

فرصت‌ها و چالش‌های هوش مصنوعی در سیاست‌گذاری عادلانه توزیع درآمد (مطالعه موردی: بهینه‌سازی مالیات با شبیه‌سازی عامل‌محور)

مسعود صالحی رزوه*

Doi: <https://doi.org/10.22096/esp.2025.2064462.1827>

[تاریخ دریافت: ۱۴۰۴/۰۴/۰۹ - تاریخ پذیرش: ۱۴۰۴/۰۹/۲۲]

چکیده

در شرایطی که نابرابری درآمدی، ثبات اجتماعی و رشد پایدار را به چالش کشانده‌اند و نظام‌های سنتی توزیع درآمد ناکارآمد بوده‌اند، هوش مصنوعی به‌مثابه پدیده‌ای دوگانه، می‌تواند ابزارهای قدرتمندی برای سیاست‌گذاری عادلانه‌تر فراهم کند و همزمان با تشدید تغییرات فناورانه مهارت‌محور، نابرابری را افزایش دهد. مقاله ضمن تبیین چهارچوب فلسفی عدالت در طراحی الگوریتم‌های منصفانه به بررسی چالش‌هایی مانند سوگیری الگوریتمی، نقض حریم خصوصی و عدم شفافیت پرداخته است و بر ضرورت ادغام اصول عدالت توزیعی در طراحی الگوریتم‌ها، مشارکت ذی‌نفعان، رویکرد انسان در حلقه و حفاظت از حریم خصوصی تأکید دارد. با تمرکز بر نقش هوش مصنوعی در جایگاه ابزاری برای سیاست‌گذاری، کاربرد آن در بهینه‌سازی بده‌بستان میان بهره‌وری و برابری از طریق مدل شبیه‌سازی عامل‌محور نشان داده می‌شود. با استفاده از آزمایشگاه اقتصادی مجازی شامل چهار عامل هوشمند (خانوارها) و برنامه‌ریز اجتماعی (دولت)، سناریوهای مالیاتی صفر، سنگین و متعادل آزمون شد. نتایج نشان داد سیاست متعادل در مقایسه با مالیات صفر با کاهش ۴۰۸۳ درصدی بهره‌وری، برابری را ۶۰۱۸ درصد بهبود بخشید؛ همچنین در مقایسه با سیاست مالیات سنگین با فدا کردن ۱۰۹ درصد از برابری، بهره‌وری را ۱۲۰۲ درصد افزایش داد. این بده‌بستان، رفاه اجتماعی را ۰۰۷۹ درصد نسبت به مالیات صفر و ۱۰۰۲۷ درصد نسبت به مالیات سنگین افزایش داد. این یافته بیانگر اهمیت هوش مصنوعی به‌منزله آزمایشگاه سیاستی برای طراحی سیاست‌های هوشمند است.

واژگان کلیدی: نابرابری درآمد؛ هوش مصنوعی؛ سیاست‌گذاری هوشمند؛ بهینه‌سازی مالیات؛ شبیه‌سازی عامل‌محور.

طبقه‌بندی موضوعی: H21, O33, C63, D63.

۱. مقدمه

در عصر حاضر، نابرابری درآمدی به یکی از پیچیده‌ترین و فراگیرترین چالش‌های جوامع بشری تبدیل شده است. درحالی‌که رشد اقتصادی جهانی به شکل بی‌سابقه‌ای شتاب گرفته، نابرابری درآمدی نیز به موازات آن گسترش یافته است. گزارش‌های سازمان‌های بین‌المللی مانند آزمایشگاه نابرابری جهانی (World Inequality Lab) حاکی از آن است که شکاف درآمدی در دو دهه گذشته به مرز هشداردهنده‌ای رسیده است به‌طوری‌که سهم ۵۰ درصد پایین جمعیت جهان از کل ثروت جهانی فقط ۲ درصد برآورد شده است؛ درحالی‌که سهم ۱۰ درصد بالا ۷۶ درصد است. بین سال‌های ۱۹۹۵ تا ۲۰۲۱، ۱ درصد بالای جامعه ۳۸ درصد از افزایش ثروت جهانی را به خود اختصاص دادند، درحالی‌که ۵۰ درصد پایین جامعه به طرز وحشتناکی ۲ درصد را به خود اختصاص دادند. سهم ثروت متعلق به ۱٪ درصد بالای جامعه در این دوره از ۷ درصد به ۱۱ درصد افزایش یافت و ثروت میلیاردرهای جهان به شدت افزایش یافت.^۱ این وضعیت، لزوم بازنگری در سازوکارهای توزیع ثروت را بیش از پیش نمایان می‌سازد.

نظام‌های سنتی توزیع درآمد، مانند سیستم‌های مالیاتی و یارانه‌ای، تاکنون نتوانسته‌اند این چالش‌ها را به‌طور مؤثر مدیریت کنند. این سیستم‌ها با مشکلاتی مانند شناسایی نادرست گروه‌های هدف، تأخیر در اجرا، هزینه‌های بالای اداری و کمبود شفافیت مواجه‌اند. در این شرایط، انقلاب دیجیتال و فناوری‌های پیشرفته‌ای نظیر هوش مصنوعی (Artificial Intelligence) (AI)، یادگیری ماشین و کلان‌داده‌ها فرصت‌های جدیدی برای مقابله با این معضل فراهم کرده‌اند. قابلیت‌های منحصر به فرد هوش مصنوعی در پردازش داده‌های بزرگ و ناهمگن، شناسایی الگوهای پیچیده و پیش‌بینی دقیق رفتارهای اقتصادی، می‌تواند پایه‌ای علمی و عادلانه برای سیاست‌گذاری‌های توزیع درآمد ایجاد کند.

