

هوش مصنوعی در ترازوی قانون، مطالعه تطبیقی الگوهای جهانی تنظیم‌گری و الزامات حقوقی ایران

رضا فرج‌پور^۱، دیوید جی. گانکل^۲

چکیده

پیشرفت شتابان هوش مصنوعی، تنظیم‌گری این فناوری را به چالشی راهبردی در سطوح ملی و بین‌المللی بدل کرده است. این پژوهش با هدف تدوین چارچوبی نظام‌مند برای تحلیل تطبیقی حکمرانی هوش مصنوعی، یازده شاخص کلیدی طراحی و وضعیت پنج نظام حقوقی اتحادیه اروپا، ایالات متحده، چین، کانادا و جمهوری اسلامی ایران را تا تیرماه ۱۴۰۵ (ژوئیه ۲۰۲۶) با روش تحلیل اسنادی و تطبیقی مقررات، لوایح و اسناد سیاستی واکاوی می‌کند. یافته‌ها نشان می‌دهد اتحادیه اروپا با قانون هوش مصنوعی همچنان پیشروترین ساختار قانونی را داراست، هرچند چالش‌های اجرایی، پیاده‌سازی الزامات پریسک را به تعویق انداخته است؛ ایالات متحده پس از تغییر دولت در سال ۲۰۲۵ رویکرد مقررات‌زدایی و رقابت‌محوری را در پیش گرفته، کانادا در پی تعلیق پارلمان فاقد قانون فدرال مانده و چین سیاستی کنترل‌گرایانه و امنیت‌محور دنبال می‌کند. در ایران نیز با وجود در جریان بودن طرح ملی هوش مصنوعی در مجلس، خلأ قانونی همچنان محسوس است. نوآوری پژوهش، معرفی شاخص ترکیبی بلوغ حکمرانی هوش مصنوعی (ARMI) است که تبدیل تحلیل‌های کیفی به رتبه‌بندی عددی و پایش طولی نظام‌های حقوقی را ممکن می‌سازد. بر پایه تحلیل شکاف‌های موجود، پیشنهادهای راهبردی مشخصی برای تدوین قانون ملی هوش مصنوعی ایران بر مبنای چارچوب یازده شاخصه ارائه می‌شود.

کلیدواژه‌ها: هوش مصنوعی، تنظیم‌گری فناوری، چارچوب یازده شاخصه، شاخص بلوغ حکمرانی، قانون‌گذاری تطبیقی.

^۱ نویسنده مسئول، گروه حقوق، واحد یاسوج، دانشگاه آزاد اسلامی، یاسوج، ایران.

r.farajpour@iaau.ac.ir

^۲ استاد گروه فلسفه و اخلاق، دانشکده فناوری اطلاعات و ارتباطات، دانشگاه ایلینوی شمالی، ایالات متحده آمریکا.

dgunkel@niu.edu

AI on the Legal Balance, A Comparative Study of Global Regulatory Models and Legal Requirements in Iran

Reza. Farajpou^{1✉} , David. J. Gunkel² 

1. Corresponding author; Department of Private Law, Yasuj Islamic Azad University, Yasuj, Iran. Email: r.farajpour@iau.ac.ir

2. Professor, Department of Philosophy and Ethics, School of Information and Communication Technology, Northern Illinois University, United States of America. Email: dgunkel@niu.edu

Abstract:

The rapid advancement of artificial intelligence (AI) has turned its regulation into a strategic challenge at both national and international levels. To develop a systematic framework for the comparative analysis of AI governance, this study designs eleven key indicators and examines five legal jurisdictions, namely the European Union, the United States, China, Canada, and the Islamic Republic of Iran, up to July 2026, using documentary and comparative analysis of regulations, bills, and policy documents. The findings show that the European Union, through its AI Act, retains the most advanced legal structure, although implementation challenges have delayed the enforcement of high-risk requirements. Following the 2025 change of administration, the United States has adopted a deregulatory, competition-oriented approach; Canada remains without a federal AI law in the wake of parliamentary suspension; and China pursues a control-oriented, security centric policy. In Iran, despite the national AI bill under deliberation in parliament, a legislative vacuum persists. The innovation of this research is the AI Regulation Maturity Index (ARMI), a composite index that converts qualitative analyses into numerical rankings and enables longitudinal monitoring of legal systems. Based on a gap analysis, the paper offers concrete strategic recommendations for drafting Iran's national AI legislation under the eleven-indicator framework.

Keywords: Artificial Intelligence, Technology Regulation, Eleven Indicator Framework, Regulatory Maturity Index (ARMI), Comparative Legislation.

تحولات پرشتاب هوش مصنوعی، الگوهای سنتی حقوق و تنظیم‌گری را با پرسش‌های بنیادین مواجه ساخته است. این فناوری با پردازش انبوه داده‌ها، تصمیم‌گیری خودکار و یادگیری مستمر، مرزهای میان انسان و ماشین را بازتعریف کرده و آثار آن به حوزه‌های اجتماعی، اخلاقی و حقوقی گسترش یافته است (تخشید، ۱۴۰۰)؛ تأخیر در قانون‌گذاری می‌تواند حقوق بنیادین بشر، عدالت اجتماعی، امنیت ملی و اعتماد عمومی به فناوری‌های نو را تهدید کند. چالش اصلی نظام‌های حقوقی در مواجهه با هوش مصنوعی، فقدان چارچوبی جامع، عملیاتی و قابل‌سنجش برای تنظیم‌گری است. بسیاری از پژوهش‌های تطبیقی موجود یا به توصیف موردی یک کشور بسنده کرده‌اند یا مقایسه را بدون شاخص‌های صریح و قابل‌رهگیری انجام داده‌اند و در نتیجه سنجش پیشرفت یا عقب‌ماندگی نظام‌ها دشوار مانده است. این مقاله برای پاسخ به این خلأ، چارچوبی یازده شاخص تدوین و آن را به‌صورت نظام‌مند بر پنج کشور اعمال می‌کند تا هم ابزاری برای سنجش تطبیقی فراهم آورد و هم مسیری برای تدوین قانون ملی هوش مصنوعی ایران پیشنهاد دهد.

جایگاه ایران در تنظیم‌گری هوش مصنوعی نیازمند تحلیل به‌روز است: با وجود تدوین اسناد راهبردی و الزامات پراکنده، هنوز قانون الزام‌آور جامع و ساختار نظارتی مشخصی وجود ندارد؛ هرچند طرح ملی هوش مصنوعی مراحل بررسی در مجلس شورای اسلامی را طی می‌کند و برخی مواد اصلی آن از جمله تأسیس سازمان ملی هوش مصنوعی تصویب شده‌اند. برای تصریح مسیر پژوهش، این مقاله می‌کوشد به سه پرسش کلیدی پاسخ دهد:

نخست، چه شاخص‌هایی می‌توانند مبنای سنجش نظام‌مند و نظریه‌بنیان مقررات الزام‌آور هوش مصنوعی باشند؟

دوم، پنج نظام حقوقی منتخب بر پایه این شاخص‌ها در چه جایگاهی قرار دارند و این جایگاه چگونه به‌صورت عددی و قابل پایش بازنمایی می‌شود؟

سوم، از این مقایسه چه درس‌آموخته‌هایی برای تکمیل مسیر تقنینی ایران قابل استخراج است؟ ساختار مقاله بر همین اساس سامان یافته است: پس از مرور پیشینه و مبانی نظری، چارچوب یازده شاخص معرفی و بر پنج کشور اعمال می‌شود؛ سپس با معرفی شاخص ترکیبی بلوغ حکمرانی، تحلیل کیفی به سنجش عددی ارتقا می‌یابد و در پایان پیشنهادهای سیاستی و تقنینی برای ایران ارائه می‌گردد.

۱. پیشینه پژوهش

پژوهش در حوزه تنظیم‌گری هوش مصنوعی در سال‌های اخیر رشد چشمگیری داشته و به‌طور کلی در دو دسته قابل تفکیک است: پژوهش‌هایی که اصول اخلاقی و راهبردهای غیرالزام‌آور ارائه می‌دهند، و مطالعاتی که ساختارهای حقوقی الزام‌آور و مدل‌های تقنینی کشورها یا نهادهای بین‌المللی را تحلیل می‌کنند. در ادبیات خارجی، برخی پژوهش‌ها چارچوب‌های ارزیابی ریسک‌محور را پایه مقررات‌گذاری دانسته و برخی دیگر بر حقوق بشر، عدالت الگوریتمی و شفافیت فرایندهای تصمیم‌گیری متمرکز شده‌اند (Floridi, 2021; Smuha, 2021)؛ با این حال، در اغلب این پژوهش‌ها بررسی به یک کشور محدود بوده یا مقایسه تطبیقی به‌صورت توصیفی و بدون شاخص‌های صریح انجام شده است.

در ایران، پژوهش‌ها عمدتاً بر اخلاق، امنیت و پیامدهای اجتماعی هوش مصنوعی متمرکز بوده و کمتر به ساختار حقوقی و تنظیم‌گری الزام‌آور پرداخته‌اند (رجبی، ۱۳۹۸؛ سیفی و رزمخواه، ۱۴۰۱). در این میان، پژوهش‌ها نشان می‌دهند فقدان شخصیت حقوقی مستقل برای سامانه‌های هوشمند، اسناد مسئولیت را به الگوهای کلاسیک مسئولیت انسانی و شرکتی محدود می‌سازد (شریفی و بیرمی، ۱۳۹۷؛ حکمت‌نیا، محمدی و واثقی، ۱۳۹۸؛ گندم‌کار، صالحی مازندرانی و حمیدی، ۱۴۰۰)؛ فرج‌پور (Farajpour, 2025) در تحلیلی تطبیقی از نظام‌های آمریکا، اتحادیه اروپا، آلمان و ایران نشان می‌دهد رویکردهای مسئولیت مدنی هوش مصنوعی هنوز به الگویی یکپارچه نرسیده‌اند و حقوق ایران کماکان بر شخصیت حقوقی انسانی و شرکتی استوار است.

همچنین پژوهش‌هایی درباره پاسخ‌گویی الگوریتمی در حوزه‌های تخصصی‌تر مانند داوری خودکار در ورزش نشان می‌دهند مسئله شفافیت و امکان اعتراض به تصمیمات خودکار،^۱ محدود به حکمرانی عمومی نیست (Farajpour, Amerinia, & Pourjavaheri, 2025). با این حال، مطالعه نظام‌مند و شاخص‌محور مقررات رسمی کشورهای پیشرو و استخراج درس‌آموخته برای حقوق ایران همچنان کمتر موردتوجه بوده است؛ خلئی که مقاله حاضر با طراحی چارچوب یازده شاخصه و اعمال آن بر پنج کشور منتخب می‌کوشد پاسخ دهد.

۲. چارچوب نظری پژوهش

چارچوب نظری پژوهش بر بنیان‌های حقوقی، سیاست‌گذاری عمومی و فلسفه فناوری استوار است و از نظریه‌های تنظیم‌گری ریسک، حقوق دیجیتال و حاکمیت الگوریتمی بهره می‌گیرد. در ادبیات حقوقی اروپا، نظریه تنظیم‌گری ریسک مبنای اصلی مقررات‌گذاری است و فناوری‌های هوش مصنوعی را بر اساس میزان خطر برای حقوق بشر، ایمنی عمومی و ارزش‌های بنیادین طبقه‌بندی می‌کند (Floridi, 2021). در آمریکای شمالی، رویکرد حقوق دیجیتال فردمحور و حمایت از آزادی نوآوری غالب است. مفهوم حاکمیت الگوریتمی^۲ و نگرانی از اقتدار غیرشفاف سیستم‌های هوشمند نیز جدی گرفته شده است؛ فرانک پاسکواله^۳ و دیوید گانکل^۴ بر لزوم شفافیت، پاسخ‌گویی و بازنگری در مفاهیم مسئولیت و شخصیت حقوقی در برابر عاملان غیرانسانی تأکید دارند (Pasquale, 2015; Gunkel, 2018, 2022) و کوکلیبرگ نسبت به حذف دیدگاه انسانی از قانون‌گذاری در صورت اتکای صرف به بازار و خودتنظیمی هشدار می‌دهد (Coeckelbergh, 2020).

