مانند پیش‌بینی احتمال تکرار جرم یا کمک به صدور حکم، به دلیل نقص فنی یا الگوریتمی نادرست، منجر به صدور حکمی ظالمانه و ورود خسارت به فردی شود، مسئولیت دولت در این زمینه نیز قابل بررسی خواهد بود. این موضوع نشان می‌دهد که مفاهیم سنتی مسئولیت مدنی در حوزه قضایی، با ورود فناوری‌های جدید، نیازمند بازنگری و تطبیق هستند.

مبانی مسئولیت مدنی دولت در قبال خسارات ناشی از هوش مصنوعی

مبانی فقهی

الف) اصل عدالت

عدالت، ستون فقرات هر جامعه پویا و الهام‌بخش آموزه‌های شرعی، به‌ویژه در فقه اسلامی، جایگاهی والا دارد (Gorji, 1997).

قانون، دو وظیفه اصلی بر عهده دارد: تبیین حقوق متقابل افراد در جامعه و پیش‌بینی سازوکارهای لازم برای حمایت از این حقوق. همچنین، قانون‌گذار باید کرامت انسانی را پاس داشته و با هرگونه ستمگری مبارزه کند تا استقلال سیاسی، اقتصادی، قضایی و فرهنگی جامعه تضمین گردد. در پرتو این مبانی، مسئولیت مدنی دولت در قبال خسارات ناشی از هوش مصنوعی، با تکیه بر اصل عدالت، قابل تحلیل است. اگر استفاده از هوش مصنوعی در دستگاه‌های دولتی، به دلیل نقص در طراحی، اجرا یا نظارت، منجر به تضییع حقوق شهروندان یا ورود خسارت به آنان شود، دولت به‌عنوان متولی اجرای عدالت، مسئول جبران این خسارات خواهد

ج) قاعده لاضرر

قاعده «لاضرر» به عنوان یکی از مستحکم ترین قواعد فقهی، می تواند مبنایی روشن برای تبیین مسئولیت مدنی دولت در قبال خسارات ناشی از هوش مصنوعی تلقی شود. مستند این قاعده، کلام پیامبر اکرم (ص) است: «لا ضرر و لا ضرار فی الدین یا فی الإسلام» (Rashti Najafi) و حتی ادعای تواتر آن نیز مطرح شده است (Hosseini Yazdi, 1980). مفاد این اصل آن است که در نظام حقوقی اسلام، حکمی که مستلزم ضرر باشد جعل نشده و هر حکمی که به زیان اشخاص بینجامد، نفی می شود (Ansari; Jafari Tabrizi). چنان که در نمونه هایی مانند وجوب روزه برای مریض یا لزوم بیع غبنی، حکم ضرری برداشته می شود (Hosseini Yazdi, 1980). ضرر به معنای نقص در شیء است (Kompani, 1954) و نتیجه قاعده، لزوم جبران خسارت در صورت تحقق زیان است (Mousavi Khomeini, 1989). نفی حکم ضرری نیز گاه به نفی حقیقت ضرر و گاه به نفی جعل ابتدایی آن بازمی گردد (Hosseini Yazdi, 1980). از این رو، به نظر می رسد اطلاق حدیث، هم منع ایجاد زیان و هم رفع خسارت موجود را دربرمی گیرد (Marashi, 2006). در حوزه هوش مصنوعی، چنانچه دولت در تنظیم گری، نظارت یا اصلاح کاستی های سامانه های هوشمند قصور ورزد و این ترک فعل به ورود زیان بینجامد، قاعده لاضرر با توسعه شمول خود به امور عدمی، ضمان را اثبات می کند؛ همان گونه که در خسارات زیست محیطی نیز چنین برداشتی پذیرفته شده است (Tarimoradi & Fakhelai, 2005).

