

<http://doi.org/10.22133/mtlj.2025.475570.1369>

The European Union Artificial Intelligence Act: A Step Towards Human Oversight of Artificial Intelligence

Abbasali Kadkhodaei¹, Abolfazl Khakzad Bajestani^{2*}, Amir Mohammad Majidi³

¹ Prof. Faculty of Law and Political Science, Tehran University, Tehran, Iran.

² MA. Student in International law, Faculty of Law and Political Science, Tehran University, Tehran, Iran.

³ MA. Student in International law, Faculty of Law and Political Science, Tehran University, Tehran, Iran.

Article Info	Abstract
Original Article Received: 28-08-2024 Accepted: 23-02-2025 Keywords: European Union Brussels effect Risk-Based Approach Artificial Intelligence *Corresponding author e-mail: abolfazl.khakzad@ut.ac.ir	In the contemporary era, artificial intelligence (AI) undeniably plays a pivotal role in human life, much as digital technology has evolved into a cornerstone of societal existence. Consequently, fostering societal trust in these technologies is deemed imperative. To cultivate such trust, the imposition of human oversight on these technologies through legislative measures stands as an effective instrument. The European Union Artificial Intelligence Act (EU AI Act), the world's first comprehensive legal framework designed to safeguard human values, exemplifies this approach. Initiated in 2017 by the European Commission and formally adopted by the Council of the European Union on 21 May 2024, this legislation positions the EU as a potential global leader in AI governance. This article addresses two primary questions: First, to what extent can the EU AI Act serve as a global regulatory model via the Brussels Effect, and what implications arise from this mechanism? Second, what innovative contributions does the Act offer in advancing legal oversight of AI? Employing a descriptive-analytical methodology and relying on a literature review, the study concludes that the EU has struck a balance between fostering AI innovation and protecting fundamental citizen rights. By leveraging the Brussels Effect, the EU aims to establish global leadership in AI regulation. Furthermore, the Act introduces pioneering measures, such as a precise definition of AI, the establishment of real jurisdiction in the scope of the Act, and a risk-based regulatory approach. These innovations not only enhance legal clarity but also set a precedent for transnational AI governance.
How to Cite: Kakhodaei, A., Khakzad Bajestani, A., & Majidi, A.M. (2026). The European Union Artificial Intelligence Act: A Step Towards Human Oversight of Artificial Intelligence. <i>Modern Technologies Law</i>, 7(13), 171-190.	Published by University of Science and Culture https://www.usc.ac.ir Online ISSN: 2783-3836

1. Introduction

In the contemporary world, artificial intelligence (AI) has become an inseparable component of human life, permeating all facets of modern society, much like digital technologies. While offering unprecedented opportunities for innovation and progress, AI also presents challenges that demand rigorous oversight and legal regulation. Public trust in AI can only emerge when humans retain the ability to monitor its operations and ensure alignment with ethical principles and fundamental rights. In this context, legislation serves as a critical instrument for establishing human oversight over AI. The European Union's Artificial Intelligence Act (EU AI Act), the world's first comprehensive legal framework, aims to safeguard human values and citizens' rights against potential risks posed by this technology. Initiated in 2017 and formally adopted on May 21, 2024, this legislation positions the EU as a global leader in AI governance. The Act seeks to balance harnessing AI's capabilities with managing its risks, fostering innovation while protecting citizens' health, safety, and fundamental rights. This article examines the EU AI Act, its potential to become a global model through the "Brussels Effect," and its legal innovations in AI governance. Although numerous articles and books exist on artificial intelligence, no specific study has been published regarding the EU AI Act and how it can leverage Brussels Effect to set a global standard.

Although research on artificial intelligence in Iran has been conducted regarding its ethical and social aspects, along with forays into its technical facets, there is no full-fledged legal examination of the European Union AI Act and what it means for Iran and at the international level. In the studies that have been conducted by global researchers, most of the research has been centered on the societal and ethical implications of artificial intelligence, developing legal and regulatory frameworks, and the challenges of AI regulation.

2. Methodology

This study employs a descriptive-analytical methodology, using library-based data collection. To analyze the EU AI Act, official EU documents—including European Commission reports, European Parliament and Council resolutions—and scholarly literature on the subject were examined. This approach enabled a comprehensive review of the Act's background, content, and implications. Additionally, legal and political analyses were conducted to assess the Act's potential to emerge as a global standard via the "Brussels Effect," which identifies its strengths, weaknesses, and challenges. This methodology allowed the researcher to critically analyze the Act's innovative features and impacts while providing a detailed description of its structure.

The research structure follows a logical progression, beginning with a detailed analysis of the key components of the EU AI Act. This includes:

- An in-depth exploration of the historical development and intricacies of the EU AI Act and its associated regulatory framework. This encompasses a thorough examination of the Act's definition of AI, its scope of application, and its risk-based approach.
- An evaluation of the EU AI Act's potential to become a global standard through the "Brussels Effect." This analysis assesses the mechanisms for extraterritorial application of the Act and explores both the opportunities and challenges associated with achieving this influence.
- An investigation of the innovative contributions of the EU AI Act. This includes an analysis of how the Act's unique features, such as its definition of AI, scope of application, and risk-based approach, may be adopted or adapted by other jurisdictions.

This methodological approach ensured a thorough and comprehensive examination of the EU AI Act and its potential global impact.

3. Results and Discussion

The development of the EU AI Act commenced in 2017 with the European Commission's initiative to establish a regulatory framework for AI. This was followed by the European Council's call for a unified European approach to AI, emphasizing data protection, digital rights, and ethical standards. The Commission's 2018 "European Strategy on AI" highlighted the necessity of a robust ethical and legal framework grounded in EU values and the Charter of Fundamental Rights. The subsequent establishment of a High-Level Expert Group on AI in 2018 led to the development of ethical guidelines, incorporating key principles such as human oversight, technical robustness, privacy, transparency, diversity, and accountability. The Commission's 2020 White Paper on AI further emphasized the need to address both the opportunities and risks associated with AI. The AI Act was formally proposed in April 2021 and, following amendments by the European Parliament, was finally approved by the European Council on May 21, 2024.

The EU aspires to establish itself as a global leader in AI regulation through the "Brussels Effect," which signifies the EU's inherent power to influence global markets through its regulatory actions. This is achieved by setting standards that companies worldwide must comply with in order to access the EU market. Key factors contributing to the Brussels Effect include the EU's substantial market size, its robust regulatory capacity, and the potential for non-regression of standards. The EU has previously demonstrated this effect in areas such as data protection, with the General Data Protection Regulation (GDPR) significantly influencing global data protection practices. The AI Act's broad scope, encompassing various AI systems utilized in regulated sectors, may compel global providers to align with EU standards. The Act's territorial reach extends to any AI system that impacts the EU market, regardless of the provider's location. However, challenges remain, such as the absence of comprehensive rules for systems that are neither banned nor classified as high-risk, and the potential for providers to offer different versions of their products to circumvent compliance. Furthermore, the global dissemination of AI safety standards based on the EU AI Act may not adequately safeguard fundamental rights, democracy, and the rule of law in all jurisdictions. This raises concerns regarding the legitimacy of the EU's unilateral regulatory actions and the potential for undermining the sovereignty of other nations.

The Act's definition of AI has undergone refinement since the initial draft, now emphasizing the capacity for "inference" and "autonomy." This definition encompasses a wide range of technologies, including those with machine-learning capabilities, while differentiating AI from simpler software based on its capacity for learning and adaptation. This broad definition ensures the Act's continued relevance amidst rapid technological advancements.

The Act possesses a broad, extra-territorial reach, encompassing providers, developers, importers, and distributors of AI systems, even those located outside the EU. The law prioritizes the impact of AI systems on the EU market rather than the location of the providers. This approach safeguards individuals within the EU who may be affected by AI systems. While comprehensive, the law includes certain exceptions for military, research, personal use, and other activities already subject to existing legislation.

The Act employs a risk-based approach to classify AI systems into four distinct categories:

- **Unacceptable Risk:** AI systems posing an unacceptable risk to human rights and EU values are prohibited. This includes systems that classify individuals based on sensitive data or manipulate human behavior.
- **High-Risk:** These systems, identified in Annex III, require strict adherence to risk management and quality standards. While not prohibited, they must comply with stringent requirements.
- **Limited Risk:** Systems in this category, such as chatbots, are subject to minimal transparency obligations.
- **Minimal/No Risk:** AI systems that do not fall into the above categories are exempt from specific obligations.

This risk-based approach acknowledges the inherent limitations of achieving complete safety and aims to strike a balance between AI innovation and the protection of individual rights and values.

4. Conclusions and Future Research

The EU AI Act represents a groundbreaking piece of legislation aimed at striking a balance between technological advancement, ethical considerations, and the protection of fundamental rights. Its primary objective is to foster trustworthy AI within the European Union, ensuring that AI systems comply with ethical principles. Following an extensive legislative process initiated in 2017, the Act was finalized in May 2024. The EU aspires to set a global standard for AI governance through the "Brussels Effect," potentially leading to widespread adoption of its regulatory framework. However, challenges such as the lack of specific requirements for certain AI categories and potential conflicts with national sovereignty persist. The Act's definition of AI is broad and adaptable, focusing on the core concepts of inference and autonomy. Its scope is extensive, encompassing both EU-based and non-EU-based companies targeting the EU market, with certain exceptions. The risk-based classification of AI systems offers a nuanced approach to regulating this complex and rapidly evolving technology.

In conclusion, the EU AI Act marks a significant milestone in establishing a regulatory framework for AI. Its impact will be closely monitored by the international community as nations navigate the challenges and opportunities presented by this transformative technology.

Future research should focus on the Act's practical implementation, its impact on corporate innovation and competitiveness, and implications for non-EU countries like Iran. Comparative studies of AI regulations across regions and analyses of the Brussels Effect's global influence could deepen understanding of the EU's role in AI governance. Additionally, legal challenges arising from the Act's extraterritorial application particularly regarding international law principles such as state sovereignty warrant further exploration.





قانون هوش مصنوعی اتحادیه اروپایی: گامی به سوی نظارت انسانی بر هوش مصنوعی

عباسعلی کدخدایی الیادرائی^۱، ابوالفضل خاکزاد بجستانی^۲، امیرمحمد مجیدی^۳

^۱ استاد، گروه حقوق عمومی، دانشکده حقوق و علوم سیاسی، دانشگاه تهران، تهران، ایران.

^۲ دانشجوی کارشناسی ارشد حقوق بین الملل، گروه حقوق عمومی، دانشکده حقوق و علوم سیاسی، دانشگاه تهران، تهران، ایران.

^۳ دانشجوی کارشناسی ارشد حقوق بین الملل، گروه حقوق عمومی، دانشکده حقوق و علوم سیاسی، دانشگاه تهران، تهران، ایران.