هوش مصنوعی با تحلیل داده‌های چندبعدی مانند درآمد، الگوی مصرف، دارایی‌ها، سطح تحصیلات و رفتارهای مالی، تصویری دقیق از وضعیت اقتصادی خانوارها ترسیم کرده است که می‌تواند مبنایی برای تخصیص بهینه و عادلانه منابع محدود باشد. علیرغم این پتانسیل، با چالش‌های جدی نظیر حفظ حریم خصوصی، امنیت داده‌ها، شفافیت الگوریتمی، سوگیری‌های سیستماتیک و جلب اعتماد عمومی روبه‌رو است. مقاله ضمن

1. Lucas Chancel et al., "World Inequality Report 2022," *World Inequality Lab* (2022): 3.

بررسی این فرصت‌ها و چالش‌ها به امکان‌سنجی طراحی سیاست‌گذاری عادلانه توزیع درآمد می‌پردازد و با تمرکز بر مطالعه موردی به این پرسش محوری می‌پردازد که هوش مصنوعی چگونه می‌تواند بده‌بستان میان بهره‌وری اقتصادی و برابری درآمد را از راه سیاست‌گذاری مالیاتی بهینه سازد. استدلال اصلی آن است که شناسایی سیاست مالیاتی متعادل از طریق هوش مصنوعی با بهینه‌سازی این بده‌بستان، سطح رفاه اجتماعی را در مقایسه با سیاست‌های افراطی (مالیات صفر یا سنگین) به حداکثر می‌رساند.

در ادامه، بخش دوم به مرور چهارچوب نظری پژوهش و مطالعات تجربی مرتبط می‌پردازد. بخش سوم معرفی الگو و روش‌شناسی پژوهش را شرح می‌دهد. در بخش چهارم ضمن جمع‌بندی یافته‌ها پیشنهادها، سیاستی اجرایی و مسیرهای پژوهشی آینده ارائه می‌شود.

۲. ادبیات نظری پژوهش

۲-۱. فرصت‌ها و چالش‌های استفاده از هوش مصنوعی در حوزه عدالت توزیعی

استفاده از هوش مصنوعی در حوزه عدالت توزیعی به معنای بهره‌گیری از الگوریتم‌ها و فناوری‌های هوش مصنوعی برای تخصیص منصفانه منابع و کاهش نابرابری‌ها در جامعه است. این فناوری، پتانسیل بالایی برای بهبود سیستم‌های توزیع منابع دارد؛ در عین حال با چالش‌های مهمی نیز مواجه است. در ادامه به برخی فرصت‌ها و چالش‌های مهم استفاده از هوش مصنوعی در عدالت توزیعی در قالب جداول (۱) و (۲) اشاره می‌شود.

جدول (۱): فرصت‌های هوش مصنوعی در عدالت توزیعی

عنوان	شرح
دسترسی عادلانه به اطلاعات	هوش مصنوعی می‌تواند با تجمیع و ساده‌سازی اطلاعات، هزینه‌های تحقیق را کاهش دهد و دسترسی به دانش را برای گروه‌های کم‌برخوردار برابرتر کند؛ همچنین با شناسایی منابع معتبر و استخراج اجماع در مورد هر موضوع، کمبود تخصص را جبران و به کاربران و کسب‌وکارها در محیط‌های کم‌منبع کمک می‌کند تا به

اطلاعاتی دست یابند که معمولاً در محیط‌های با منابع بالا در دسترس است. ^۲	
با تحلیل داده‌های کلان، هوش مصنوعی به شناسایی الگوهای پنهان نابرابری درآمدی کمک و امکان ارزیابی اثربخشی سیاست‌ها را برای طراحی مداخلات مؤثرتر فراهم می‌کند. ^۳	تصمیم‌گیری مبتنی بر داده (Data-Driven Decision Making)
الگوریتم‌های هوش مصنوعی می‌توانند تخصیص منابع محدود (پول، انرژی و زمان) را بهینه‌سازی کنند و با افزایش کارایی به ارائه خدمات بهتر با هزینه کمتر منجر شوند. ^۴	بهینه‌سازی تخصیص منابع
دولت‌ها می‌توانند از یادگیری ماشین برای بهبود فرایندهای نظارتی، شناسایی تخلفات قانونی و اجرای بهینه مقررات استفاده کنند. ^۵	تنظیم‌گری هوشمند

جدول (۲): چالش‌های هوش مصنوعی در عدالت توزیعی

عنوان	شرح
سوگیری الگوریتمی (Algorithmic Bias) و تبعیض	سیستم‌های هوش مصنوعی با یادگیری از داده‌های تاریخی که خود منعکس‌کننده نابرابری‌های گذشته هستند، ممکن است این نابرابری‌ها را بازتولید و تشدید کنند. ^۶
عدم شفافیت (Lack of Transparency)	عملکرد درونی بسیاری از مدل‌های پیشرفته هوش مصنوعی مانند جعبه سیاه است. این عدم شفافیت، درک نحوه تصمیم‌گیری و اعتمادسازی را دشوار می‌کند و پاسخگویی را کاهش می‌دهد. ^۷
کیفیت داده‌ها	عملکرد الگوریتم‌ها به شدت به کیفیت داده‌های آموزشی وابسته است. داده‌های ناقص، جانبدارانه یا غیردقیق منجر به نتایج و پیش‌بینی‌های ناعادلانه خواهد شد.

2. Valerio Capraro et al., "The Impact of Generative Artificial Intelligence on Socioeconomic Inequalities and Policy Making," *PNAS Nexus* 3, no. 6 (2024): 6.
3. Sekiswa Peter and Namuyanga Rebecca, "Data-Driven Approaches to Addressing Income Inequality: A Case Study of Nansana Municipality, Uganda," *Metropolitan Journal of Business & Economics (MJBE)* 3, no. 10 (2024): 22.
4. Stuart Russell and Peter Norvig, *Artificial Intelligence: A Modern Approach*, 3rd edition (Hoboken, NJ: Pearson, 2009), 401-410.
5. Capraro et al., "The Impact of Generative Artificial," 12.
6. Maximilian Kasy and Rediet Abebe, "Fairness, Equality, and Power in Algorithmic Decision-Making", in *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency (FAccT '21)* (New York: Association for Computing Machinery, 2021), 578.
7. Zachary C. Lipton, "The Mythos of Model Interpretability: In Machine Learning, the Concept of Interpretability is Both Important and Slippery," *Queue* 16, no. 3 (2018): 42-45.