د) قاعده اتلاف

قاعده اتلاف از بنیادی ترین اصول فقهی در اثبات ضمان قهری است و می تواند به عنوان مبنایی استوار برای تعیین مسئولیت مدنی دولت در قبال خسارات ناشی از هوش مصنوعی تلقی شود. در مفهوم اصطلاحی، اتلاف به معنای از بین بردن مال دیگری است

بود. این مسئولیت، ناشی از وظیفه حاکمیت در برقراری نظم عادلانه و حمایت از حقوق افراد در برابر خطاهای احتمالی فناوری های نوین است. این امر، پاسخی به ضرورت اجتماعی امروز برای تحقق عدالت در تمامی عرصه ها، از جمله حوزه فناوری و حقوق شهروندی است.

ب) قاعده مصلحت

قاعده مصلحت، هنگامی که در مقام تحلیل مسئولیت مدنی دولت در برابر خسارات ناشی از هوش مصنوعی به کار گرفته می شود، چارچوبی پویا و انعطاف پذیر برای پاسخ دهی به چالش های نوپدید ارائه می کند. مصلحت که در ادبیات فقهی با مفاهیمی همچون عدم مفسده، ضرورت، منفعت، استحسان، سدّ و فتح ذریعه، عرف و مقاصد شریعت پیوند خورده است (Dehkhoda, 1998; Ibn Manzur, 1988; Shartuni, 1983)، ظرفیت دارد تا مبنای تنظیم تکالیف حاکمیتی در مواجهه با فناوری های پیچیده ای چون هوش مصنوعی قرار گیرد. از آنجا که مصلحت در معنای مصدری به «جلب منفعت و دفع مفسده» و در معنای اسم مصدری به «خیر، صواب و شایستگی» اشاره دارد (Dehkhoda, 1998; Feyz, 1994)، می تواند نقشی راهبردی در توجیه الزام دولت به پیشگیری از آسیب های ناشی از عملکرد سامانه های هوشمند ایفا کند. با توجه به اینکه این فناوری ها توان ایجاد خسارات گسترده، سریع و بعضاً غیرقابل پیش بینی را دارند، رعایت مصلحت عمومی اقتضا می کند دولت با اتخاذ سازوکارهای نظارتی، تقنینی و حمایتی، زمینه کاهش خطا و تضمین جبران خسارت را فراهم آورد. در واقع، مصلحت با ناظر قرار دادن مقاصد شریعت و ضرورت حفظ نظم و امنیت اجتماعی، دولت را مکلف می سازد که مسئولیت مدنی خود را در قبال خطرات ناشی از هوش مصنوعی با رویکردی فعال، پیشگیرانه و سازگار با تحولات فناوری تنظیم کند.

انتفاعی، نارسایی مهمی ایجاد می‌کند؛ زیرا بسیاری از فعالیت‌های دولت در حوزه‌هایی چون امنیت، سلامت یا فناوری، غیرانتفاعی‌اند و نظریه قادر به تبیین مسئولیت بدون تقصیر در آن‌ها نیست (Zargoosh, 2010). افزون بر این، استدلال متقابل مبنی بر اینکه مستخدم نیز از اشتغال منتفع است، چالش دیگری برای این نظریه پدید می‌آورد (Zargoosh, 2010). به نظر می‌رسد که نظریه خطر ابزار مؤثری در اثبات مسئولیت مدنی دولت در قبال خسارات ناشی از هوش مصنوعی است، هرچند به نظر می‌رسد انتقاداتی چون افزایش ریسک اقتصادی نیز همچنان بر آن وارد باشد.

