

The EU's Strategic Autonomy and Good Artificial Intelligence

Mostafa Kaka 

Master Student, International Relations, University of Mazandaran, Babolsar, Iran

Mokhtar Salehi * 

Assistant Professor, International Relations, Department of Political Sciences and International Relations, University of Mazandaran, Babolsar, Iran

Introduction

Artificial intelligence (AI) can be considered one of the key indicators and outcomes of the Fifth Industrial Revolution. This transformation has led many countries to consider adopting appropriate strategies to manage its consequences. The present study aimed to answer the following question: Why has the European Union (EU) pursued policymaking in the field of artificial intelligence? The study is based on the hypothesis that the EU seeks to establish appropriate and sustainable laws and policies for AI in order to manage its consequences, while also creating conditions to achieve reliable and trustworthy AI in line with its strategic autonomy.

Literature Review

Previous studies have addressed various aspects of AI and its impact on security and governance. For example, in the report titled *Artificial Intelligence and Life in 2030*, Stone et al. (2016) emphasized the role

* Corresponding Author: mo.salehi@umz.ac.ir

How to Cite: Kaka, M & Salehi, M., (2025), "The EU's Strategic Autonomy and Good Artificial Intelligence", *Political Strategic Studies*, Vol. 14, No. 54, 313-349. Doi: [10.22054/QPSS.2024.80662.3434](https://doi.org/10.22054/QPSS.2024.80662.3434).

of artificial intelligence in enhancing strategic autonomy and its applications across sectors such as defense, agriculture, and healthcare. Furthermore, the 2022 report by the Centre for European Economic Policy (CEEP) analyzed the EU's strategies in the field of AI and its impact on global competitiveness. Building on these studies, the present research employed existing theoretical frameworks and new data to provide a more detailed examination of the topic.

Materials and Methods

The current study focused on analyzing and identifying the main themes within relevant texts. The data was collected through both library and online sources, including official European Union documents, research papers, and reports related to AI and public policy. The method of thematic analysis method was used to analyze the data.

Results and Discussion

The thematic analysis was applied to examine European Union documents, policies, and regulations related to AI. This resulted in the identification of key themes that reflect the EU's strategy toward achieving strategic autonomy and developing trustworthy AI. The findings indicate that the EU aims to strengthen its digital sovereignty by emphasizing factors such as reducing dependency on global actors, ensuring independent decision-making in critical areas, implementing comprehensive regulations, and protecting citizens' data and privacy. The identified themes were organized into seven major categories: 1) strategic autonomy, 2) AI governance, 3) data transparency and security, 4) risk management and technological ethics, 5) public trust, 6) human oversight, and 7) security and military applications of AI. Each of theme was supported by various open codes, reflecting the

EU's focus on balancing technological innovation with human rights principles. Moreover, the EU tends to classify AI systems according to risk levels and impose strict requirements on high-risk technologies, thus seeking to establish a legal, ethical, and accountable framework that ensures both societal security and public acceptance of AI. In addition, relying on comprehensive policies on transparency, accountability, and system explainability, the EU is actively working to strengthen public trust in this technology. According to the findings, in response to technological transformations and the rise of AI, the EU has assumed a proactive role in regulation and policymaking. Leveraging its legal and institutional capacities, the EU seeks to establish a framework for the safe, ethical, and trustworthy use of AI—one that addresses internal concerns related to human rights, security, and public trust, while also serving as a viable model for other countries. The analysis also focused on the concept of good AI, which translates the EU's core values into the digital domain. By emphasizing principles such as transparency, human oversight, privacy, and accountability, the EU seeks to align technological development with human dignity and democratic values. This approach, in contrast to purely technocratic views, demonstrates that the future of AI depends not only on technical progress but also heavily on governance and policy choices. Moreover, the thematic analysis indicated that the EU, as a normative power, seeks to extend its influence in the global technology market by implementing strict regulations, thereby encouraging other international actors to align with these standards. This approach strengthens the EU's regulatory power on the global stage and gradually positions it as a global authority in the governance of emerging technologies such as AI. Overall, the current analysis demonstrated that the EU's AI policy is not only designed to prevent potential threats but also to embody a forward-looking and normative approach that integrates innovation with ethics, and security with

fundamental human freedoms. While this path comes with challenges, it can serve as a model for responsible governance in the digital transformation era.

Conclusion

The results underscored the EU's commitment to establishing a regulatory framework that balances technological advancement with ethical considerations and public safety. By focusing on transparency, data protection, and human oversight, the EU is setting a global example for AI governance that aligns with democratic values and human rights. However, challenges remain, particularly in the rapid evolution of AI technologies, international competition, and the need for cross-border collaboration to address global AI challenges. Future research should explore the effectiveness of the EU's regulatory policies in practice, assess their impact on innovation, and examine how these policies can adapt to emerging technological developments.

Keywords: Good AI, Trustworthy AI, European Strategic Autonomy, AI Security, AI Governance



----- فصلنامه علمی پژوهش‌های راهبردی سیاست -----

دوره ۱۴، شماره ۵۴، پاییز ۱۴۰۴، ۳۴۹-۳۱۳

qpss.atu.ac.ir

DOI: [10.22054/QPSS.2024.80662.3434](https://doi.org/10.22054/QPSS.2024.80662.3434)

راهبرد استقلال استراتژیک اتحادیه اروپا و هوش مصنوعی خوب

دانشجوی کارشناسی ارشد روابط بین‌الملل دانشگاه مازندران، بابلسر، ایران

مصطفی کا کا

استادیار روابط بین‌الملل، گروه علوم سیاسی و روابط بین‌الملل دانشگاه مازندران،

مختار صالحی *

بابلسر، ایران

چکیده

هوش مصنوعی را می‌توان از شاخص‌ها و نتایج مهم انقلاب صنعتی پنجم دانست. تحولی که بسیاری از کشورها را به فکر اتخاذ یک راهبرد مناسب برای مدیریت عواقب آن انداخته است. هدف مقاله پاسخ به این پرسش است که چرا اتحادیه اروپا اقدام به سیاستگذاری در حوزه فناوری هوش مصنوعی نموده است؟ بنا بر فرضیه، اتحادیه اروپا بنا دارد قوانین و سیاستگذاری مناسب و پایداری در حوزه هوش مصنوعی وضع نماید تا ضمن مدیریت عواقب آن، شرایط دستیابی به هوش مصنوعی خوب و قابل اعتماد را در راستای استقلال استراتژیک خود فراهم نماید. فرضیه فوق از دریچه نظریه سیاستگذاری عمومی و تنظیم‌گری به بحث گذاشته می‌شود. روش جمع‌آوری اطلاعات در قالب استفاده از منابع کتابخانه‌ای و اینترنتی است. سپس اطلاعات جمع‌آوری شده به روش کیفی در قالب روش تحلیل مضمون مورد تجزیه و تحلیل قرار گرفته است. یافته‌های تحقیق نشان می‌دهد اتحادیه اروپا در تلاش است با وضع قوانین و دستورالعمل‌ها، ریسک پذیرش تکنولوژی جدید را به حداقل رسانده و الزامات امنیتی و حکمرانی را متناسب با هوش مصنوعی بازتنظیم نماید.

واژگان کلیدی: هوش مصنوعی خوب، هوش مصنوعی قابل اعتماد، استقلال استراتژیک اروپا،

امنیت هوش مصنوعی، حکمرانی هوش مصنوعی.

مقدمه

فناوری‌ها همیشه علاوه بر ایجاد فرصت‌های جدید، چالش‌هایی را ایجاد می‌کنند که تنظیم‌کننده‌ها را ملزم می‌کنند تعادلی بین نوآوری‌ها و کاهش خطرات برقرار کنند. در طول تاریخ، فناوری‌ها برای دستیابی و خدمت به اهداف مختلف مورد استفاده قرار گرفته‌اند و از یک سو مراحل جدیدی را برای رشد اجتماعی و از سوی دیگر زیرسوال بردن وضعیت موجود فراهم کرده‌اند. فناوری‌های دیجیتال و به شکل خاص هوش مصنوعی نیز از این قاعده مستثنی نیستند. این پژوهش تلاش می‌کند هوش مصنوعی را در اتحادیه اروپا مورد بررسی قرار دهد. از این جهت که هوش مصنوعی در اتحادیه اروپا در چه وضعیتی قرار دارد؛ چه تصمیماتی در این اتحادیه برای توسعه و پیشرفت هوش مصنوعی گرفته شده است. همچنین تلاش می‌شود با بررسی قوانین هوش مصنوعی در اتحادیه اروپا، آینده اخلاقی این فناوری را در جامعه اروپایی و راهبردهای امنیتی این اتحادیه مورد بررسی قرار گیرد. پرسش اصلی این است که چرا اتحادیه اروپا اقدام به سیاست‌گذاری در حوزه فناوری هوش مصنوعی نموده است؟ ضمن اینکه به این پرسش تکمیلی نیز پرداخته شده است که اتحادیه اروپا چگونه نسبت به مدیریت عواقب هوش مصنوعی اقدام کرده است؟

ظهور نسل جدید تکنولوژی‌های دیجیتال و هوش مصنوعی و تاثیرگذاری آن بر جوامع، انسان‌ها، سیاست و روابط میان کشورها این ضرورت را بیان می‌کند که می‌بایست تکنولوژی‌های جدید را بیشتر شناخته و همچنین به سیاست‌گذاری دیگر بازیگران و کشورها توجه شود. سیاست‌هایی که می‌توانند شرایط جوامع انسانی، کسب و کار و امنیت را مورد تاثیر قرار دهند و حتی شیوه تعامل دیگران با آنها را مطابق دستورکارهای جدید بازتعریف نمایند. از این رو در این پژوهش تلاش می‌شود در دو بخش شاخص‌ها و معیارهای هوش مصنوعی خوب، امنیت اتحادیه اروپا، قوانین، ریسک‌ها، مسائل تجاری، امنیت و حکمرانی تحت تاثیر فناوری هوش مصنوعی به بحث گذاشته شوند.

در این مقاله برای تحلیل موضوع استقلال استراتژیک اتحادیه اروپا در حوزه هوش مصنوعی از روش تحلیل مضمون استفاده شده است. بکارگیری روش فوق، این امکان را

داده است تا از داده‌ها و منابع مختلف، مضامین اصلی مرتبط با این حوزه شناسایی و سازماندهی شوند و در نهایت با ایجاد کدهای مناسب باز و محوری موضوع مورد تجزیه و تحلیل قرار گیرد.