شد. ^۸	
سوء استفاده از داده‌ها و حریم خصوصی	جمع‌آوری حجم عظیم داده‌های شخصی می‌تواند به نقض حریم خصوصی و پدیده‌ای تحت عنوان سرمایه‌داری نظارتی ^۹ منجر شود که در آن شرکت‌ها از داده‌ها داده‌ها برای کسب مزیت‌های ضد رقابتی استفاده می‌کنند. ^{۱۰}
تقویت نابرابری‌های موجود با خطا در طراحی	اگر در طراحی الگوریتم‌ها به مفاهیم عمیق عدالت اجتماعی توجه نشود، این فناوری می‌تواند ناخواسته نابرابری‌ها را تقویت کند. برای تضمین انصاف در تصمیم‌گیری‌های الگوریتمی، جنبه‌های کلیدی زیر باید مد نظر قرار گیرند: انتخاب معیار مناسب برای عدالت که فراتر از توزیع منابع مادی به تفاوت‌های فردی و اجتماعی (معلولیت‌ها) حساس باشد؛ جلوگیری از بازتولید یا تشدید نابرابری‌های ساختاری و تاریخی ریشه‌دار در داده‌های آموزشی، توجه به نابرابری‌های درون‌گروهی و پیچیدگی‌های تقاطع‌گرایی هویت‌ها، شفاف‌سازی نقش قدرت و کنترل نهادهای مسئول در طراحی و تنظیم الگوریتم‌ها برای جلوگیری از اعمال منافع خاص، مشارکت و مشورت فعال با ذینفعان به‌ویژه گروه‌های آسیب‌پذیر برای لحاظ کردن نیازها و نگرانی‌های آنها، ملاحظات جدی حریم خصوصی و حفاظت از داده‌ها با تکنیک‌هایی مانند حریم خصوصی تفاضلی ^{۱۱} برای جلوگیری از سوء استفاده و تبعیض و در نهایت، ایجاد قابلیت پاسخگویی و مسئولیت‌پذیری در صورت بروز خطا یا تبعیض در عملکرد الگوریتم‌ها. ^{۱۲}

۲-۲. تأثیر هوش مصنوعی مولد بر نابرابری اجتماعی-اقتصادی

هوش مصنوعی به‌ویژه هوش مصنوعی مولد (Generative AI)، تأثیرات دوگانه بر سه حوزه

8. Ninareh Mehrabi et al., "A Survey on Bias and Fairness in Machine Learning," *ACM Computing Surveys* 54, no. 6 (2021): 3.

۹. سرمایه‌داری نظارتی (Surveillance Capitalism) به شیوه‌ای اشاره می‌کند که در آن شرکت‌ها و سازمان‌ها از داده‌های شخصی کاربران برای کسب سود استفاده می‌کنند. در این مدل، شرکت‌ها (غول‌های فناوری) اطلاعات کاربران را از طریق فعالیت‌های آنلاین، جست‌وجوها، خریده‌ها و حتی تعاملات روزمره جمع‌آوری می‌کنند؛ سپس این داده‌ها تحلیل و برای پیش‌بینی رفتار کاربران یا حتی تأثیرگذاری بر تصمیمات آنها استفاده می‌شود. این پیش‌بینی‌ها معمولاً به تبلیغ‌کنندگان یا سایر شرکت‌ها فروخته می‌شود. این موضوع نگرانی‌هایی درباره حریم خصوصی، آزادی فردی و تأثیرات اجتماعی ایجاد کرده است.

10. Capraro et al., "The Impact of Generative Artificial," 2.

۱۱. حریم خصوصی تفاضلی (Differential Privacy) روشی است که با افزودن نویز تصادفی به داده‌ها یا نتایج تحلیل‌ها از شناسایی اطلاعات فردی در داده‌های آماری جلوگیری می‌کند. این روش تضمین می‌کند که حضور یا عدم حضور فرد در داده‌ها تأثیر قابل‌توجهی بر خروجی تحلیل نداشته باشد و در نتیجه حریم خصوصی افراد حفظ شود.

12. Maximilian Kasy and Rediet Abebe, "Fairness, equality, and power in algorithmic," 578.

کلیدی اجتماعی-اقتصادی شامل بازار کار، آموزش و بهداشت دارد. در بازار کار درحالی‌که فناوری‌های دیجیتال پیشین به نفع کارگران با مهارت بالا و به ضرر کارگران کم‌مهارت عمل کرده‌اند^{۱۳} (تغییر فناوری مهارت‌محور (Skill-biased Technological Change))، شواهد اولیه نشان می‌دهد که هوش مصنوعی مولد می‌تواند بهره‌وری کارگران کم‌تجربه یا کم‌مهارت را به شکل چشمگیری افزایش دهد که این پدیده سوگیری معکوس مهارتی (inverse skill-bias) شناخته می‌شود. در حوزه آموزش، هوش مصنوعی مولد فرصت‌های بی‌مانند برای شخصی‌سازی یادگیری و کاهش شکاف‌های آموزشی ایجاد می‌کند، اما چالش دسترسی نابرابر به فناوری (شکاف دیجیتال) می‌تواند این مزایا را محدود کند. در نهایت در بخش بهداشت، هوش مصنوعی پتانسیل افزایش دسترسی به مراقبت‌ها و بهبود اثربخشی خدمات را دارد؛ درحالی‌که حفظ حریم خصوصی بیماران و مسائل اخلاقی همچنان چالش‌برانگیز هستند. در ادامه به فرصت‌ها و چالش‌های هوش مصنوعی در بازارهای مختلف در قالب جدول (۳) می‌پردازیم:

جدول (۳): چالش‌ها و فرصت‌های هوش مصنوعی در بازارهای مختلف

فرصت‌ها	چالش‌ها	حوزه
۱. افزایش بهره‌وری کارگران کم‌مهارت (سوگیری معکوس مهارتی): ابزارهای هوش مصنوعی مانند دستیارهای چت و برنامه‌نویسی می‌تواند بهره‌وری کارگران کم‌تجربه یا کم‌مهارت را به‌طور چشمگیری افزایش دهد.	۱. تشدید نابرابری (SBTC): فناوری‌های دیجیتال پیشین مانند رایانه‌ها و ربات‌های صنعتی بیشتر به نفع کارگران با مهارت بالا و به ضرر کارگران کم‌مهارت عمل کرده‌اند.	بازار کار
۲. نقیض مکمل و ایجاد مشاغل ارزش‌آفرین: هوش مصنوعی مانند طراحی هوشمند می‌تواند مکمل نیروی انسانی باشد، به افزایش کیفیت محصول کمک برساند و با توانمندسازی کارگران، وظایف و مشاغل جدید با درآمد مناسب ایجاد کند. ^{۱۴}	۲. خطر اتوماسیون بدون ایجاد شغل: هوش مصنوعی ممکن است روند اتوماسیون را تسریع کند، بدون اینکه شغل‌های جدید و باکیفیت برای کارگران بدون تحصیلات دانشگاهی ایجاد کند.	