۲-۱. از نظریه تا شاخص: پیوند مبانی نظری با چارچوب یازده شاخصه

نوآوری اصلی پژوهش در گذار از سطح نظری به سطح سنجش‌پذیر و عملیاتی نهفته است: هر یک از یازده شاخص مستقیماً از یکی از سه جریان نظری فوق مشتق شده و همین پیوند صریح، بر خلاف بسیاری از مطالعات پیشین، امکان دفاع نظری از هر شاخص را فراهم می‌آورد. برای نمونه، شاخص‌های طبقه‌بندی ریسک و ممیزی و گزارش‌دهی از نظریه تنظیم‌گری ریسک‌محور (Floridi, 2021)، شاخص‌های شفافیت الگوریتمی و حریم خصوصی و داده از ادبیات حقوق دیجیتال و حق توضیح‌پذیری (Wachter et al., 2017)، شاخص‌های مسئولیت مدنی، مسئولیت کیفری و تعریف مفاهیم بنیادین از مناظرات فلسفی شخصیت حقوقی و حاکمیت الگوریتمی (Pasquale, 2015; Gunkel, 2018, 2022)، و شاخص‌های نهاد ناظر مستقل، ضمانت اجرا، مشارکت عمومی و هماهنگی بین‌المللی از ادبیات سیاست‌گذاری عمومی و حکمرانی نهادی (Cath et al., 2018; Ulnicane et al., 2021) اخذ شده‌اند. این نگاشت میان نظریه و شاخص در جدول زیر قابل ردیابی است.

^۱ برای نمونه‌ای از چالش پاسخ‌گویی حقوقی تصمیمات خودکار در حوزه‌ای غیر از حکمرانی عمومی، ر.ک. تحلیل داوری ورزشی مبتنی بر هوش مصنوعی (فناوری کمک‌داور ویدیویی و Hawk-Eye) که ابهام در شفافیت الگوریتمی و امکان اعتراض را برجسته می‌کند.

^۲ Algorithmic Governance

^۳ Frank Pasquale

^۴ David Gunkel

جدول ۱. نگاشت جریان‌های نظری به شاخص‌های یازده‌گانه

جریان نظری	شاخص‌های مرتبط	منبع نظری کلیدی
تنظیم‌گری ریسک محور	طبقه‌بندی ریسک؛ ممیزی و گزارش‌دهی؛ ضمانت اجرا	Floridi (2021)
حقوق دیجیتال و توضیح‌پذیری	شفافیت الگوریتمی؛ حریم خصوصی و داده	Wachter, Mittelstadt, & Floridi (2017)
حاکمیت الگوریتمی و شخصیت حقوقی	تعریف مفاهیم بنیادین؛ مسئولیت مدنی؛ مسئولیت کیفری	Pasquale (2015)؛ Gunkel (2018, 2022)
سیاست‌گذاری عمومی و حکمرانی نهادی	نهاد ناظر مستقل؛ مشارکت عمومی؛ هماهنگی بین‌المللی	Ulnicane et al. (2021)؛ Cath et al. (2018)

۳. روش‌شناسی پژوهش

۳-۱. نوع و روش پژوهش

این پژوهش از نوع کیفی و از حیث روش، تحلیل اسنادی^۱ توأم با رویکرد تطبیقی^۲ است. واحد تحلیل، متن رسمی مقررات، لوایح، دستورهای اجرایی و اسناد سیاستی الزام‌آور یا در دست تصویب در هر یک از پنج نظام حقوقی مورد مطالعه است؛ اسناد راهبردی غیرالزام‌آور (مانند بیانیه‌های اصولی) تنها به‌عنوان داده تکمیلی برای تشخیص جهت‌گیری آتی هر نظام به کار رفته‌اند، نه به‌عنوان شاهی بر وجود مقررات الزام‌آور.

۳-۲. معیار انتخاب کشورهای مورد مطالعه

انتخاب پنج نظام حقوقی بر دو معیار استوار است: نخست، هر یک از چهار نظام غیرایرانی نماینده یکی از فلسفه‌های تقنینی متمایز بخش هفتم است؛ اتحادیه اروپا نماینده رویکرد حقوق‌محور، ایالات متحده نماینده رویکرد بازارمحور و چین نماینده رویکرد اقتدارمحور است و کانادا به‌عنوان نمونه تلاش برای مدلی میانی و ترکیبی برگزیده شده است. دوم، ایران کانون اصلی تحلیل و کشور هدف پیشنهادهای پژوهش است. این چهار نظام از حیث حجم تولید و کاربرد هوش مصنوعی نیز در شمار پیشتازان جهان‌اند و بیشترین اسناد رسمی تنظیم‌گری را در دسترس قرار می‌دهند؛ معیاری که امکان مقایسه با نمونه‌های دارای داده کافی را تضمین می‌کند.

۳-۳. فرایند ساخت و اعتباریابی چارچوب یازده شاخصه

چارچوب یازده شاخصه از طریق کدگذاری موضوعی محتوای اسناد رسمی پنج نظام حقوقی و مرور نظام‌مند ادبیات نظری استخراج شده است: در گام نخست، مؤلفه‌های تکرارشونده در متن AI Act اتحادیه اروپا، چارچوب مدیریت ریسک NIST، مقررات هوش مصنوعی مولد چین و لایحه AIDA کانادا فهرست و دسته‌بندی شدند؛ در گام دوم، این مؤلفه‌ها با ادبیات نظری تطبیق داده شدند؛ و در گام سوم، مؤلفه‌های همپوشان ادغام و موارد حاشیه‌ای حذف شدند تا فهرست نهایی به یازده شاخص غیرهم‌پوشان و قابل کدگذاری برسد.^۳ برای هر شاخص در هر کشور، وضعیت بر اساس متن رسمی سند مربوط در یکی از چهار سطح قوی، متوسط، ضعیف یا فاقد رتبه‌بندی شده است.

برای فاصله‌گرفتن این رتبه‌بندی از قضاوت شهودی، هر رتبه با معیاری عینی و سه مؤلفه‌ای تعریف شده است: (الف) وجود مبنای هنجاری صریح؛ (ب) لازم‌الاجرا بودن آن مبنا در زمان نگارش؛ و (ج) پیش‌بینی ضمانت اجرا و نهاد مجری مشخص. رتبه قوی به معنای احراز هر سه مؤلفه، متوسط دو مؤلفه، ضعیف تنها یک مؤلفه یا وجود صرفاً اسناد غیرالزام‌آور، و فاقد به معنای عدم احراز هیچ‌یک است. برای نمونه در شاخص تعریف مفاهیم بنیادین، ماده ۳ قانون هوش مصنوعی اتحادیه اروپا تعاریف صریح، لازم‌الاجرا و دارای ضمانت اجرا ارائه می‌کند (رتبه قوی)؛ مقررات هوش مصنوعی مولد چین تعریف صریح و لازم‌الاجرا اما محدود

^۱ Documentary Analysis

^۲ Comparative Legal Analysis

^۳ این یازده شاخص برای اولین بار در همین پژوهش، بر مبنای مرور نظام‌مند اسناد رسمی پنج نظام حقوقی مورد مطالعه و ادبیات نظری بخش دوم و سوم مقاله، استخراج و کدگذاری شده‌اند؛ به این معنا که چارچوب از داده‌ها استخراج شده (data-driven) نه آنکه پیشاپیش از یک نظریه واحد وام گرفته شده باشد.

به یک زیرشاخه فناوری به دست می‌دهد (رتبه متوسط)؛ و در ایالات متحده پس از لغو دستور اجرایی ۱۴۱۱۰، تعاریف عمدتاً در اسناد غیرالزام‌آوری چون چارچوب NIST باقی‌مانده‌اند (رتبه ضعیف).

برای شفافیت فرایند کدگذاری، سه نمونه عملی از زنجیره (متن سند ← کد اولیه ← شاخص نهایی) ارائه می‌شود: الزام ارائه‌دهندگان به ثبت الگوریتم نزد اداره فضای سایبری در مقررات هوش مصنوعی مولد چین، ذیل شاخص ممیزی و گزارش‌دهی طبقه‌بندی شد؛ سازوکار ماده ۶ و ضمیمه سوم قانون هوش مصنوعی اتحادیه اروپا در تفکیک سامانه‌های پرریسک، ذیل شاخص طبقه‌بندی ریسک قرار گرفت؛ و ماده ناظر بر تأسیس سازمان ملی هوش مصنوعی در طرح ملی ایران ذیل شاخص نهاد ناظر مستقل کدگذاری شد که چون کلیت طرح هنوز لازم‌الاجرا نشده و نهاد عملاً تشکیل نگردیده، تنها مؤلفه نخست را احراز می‌کند. این نمونه‌ها امکان ردیابی و ارزیابی مستقل منطق کدگذاری را فراهم می‌آورند.

۴-۳. محدودیت‌های روش‌شناختی

این پژوهش با دو محدودیت اصلی مواجه است: نخست، ماهیت به‌شدت پویای تنظیم‌گری هوش مصنوعی در سال‌های ۲۰۲۵ و ۲۰۲۶ که ممکن است برخی یافته‌ها (به‌ویژه درباره آمریکا و کانادا) را در فاصله کوتاهی پس از انتشار دگرگون کند؛ نگارندگان با ارجاع صریح به تاریخ هر رویداد کوشیده‌اند این پویایی را مدیریت کنند. دوم، رتبه‌بندی چهارسطحی هرچند بر معیار عینی سه‌مؤلفه‌ای بخش ۳-۳ استوار است، توسط نگارندگان و نه دو کدگذار مستقل انجام شده و از این رو ضریب پایایی بین کدگذار گزارش نمی‌شود؛ پژوهش‌های آتی می‌توانند با دو کدگذار مستقل، پایایی این چارچوب را به‌صورت کمی بیازمایند.

۴. چارچوب یازده شاخصه تحلیلی تطبیقی مقررات هوش مصنوعی

جدول زیر یازده شاخص پیشنهادی این پژوهش را به همراه پرسش ارزیابی مربوط به هر شاخص ارائه می‌دهد. این چارچوب در بخش ششم بر پنج نظام حقوقی مورد مطالعه اعمال شده است.

جدول ۲. شاخص‌های یازده‌گانه و پرسش‌های ارزیابی

ردیف	شاخص	پرسش ارزیابی
۱	تعریف مفاهیم بنیادین	آیا هوش مصنوعی، عامل خودکار، تصمیم‌گیری الگوریتمی و یادگیری ماشین در متن قانونی تعریف شده‌اند؟
۲	طبقه‌بندی ریسک سامانه‌ها	آیا سامانه‌های هوش مصنوعی بر مبنای سطح ریسک طبقه‌بندی و مقررات متناسب با هر سطح اعمال شده است؟
۳	مسئولیت مدنی	آیا سازوکار مشخصی برای جبران خسارت ناشی از سامانه‌های هوشمند (توسعه‌دهنده، بهره‌بردار یا محصول) وجود دارد؟
۴	مسئولیت کیفری	آیا نظام حقوقی تکلیف مسئولیت کیفری در تصمیم‌گیری خودکار را روشن کرده است؟
۵	نهاد ناظر مستقل	آیا نهاد ناظر مستقل با اختیار صدور مجوز، ثبت و اعمال ضمانت اجرا تأسیس شده است؟
۶	شفافیت الگوریتمی	آیا الزام به افشای منطق تصمیم‌گیری و حق اعتراض کاربر پیش‌بینی شده است؟
۷	ممیزی و گزارش‌دهی	آیا ممیزی مستقل دوره‌ای و گزارش‌دهی الزامی سامانه‌های پرریسک وجود دارد؟
۸	حریم خصوصی و داده	آیا حمایت از داده‌های شخصی در برابر پردازش الگوریتمی به طور خاص تضمین شده است؟
۹	ضمانت اجرا	آیا نظام جریمه یا تعزیرات مشخص و بازدارنده برای نقض مقررات وجود دارد؟
۱۰	مشارکت عمومی	آیا سازوکار رسمی برای مشارکت ذی‌نفعان (جامعه مدنی، صنعت، دانشگاه) در تدوین مقررات پیش‌بینی شده

	است؟	
۱۱	هماهنگی و چابکی بین‌المللی	آیا نظام حقوقی سازوکاری برای هماهنگی با استانداردهای بین‌المللی و به‌روزرسانی چابک مقررات دارد؟