ب) اصل جبران کامل خسارت

اصل جبران کامل خسارت، به‌عنوان یکی از مبانی کلیدی حقوق مسئولیت مدنی، بر لزوم جبران تمامی زیان‌های ناروایی که بر اشخاص وارد می‌شود، تأکید دارد. این اصل که قدمتی کمتر از یک قرن دارد، تحول شگرفی را در حقوق مسئولیت مدنی، به‌ویژه پس از انقلاب صنعتی، رقم زده است. پیش از این تحولات، تنها خسارات مالی و جانی قابل جبران شناخته می‌شدند، اما امروزه دامنه شمول این اصل گسترده‌تر شده و خسارات معنوی، اقتصادی، عدم‌النفع، و حتی لطمه به زیبایی و امکان تفریح نیز در برخی نظام‌های حقوقی به رسمیت شناخته شده‌اند (Babaei, 2005). به نظر می‌رسد، در پرتو گسترش روزافزون کاربرد هوش مصنوعی و پیامدهای احتمالی آن، اصل جبران کامل خسارت، اهمیتی مضاعف می‌یابد. این اصل، دولت را مکلف می‌سازد تا در صورت بروز هرگونه زیان ناروا، اعم از مالی، معنوی یا جانی، ناشی از عملکرد سامانه‌های هوشمند، نسبت به جبران آن اقدام نماید.

بر اساس این اصل، کفایت عرف برای تشخیص قابلیت جبران خسارت، و عدم لزوم تصریح مقنن به آن، یکی از ویژگی‌های برجسته آن است. به عبارت دیگر، قاضی مسئولیت مدنی باید با اتکا به نظر عرف، حکم به جبران خسارت دهد و عدم تصریح

و تحقق آن مستلزم فعل هدفمند و هدایت‌شده انسان است (Jafari Langroudi, 2006). با این حال، در ماده ۳۲۸ قانون مدنی، عمد یا غیرعمد بودن در اتلاف تأثیری در تحقق ضمان ندارد و ماده ۳۸۹ نیز دامنه اتلاف را به معنای اعم در نظر گرفته است. فقه شیعه سه سبب اصلی برای ضمان قهری برمی‌شمارد: ید، اتلاف و تسبیب. آموزه قرآنی «فَمَنْ اَعْتَدَىٰ عَلَیْكُمْ فَاَعْتَدُوا عَلَیْهِ بِمِثْلِ مَا اَعْتَدَىٰ عَلَیْكُمْ» (بقره، ۱۹۴) زیربنای این قاعده را فراهم می‌کند و اصل «مَنْ اتلف مال الغير فهو له ضامن» صریح‌ترین بیان فقهی آن است (Najafi, 1988). اتلاف گاه حقیقی است، مانند سوزاندن مال، و گاه حکمی، مانند جلوگیری از انتفاع یا افت قیمت مال. به نظر می‌رسد در مواجهه با هوش مصنوعی، اتلاف حکمی مصداق روشن‌تری دارد؛ زمانی که دولت در نظارت یا اصلاح الگوریتم‌ها قصور ورزد و این بی‌توجهی موجب تلف مال یا حق اشخاص گردد، ضمان قهری بر او بار می‌شود. بر اساس بنای عقلایی نیز متلف همواره ضامن تلقی می‌گردد (Sami, 2021). بنابراین، دولت در قبال خسارات مستقیم یا حکمی ناشی از عملکرد هوش مصنوعی، مکلف به جبران است؛ خواه از راه مثل، خواه از طریق پرداخت قیمت، به نحوی که عدالت و جبران کامل محقق گردد.

مبانی حقوقی

الف) نظریه خطر

نظریه خطر توسط «ژوسران» و «سالی» بنیان‌گذاری شد و بر این مبنا استوار است که هرکس با فعالیت خود محیطی خطرزا ایجاد کند، باید مسئولیت پیامدهای آن را بپذیرد (Bahrami, 2007). مطابق این دیدگاه، مسئولیت بر دوش کسی قرار می‌گیرد که از فعالیت یا ابزار منتفع می‌شود، زیرا خطر نیز از همان منفعت‌بردن سرچشمه می‌گیرد و نمونه بارز آن، مسئولیت کارفرما در حوادث کار است (Zargoosh, 2010). با این حال، به نظر می‌رسد محدود کردن این نظریه به فعالیت‌های