اطلاعات مقاله	چکیده
مقاله پژوهشی	در عصر حاضر، هوش مصنوعی نقشی انکارناپذیر در زندگی بشر ایفا می کند؛ همان گونه که فناوری دیجیتال به رکن اصلی حیات جوامع بشری بدل شده است. بنابراین، اعتماد جوامع به این فناوری امری ضروری تلقی می شود. برای ایجاد چنین اعتمادی، اعمال نظارت انسانی بر این فناوری ها با استفاده از قانون گذاری یکی از ابزارهای مؤثر خواهد بود. قانون هوش مصنوعی اتحادیه اروپایی اولین چهارچوب قانونی است که به دنبال صیانت از ارزش های انسانی است. نقطه شروع این قانون در سال ۲۰۱۷ با تلاش های کمیسیون اروپایی بوده و سرانجام در تاریخ ۲۱ می ۲۰۲۴ توسط شورای اروپایی تصویب شده است. این قانون به دنبال حکمرانی اتحادیه اروپایی بر مقررات گذاری هوش مصنوعی به عنوان یک الگوست؛ بنابراین سؤال اصلی مقاله این است که اولاً این قانون قابلیت تبدیل شدن به الگوی جهانی از طریق سازوکار اثر بروکسل را دارد چه آثاری بر آن مترتب است؟ ثانیاً این قانون چه آورده های نوآورانه ای در مسیر نظارت حقوقی بر هوش مصنوعی ارائه کرده است؟ بدین منظور سعی شده است با روش توصیفی - تحلیلی و شیوه گردآوری اطلاعات به صورت کتابخانه ای به این سؤالات پرداخته شود. در پایان نویسندگان به این نتیجه می رسند اتحادیه اروپایی با تصویب اولین قانون جامع هوش مصنوعی جهان، تعادلی بین پیشرفت فناوری هوش مصنوعی و محافظت از حقوق اساسی شهروندان ایجاد کرده است و می خواهد از طریق سازوکار اثر بروکسل، رهبری جهانی در تنظیم مقررات هوش مصنوعی را به دست گیرد. همچنین این قانون در این مسیر نوآوری هایی همچون تعریف هوش مصنوعی، ایجاد صلاحیت واقعی در قلمرو قانون و رویکرد مبتنی بر خطر ارائه داده است. این نوآوری ها علاوه بر افزایش شفافیت حقوقی، به شکل گیری رویه ای برای اعمال حاکمیت فراملی در حوزه هوش مصنوعی کمک می کنند.
تاریخ دریافت: ۱۴۰۳/۰۶/۰۷	
تاریخ پذیرش: ۱۴۰۳/۱۲/۰۵	
واژگان کلیدی: اتحادیه اروپایی اثر بروکسل رویکرد مبتنی بر خطر هوش مصنوعی	
*نویسنده مسئول رایانامه: abolfazl.khakzad@ut.ac.ir	

نحوه استناددهی:

کدخدایی، عباسعلی، خاکزاد بجستانی، ابوالفضل، و مجیدی، امیرمحمد (۱۴۰۵). قانون هوش مصنوعی اتحادیه اروپایی: گامی به سوی نظارت انسانی بر هوش مصنوعی. حقوق فناوری های نوین، ۷(۱۳)، ۱۷۱-۱۹۰.

ناشر: دانشگاه علم و فرهنگ <https://www.usc.ac.ir>

شاپای الکترونیکی: ۲۷۸۳-۳۸۳۶

استفاده روزافزون از هوش مصنوعی^۱ و سیستم‌های مبتنی بر آن، برای کمک یا تقویت تصمیم‌گیری انسان به روش‌های مختلف، طبیعتاً موضوعی داغ در بحث‌های عمومی و دانشگاهی است. این مسئله به گونه‌ای است که حتی وکلا امروزه برای امور حقوقی خود از هوش مصنوعی مدد می‌گیرند.^۲ یکی از معضلات اساسی مرتبط با این مباحث این است که چگونه می‌توان از فناوری‌های مبتنی بر هوش مصنوعی به طور کامل، بدون خطر و آسیب رساندن به جامعه و شهروندان، استفاده کرد؟ در مواجهه با هوش مصنوعی، چالش اصلی جوامع انسانی این است که چگونه می‌توان اطمینان حاصل کرد کنترل مؤثر انسانی در جایی که نیاز است اعمال شود. در این زمینه، به «نظارت انسانی»^۳ به منزله یک معیار، بسیار تأکید شده است. این معیار نشان‌دهنده نوعی فیلتر متعادل‌کننده بین اهداف سیستم مبتنی بر هوش مصنوعی و استدلال معتبر انسانی است.^۴ (Clyde, 2021).

نظارت انسانی به سازوکاری کلیدی برای مدیریت هوش مصنوعی تبدیل شده است. مقررات‌گذاری حقوقی یکی از مهم‌ترین جلوه‌های نظارت انسانی است که قرار است دقت و ایمنی سیستم‌های هوش مصنوعی را افزایش دهد، از ارزش‌های انسانی حمایت کند و باعث اعتماد به فناوری‌های نوظهور شود. مطابق بند یک ماده یک قانون هوش مصنوعی اتحادیه اروپایی، هدف از این قانون، بهبود عملکرد بازار داخلی و ارتقای استفاده از هوش مصنوعی انسانی محور^۵ و قابل اعتماد است. در عین حال، این قانون معیاری برای حفظ سطح بالایی از سلامت، ایمنی و حقوق بشر نیز به شمار می‌آید. حقوقی که در این قانون مدنظر قرار گرفته عبارت است از: دموکراسی، حاکمیت قانون و حفاظت از محیط زیست. از این موارد باید در برابر تأثیرات منفی سیستم‌های هوش مصنوعی در اتحادیه اروپا محافظت شود و در عین حال از نوآوری‌ها نیز حمایت شود، حفاظتی که ایجاد آن را قانون مذکور حاصل می‌کند. در واقع این قانون میان بهره‌وری از هوش مصنوعی و حقوق اساسی انسان‌ها توازن برقرار می‌کند و به سلامت و ایمنی جامعه آسیب نمی‌رساند.^۶ پیشرفت‌های اخیر در حوزه هوش مصنوعی نشان داده است که این فناوری قدرتمند قادر است تحولات عمیقی در جوامع انسانی به وجود آورد. با این حال، همراه با این پیشرفت‌ها، نگرانی‌هایی نیز درباره تأثیرات منفی و پتانسیل خطرناک هوش مصنوعی بر جوامع و حقوق انسانی به وجود آمده است. شفاف نبودن چگونگی و چرایی نتیجه‌گیری سیستم‌های مبتنی بر هوش مصنوعی، اعمال نفوذ در بُعد اجتماعی از طریق الگوریتم‌های هوش مصنوعی^۷، نظارت اجتماعی به عنوان یک نگرانی از کاربرد هوش مصنوعی^۸ و نقض حریم خصوصی داده‌ها با استفاده از ابزارهای هوش مصنوعی، از مهم‌ترین مشکلات تهدیدکننده حقوق جوامع انسانی است. به منظور مقابله با این مشکلات و تضمین استفاده اخلاقی و مسئولانه از هوش مصنوعی، اتحادیه اروپایی^۹ اقدام به تصویب «قانون هوش مصنوعی»^{۱۰}

1. Artificial Intelligence (AI).

۲. بر اساس نظرسنجی اخیر ABA، ۲۶ درصد وکلای دادگستری در شرکت‌هایی با بیش از ۱۰۰ وکیل از فناوری مبتنی بر هوش مصنوعی استفاده می‌کنند. Eckart, 2022, p. (275).

3. Human oversight.

۴. رجوع کنید به: <https://www.techpolicy.press/human-in-the-loop-systems-are-no-panacea-for-ai-accountability>.

5. Human-Centered AI.

۶. بند یک ماده یک قانون هوش مصنوعی در این زمینه مقرر می‌کند: «هدف این مقررات بهبود عملکرد بازار داخلی و ارتقای جذب هوش مصنوعی انسان‌محور و قابل اعتماد و در عین حال تضمین سطح بالایی از حفاظت از سلامت، ایمنی و حقوق اساسی مندرج در منشور بنیادین حقوق است. این حقوق شامل دموکراسی، حاکمیت قانون و حفاظت از محیط زیست، در برابر اثرات مضر سیستم‌های هوش مصنوعی در اتحادیه و حمایت از نوآوری‌ها است».

۷. رسانه و شبکه‌های اجتماعی مملو از تصاویر و ویدیوهای تولید شده توسط هوش مصنوعی، تغییر صداها، هوش مصنوعی و همچنین دیپ فیک، که در حوزه‌های سیاسی و اجتماعی نفوذ کرده‌اند، شده‌اند. با استفاده از این فناوری‌ها به راحتی می‌توان تصویر یک چهره با چهره دیگر در یک تصویر یا ویدیوی موجود را جایگزین کرد؛ در نتیجه شرایطی ایجاد می‌شود که در آن تشخیص اخبار معتبر از نامعتبر، تقریباً غیرممکن است. هیچ‌کس نمی‌داند واقعیت چیست و این یک مسئله بزرگ خواهد بود.

۸. نمونه بارز آن، استفاده چین از فناوری تشخیص چهره است. برای مطالعه بیشتر رجوع کنید به: <https://digiato.com/article/2022/06/29/how-china-monitoring-systems-sees-the-future>.

9. European Union.

10. Artificial Intelligence Act.

کرده است. این قانون و الزامات نظارتی آن نه تنها در قوانین اروپا، بلکه در جامعه بین‌المللی تازگی دارد.^۱ به بیانی دیگر، این قانون در پس‌زمینه‌ای معرفی می‌شود که در آن، هیچ تعریف قانونی یا هرگونه رویکرد انسان‌محور، که به‌طور خاص به هوش مصنوعی دلالت کند، نه از منظر اجتماعی و نه از منظر نظارتی، وجود ندارد^۲ (Enqvist, 2023, p. 3).

این قانون، که به‌عنوان «قانون هوش مصنوعی» شناخته می‌شود، به‌منظور تنظیم و کنترل استفاده از هوش مصنوعی در اتحادیه اروپایی ارائه شده است. مقررات این قانون برای توسعه، استقرار و استفاده از هوش مصنوعی در اتحادیه اروپایی یا زمانی که افراد در اتحادیه اروپایی تحت تأثیر هوش مصنوعی قرار می‌گیرند، اعمال می‌شود.

بنابراین در این مقاله سعی شده است با روش توصیفی - تحلیلی و شیوه گردآوری اطلاعات به‌صورت کتابخانه‌ای، ابعاد مختلف قانون هوش مصنوعی اتحادیه اروپایی بررسی شود؛ از این‌رو ابتدا به پیشینه تصویب این قانون، سپس به اهمیت بین‌المللی قانون و قابلیت تبدیل به الگوی جهانی و در نهایت به بررسی تحلیلی - انتقادی آورده‌های نوآورانه قانون هوش مصنوعی اتحادیه اروپایی پرداخته می‌شود.

۱. پیشینه قانون هوش مصنوعی

قانون هوش مصنوعی اتحادیه اروپایی اولین چهارچوب قانونی در زمینه هوش مصنوعی است که به خطرات هوش مصنوعی می‌پردازد و اروپا را در عرصه رقابت برای مقررات‌گذاری و تنظیم حقوقی هوش مصنوعی به بازیگری پیش‌رو در سطح جهانی تبدیل می‌کند. هدف قانون هوش مصنوعی ارائه الزامات و تعهدات روشن به توسعه‌دهندگان هوش مصنوعی درباره کاربردهای خاص هوش مصنوعی است. برای سال‌ها، کمیسیون اروپایی^۳ تلاش‌های فراوانی برای چهارچوب‌بندی و نظارت انسانی بر هوش مصنوعی به‌منظور تقویت رقابت‌پذیری و تضمین ارزش‌های اتحادیه اروپایی انجام داده است.

اولین جرعه در می ۲۰۱۷ زده شد؛ جایی که کمیسیون اروپایی بررسی میان‌مدت خود را از «راهبرد بازار واحد دیجیتال»^۴ منتشر کرد. این بیانیه بر اهمیت اتکا بر نقاط قوت علمی و صنعتی اروپا و همچنین استراتیپ‌های نوآورانه برای قرار گرفتن در موقعیت پیش‌رو در توسعه فناوری‌ها، پلتفرم‌ها و برنامه‌های کاربردی هوش مصنوعی تأکید کرد (EU Commission, 2017, p. 17).