۵. تحلیل تطبیقی مقررات هوش مصنوعی در کشورهای منتخب

۵-۱. اتحادیه اروپا

اتحادیه اروپا طی سال‌های گذشته پیشروترین بازیگر جهانی تنظیم‌گری هوش مصنوعی بوده است. قانون هوش مصنوعی اتحادیه^۱ که در ژوئن ۲۰۲۴ به تصویب نهایی رسید، الگویی ساختارمند و ریسک‌محور ارائه می‌دهد (European Commission, 2021; Veale & Zuiderveen Borgesius, 2024) و سامانه‌ها را به چهار دسته ریسک غیرقابل قبول، بالا، محدود و پایین طبقه‌بندی می‌کند؛ برای نمونه امتیازدهی اجتماعی دولتی یا شناسایی چهره در فضاهای عمومی بدون رضایت ممنوع شده است (Hacker et al., 2023). اجرای عملی این قانون در مرحله کنونی با تعدیل زمانی همراه شده است (صادقی، رشوند بوکانی و ناصر، ۱۴۰۲)^۲ و کمیسیون اروپا پیش‌نویس دستورالعمل مسئولیت مدنی هوش مصنوعی^۳ را که قرار بود مکمل AI Act در حوزه جبران خسارت باشد، رسماً کنار گذاشت.^۴

سازوکارهای اجرایی قانون چندلایه‌اند: در سطح ملی نهادهای نظارتی کشورهای عضو و در سطح اتحادیه، دفتر هوش مصنوعی اروپا^۵ وظیفه هماهنگی و راهبری را بر عهده دارند.^۶ ضمانت اجراها شامل جریمه تا سقف ۳۵ میلیون یورو یا هفت درصد گردش مالی جهانی سالانه (هر کدام بیشتر باشد) است؛ رویکردی در امتداد استدلال‌های پیشین درباره لزوم شفافیت و توضیح‌پذیری الگوریتمی (Wachter et al., 2017). مشارکت گسترده ذی‌نفعان نیز از ویژگی‌های بارز تدوین این قانون بود. در بُعد هماهنگی بین‌المللی نیز اتحادیه اروپا فعال‌ترین بازیگر میان پنج نظام مورد مطالعه است: در سطح درونی، هیئت اروپایی هوش مصنوعی^۷ اجرای یکپارچه قانون را میان بیست‌وهفت دولت عضو هماهنگ می‌کند و در سطح بیرونی، اتحادیه هم در تدوین توصیه‌نامه هوش مصنوعی سازمان همکاری و توسعه اقتصادی مشارکت فعال داشته (OECD, 2019) و هم از نخستین امضاکنندگان کنوانسیون چارچوب شورای اروپا^۸ نخستین معاهده الزام‌آور بین‌المللی این حوزه بوده است (Council of Europe, 2024).

^۱ AI Act

^۲ در ۷ مه ۲۰۲۵ نهادهای اتحادیه اروپا بر سر بسته‌ی اصلاحی دیجیتال (Digital Omnibus on AI) به توافق سیاسی موقت رسیدند که موعد الزامات مربوط به سامانه‌های پرریسک ضمیمه‌ی سوم را از ۲ اوت ۲۰۲۶ به ۲ دسامبر ۲۰۲۷ و الزامات سامانه‌های ضمیمه‌ی یکم را به ۲ اوت ۲۰۲۸ موکول می‌کند؛ در مقابل، ممنوعیت‌های ماده‌ی ۵ (از جمله سامانه‌های امتیازدهی اجتماعی) و الزامات مربوط به مدل‌های هوش مصنوعی عمومی منظوره (General-Purpose AI) که پیش‌تر لازم‌الاجرا شده بودند، بدون تغییر باقی می‌مانند.

^۳ AI Liability Directive

^۴ کمیسیون اروپا پیش‌نویس دستورالعمل مسئولیت مدنی هوش مصنوعی (AI Liability Directive) را که از سال ۲۰۲۲ در دست بررسی بود، در برنامه‌ی کاری ۲۰۲۵ خود کنار گذاشت و برداشتن رسمی آن در ۱۱ فوریه‌ی ۲۰۲۵ در روزنامه‌ی رسمی اتحادیه اروپا منتشر شد. در غیاب این سند، خسارات ناشی از محصولات هوشمند معیوب عمدتاً تابع دستورالعمل بازنگری‌شده‌ی مسئولیت محصول (Product Liability Directive (EU) 2024/2853) و مقررات داخلی کشورهای عضو باقی می‌ماند.

^۵ European AI Office

^۶ دفتر هوش مصنوعی اروپا، مستقر در کمیسیون اروپا، مسئولیت نظارت بر مدل‌های هوش مصنوعی عمومی منظوره (GPAI)، تدوین آیین‌نامه‌های عملی (Codes of Practice) و هماهنگی با نهادهای ناظر ملی را بر عهده دارد؛ نخستین آیین‌نامه‌ی عملی GPAI در سال ۲۰۲۵ با مشارکت بیش از ۱۰۰۰ ذی‌نفع نهایی شد.

^۷ هیئت اروپایی هوش مصنوعی (European Artificial Intelligence Board)، نهاد هماهنگی متشکل از نمایندگان کشورهای عضو، موضوع ماده‌ی ۶۵ قانون هوش مصنوعی اتحادیه اروپا.

^۸ کنوانسیون چارچوب شورای اروپا درباره‌ی هوش مصنوعی و حقوق بشر، دموکراسی و حاکمیت قانون (Council of Europe Framework Convention on Artificial Intelligence, CETS No. 225)، گشوده شده برای امضا از ۵ سپتامبر ۲۰۲۴؛ نخستین معاهده‌ی بین‌المللی الزام‌آور در حوزه‌ی هوش مصنوعی.

همین ترکیب از هماهنگی درونی و معاهده‌سازی بیرونی، بر پایه معیار سه‌مؤلفه‌ای بخش ۳-۳، رتبه قوی اتحادیه را در شاخص هماهنگی بین‌المللی توجیه می‌کند.

۵-۲. ایالات متحده آمریکا

ایالات متحده به‌رغم جایگاه پیشرو خود در توسعه فناوری‌های هوش مصنوعی، فاقد یک نظام قانون‌گذاری جامع، واحد و الزام‌آور در سطح فدرال است. سیاست آمریکا در این حوزه در فاصله کوتاهی به‌شدت دگرگون شده است. در اکتبر ۲۰۲۳ دولت بایدن دستور اجرایی ۱۴۱۱۰ را با هدف توسعه مسئولانه و ایمن سامانه‌های هوش مصنوعی صادر کرد، اما این دستور در نخستین روز دور دوم دولت ترامپ لغو و با رویکردی کاملاً متفاوت جایگزین شد.^۱

برنامه اقدام هوش مصنوعی آمریکا: برنده شدن در مسابقه^۲ (منتشر شده در ۲۳ ژوئیه ۲۰۲۵) بیش از ۹۰ اقدام فدرال را حول سه محور تسریع نوآوری، ساخت زیرساخت و رهبری بین‌المللی معرفی و به‌صراحت بر کاهش موانع مقرراتی، حذف ارجاعات عدالت الگوریتمی و تنوع از چارچوب NIST و محدودسازی مقررات ایالتی تمرکز می‌کند (The White House, Office of Science and Technology Policy, 2025). در دسامبر ۲۰۲۵ نیز دستور اجرایی دیگری وزارت بازرگانی را مکلف به ارزیابی قوانین ایالتی ناسازگار با سیاست ملی و مشروط کردن برخی کمک‌های فدرال کرد؛ حال آنکه ایالت‌هایی چون کالیفرنیا و کلرادو مقررات خاص خود را دنبال می‌کنند و خطر چندپارگی حقوقی را افزایش می‌دهند. در نتیجه، آمریکا در سال ۲۰۲۶ بیش از هر زمان بر خودتنظیمی صنعت و رقابت‌پذیری تکیه دارد و توصیف پیشین نظام رو به تکامل به‌سوی چارچوب فدرال منسجم عملاً از اعتبار ساقط شده است؛ تحلیل‌گرانی چون آلوندرا نلسون این وضعیت را نه فقدان تنظیم‌گری، بلکه بازآرایی آن از طریق سیاست صنعتی و محدودیت‌های تجاری و مهاجرتی توصیف کرده‌اند.

۵-۲-۱. مطالعه موردی: فراز و فرود قانون هوش مصنوعی کلرادو به‌عنوان آزمایشگاه فدرالیسم آمریکایی

سرگذشت قانون هوش مصنوعی ایالت کلرادو نمونه‌ای گویا از تنش میان تنظیم‌گری ایالتی و سیاست فدرال در آمریکای پس از ۲۰۲۵ است.^۳ این تجربه اهمیت تحلیلی دوگانه دارد: نخست نشان می‌دهد حتی در غیاب قانون فدرال، برخی ایالت‌ها (کالیفرنیا، کلرادو) کوشیده‌اند مدل ریسک‌محور شبیه AI Act اروپا را در سطح محلی پیاده کنند؛ و دوم، این تلاش‌های ایالتی اکنون با فشار مستقیم دولت فدرال از جمله دخالت وزارت دادگستری در دعاوی علیه قوانین ایالتی با موانع جدی روبه‌رو شده‌اند. این الگو مصداق عینی تکه‌تکه شدن سیاستی است که نه‌فقط میان کشورها، بلکه میان ایالت‌های یک کشور واحد نیز رخ می‌دهد.

۵-۳. جمهوری خلق چین

چین در سال‌های اخیر به یکی از تأثیرگذارترین بازیگران توسعه و تنظیم‌گری هوش مصنوعی بدل شده و سیاست آن مبتنی بر دولت‌محوری، تمرکزگرایی و الزام‌آوری فوری است. برجسته‌ترین سند این حوزه، مقررات اجرایی خدمات هوش مصنوعی مولد^۴ است که از اول آگوست ۲۰۲۳ لازم‌الاجرا شده و ارائه‌دهندگان مدل‌های زبانی و تولید محتوا را ملزم به رعایت الزامات ایمنی،

^۱ دستور اجرایی ۱۴۱۱۰ در نخستین روز دور دوم ریاست‌جمهوری ترامپ (۲۰ ژانویه ۲۰۲۵) لغو و با دستور اجرایی ۱۴۱۷۹ با عنوان رفع موانع رهبری آمریکا در هوش مصنوعی (Removing Barriers to American Leadership in AI، ۲۳ ژانویه ۲۰۲۵) جایگزین شد. در ادامه، برنامه‌ی اقدام هوش مصنوعی آمریکا (۲۳ ژوئیه ۲۰۲۵) با محوریت مقررات‌زدایی و رقابت جهانی منتشر گردید و در دسامبر ۲۰۲۵ دستور اجرایی دیگری برای محدود کردن مقررات ایالتی هوش مصنوعی (از جمله کالیفرنیا و کلرادو) صادر شد.

^۲ Winning the Race: America's AI Action Plan

^۳ کلرادو در ۱۷ مه ۲۰۲۴ با تصویب SB 24-205 نخستین ایالت آمریکایی بود که چارچوبی جامع و ریسک‌محور، با الهام آشکار از AI Act اروپا، برای سامانه‌های پریسک وضع کرد. اجرای این قانون دوبار (از فوریه به ژوئن ۲۰۲۶) به تعویق افتاد؛ در آوریل ۲۰۲۶ شرکت xAI با استناد به اصل اول متمم قانون اساسی و بند تجارت خفته، از این قانون شکایت کرد و وزارت دادگستری آمریکا برای نخستین بار در تاریخ بر پایه‌ی دستور اجرایی دسامبر ۲۰۲۵ رئیس‌جمهور، به نفع خواهان وارد دعا شد. سرانجام در ۱۲ مه ۲۰۲۶ فرماندار کلرادو با امضای SB ۱۸۹-۲۶، قانون اصلی را با چارچوبی سبک‌تر و متمرکز بر افشا و شفافیت (نه مدیریت ریسک الزام‌آور) جایگزین کرد؛ این فرجام، نمونه‌ای صریح از فشار فدرال برای همگون‌سازی رو به پایین مقررات ایالتی آمریکاست.