13. Giovanni L. Violante, "Skill-Biased Technical Change", in *The New Palgrave Dictionary of Economics*, eds. by Steven N. Durlauf and Lawrence E. Blume (London: Palgrave Macmillan, 2008), 2.

14. Capraro et al., "The Impact of Generative Artificial," 4-6.

آموزش	شکاف دیجیتال و دسترسی نابرابر: دسترسی نابرابر به فناوری‌های پیشرفته و اینترنت پرسرعت (به‌ویژه در مناطق کم‌برخوردار)، می‌تواند نابرابری‌های آموزشی موجود را تشدید کند.	۱. آموزش شخصی‌سازی‌شده و سازگار: هوش مصنوعی مولد (چت‌بات‌ها) می‌توانند در جایگاه مربیان مجازی شخصی‌سازی‌شده عمل کنند و یادگیری را با نیازهای فردی سازگار سازند. ۲. تقویت مهارت‌ها و کاهش شکاف‌ها: مهارت‌هایی مانند مهندسی فرامین، تفکر انتقادی و ایده‌پردازی خلاقانه و کمک به کاهش شکاف‌های آموزشی به‌ویژه در کلاس‌های بزرگ تقویت می‌شود.^{۱۵}
مراقبت‌های بهداشتی	۱. چالش حریم خصوصی در داده‌های پزشکی: حفظ حریم خصوصی بیماران و حفاظت از اطلاعات حساس پزشکی آنان، یکی از مشکلات اساسی و پیچیده در این حوزه محسوب می‌شود. ۲. نیاز به داده‌های باکیفیت: دسترسی به داده‌های باکیفیت و جامع، پیش‌نیازی اساسی برای آموزش مؤثر و عملکرد دقیق مدل‌های هوش مصنوعی به شمار می‌رود. ۳. مسائل اخلاقی تصمیم‌گیری خودکار: تصمیم‌گیری‌های خودکار توسط هوش مصنوعی در حوزه‌های حیاتی سلامت، چالش‌های اخلاقی بنیادینی را مطرح می‌کند که می‌تواند به تبعیض سیستماتیک در تخصیص درمان منجر شود.	۱. افزایش دسترسی و اثربخشی: هوش مصنوعی می‌تواند دسترسی به مراقبت‌های بهداشتی را افزایش و خدمات پزشکی را مؤثرتر و مقرون‌به‌صرفه‌تر کند. ۲. تقویت توانایی‌های پزشکان: هوش مصنوعی با کمک به پزشکان در فرایندهای کلیدی نظیر تشخیص، غربالگری، پیش‌بینی بیماری‌ها و اولویت‌بندی مراقبت‌ها به آنان «هدیه زمان» می‌بخشد؛ این امر به پزشکان اجازه می‌دهد تا تمرکز خود را از وظایف تکراری به مراقبت‌های پیچیده و تعامل انسانی با بیمار معطوف کند. ۳. مدیریت فعال سلامت بیماران: هوش مصنوعی به بیماران امکان می‌دهد از طریق اپلیکیشن‌ها به‌صورت فعال سلامت خود را مدیریت کنند و با ساده‌سازی اصطلاحات پیچیده پزشکی، به درک عمیق‌تری از وضعیت خود دست یابند.^{۱۶}

15. Capraro et al., "The Impact of Generative Artificial," 6.

16. Capraro et al., "The Impact of Generative Artificial," 8.

۲-۳. امکان‌سنجی طراحی سیاست‌گذاری عادلانه توزیع درآمد

یکی از چالش‌های کلیدی در هوش مصنوعی، عدالت در یادگیری ماشین است که به طراحی، توسعه و پیاده‌سازی مدل‌های یادگیری ماشینی برای جلوگیری از تبعیض و نابرابری در فرایندهای سیاست‌گذاری می‌پردازد. بی‌عدالتی الگوریتمی (algorithmic unfairness) زمانی رخ می‌دهد که سیستم‌های مبتنی بر هوش مصنوعی به‌طور ناعادلانه با گروه‌ها یا افراد مختلف رفتار کنند. برای مثال، این پدیده می‌تواند در قالب یک الگوریتم استخدامی ظهور کند که به دلیل تغذیه از داده‌های تاریخی جانبدارانه، نامزدهای مرد را بر زن ترجیح می‌دهد؛ چنانکه اخیراً فاش شد سیستم استخدام خودکار شرکت آمازون علیه نامزدهای زن، به‌ویژه برای موقعیت‌های توسعه نرم‌افزار و فنی، تبعیض قائل می‌شد. یکی از دلایل محتمل این سوگیری آن بود که داده‌های آموزشی سیستم، عمدتاً بر اساس رزومه‌های موفق گذشته شکل گرفته بود که در آنها اکثریت توسعه‌دهندگان نرم‌افزار را مردان تشکیل می‌دادند.^{۱۷} در پاسخ به این چالش، حوزه یادگیری ماشین منصفانه (Fair Machine Learning (Fair ML)) شکل گرفته است که هدف آن عملیاتی کردن نظریه‌های عدالت توزیعی در سیاست‌گذاری‌های مبتنی بر یادگیری ماشین است.

۲-۳-۱. امکان‌پذیری فنی طراحی سیستم توزیع درآمد با هوش مصنوعی

پیشرفت‌های قابل توجه در حوزه هوش مصنوعی از جمله مدل‌های پیشرفته یادگیری عمیق و تحلیل داده‌های بزرگ، امکان طراحی ابزارهایی را فراهم کرده‌اند که می‌توانند الگوهای پیچیده اقتصادی و اجتماعی را شناسایی و به بهینه‌سازی سیاست‌های توزیع منابع کمک کنند؛^{۱۸} همچنین توسعه فناوری‌های هوش مصنوعی در زمینه پردازش زبان طبیعی و تحلیل رفتار کاربران، امکان تعامل بهتر با افراد و پاسخ‌دهی به نیازهای متنوع را تسهیل می‌کند.^{۱۹} در ادامه برخی از این رویکردها بررسی می‌شود.