شورای اروپایی^۵ در اکتبر ۲۰۱۷ اعلام کرد که اتحادیه اروپایی برای رسیدگی به پدیده‌های نوظهور مانند هوش مصنوعی به‌منظور دستیابی به سطح بالایی از حفاظت داده‌ها، حقوق دیجیتال و استانداردهای اخلاقی نیاز به فوریت دارد و از کمیسیون اروپایی برای ارائه پیشنهاد درمورد رویکرد اروپایی به هوش مصنوعی دعوت کرد (EU Council, 2017, pp. 1-5).

در سال ۲۰۱۸ کمیسیون اروپایی «اهبرد اروپایی درمورد هوش مصنوعی»^۶ را منتشر کرد. این برنامه در راستای بیان اهداف اتحادیه اروپایی برای تصویب قانون هوش مصنوعی مقرر می‌کند «کسب اطمینان از یک چهارچوب اخلاقی و قانونی مناسب براساس ارزش‌های اتحادیه در

۱. گفتنی است در تاریخ ۲۱ مارس ۲۰۲۴ مجمع عمومی سازمان ملل متحد نخستین قطعنامه درمورد هوش مصنوعی را تصویب کرد:

<https://news.un.org/en/story/2024/03/1147831>

۲. در این زمینه مارگرت وستاگر، معاون اجرایی اروپای سازگار با عصر دیجیتال، می‌گوید: «درمورد هوش مصنوعی، اعتماد یک امر ضروری است. با این قوانین مهم، اتحادیه اروپایی درحال پیشبرد توسعه هنجارهای جهانی جدید است تا اطمینان حاصل شود که هوش مصنوعی قابل اعتماد است. با تنظیم استانداردها، می‌توانیم راه را برای فناوری اخلاقی در سراسر جهان هموار کنیم و اطمینان حاصل کنیم که اتحادیه اروپایی در طول مسیر رقابتی باقی می‌ماند». همچنین تیری برتون، کمیسر بازار داخلی خاطرنشان کرد: «هوش مصنوعی یک وسیله است، نه یک هدف. ده‌ها سال است که وجود داشته است، اما به ظرفیت‌های جدیدی رسیده است که توسط قدرت محاسباتی تأمین می‌شود. این امر پتانسیل بسیار زیادی را در حوزه‌های متنوعی مانند سلامت، حمل‌ونقل، انرژی، کشاورزی، گردشگری یا امنیت سایبری ارائه می‌کند. همچنین مجموعه‌ای از خطرات را به همراه دارد. هدف پیشنهادات امروز، تقویت موقعیت اروپا به عنوان یک مرکز جهانی تعالی در هوش مصنوعی است؛ همچنین از این‌که هوش مصنوعی در اروپا به ارزش‌ها و قوانین ما احترام می‌گذارد، اطمینان حاصل شود (European Commission, 2021).

3. European Commission.

4. The Mid-Term Review on the implementation of the Digital Single Market.

5. European Council.

6. Strategy on AI for Europe.

راستای منشور حقوق اساسی اروپا؛ این مسئله شامل راهنمایی‌های آتی در مورد قوانین مسئولیت محصولات، تجزیه و تحلیل دقیق مشکلات نوظهور و همکاری با سهام‌داران برای توسعه دستورالعمل‌های اخلاقی هوش مصنوعی است» (EU Commission, 2018, p. 3). پس از اطلاعیه کمیسیون اروپایی تحت عنوان «راهبرد اتحادیه اروپا در زمینه هوش مصنوعی»، کشورهای عضو (به‌منزله بخشی از گروه دیجیتالی کردن صنعت اروپا و هوش مصنوعی)،^۱ نروژ و سوئیس، طی جلسات متعددی بین ژوئن و نوامبر ۲۰۱۸، یک طرح برای هوش مصنوعی آماده کردند. در طول این جلسات، کشورهای عضو و کمیسیون مجموعه‌ای از اقدامات مشترک را برای افزایش سرمایه‌گذاری، گردآوری داده‌ها و پرورش استعدادها و تضمین از ایجاد اعتماد در راهبرد اروپایی را شناسایی کردند. آن‌ها حوزه‌های محبوب عموم، مانند مراقبت‌های بهداشتی، حمل‌ونقل، امنیت و انرژی و همچنین بخش‌های مهم اقتصادی مانند تولید و خدمات مالی را در اولویت قرار دادند (EU Commission, 2018, p. 2). نتیجه این کار مشترک یک «برنامه هماهنگ»^۲ در مورد هوش مصنوعی است که طی این برنامه هماهنگ، قرار است اقداماتی در سال‌های ۲۰۱۹ و ۲۰۲۰ آغاز و به‌صورت سالانه به‌روز شود.

در سال ۲۰۱۸ کمیسیون اروپایی اطلاعیه‌ای مبنی بر نظارت انسانی بر هوش مصنوعی در راستای اعتماد بر آن، خطاب به پارلمان اروپا، شورای اروپایی، کمیته اقتصادی و اجتماعی اروپا و کمیته مناطق منتشر کرد. هم‌زمان کمیسیون یک «کارگروه تخصصی سطح بالا در زمینه هوش مصنوعی»^۳ به نمایندگی از طیف گسترده‌ای از ذی‌نفعان تشکیل داد و آن را موظف به تهیه پیش‌نویس دستورالعمل‌های اخلاقی هوش مصنوعی و همچنین تهیه مجموعه‌ای از توصیه‌ها برای سیاست‌های گسترده‌تر هوش مصنوعی کرد.^۴

کارگروه تخصصی سطح بالا در زمینه هوش مصنوعی اولین پیش‌نویس دستورالعمل‌های اخلاقی را در دسامبر ۲۰۱۸ منتشر کرد. این دستورالعمل مقرر می‌کند برای دستیابی به هوش مصنوعی اطمینان‌بخش وجود سه جزء ضروری است: (۱) مطابقت با قوانین اتحادیه اروپایی؛ (۲) رعایت اصول اخلاقی؛ (۳) داشتن قدرت الزام‌آوری.

براساس این سه مؤلفه و ارزش‌های اروپایی، این دستورالعمل هفت الزام کلیدی را مشخص می‌کند که برنامه‌های هوش مصنوعی باید به آن‌ها احترام بگذارند تا قابل اعتماد تلقی شوند. همچنین این دستورالعمل شامل یک فهرست ارزیابی است که بررسی کند آیا این الزامات برآورده شده‌اند یا خیر؟ (EU Commission, 2019, p. 3)

هفت الزام کلیدی عبارت‌اند از: نظارت انسانی، استحکام فنی و ایمنی، حریم خصوصی و حاکمیت داده، شفافیت، تنوع و عدم تبعیض و انصاف، رفاه اجتماعی و محیطی و مسئولیت‌پذیری. این الزامات برای تمامی سیستم‌های هوش مصنوعی، صرف‌نظر از نوع آن‌ها، در نظر گرفته شده‌اند. با این حال، رویکردی مبتنی بر ارزیابی میزان تأثیرگذاری این سیستم‌ها بر افراد و جامعه بر آن‌ها اعمال می‌شود. برای مثال یک سیستم هوش مصنوعی، که کتابی نامناسب را به کاربر انسانی پیشنهاد می‌دهد، به نسبت آن سیستم، که بیماری سرطان را به اشتباه تشخیص می‌دهد، خطر کمتری دارد؛ در نتیجه ممکن است به نظارت کمتری نیاز داشته باشد (EU Commission, 2019, p. 3).

این دستورالعمل را کارگروه تخصصی سطح بالا در زمینه هوش مصنوعی غیرالزام‌آور تهیه کرده است و به این ترتیب هیچ تعهد قانونی ایجاد نمی‌کند. با این حال، بسیاری از مقررات این دستورالعمل قبلاً در قوانین اتحادیه اروپایی در زمینه قوانین ایمنی، حفاظت از داده‌های شخصی، حریم خصوصی و قوانین حفاظت از محیط‌زیست منعکس شده بود (EU Commission, 2019, p. 4).

در سال ۲۰۲۰ «اوراق سفید»^۵ کمیسیون اروپایی در مورد هوش مصنوعی منتشر شد. سند مذکور چشم‌انداز روشنی را برای هوش مصنوعی در اروپا ارائه می‌کند. در بخش پنجم این سند، ضرورت چهارچوب‌بندی هوش مصنوعی به این صورت بیان شده است: «مانند هر فناوری جدید

1. Group on Digitising European industry and AI.

2. Coordinated Plan.

3. High-level expert group on AI(HLEG).

۴. رجوع کنید به: <https://digital-strategy.ec.europa.eu/en/policies/expert-group-ai>

۵. White Paper: وایت پیپر، که در فارسی تحت‌عنوان «اوراق سفید» ترجمه شده است، درحقیقت به گزارشی جامع و کامل از عملکرد یک طرح یا پروژه به‌منظور ارائه راهکار یا توضیح یک مسئله تهیه می‌شود. این مستند معمولاً شامل تحلیل‌ها، داده‌ها، استدلال‌ها و پیشنهادات مربوط به موضوع مدنظر است.

استفاده از هوش مصنوعی، فرصت‌ها و خطراتی را به همراه دارد. هوش مصنوعی همان‌طور که می‌تواند به محافظت از شهروندان کمک کند، می‌تواند اثرات ناخواسته‌ای داشته باشد یا حتی برای اهداف مخرب استفاده شود؛ این نگرانی‌ها باید برطرف شود. علاوه بر فقدان سرمایه‌گذاری و مهارت، عدم اعتماد عامل اصلی برای جلوگیری از جذب گسترده هوش مصنوعی است» (EU Commission, 2020, p. 9).

در نهایت در آوریل ۲۰۲۱، کمیسیون پیش‌نویس قانون هوش مصنوعی را ارائه کرد و در دسامبر ۲۰۲۳، قانون‌گذاران اتحادیه اروپایی به یک توافق سیاسی در مورد هوش مصنوعی دست یافتند. در تاریخ ۱۴ ژوئن ۲۰۲۳، اصلاحاتی بر پیشنهاد کمیسیون توسط پارلمان اروپایی تصویب شد و در نهایت در تاریخ ۱۳ مارس ۲۰۲۴، قطعنامه قانون هوش مصنوعی توسط پارلمان اروپایی و در تاریخ ۲۱ می ۲۰۲۴، توسط شورای اروپایی تصویب شد.

۲. الگوی بین‌المللی: قابلیت ایجاد اثر بروکسل در جهان

با تصویب قانون جامع هوش مصنوعی، پیش‌بینی می‌شود اتحادیه اروپایی به‌عنوان پیشگام جهانی در تنظیم مقررات این فناوری نوین شناخته شود. یکی از سازوکارهایی که برای نفوذ جهانی قانون هوش مصنوعی در نظر گرفته شده «اثر بروکسل»^۱ است؛ از این‌رو در بخش نخست به مفهوم‌شناسی اثر بروکسل، در بخش دوم به امکان‌سنجی ایجاد اثر بروکسل در قانون هوش مصنوعی، در بخش سوم به مشروعیت حقوق بین‌المللی سازوکار اثر بروکسل و در پایان، به پیامدهای سازوکار اثر بروکسل بر ایران پرداخته می‌شود.