^۴ Generative AI Measures

شفافیت الگوریتمی، فیلترینگ محتوا و ثبت الگوریتم نزد اداره فضای سایبری چین^۱ کرده است.^۲ رویکرد چین نه تنها فناوری محور بلکه ایدئولوژی محور است و نظارت اغلب پیشگیرانه و پیش از عرضه انجام می گیرد. این ترکیب از الزام آوری فوری، تمرکز نهادی و کنترل محتوایی، چین را در شاخص های نهاد ناظر و ضمانت اجرا نسبتاً قوی اما در مشارکت عمومی و حریم خصوصی و داده ضعیف قرار می دهد؛ ترکیبی که آن را از الگوی اروپایی متمایز می سازد.

۴-۵. کانادا

کانادا در سال ۲۰۱۷ نخستین راهبرد ملی هوش مصنوعی جهان را تدوین کرد و در سال ۲۰۲۲ لایحه C-۲۷ حاوی قانون هوش مصنوعی و داده ها^۳ را به پارلمان ارائه داد. این لایحه قرار بود نخستین تلاش رسمی دولت فدرال کانادا برای وضع مقررات الزام آور در حوزه هوش مصنوعی باشد و کاربردهای پرریسک را بر مبنای طبقه بندی ریسک تنظیم کند. با این حال، برخلاف آنچه در بسیاری از مطالعات تطبیقی پیشین آمده، این لایحه هرگز به تصویب نهایی نرسید.^۴

در نتیجه، کانادا در تیرماه ۱۴۰۵ (ژوئیه ۲۰۲۶) بدون قانون فدرال هوش مصنوعی است و آنچه موجود است، صرفاً برنامه های نهادی نظیر مؤسسه ایمنی هوش مصنوعی کانادا^۵ و مقررات پراکنده استانی از جمله لایحه ۱۹۴ انتاریو برای بخش عمومی است. بدین ترتیب کانادا، بر خلاف جایگاه الگوی میانی و متعادل در ادبیات پیشین، اکنون از حیث تصویب قانون الزام آور فدرال، عقب مانده ترین جایگاه را میان پنج نظام مورد مطالعه دارد.

۵-۵. جمهوری اسلامی ایران

جمهوری اسلامی ایران تا زمان نگارش این مقاله فاقد قانون جامع، الزام آور و نظام مند در زمینه ی هوش مصنوعی است. مهم ترین سند راهبردی موجود، سند ملی هوش مصنوعی مصوب شورای عالی انقلاب فرهنگی (۱۴۰۳) است که فاقد ضمانت اجرا و ساختار حقوقی الزام آور است. با این حال، در کنار این سند راهبردی، مسیر تقنینی موازی دیگری نیز در حال طی شدن است. در اردیبهشت ۱۴۰۴ کلیات طرح ملی هوش مصنوعی در مجلس شورای اسلامی به تصویب رسید و در ماه های بعد، موادی از آن از جمله تشکیل شورای ملی راهبری هوش مصنوعی به ریاست رئیس جمهور و تأسیس سازمان ملی هوش مصنوعی زیر نظر ایشان در بررسی جزئیات به تصویب رسیدند.^۶

این طرح در صورت تکمیل و لازم الاجرا شدن می تواند نخستین گام رسمی ایران در تدوین ساختار نهادی حکمرانی هوش مصنوعی باشد؛ اما تا تیرماه ۱۴۰۵ هنوز نهاد ناظر مستقل، تعریف حقوقی مفاهیم بنیادین، طبقه بندی ریسک، نظام مسئولیت مدنی و کیفری مشخص و الزامات شفافیت الگوریتمی به طور صریح در متن قانونی لازم الاجرا پیش بینی نشده اند و ایران در اغلب شاخص های یازده گانه با خلأ ساختاری مواجه است. فرج پور و همکاران (۱۴۰۴) نشان می دهند فقدان شخصیت حقوقی مستقل برای سامانه های هوشمند در حقوق ایران، امکان طرح مستقیم دعاوی علیه این سامانه ها را از بین برده و مسئولیت را یا متوجه اشخاص انسانی می کند یا در فضای مبهم رها می سازد؛ یافته ای هم راستا با هشدارهای گانکل درباره بحران مفهومی نظام های حقوقی کلاسیک (Gunkel, 2022). نمونه های کاربرد بدون چارچوب در ایران، از تشخیص چهره در حمل و نقل عمومی

^۱ CAC (Cyberspace Administration of China)

^۲ تا پایان سال ۲۰۲۵ بیش از هزار الگوریتم توصیه گر و مولد نزد اداره ی فضای سایبری چین (CAC) به ثبت رسیده اند؛ عدم ثبت یا تخلف از الزامات محتوایی می تواند به تعلیق یا لغو مجوز بهره برداری از سامانه بینجامد.

^۳ AIDA

^۴ با استعفای نخست وزیر جاستین ترودو و تعلیق کار پارلمان کانادا در ۶ ژانویه ۲۰۲۵، لایحه ی C-27 و به تبع آن قانون هوش مصنوعی و داده ها (AIDA) (AIDA) به طور کامل از دستور کار خارج و عملاً منتفی شد. وزیر نوآوری کانادا در ژوئن ۲۰۲۵ تصریح کرد که AIDA در شکل پیشین باز نمی گردد و رویکرد آتی سبک، دقیق و متناسب (light, tight, right) خواهد بود؛ ایالت انتاریو نیز با لایحه ی ۱۹۴ مقرراتی برای بخش عمومی وضع کرده است.

^۵ CAISI

^۶ کلیات طرح ملی هوش مصنوعی در ۲۶ اردیبهشت ۱۴۰۴ در صحن علنی مجلس تصویب و برای بررسی جزئیات به کمیسیون صنایع و معادن ارجاع شد. در ادامه، ماده ی ۳ (تأسیس سازمان ملی هوش مصنوعی زیر نظر رئیس جمهور) در ۴ آبان ۱۴۰۴ و ماده ی ۲ (تشکیل شورای ملی راهبری هوش مصنوعی به ریاست رئیس جمهور) در ۲۳ مهر ۱۴۰۴ به تصویب رسیدند؛ با این حال، تا زمان نگارش این سطور، طرح هنوز کلیت خود را در قالب قانون لازم الاجرا طی نکرده و برخی مواد آن کماکان در نوبت بررسی شور دوم است.

و الگوریتم‌های پیش‌بینی در بانک و بیمه تا داده کاوی در حوزه بهداشت، فاقد سازوکار شفاف‌سازی یا حق اعتراض کاربران‌اند. بر همین اساس و با کاربست ضابطه سه مؤلفه‌ای بخش ۳-۳، ایران در شاخص نهاد ناظر مستقل رتبه ضعیف و نه بالاتر می‌گیرد: تصویب ماده تأسیس سازمان ملی هوش مصنوعی صرفاً مؤلفه نخست (مبنای هنجاری صریح) را احراز می‌کند و تا لازم‌الاجرا شدن کلیت طرح و تشکیل عملی نهاد، دو مؤلفه دیگر مفقودند؛ اتخاذ رتبه بالاتر، خلط اراده تقنینی با واقعیت نهادی می‌بود.

۵-۱. جدول جمع‌بندی: ارزیابی تطبیقی پنج نظام حقوقی بر مبنای چارچوب یازده شاخصه

مقیاس رتبه‌بندی به شرح زیر است^۱:

جدول ۳. ارزیابی تطبیقی پنج نظام حقوقی بر مبنای چارچوب یازده شاخصه

شاخص	اتحادیه اروپا	آمریکا	چین	کانادا	ایران
۱. تعریف مفاهیم بنیادین	قوی	ضعیف	متوسط	ضعیف	ضعیف
۲. طبقه‌بندی ریسک	قوی	فاقد	متوسط	ضعیف	فاقد
۳. مسئولیت مدنی	متوسط	متوسط	ضعیف	ضعیف	ضعیف
۴. مسئولیت کیفری	ضعیف	فاقد	متوسط	فاقد	فاقد
۵. نهاد ناظر مستقل	قوی	ضعیف	قوی	متوسط	ضعیف
۶. شفافیت الگوریتمی	قوی	ضعیف	ضعیف	ضعیف	فاقد
۷. ممیزی و گزارش‌دهی	قوی	متوسط	متوسط	ضعیف	فاقد
۸. حریم خصوصی و داده	قوی	ضعیف	ضعیف	ضعیف	ضعیف
۹. ضمانت اجرا	قوی	ضعیف	قوی	ضعیف	فاقد
۱۰. مشارکت عمومی	قوی	متوسط	ضعیف	متوسط	ضعیف
۱۱. هماهنگی بین‌المللی	قوی	ضعیف	ضعیف	ضعیف	فاقد

چنان‌که جدول نشان می‌دهد، اتحادیه اروپا در هشت شاخص از یازده شاخص وضعیت قوی دارد، حال آنکه ایران در هیچ شاخصی به سطح قوی و در اغلب موارد حتی متوسط نرسیده است. نکته درخور توجه، افت جایگاه کانادا در پی توقف لایحه AIDA است؛ به گونه‌ای که در بیشتر شاخص‌ها به سطح ضعیف یا فاقد سقوط کرده و در مجموع به آمریکا نزدیک شده است.

۵-۶. شاخص ترکیبی بلوغ حکمرانی هوش مصنوعی (نواوری روش‌شناختی پژوهش)

برای فراتر رفتن از توصیف کیفی و ارائه ابزاری قابل‌رهگیری، این پژوهش رتبه‌بندی کیفی جدول پیشین را به شاخص ترکیبی بلوغ حکمرانی هوش مصنوعی (ARMI)^۲ تبدیل می‌کند: به هر رتبه امتیازی عددی اختصاص می‌یابد (قوی = ۳، متوسط = ۲، ضعیف = ۱، فاقد = ۰) و جمع امتیازات یازده شاخص برای هر کشور، عددی میان صفر تا سی‌وسه به دست می‌دهد که سنجه‌ای خلاصه، هرچند غیردقیق از حیث آماری، برای رتبه‌بندی نسبی نظام‌های حقوقی است. این شاخص بر خلاف رتبه‌بندی‌های تجاری مبتنی بر ادراک عمومی یا شاخص‌های کلان اقتصادی، مستقیماً از تحلیل متن رسمی مقررات مشتق شده است.

جدول ۴. شاخص ترکیبی بلوغ حکمرانی هوش مصنوعی (ARMI) در پنج نظام حقوقی

سطح بلوغ تفسیری	نسبت به سقف (از ۳۳)	امتیاز خام	نظام حقوقی
-----------------	---------------------	------------	------------

^۱ مقیاس ارزیابی بر پایه‌ی معیار سه مؤلفه‌ای بخش ۳-۳ قوی احراز هر سه مؤلفه (مبنای هنجاری صریح، لازم‌الاجرا بودن، ضمانت اجرا و نهاد مجری)؛ متوسط احراز دو مؤلفه؛ ضعیف احراز تنها یک مؤلفه یا وجود صرفاً اسناد غیرالزام‌آور؛ فاقد عدم احراز هیچ مؤلفه‌ای در اسناد رسمی تا زمان نگارش.

^۲ AI Regulatory Maturity Index

اتحادیه اروپا	۳۰	۳۳/۳۰	بالا
چین	۱۹	۳۳/۱۹	متوسط روبه بالا
ایالات متحده	۱۲	۳۳/۱۲	پایین (روبه کاهش)
کانادا	۱۲	۳۳/۱۲	پایین (روبه کاهش شدید)
ایران	۵	۳۳/۵	بسیار پایین

محاسبه عددی نشان می‌دهد شکاف میان اتحادیه اروپا (۳۰ از ۳۳) و ایران (۵ از ۳۳) بزرگ‌ترین فاصله میان دو سر طیف پنج نظام مورد مطالعه است. نکته درخور توجه دیگر، برابری دقیق امتیاز آمریکا و کانادا (هر یک ۱۲) است؛ دو کشوری که در ادبیات پیشین الگوی بازارمحور باثبات و الگوی میانی متعادل توصیف می‌شدند، در تیرماه ۱۴۰۵ عملاً به یک سطح بلوغ رسیده‌اند؛ یافته‌ای که بدون این شاخص ترکیبی به‌سختی از مقایسه صرفاً کیفی قابل استخراج بود. چین نیز با ۱۹ امتیاز در رتبه دوم است؛ نتیجه‌ای که نشان می‌دهد الزام‌آوری و سرعت اجرا لزوماً با معیارهای حقوق بشری و مشارکتی هم‌بسته نیست و این دو بُعد باید در ارزیابی حکمرانی از هم تفکیک شوند. ARMI را می‌توان در پژوهش‌های آتی با وزن‌دهی متفاوت به شاخص‌ها بازتولید کرد؛ این پژوهش عمداً وزن یکسان به کار برده تا از سوگیری پیشینی در اولویت‌بخشی به بُعدی خاص پرهیز شود.