ضرورت و پیشینه پژوهش

با توجه به اینکه این موضوع در طی دو سال اخیر بیشتر مورد توجه قرار گرفته است و یک موضوع بسیار جدید است لذا منابع نیز متناسب با آن منتشر شده است. در این نوشتار تلاش می‌شود به برخی از مهمترین منابع که می‌توانند در این تحقیق کمک کنند اشاره شود و ضمن اینکه تمایز موضوع فوق با موضوعات صورت گرفته در این است که این موضوع در اتحادیه اروپا نیز موضوعی جدید است و تاکنون منابع زیادی مستقلاً به این موضوع نپرداخته‌اند. ضمن اینکه این پژوهش می‌تواند به عنوان راهنمایی برای الگوبرداری از تجربه دیگر کشورها در سطح داخل مورد استفاده قرار گیرد. در ادامه به برخی از مهمترین منابع اشاره می‌شود.

استیکس^۱ در سال ۲۰۲۳ در رساله دکتری خود با عنوان «به سوی هوش مصنوعی ایمن، اخلاقی و سودمند در اتحادیه اروپا و فراتر از آن»^۲ در خصوص آینده هوش مصنوعی و ضرورت ایجاد قوانین و مقررات در خصوص آینده و رشد و توسعه هوش مصنوعی و در نظرگیری شرایط انسانی و ایجاد یک اعتماد عمومی در جهت آینده هوش مصنوعی و نداشتن خطرات متقابل برای زندگی انسان اشاره می‌کند. شارلوت^۳ و مارکوس^۴ در سال ۲۰۲۲ در مقاله خود با عنوان «اثر بروکسل و هوش مصنوعی: چگونه مقررات اتحادیه اروپا بر بازار جهانی هوش مصنوعی تأثیر می‌گذارد»^۵ موضوع هوش مصنوعی در اروپا را از بُعد قوانین تنظیم شده در نشست بروکسل بررسی می‌کنند. قوانین و بازار فروش هوش مصنوعی

1. Styx

2. Towards Safe, Ethical and Beneficial Artificial Intelligence in the European Union and Beyond

3. Charlotte.

4. Marcus.

5. The Brussels Effect and Artificial Intelligence: How EU Regulation Will Impact the Global AI Market

و همچنین میزان تاثیر گذاري قوانين بر جامعه و بازار است. نيز به بررسي تاثير قوانين هوش مصنوعي بر كشورهاي پيشرو مانند چين و آمريكا پرداخته شده است.

هوورث^۱ در سال ۲۰۱۹ در مقاله خود با عنوان «استقلال استراتژيك و امنيت اتحاديۀ اروپا»^۲ بيان مي‌كند استقلال استراتژيك به معنای توانایی يك منطقه يا كشور در تصميم گيري و اقدام مستقل از نفوذ و فشارهاي خارجي است. اتحاديۀ اروپا به دنبال کاهش وابستگي خود به ديگر قدرت‌هاي جهاني خصوصاً ايالات متحده و چين در حوزه‌هاي كليدي همچون دفاع، انرژي، فناوري و اقتصاد است. استون^۳ و همكاران در سال ۲۰۲۰، در مقاله‌اي با عنوان «هوش مصنوعي و زندگي در سال ۲۰۳۰»^۴ بيان مي‌كنند هوش مصنوعي به عنوان يكي از فناوري‌هاي پيشرفته، نقش حياتي در تقويت استقلال استراتژيك اتحاديۀ اروپا ايفا مي‌كند. اين فناوري مي‌تواند در حوزه‌هاي مختلفي از جمله دفاع و امنيت، بهداشت، كشاورزي و حمل و نقل کاربرد داشته باشد.

مرکز اروپا برای اقتصاد سياسي بين‌المللي (ECIPE)^۵ ۲۰۲۲، در گزارش خود با نام «استقلال استراتژيك و هوش مصنوعي: اتحاديۀ اروپا در مسابقه جهاني»^۶ به بررسي استراتژي‌هاي استقلال استراتژيك اتحاديۀ اروپا در حوزه هوش مصنوعي مي‌پردازد و به تاثيرات اين استراتژي‌ها بر توسعه فناوري و رقابت پذيري منطقه‌اي مي‌پردازد. ميهاي^۷ و سيمونا^۸ در سال ۲۰۲۰ در مقاله‌اي با نام «استقلال استراتژيك در سياست هوش مصنوعي اتحاديۀ اروپا»^۹ به بررسي استراتژي‌هاي استقلال استراتژيك اتحاديۀ اروپا در سياست‌هاي هوش مصنوعي مي‌پردازند و تاثيرات آن را بر رقابت پذيري و امنيت منطقه‌اي بررسي مي‌كنند.

1. Haworth.

2. Strategic Autonomy and EU Security.

3. Stone.

4. Artificial Intelligence and Life in 2030.

5. European Centre for International Political Economy.

6. Strategic Autonomy and Artificial Intelligence: European Union in the Global Race

7. Mihai Macovei

8. Simona Mihai Yiannaki

9. Strategic Autonomy in the EU's AI Policy

چارچوب نظری و روش‌شناسی پژوهش: سیاستگذاری عمومی و تنظیم‌گری^۱

نظریه تنظیم‌گری به‌عنوان یکی از شاخه‌های مهم در سیاستگذاری عمومی و علوم سیاسی، به بررسی نقش دولت و نهادهای عمومی در تنظیم و کنترل فعالیت‌های اقتصادی، اجتماعی و زیست‌محیطی می‌پردازد. این نظریه در طول تاریخ تکامل یافته و تلاش می‌کند توضیح دهد چگونه دولت‌ها با استفاده از سیاست‌ها و قوانین، رفتارهای بازار و افراد را هدایت کرده و از پیامدهای منفی جلوگیری می‌کنند. در اوایل قرن بیستم افرادی مانند آرتور پیگو نظریه شکست بازار را مطرح کردند. وی معتقد است در برخی موارد بازارها به تنهایی قادر به تخصیص بهینه منابع نیستند. انحصارها، اطلاعات ناقص و آثار جانبی (مانند آلودگی محیط زیست) نمونه‌هایی از شکست بازار هستند که نیاز به دخالت دولت و تنظیم‌گری دارند (Pigou, 2017: 70).

در دهه ۱۹۷۰، نظریه تسخیر تنظیم‌گر توسط اقتصاددانانی مانند جرج استیگلر مطرح شد. این نظریه بیان می‌کند نهادهای تنظیم‌گر که وظیفه نظارت و کنترل صنایع و بازارها را برعهده دارند، ممکن است به تدریج تحت تاثیر لابی‌ها و گروه‌های ذینفع قرار گیرند و به‌جای خدمت به منافع عمومی، به نفع صنایع و شرکت‌هایی که قرار است تنظیم شوند، عمل کنند (Stigler, 2021: 73). تنظیم‌گری به‌عنوان یکی از وظایف اصلی دولت‌ها و نهادهای عمومی، فرآیند تدوین و اجرای قوانین و مقررات برای هدایت رفتار اقتصادی، اجتماعی و زیست‌محیطی افراد و سازمان‌ها است. چارچوب نظریه تنظیم‌گری به مجموعه‌ای از اصول، مفاهیم و نظریه‌ها اشاره دارد که در تحلیل و تدوین این مقررات به کار گرفته می‌شود. این چارچوب به‌ویژه در سیاست‌گذاری‌های عمومی، هدفمند کردن نظارت‌ها و کنترل رفتارهای بازار اهمیت دارد. سیاست تنظیم‌گری دارای زیرشاخه‌های متعددی است که هر کدام از منظر متفاوتی به موضوع نگاه می‌کنند:

نظریه شکست بازار^۲: این نظریه اساس بسیاری از سیاست‌های تنظیم‌گری را تشکیل می‌دهد و بر این ایده استوار است که تنظیم‌گری زمانی ضروری است که بازارها به تنهایی

1. Regulatory Theory
2. Market Failure Theory

نتوانند تخصیص منابع بهینه را محقق کنند. از جمله موارد شکست بازار می‌توان به وجود انحصار، اطلاعات نامتقارن و آثار جانبی منفی اشاره کرد. تنظیم‌گری در این موارد به منظور اصلاح این نقص‌ها و افزایش کارایی اجتماعی صورت می‌گیرد (Posner, 2001: 241).

نظریه تسخیر تنظیم‌گر^۱: این نظریه پیشنهاد می‌کند که نهادهای تنظیم‌گر ممکن است به تدریج تحت تأثیر گروه‌های ذینفع قرار بگیرند و به جای عمل به نفع عموم، به نفع صنایع و شرکت‌هایی که قرار است تنظیم شوند، عمل کنند. این پدیده می‌تواند منجر به کاهش اثربخشی مقررات و حتی تشدید مشکلاتی شود که تنظیم‌گری برای حل آن‌ها ایجاد شده است (Baldwin, Cave & Lodge, 2011: 16).

استفاده از چارچوب نظری تنظیم‌گری به سیاست‌گذاران و محققان اجازه می‌دهد رفتار و اثرات نهادهای تنظیم‌گر را بهتر درک کنند و به بهبود طراحی و اجرای مقررات بپردازند. به عنوان مثال از دریچه نظریه فوق می‌توان در حوزه‌هایی چون تنظیم‌گری اقتصادی (مانند کنترل انحصارها)، تنظیم‌گری محیط زیستی (مانند قوانین مربوط به کاهش آلودگی) و تنظیم‌گری فناوری (مانند مقررات هوش مصنوعی و داده‌های شخصی)، به شناسایی و تحلیل چالش‌ها و فرصت‌ها اقدام کرد (Stigler, 2021: 70). هوش مصنوعی به عنوان یکی از تکنولوژی‌های مهم قرن بیست و یکم، تحولات گسترده‌ای در صنایع مختلف ایجاد کرده است. این تکنولوژی علاوه بر فرصت‌های فراوان، چالش‌هایی نیز به همراه دارد که از جمله آن‌ها می‌توان به مسائل اخلاقی، حفظ حریم خصوصی و تبعیض‌های ناشی از الگوریتم‌ها اشاره کرد. به همین دلیل تنظیم‌گری در حوزه هوش مصنوعی اهمیت ویژه‌ای پیدا کرده است. اتحادیه اروپا به عنوان یکی از پیشگامان تنظیم‌گری هوش مصنوعی، چارچوب‌های قانونی و سیاستی متعددی را تدوین کرده است که این فرآیند را هدایت می‌کنند. اتحادیه اروپا با استفاده از چارچوب نظری تنظیم‌گری، تلاش می‌نماید به طور موثر به چالش‌های ناشی از هوش مصنوعی پاسخ دهد. این چارچوب به ویژه بر مبنای نظریه‌های شکست بازار و نظریه نهادگرایی قرار دارد.

نظریه شکست بازار و تنظیم‌گری هوش مصنوعی: نظریه شکست بازار پیشنهاد می‌کند که دولت‌ها باید زمانی که بازار قادر به تخصیص منابع بهینه نیست، مداخله کنند. در مورد هوش مصنوعی، خطرات بالقوه مانند تبعیض الگوریتمی، اطلاعات ناقص و تهدیدات حریم خصوصی نشان‌دهنده شکست‌های احتمالی بازار هستند. اتحادیه اروپا از طریق تدوین مقرراتی مانند "قانون هوش مصنوعی" تلاش کرده است شکست‌های بازار را اصلاح کند و استفاده از هوش مصنوعی را در چارچوبی امن و قابل اعتماد قرار دهد (European Commission, 13 April, 2021).