17. Xiaomeng Wang, Yishi Zhang and Ruilin Zhu, "A Brief Review on Algorithmic Fairness," *MSE* 1, no. 7 (2022): 1.

18. M. I. Jordan and T. M. Mitchell, "Machine learning: Trends, perspectives, and prospects," *Science* 349, no. 6245 (2015): 255.

19. Ian Goodfellow, Yoshua Bengio, and Aaron Courville, *Deep Learning* (Cambridge: The MIT Press, 2016), 478.

جدول (۴): روش‌های فنی طراحی سیستم‌های توزیع درآمد با هوش مصنوعی

عنوان بخش	شرح
سیستم‌های ترکیبی انسان در حلقه Human-in-(the-Loop (HITL)	این سیستم‌ها انسان را در جایگاه عنصری کلیدی و فعال در فرایند تصمیم‌گیری الگوریتم‌های هوش مصنوعی دخیل می‌کنند تا قضاوت و نظارت انسانی تضمین شود. ^{۲۰} در زمینه توزیع درآمد، HITL برای بهینه‌سازی انصاف، شفافیت و پاسخگویی طراحی شده است. هوش مصنوعی با تحلیل داده‌ها پیشنهادهایی برای تخصیص منابع ارائه می‌دهد، اما تصمیم نهایی با مشارکت انسان (کارشناسان یا نهادهای نظارتی) گرفته می‌شود. ^{۲۱}
درآمد پایه همگانی Universal Basic Income (UBI)	این سیاست، نسخه پیشرفته و فناوری‌محور UBI سنتی است که در آن پرداخت نقدی ثابت و منظمی بدون شرط به هر شهروند ارائه می‌شود. ^{۲۲} ویژگی هوشمند به معنای بهره‌گیری از هوش مصنوعی، تحلیل داده‌ها و فناوری‌های دیجیتال برای بهبود کارایی، عدالت و شخصی‌سازی این طرح است. هوش مصنوعی با تحلیل داده‌های واقعی (درآمد، هزینه، تورم، نیازهای منطقه‌ای) مبلغ UBI را پویا و متناسب با شرایط اقتصادی تنظیم می‌کند.
مالیات و یارانه‌های هدفمند	الگوریتم‌های هوش مصنوعی با بهره‌گیری از تکنیک‌هایی نظیر تحلیل داده، یادگیری ماشین و مدل‌سازی پیش‌بینانه، فرایند شناسایی افراد یا صنایع نیازمند حمایت را برای سیاست‌گذاران تسهیل می‌کنند. به این ترتیب با تنظیم مالیات‌ها و یارانه‌های هدفمند، منابع به‌صورت بهینه توزیع می‌شوند تا کمک‌ها دقیقاً به دست مؤثرترین گروه‌های هدف برسد و بیشترین تأثیر را داشته باشد. ^{۲۳}

20. Fabio Massimo Zanzotto, "Viewpoint: Human-in-the-Loop Artificial Intelligence," *Journal of Artificial Intelligence Research* 64 (2019): 245.

21. Raphael Koster et al., "Human-centred Mechanism Design with Democratic AI," *Nature Human Behaviour* 6 (2022): 1399.

22. Daniel Raventós, *Basic income: The material conditions of freedom* (London: Pluto Press, 2007), 11.

23. Andini Monica et al., "Targeting with machine learning: An application to a tax rebate program in Italy," *Journal of Economic Behavior and Organization* 156, no 1 (2018): 87.

این سیاست، هوش مصنوعی را برای تحلیل داده‌های اجتماعی-اقتصادی (درآمد، پوشش بیمه و منطقه جغرافیایی) به کار گرفته است و قیمت‌گذاری کالاها و خدمات ضروری (دارو) را به گونه‌ای تنظیم می‌کند که برای هر گروه از جامعه متناسب و قابل دسترس باشد. هدف این است که هم سودآوری کسب‌وکار حفظ شود و هم نابرابری‌های اقتصادی کاهش یابد و دسترسی عادلانه به کالاها و خدمات تضمین شود. برای مثال، ممکن است یک داروی سرطان برای بیماران در گروه‌های درآمدی مختلف یا مناطق جغرافیایی متفاوت با قیمت‌های متفاوتی عرضه شود تا اطمینان حاصل شود افرادی که بیشترین نیاز را دارند، توانایی خرید آن را داشته باشند. ^{۲۴} چالش‌هایی مانند بی‌طرفی و شفافیت هوش مصنوعی در این فرایند همچنان وجود دارد.	قیمت‌گذاری اخلاقی و رفاه اجتماعی
---	----------------------------------

۲-۳-۲. طراحی الگوریتم‌ها در یادگیری ماشین براساس نظریه‌های عدالت توزیعی

طراحی الگوریتم‌های یادگیری ماشین مبتنی بر نظریه‌های عدالت توزیعی، فرایندی است که در آن اصول فلسفی توزیع عادلانه منابع و فرصت‌ها به صورت مستقیم در ساختار و تابع هدف الگوریتم‌ها گنجانده می‌شود. برای عملیاتی کردن این مفهوم، دو رویکرد اصلی قابل طرح است: نخست، رویکرد مبتنی بر منابع که توسط فیلسوفانی مانند جان رالز، رونالد دورکین و توماس پوگ (John Rawls, Ronald Dworkin, and Thomas Pogge) مطرح شده است. دوم، رویکرد مبتنی بر قابلیت‌ها که توسط فیلسوفانی چون آمارتیا سن، مارتا نوسبام و الیزابت اندرسون (Amartya Sen, Martha Nussbaum, and Elizabeth Anderson) حمایت می‌شود. مسئله محوری این است که چگونه می‌توان این چهارچوب‌های نظری را در معماری سیستم‌های هوشمند پیاده‌سازی کرد.