افزون بر کارکرد مقایسه‌ای، شاخص ARMI امکان پایش طولی نیز فراهم می‌آورد؛ چون عددی است، می‌توان آن را در فواصل زمانی معین بازمحاسبه و مسیر صعود یا نزول هر نظام حقوقی را در طول زمان ترسیم کرد؛ کاری که توصیف کیفی صرف ممکن نمی‌سازد. برای نمونه، اگر ایران طرح ملی هوش مصنوعی را با گنجاندن صریح طبقه‌بندی ریسک و الزامات شفافیت به قانون لازم‌الاجرا تبدیل کند، امتیاز آن دست‌کم در سه شاخص ارتقا می‌یابد و این جهش به‌صورت کمی قابل نمایش و دفاع خواهد بود؛ بدین‌سان ARMI نه صرفاً عکسی لحظه‌ای از وضع موجود، بلکه ابزاری برای رصد پویایی سیاست‌گذاری است.

۵-۷. تطبیق مفهوم محور: مسئولیت مدنی و مفهوم خسارت در حقوق ایران و اتحادیه اروپا

مقایسه کشور به کشور، هرچند برای ترسیم تصویر کلان ضروری است، در معرض این خطر است که به توصیف موازی نظام‌ها فروکاسته شود؛ تطبیق واقعی زمانی رخ می‌دهد که یک مفهوم حقوقی واحد در دو نظام، رودررو و در سطح مبانی تحلیل شود. این بخش شاخص مسئولیت مدنی را که هم در حقوق ایران پیشینه فقهی و قانونی غنی دارد و هم در اتحادیه اروپا در سال‌های ۲۰۲۴ و ۲۰۲۵ دستخوش تحول اساسی شده، در سه محور مبنای مسئولیت، مفهوم خسارت و بار اثبات و می‌کاود.

در محور نخست، مبنای مسئولیت، دو نظام از نقطه عزیمت متفاوت اما با منطقی همگرا حرکت می‌کنند. حقوق ایران بر دوگانه اتلاف و تسبیب (مواد ۳۲۸ و ۳۳۱ قانون مدنی) استوار است؛ در اتلاف، مباشر تلف بدون نیاز به اثبات تقصیر ضامن است و در تسبیب، مسئولیت غالباً منوط به احراز تقصیر است؛ ماده ۱ قانون مسئولیت مدنی ۱۳۳۹ نیز قاعده عام تقصیرمحور را مقرر می‌کند (ذاکری‌نیا، ۱۴۰۲). در اتحادیه اروپا، دستورالعمل جدید مسئولیت محصولات معیوب^۱ نرم‌افزار و سامانه‌های هوش مصنوعی را صریحاً محصول شمرده و مسئولیت محض تولیدکننده را برقرار کرده است (European Parliament & Council of the European Union, 2024)؛ هم‌زمان کمیسیون اروپا پیش‌نویس دستورالعمل خاص مسئولیت هوش مصنوعی را در فوریه ۲۰۲۵ از دستور کار خارج کرد. وجه همگرایی آن است که هر دو نظام برای فناوری خارج از فهم زیان‌دیده از تقصیرمحوری فاصله می‌گیرند؛ اما وجه افتراق بنیادین است: تسری اتلاف به زیان‌های نرم‌افزاری فاقد مباشرت متعارف در حقوق ایران محل تردید جدی است و رویه قضایی روشنی ندارد (تخشید، ۱۴۰۰؛ قلی‌نیا و مشهدی‌زاده، ۱۴۰۱)، حال آنکه قانون‌گذار اروپایی این تردید را با تصریح قانونی به محصول بودن نرم‌افزار منتفی کرده است (Farajpour, 2025b).

^۱ دستورالعمل (اتحادیه اروپا) ۲۸۵۳/۲۰۲۴ درباره‌ی مسئولیت ناشی از محصولات معیوب (Product Liability Directive (EU) 2024/2853)، لازم‌الاجرا از دسامبر ۲۰۲۴؛ دولت‌های عضو تا ۹ دسامبر ۲۰۲۶ برای انطباق حقوق داخلی مهلت دارند.

در محور دوم، مفهوم خسارت، فاصله دو نظام آشکارتر است: در حقوق ایران خسارت قابل جبران علی‌الاصول مادی، مسلم و مستقیم است؛ جبران خسارت معنوی با تردیدهای رویه‌ای همراه است و تبصره ۲ ماده ۵۱۵ قانون آیین دادرسی مدنی، عدم‌النفع را قابل مطالبه نمی‌داند. در برابر، دستورالعمل ۲۰۲۴ اتحادیه اروپا افزون بر مرگ و صدمه بدنی، آسیب روان‌شناختی از نظر پزشکی تأیید شده و مهم‌تر از همه، نابودی یا تخریب داده‌های غیرحرفه‌ای زیان‌دیده را صریحاً خسارت قابل جبران شناخته است. این تفاوت بازتاب دو درک متفاوت از موضوع حمایت است: حقوق ایران هنوز مال را در قالب سنتی می‌بیند، حال آنکه حقوق اروپا داده را دارایی مستقل شایسته حمایت شناخته است؛ خلئی که در دعاوی ناشی از تصمیمات خودکار، بخش بزرگی از زیان واقعی کاربران ایرانی را جبران‌ناپذیر می‌گذارد.

در محور سوم، بار اثبات، تقابل دو نظام سرنوشت‌ساز است: در حقوق ایران بر پایه قاعده البینه علی المدعی، اثبات عیب، ضرر و رابطه سببیت بر عهده زیان‌دیده است؛ الزامی که در برابر سامانه‌های الگوریتمی تاریک عملاً حق جبران را بی‌اثر می‌کند (Pasquale, 2015). قانون‌گذار اروپایی برای همین مسئله دو ابزار پیش‌بینی کرده است: تکلیف خواننده به افشای مستندات فنی به دستور دادگاه، و اماره قابل رد عیب و سببیت به سود زیان‌دیده در فرض خودداری از افشا یا پیچیدگی فنی نامتعارف دعا. حاصل این تطبیق سه‌محوری، درس تقنینی مشخصی است که مستقیماً به پیشنهاد ۸-۳ همین مقاله راه می‌برد: نظام مسئولیت مدنی هوش مصنوعی ایران بدون سه اصلاح به‌هم‌پیوسته، یعنی تصریح به شمول زیان‌های نرم‌افزاری و داده‌ای در مفهوم خسارت، پیش‌بینی اماره سببیت به سود زیان‌دیده و الزام توسعه‌دهنده به افشای مستندات فنی، از حل مسئله اثبات در دعاوی هوش مصنوعی ناتوان خواهد ماند.

۶. فلسفه‌های تقنینی در تنظیم‌گری هوش مصنوعی: حقوق محور، بازار محور، اقتدار محور

نحوه قانون‌گذاری در حوزه هوش مصنوعی تنها تابع سطح توسعه فناوری یا زیرساخت‌های اقتصادی نیست، بلکه ریشه در بنیان‌های فلسفی، اجتماعی، حقوقی و تاریخی هر جامعه دارد. در بررسی تطبیقی تجربه‌های قانون‌گذاری جهانی، سه رویکرد عمده قابل تمایزند: رویکرد حقوق محور^۱، رویکرد بازار محور^۲ و رویکرد اقتدار محور^۳.

۶-۱. رویکرد حقوق محور: تقدم کرامت انسانی و عدالت دیجیتال

در این رویکرد که بیش از همه در اتحادیه اروپا تجلی یافته، انسان و حقوق بنیادین او محور اصلی قانون‌گذاری است. هدف، نه صرفاً رشد فناوری، بلکه مهار آن در چارچوب اصولی چون کرامت انسانی، آزادی، حریم خصوصی، عدم تبعیض و شفافیت است (سیفی و رزمخواه، ۱۴۰۱). قانون هوش مصنوعی اتحادیه اروپا به‌عنوان نمونه بارز این رویکرد، استفاده از فناوری‌های پرخطر مانند امتیازدهی اجتماعی و شناسایی بیومتریک در فضاهای عمومی را ممنوع می‌کند.

۶-۲. رویکرد بازار محور: اولویت نوآوری و حداقل‌گرایی در مداخله قانونی

ایالات متحده آمریکا فلسفه‌ای مبتنی بر نوآوری آزاد، رقابت و اعتماد به بازار را دنبال می‌کند. در این نگاه، مقررات باید به حداقل برسند تا مانعی برای رشد سریع فناوری‌های نو ایجاد نشود؛ دولت آمریکا به جای وضع قوانین سخت‌گیرانه، ابزارهایی چون دستور اجرایی و استانداردهای داوطلبانه را به کار می‌گیرد. رویکرد برنامه اقدام هوش مصنوعی ۲۰۲۵ آمریکا، با محوریت رفع موانع مقرراتی، مصداق روشن این فلسفه است.

۶-۳. رویکرد اقتدار محور: فناوری در خدمت نظم سیاسی و امنیت عمومی

در جمهوری خلق چین، فلسفه تنظیم‌گری هوش مصنوعی شدیداً متأثر از مدل دولت‌محور و کنترل‌گر است. قوانین در این کشور باهدف حفظ امنیت ملی، انسجام اجتماعی و پایداری حکمرانی طراحی می‌شوند و شرکت‌ها ملزم به ثبت الگوریتم و رعایت انطباق محتوایی با ارزش‌های رسمی هستند.

¹ Human-Centric

² Market-Oriented

³ State-Centric

۴-۶. جایگاه ایران در میان سه رویکرد تقنینی: تردید میان مسیرها

ایران تاکنون فلسفه تقنینی مشخصی درباره هوش مصنوعی ارائه نداده است: برخی نهادهای سیاست‌گذار به مدل بازارمحور و رشد استارت‌آپ‌ها گرایش داشته‌اند؛ نهادهای امنیتی و حاکمیتی، به‌ویژه در حوزه زیرساخت و داده، رویکردی اقتدارگرایانه اتخاذ کرده‌اند؛ و برخی دانشگاهیان الگوی حقوق‌محور و کرامت‌مدار را پیشنهاد می‌دهند که هنوز در سیاست‌گذاری رسمی بازتاب چندانی نیافته است. پیش از تدوین هر متن قانونی، قانون‌گذار ایرانی باید به این پرسش بنیادی پاسخ دهد که چه نوع رابطه‌ای میان انسان، دولت و فناوری در ایران می‌خواهیم؛ بدون این بستر مفهومی، هر متنی که تحت عنوان قانون تصویب شود، فاقد عمق و انسجام کافی خواهد بود.

۷. تکه‌تکه بودن مقررات جهانی و چالش فناوری‌های بدون قانون در ایران

یکی از اساسی‌ترین موانع در راه حکمرانی مؤثر بر فناوری‌های نوظهور، پدیده‌ای است که در ادبیات حقوق تطبیقی از آن با عنوان تکه‌تکه‌شدن سیاستی^۱ یاد می‌شود (Ulnicane et al., 2021). در سطح جهانی، شکاف میان قدرت‌های تنظیم‌گر عمده اتحادیه اروپا، ایالات متحده و چین به الگویی آشکار از این تکه‌تکه بودن بدل شده است؛ اروپا با رویکردی سخت‌گیرانه و اخلاق‌محور، آمریکا با تکیه بر خودتنظیمی بخش خصوصی، و چین با نگاه اقتدارگرایانه، هر یک مسیر جداگانه‌ای را دنبال می‌کنند (اسکندری خوشگو و ابوالقاسمی، زیر چاپ).