نظریه تسخیر تنظیم‌گر و چالش‌های هوش مصنوعی: در حالی که چارچوب تنظیم‌گری اتحادیه اروپا برای هوش مصنوعی تلاش می‌کند به نفع عموم عمل نماید، نظریه تسخیر تنظیم‌گر یادآور می‌شود که همواره خطر نفوذ گروه‌های ذینفع وجود دارد. برای مثال شرکت‌های بزرگ فناوری ممکن است تلاش کنند مقررات به گونه‌ای تدوین شود که منافع آن‌ها حفظ شود. اتحادیه اروپا برای مقابله با این چالش، فرآیندهای شفاف‌سازی و مشارکت عمومی در تدوین مقررات هوش مصنوعی را تقویت کرده است (European Commission, 13 April, 2021). چارچوب نظری تنظیم‌گری به عنوان ابزاری مفهومی به اتحادیه اروپا کمک کرده است به چالش‌ها و فرصت‌های ناشی از هوش مصنوعی به طور جامع و هماهنگ پاسخ دهد. با بهره‌گیری از نظریه‌های شکست بازار، نهادگرایی و تسخیر تنظیم‌گر، اتحادیه اروپا توانسته است مقرراتی تدوین کند که نه تنها به مدیریت خطرات مرتبط با هوش مصنوعی بپردازد، بلکه نوآوری و رقابت را نیز تشویق کند.

روش‌شناسی پژوهش: روش تحلیل مضمون

در این نوع پژوهش تمرکز بر تحلیل و استنباط مضامین اصلی از متون مرتبط است. گردآوری اطلاعات در این تحقیق از طریق منابع کتابخانه‌ای و اینترنتی صورت گرفته است. اسناد رسمی اتحادیه اروپا، مقالات پژوهشی و گزارش‌های مرتبط با حوزه هوش مصنوعی و سیاست‌گذاری عمومی از جمله منابع مورد استفاده بوده‌اند. برای تحلیل داده‌ها

از روش تحلیل مضمون استفاده شده است. این روش به عنوان یکی از تکنیک‌های متداول در تحقیقات کیفی، به محققان اجازه می‌دهد از میان محتوای متنی، مضامین اصلی و فرعی را شناسایی و دسته‌بندی کنند. در این پژوهش، تحلیل مضمون کمک کرده است مفاهیم کلیدی همچون استقلال استراتژیک، حکمرانی هوش مصنوعی و شفافیت و امنیت داده‌ها را استخراج و ساختار مقاله بر اساس این مضامین تدوین شود:

مراحل تحلیل مضمون در این پژوهش به صورت زیر انجام شده است:

- **کدگذاری مضامین:** ابتدا جملات و بخش‌های مرتبط با استقلال استراتژیک و حکمرانی هوش مصنوعی کدگذاری شدند.
- **دسته‌بندی کدها:** مضامین اصلی به زیرمضامین مرتبط مانند کاهش وابستگی، حاکمیت دیجیتال و شفافیت تقسیم‌بندی شدند.
- **شناسایی الگوهای اصلی:** در نهایت، ارتباط میان مضامین و تأثیرات آن‌ها بر سیاست‌های اتحادیه اروپا در حوزه هوش مصنوعی در این مقاله مورد بررسی قرار گرفت.

توضیحات	کدهای باز
این موضوع یکی از شاخص‌های کلیدی استقلال استراتژیک اتحادیه اروپا است که هدف آن خودکفایی در تصمیم‌گیری‌ها و توسعه فناوری‌های مهم است.	کاهش وابستگی به قدرت‌های جهانی
جلوگیری از وابستگی به کشورهای دیگر در زمینه سیاست‌گذاری‌های داخلی، اقتصادی و تکنولوژیک است.	تصمیم‌گیری مستقل در حوزه‌های کلیدی
به مقررات سخت‌گیرانه‌ای اشاره دارد که اتحادیه اروپا برای هوش مصنوعی وضع کرده است. این قوانین به منظور افزایش اعتماد عمومی و اطمینان از استفاده اخلاقی و شفاف از هوش مصنوعی تدوین شده‌اند.	تدوین و اجرای مقررات جامع
بر اولویت دادن به منافع عمومی در قانون‌گذاری‌های مربوط به هوش مصنوعی تأکید دارد. اتحادیه اروپا سعی دارد مقرراتی را وضع کند که در خدمت منافع جامعه باشد و از کاربردهایی که ممکن است برای عموم خطرآفرین باشد، جلوگیری کند.	حفظ منافع عمومی

توضیحات	کدهای باز
به حفاظت از داده‌ها و حریم خصوصی شهروندان اشاره دارد که در چارچوب سیاست‌های هوش مصنوعی اتحادیه اروپا بسیار مورد توجه است. این امر با توجه به خطرات مرتبط با امنیت داده‌ها اهمیت دارد و از جمله اهداف کلیدی اتحادیه در حکمرانی هوش مصنوعی است.	حفاظت از داده‌ها و حریم خصوصی
بیانگر تلاش‌های اتحادیه اروپا برای شفافیت و توضیح‌پذیری در سیستم‌های هوش مصنوعی است. این ویژگی به کاربران اجازه می‌دهد تا بدانند هوش مصنوعی چگونه به تصمیم‌گیری می‌پردازد و از اعتماد عمومی پشتیبانی می‌کند.	شفافیت و قابل توضیح بودن سیستم
این کد به تلاش اتحادیه اروپا برای افزایش اعتماد عمومی به هوش مصنوعی از طریق ایجاد شفافیت، حفظ حریم خصوصی و اجرای اصول اخلاقی اشاره دارد. هدف این است که شهروندان اطمینان داشته باشند که هوش مصنوعی با رعایت حقوق و ارزش‌های انسانی به کار گرفته می‌شود. اتحادیه اروپا با وضع مقررات سخت‌گیرانه در حوزه امنیت داده‌ها و شفافیت سیستم‌ها، به دنبال جلب اعتماد جامعه و تسهیل پذیرش هوش مصنوعی در جامعه است.	افزایش اعتماد عمومی
به اصول اخلاقی و ارزش‌های انسانی که باید در طراحی و استفاده از هوش مصنوعی رعایت شوند، اشاره دارد. اتحادیه اروپا بر این باور است که هوش مصنوعی نباید باعث تبعیض یا نقض کرامت انسانی شود و باید به گونه‌ای طراحی شود که در خدمت جامعه باشد.	اخلاق و اصول انسانی
به فرآیندهای پیش‌بینی و مدیریت ریسک‌های مرتبط با هوش مصنوعی اشاره دارد. اتحادیه اروپا با در نظر گرفتن ریسک‌های احتمالی، مقرراتی را تدوین کرده است که به مدیریت بهتر این فناوری کمک می‌کند.	پیش‌بینی و مدیریت ریسک‌ها
به لزوم نظارت انسانی بر سیستم‌های هوش مصنوعی اشاره دارد. اتحادیه اروپا باور دارد که در تصمیمات مهم و حساس (به‌ویژه در زمینه‌هایی مانند پزشکی، قضایی و امنیتی)، حضور و نظارت انسانی ضروری است تا از بروز خطاها یا تصمیمات غیرعادلانه جلوگیری شود. نظارت انسانی تضمین می‌کند که سیستم‌های هوش مصنوعی مطابق با اصول اخلاقی و حقوق بشر عمل کنند و تصمیمات مهم بدون مداخله انسان انجام نشوند.	نظارت انسانی بر سیستم‌ها

توضیحات	کدهای باز
این کد به کاربردهای هوش مصنوعی در امنیت سایبری و دفاعی اتحادیه اروپا اشاره دارد. اتحادیه تلاش می‌کند از هوش مصنوعی به عنوان ابزاری برای مقابله با تهدیدات سایبری و افزایش امنیت ملی استفاده کند.	هوش مصنوعی و امنیت سایبری
این کد به کاربردهای هوش مصنوعی در حوزه دفاع از امنیت ملی اشاره دارد. اتحادیه اروپا از هوش مصنوعی به عنوان ابزاری استراتژیک برای تقویت امنیت ملی و دفاعی استفاده می‌کند. این شامل توسعه سیستم‌های هوش مصنوعی برای پیش‌بینی تهدیدات، تحلیل سریع اطلاعات و اجرای عملیات دفاعی است. هوش مصنوعی همچنین در سیستم‌های نظارتی و تحلیل اطلاعات مورد استفاده قرار می‌گیرد تا از امنیت و ثبات ملی حمایت کند.	دفاع از امنیت ملی

کدهای اصلی انتخاب‌شده در این پژوهش بر اساس مفاهیم و مضامینی است که در اسناد رسمی و سیاست‌گذاری‌های اتحادیه اروپا درباره هوش مصنوعی و استقلال استراتژیک آمده‌اند. در ادامه به بررسی اسنادی پرداخته شده است که هر یک از این مضامین در آن‌ها مطرح شده‌اند و سپس در ادامه سیاست‌های اتحادیه اروپا در قالب مضامین استخراج شده به بحث گذاشته شده است.

کدهای باز (اولیه)	کدهای محوری (اصلی)
کاهش وابستگی به قدرت‌های جهانی	استقلال استراتژیک اتحادیه اروپا ^۱
تصمیم‌گیری مستقل در حوزه‌های کلیدی	
تدوین و اجرای مقررات جامع	حکمرانی هوش مصنوعی ^۲
حفظ منافع عمومی	

1. European Commission, (2018b), "Artificial Intelligence for Europe", *European Commission, Brussels*.
2. European Commission, (2018b), "Where Will the EU's Strategic Compass Point?", *European Parliamentary Research Service*.

کدهای محوری (اصلی)	کدهای باز (اولیه)
شفافیت و امنیت داده ^۱	حفاظت از داده ها و حریم خصوصی
	شفافیت و قابل توضیح بودن سیستم
	افزایش اعتماد عمومی
مدیریت ریسک و اخلاق هوش مصنوعی ^۲	اخلاق و اصول انسانی
	پیش‌بینی و مدیریت ریسک‌ها
	نظارت انسانی بر سیستم‌ها
کاربردهای امنیتی و نظامی هوش مصنوعی ^۳	هوش مصنوعی و امنیت سایبری
	دفاع از امنیت ملی

استقلال استراتژیک اتحادیه اروپا

اتحادیه اروپا در زمینه حاکمیت دیجیتال، اغلب زیر چتر یک ابتکار گسترده‌تر در مورد «استقلال استراتژیک» مورد توجه قرار می‌گیرد (Timmers, 2019: 636). مفهوم خودمختاری استراتژیک در استراتژی جهانی اتحادیه اروپا در سال ۲۰۱۶ معرفی شد که در آن حاکمیت، عمدتاً در چارچوب موضوعات نظامی، امنیتی و دفاعی مورد استفاده قرار می‌گرفت (European External Action Services, 2016). با این حال، حاکمیت دیجیتال معنای بسیار گسترده‌تری دارد؛ علی‌رغم فقدان یک تعریف منسجم و روشن، تلاش برای دسترسی به حاکمیت دیجیتالی مبتنی بر وابستگی کامل اتحادیه اروپا در قبال دیجیتال شرکت‌های فناوری خارجی است.