قدرتمند برای اثبات ضمان دولت در صورت قصور در نظارت یا اصلاح الگوریتم‌ها فراهم می‌آورد. در کنار مبانی فقهی، نظریه خطر، علی‌رغم برخی انتقادات، به‌عنوان یک مبانی حقوقی، دولت را مسئول پیامدهای فعالیت‌های بالقوه خطرناک هوش مصنوعی می‌شناسد، حتی اگر انتفاع مادی یا تقصیری متوجه آن نباشد. همچنین، اصل جبران کامل خسارت، با گسترش دامنه شمول زیان‌های قابل جبران به معنوی و اقتصادی، بر لزوم پوشش تمامی خسارات وارده، حتی بدون تصریح مقنن، تأکید می‌ورزد. در مجموع، این مبانی فقهی و حقوقی، لزوم رویکردی مسئولانه، پیشگیرانه و جبران‌گرانه را از سوی دولت در مواجهه با چالش‌های ناشی از هوش مصنوعی، ایجاب می‌کنند.

مشارکت نویسندگان

در نگارش این مقاله تمامی نویسندگان نقش یکسانی ایفا کردند.

تعارض منافع

در انجام مطالعه حاضر، هیچ‌گونه تضاد منافی وجود ندارد.

EXTENDED ABSTRACT

The rapid expansion of artificial intelligence has created a new field of inquiry in public law, civil liability, and Islamic jurisprudence, especially where intelligent systems are used by governmental institutions in public administration, judicial assistance, environmental monitoring, urban management, security services, welfare allocation, and regulatory decision-making. Artificial intelligence is no longer limited to technical laboratories or private commercial applications; rather, it has entered the sphere of public power, where automated data processing, machine learning, expert systems, machine vision, natural language processing, robotics, neural networks, and genetic algorithms may directly or indirectly affect citizens' rights, public interests, and legally protected values. The distinctive feature of artificial intelligence lies in its capacity to

قانون‌گذار را مانعی برای جبران تلقی نکند. مبانی متعددی، از جمله اصل ۴۰ قانون اساسی، ماده ۲۲۱ قانون مدنی و ماده ۱ قانون مسئولیت مدنی، بر این اصل صحه می‌گذارند و بر لزوم تحقق جبران کامل خسارت در نظام حقوقی تأکید دارند.

نتیجه‌گیری

اصول عدالت، مصلحت و لاضرر، چارچوب‌های مستحکمی را برای تضمین حقوق شهروندان فراهم می‌آورند. اصل عدالت، دولت را به‌عنوان مجری نظم عادلانه، مکلف به جبران هرگونه تضییع حقوق یا ورود خسارت ناشی از خطاهای هوش مصنوعی می‌سازد. قاعده مصلحت، با رویکردی پویا، دولت را به پیشگیری و دفع مفاسد ناشی از این فناوری‌ها ترغیب کرده و ضرورت اتخاذ سازوکارهای نظارتی و حمایتی را گوشزد می‌کند. همچنین، قاعده لاضرر، با نفی هرگونه حکم ضرری و تأکید بر لزوم جبران زیان وارده، مانع از تضییع حقوق افراد در اثر قصور دولت در مدیریت هوش مصنوعی می‌شود. نیز می‌توان گفت که قاعده اتلاف، با تأکید بر لزوم جبران هرگونه تلف مال، چه حقیقی و چه حکمی، مبنایی

process large volumes of data, learn from previous inputs, generate predictive outputs, and perform quasi-autonomous functions that may not always be fully predictable by human operators. This characteristic raises a central legal question: when damage results from the operation, malfunction, deficient supervision, or improper regulatory use of artificial intelligence by the state, on what basis can governmental civil liability be established? The importance of this question becomes clearer when artificial intelligence is understood not merely as a technical tool, but as a decision-support or decision-making mechanism capable of producing consequences in areas traditionally governed by human discretion and legal accountability (Aghajani, 2021; Ahmadi et al., 2013; Ghaeminia, 2006).