۱-۲. مفهوم اثر بروکسل

اتحادیه اروپایی در دهه‌های اخیر، با وضع مقررات پیشرو در حوزه‌هایی مانند ایمنی مواد غذایی، استانداردهای شیمیایی، حقوق رقابت و حریم خصوصی، جایگاهی بی‌بدیل در شکل‌دهی به قواعد جهانی به‌دست آورده است. جالب آن‌که تأثیر این مقررات فراتر از مرزهای اروپاست و زندگی روزمره میلیاردها نفر در سراسر جهان را تحت تأثیر قرار می‌دهد. برای نمونه، بسیاری از شهروندان آمریکایی از این واقعیت غافل‌اند که استانداردهای اتحادیه اروپایی بر لوازم آرایشی مصرفی، ترکیبات غلات صبحانه، و حتی نرم‌افزارهای نصب شده بر رایانه‌های شخصی آنان حاکم است (Bradford, 2012, p. 3). در یک تعریف کوتاه، اثر بروکسل^۲ یعنی «قدرت یک‌جانبه اتحادیه اروپایی برای تنظیم بازار جهانی». جهانی شدن نظارت یک‌جانبه زمانی اتفاق می‌افتد که یک حوزه قضایی بتواند قوانین و مقررات خود را در خارج از مرزهای خود از طریق سازوکارهای بازار اعمال کند تا از این طریق، استانداردهای مطلوب یک حوزه قضایی، جهانی شوند (Bradford, 2012, p. 3).

این پدیده زمانی محقق می‌شود که یک بازیگر منطقه‌ای، نظیر اتحادیه اروپایی، بدون توسل به معاهدات و موافقت‌نامه‌ها یا ابزارهایی نظیر تحریم‌ها، صرفاً از طریق جذابیت بازار داخلی خود، شرکت‌ها و دولت‌ها را به تطبیق با استانداردهایش وادار کند (Bradford, 2012, p. 4). بنابراین، اتحادیه اروپایی بدون توسل به نهادهای بین‌المللی برای همکاری، اعمال زور و تحریم یک‌جانبه، مقرراتی را منتشر می‌کند که در چهارچوب‌های قانونی بازارهای توسعه‌یافته و در حال توسعه به‌طور یکسان تثبیت شده و منتهی به اروپایی‌سازی جنبه‌های مهم تجارت جهانی و حکمرانی اتحادیه اروپایی می‌شود. برای اعمال اثر بروکسل، شرایطی در نظر گرفته می‌شود؛ از قبیل قدرت بازار^۳، ظرفیت نظارتی^۴، وجود مقررۀ قاطع، پتانسیل برای اهداف غیرارتجاعی^۵ و تقسیم‌ناپذیری استانداردها^۶ که پرداختن به این شرایط، می‌تواند موضوع نوشتاری جداگانه قرار گیرد

1. Brussels Effect.

۲. در ادبیات حقوقی، پژوهشگران حوزه حقوق اتحادیه اروپایی، این تأثیر نظارتی اتحادیه اروپایی از طریق اثر بروکسل محقق می‌شود. اصطلاحی که آنو برادفورد (Anu Bradford) استاد دانشکده حقوق دانشگاه کلمبیا، در مقاله سال ۲۰۱۲ خود برای نخستین بار به کار برد. این کاربرد به ریشه‌های قدرت نظارتی - مقرراتی اتحادیه اروپایی، که از نهادهای مستقر در بروکسل نشئت می‌گیرد، اشاره دارد تا توضیح دهد که چگونه جهانی شدن نظارتی یک‌جانبه اتحادیه اروپایی، از طریق سازوکارهای بازار رخ داده و اتحادیه اروپایی، قوانین خود را، که دربرگیرنده استانداردهای سخت‌گیرانه است، در خارج از مرزهای خود اعمال می‌کند.

3. Market Power.

4. Regulatory Capacity.

5. Predisposition to Regulate Inelastic Targets.

6. Nondivisibility of Standards.

(Bradford, 2012, pp. 10-18). همچنین اثر بروکسل هم به صورت «دوفاکتو»^۱ و هم به صورت «دوژوره»^۲ اعمال می‌شود^۳ (Siegmann & Anderljung, 2022, p. 3).

۲-۲. امکان‌سنجی ایجاد اثر بروکسل در قانون هوش مصنوعی

در سال ۲۰۱۹ دو تن از قدرتمندترین سیاست‌مداران اروپایی، یعنی اوسولا فون درلاین^۴ و آنگلا مرکل^۵ بعد از تجربه موفق اتحادیه اروپایی در تدوین «مقررات عمومی مربوط به حفاظت از داده‌ها»^۶ این بار از اتحادیه اروپایی خواستار تدوین مقررات حفاظت از داده‌ها در زمینه هوش مصنوعی شدند. «مقررات عمومی حفاظت از داده‌ها» یکی از مؤثرترین مقررات قانون‌گذاری اتحادیه اروپایی در دهه گذشته است. این مقررات نه تنها رویه‌های تجاری در اتحادیه اروپایی را تغییر داد، بلکه موجب ایجاد اثر بروکسل خارج از مرزهای اتحادیه اروپایی شد (Siegmann & Anderljung, 2022, p. 3).

اتحادیه اروپایی با تصویب مقرراتی در حوزه حفاظت از داده، مانند «مقررات عمومی حفاظت از داده‌ها» و «دستورالعمل حفظ حریم خصوصی الکترونیکی»^۸ توانست به عنوان یکی از مؤثرترین نهادهای فعال در عرصه مقررات دیجیتال ظاهر شود و از این راه، دیگر کشورها را نیز به اجرای مقررات خود وادارد؛ به نحوی که در حال حاضر بسیاری از وبسایت‌های جهان، الزامات اتحادیه اروپایی را دنبال می‌کنند؛ این اقتدار اروپا در عرصه داده‌ها یکی از جلوه‌های اثر بروکسل است که جریان‌ساز رویه‌ای است که در آن شرکت‌ها در ابعاد جهانی به خواست خود از قوانین اتحادیه اروپایی پیروی می‌کنند تا یک محصول، خدمات یا فرایندهای تجاری را مطابق با استانداردهای اتحادیه اروپایی طراحی کنند. در ادامه دلایلی بیان می‌شود که امکان جهانی‌سازی قانون هوش مصنوعی را از طریق اثر بروکسل فراهم می‌کند:

الف) اتحادیه اروپایی می‌تواند از طریق جذابیت بازار ۴۵۰ میلیونی خود، قانون هوش مصنوعی را به استانداردی جهانی تبدیل کند. قانون هوش مصنوعی اتحادیه اروپایی، طیف گسترده‌ای از سیستم‌های هوش مصنوعی استفاده شده در بخش‌های از قبل تنظیم شده، مانند هوانوردی، وسایل نقلیه خودروبی، آسانسورها، دستگاه‌های پزشکی، ماشین‌آلات صنعتی و موارد دیگر را پوشش می‌دهد. از آنجاکه انطباق با قانون هوش مصنوعی پیش‌نیاز فروش و عرضه سیستم‌های هوش مصنوعی در کشورهای عضو اتحادیه است، ارائه‌دهندگانی که در سطح جهانی فعالیت می‌کنند باید خود را با این استانداردها هماهنگ کنند (Siegmann & Anderljung, 2022, pp. 28-30).

ب) از آنجاکه هوش مصنوعی فناوری جدیدی است که در چند سال گذشته پیشرفت‌های چشمگیری داشته است، دانش موجود در مورد نحوه تنظیم حقوقی سیستم‌های مبتنی بر هوش مصنوعی محدود است؛ اما اتحادیه اروپایی با اصلاح قانون حفاظت از داده‌های اتحادیه اروپایی بر محور هوش مصنوعی در سال ۲۰۱۰ و استخدام یک کارگروه تخصصی سطح بالا در زمینه هوش مصنوعی، متشکل از افراد صنعتی و دانشگاهی، این نگرانی را از بین برده است (Almada & Radu, 2024, p. 10). اتکا به قانون هوش مصنوعی اتحادیه اروپایی اجازه می‌دهد که دهه‌ها تخصص و تجربه اتحادیه به این قانون منتقل شود؛ نه این‌که مجموعه‌ای کاملاً جدید از نهادها و شیوه‌ها در این حوزه ایجاد شود.

1. De facto.

2. De jure.

۳. درجایی که شرکت‌های بین‌المللی، با هدف ورود به بازار واحد، تولیدات و فعالیت‌های خود را در سطح جهانی با مقررات و استانداردهای اتحادیه اروپایی هماهنگ می‌سازند اثر بروکسل به صورت دوفاکتو، و در جایی که دولت‌های ثالث غیر عضو اتحادیه اروپایی با الگوبرداری از این مقررات و استانداردها، قوانین داخلی خود را اصلاح می‌کنند، اثر بروکسل به صورت دوژوره حاصل می‌گردد (Bernat & Spera, 2024, p. 82).

4. Ursula von der Leyen.

5. Angela Merket

6. General Data Protection Regulation (GDPR).

۷. هدف این مقرره حمایت از داده‌ها و اطلاعات شخصی است که در سال ۲۰۱۸ برای همه کشورهای عضو اتحادیه اروپایی لازم‌الاجرا شده است. این قانون با وضع مقررات سخت از جمله جرمه سنگین و اعتباربخشی به خسارت معنوی، برای شرکت‌هایی که با اتحادیه مراد و تبادل اطلاعات می‌کنند، دستورالعمل‌هایی را مشخص کرده است (Farahmand, 2024, p. 19).

8. Privacy Directive.

ج) نگرانی در مورد هوش مصنوعی این است که ارائه دهندگان بتوانند به راحتی سیستم های خود را در خارج از اتحادیه اروپایی ارائه دهند و در این صورت، انگیزه ای برای پیروی از چهارچوب سختگیرانه اتحادیه اروپایی نداشته باشند؛ اما قانون هوش مصنوعی اتحادیه اروپایی این نگرانی را از طریق سازوکار گسترش سرزمینی مقررات خود حل کرده است؛ به این معنا که مفاد قانون هوش مصنوعی برای هر سیستم هوش مصنوعی که خروجی هایش در اتحادیه قرار دارد قابل اجراست؛ حتی اگر ارائه دهنده یا کاربر در کشور ثالث مستقر باشند^۱.

از سوی دیگر، بر سر راه جهانی سازی قانون هوش مصنوعی اتحادیه اروپایی موانعی وجود دارد:

۱. همان طور که قبلاً ذکر شد، یکی از شرایط اعمال اثر بروکسل این است که آن مقرره باید قاطع باشد. قانون هوش مصنوعی در این خصوص موقعیت شکننده ای دارد؛ به این معنا که در مورد سیستم های هوش مصنوعی، که نه ممنوع هستند و نه به منزله سیستم پرخطر طبقه بندی می شوند، قانون هوش مصنوعی چهارچوب جامعی را ایجاد نمی کند و بیشتر خود را به تعهدات شفافیت برای سیستم های مبتنی بر هوش مصنوعی خاص، موضوع ماده ۵۲ قانون هوش مصنوعی را محدود می کند^۲ (Almada & Radu, 2024, p. 10)؛

۲. یکی از پیش شرط های تحقق «اثر بروکسل»، عدم تقسیم پذیری استاندارد هاست. این مفهوم به معنای آن است که برای تحقق اثر بروکسل، ارائه دهندگان فناوری مجاز به آن نباشند که بتوانند محصولات یا خدمات خود را به صورت تفکیک شده یعنی نسخه ای خاص برای بازار اتحادیه اروپایی و نسخه ای دیگر برای بازارهای غیراروپایی طراحی کنند. در صورت امکان چنین تقسیم بندی ای، انگیزه ای برای رعایت استانداردهای اروپایی در خارج از حوزه قضایی اتحادیه وجود نخواهد داشت. در مورد سیستم های ممنوعه هوش مصنوعی این اصل نقض می شود. قانون هوش مصنوعی اتحادیه فاقد سازوکار الزام آور برای جلوگیری از تقسیم پذیری است؛ در نتیجه ارائه دهندگان می توانند با حفظ دو رویکرد موازی، یعنی تطبیق محصولات خود با استانداردهای اروپایی برای بازار داخلی اتحادیه، و در کنار آن، تجاری سازی همان سیستم های ممنوعه در حوزه های خارج از اتحادیه که فاقد محدودیت های مشابه هستند، به فعالیت خود ادامه دهند. این وضعیت، چالش «دوگانه تنظیمگری»^۳ را ایجاد می کند که در آن استانداردهای سختگیرانه اتحادیه اروپایی تنها به صورت جغرافیایی محدود می شوند و تأثیر جهانی مدنظر را محقق نمی سازند (Bradford, 2012, pp. 17-19)؛

۳. اثر بروکسل برای قانون هوش مصنوعی یک عارضه جانبی شایان توجه ایجاد می کند و آن عبارت است از این که انتشار جهانی استانداردهای ایمنی هوش مصنوعی براساس قانون هوش مصنوعی، محافظت کافی از ارزش هایی مانند حقوق اساسی، دموکراسی و حاکمیت قانون که اتحادیه اروپایی بر آن بنا شده است ارائه نمی دهد. به عبارت دیگر، استانداردهای مبتنی بر قانون هوش مصنوعی می توانند با تحمیل هنجارهایی با نگاه محدودکننده خطرات جدی برای این ارزش ها ایجاد کنند. اگر کاستی های قانون هوش مصنوعی در طول فرایند آن برطرف نشود، اثر بروکسل می تواند به تضعیف ارزش های اتحادیه اروپایی منجر شود (Almada & Radu, 2024, p. 13).