در رویکرد نخست، جان رالز دو اصل بنیادین برای توزیع عادلانه منابع مطرح می‌کند: اول، اصل آزادی‌های اساسی برابر که بیان می‌کند هر فرد باید از حقوق و آزادی‌های اساسی برابر برخوردار باشد. دوم، اصل تفاوت که بیان می‌کند اگر توزیع منابع به گونه‌ای باشد که نابرابری ایجاد کند، این نابرابری باید بیشترین برخورداری و بهره‌مندی را برای محروم‌ترین

24. Karthik Ramakrishnan, "AI-driven ethical pricing: Balancing profitability and social responsibility," *International Journal of Scientific Research in Computer Science, Engineering and Information Technology* 11, no. 1 (2025): 187.

اعضای جامعه داشته باشد.^{۲۵}

در رویکرد مبتنی بر منابع، درآمد و ثروت به‌منزله شاخص‌های اصلی رفاه در نظر گرفته می‌شوند؛ به‌طوری که معمولاً افراد با درآمد و ثروت بیشتر، وضع زندگی بهتری دارند. با این حال، صرف ثروت مالی لزوماً تضمین‌کننده رفاه واقعی نیست، زیرا تبدیل منابع به رفاه به شرایط فیزیولوژیکی، روانی، اجتماعی، سیاسی و اقتصادی فرد وابسته است. این ناهمگونی‌ها می‌توانند بر میزان بهره‌مندی فرد از منابع اثر بگذارند؛ بنابراین معیار عدالت باید این تفاوت‌ها را در نظر بگیرد. بسیاری از این تفاوت‌ها با معیارهای مالی قابل سنجش نیستند، از این رو معیار عدالت باید فراتر از درآمد و ثروت صرف باشد.^{۲۶}

رالز در این زمینه از معیار کالاهای اولیه (primary goods) به‌عنوان معیاری جامع‌تر برای سنجش رفاه در رویکرد منابع دفاع می‌کند. برخلاف معیار ساده درآمد و ثروت، برداشت رالز از منابع گسترده‌تر و همه‌جانبه‌تر است و شامل منابع مادی (درآمد و ثروت) و منابع غیرمادی (فرصت‌ها، حقوق و آزادی‌ها) می‌شود.

با این وجود، آمارتیا سن معتقد است که معیار رالز به دلیل نادیده گرفتن ناهمگونی‌های فردی، برای اصلاح نابرابری‌ها کافی نیست. او در مقابل بر رویکرد قابلیت‌ها تأکید می‌ورزد؛ رویکردی که عدالت را نه در توزیع منابع، بلکه در فراهم‌سازی «قابلیت‌ها» (فرصت‌ها و امکانات واقعی) و «کارکردها» (دستاوردهای واقعی زندگی) جست‌وجو می‌کند تا اطمینان حاصل شود افراد توانایی واقعی برای دستیابی به زندگی ارزشمند را دارند.^{۲۷}

در حوزه یادگیری ماشین، می‌توان با الهام از رویکرد مبتنی بر منابع، خروجی مدل (دقت پیش‌بینی) را به‌عنوان منبعی در نظر گرفت و آن را براساس اصل تفاوت رالز توزیع کرد. برای مثال در سیستم‌های تشخیص گفتار هوشمند، گویش‌وران اقلیت‌های زبانی یا قومی معمولاً با نرخ خطای بالاتری مواجه می‌شوند که دسترسی آنها را به خدمات دیجیتال محدود می‌کند. با اعمال اصل تفاوت، الگوریتم به‌گونه‌ای تنظیم می‌شود که منابع محاسباتی و وزن‌های یادگیری به نفع این گروه‌های آسیب‌پذیر تغییر یابد تا دقت پیش‌بینی برای آنها افزایش یافته و نابرابری عملکردی میان گویش معیار و گویش‌های اقلیت کاهش یابد.

۲۵. سید هادی عربی، نظریات عدالت توزیعی، چاپ دوم (قم: پژوهشگاه حوزه و دانشگاه، ۱۳۹۹)، ۱۵۹.

26. Alan Lundgard, "Measuring Justice in Machine Learning," *arXiv*. 2009.10050 (2020): 6.

27. Lundgard, "Measuring Justice in Machine Learning," 19.

با این حال، این رویکرد فنی که بر عدالت تخصیصی تمرکز دارد به‌تنهایی کافی نیست؛ زیرا افزایش صرف دقت آماری برای گروه‌های آسیب‌پذیر، اگرچه نابرابری در خروجی را کاهش می‌دهد، اما قادر به درک و اصلاح ریشه‌های عمیق‌تری عدالتی نیست. در اینجا با چالشی به نام نابرابری در بازنمایی (Representation Disparity) مواجه می‌شویم. این نوع نابرابری زمانی رخ می‌دهد که سیستم‌های هوشمند، هویت، فرهنگ یا شأن گروه‌های خاصی را به‌درستی بازتاب نمی‌دهند یا آنها را در قالب کلیشه‌های منفی بازتولید می‌کنند. برای مثال، اگر یک مدل زبانی، کلمات مربوط به مشاغل مدیریتی را بیشتر با مردان و مشاغل خدماتی را با زنان مرتبط بداند، حتی اگر دقت پیش‌بینی آن بالا باشد، همچنان در حال بازتولید کلیشه‌ای جنسیتی و ایجاد آسیب بازنمایی است.

آسیب‌های ناشی از تخصیص ناعادلانه منابع را می‌توان با توزیع مجدد منابع کاهش داد، اما آسیب‌های بازنمایی را نمی‌توان تنها با تخصیص منابع برطرف کرد. این آسیب‌ها ریشه در نابرابری‌های ساختاری و منزلت اجتماعی دارند که کمی‌سازی آنها دشوار است و معیارهای فنی معمولاً فاقد حساسیت لازم برای درک آنها هستند.^{۲۸}

در اینجا رویکرد قابلیت‌ها چهارچوب غنی‌تری ارائه می‌دهد. این رویکرد به جای تمرکز صرف بر خروجی فنی، بررسی می‌کند که آیا سیستم هوشمند فرصت واقعی برای شکوفایی را در اختیار افراد قرار می‌دهد یا خیر. تحقق این امر نیازمند فرایندهای دموکراتیک و مشارکتی است تا سیستم‌ها به جای بازتولید نابرابری، قابلیت‌های انسانی را گسترش دهند.