در همین راستا کمیسیون اروپا پیشنهاد وضع مقررات و پارلمان اروپا به همراه شورای اروپا وضع قوانین هماهنگ در مورد هوش مصنوعی (قانون هوش مصنوعی) و اصلاح برخی از قوانین قانونی اتحادیه را از ۲۱ آوریل ۲۰۲۱ با اهداف خاص زیر ارائه کرده است:

1. European Commission, (2019), "Communication from the Commission to the European Parliament, The Council, the European Economic and Social Committee and the Committee of the Regions - Building Trust in Human", Com, 168.
2. European Parliament, (2022), "Artificial Intelligence: Meps Want the Eu to be a Global Standard-Setter".
3. European Parliament, (2021), "Guidelines for Military and Non-Military Use of Artificial Intelligence".

نخست، این اطمینان حاصل شود که سیستم‌های هوش مصنوعی که در بازار اتحادیه مورد استفاده قرار می‌گیرند، ایمن هستند و به قوانین موجود در مورد حقوق اساسی و ارزش‌های اتحادیه احترام می‌گذارند؛ دوم، اطمینان به چارچوب قانونی برای تسهیل سرمایه‌گذاری و نوآوری در هوش مصنوعی در دستور کار قرار گیرد؛ سوم، افزایش حاکمیت و اجرای موثر قوانین موجود در مورد حقوق اساسی و الزامات ایمنی قابل اعمال در سیستم‌های هوش مصنوعی و چهارم، توسعه یک بازار واحد برای برنامه‌های کاربردی هوش مصنوعی قانونی، ایمن و قابل اعتماد و جلوگیری از تکه‌تکه شدن بازار (COM.E, 24 April 2021).

پیشنهادات فوق یک چارچوب قانونی قوی اما در عین حال منعطف را با مکانیسم‌هایی تعیین می‌کند که آن را قادر می‌سازد مطابق با تکامل فناوری و موقعیت‌هایی که در آینده پدیدار می‌شوند، سازگار شود. این چارچوب قانونی در انتخاب‌های نظارتی خود به‌ویژه در تمام الزامات مبتنی بر اصولی که سیستم‌های هوش مصنوعی باید از آنها پیروی کنند، کاملاً وجود دارد. از منظر اتحادیه اروپا، ایجاد یک سیستم نظارتی متناسب با محوریت یک رویکرد نظارتی ریسک‌پذیر به‌خوبی تعریف شده است. لذا از منظر کمیسیون اتحادیه اروپا، مداخله قانونی در استفاده از هوش مصنوعی در شرایط مشخصی صورت می‌گیرد که دلیل موجهی برای نگرانی وجود داشته باشد یا در آینده نزدیک بتوان چنین نگرانی را به طور منطقی پیش‌بینی کرد (European Commission, 21 April 2021).

اولویت کمیسیون، یک چارچوب نظارتی فقط برای سیستم‌های هوش مصنوعی پرخطر است. همه ارائه‌دهندگان سیستم‌های هوش مصنوعی کم‌خطر این امکان را دارند که از یک کد رفتار پیروی کنند. الزامات مربوط به کیفیت داده‌ها، اسناد و قابلیت ردیابی، ارائه اطلاعات و شفافیت، نظارت و استحکام انسانی، دقت و امنیت سایبری مورد توجه قرار می‌گیرد و برای سیستم‌های هوش مصنوعی پرخطر الزامی است. همه شرکت‌های فناوری اطلاعات که ممکن است کدهای رفتاری را برای سایر سیستم‌های هوش مصنوعی معرفی کنند، داوطلبانه این کار را انجام می‌دهند (European Commission, 2021: 9).

نکته قابل توجه این است که اتحادیه اروپا این امکان را دارد که تمام فناوری هوش مصنوعی خارجی را که ممکن است شهروندان اروپایی را تحت تأثیر قرار دهد و می تواند اقداماتی را علیه آن انجام دهد و حتی آن را ممنوع کند به عنوان فناوری هوش مصنوعی پرخطر طبقه بندی کند. به عبارت دیگر، سیستم های هوش مصنوعی که توانایی تحریف رفتار انسان، دستکاری، بهره برداری و استفاده از شیوه های کنترل اجتماعی و امتیازدهی اجتماعی را دارند از منظر کمیسیون اتحادیه اروپا ممنوع اعلام شده است (European Commission, 2021: 20-21). نکته مهم در سیاست فوق این است که ارائه دهندگان خارجی فناوری هوش مصنوعی اگر می خواهند محصولات خود را در بازار اتحادیه اروپا بفروشند باید به الزامات قانون هوش مصنوعی احترام بگذارند و این بدان معنی است که اتحادیه اروپا با اعمال مقررات خود قدرت تأثیر گذاری و تغییر فناوری های هوش مصنوعی را دارد.

حکمرانی هوش مصنوعی

حاکمیت قانون یکی از ارکان اتحادیه اروپا است که در ماده ۲ معاهده اتحادیه اروپا به عنوان یک ارزش بنیادی ذکر شده است. بهره مندی از سایر ارزش ها و حقوق، ناشی از تبعیت از اصول حاکمیت قانون است. یکی از اصول مهم حاکمیت قانون شامل پاسخگویی، شفافیت، اطمینان قانونی، احترام به حقوق اساسی و یکسانی در برابر قانون است. همه این اصول ممکن است با جایگزینی سیستم های هوش مصنوعی بر نظارت انسانی تأثیر منفی بگذارد. با این حال زمانی که سیستم های هوش مصنوعی مطابق با حاکمیت قانون مستقر شوند، می توانند دستاوردهای کارآمد قابل توجهی در مدیریت عمومی، امنیت و اجرای عدالت ایجاد کنند (Emily, 2022: 2).

در حالی که قانون هوش مصنوعی توسط متخصصانی همچون ساندرا واچر^۱ با تأکید بر "سیستم هماهنگ" مورد توجه قرار گرفته است، منتقدان به ناقص بودن قانون هوش مصنوعی اشاره می نمایند که می تواند در اتحادیه اروپا و فراتر از آن اعمال شود

1. Sandra Wachter

(Wachter, 2024: 680). مقررات جدید هوش مصنوعی با این هدف وضع می‌شود که اروپایی‌ها می‌توانند به آنچه هوش مصنوعی ارائه می‌دهد اعتماد کنند. قوانین مناسب و منعطف به خطرات خاص ناشی از سیستم‌های هوش مصنوعی می‌پردازد و بالاترین استاندارد را در سراسر جهان تعیین می‌کند. برنامه هماهنگی، تغییرات سیاست و سرمایه‌گذاری لازم را در سطوح کشورهای عضو برای تقویت موقعیت پیشرو اروپا در توسعه هوش مصنوعی انسان‌محور، پایدار، ایمن، فراگیر و قابل اعتماد ترسیم می‌کند (COM.E, 21Apr 2021). فناوری‌هایی مانند یادگیری ماشینی، ربات‌های چت هوشمند و تشخیص چهره و گفتار به مرحله جدیدی تبدیل شده است که در سطح جهانی بین دولت‌های سراسر جهان مشترک است. هوش مصنوعی نشان دهنده فناوری ایده‌آلی برای اعمال در زمینه بخش عمومی است، جایی که تنظیمات محیطی دائماً در حال تغییر هستند و برنامه‌ریزی‌های قبلی نمی‌تواند همه موارد ممکن را در نظر بگیرند (Sun and Medaglia, 2019: 370).

شفافیت و امنیت داده

کمیسیون اشتراک‌گذاری داده‌ها را آسان‌تر می‌کند و داده‌های بیشتری را که عنوان ماده خام برای هوش مصنوعی جمع‌آوری می‌شود، از صنایع، تحقیقات عمومی و علمی تجزیه می‌کند. در این سند، اتحادیه اروپا بر ارزش‌های اروپایی مانند «مقررات عمومی حفاظت از داده‌ها» تاکید می‌کند و اعتماد و خلاقیتی را که در بلندمدت برای مردم و شرکت‌ها ضروری است، تضمین می‌کند. این ارتباطات یک ابتکار اروپایی را برای تقویت ظرفیت‌های فنی و صنعتی، آماده‌سازی چارچوب اخلاقی و قانونی و آمادگی در برابر تغییرات اجتماعی - اقتصادی مانند سیستم‌های آموزشی، بازار کار و حمایت‌های اجتماعی ارائه می‌کند (European Commission, 25 April 2018b).

بر این اساس، اتحادیه اروپا قصد دارد شبکه مراکز برتر هوش مصنوعی، مراکز پشتیبانی برای اشتراک‌گذاری داده‌ها، ایجاد تسهیلات آزمایشی و پیشرفته‌تر در حمل و نقل، مراقبت‌های بهداشتی، مواد غذایی کشاورزی، تولید و حمایت از پذیرش محصولات هوش مصنوعی را ارتقا دهد. خدمات در تمامی بخش‌های کمیسیون مصمم به بهبود پذیرش

دیجیتالی شدن در این بخش‌ها است و می‌خواهد منافع عمومی را افزایش دهد و شهروندان را وادار کند که مهارت‌ها و دانش‌های مورد نیاز خود را با شروع تغییر یا ناپدید شدن مشاغل به دلیل فرایند دیجیتالی‌سازی و اتوماسیون به‌دست آورند (European Commission, 25 April 2018b).

برای هوش مصنوعی قابل اعتماد برخی الزامات وجود دارد که لازم است به آنها توجه شود. کمیسیون از الزامات کلیدی زیر برای هوش مصنوعی قابل اعتماد پشتیبانی می‌کند که بر اساس ارزش‌های اروپایی هستند:

۱. نمایندگی و نظارت انسانی: سیستم هوش مصنوعی باید از افراد در انتخاب‌های بهتر و آگاهانه‌تر مطابق با اهدافش حمایت کند؛

۲. استحکام فنی و ایمنی: هوش مصنوعی قابل اعتماد به الگوریتم‌هایی نیاز دارد که برای مقابله با خطاها یا ناسازگاری‌ها، در تمام مراحل چرخه حیات سیستم هوش مصنوعی، ایمن، قابل اعتماد و قوی باشند و به اندازه کافی با نتایج نادرست مقابله کنند؛

۳. حریم خصوصی و حاکمیت داده: حفظ حریم خصوصی و داده‌ها باید در تمامی مراحل چرخه حیات سیستم هوش مصنوعی تضمین شود؛

۴. شفافیت: ردیابی سیستم‌های هوش مصنوعی باید تضمین شود. ثبت و مستندسازی تصمیمات اخذ شده توسط سیستم‌ها و همچنین کل فرایند که منجر به تصمیمات می‌شود مهم است؛

۵. رفاه اجتماعی و محیطی: برای قابل اعتماد بودن هوش مصنوعی، تاثیر آن بر محیط و سایر موجودات باید در نظر گرفته شود؛

۶. مسئولیت: باید مکانیسم‌هایی برای اطمینان از مسئولیت و پاسخگویی سیستم‌های هوش مصنوعی و نتایج آنها، چه قبل و چه بعد از اجرای آنها ایجاد شود (COM.E, 2019: 168).