۲-۳. مشروعیت حقوق بین المللی سازوکار اثر بروکسل

پدیده اثر بروکسل و گسترش سازوکارهای اعمال «فراسرزمینی»^۴ مقررات اتحادیه اروپایی، پرسش های بنیادینی را در خصوص مشروعیت این سازوکارها مطرح می سازد. از یک سو، این فرایند به اتحادیه اروپایی امکان می دهد بر فعالیت هایی که لزوماً در قلمرو جغرافیایی آن صورت نمی گیرد نظارت داشته باشد و به این ترتیب، بر محتوای قوانین کشورهای ثالث و حتی بر شکل گیری هنجارهای حقوق بین الملل تأثیرگذار باشد. این امر، ضمن گسترش حوزه نفوذ اتحادیه اروپایی، چالش هایی جدی را برای اصول بنیادین حقوق بین الملل، از جمله اصل «حاکمیت

1. Article 2(1)(b) AI Act.

۲. ماده ۵۲ قانون هوش مصنوعی مقرر می کند: ارائه دهندگان سیستم های هوش مصنوعی باید اطمینان حاصل کنند سیستم هایی که برای تعامل مستقیم با اشخاص حقیقی در نظر گرفته شده اند؛ به گونه ای طراحی و توسعه یابند که افراد حقیقی مربوطه از این مسئله که در حال تعامل با یک سیستم هوش مصنوعی هستند، مطلع شوند.

3. Regulatory Dilemma.

4. Extra-territorial.

دولت‌ها»^۱ و اصل «عدم مداخله»^۲ در امور داخلی کشورها، ایجاد می‌کند. از سوی دیگر، فقدان رضایت دولت‌های ثالث از این رویکرد یکجانبه‌گرایانه، مشروعیت اقدامات اتحادیه اروپایی را به طور جدی زیر سؤال می‌برد و پرسش‌هایی را در خصوص توازن قدرت در نظام بین‌الملل و مشروعیت دموکراتیک تصمیم‌گیری‌های اتحادیه اروپایی مطرح می‌سازد (Hadjiyianni, 2021, pp. 261-262).

به بیان دیگر، گسترش نفوذ قوانین اتحادیه اروپایی فراتر از مرزهای جغرافیایی آن، ضمن ایجاد فرصت‌هایی برای تنظیمگری مؤثرتر در مسائل جهانی، برای نظم حقوقی بین‌الملل چالش‌های جدی ایجاد می‌کند. این چالش‌ها مستلزم آن است که اتحادیه اروپایی در تدوین و اجرای قوانین با قابلیت اعمال فراسرزمینی خود، به طور جدی به اصول حقوق بین‌الملل و منافع دولت‌های ثالث توجه کرده و تلاش کند تا توازن مناسبی بین اهداف خود و احترام به حاکمیت سایر کشورها برقرار سازد.

۲-۴. پیامدهای سازوکار اثر بروکسل بر ایران

اتحادیه اروپایی با استفاده از اثر بروکسل، که از طریق آن توانایی اعمال استانداردهای خود در سطح جهانی را دارد، به تدریج مقررات هوش مصنوعی را در سراسر جهان شکل می‌دهد. مقررات سختگیرانه اتحادیه اروپایی، که بر استانداردهای جهانی تأثیر می‌گذارد، نه تنها نشان‌دهنده قدرت نرم فزاینده اتحادیه اروپایی در فناوری است، بلکه باعث شده است شرکت‌ها و دولت‌ها برای اجتناب از موانع ورود به بازار اتحادیه، به رعایت قوانین آن تمایل بیشتری داشته باشند. تأثیرات چندوجهی این موضوع بر کشورهای غیرعضو مانند ایران نیز شایان توجه است. شرکت‌های توسعه‌دهنده سیستم‌های هوش مصنوعی در ایران برای ورود به بازار اتحادیه اروپایی، به رعایت قانون هوش مصنوعی اتحادیه اروپایی نیازمند خواهند بود. این امر ممکن است به پذیرش ضمنی^۳ استانداردهای اتحادیه اروپایی (بدون هماهنگ‌سازی رسمی حقوقی)^۴ در ایران منجر شود. رعایت استانداردهای اتحادیه اروپایی می‌تواند هزینه‌ها را برای شرکت‌های ایرانی فعال در عرصه هوش مصنوعی افزایش دهد. با این حال، رعایت این استانداردها همچنین می‌تواند رقابت‌پذیری و جذابیت جهانی محصولات و خدمات هوش مصنوعی ایران را افزایش دهد و در نتیجه، بازارهای جدیدی را برای آن‌ها بگشاید. مقررات عمومی حفاظت از داده‌های اتحادیه اروپایی، پیش از این تأثیر بسزایی بر حفاظت از داده‌ها در سطح جهانی داشته است. مقررات مشابه در قانون هوش مصنوعی اتحادیه اروپایی، به‌ویژه در مورد استفاده از داده‌ها و شفافیت الگوریتمی، می‌تواند بر قوانین حفاظت از داده‌های ایران، به‌ویژه با توجه به اقتصاد دیجیتال رو به رشد این کشور، تأثیرگذار باشد.

۳. بررسی تحلیلی - انتقادی آورده‌های نوآورانه قانون هوش مصنوعی

قانون هوش مصنوعی مصوب پارلمان اروپایی دارای ۱۳ فصل، ۱۱۳ ماده و ۱۳ ضمیمه است. قانون هوش مصنوعی موضوعات مختلفی را در خود جای داده است؛ اما در اینجا به سه موضوع بنیادین آن پرداخته می‌شود. از این رو، در قسمت اول فرایند تعریف هوش مصنوعی بیان می‌شود، سپس به قلمرو قانون هوش مصنوعی و در نهایت به نوآوری این قانون در طبقه‌بندی سیستم‌های مبتنی بر هوش مصنوعی، یعنی رویکرد مبتنی بر خطر، پرداخته می‌شود.

1. Principle of Sovereignty.
2. Principle of Non- Intervention.
3. The de facto adoption.
4. formal legal harmonization.

۵. هماهنگ‌سازی غیررسمی حقوقی به همسودن و سازگاری متقابل قوانین، رویه‌ها و مفاهیم حقوقی میان حوزه‌های قضایی مختلف بدون وجود توافقات رسمی یا مقررات الزام‌آور

اشاره دارد.

۱-۳. تعریف هوش مصنوعی در قانون: فرایندی متغیر

بحث در مورد هوش مصنوعی همیشه با سؤالات زیادی همراه بوده است (Stark & Hutson, 2021, p. 933). جای تعجب نیست که فقدان یک تعریف دقیق و پذیرفته شده جهانی از هوش مصنوعی، احتمالاً به رشد، شکوفایی و پیشرفت این رشته با سرعتی فزاینده کمک کرده است (Stone et al., 2016, p. 12).

نکته اساسی در شناخت هوش مصنوعی، مقوله «هوشمندی»^۱ است. هوش مصنوعی فعالیتی است که به هوشمندسازی ماشین‌ها اختصاص می‌یابد و هوش کیفیتی است که موجودی را قادر می‌سازد به درستی و با آینده‌نگری در محیط خود عمل کند (Nilsson, 2009, p. 13).

یکی از پراهمیت‌ترین مسائل در قانون هوش مصنوعی، «تعریف» آن است. با توجه به تغییرات گسترده و پرسرعت در عرصه فناوری، این نگرانی وجود دارد که تعاریف ارائه شده نمی‌توانند در برابر سرعت پیشرفت فناوری هوش مصنوعی مقاومت کنند و در عرصه‌های جدید قابلیت اجرا ندارند؛ در نتیجه هر قانونی که در تلاش برای تنظیم‌گری و نظارت بر هوش مصنوعی است، در کمترین زمان ممکن، فایده عملی خود را از دست می‌دهد. این نگرانی تأثیر خود را در فرایند تدوین قانون مدنظر اتحادیه اروپایی درخصوص هوش مصنوعی نشان داده است. با مقایسه تعریف اولیه از هوش مصنوعی در نسخه پیشنهادی کمیسیون اروپایی^۲ و تعریف نهایی، که در ۱۳ مارچ تصویب شده است^۳، به تغییرات اساسی در عناصر تعریف هوش مصنوعی برمی‌خوریم. در طرح پیشنهادی اولیه کمیسیون اروپایی در ماده ۳ قسمت اول، هوش مصنوعی این‌گونه تعریف شده است: «سیستم هوش مصنوعی به معنای یک سیستم ماشین‌محور است که برای اقدام با سطوح مختلفی از خودمختاری طراحی شده است، که بعد از توسعه می‌تواند تطبیق‌پذیری را به نمایش بگذارد و از داده‌هایی که دریافت می‌کند، برای اهداف صریح یا ضمنی، استنباط کند که چگونه داده‌های خروجی مانند پیش‌بینی‌ها، محتوا، توصیه‌ها یا تصمیماتی که محیط مجازی یا فیزیکی را متأثر می‌کند، تولید کند.» اما در نسخه نهایی، ارکان تعریف هوش مصنوعی تغییرات اساسی را تجربه کرده است. در ماده ۳ قسمت یک قانون هوش مصنوعی، یک سیستم مبتنی بر هوش مصنوعی را به شرح زیر تعریف می‌کند: «سیستم هوش مصنوعی به معنای یک سیستم مبتنی بر ماشین است که برای عملکرد با سطوح مختلف استقلال و خودمختاری^۴ طراحی شده است، که ممکن است پس از استقرار نوعی سازگاری نشان دهد و برای اهداف صریح یا ضمنی و استنتاج از ورودی دریافتی‌اش، روی خروجی‌هایی مانند پیش‌بینی، محتوا، توصیه‌ها یا تصمیماتی که می‌تواند بر محیط‌های فیزیکی یا مجازی تأثیر بگذارد، مدیریت داشته باشد»^۵.