مفهوم مکمل تکه‌تکه‌شدن سیاستی در سطح ملی، وضعیت فناوری‌های بدون قانون یا خلأهای تنظیم‌گری^۲ است؛ وضعیتی که به‌طور خاص در ایران به چشم می‌خورد، جایی که فناوری‌هایی از پهنادهای غیرنظامی و دوربین‌های هوشمند تشخیص چهره تا هوش مصنوعی مولد بدون قانون، آیین‌نامه یا نهاد ناظر مستقل توسعه یافته و به کار گرفته می‌شوند. از منظر حقوق عمومی، فقدان مقررات شفاف به گسترش نظارت‌های بدون کنترل و نقض اصل حریم خصوصی و اصل قانونی‌بودن جرم و مجازات انجامیده است؛ برون‌رفت از این وضعیت نیازمند تکمیل مسیر تقنینی طرح ملی هوش مصنوعی در مجلس شورای اسلامی است.

۸. جمع‌بندی تطبیقی و پیشنهادها برای آینده قانون‌گذاری هوش مصنوعی در ایران

پیشنهاد‌های این بخش، بر خلاف فهرست‌های توصیه‌ای عام، مستقیماً از خروجی چارچوب یازده شاخصه استخراج و بر پایه رتبه ایران در جدول ۳ اولویت‌بندی شده‌اند: ایران در شش شاخص طبقه‌بندی ریسک، مسئولیت کیفری، شفافیت الگوریتمی، ممیزی و گزارش‌دهی، ضمانت اجرا و هماهنگی بین‌المللی در وضعیت فاقد است و این شش شاخص به دلیل خلأ مطلق هنجاری، اولویت نخست تقنینی‌اند؛ پنج شاخص تعریف مفاهیم بنیادین، مسئولیت مدنی، نهاد ناظر مستقل، حریم خصوصی و داده و مشارکت عمومی که در وضعیت ضعیف قرار دارند، اولویت دوم را تشکیل می‌دهند، زیرا در آن‌ها مبنای هنجاری موجود باید تکمیل شود، نه ساخته از صفر. کاربست شاخص ترکیبی ARMI این اولویت‌بندی را کمی و پیامد آن را سنجش‌پذیر می‌کند: اگر ایران تنها شش شاخص فاقد را در قالب تکمیل طرح ملی هوش مصنوعی به سطح متوسط ارتقا دهد، امتیاز آن از ۵ به ۱۷ از ۳۳ جهش می‌کند و از قعر جدول به جایگاهی بالاتر از وضعیت کنونی آمریکا و کانادا (هر یک ۱۲) می‌رسد. این محاسبه نشان می‌دهد فاصله ایران با میانه جدول، نه محصول عقب‌ماندگی جبران‌ناپذیر فناورانه، بلکه خلأ تقنینی قابل رفع در یک چرخه قانون‌گذاری است. بر این اساس ده پیشنهاد زیر ارائه می‌شود و در عنوان هر یک، شاخص متناظر، وضعیت کنونی ایران و رتبه اولویت درج شده است.

۸-۱. تدوین نظام طبقه‌بندی ریسک متناسب با شرایط بومی (شاخص ۲؛ وضعیت ایران: فاقد؛ اولویت نخست)

در چشم‌انداز جهانی تنظیم‌گری هوش مصنوعی، ایران باید به‌جای تکیه صرف بر رشد فناورانه، رویکردی انسان‌محور و مبتنی بر طبقه‌بندی ریسک اتخاذ کند. طرح ملی هوش مصنوعی که کلیات آن در اردیبهشت ۱۴۰۴ در مجلس شورای اسلامی تصویب

¹ Policy Fragmentation

² Regulation Vacuums

شد، بستری مناسب برای ادغام چنین رویکردی فراهم می‌کند. پیشنهاد می‌شود سه سطح ریسک به‌صراحت در متن نهایی این طرح گنجانده شود:

سطح پایین: سامانه‌هایی با کاربرد محدود و غیرحساس، مانند چت‌بات‌های خدماتی ساده؛ الزامات این سطح باید محدود به شفافیت عملکرد و نظارت پسینی باشد تا نوآوری مختل نشود.

سطح بالا: سامانه‌هایی که بر حقوق اساسی یا تصمیم‌گیری‌های مهم تأثیرگذارند، مانند سامانه‌های استخدام یا تشخیص پزشکی؛ این سطح نیازمند ارزیابی پیشینی، آزمایش‌های ایمنی و گزارش‌دهی منظم است.

سطح غیرقابل قبول: سامانه‌هایی که مستقیماً کرامت انسانی، آزادی‌های مدنی یا امنیت ملی را تهدید می‌کنند، مانند تشخیص چهره برای نظارت گسترده بدون رضایت یا امتیازدهی اجتماعی؛ این سامانه‌ها باید به طور کامل ممنوع یا تنها با نظارت شدید قضایی مجاز شوند.

۸-۲. تعریف شفاف مفاهیم بنیادین (شاخص ۱؛ وضعیت ایران: ضعیف؛ اولویت دوم)

فقدان تعریف دقیق مفاهیمی چون هوش مصنوعی، عامل خودکار، تصمیم‌گیری الگوریتمی، یادگیری ماشین و تبعیض الگوریتمی موجب ابهام، تفاسیر متناقض و ناکارآمدی قانون‌گذاری می‌شود. ضروری است فرایند تدوین این تعاریف در ایران با همکاری میان‌رشته‌ای حقوق‌دانان فناوری، متخصصان علوم رایانه، فیلسوفان اخلاق و جامعه‌شناسان دنبال شود تا ضمن تضمین دقت علمی و حقوقی، ملاحظات فرهنگی و ارزشی نیز لحاظ گردد.

۸-۳. تعیین مسئولیت مدنی و کیفری توسعه‌دهندگان (شاخص‌های ۳ و ۴؛ وضعیت ایران: ضعیف و فاقد؛ اولویت نخست)

یکی از اساسی‌ترین چالش‌های تنظیم‌گری هوش مصنوعی، تعیین حدود مسئولیت مدنی و کیفری توسعه‌دهندگان است (خلیلی پاجی و احمدی، ۱۴۰۴؛ ذاکری‌نیا، ۱۴۰۲). فرج‌پور (Farajpour, 2025) با مطالعه تطبیقی نظام‌های آمریکا، اتحادیه اروپا، آلمان و ایران، مدل‌های گوناگون از مسئولیت توسعه‌دهنده و بهره‌بردار تا مسئولیت محض، بیمه اجباری و شخصیت حقوقی محدود را تحلیل و نشان می‌دهد در حقوق ایران، مسئولیت مدنی همچنان بر شخصیت حقوقی انسانی و شرکتی استوار است و هوش مصنوعی فاقد شخصیت حقوقی مستقل است. فرج‌پور و همکاران (۱۴۰۴) نیز پیشنهاد می‌دهند حقوق ایران می‌تواند از الگوی مسئولیت محض اتحادیه اروپا و مدل مسئولیت شرکتی آلمان الهام بگیرد، بی‌آنکه لزوماً شخصیت حقوقی کامل برای سامانه‌های هوشمند به رسمیت شناخته شود.

ایران می‌تواند بر پایه این تجارب، نظامی ترکیبی و بومی ارائه دهد: مسئولیت مستقیم توسعه‌دهنده در خطاهای ناشی از طراحی معیوب یا غفلت؛ تقسیم نسبی مسئولیت میان توسعه‌دهنده و بهره‌بردار در نقص‌های غیرقابل‌پیش‌بینی ناشی از خودآیینی الگوریتم؛ و پیش‌بینی بیمه اجباری برای سامانه‌های پرریسک، مشابه بیمه اجباری وسایل نقلیه موتوری. سه سازوکاری که تطبیق مفهوم محور بخش ۵-۷ به دست داد، باید ستون فقرات این نظام باشد: تصریح قانونی به شمول زیان‌های نرم‌افزاری و خسارت داده‌ای در مفهوم خسارت قابل جبران؛ پیش‌بینی اماره قانونی قابل رد سببیت به سود زیان‌دیده در دعاوی سامانه‌های پرریسک؛ و الزام توسعه‌دهنده به افشای مستندات فنی به دستور دادگاه، به‌عنوان پیش‌شرط توزیع عادلانه بار اثبات.

۸-۴. تشکیل و تکمیل اختیارات نهاد ناظر مستقل (شاخص ۵؛ وضعیت ایران: ضعیف؛ اولویت دوم)

ایران در مسیر تأسیس سازمان ملی هوش مصنوعی زیر نظر رئیس‌جمهور است که ماده تأسیس آن در آبان ۱۴۰۴ در مجلس شورای اسلامی تصویب شد. برای ایفای نقش مؤثر، ضروری است اختیارات این نهاد شامل صدور مجوز فعالیت سامانه‌های پرریسک، ارزیابی مستمر ریسک‌ها، اعمال ضمانت اجرا در موارد تخلف و سازوکار رسیدگی به شکایات کاربران، با استقلال مالی و سازمانی کافی در متن نهایی قانون تثبیت شود تا از تجربه چندپارگی نهادی مشاهده‌شده در جریان تصویب این طرح (میان شورای عالی انقلاب فرهنگی و مجلس) در آینده پرهیز گردد.

۸-۵. شفافیت، گزارش‌دهی و قابلیت ممیزی الگوریتم‌ها (شاخص‌های ۶ و ۷؛ وضعیت ایران: فاقد؛ اولویت نخست)

توسعه‌دهندگان باید موظف باشند مراحل طراحی الگوریتم، داده‌های آموزشی و معیارهای تصمیم‌گیری را به‌صورت مستند ارائه دهند تا امکان نظارت و پاسخ‌گویی فراهم شود (حسینی، عبدخدائی و شریف‌خانی، ۱۴۰۲). الگوریتم‌های غیرشفاف می‌توانند آثار گسترده‌ای بر عدالت اجتماعی و حقوق بنیادین افراد بر جای بگذارند (Binns et al., 2018).

۸-۶. صیانت از حقوق بشر دیجیتال و حریم خصوصی (شاخص ۸؛ وضعیت ایران: ضعیف؛ اولویت دوم)

ایران نیازمند آن است که صیانت از حقوق بشر دیجیتال و حریم خصوصی را به‌عنوان یکی از ارکان اصلی تنظیم‌گری خود قرار دهد؛ از جمله از طریق تضمین اصل رضایت آگاهانه در جمع‌آوری و پردازش داده‌ها، ممنوعیت استفاده تبعیض‌آمیز از الگوریتم‌ها، و ایجاد حق دسترسی و کنترل فرد بر داده‌های شخصی خویش.

۸-۷. مشارکت عمومی و تنظیم‌گری مشارکتی (شاخص ۱۰؛ وضعیت ایران: ضعیف؛ اولویت دوم)

تجربه اتحادیه اروپا در تدوین AI Act نشان می‌دهد که مشورت گسترده با متخصصان فنی و حقوقی، انجمن‌های حقوق بشر و سازمان‌های غیردولتی، به قانونی منجر می‌شود که فراتر از ایمنی فنی، به تبعیض الگوریتمی، حریم خصوصی و عدالت اجتماعی نیز توجه می‌کند. در ایران نیز می‌توان شورای مشورتی چنددلی نفعی متشکل از نمایندگان جامعه مدنی، اساتید دانشگاه، متخصصان علوم داده و نمایندگان گروه‌های آسیب‌پذیر تشکیل داد.

۸-۸. آموزش، سواد دیجیتال و اخلاق عمومی (پیشنهاد مکمل فرا شاخصی؛ بستر ساز تحقق شاخص‌های ۸ و ۱۰)

افزایش سواد هوش مصنوعی در میان شهروندان، دانش‌آموزان، قضات و دیگر گروه‌های تصمیم‌ساز، از ارکان اصلی حکمرانی مسئولانه است؛ تجربه فنلاند با پروژه (Elements of AI) نشان می‌دهد آموزش‌های ساده و قابل فهم می‌تواند درک جامعه از فناوری را به‌طور چشمگیر افزایش دهد. در ایران می‌توان طرحی ملی با عنوان سواد هوش مصنوعی برای مدارس طراحی و هم‌زمان کارگاه‌های تخصصی برای قضات و وکلا برگزار کرد تا نظام قضایی توان مواجهه با پرونده‌های تبعیض الگوریتمی یا خطای سامانه‌های هوشمند را بیابد (حاجی ده‌آبادی، خالقی و ثقفی، ۱۴۰۳).