مدیریت ریسک و اخلاق هوش مصنوعی

اتحادیه اروپا با اتخاذ رویکرد مبتنی بر ریسک نسبت به هوش مصنوعی، اعتماد به هوش مصنوعی را از نظر قابل قبول بودن خطرات آن مورد توجه قرار داده است. این آمیختگی

قابلیت اعتماد با قابل قبول بودن ریسک، نیاز به تامل بیشتری دارد. کمیسیون اروپا به عنوان بخشی از تلاش جهانی در جهت اعتمادسازی، چارچوب قانونی برای هوش مصنوعی پیشنهاد کرده است. قانون هوش مصنوعی به صراحت هدف دوگانه ترویج جذب فناوری و پرداختن به خطرات مرتبط با استفاده از هوش مصنوعی را دنبال می‌کند (Laux, 2023: 15). بارزترین علامت تجاری رویکرد اتحادیه اروپا به حکمرانی هوش مصنوعی، احتمالاً مفهوم "هوش مصنوعی قابل اعتماد" است. طبق دستورالعمل‌های اخلاقی، قابل اعتماد بودن در هوش مصنوعی به این معنا است که یک سیستم هوش مصنوعی باید: قانونمند به معنی رعایت کلیه قوانین و مقررات قابل اجرا؛ اخلاقی به معنای رعایت اصول و ارزش‌های اخلاقی و مستحکم به معنای فنی و اجتماعی باشد (Now, 2020: 6).

قوانین تنظیم شده هوش مصنوعی توسط کمیسیون اروپا رویکردی مبتنی بر ریسک را برای تنظیم هوش مصنوعی پیشنهاد می‌کند و چهار دسته از ریسک را ترسیم می‌کند: غیرقابل قبول، قبول، پرخطر، ریسک محدود و حداقل / بدون ریسک. سیستم‌هایی که خطر غیرقابل قبول دارند از جمله موارد امتیازدهی اجتماعی و سیستم‌های دستکاری ناخودآگاه ممنوع خواهد بود. هوش مصنوعی پرخطر شامل سیستم‌هایی می‌شود که اجزای حیاتی ایمن هستند و آن‌هایی که خطرات خاصی برای حقوق اساسی دارند. برای این سیستم‌ها تعهدات خاص برای ارائه دهندگان، وارد کنندگان، توزیع کنندگان، کاربران و نمایندگان مجاز مشخص شده است (Roberts et al., 2021: 6).

مقررات مبتنی بر ریسک ذاتاً به دنبال ایجاد یک تبعیض بین موضوعات قانون است. بنابراین رژیم حقوقی حاکم بر آنها را دقیقاً بر اساس عامل خطر متمایز می‌کند. از این نظر، مقررات مبتنی بر ریسک، دنبال کردن اهداف مشابه آنچه آدرین ورمول^۱ به عنوان «بهنه سازی قانون اساسی» یا «موقعیت بالغ» برای تنظیم ریسک (قانون) تعریف کرده است. ورمول در واقع تمایزی بین «مشروطیت احتیاطی» و «مشروطیت بهینه سازی» اعمال می‌کند.

1. Adrian Vermeule

در مقررات مبتنی بر ریسک، چنین عملیات متعادل سازی تا حدی مستقیماً به صلاح حدید «تنظیم گر» واگذار می شود (Dunn and De Gregorio, 2022: 2-3).

رویکرد مبتنی بر ریسک بر یک سیستم متمرکز متکی است، زیرا این به کمیسیون اروپا بستگی دارد که تعیین کند چه چیزی یک ریسک بالاست و چه چیزی نه. این شامل سیستم هایی برای مدیریت و بهره برداری از زیرساخت های حیاتی است. علاوه بر این کمیسیون می تواند فهرست پرخطر را در قوانین خود اصلاح کند. در صورتی که تشخیص دهد سیستم هوش مصنوعی دیگری که هنوز در نظر گرفته نشده است، خطری برای آسیب، ایمنی و حقوق اساسی دارد و به اندازه ای پرخطر است که لازم است در ضمیمه گنجانده شود (Gellert, 2021: 6).

نمودار: هرم خطرات



Data Source: European Commission

آنچه نگران کننده است این است که سیستم های هوش مصنوعی به اصطلاح «مستقل» فقط به کنترل داخلی نیاز دارند. طبق ماده ۴۳ قانون هوش مصنوعی، شرکت یک سازمان آگاه فقط برای تعداد محدودی از سیستم های هوش مصنوعی ضروری است. به عنوان مثال سیستم های بیومتریک خاص. با این رویکرد، این پیشنهاد به ارائه دهندگان هوش مصنوعی

در مورد کاربرد و اجرای قوانین هوش مصنوعی یک محدوده اختیارات نابجا می‌دهد. قانون هوش مصنوعی تلاش کرده است الزامات ضروری برای سیستم‌های هوش مصنوعی پرخطر را اجباری کند؛ با این حال بسیاری از الزامات به روشی نسبتاً گسترده بیان شده و اختیارات زیادی را به ارائه دهندگان هوش مصنوعی می‌دهد (Ebers et al., 2021: 595).

نمایندگان پارلمان اتحادیه اروپا همچنین در مورد تهدیدات علیه حقوق بشر و حاکمیت دولتی ناشی از استفاده از فناوری هوش مصنوعی در نظارت سطح وسیع در دو بعد غیرنظامی و نظامی هشدار می‌دهند. آنها خواستار ممنوعیت مقامات دولتی در استفاده از "برنامه‌های امتیازدهی اجتماعی بسیار مزاحم" (برای نظارت و رتبه‌بندی شهروندان) هستند. این گزارش همچنین نگرانی‌هایی را درباره فناوری‌های جعلی که پتانسیل بی‌ثبات‌سازی کشورها، انتشار اطلاعات نادرست و تاثیرگذاری بر انتخابات را دارند، افزایش می‌دهد. پدیدآورندگان قوانین هوش مصنوعی باید موظف شوند که چنین مطالبی را به‌عنوان غیراصلی برجسب‌گذاری کنند و برای مقابله با این پدیده باید تحقیقات بیشتری در زمینه فناوری انجام شود (European Parliament, 20 January 2021).

کمیسیون اتحادیه اروپا با ادعای ایجاد یک چارچوب هوش مصنوعی انسان‌محور که مردم را در اولویت قرار می‌دهد، شروع به کار کرد. نکته مهم‌تر این است که پیشنهاد فعلی به‌طور کامل دیدگاه کسانی را که تحت تاثیر خروجی سیستم‌های هوش مصنوعی قرار گرفته‌اند نادیده می‌گیرد. اگر چنین سیستم‌هایی بر زندگی مردم تاثیر منفی داشته باشند، نه تنها باید شفافیت در مورد استقرار این سیستم‌ها را داشته باشند، بلکه باید امکان به چالش کشیدن نتایج را نیز داشته باشند. بنابراین باید برای افراد و گروه‌های آسیب‌دیده، گزینه‌هایی به آسانی دسترسی و تضمین شده وجود داشته باشد تا بتوانند چنین تصمیم‌هایی را به چالش بکشند و در صورت لزوم خواستار لغو، بررسی مجدد و تفاوت جبران خسارت شوند (European Commission, 21 April 2021).

ایده مقامات اتحادیه اروپا برای تنظیم هوش مصنوعی، مبتنی بر مطالعات تحقیقاتی و داده‌های آماری است که نشان دهنده نظر شهروندان اتحادیه اروپا در زمینه هوش مصنوعی

و ربات‌ها است. در می سال ۲۰۱۷، کمیسیون اروپا یک نظرسنجی یوروبارومتر^۱ را منتشر کرد که در آن نظرات شهروندان اروپایی در مورد تاثیر دیجیتالی شدن و اتوماسیون بر زندگی روزمره ارائه شد (European Commission, 10 May 2017). در مارس ۲۰۱۷، زمانی که نظرسنجی یوروبارومتر انجام شد، شهروندان اروپایی باید به این پرسش پاسخ می‌دادند: «در ۱۲ ماه گذشته، آیا چیزی در مورد هوش مصنوعی شنیده‌اید، خوانده‌اید یا دیده‌اید؟» از مجموع پاسخ دهندگان، ۵۲ درصد پاسخ دادند «نه»، ۴۷ درصد گفتند «بله» و ۱ درصد پاسخ دادند، نمی‌دانم (Special Eurobarometer, 2017: 55). نظر نگرشی که شهروندان اروپایی نسبت به ربات‌ها و هوش مصنوعی دارند، به نظر می‌رسد اکثریت، دیدگاه مثبتی نسبت به این فناوری‌های انقلابی دارند. این نظرسنجی و به‌ویژه پرسش فوق، نقطه شروع این ایده است که هوش مصنوعی باید تنظیم شود. این ایده منجر به تولد قانون هوش مصنوعی اتحادیه اروپا شد که اولین مقررات مربوط به هوش مصنوعی بود.

پتانسیل هوش مصنوعی برای افزایش مزایای اجتماعی و رشد اقتصادی در بسیاری از مقالات تحقیقاتی و اسناد سیاسی مورد تاکید قرار گرفته است. هدف دولت‌ها در سراسر جهان برای آماده‌سازی کشورشان برای معرفی هوش مصنوعی و کشور پیشرو در هوش مصنوعی می‌باشد. در این راستا دولت‌ها، سیاست‌ها و استراتژی‌هایی را برای تسهیل تحقیقات در مورد هوش مصنوعی، توسعه راه‌حل‌های جدید و اتخاذ این فناوری‌ها در اقتصاد و جامعه خود فراهم کرده‌اند (Guenduez & Mettler, January 2022).

کاربردهای امنیتی و نظامی هوش مصنوعی

هوش مصنوعی خوب به معنای توسعه و استفاده از فناوری‌های هوش مصنوعی به گونه‌ای است که با اصول اخلاقی و حقوق بشر همخوانی داشته باشد. اتحادیه اروپا در راستای

۱. یوروبارومتر (Eurobarometer) نظرسنجی‌های عمومی کمیسیون اروپا در کشورهای عضو اتحادیه اروپا است که از سال ۱۹۷۳ اجرا می‌شود. این نظرسنجی‌ها نظرات شهروندان درباره موضوعات مختلف سیاسی، اجتماعی، اقتصادی و فرهنگی را بررسی می‌کنند و به تصمیم‌گیری‌های سیاسی و راهبردی کمک می‌کنند. نتایج به طور منظم منتشر و در دسترس عموم قرار می‌گیرند.