قانون هوش مصنوعی تأکید می‌کند که ویژگی کلیدی متمایزکننده سیستم‌های هوش مصنوعی از سیستم‌های نرم‌افزاری ساده‌تر و سنتی، توانایی آن‌ها در استنتاج است. تکنیک‌هایی که در حین طراحی یک سیستم مبتنی بر هوش مصنوعی استنتاج را امکان‌پذیر می‌کنند شامل

1. Intelligence

2. Article 3(1) AI Act(proposal): <https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=CELEX:52021PC0206>

3. Article 3(1) AI Act

4. Autonomy

۵. Article 3(1) AI Act؛ لازم به ذکر است که «هوش مصنوعی» در سند ملی هوش مصنوعی جمهوری اسلامی ایران مصوب شورای عالی انقلاب فرهنگی تعریف شده است که نخستین تعریف از یک سیستم مبتنی بر هوش مصنوعی در چهارچوب یک سند حقوقی در ایران است. بند اول ماده یک این سند، هوش مصنوعی را این‌طور تعریف کرده است: «به توانایی ماشین برای انجام عملکردهای خودکار و نظام‌مند از جمله یادگیری، درک، استنتاج، حل مسأله، پیش‌بینی، تصمیم‌گیری و اقدام از طریق به کارگیری دانش و اطلاعات و پردازش داده گفته می‌شود که منشأ اثرگذاری‌های گسترده بر انسان و روابط انسانی در محیط فیزیکی یا مجازی و همچنین بازتاب‌های زیست محیطی است. هوش مصنوعی ماهیتی داده‌ای، شبکه‌ای، الگوریتمی، خوشه‌ای، لایه‌ای و یکپارچه، مبتنی بر منطق‌های کلاسیک و سایر منطق‌های نوین دارد.» نکته حائز اهمیت در این تعریف آن است که به مؤلفه استنتاج و تصمیم‌گیری (که به طور ضمنی خودمختاری را هم شامل می‌شود) اشاره شده است. این امر در راستای جلوگیری از متروک شدن تلاش‌های تنظیم‌گری در چهارچوب حقوقی در اثر پیشرفت افسارگسیخته فناوری هوش مصنوعی است.

رویکردهای «یادگیری ماشینی»^۱ است که یاد می‌گیرند چگونه از اطلاعات به هدفی خاص دست یابند. ظرفیت یک سیستم هوش مصنوعی برای استنتاج فراتر از پردازش داده‌های اولیه، امکان یادگیری، استدلال یا مدل‌سازی است.^۲

با مقایسه آخرین نسخه تعریف هوش مصنوعی و نسخه پیشنهادی کمیسیون اروپایی، به ظهور عناصر جدید و کلیدی، یعنی «استنتاج»^۳ و «خودمختاری»^۴، در نسخه نهایی برمی‌خوریم که به‌وضوح یک سیستم هوش مصنوعی را از هر نرم‌افزار دیگری که خروجی آن از طریق یک الگوریتم از پیش تعیین شده است (اگر x آن‌گاه y) متمایز می‌کند (Fernhout & Duquin, 2024).^۵

تعریف ارائه شده در این قانون، به‌منظور آینده‌نگری و انطباق با تحولات سریع فناوری هوش مصنوعی، به‌صورت گسترده و فراگیر تدوین شده است. برخلاف تعاریف سنتی، که هوش مصنوعی را به فناوری‌ها و روش‌های مشخصی محدود می‌کنند، این تعریف رویکردی فراگیر و مستقل از فناوری را اتخاذ کرده است. این به آن علت است که قانون پیش رو در برابر تحولات بی‌وقفه هوش مصنوعی، به قانونی منسوخ تبدیل نشود.

۲-۳. قلمرو قانون هوش مصنوعی

قانون هوش مصنوعی اتحادیه اروپایی، با اتخاذ رویکردی فراگیر، قلمرو شمول خود را به‌گونه‌ای طراحی کرده است که حداکثر پوشش را روی سیستم‌های هوش مصنوعی تأثیرگذار در بازار اتحادیه داشته باشد. این قانون، با الهام از قوانین پیشین همچون مقررات عمومی حفاظت از داده، به‌دنبال وضع چهارچوبی قانونی جامع و مؤثر برای تنظیم توسعه و کاربرد هوش مصنوعی است.

ویژگی‌های بارز قلمرو شمول این قانون عبارت‌اند از:

۱. رویکرد فرامرزی: قانون نه تنها توسعه‌دهندگان^۶ و ارائه‌دهندگان^۷ مستقر در اتحادیه اروپایی، بلکه تمامی بازیگران کلیدی در زنجیره ارزش هوش مصنوعی را، که محصولات یا خدمات آن‌ها به‌نوعی در بازار اتحادیه عرضه می‌شود، تحت پوشش قرار می‌دهد. این رویکرد، اهمیت فزاینده هوش مصنوعی در اقتصاد جهانی و ضرورت تنظیم آن در سطح بین‌المللی را به‌خوبی نشان می‌دهد؛

۲. تمرکز بر تأثیرات: قانون به‌جای تمرکز صرف بر محل تأسیس یا اقامتگاه بازیگران، بر تأثیرات سیستم‌های هوش مصنوعی در بازار اتحادیه اروپایی تأکید دارد. این رویکرد، تضمین می‌کند که حتی اگر یک سیستم هوش مصنوعی در خارج از اتحادیه توسعه یافته باشد، اما تأثیراتش در داخل اتحادیه مشهود باشد، مشمول مقررات این قانون خواهد بود.

۳. شمول بازیگران مختلف: قانون طیف وسیعی از بازیگران از جمله توسعه‌دهندگان، ارائه‌دهندگان، واردکنندگان^۸، توزیع‌کنندگان^۹، تولیدکنندگان^{۱۰} و نمایندگان دارای مجوز^{۱۱} را شامل می‌شود. این امر نشان می‌دهد که قانون‌گذار اتحادیه اروپایی به پیچیدگی زنجیره ارزش هوش مصنوعی واقف است و کوشیده است تمامی حلقه‌های این زنجیره را تحت پوشش قرار دهد؛

۴. محافظت از افراد آسیب‌دیده^{۱۲}: قانون به‌طور خاص به حقوق افرادی که از اعمال سیستم‌های هوش مصنوعی آسیب می‌بینند توجه کرده است. این امر نشان می‌دهد که قانون‌گذار اتحادیه اروپایی، علاوه بر تنظیم صنعت، به‌دنبال حفاظت از حقوق شهروندان نیز هست؛

1. Machine Learning Approaches.

2. https://assets.ey.com/content/dam/ey-sites/ey-com/en_gl/topics/ai/ey-eu-ai-act-political-agreement-overview-february-2024.

3. Inference.

4. Autonomy.

۵. رجوع کنید به: <https://www.stibbe.com/publications-and-insights/the-eu-artificial-intelligence-act-our-16-key-takeaways>

6. Developer.

7. Provider.

8. Importer.

9. Distributor.

10. Productor.

11. Authorised Representatives.

12. Affected Person.

۵. انعطاف‌پذیری در مورد نرم‌افزارهای منبع باز: قانون در مورد سیستم‌های هوش مصنوعی، که تحت مجوزهای آزاد و منبع باز منتشر می‌شوند، رویکردی منعطف‌تر اتخاذ کرده است. با این حال، این سیستم‌ها نیز در صورتی که به عنوان سیستم‌های پرخطر طبقه‌بندی شوند، مشمول مقررات خاص خواهند بود^۱؛

قانون هوش مصنوعی اتحادیه اروپایی نقش فعالی به عرضه و ارائه در سرزمین اتحادیه اروپایی داده است و محل استقرار و تأسیس ارائه‌دهنده محل مانور نیست. به عبارت دیگر، به صرف این که سیستم‌های مبتنی بر هوش مصنوعی مستقل یا کالایی که همراه آن سیستم‌های هوش مصنوعی است ارائه شود، مشمول قانون هوش مصنوعی می‌شود و محل تأسیس و اقامتگاه آن نقشی نخواهند داشت. در واقع قانون با پذیرش اصل صلاحیت واقعی^۲، به این معنا که صرف آسیب دیدن افراد در اتحادیه اروپایی، فارغ از محل استقرار توسعه‌دهنده، امکان اعمال قانون را فراهم می‌کند، گامی مهم در راستای حمایت از حقوق شهروندان اتحادیه برداشته است. همچنین قانون با برقراری تعادل بین تشویق به نوآوری در حوزه هوش مصنوعی و حفاظت از حقوق شهروندان رویکردی متوازن اتخاذ کرده است.

از سوی دیگر، قانون هوش مصنوعی اتحادیه اروپایی ضمن تعریف قلمرو و شمول خود برای برخی موارد، به صراحت استثنا قائل شده است. این استثنائات با هدف ایجاد توازن بین تنظیم‌گری حقوقی فناوری هوش مصنوعی و حمایت از فعالیت‌های تحقیقی، توسعه و کاربردهای خاص آن در نظر گرفته است. از جمله این استثنائات می‌توان به فعالیت‌های نظامی و امنیتی^۳، فعالیت‌های تحقیقی و توسعه‌ای، کاربردهای شخصی و موارد تحت پوشش قوانین دیگر اشاره کرد^۴.

۳-۳. رویکرد مبتنی بر خطر

از زمان راه‌اندازی «استراتژی بازار واحد دیجیتال» اتحادیه اروپایی به طور فزاینده‌ای بر رویکرد مبتنی بر خطر^۵ تکیه کرده است (Dunn & De Gregorio, 2022, pp. 475-500). اتحادیه اروپایی به جای تعیین حقوق و پادمان‌های جدید، سعی کرده است با افزایش مسئولیت‌پذیری بازیگران دولتی و خصوصی نسبت به خطرات و اثرات جانبی احتمالی ناشی از فعالیت‌های آن‌ها، خطرهای ممکن را تنظیم کند. ظهور رویکرد مبتنی بر خطر در سیاست‌های دیجیتال اروپا، به ویژه با توجه به تحولات قانونی اخیر در زمینه داده‌ها، محتوای آنلاین و هوش مصنوعی مشهود است (Dunn & De Gregorio, 2022, p. 3).

مانند مقررات عمومی حفاظت از داده‌ها، قانون هوش مصنوعی نیز بر رویکرد مبتنی بر خطر استوار شده است (Gellert, 2021, p. 3). این رویکرد هم بر اساس اسناد راهبردی کمیسیون اروپایی در مورد هوش مصنوعی است که چهارچوب نظارتی آتی این فناوری را در اتحادیه اروپایی ترسیم می‌کند (EU Commission, 2020, p. 10) و هم بر اساس مصوبه‌ای است که پارلمان اروپایی تصویب کرده است. در قانون هوش مصنوعی، خطر به معنای «ترکیبی از احتمال وقوع آسیب و شدت آن آسیب» تعریف شده است^۶. باید توجه کرد مقررات‌گذاری مبتنی بر خطر انواع مختلفی دارد که رویکردشان یکسان نیست. یکی از انواع مقررات‌گذاری مبتنی بر خطر، تمرکز بر خطراتی است که ممکن است از یک حوزه خاص به جامعه، محیط‌زیست یا بهداشت عمومی برسد. در این مدل، خطراتی که ممکن است از یک حوزه ناشی شود به دلیلی برای تنظیم‌گری آن حوزه و چگونگی آن تبدیل می‌شود (Black, 2010, p. 187).

1. Article 2(1).

۲. Protective Jurisdiction: در برخی موارد، اعمال صلاحیت از سوی دولت‌ها، به دلیل رفتارهایی است که در خارج از سرزمین رخ داده است؛ اما آثار آن متوجه دولت است.

این صلاحیت، که تحت عنوان «نظریه اثر» نیز شناخته می‌شود، در واقع به گونه‌ای گسترش صلاحیت سرزمینی نیز هست (Beigzadeh, 2023, p. 615).