From a conceptual perspective, artificial intelligence can be understood as a set of

computational systems designed to simulate or reproduce aspects of human intelligence, including perception, learning, reasoning, problem-solving, prediction, and decision-making. Its various branches demonstrate the breadth of this technological domain. Machine learning enables systems to learn from data and improve their outputs over time, while machine vision allows computers and robots to interpret visual information and transform numerical image data into meaningful patterns. Natural language processing connects computer science, linguistics, and artificial intelligence by enabling machines to analyze, interpret, and generate human language. Robotics extends artificial intelligence into physical environments, where machines interact with the external world through sensors, actuators, and programmed responses. Expert systems simulate the reasoning of specialists in particular fields, whereas artificial neural networks imitate selected features of biological nervous systems and are especially powerful in pattern recognition and prediction. Genetic algorithms, inspired by natural selection, are used to search for optimal solutions among complex sets of possibilities. These technical capacities explain why governmental reliance on artificial intelligence may generate both efficiency and risk: the same systems that enhance public decision-making may also cause damage when they are poorly designed, insufficiently trained, inadequately supervised, or deployed without proper legal safeguards (Biranvand & Sahraeyan, 2018; Ebrahimi & Pasha, 2011; Jafari et al., 2016; Katouz et al., 2017; Razavi Salestani et al., 2015; Sarabadi & Amir, 2018; Sivandian & Rastgari, 2014).

The manifestations of governmental civil liability for damages caused by artificial intelligence can be observed in several public domains, among which environmental protection, cultural heritage, and judicial

systems are particularly significant. In the environmental field, artificial intelligence may be used to monitor pollution, manage water resources, predict natural hazards, supervise industrial emissions, or analyze ecological data. If an algorithmic error, defective software code, insufficient updating of monitoring systems, or negligent governmental supervision causes environmental harm, the state's responsibility may arise from its failure to prevent foreseeable damage or from its improper reliance on automated systems. Environmental harm is especially important because it may involve widespread, cumulative, and sometimes irreversible damage to ecosystems and public health. In the field of cultural heritage, intelligent surveillance systems, satellite-image analysis tools, or automated monitoring technologies may be used to detect illegal excavation, unauthorized construction, or damage to historical sites. If these systems fail due to defective training, limited data, or insufficient governmental oversight, damage may be caused to cultural and historical assets that represent collective identity and intergenerational value. In the judicial sphere, artificial intelligence may assist in risk assessment, case analysis, prediction of recidivism, or decision support. If such systems lead to unjust decisions, discriminatory outcomes, or erroneous judicial consequences, the issue of state liability becomes directly connected to the protection of due process, human dignity, and the right to effective remedy (Arzhang et al., 2019; Ghasemi, 2012; Salehabadi et al., 2019; Sarikhan, 2005; Taghizadeh Ansari, 1995).

The jurisprudential foundations of governmental civil liability in relation to artificial intelligence can be explained through several major principles, including justice, public interest, the no-harm rule, and the rule of destruction. The principle of justice

occupies a central place in Islamic jurisprudence and provides a normative basis for requiring the state to protect individuals from unjust harm, especially when such harm results from technologies deployed within public institutions. If the government uses artificial intelligence in a manner that causes deprivation of rights, discriminatory treatment, loss of property, or unlawful restriction of liberty, justice requires that the injured party not be left without remedy. The principle of public interest also offers a dynamic framework for addressing new technological risks. Since public interest is linked to the prevention of corruption, the realization of benefit, and the preservation of social order, it can justify proactive governmental obligations in regulating, supervising, and correcting artificial intelligence systems before they cause harm. This principle is particularly suitable for emerging technologies because it permits legal reasoning to respond to new realities without abandoning the higher purposes of law and religion. Therefore, public interest requires the state to establish preventive mechanisms, regulatory standards, technical audits, and compensatory procedures for damages caused by intelligent systems (Dehkhoda, 1998; Feyz, 1994; Gorji, 1997; Ibn Manzur, 1988; Shartuni, 1983).