اطمینان از امنیت و اخلاق در استفاده از هوش مصنوعی، راهنمایی‌ها و توصیه‌هایی را برای توسعه‌دهندگان و کاربران این فناوری ارائه کرده است. این راهنمایی‌ها بر اهمیت شفافیت، جلوگیری از تبعیض، حفظ حریم خصوصی و تضمین امنیت داده‌ها تأکید دارند. هدف اتحادیه اروپا از این اقدامات، ایجاد اعتماد عمومی به هوش مصنوعی و اطمینان از این است که این فناوری به نفع جامعه و با احترام به ارزش‌های انسانی به کار گرفته شود. البته نکته مهم در این است که تا زمان حاضر فقدان شرکت‌های فناوری بزرگ اتحادیه اروپا مانع از آن شده که اتحادیه اروپا و کشورهای عضو آن در رقابت برای رهبری فناوری جهانی فعال باشند. در نتیجه، استراتژی اتحادیه اروپا برای دستیابی به حاکمیت دیجیتال نه تنها با هدف ایجاد شرایط رهبری اتحادیه اروپا، بلکه با اجرای ابتکارات حمایت‌گرایانه در دستور کار قرار گرفته است (Csernaton 2022: 396, Lambach & Monsees 2022: 379). لذا حاکمیت دیجیتال اتحادیه اروپا بیشتر با ایده اتحادیه اروپا به عنوان یک بازیگر تنظیم‌کننده در محیط دیجیتال بین‌المللی گره خورده است (Micklitz et al., 2021: 221).

بعد امنیتی هوش مصنوعی در دستور کار سازمان‌های بین‌المللی مختلف مانند گروه هفت، سازمان همکاری اقتصادی و توسعه و همکاری اقتصادی آسیا و اقیانوسیه قرار گرفته است. با این حال در چارچوب سازمان ملل متحد با «کنوانسیون برخی از سلاح‌های متعارف در مورد سیستم‌های تسلیحات خودمختار مرگبار» است که می‌تواند یکی از نخستین تلاش‌ها برای در نظر گرفتن نقش هوش مصنوعی در زمینه نظامی در همکاری‌های بین‌المللی محسوب کرد (High Contracting Parties CCW, 2019: 14).

هوش مصنوعی بر امنیت در چندین حوزه تأثیر می‌گذارد. در زمینه امنیت سایبری، هوش مصنوعی می‌تواند برای کشف و بهره‌برداری از آسیب‌پذیری‌ها در حین اصلاح آسیب‌پذیری‌های خود استفاده شود. بنابراین، دفاعی در برابر حملات سایبری خارجی ایجاد می‌کند (Brown, 2019: 3). هوش مصنوعی نگرانی‌هایی را در زمینه کمپین‌های اطلاعات نادرست به خود جلب کرده است (Marsden et al., April 2020). در حالی که هوش مصنوعی ابزارهایی را برای از بین بردن اعتماد اجتماعی از طریق انتشار سریع و

گسترده اطلاعات نادرست معتبر و جعلی عمیق فراهم می‌کند، یادگیری ماشینی می‌تواند برای شناسایی، تجزیه و تحلیل و ازبین بردن محتوای تبلیغاتی ناخواسته استفاده شود (Whyte, 2020: 202). هوش مصنوعی به عنوان یک ابزار اقتصادی و مالی دولت سازی، می‌تواند برای تقویت عملیات مقابله با سرمایه‌های غیرقانونی حامی تروریسم و سلاح‌های کشتار جمعی مورد استفاده قرار گیرد (Kirkos et al., 2007: 909).

در زمینه دفاع، سیستم‌های هوش مصنوعی نیز ممکن است با تسهیل تجزیه و تحلیل داده‌ها برای پیش‌بینی مولفه‌های مختلف مانند تعمیر و نگهداری تجهیزات، عملکرد اعضای سرویس و موارد دیگر در لجستیک مورد استفاده قرار گیرد (Taddeo, 2019: 189). علاوه بر این با توجه به پذیرش بالقوه هوش مصنوعی در نبرد، سیستم‌های خودمختار به طور فزاینده‌ای در سیستم‌های نظامی در سراسر جهان با سیستم‌های تسلیحات خودکار مرگبار که قادر به شناسایی و از بین بردن هدف به طور مستقیم و بدون دخالت انسانی هستند در حال گسترش می‌باشد (Haner & Garcia, 2019: 335).

فراخوان برای حاکمیت دیجیتال اتحادیه اروپا که نیاز به ایجاد فضای سایبری آزاد و ایمن و در عین حال انعطاف‌پذیر را شناسایی می‌کند، به طور فزاینده‌ای در مرکز دستور کار اتحادیه اروپا قرار دارد. این موضوعات با بلندپروازی گسترده اتحادیه اروپا برای دستیابی به «استقلال استراتژیک» که برای اولین بار سال ۲۰۱۳ در رابطه با صنعت دفاعی اعلام شد سازگار است (European Council, 2013: 1). امری که بعداً برای اهداف دفاعی و امنیتی گسترده‌تر با راه‌اندازی اتحادیه جهانی اتحادیه اروپا به تصویب رسید.

ادعای نوظهور حاکمیت دیجیتال با بحث‌های پرجنب و جوش در مورد نقش هوش مصنوعی از جمله در مورد داده کاوی، پاسخگویی الگوریتم‌ها و رهبری در صنعت فناوری همراه شده است که به یکی از ویژگی‌های برجسته در تلاش اروپا برای دستیابی به استقلال استراتژیک تبدیل شده است. قطب‌نمای استراتژیک ۲۰۲۲، جاه‌طلبی اتحادیه اروپا، برای تبدیل شدن به یک رهبر جهانی در هوش مصنوعی را با کاهش وابستگی به بازیگران خارجی برای فناوری‌های نوظهور، افزایش تولید پردازنده‌های رایانه‌ای با کارایی بالا و ایجاد یک فضای داده مستقل مشخص می‌کند. به علاوه، این سند نشان می‌دهد که چگونه

هوش مصنوعی در اتحادیه اروپا جزء کلیدی استراتژی امنیتی و دفاعی جدید اروپا محسوب می‌شود و حتی می‌تواند به عنوان بخشی از سیستم‌های تسلیحاتی آینده اثرگذار باشد (European External Action Service, 24 March 2022). اتحادیه اروپا در زمینه حاکمیت دیجیتال، اغلب زیر چتر یک ابتکار گسترده‌تر در مورد «استقلال استراتژیک» خود را تعریف می‌کند (Timmers, 2019: 636). بطور مثال مفهوم خودمختاری استراتژیک در استراتژی جهانی اتحادیه اروپا در سال ۲۰۱۶ معرفی شد که در آن حاکمیت عمدتاً در چارچوب موضوعات نظامی، امنیتی و دفاعی مورد استفاده قرار گرفت (European External Action Services, 2016: 18).

تا سال ۲۰۲۱ حدود ۴۸ کشور استراتژی‌ها یا برنامه‌های ملی رسمی برای هوش مصنوعی وضع کرده‌اند که ۲۱ مورد از آنها عضو اتحادیه اروپا هستند (OECD.AI, 2021). در حالی که هر دولتی قصد دارد توانایی‌های هوش مصنوعی خود را ایجاد کند و رشد اقتصادی را از طریق یک محیط امن و اخلاقی تقویت کند، کاربرد هوش مصنوعی در جهت ادغام در سیستم‌های امنیت ملی و دفاع ملی ضروری است. در زمان تصویب معاهده اتحادیه اروپا، کشورهای عضو اتحادیه اروپا بر لزوم تعهد جمعی به ایجاد یک سیاست دفاعی منسجم توافق کرده‌اند. با این حال از نظر تاریخی، نهادهای اتحادیه اروپا دسترسی محدودی به مسائل دفاعی داشته‌اند. در حالی که برخی از اعضای اتحادیه اروپا به دنبال همکاری و تعامل بیشتر در مورد نقش هوش مصنوعی برای دفاع از طریق راهبردها و ابتکارات ملی خود هستند، برخی دیگر نسبت به مشخص کردن موضع خود در مورد هوش مصنوعی و استفاده دفاعی از آن در بودجه ملی و برنامه‌ریزی دفاعی محتاط هستند. در همین حال در اوایل سال ۲۰۲۲ کمیسیون اروپا نقشه راه فناوری‌های حیاتی را منتشر کرد که معتقد است توسعه آنها چقدر در تاثیرگذاری بر چشم‌انداز امنیت جهانی مهم است (European Commission, 14 December 2022a).

رویکرد مبهم ارائه شده نسبت به هوش مصنوعی در حوزه امنیت، بازتابی از اختلاف نظرهای دیرینه درباره ضرورت تدوین یک راهبرد دفاعی مشترک و منحصر به فرد برای اتحادیه اروپا از سوی کشورهای عضو است. از نظر تاریخی، هدف اتحادیه اروپا

بهبود اقتصاد کشورهای عضو برای دستیابی به سطحی از وابستگی متقابل است که اجازه توسعه آرزوهای نظامی فردی را نمی‌دهد. در نتیجه، توسعه فناوری و جاه‌طلبی‌ها برای ایجاد یک قابلیت فناوری فراملی در زمینه اتحادیه اروپا از فقدان دفاع یکپارچه اروپا در میان کشورهای عضو آن بسیار ضربه خورده است (Citi, 2014: 145). لذا به نظر می‌رسد فقدان رویکرد اتحادیه اروپا در زمینه هوش مصنوعی، شکاف قابلیت‌های فناوری بین کشورهای عضو اتحادیه اروپا را افزایش می‌دهد.

دیدگاه اتحادیه اروپا در مورد مقررات و حمایت از حقوق فردی به عنوان مانعی در برابر توسعه هوش مصنوعی در نظر گرفته شده است، چرا که نوآوری داده‌ها را با مشکل مواجه می‌کند (Brattberg et al., 9 July 2020). اگر اتحادیه اروپا مقررات خود را با وضوح و راهنمایی خوب اجرا کند روند فوق قابلیت اصلاح خواهد داشت (Roberts et al., 2021). پارلمان اروپا به طور فزاینده‌ای بر اهمیت همکاری با دیگر بازیگران همفکر در پاسخ به جنبه‌های نظامی و امنیتی هوش مصنوعی تاکید می‌کند (European Parliament, 03 May 2022). این امر به ویژه در حوزه امنیت سایبری مشهود است، جایی که اتحادیه اروپا در حال افزایش هدف خود برای ایفای نقش محوری در همکاری‌های بین‌المللی است (Calderaro, 2021: 400).