۳. بسیاری نگران هستند که عدم ممنوعیت کامل نظارت گسترده عمومی از طریق هوش مصنوعی به دولت‌ها اجازه می‌دهد حریم خصوصی مردم را نقض کرده و از هوش مصنوعی برای نظارت بر آن‌ها استفاده کنند. همچنین، این تصمیم ممکن است به سایر کشورها نیز این پیام را بدهد که نظارت گسترده پذیرفتنی است.

4. Article 2(3).

5. Risk-Based Approach

6. Article 3(2) AI Act.

این نوع مقررات‌گذاری، که براساس ارزیابی خطر صورت می‌گیرد، همان روشی است که در قانون هوش مصنوعی اتحادیه اروپا به کار رفته است. رویکرد مبتنی بر خطر به عنوان روشی ظاهراً دقیق و منطقی معرفی می‌شود که قوانین را با توجه به شدت و وسعت خطراتی تنظیم می‌کند که سیستم‌های هوش مصنوعی می‌توانند ایجاد کنند (Mahler, 2021, p. 3).

قانون هوش مصنوعی اتحادیه اروپایی با هدف اتخاذ رویکردی هدفمند و متناسب برای مدیریت خطرات ناشی از سیستم‌های هوش مصنوعی، روش‌ها و سیستم‌های مبتنی بر هوش مصنوعی را بسته به موارد استفاده به چهار دسته مختلف طبقه‌بندی می‌کند:

نخست «سیستم‌های ممنوعه هوش مصنوعی»^۱، سیستم‌هایی که خطر ناپذیرفتنی ایجاد می‌کنند، با ارزش‌های اتحادیه اروپایی مغایرت دارند و تهدیدی آشکار برای حقوق اساسی تلقی می‌شوند در این قانون ممنوع خواهند شد؛ زیرا بیش از حد خطرناک تلقی می‌شوند.^۲ سیستم‌های هوش مصنوعی، که خطر نامقبولی دارند، مشمول بیشترین مجازات‌ها هستند و ممکن است شرکت‌ها را در معرض جریمه‌های ۳۵ میلیون یورویی یا ۷ درصد از درآمد سالانه یک شرکت قرار دهند؛ هر کدام بیشتر باشد.^۳ این ممنوعیت‌ها شامل سیستم‌هایی می‌شود که از ویژگی‌های حساس افراد (زیست‌سنجی)^۴ برای طبقه‌بندی آن‌ها استفاده می‌کنند (مانند سیستم‌هایی که افراد را براساس نژاد، دین یا گرایش جنسی دسته‌بندی می‌کنند)، سیستم‌هایی که پایگاه‌های داده تشخیص چهره^۵ ایجاد می‌کنند، سیستم‌هایی که احساسات افراد را در محیط‌های کاری و آموزشی تشخیص می‌دهند، سیستم‌هایی که به افراد براساس رفتار یا ویژگی‌های شخصی امتیاز می‌دهند^۶ و سیستم‌هایی که برای دستکاری رفتار انسان و پیش‌بینی جرم طراحی شده‌اند. همچنین استفاده از سیستم‌های شناسایی زیست‌سنجی بلادرنگ^۷ برای اهداف اجرای قانون، با برخی استثنائات، ممنوع شده است. هدف از این محدودیت‌ها جلوگیری از سوءاستفاده از هوش مصنوعی و تضمین این موضوع است که این فناوری به نفع همه انسان‌ها باشد.^۸

گروه دوم «سیستم‌های پرخطر»^۹ است^{۱۰} که بیشتر آن‌ها با فهرستی که در ضمیمه سه موجود است شناسایی می‌شوند و توسط کمیسیون اروپایی براساس طیف وسیعی از معیارهای تعیین‌شده اصلاح می‌شوند.^{۱۱} سیستم‌های هوش مصنوعی پرخطر باید با مجموعه‌ای از الزامات طولانی و گسترده مطابقت داشته باشند. در واقع، ارائه‌دهندگان و کاربران آن سیستم‌ها باید یک سیستم مدیریت خطر را با هدف اتخاذ تدابیر مناسب برای رویارویی با هر خطر شناخته شده یا قابل پیش‌بینی ایجاد، پیاده‌سازی، مستندسازی و نگهداری کنند.^{۱۲} علاوه بر این، ارائه‌دهندگان سیستم‌های هوش مصنوعی با خطر بسیار ملزم به ایجاد یک سیستم مدیریت کیفیت برای اطمینان از انطباق با کل مقررات هستند.^{۱۳} این دسته بسیار گسترده‌تر از دسته «سیستم‌های ممنوعه» است. سیستم‌های هوش مصنوعی با آسیب بالقوه درخور توجه به سلامت، ایمنی، حقوق اساسی، محیط‌زیست، دموکراسی و حاکمیت قانون به عنوان پرخطر طبقه‌بندی می‌شوند.^{۱۴} نکته حائز اهمیت این است که سیستم‌های هوش مصنوعی

1. Prohibited AI System.

2. Article 5 AI Act.

3. Article 99 AI Act.

۴. Biometric Categorisation Systems: سیستم‌های طبقه‌بندی زیست‌سنجی از ویژگی‌های فیزیولوژیکی یا رفتاری یک فرد برای احراز هویت و تشخیص افراد استفاده می‌کنند.

این سیستم‌ها براساس ویژگی‌هایی مانند اثر انگشت، تصویر چهره، ساختار عنبیه چشم، اثر صدا، الگوی نوشتار و ... افراد را تشخیص می‌دهند. اطلاعات بیولوژیکی هر فرد به عنوان الگویی منحصر به فرد برای او شناخته می‌شود و در سیستم‌های طبقه‌بندی زیست‌سنجی بهره‌برداری از این الگوها برای شناسایی و احراز هویت افراد انجام می‌شود.

5. Facial Recognition Databases

6. Social scoring

7. Real-time Biometric Identification

۸. Article 5(1); لازم به ذکر است چنین داده‌هایی در دسته داده‌های شخصی قرار می‌گیرند که اتحادیه اروپایی در قوانین مختلفی از جمله مقررات عمومی حفاظت از داده‌های

شخصی، حمایت کرده است تا از حریم و زندگی خصوصی افراد حمایت شود (Dashti And Motamednejad, 2024, p. 13)

9. High-risk

10. Article 6 AI Act

11. Article 7 AI Act

12. Article 9 AI Act

13. Article 17 AI Act

۱۴. نمونه‌هایی از این سیستم‌ها عبارتند از: سیستم‌هایی که در زیرساخت‌های حیاتی مانند آب، گاز و برق استفاده می‌شوند، سیستم‌هایی که در اجرای عدالت و فرایندهای دموکراتیک

نقش دارند، سیستم‌هایی که در آموزش و اشتغال کاربرد دارند، سیستم‌های شناسایی زیستی و سیستم‌هایی که برای مدیریت مهاجرت و مرزها استفاده می‌شوند (Annex III).

پرخاطر ممنوع نیستند؛ اما به رعایت تعهدات سخت‌گیرانه نیاز دارند. در صورت استقرار یا استفاده از یک سیستم هوش مصنوعی پرخاطر، قانون هوش مصنوعی تعهداتی از قبیل ارزیابی خطر، آموزش و پشتیبانی و نگهداری گزارش‌ها را تحمیل می‌کند.^۱

سومین دسته از سیستم‌های هوش مصنوعی، سیستم‌هایی هستند که خطر محدودی دارند.^۲ این دسته شامل سیستم‌هایی می‌شود که برای تعامل با انسان‌ها طراحی شده‌اند؛ مانند چت‌بات‌ها و تولیدکنندگان محتوای جعلی (دیپ‌فیک)^۳ مشمول حداقل الزامات شفافیت هستند. این بدان معناست که افراد باید آگاه باشند که در حال تعامل با یک سیستم هوش مصنوعی هستند. هدف از این الزام، جلوگیری از فریب خوردن افراد و اطمینان از آن است که افراد بدانند با چه کسی در ارتباط هستند. براساس منطق قانون هوش مصنوعی، با توجه به تعهدات محدودتر این دسته از سیستم‌ها، می‌توان نتیجه گرفت که سطح خطر آن‌ها به نسبت سیستم‌های پرخاطر کمتر است. این بدان معناست که اگرچه این سیستم‌ها نیز خطراتی را به همراه دارند؛ اما به اندازه خطرات ناشی از سیستم‌های پرخاطر جدی نیستند.

در نهایت تمام سیستم‌های هوش مصنوعی دیگر، که تحت یکی از سه دسته خطر اصلی قرار نمی‌گیرند، مانند سیستم‌های توصیه‌کننده فعال با هوش مصنوعی یا فیلترهای هرزنامه، به مثابه حداقل/ بدون خطر^۴ طبقه‌بندی می‌شوند. سیستم‌های هوش مصنوعی، که در هیچ‌کدام از سه دسته مذکور قرار نمی‌گیرند، مشمول هیچ وظیفه یا تعهد خاصی نیستند؛ اگرچه کمیسیون و کشورهای عضو باید تدوین آیین‌نامه‌های رفتاری را با هدف تقویت اعمال داوطلبانه الزامات تعیین شده برای سیستم‌های پرخاطر، تشویق و تسهیل کنند.^۵

قانون هوش مصنوعی اتحادیه اروپایی با اتخاذ رویکرد مبتنی بر خطر، دیدگاه «لومان»^۶ را تأیید می‌کند که معتقد است نمی‌توان مفاهیم خطر و ایمنی را کاملاً از هم جدا کرد. طبق نظر لومان، دستیابی کامل به ایمنی غیرممکن است.^۷ با توجه به این که دامنه خطرات اکنون شامل حقوق بنیادین نیز شده است، مفهوم ایمنی نه تنها مهم تلقی می‌شود، بلکه در کنار ارزش‌هایی مانند برابری و خودمختاری قرار گرفته و ممکن است تحت تأثیر برخی سیستم‌های هوش مصنوعی قرار گیرد (Luhmann, 2017, p. 23).

قانون هوش مصنوعی اتحادیه اروپایی با پذیرش دیدگاه لومان، به این نتیجه می‌رسد که ایمنی ارزشی است که در کنار سایر ارزش‌ها باید مورد توجه قرار گیرد. گسترش دامنه خطرات به حقوق بنیادین نشان می‌دهد که ایمنی نه تنها مسئله‌ای فنی، بلکه مسئله‌ای اخلاقی و اجتماعی نیز هست. رویکرد مبتنی بر ریسک در این قانون، به دنبال برقراری تعادل بین ترویج نوآوری در حوزه هوش مصنوعی و محافظت از حقوق افراد است.

نتیجه‌گیری

۱. قانون هوش مصنوعی اتحادیه اروپایی اولین چهارچوب قانونی جامع هوش مصنوعی در سراسر جهان است. هدف این قانون تقویت هوش مصنوعی قابل اعتماد در اروپا و اطمینان از این است که سیستم‌های هوش مصنوعی به حقوق اساسی و اصول اخلاقی احترام می‌گذارند؛

۲. این قانون تضمین می‌کند که اروپایی‌ها می‌توانند به آنچه هوش مصنوعی ارائه می‌دهد اعتماد کنند؛ زیرا بیشتر سیستم‌های هوش مصنوعی نه تنها هیچ خطری ندارند، بلکه به حل بسیاری از چالش‌های اجتماعی کمک می‌کنند. ناگفته نماند برخی از سیستم‌های هوش مصنوعی خطراتی را ایجاد می‌کنند که باید برای جلوگیری از نتایج نامطلوب به آن‌ها رسیدگی کرد؛

1. Article 27 AI Act

2. Limited-risk.

۳. Deep Fake؛ دیپ‌فیک عبارت است از یک نرم‌افزار یا سیستم هوش مصنوعی که از تکنیک‌های یادگیری عمیق (Deep learning) برای تولید تصاویر یا ویدئوهای جعلی استفاده می‌کند. این سیستم‌ها معمولاً از شبکه‌های مولد تصادفی (GAN) برای تولید تصاویر و ویدیوهایی که به نظر واقعی می‌رسند استفاده می‌کنند. استفاده از Deepfake generator ممکن است به ایجاد تصاویر جعلی از افراد، ویدئوهای تقلبی و اخباری منجر شود که ممکن است برای تبلیغات تقلبی، تبلیغات سیاسی، یا فریب مخاطبان به کار رود.