واکاوی همگرایی میان نظریه‌های عدالت توزیعی و الزامات طراحی سیستم‌های هوشمند، دلالت‌های راهبردی زیر را آشکار می‌سازد: اول، بهره‌گیری از نظریه‌های کلاسیک عدالت همچون رالز و سن به‌منزله چهارچوب‌های مفهومی در طراحی الگوریتم‌های منصفانه، بیانگر پیوندی مستحکم با پیشینه فلسفی این حوزه است، اما توجه همزمان به محدودیت‌ها و نقدهای معاصر این نظریه‌ها نیز ضروری است. این رویکرد تعادلی امکان‌بازنگری و به‌روزرسانی مفاهیم عدالت را در پرتو فناوری‌های نوین فراهم می‌کند. دوم، تشخیص تمایز میان عدالت توزیعی و عدالت بازنمایی، ضرورت گذار از معیارهای صرفاً فنی (دقت) به سمت طراحی‌های جامع‌نگر را برجسته می‌کند. در این راستا طراحی الگوریتم‌های منصفانه باید فراتر از تخصیص منابع، ابعاد روانی، اجتماعی و فرهنگی را نیز در برگیرد تا از آسیب‌هایی

28. Lundgard, "Measuring Justice in Machine Learning," 15.

فرصت‌ها و چالش‌های هوش مصنوعی در سیاست‌گذاری عادلانه ... / صالحی رزوه ۶۷

مانند بازتاب نادرست هویت و بازتولید کلیشه‌های منفی جلوگیری شود. این دیدگاه، ضمن روشن ساختن مسیر پژوهش‌های آینده در زمینه هوش مصنوعی اخلاق‌مدار، ضرورت توسعه چهارچوب‌های نظری و عملی جامع‌تر را برای تحقق عدالت واقعی در فناوری‌های نوین مورد تأکید قرار می‌دهد.

۴-۲. پیشینه تحقیق

پژوهش‌های بسیاری به بررسی جنبه‌های مختلف عدالت در هوش مصنوعی پرداخته‌اند. در ادامه به برخی از این مطالعات در قالب جدول (۵) اشاره می‌شود.

جدول (۵): یافته‌های تجربی

محقق یا محققین	کشور یا کشورهای مورد مطالعه (دوره زمانی)	روش و تکنیک مورد استفاده	یافته‌های تجربی
گبرو و همکاران	-	چهارچوب «برگه‌های مشخصات داده» (Datasheets for Datasets)	افزایش شفافیت و پاسخگویی در مجموعه داده‌های یادگیری ماشین با مستندسازی دقیق فرآیندهای جمع‌آوری و ویژگی‌های داده‌ها، امکان شناسایی تبعیض‌ها و مسائل اخلاقی را فراهم و بهبود عدالت در سیستم‌های هوش مصنوعی را تسهیل می‌کند. ^{۲۹}
مونیکا آندینی و همکاران	ایتالیا (۲۰۱۴)	الگوریتم‌های یادگیری ماشین (ML) و تکنیک‌های پیش‌بینی‌کننده (بررسی برنامه)	با انتخاب دریافت‌کنندگان کمک مالی (پاداش ۸۰ یورویی) براساس ML، مصرف کلی مواد غذایی ۴۱.۸٪ (معادل ۷۶۰ میلیون یورو) افزایش و ۲۹.۵٪ از بودجه برنامه (۲ میلیارد یورو) صرفه‌جویی می‌شد. ^{۳۰}

29. Gebru et al., "Datasheets for Datasets," *Communications of the ACM* 64, no.12 (2021): 86-92.

30. Monica et al., "Targeting with machine learning," 101.

	بازپرداخت (مالیاتی)		
سیاست‌های مالیاتی مبتنی بر هوش مصنوعی قادرند تعادل بین برابری و بهره‌وری اقتصادی را تا ۱۶٪ نسبت به سیاست‌های سنتی (چهارچوب مالیاتی سائز) بهبود بخشند. ^{۳۱}	سیاست‌های مالیاتی مبتنی بر هوش مصنوعی (طراحی نرخ‌های مالیاتی هوشمندانه‌تر)	-	استفان ژنگ و همکاران
محققان در مطالعه خود، روشی دقیق، ارزان و مقیاس‌پذیر برای تخمین مصرف خانوار و ثروت از تصاویر ماهواره‌ای با وضوح بالا ارائه دادند. این مدل با استخراج ویژگی‌های تصویری مرتبط با رفاه اقتصادی، توانایی بالایی در پیش‌بینی مصرف (۳۷٪ تا ۵۵٪) و آریانس (۵۵٪ تا ۷۵٪) و ثروت (۵۵٪ تا ۷۵٪) و آریانس) در سطح خوشه‌ها نشان داد و به‌طور قابل توجهی بهتر از استفاده صرف از نورهای شب عمل کرد. ^{۳۲}	شبکه عصبی پیچشی (Convolutional Neural Network) و (CNN) و یادگیری انتقال (تحلیل تصاویر ماهواره‌ای با وضوح بالا)	نیجریه، تانزانیا، اوگاندا، مالاوی، رواندا	نیل جین و همکاران
رگرسیون خطی محبوب‌ترین روش مورد استفاده است و پس از آن جنگل تصادفی (Random Forest) و یادگیری عمیق (Deep Learning) و شبکه‌های عصبی پیچشی قرار دارند؛ استفاده از داده‌های ماهواره‌ای به‌طور فزاینده‌ای اهمیت یافته و	مرور نظام‌مند (Systematic Review)؛ روش‌های هوش مصنوعی در سنجش فقر	مرور نظام‌مند (۲۰۰۷ تا اوایل ۲۰۱۹)	روسنیتا ایسنین حمدان و همکاران

31. Stephan Zheng et al., "The AI economist: Improving equality and productivity with AI-driven tax policies," *arXiv: General Economics* (2020): 3.

32. Neal Jean et al., "Combining Satellite Imagery and Machine Learning to Predict Poverty," *Science* 353, no. 6301 (2016): 790-794.