در کل، اتحادیه اروپا با گسترش رویکرد دیپلماتیک سایبری خود فراتر از حوزه امنیت، فضای مشارکت بیشتر را تسهیل می‌کند. در این زمین، یکی از همکاری مهم اتحادیه اروپا در زمینه هوش مصنوعی، توسعه روابط با برزیل در زمینه اقتصاد دیجیتال از طریق گفتگوی اقتصاد دیجیتال است (European Commission, 21 March 2021d). همکاری اروپا با سنگاپور نیز حکایت از توجه طرفین به نقش هوش مصنوعی در تعاملات اقتصادی متقابل دارد (European Commission, 14 February 2022b). مثال مهم دیگر از توجه اروپا به همکاری میان کشورها، مشارکت دیجیتال رسمی این اتحادیه با ژاپن است. طرفین بر این نکته توافق کردند که همکاری بین دو بازیگر در زمینه هوش مصنوعی، باید بر کاربردهای ایمن و اخلاقی این فناوری متمرکز شود (European Commission, 12 May 2022c).

نتیجه‌گیری

پیشنهاد کمیسیون اروپا برای تنظیم و قانون‌گذاری در حوزه هوش مصنوعی در ۲۱ آوریل ۲۰۲۱ که به عنوان "قانون هوش مصنوعی" شناخته می‌شود، نشان‌دهنده یک تلاش جامع و نظام‌مند برای مدیریت مخاطرات و بهره‌گیری از فرصت‌های این فناوری است. نتایج این پژوهش حاکی از آن است که اتحادیه اروپا با تکیه بر قدرت هنجاری خود درصدد است به سیاستگذاری در حوزه هوش مصنوعی در قالب ارزش‌هایی همچون حقوق بشر، دموکراسی، حاکمیت قانون و توسعه پایدار اقدام نماید. اتحادیه اروپا از ابزارهای متنوعی نظیر دیپلماسی، سیاست‌های توسعه‌ای، توافقات تجاری و شرایط الحاق برای ترویج این هنجارها و استانداردها در سطح بین‌المللی بهره می‌گیرد. قوانین پیشنهادی در زمینه هوش مصنوعی در اتحادیه اروپا بر اصولی نظیر شفافیت، پاسخگویی، نظارت انسانی، ایمنی و حفاظت از داده‌ها تأکید دارند. این قوانین، چارچوبی قانونی و جامع را برای ارزیابی، نظارت و تضمین انطباق با استانداردهای اخلاقی و حقوقی فراهم می‌آورند. چنین رویکردی، نه تنها به کاهش ریسک‌های مرتبط با هوش مصنوعی کمک می‌کند، بلکه زمینه‌ساز افزایش اعتماد عمومی به این فناوری می‌شود.

در نهایت می‌توان نتیجه گرفت که اتحادیه اروپا با تدوین و اجرای این قوانین به دنبال تبدیل شدن به یک قدرت هنجاری در عرصه بین‌المللی است. اتحادیه اروپا قصد دارد با تنظیم مقررات دقیق و جامع، هم امنیت و حکمرانی دیجیتال را تقویت کند و هم از مزایای بالقوه هوش مصنوعی بهره‌مند شود. رویکرد اتحادیه اروپا می‌تواند الگویی برای سایر کشورها و مناطق جهان باشد و به استانداردسازی و مدیریت بهتر هوش مصنوعی در سطح جهانی کمک کند. این پژوهش نشان می‌دهد که با اتخاذ راهبردهای مناسب و متناسب با شرایط، می‌توان از فرصت‌های هوش مصنوعی بهره‌برداری کرد و در عین حال، مخاطرات و چالش‌های آن را نیز به حداقل رساند. رویکرد اتحادیه اروپا در زمینه سیاستگذاری در خصوص هوش مصنوعی، نمونه‌ای از تلاش برای ایجاد تعادل بین نوآوری و امنیت و ترویج ارزش‌های اخلاقی و حقوقی در استفاده از فناوری‌های نوین است.

تعارض منافع
تعارض منافع ندارم.

ORCID

Mostafa Kaka



<https://orcid.org/0009-0007-6505-2230>

Mokhtar Salehi



<https://orcid.org/0000-0003-2611-4420>

References

- Baldwin, R., Cave, M., & Lodge, M, (2011), “Understanding Regulation: Theory, Strategy, and Practice”, *Oxford University Press*, [https:// doi. org/ 10. 1093/ acprof: osobl/ 9780199576081. 001.0001](https://doi.org/10.1093/acprof:osobl/9780199576081.001.0001).
- Brattberg, E., Rugova, V & Csernaton, R, (2020), *Europe and AI: Leading, lagging Behind, or Carving Its Own Way?* (Vol. 9), Washington, DC, USA: Carnegie endowment for International Peace, [https:// carnegieendowment. org/ research/ 2020/ 07/europe-and-ai-leading-lagging-behind-or-carving-its-own-way?lang=en](https://carnegieendowment.org/research/2020/07/europe-and-ai-leading-lagging-behind-or-carving-its-own-way?lang=en).
- Brown, M, (2019), “Statement by Michael Brown, Director of the Defense Innovation Unit”, before the Senate Armed Services Committee Subcommittee on Emerging Threats and Capabilities Hearing On “Artificial Intelligence Initiatives within the Defense Innovation Unit”, *National Security Archive*, [https:// nsarchive. gwu. edu/ document/ 20136-national-security-archive-170-statement](https://nsarchive.gwu.edu/document/20136-national-security-archive-170-statement)
- Calderaro, A., (2021), *Diplomacy and Responsibilities in the Transnational Governance of the Cyber Domain*. In: H. Hansen-Magnusson, and A. Vetterlein, Eds. the Routledge Handbook of Responsibility in World Politics (Pp. 394–405). London: Routledge.

- Citi, M., (2014), “Revisiting Creeping Competences in the Eu: The Case of Security R&D Policy”, *Journal of European Integration*, Vol. 36, No. 2, (135–151).
- Com, E, (2021), “Laying Down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act) and Amending Certain Union Legislative Acts”, *Proposal for a Regulation of the European Parliament and of the Council*, [https://eur-lex.europa.eu/legal-content/ EN/ TXT/ ?uri= celex %3 A52021 PC0206](https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=celex%3A52021PC0206).
- Csernaton, R., (2022), “The Eu’s Hegemonic Imaginaries: From European Strategic Autonomy in Defence to Technological Sovereignty”, *European Security*, Vol. 31, No. 3, 395–414.
- Dunlop, C., (2023), *Policy Briefing - A Eu Ai Act That Works for People and Society*, Ada Lovelace Institute, United Kingdom, available at: Hyperlink Reference not valid. On29may2024. Cid: 20. 500. 12592/ S4c5xj.
- 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*. https://eur-ws.org/Vol-3221/IAIL_paper7.pdf.
- Ebers, M., Hoch, V. R., Rosenkranz, F., Ruschemeier, H., & Steinrötter, B, (2021), *The European Commission’s Proposal for an Artificial Intelligence Act—A Critical Assessment by Members of the Robotics and Ai Law Society (Rails). J*, Vol. 4, No. 4, 589-603.
- Emily Binchy, (2022), “Advancement or Impediment? Ai and the Rule of Law”, *The Institute of International and European Affairs*, [https:// www. iiea. com/ images/ uploads/ resources/ Advancement- or- Impediment-AI-and-the-Rule-of-Law.pdf](https://www.iiea.com/images/uploads/resources/Advancement-or-Impediment-AI-and-the-Rule-of-Law.pdf).
- European Centre for International Political Economy (Ecipe), (2022), “Strategic Autonomy and Artificial Intelligence:

European Union in the Global Race”, available at: [https:// Ecipe. Org/ Publications/ Strategic- Autonomy- And- Artificial- Intelligence- European-Union-In-The-Global-Race](https://ecipe.org/publications/strategic-autonomy-and-artificial-intelligence-european-union-in-the-global-race).

- European Commission, (2018b), “Artificial Intelligence for Europe”, *European Commission, Brussels*, [https:// eur- lex. Europa, eu/ legal-content/ EN/ TXT/ ? uri= COM%3 A2018%3 A237%3AFIN](https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=COM%3A2018%3A237%3AFIN).
- European Commission, (2022b), “Joint Statement: Eu and Singapore Agree to Accelerate Steps Towards a Comprehensive Digital Partnership”, *European Commission*, [https:// ec. europa. eu/ commission/ presscorner/ detail/ en/ statement_ 22_ 1024](https://ec.europa.eu/commission/presscorner/detail/en/statement_22_1024).
- European Commission, (2017), “Attitudes towards the Impact of Digitisation and Automation on Daily Life”, [https:// Digital- Strategy, Ec. Europa. Eu/ En/ News/ Attitudes- Towards-Impact- Digitisation-And-Automation-Daily-Life](https://digital-strategy.ec.europa.eu/en/news/attitudes-towards-impact-digitisation-and-automation-daily-life).
- European Commission, (2019), “Artificial Intelligence”, 7 January 2019 (25 February 2019), [https:// Digital- Strategy. Ec. Europa. Eu/ En/ Policies/ European- Approach- Artificial- Intelligence# %22](https://digital-strategy.ec.europa.eu/en/policies/european-approach-artificial-intelligence#%22).
- European Commission, (2021c), “Europe Fit for the Digital Age: Commission Proposes New Rules and Actions for Excellence and Trust in Artificial Intelligence”, Retrieved 11 January, 2022, available at: [https:// Ec. Europa.Eu/ Commi Ssion/ Press Corner/ Detail/ En/ Ip_ 21_ 1682](https://ec.europa.eu/commission/press-corner/detail/en/ip_21_1682).
- European Commission, (2021d), “Eu and Brazil to Reinforce Cooperation Ahead of 12th Digital Economy Dialogue”, *Shaping Europe’s Digital Future*, available at: [https:// digital- strategy. ec. europa. eu/ en/ news/ eu- and- brazil- strengthen- their-digital-cooperation](https://digital-strategy.ec.europa.eu/en/news/eu-and-brazil-strengthen-their-digital-cooperation).
- European Commission, (2022a), “Roadmap on Critical Technologies for Security and Defence”, Strasbourg, Text,

available at: <https://eurlex.europa.eu/legalcontent/en/txt/?uri=celex%3a52022dc0061>.