۸-۹. تکمیل زنجیرهٔ تقنینی طرح ملی هوش مصنوعی (پیش‌نیاز عرضی تحقق هر یازده شاخص)

برخلاف رویکردهای پیشنهادی که تشکیل نهاد تازه یا تدوین سند جدید را توصیه می‌کردند، وضعیت کنونی ایران نشان می‌دهد که ظرفیت نهادی لازم (طرح ملی هوش مصنوعی، شورای ملی راهبری، سازمان ملی هوش مصنوعی) در حال شکل‌گیری است؛ اولویت اصلی، تکمیل شور دوم این طرح در مجلس با گنجاندن صریح شاخص‌های یازده‌گانهٔ این پژوهش به‌ویژه طبقه‌بندی ریسک، مسئولیت مدنی و کیفری، و الزامات شفافیت است، نه آغاز فرایندی موازی و جدید که خطر تکرار چندپارگی نهادی مشاهده‌شده در تجربهٔ اخیر ایران را در پی داشته باشد.

۸-۱۰. پیش‌بینی سازوکار تنظیم‌گری چابک و پایش ادواری (شاخص ۱۱؛ وضعیت ایران: فاقد؛ اولویت نخست)

تجربه اتحادیه اروپا در سال ۲۰۲۵ که پیشرفته‌ترین نظام حقوقی هوش مصنوعی جهان را ناگزیر از تعویق و بازنگری الزامات پرریسک کرد، و تجربه کلرادو که ظرف کمتر از دو سال از قانونی ریسک‌محور به چارچوبی سبک مبتنی بر افشا و شفافیت بازنگری شد، نشان می‌دهد هیچ متن قانونی درباره هوش مصنوعی نمی‌تواند ایستا باشد. از این رو پیشنهاد می‌شود قانون ملی هوش مصنوعی ایران از همان ابتدا دو سازوکار نهادی را بگنجانند: نخست، قرارگاه آزمایش تنظیمی^۱ تحت نظارت سازمان ملی هوش مصنوعی که به شرکت‌های دانش‌بنیان اجازه می‌دهد سامانه‌های خود را پیش از اعمال کامل الزامات، در محیطی کنترل‌شده آزمایش کنند؛ دوم، الزام قانونی به بازنگری ادواری متن قانون (برای نمونه هر سه سال یک‌بار) توسط کمیته‌ای میان‌رشته‌ای از نمایندگان سازمان ملی هوش مصنوعی، مجلس و دانشگاه‌ها. نبود چنین سازوکاری در بسیاری از قوانین سخت‌گیرانه اولیه، از جمله نسخه نخست قانون کلرادو، از دلایل اصلی بازنگری‌های ناگهانی و پرهزینه بوده است.

^۱ Regulatory Sandbox

۹. آینده پژوهی قانون گذاری هوش مصنوعی و جایگاه ایران

روند پرشتاب تحولات هوش مصنوعی، نظام های حقوقی جهان را با دو سناریوی متقابل روبه رو کرده است: نخست، حرکت به سوی جهانی سازی مقررات و پذیرش اصول بنیادین اخلاقی مشترک؛ رویکردی مورد حمایت یونسکو و OECD با تأکید بر شفافیت، عدم تبعیض و امنیت داده ها (OECD, 2019). در نقطه مقابل، سناریوی گسترش شکاف های تقنینی جهانی قرار دارد که در آن نبود اجماع بر اولویت های اخلاقی، اقتصادی و امنیتی به قانون گذاری های ناهمگون می انجامد؛ واگرایی ای که از دیرباز در ادبیات تطبیقی مشاهده شده (Cath et al., 2018) و تحولات آمریکا، اروپا و کانادا در سال های ۲۰۲۵ و ۲۰۲۶ آن را تشدید کرده است. در این میان، ایران هنوز به هیچ سناریویی به طور شفاف نپیوسته و همین امر خطر قانون گریزی یا وابستگی به فناوری های خارجی را افزایش می دهد؛ راهکار برتر، طراحی سناریویی بومی شده، چندسطحی و متوازن است که از اصول بین المللی بهره گیرد و به اقتضائات حقوقی، فرهنگی و سیاسی کشور نیز توجه کند.

۱۰. محدودیت های پژوهش و پیشنهادهای پژوهش آتی

علاوه بر محدودیت های روش شناختی مطرح شده در بخش چهارم، این پژوهش از چند جهت دیگر نیز محدودیت دارد که شناخت آن ها برای ارزیابی صحیح یافته ها و برای طراحی پژوهش های آتی ضروری است.

نخست، شاخص ترکیبی بلوغ حکمرانی هوش مصنوعی^۱ با وزن یکسان برای هر یازده شاخص محاسبه شده است؛ انتخابی که هرچند از سوگیری پیشینی جلوگیری می کند، به معنای اهمیت سیاستی یکسان همه ابعاد تنظیم گری نیست. پژوهش های آتی می توانند با روش های وزن دهی مبتنی بر نظر سنجی از خبرگان (مانند فرایند تحلیل سلسله مراتبی یا AHP) وزن های متفاوتی به هر شاخص اختصاص دهند و نتایج را با یافته های این پژوهش مقایسه کنند.

دوم، این پژوهش تنها پنج نظام حقوقی را پوشش می دهد. گسترش چارچوب یازده شاخص به کشورهای دیگر منطقه خاورمیانه و آسیای غربی به ویژه امارات متحده عربی که با تأسیس وزارت هوش مصنوعی و راهبردهای بلندپروازانه، در ادبیات سیاستی منطقه ای اغلب به عنوان الگوی رقیب ایران مطرح می شود می تواند تصویر تطبیقی کامل تری از جایگاه منطقه ای ایران به دست دهد.

سوم، به دلیل ماهیت به شدت پویای تنظیم گری هوش مصنوعی، بسیاری از یافته ها به ویژه درباره لایحه های ایالتی آمریکا و روند بررسی طرح ملی هوش مصنوعی در مجلس، ماهیتی موقت دارند. پژوهش های آتی می توانند این چارچوب را در فواصل زمانی منظم (برای نمونه سالانه) بر همین پنج کشور یا نمونه ای گسترده تر اعمال کنند تا سیر تحول ARMI هر کشور در طول زمان رصد شود؛ همان کارکردی که یک شاخص ترکیبی عددی، بر خلاف توصیف کیفی صرف، ممکن می سازد.

چهارم، این پژوهش بر تحلیل متن رسمی اسناد متمرکز است و ارزیابی اثربخشی اجرایی مقررات (میزان انطباق عملی نهادها و شرکت ها با الزامات مصوب) را در برنمی گیرد؛ مطالعات میدانی و تجربی درباره میزان اجرای واقعی قانون هوش مصنوعی اتحادیه اروپا یا مقررات چین، مکمل ضروری برای یافته های این پژوهش خواهد بود.

^۱ ARMI (AI Regulatory Maturity Index)

نتیجه‌گیری

این مقاله با طراحی و اعمال چارچوبی یازده شاخصه، مقررات هوش مصنوعی را در اتحادیه اروپا، ایالات متحده، چین، کانادا و ایران به صورت تطبیقی تحلیل کرد. یافته‌ها نشان می‌دهد اتحادیه اروپا همچنان منسجم‌ترین ساختار را دارد، هرچند با تعدیل زمانی الزامات پریسک و کنارگذاشتن دستورالعمل مسئولیت مدنی هوش مصنوعی مواجه شده است؛ ایالات متحده پس از ژانویه ۲۰۲۵ مسیر مقررات‌زدایی را در پیش گرفته، کانادا بر خلاف تصور رایج بدون قانون فدرال مانده و چین با تأکید بر نظم سیاسی و امنیت جمعی بر توسعه قواعد سخت‌گیرانه تمرکز کرده است.

نوآوری اصلی پژوهش دووجهی است: در سطح نظری، پیوند صریح هر شاخص با یک جریان نظری مشخص، چارچوبی دانش‌بنیان و قابل دفاع پدید آورده است؛ و در سطح روش‌شناختی، تبدیل رتبه‌بندی کیفی به شاخص ترکیبی بلوغ حکمرانی هوش مصنوعی امکان مقایسه عددی و پایش طولی نظام‌های حقوقی را فراهم آورده است؛ سنجه‌ای که نشان می‌دهد شکاف اتحادیه اروپا (۳۰ از ۳۳) و ایران (۵ از ۳۳) بزرگ‌ترین فاصله میان پنج نظام مورد مطالعه است و آمریکا و کانادا در تیرماه ۱۴۰۵ به یک سطح بلوغ رسیده‌اند. وجه تمایز این نتیجه‌گیری از جمع‌بندی‌های متعارف آن است که توصیه‌های مقاله خروجی مستقیم همین چارچوب است و از سه مسیر استخراج شده است: نخست، رتبه‌های جدول ۳ مبنای اولویت‌بندی دوسطحی شد؛ شش شاخصی که ایران در آن‌ها فاقد است اولویت نخست و پنج شاخص ضعیف اولویت دوم را ساختند. دوم، شاخص ARMI پیامد این اولویت‌بندی را کمی کرد: ارتقای شش شاخص فاقد به سطح متوسط، امتیاز ایران را از ۵ به ۱۷ از ۳۳ می‌رساند و از قعر جدول به بالاتر از وضعیت کنونی آمریکا و کانادا منتقل می‌کند؛ گزاره‌ای سنجش‌پذیر که در بازنگری‌های سالانه آتی قابل راستی‌آزمایی است. سوم، تطبیق مفهوم‌محور مسئولیت مدنی در بخش ۵-۷ سه سازوکار مشخص، یعنی شناسایی خسارت داده‌ای، اماره سببیت به سود زیان‌دیده و تکلیف افشای مستندات فنی را برای قانون‌گذار ایرانی به دست داد. بدین‌سان چارچوب یازده شاخصه صرفاً ابزار توصیف وضع موجود نیست؛ موتور تولید توصیه‌های متمایز، اولویت‌بندی‌شده و قابل پایش این پژوهش است. ایران هرچند در آغاز راه قانون‌گذاری هوش مصنوعی است، با وجود طرح ملی در حال بررسی در مجلس شورای اسلامی، دیگر از نقطه صفر آغاز نمی‌کند؛ اولویت سیاست‌گذار، نه ابداع نهادهای موازی جدید، بلکه تکمیل و انسجام‌بخشی به همین مسیر تقنینی بر پایه شاخص‌های یازده‌گانه پیشنهادی است: تعریف صریح مفاهیم، طبقه‌بندی ریسک، تعیین مسئولیت مدنی و کیفری، تثبیت اختیارات نهاد ناظر، الزام به شفافیت و ممیزی، صیانت از حریم خصوصی، تضمین ضمانت اجرا و نهادینه‌سازی مشارکت عمومی. مسیر پیش‌رو نه صرفاً ضرورتی حقوقی، بلکه آزمونی راهبردی برای توان حکمرانی ایران در عصر فناوری‌های خودآیین است؛ آزمونی که موفقیت در آن در گرو گذار از سیاست‌گذاری واکنشی و پراکنده به سیاست‌گذاری نظام‌مند، پیش‌نگر و قابل پایش خواهد بود.