4. Minimal/no-risk

5. Article 95 AI Act

6. Luhmann

۷. لومان، جامعه‌شناس آلمانی، معتقد بود که مفاهیم ریسک و ایمنی به‌طور جدایی‌ناپذیری با یکدیگر مرتبط هستند و نمی‌توان آن‌ها را به‌طور کامل از هم جدا کرد. او بر این باور بود که ایمنی هدفی مطلق نیست و همیشه در معرض تهدید قرار دارد.

۳. فرایند تصویب این قانون در سال ۲۰۱۷ با تلاش‌های کمیسیون اروپایی شروع شد و در دسامبر ۲۰۲۳ قانون‌گذاران اتحادیه اروپایی به توافق سیاسی دست یافتند. در تاریخ ۱۴ ژوئن ۲۰۲۳ پارلمان اروپایی اصلاحاتی را بر پیش‌نویس کمیسیون اروپایی اعمال کرد. این اصلاحات در تاریخ ۱۳ مارس ۲۰۲۴ نهایی شد و در تاریخ ۲۱ می ۲۰۲۴ به تأیید شورای اروپایی رسید. این فرایند نشان‌دهنده حرکتی جدید به سمت تنظیم چهارچوب نظارتی برای هوش مصنوعی است که به تعادل عادلانه پیشرفت فناوری و ملاحظات اخلاقی دست می‌یابد؛

۴. با تصویب این قانون، اتحادیه اروپایی به دنبال حکمرانی جهانی در زمینه هوش مصنوعی از طریق اثر بروکسل است. اثر بروکسل فرایندی است که طی آن اتحادیه اروپایی به صورت یک‌جانبه بر بازار جهانی نظارت می‌کند و قوانین و مقررات خود را بر جامعه بین‌المللی اعمال می‌کند. جذابیت بازار ۴۵۰ میلیونی، سابقه تصویب قوانین مشابه و وجود چند دهه تجربه و تخصص و از بین بردن نگرانی‌ها در مورد دور زدن قانون از طریق ارائه‌دهندگان سیستم‌های هوش مصنوعی، دلایلی است که به قابلیت اعمال اثر بروکسل کمک می‌کند؛ البته اتحادیه اروپایی برای حکمرانی جهانی در زمینه هوش مصنوعی با موانعی روبه‌روست؛ از جمله عدم قاطعیت ادعایی قانون هوش مصنوعی، فقدان مقرراتی در مورد تقسیم‌ناپذیری سیستم‌های هوش مصنوعی و فقدان حمایت کافی قانون هوش مصنوعی از ارزش‌های اتحادیه اروپایی؛ البته این فرایند جهانی‌سازی می‌تواند چالش‌های حقوقی همچون تعارض با حاکمیت کشورها و اصل عدم مداخله به وجود آورد و مشکلاتی برای شرکت‌های فعال در عرصه هوش مصنوعی از جمله شرکت‌های ایرانی به وجود آورد.

۵. تعریف هوش مصنوعی در فرایند تدوین قانون هوش مصنوعی به دلیل پیشرفت‌های این فناوری متحول شده است. در تعریف اخیر، عناصری مثل خودمختاری و استنتاج اضافه شده است. همچنین این تعریف دارای مصادیق گسترده‌ای است تا بتواند در برابر پیشرفت این فناوری مقاومت کند؛

۶. قانون هوش مصنوعی بر ارائه‌دهندگان و توسعه‌دهندگانی که در اتحادیه اروپایی مستقرند یا خروجی محصولاتشان در اتحادیه اروپایی استفاده می‌شود، واردکنندگان و توزیع‌کنندگان سیستم‌های هوش مصنوعی، نمایندگان دارای صلاحیت آنان و افراد آسیب‌دیده، که در اتحادیه اروپایی مستقرند، اعمال می‌شود. همچنین این قانون واجد استثنائاتی از قبیل اهداف نظامی و امنیتی، آموزشی و تحقیقی و آزمایشگاهی است. این قانون با پذیرش صلاحیت واقعی، گامی بنیادین برای حمایت از شهروندان اتحادیه اروپایی برداشته است؛

۷. در نهایت قانون هوش مصنوعی، سیستم‌های هوش مصنوعی را با توجه به درجه خطر، به چهار دسته سیستم‌های ممنوعه، سیستم‌های با خطر بالا، سیستم‌های با خطر محدود و سیستم‌هایی که در هیچ‌کدام از این دسته‌ها قرار نمی‌گیرند، طبقه‌بندی کرده است. سیستم‌های دسته اول ممنوع هستند و اعمال آن مجازات به همراه خواهد داشت. دسته دوم ممنوع نیستند، اما باید در موردشان ارزیابی‌های نظارتی صورت گیرد. دسته سوم سیستم‌هایی هستند که درجه نظارتی‌شان خفیف است و دسته چهارم که مقررات خاصی در مورد آن‌ها وجود ندارد. باید دید بعد از اجرایی شدن این قانون و حکمرانی آن بر فناوری افسارگسیخته هوش مصنوعی، آیا نظارت حقوقی بر این فناوری محقق می‌شود یا این که سرعت پیشرفت این فناوری مجال تحقق هرگونه نظارت انسانی در قالب مقررات‌گذاری را نمی‌دهد.

منابع

- بیگزاده، ابراهیم (۱۴۰۱). حقوق بین الملل. میزان
- دشتی، طاهره سادات، و محمدرزاد، رویا (۱۴۰۳). جایگاه هوش مصنوعی در قانونگذاری اتحادیه اروپا. علوم خبری، ۱۳(۱)، ۲۰-۱.
- سند ملی هوش مصنوعی جمهوری اسلامی ایران (۱۴۰۳). مصوب شورای عالی انقلاب فرهنگی. دسترسی در: [لینک](#)
- فرهمند، فرید (۱۴۰۳). طرح حمایت و حفاظت از داده و اطلاعات شخصی با مطالعه تطبیقی مقررات عمومی حفاظت از داده. حقوق فناوری های نوین، ۵(۹)، ۱۷-۲۵. <https://doi.org/10.22133/mtlj.2023.383528.1164>
- Almada, M., & Radu, A. (2024). The Brussels side-effect: how the AI act can reduce the global reach of EU policy. *German Law Journal*, 25(4), 646–663. <https://doi.org/10.1017/glj.2023.108>
- Bernat, M., & Spera, F. (2024). The De Jure “Brussels Effect”: Current Legal Trends in the EU’s Unilateral Regulatory Globalisation. *Studia Europejskie-Studies in European Affairs*, 28(4), 77–105. <https://doi.org/10.33067/SE.4.2024.4>
- Black, J. (2010). *Risk-based regulation: Choices, practices and lessons being learnt*. [Link](#)
- Bradford, A. (2012). *The Brussels effect: How the European Union rules the world*. Oxford University Press.
- Clyde, A. (2021, March 20). *Human-in-the-Loop Systems Are No Panacea for AI Accountability*. Tech Policy Press. [Link](#)
- Dunn, P., & De Gregorio, G. (2022). The Ambiguous Risk-Based Approach of the Artificial Intelligence Act: Links and Discrepancies with Other Union Strategies. In *IAIL@ HHAI*.
- Eckart, A. N. (2022). Transactional Artificial Intelligence. *Legal Writing: The Journal of the Legal Writing Institute*, 26, 273. [Link](#)
- Enqvist, L. (2023). ‘Human oversight’ in the EU artificial intelligence act: what, when and by whom?. *Law, Innovation and Technology*, 15(2), 508-535. <https://doi.org/10.1080/17579961.2023.2245683>
- European Commission. (2019). *High-level expert group on artificial intelligence*. (visited:21/03/2024). Retrieved from [Link](#)
- European Commission. (2021). Europe fit for the Digital Age. (visited:22/03/2024) at: [Link](#)
- Fernhout, F., & Duquin, T. (2024). The EU Artificial Intelligence Act: Our 16 Key Takeaways. *Stibbe*. [Link](#)
- Gellert, R. M. (2021). The role of the risk-based approach in the General data protection Regulation and in the European Commission’s proposed Artificial Intelligence Act. Business as usual?. <https://doi.org/10.14658/pupj-jelt-2021-2-2>
- Hadjiyianni, I. (2021). The European Union as a global regulatory power. *Oxford Journal of Legal Studies*, 41(1), 243-264. <https://doi.org/10.1093/ojls/gqaa042>
- Luhmann, N. (2018). *Trust and power*. John Wiley & Sons, New York.
- Mahler, T. (2021). Between risk management and proportionality: The risk-based approach in the EU’s Artificial Intelligence Act Proposal. *Nordic Yearbook of Law and Informatics*. [Link](#)
- Siegmann, C., & Anderljung, M. (2022). The Brussels effect and artificial intelligence: How EU regulation will impact the global AI market. *arXiv preprint arXiv:2208.12645*. [Link](#)

- Stark, L., & Hutson, J. (2021). Physiognomic artificial intelligence. *Fordham Intell. Prop. Media & Ent. LJ*, 32, 922. [Link](#)
- Stone, P., Brooks, R., Brynjolfsson, E., Calo, R., Etzioni, O., Hager, G., ... & Teller, A. (2022). Artificial intelligence and life in 2030: the one hundred year study on artificial intelligence. *arXiv preprint arXiv:2211.06318*. [Link](#)
- EU Commission. (2017). *Mid-Term Review on the implementation of the Digital Single Market Strategy A Connected Digital Single Market for All* (COM/2017/0228 final). European Commission. [Link](#)
- EU Commission. (2018). Artificial Intelligence for Europe (COM/2018/0237 final). European Commission. [Link](#)
- EU Commission. (2018). Coordinated Plan on Artificial Intelligence (COM/2018/0795 final). European Commission. [Link](#)
- EU Commission. (2020). White Paper on Artificial Intelligence A European Approach to Excellence and Trust (COM/2020/0065 final). European Commission. [Link](#)
- EU Commission. (2019). Building Trust in Human-Centric Artificial Intelligence (COM/2019/0168 final). European Commission. [Link](#)
- Council of the European Union. (2017). *Proposal for a Directive of the European Parliament and of the Council on copyright in the Digital Single Market Revised Presidency compromise proposal on Articles 2 and 13 and relevant recitals* (Document ST 14 2017 INIT). Retrieved from [Link](#)

[In persian]

- Beigzadeh, E. (2023). *International Law*. Mizan Publications.
- Dashti, T. Motanednejad, R. (2024). The role of artificial intelligence in EU legislation. *News Science*, 11(3). 1-20. <https://doi.org/10.22034/lrsi.2024.454812.1175>
- Farahmand, F. (2024). The Bill of Data Protection and GDPR 2018: A Comparative Study. *Modern Technologies Law*, 5(9), 17-25. <http://doi.org/10.22133/mtlj.2023.383528.1164>.
- National Artificial Intelligence Document of the Islamic Republic of Iran, approved by the Supreme Council of the Cultural Revolution, accessible at: [Link](#)

COPYRIGHTS

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