- European Commission, (2022c), “Eu-Japan Summit: Strengthening Our Partnership”, *European Commission*, available at: <https://digital-strategy.ec.europa.eu/en/news/eu-japan-summit-strengthening-our-partnership>.
- European Commission, (2019), “Communication from the Commission to the European Parliament, the Council, the European Economic and Social Committee and the Committee of the Regions - Building Trust in Human”, Com, 168, available at: <https://digital-strategy.ec.europa.eu/en/library/communication-building-trust-human-centric-artificial-intelligence>.
- European Council, (2013), “Conclusions of the European Council of 19/20 December 2013”, *Euco* 217/13, <https://data.consilium.europa.eu/doc/document/ST-217-2013-INIT/en/pdf>.
- European External Action Service, (2022), “A Strategic Compass for Security and Defence: for a European Union That Protects Its Citizens, Values and Interests and Contributes to International Peace and Security”, Brussels: the European External Action Service, <https://data.consilium.europa.eu/doc/document/ST-7371-2022-INIT/en/pdf>.
- European External Action Services. (2016). “Shared Vision, Common Action: A Stronger Europe, A Global Strategy for the European Union’s Foreign and Security Policy”, https://eeas.europa.eu/sites/eeas/files/eugs_review_web_0.pdf.
- European Parliament, (2021), “Guidelines for Military and Non-Military Use of Artificial Intelligence”, available at: <https://www.europarl.europa.eu/news/en/press-room/20210114ipr95627/guidelines-for-military-and-non-military-use-of-artificial-intelligence>.

- European Parliament, (2022), “Artificial Intelligence: Meps Want the Eu to be a Global Standard-Setter”, available at: <https://Www.Europarl.Europa.Eu/News/En/Press-Room/20220429ipr28228/Artificial-Intelligence-Meps-Want-The-Eu-To-Be-A-Global-Standard-Setter>.
- Farrand, B & Carrapico, H., (2022), “Digital Sovereignty and Taking Back Control: from Regulatory Capitalism to Regulatory Mercantilism in Eu Cybersecurity”, *European Security*, Vol. 31, No. 3, 435–453.
- Floridi, L et al., (2018), “Ai4people— an Ethical Framework for a Good Ai Society: Opportunities, Risks, Principles, and Recommendations”, *Minds and Machines*, Vol. 28, No. 4, 689–707. Doi: 10.1007/S11023-018-9482-5.
- 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?”, available at: [https:// repository. ubn. ru. nl/ bitstream/ handle/ 2066/ 241998/241998.pdf](https://repository.ubn.ru.nl/bitstream/handle/2066/241998/241998.pdf).
- Haner, J., & Garcia, D., (2019). “The Artificial Intelligence Arms Race: Trends and World Leaders in Autonomous Weapons Development”, *Global Policy*, Vol. 10, No. 3, PP. 331–337.
- High Contracting Parties Ccw, (2019), “Meeting of the High Contracting Parties to the Convention on Prohibitions or Restrictions on the use of Certain Conventional Weapons Which May Be Deemed to Be Excessively Injurious or to Have Indiscriminate Effects ,*United Nations*, Geneva, 13-15, available at: [https://unoda-documents-library. s3. amazonaws. com/ Convention_ on_ Certain_ Conventional_ Weapons_ -_Fifth_ Review_ Conference_\(2016\)/Sri%2BLanka.pdf](https://unoda-documents-library.s3.amazonaws.com/Convention_on_Certain_Conventional_Weapons_-_Fifth_Review_Conference_(2016)/Sri%2BLanka.pdf).
- Kirkos, E., Spathis, C & Manolopoulos, Y., (2007), “Data Mining Techniques for the Detection of Fraudulent Financial

Statements”, *Expert Systems with Applications*, Vol. 32, No. 4, PP. 995–1003.

- Lambach, D & Monsees, L., (2022), “Digital Sovereignty, Geopolitical Imaginaries, and the Reproduction of European Identity”, *European Security*, Vol. 31, No. 3, PP. 377–394.
- Laux, J., Wachter, S & Mittelstadt, B, (2024), “Trustworthy Artificial Intelligence and the European Union Ai Act: On the Conflation of Trustworthiness and Acceptability of Risk”, *Regulation & Governance*, Vol. 18, No. 1, PP. 3-32.
- Macovei, M., & Yiannaki, S. M, (2020), “Strategic Autonomy in the Eu’s Ai Policy”, *European Papers*. Retrieved From available at: [https:// Www. Europeanpapers. Eu/ En/ Europeanforum/ Strategic- Autonomy-Eu%E2%80%99s-Ai-Policy](https://www.europeanpapers.eu/en/europeanforum/Strategic-Autonomy-Eu%E2%80%99s-Ai-Policy).
- Manners, I, (2002), “Normative Power Europe: A Contradiction in Terms?”, *Jcms: Journal of Common Market Studies*, Vol. 40, No. 2, PP. 235-258.
- Marsden, C., Meyer, T & Brown, I., (2020), “Platform Values and Democratic Elections: How Can the Law Regulate Digital Disinformation?”, *Computer Law & Security Review*, Vol. 36, 105373. <https://doi.org/10.1016/j.clsr.2019.105373>.
- Mckinsey Global Institute, (2017), “10 Imperatives for Europe in the Age of Ai and Automation”, *Mckinsey Global Institute*, available at: [https:// Www. Mckinsey. Com/ Featured- Insights/ Europe/ ten- Imperatives- for- Europe- in- the- Age- of- Ai-and- Automation](https://www.mckinsey.com/featured-insights/Europe/ten-Imperatives-for-Europe-in-the-Age-of-Ai-and-Automation).
- Meltzer, J & Tielemans, A., (2022), *The European Union AI Act: Next steps and issues for building international cooperation in AI*, Brookings Institution. United States of America, available at: <https://coilink.org/20.500.12592/tpb2j8>.

- Micklitz, H. W., et al., (2021), *Constitutional Challenges in the Algorithmic Society*, Cambridge: Cambridge University Press, available at: [https:// cadmus. eui. eu/ bitstream/ handle/ 1814/ 74296/ Constitutional_ Challenges_ in_ the_ Algorithmic_ Society. pdf](https://cadmus.eui.eu/bitstream/handle/1814/74296/Constitutional_Challenges_in_the_Algorithmic_Society.pdf).
- Now, A, (2020), Europe's Approach to Artificial Intelligence: How Ai Strategy is Evolving. available at: [https:// www. accessnow. org/ wp- content/ uploads/ 2020/12/ Europes- approach- to-AI-strategy-is-evolving.pdf](https://www.accessnow.org/wp-content/uploads/2020/12/Europes-approach-to-AI-strategy-is-evolving.pdf).
- Oecd.Ai, (2021), "Database of National Ai Policies", available at: <https://Www.Oecd.Ai/Dashboards>.
- Pigou, A, (2002), *The Economics of Welfare* (1st ed.), Routledge, New York: Taylor and Francis. [https:// doi. org/ 10. 4324/ 9781351304368](https://doi.org/10.4324/9781351304368).
- Posner, R. A, (1974), *Theories of Economic Regulation* (No. w0041). National Bureau of Economic Research", New York: [https:// www. nber. org/ system/ files/ working_ papers/ w0041/ w0041.pdf](https://www.nber.org/system/files/working_papers/w0041/w0041.pdf).
- Roberts, H et al., (2021), "Achieving a "Good Ai Society": Comparing the Aims and Progress of the Eu and the Us", *Science and Engineering Ethics*, Vol. 27, No. 6, PP. 1–25, Doi: 10.1007/S11948-021-00340-7.
- Robles, P & Mallinson, D. J, (2023), "Artificial Intelligence Technology, Public Trust, and Effective Governance", *Review of Policy Research*, Vol. 42, No. 1, 11-28. [https:// Doi. Org/ 10.1111/Ropr.12555](https://Doi.Org/10.1111/Ropr.12555).
- Sovrano F et al., (2022), "Explainability and the European AI Act Proposal", *J — Multidisciplinary Scientific Journal*, Vol. 5, No. 1, 126-138. <https://doi.org/10.3390/j5010010>.

- Special Eurobarometer, (2017), “Attitudes towards the Impact of Digitization and Automation on Daily Life”, Special Eurobarometer 460 – Wave EB87.1 – TNS opinion & social. <https:// europa.eu/eurobarometer/surveys/detail/2160>.
- Stigler, G. J, (2021), “The Theory of Economic Regulation”, In *the Political Economy: Readings in the Politics and Economics of American Public Policy*, Chicago: Routledge, PP. 67-81.
- Straus, J, (2021), “Artificial Intelligence–Challenges and Chances for Europe”, *European Review*, Vol. 29, No. 1, PP. 142-158. doi: 10.1017/S1062798720001106.
- Stone, P et al., (2022), “Artificial Intelligence and Life in 2030: the one Hundred Year Study on Artificial Intelligence”, Cornell University, <https:// doi. org/ 10.48550/arXiv.2211.06318>
- Sun Tq & Medaglia R, (2019), “Mapping the Challenges of Artificial Intelligence in the Public Sector: Evidence from Public Healthcare”, *Government Information Quarterly*, Vol. 36, No. 2, PP. 368–383. Doi: 10.1016/J.Giq.2018.09.008.
- Russell, S. J & Norvig, P, (2016), “Artificial Intelligence: A Modern Approach”, Pearson, *Pearson Education*, Inc., Upper Saddle River, New Jersey 07458. https:// people. engr. tamu. edu/ guni/ csce642/ files/ AI_Russell_Norvig.pdf.
- Szczepanski, M., (2019), “Economic Impacts of Artificial Intelligence (Ai), Eprs: European Parliamentary Research Service Belgium”, availabled at: <https:// Policycommons. Net/ Artifacts/ 1334867/ Economic- Impacts- of- Artificial- Intelligence- Ai/1940719/ On 01 May 2024. Cid: 20.500.12592/0phht4>.
- Taddeo, M., (2019), “Three Ethical Challenges of Applications of Artificial Intelligence in Cybersecurity”, *Minds and Machines*, Vol. 29, No. 2, PP. 187–191.

- Timmers, P., (2019), "Ethics of Ai and Cybersecurity When Sovereignty is at Stake", *Minds and Machines*, Vol. 29, No. 4, PP. 635–645.
- Wachter, S, (2024), "Limitations and Loopholes in the EU AI Act and AI Liability Directives: What this Means for the European Union, the United States, and Beyond", *Yale Journal of Law & Technology*, 26(3), 671–705. [https:// law. yale. edu/ sites/ default/ files/ area/ center/ isp/ documents/ wachter_26yalejltech671.pdf](https://law.yale.edu/sites/default/files/area/center/isp/documents/wachter_26yalejltech671.pdf)
- Whyte, C., (2020), "Deepfake News: Ai-Enabled Disinformation as a Multi-Level Public Policy Challenge", *Journal of Cyber Policy*, Vol. 5, No. 2, PP. 199–217.

استناد به این مقاله: کا کا، مصطفی و صالحی، مختار، (۱۴۰۴)، «راهبرد استقلال استراتژیک اتحادیه اروپا و هوش مصنوعی خوب»، *پژوهش‌های راهبردی سیاست*، دوره ۱۴، شماره ۵۴، ۳۱۳–۳۴۹.

Doi: 10.22054/QPSS.2024.80662.3434



Quarterly of Political Strategic Studies is licensed under a Creative Commons Attribution-Non Commercial 4.0 International License