منابع

الف) منابع فارسی

- نوروزی، میثم؛ اسکندری خوشگو، مهدی و ابوالقاسمی، ساناز. (۱۴۰۳). جایگاه حقوق بین الملل فناوری در توسعه و نظم دهی هوش مصنوعی: تحلیلی بر نخستین قطعنامه مجمع عمومی سازمان ملل متحد در خصوص هوش مصنوعی. (e208444). پژوهش های حقوقی. <https://doi.org/10.48300/jlr.2024.472127.2711>
- تخشید، زهرا (۱۴۰۰). مقدمه ای بر چالش های هوش مصنوعی در حوزه مسئولیت مدنی. حقوق خصوصی، ۱۸(۱)، ۲۲۷-۲۵۰. <https://doi.org/10.22059/jolt.2021.319529.1006965>
- حاجی ده آبادی، محمدعلی؛ خالقی، ابوالفتح؛ ثقفی، پریسا (۱۴۰۳). ارزش اثباتی ادله هوشمند با نگاهی به ادله موضوعی و طریقی. حقوق فناوری های نوین، ۵(۱۰)، ۱۸۳-۱۹۷. <https://doi.org/10.22133/mtlj.2024.421960.1259>
- حسینی، سید احمد؛ عبدخدائی، زهره؛ شریف خانی، محمد (۱۴۰۲). کاربرد هوش مصنوعی در رسیدگی های قضائی، چالش شفافیت و راهکارهای آن. دیدگاه های حقوق قضائی، ۲۸(۱۰۱)، ۶۷-۹۰. <http://jlvviews2.ujsas.ac.ir/article-1-2139-fa.html>
- حکمت نیا، محمود؛ محمدی، مرتضی؛ واتقی، محسن (۱۳۹۸). مسئولیت مدنی ناشی از تولید ربات های مبتنی بر هوش مصنوعی خودمختار. حقوق اسلامی، ۱۵(۶۰)، ۲۳۱-۲۵۸. http://hoquq.iict.ac.ir/article_36476.html
- خلیلی پاچی، عارف؛ احمدی، امیرحسین (۱۴۰۴). کاربرد هوش مصنوعی در حقوق کیفری بین المللی؛ از مسئولیت گریزی تا عدالت الگوریتمی. پژوهش های حقوقی (مقالات آماده انتشار). <https://doi.org/10.48300/jlr.2025.527092.2955>
- ذاکری نیا، حانیه (۱۴۰۲). ماهیت و مبنای مسئولیت مدنی ناشی از هوش مصنوعی در حقوق ایران و کشورهای اتحادیه اروپا. حقوق خصوصی، ۲۰(۱)، ۱۳۵-۱۵۲. <https://doi.org/10.22059/jolt.2023.356703.1007186>
- رجبی، عبدالله (۱۳۹۸). ضمان در هوش مصنوعی. مطالعات حقوق تطبیقی، ۱۰(۲)، ۴۴۹-۴۶۶. <https://doi.org/10.22059/jcl.2019.274782.633787>
- سیفی، آناهیتا؛ رزمخواه، نجمه (۱۴۰۱). هوش مصنوعی و چالش های پیش رو در قلمرو حقوق بین الملل بشر، با رویکردی بر حق کار. دوفصلنامه بین المللی حقوق بشر، ۱۷(۱)، ۵۵-۷۴. <https://doi.org/10.22096/hr.2021.524039.1289>
- شریفی، سید الهام الدین؛ بیرمی، گلناز (۱۳۹۷). ماهیت حقوقی نمایندگان هوشمند در عرصه قراردادهای الکترونیکی. پژوهش های حقوقی، ۱۷(۳۳)، ۲۷-۵۴. <https://doi.org/10.48300/jlr.2018.65509>
- صادقی، حسین؛ رشوند بوکانی، مهدی؛ ناصر، مهدی (۱۴۰۲). چالش های اجرای نظام جدید مسئولیت مدنی ناشی از نقض قواعد حقوق رقابت در اتحادیه اروپا. پژوهش های حقوقی، ۲۲(۵۳)، ۱۷۱-۱۹۴. https://jlr.sdil.ac.ir/article_170537.html
- قلی نیا، رضا؛ مشهدی زاده، علیرضا (۱۴۰۱). مسئولیت مدنی کاربر در به کارگیری سیستم هوش مصنوعی در خودرو. پژوهش های حقوقی، ۲۱(۵۰)، ۳۰۵-۳۳۱. <https://doi.org/10.48300/jlr.2021.285755.1653>
- گندم کار، رضاحسین؛ صالحی مازندرانی، محمد؛ حمیدی، محمدمهدی (۱۴۰۰). بررسی تطبیقی امکان وجود شخصیت حقوقی برای سامانه های هوشمند در فقه امامیه، حقوق ایران و حقوق غرب. پژوهش تطبیقی حقوق اسلام و غرب، ۸(۴)، ۲۳۵-۲۶۶. <https://doi.org/10.22091/CSIW.2021.5944.1903>

ب) منابع انگلیسی

- Binns, R., Veale, M., Van Kleek, M., & Shadbolt, N. (2018). It's reducing a human being to a percentage: Perceptions of justice in algorithmic decisions. In Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems (pp. 1-14). ACM. <https://doi.org/10.1145/3173574.3173951>
- Buiten, M. C. (2020). Regulating emerging technologies in Europe: Artificial intelligence. Common Market Law Review, 57(5), 1143-1180. <https://doi.org/10.54648/COLA2020426>
- Cath, C., Wachter, S., Mittelstadt, B., Taddeo, M., & Floridi, L. (2018). Artificial intelligence and the "good society": The US, EU, and UK approach. Science and Engineering Ethics, 24(2), 505-528. <https://doi.org/10.1007/s11948-017-9901-7>

- Coeckelbergh, M. (2020). AI ethics. MIT Press. URL: <https://mitpress.mit.edu/9780262538190/ai-ethics/>
- Council of Europe. (2024). Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law (CETS No. 225). URL: <https://coe.int/en/web/artificial-intelligence/the-framework-convention-on-artificial-intelligence>
- European Commission. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). Official Journal of the European Union, L 2024/1689. URL: <http://data.europa.eu/eli/reg/2024/1689/oj>
- European Parliament & Council of the European Union. (2024). Directive (EU) 2024/2853 of the European Parliament and of the Council of 23 October 2024 on liability for defective products. Official Journal of the European Union, L 2024/2853. URL: <http://data.europa.eu/eli/dir/2024/2853/oj>
- Farajpour, R. (2025a). Legal and comparative analysis of civil liability of artificial intelligence in automated decision-making. AI and Tech in behavioural and Social Sciences, 3(3), 1–9. <https://doi.org/10.61838/kman.aitech.3.3.4>
- Farajpour, R. (2025b). The role of civil liability in artificial intelligence laws from the perspective of major global legal systems. Journal of Law and Political Studies, 5(1), 182–196. <https://doi.org/10.48309/jlps.2025.518711.1353>
- Farajpour, R., Amerinia, M.B., & Pourjavaheri, A. (2025). The role of artificial intelligence in arbitration and legal challenges arising from automated decisions in sports. AI and Tech in Behavioural and Social Sciences, 3(2), 21–28. <https://doi.org/10.61838/kman.aitech.3.2.3>
- Floridi, L. (2021). AI and its new winter: From myths to realities. Philosophy & Technology, 34(1), 1–3. <https://doi.org/10.1007/s13347-020-00396-6>
- Gandomar, R., Salehi Mazandarani, M., and Hamidi, M. (2021). A Comparative Study of the Possibility of the Existence of Legal Personality for Intelligent Systems in Islamic Jurisprudence, Iranian law and Law of the West. Comparative Studies on Islamic and Western Law, 8(4), 235–266. [In Persian] <https://doi.org/10.22091/CSIW.2021.5944.1903>
- Gunkel, D. J. (2018). Robot rights. MIT Press. URL: <https://mitpress.mit.edu/9780262038621/robot-rights/>
- Gunkel, D. J. (2022). Deconstructing the moral machine: A critical analysis of AI ethics. MIT Press. URL: <https://mitpress.mit.edu/9780262544245/>
- Hacker, P., Engel, A., & Mauer, M. (2023). Regulating ChatGPT and other large generative AI models. In Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency (FAccT '23) (pp. 1112–1123). ACM. <https://doi.org/10.1145/3593013.3594067>
- Haji Deh Abadi, M. A., Khaleghi, A and Saghafi, P. (2024). The Probative Value of Smart Proofs Looking at the Thematic and Methodological Evidence. Modern Technologies Law, 5(10), 183–197. [In Persian] <https://doi.org/10.22133/mtlj.2024.421960.1259>
- Hekmatnia, M., Mohammadi, M. and Vaseghi, M. (2019). civil Liability for damages caused by robots based on autonomous artificial intelligence. Islamic Law, 16(60), 231–258. [In Persian] URL: http://hoquq.iict.ac.ir/article_36476.html
- Hosseini A, Abdekhodae, Z. and Sharifkhani, M. (2023). Application of Artificial Intelligence in Judicial Proceedings, the Challenge of Transparency and its Solutions. Judicial Law Views Quarterly (Law Views), 28 (101), 67–90. [In Persian] URL: <http://jlviews2.ujss.ac.ir/article-1-2139-en.html>
- Khalili Paji, A. and Ahmadi, A. (2025). The Application of Artificial Intelligence in International Criminal Law: From Responsibility Evasion to Algorithmic Justice. (e232863). Journal of Legal Research. [In Persian] <https://doi.org/10.48300/jlr.2025.527092.2955>
- National Institute of Standards and Technology. (2023). Artificial intelligence risk management framework (AI RMF 1.0). U.S. Department of Commerce. <https://doi.org/10.6028/NIST.AI.100-1>
- Norouzi, M., Eskandari Khoshgu, M. and Abolghasemi, S. (2024). The position of international technology law in the development and regulation of artificial intelligence: an analysis of the first resolution of the United Nations General Assembly regarding artificial intelligence. (e208444). Journal of Legal Research. [In Persian] <https://doi.org/10.48300/jlr.2024.472127.2711>
- Organisation for Economic Co-operation and Development. (2019). Recommendation of the Council on Artificial Intelligence (OECD/LEGAL/0449). <https://doi.org/10.1787/eedfee77-en>
- Pasquale, F. (2015). The black box society: The secret algorithms that control money and information. Harvard University Press. URL: <https://hup.harvard.edu/books/9780674970847>

- Qoliniya, R. and Mashhadizadeh, A. (2022). Civil Liability of The User in Using The Artificial Intelligence System in The Car. *Journal of Legal Research*, 21(50), 305–331. [In Persian] <https://doi.org/10.48300/jlr.2021.285755.1653>
- Rajabi, A. (2019). Liability of Artificial Intelligence; the Reflection of Developments in the Liability Rules. *Comparative Law Review*, 10(2), 449–466. [In Persian] <https://doi.org/10.22059/jcl.2019.274782.633787>
- Sadeghi, H., Rashvand Bukani, M. and Naser, M. (2023). The Challenges of Implementing the New Civil Liability
- Sharifi, S.E. and Beyrami, G. (2018). Intelligent Agent's Legal Nature in the Field of Electronic Contracts. *Journal of Legal Research*, 17(33), 27–54. [In Persian] <https://doi.org/10.48300/jlr.2018.65509>
- Seifi, A. and Razmkhah, N. (2022). Artificial Intelligence and Challenges Ahead on the Base of International Human Rights; with the Emphasis on Right to Work. *The Journal of Human Rights*, 17(Number 1), 55–74. [In Persian] <https://doi.org/10.22096/hr.2021.524039.1289>
- System due to Violation of Competition Rules in the European Union. *Journal of Legal Research*, 22(53), 171–194. [In Persian] URL: https://jlr.sdil.ac.ir/article_170537.html
- Smuha, N. A. (2021). Beyond a human rights-based approach to AI governance: Promise, pitfalls, plea. *Philosophy & Technology*, 34(1), 91–104. <https://doi.org/10.1007/s13347-020-00403-w>
- Takhshid, Z. (2021). An Introductory Study on the Challenges of Artificial Intelligence in Tort Law. *Private Law*, 18(1), 227–250. [In Persian] <https://doi.org/10.22059/jolt.2021.319529.1006965>
- The White House, Office of Science and Technology Policy. (2025a). winning the race: America's AI action plan. URL: <https://whitehouse.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf>
- The White House. (2025b). Removing barriers to American leadership in artificial intelligence (Executive Order No. 14179). *Federal Register*, 90(20), 8741–8743. <https://federalregister.gov/d/2025-02172>
- Ulnicane, I., Knight, W., Leach, T., Stahl, B. C., & Wanjiku, W. (2021). Framing governance for a contested emerging technology: Insights from AI policy. *Policy and Society*, 40(2), 158–177. <https://doi.org/10.1080/14494035.2020.1855800>
- Veale, M., & Zuiderveen Borgesius, F. (2021). Demystifying the draft EU Artificial Intelligence Act. *Computer Law Review International*, 22(4), 97–112. <https://doi.org/10.9785/cr-2021-220402>
- Wachter, S., Mittelstadt, B., & Floridi, L. (2017). Transparent, explainable, and accountable AI for robotics. *Science Robotics*, 2(6), eaan6080. <https://doi.org/10.1126/scirobotics.aan6080>
- 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. [In Persian] <https://doi.org/10.22059/jolt.2023.356703.1007186>