The no-harm rule and the rule of destruction further strengthen the jurisprudential basis for state liability. The no-harm rule rejects the legitimacy of legal arrangements that produce unjust harm and supports both the prevention of injury and the compensation of damage once it has occurred. In the context of artificial intelligence, if the state neglects its regulatory duties, fails to supervise high-risk systems, or continues to use defective algorithms despite knowledge of their harmful potential, the resulting injury can be understood through the logic of this rule. The rule of destruction also provides a clear basis for liability when a

person's property, right, or legally protected interest is destroyed or impaired. Although classical formulations often refer to direct destruction of property, the logic of the rule can also apply to constructive or indirect destruction, such as deprivation of benefit, loss of value, or damage caused through defective systems controlled or supervised by the state. In modern artificial intelligence cases, damage may not always occur through direct physical action; it may arise through algorithmic misclassification, wrongful denial of benefits, erroneous judicial assistance, environmental mismanagement, or automated decisions that produce material or moral harm. Accordingly, jurisprudential principles can support a broader theory of governmental liability that includes both direct and indirect damages, provided that causation, harm, and governmental responsibility can be established (Ansari; Hosseini Yazdi, 1980; Jafari Langroudi, 2006; Jafari Tabrizi; Kompani, 1954; Marashi, 2006; Mousavi Khomeini, 1989; Najafi, 1988; Rashti Najafi; Sami, 2021; Tarimoradi & Fakhilai, 2005).

In conclusion, the civil liability of the state for damages arising from artificial intelligence should be understood through a combined framework that integrates fault-based responsibility, risk-based responsibility, regulatory obligations, and the duty of full compensation. Artificial intelligence creates new forms of harm because its outputs may be complex, opaque, probabilistic, and difficult to attribute to a single human actor. For this reason, limiting liability only to traditional fault may leave many injured persons without adequate remedy. At the same time, imposing absolute liability in all cases may weaken technological development and administrative efficiency. A balanced approach requires the state to be held liable where it has failed to design, supervise, regulate, update, audit, or responsibly deploy artificial intelligence

systems, and also where the inherently risky nature of governmental use of such systems justifies compensation despite the absence of ordinary fault. The findings of this study indicate that Islamic jurisprudential principles and modern civil liability doctrines can jointly support a preventive, supervisory, and compensatory model of state responsibility. Under this model, the government is not merely a passive user of artificial intelligence, but a public authority obligated to ensure that technological innovation remains compatible with justice, public interest, protection from harm, and full reparation of unlawful damage.

References

- Aghajani, M. (2021). *Artificial Intelligence: From Introductory to Advanced*. Nasl-e Roshan.
- Ahmadi, S. A. A., Daraei, M., Salamzadeh, A., & Jafari, M. (2013). Artificial intelligence and business opportunities: Identifying the functions of artificial intelligence in creating competitive advantage for technology businesses: A study of the computer games industry. *Journal of Entrepreneurship Development*, 6(2), 7-26.
- Ansari, M. i. M. A. *The Rules of No-Harm, Hand, Validity, and Lottery (Fara'id al-Usul)*. Islamic Publications Office Affiliated with the Society of Seminary Teachers of Qom.
- Arzhang, A., Mirkhalili, S. A., & Bijani, M. (2019). Preserving or destroying historical works concerning sculpture and painting from jurisprudential and legal perspectives. *History of Islam Quarterly*, 20(78).
- Babaei, I. (2005). *Conditions of Compensable Damage*. Research work at Allameh Tabataba'i University.
- Bahrami Ahmadi, H., & Fahimi, A. (2007). The basis of environmental civil liability in jurisprudence and law. *Islamic Law Research Journal*(2).
- Biranvand, S., & Sahraeyan, K. (2018). Types of neural networks and their applications. Third International Conference on Management Research and Humanities,
- Dehkhoda, A. A. (1998). *Dehkhoda Dictionary*. University of Tehran.
- Ebrahimi, A., & Pasha, E. (2011). *Prediction Based on Artificial Neural Networks*. Tarbiat Moallem University, Faculty of Mathematics and Computer Science.
- Feyz, A. (1994). *Comparison and Adaptation in General Islamic Criminal Law*. Ministry of Culture and Islamic Guidance.
- Ghaeminia, A. (2006). Religion and artificial intelligence. *Zehn*, 7(25), 23-26.
- Ghasemi, N. (2012). *Environmental Criminal Law*. Jamal al-Haq.
- Gorji, A. (1997). *Islamic Law: Graduate Private Law Polycopy Notes*. University of Tehran.
- Hayati, A. A. (2010). *Introduction to Jurisprudence*. Mizan Legal Foundation.
- Hosseini Yazdi, M. (1980). *Jurisprudential Rules, Ijtihad, and Taqlid (Inayat al-Usul)*. Firoozabadi Bookstore.
- Ibn Manzur, M. i. M. (1988). *Lisan al-Arab*. Dar Ihya al-Turath al-Arabi.
- Jafari Langroudi, J. (2006). *Legal Terminology*. Ganj-e Danesh.
- Jafari, S., Ghanbary, A., & Sarbaz, Y. (2016). *Breast cancer diagnosis using machine vision algorithms*. University of Tabriz, Faculty of New Technologies Engineering.
- Jafari Tabrizi, M. T. *Sources of Jurisprudence*. No publisher.
- Katouz, S., Ramezani, M., & Amin Tahmasbi, H. (2017). *Examining factors affecting the application of expert systems in production planning and inventory control using ISM and DEMATEL techniques*. Rahbord Shomal Institute of Higher Education.
- Kompani, M. H. (1954). *Jurisprudential Rules, Ijtihad, and Taqlid (Nihayat al-Dirayah)*. Sayyid al-Shuhada Bookstore.
- Marashi, M. H. (2006). *New Perspectives in Law*. Mizan.
- Mousavi Khomeini, R. (1989). *Tahrir al-Wasilah*. Dar al-Kutub al-Ilmiyyah, Ismailiyan.
- Najafi, M. H. (1988). *Jawahir al-Kalam fi Sharh Shara'i al-Islam*. Dar al-Kutub al-Islamiyyah.
- Rashti Najafi, M. H. *The Book of Usurpation*. No publisher.
- Razavi Salestani, Z., Shahbahrani, A., & Mirroshandel, S. A. (2015). *Improving the efficiency of natural language processing applications using GPU*. University of Guilan, Faculty of Engineering.
- Salehabadi, R., Rasekh, M., & Talebpour, A. (2019). *Determining the limits of artificial intelligence research from the perspective of rights and public interest*. Shahid Beheshti University, Faculty of Law.
- Sami, M. (2021). Application of the rule of deception to the rules of destruction, causation, and no-harm in Iranian criminal law. *New Interdisciplinary Legal Research Quarterly*, 1(2), 1-12.
- Sarabadi, A., & Amir, F. (2018). *Artificial Intelligence and Expert Systems*. Jaliz.
- Sarikhani, A. (2005). *Crimes against Public Security and Order: Special Criminal Law 3*. University of Qom.
- Shartuni, S. i. (1983). *Aqrah al-Mawarid*. Library of Ayatollah al-Uzma Marashi Najafi.
- Sivandian, Z., & Rastgari, H. (2014). Examining the application of genetic algorithms in detecting

intrusion in computer networks. Second National Conference on Computer Science and Engineering,

Taghizadeh Ansari, M. (1995). *Environmental Law in Iran*. SAMT.

Tarimoradi, E., & Fakhlai, M. (2005). Jurisprudential foundations and rulings of the environment. *Islamic Studies*(71).

Zargoosh, M. (2010). *Civil Liability of the State*. Mizan.