			(رگرسیون خطی، جنگل تصادفی، یادگیری عمیق، شبکه‌های عصبی پیچشی)	پتانسیل هوش مصنوعی در مطالعات اجتماعی-اقتصادی به‌ویژه سنجش فقر، وسیع و در حال توسعه است. ^{۳۳}
راکال، تاواریس و پیزینلی	مقایسه تأثیرات تاریخی اتوماسیون (۱۹۸۰-۲۰۱۴) با پیش‌بینی‌های آینده هوش مصنوعی (۲۰۱۴-۲۰۴۸) با استفاده از داده‌های تجربی (۲۰۱۶- ۲۰۲۰) بریتانیا	مدل تعادل عمومی	پذیرش هوش مصنوعی با کاهش شکاف دستمزد بین کارگران، نابرابری حقوق را کاهش می‌دهد، اما از آنجا که صاحبان سرمایه (همان افراد پردرآمد) از افزایش بهره‌وری سود می‌برند، نابرابری ثروت را به‌شدت افزایش می‌دهد؛ این وضعیت با پذیرش سریع‌تر فناوری تشدید شده و سیاست‌گذاران را با یک بده‌بستان دشوار بین کنترل نابرابری و حفظ رشد اقتصادی روبه‌رو می‌کند. ^{۳۴}	

این پژوهش از چندین جنبه نوآوری و وجه تمایز قابل توجهی نسبت به مطالعات پیشین دارد:

۱. رویکرد جامع به عدالت در هوش مصنوعی: درحالی‌که بسیاری از مطالعات پیشین مانند گبرو و همکاران یا ایسنین حمدان و همکاران بر جنبه‌های فنی شفافیت داده‌ها یا ابزارهای سنجش فقر با AI تمرکز دارند، این پژوهش به ابعاد اجتماعی و فلسفی عدالت در هوش مصنوعی می‌پردازد. با تمایز قائل شدن میان عدالت توزیعی و عدالت بازنمایی و ارجاع

33. Rusnita Isnin Hamdan, Azuraliza Abu Bakar, and Nur Samsiah Sani, "Does Artificial Intelligence Prevail in Poverty Measurement," *Journal of Physics: Conference Series* 1529, no. 4 (2020): 1.

34. Emma J. Rockall, Marina Mendes Tavares, and Carlo Pizzinelli, "AI Adoption and Inequality," *IMF Working Papers*, No. 2025/068 (2025): 44-46.

به نظریات رالز و سن، چهارچوب مفهومی غنی برای درک پیچیدگی‌های نابرابری در عصر هوش مصنوعی ارائه می‌دهد که فراتر از صرف تخصیص منابع مادی است.

۲. کاربرد هوش مصنوعی به عنوان آزمایشگاه سیاستی برای بهینه‌سازی مالیات با مدل‌سازی عامل‌محور: برخلاف مطالعات پیشین که بر شناسایی یا تخمین فقر (نیل جین و همکاران)، پیش‌بینی اثرات فناوری (راکال و همکاران) یا بهبود کارایی تخصیص (مونیکا آندینی و همکاران) تمرکز دارند، این مقاله با استفاده از مدل شبیه‌سازی مبتنی بر عامل (Agent-Based Model) برای طراحی و کشف فعالانه یک سیاست مالیاتی بهینه بهره می‌برد.

سهم اصلی این مقاله در ترکیب تحلیلی مفهومی از عدالت در هوش مصنوعی با مطالعه موردی روش‌شناختی نوآورانه (مدل شبیه‌سازی مبتنی بر عامل برای بهینه‌سازی مالیات) است که راهبردهای سیاستی عملی را ارائه می‌دهد. این ترکیب، شکاف موجود در ادبیات را پر کرده است و به سیاست‌گذاران و پژوهشگران ابزاری جامع برای مواجهه با چالش نابرابری درآمدی در عصر هوش مصنوعی می‌بخشد.

۳. معرفی الگو و روش‌شناسی پژوهش

۳-۱. معرفی الگو

یکی از چالش‌های بنیادین علم اقتصاد، بده‌بستان بین بهره‌وری و برابری (Productivity-Equality Trade-off) است. از یک سو، سیاست‌های بازار آزاد با حداقل کردن مالیات، سرمایه‌گذاری و کار را تشویق می‌کند و به حداکثر بهره‌وری اقتصادی منجر می‌شوند؛ اما این کارایی اغلب با افزایش شکاف طبقاتی و نابرابری شدید اجتماعی همراه است. از سوی دیگر، سیاست‌های بازتوزیعی که از طریق مالیات‌های سنگین به دنبال افزایش برابری هستند، می‌توانند انگیزه تولید را در افراد ماهر و مولد کاهش دهند و به افت بهره‌وری کل اقتصاد منجر شوند.

یافتن نقطه بهینه این تعادل از راه آزمایش سیاست‌های مختلف در دنیای واقعی، پرهزینه و زمان‌بر است. با الهام از مطالعه استفان ژنگ و همکاران،^{۳۵} این پژوهش به دنبال طراحی یک سیاست مالیاتی بهینه است. این هدف با شبیه‌سازی اثرات سیاست‌های مالیاتی مختلف در یک آزمایشگاه اقتصادی مجازی با بهره‌گیری از عامل‌های هوشمند دنبال می‌شود تا

35. Zheng et al., "The AI economist: Improving equality," 1-46.

فرصت‌ها و چالش‌های هوش مصنوعی در سیاست‌گذاری عادلانه ... / صالحی رزوه ۷۱

بهترین تعادل میان بهره‌وری و برابری برقرار شود و رفاه اجتماعی کل به حداکثر برسد. الگوی این پژوهش به شرح زیر است:

مرحله ۱: ساخت یک آزمایشگاه اقتصادی مجازی (Simulation Environment)

ابتدا یک دنیای مجازی شبیه‌سازی شده، طراحی می‌شود که شبیه به بازی است. در این دنیا منابعی مانند سنگ و چوب وجود دارد و عاملان اقتصادی می‌توانند با جمع‌آوری این منابع و ساختن خانه، سکه (معادل پول) به دست آورند. این دنیا به گونه‌ای طراحی شده است که تفاوت در «مهارت» عاملان، به‌طور طبیعی باعث ایجاد نابرابری اقتصادی و انگیزه برای تخصص‌گرایی و تجارت می‌شود.

مرحله ۲: خلق عاملان اقتصادی هوشمند (AI Agents)

چهار عامل اقتصادی (خانوارها) در این دنیا قرار داده شدند که توسط هوش مصنوعی (یادگیری تقویتی (Reinforcement Learning)) کنترل می‌شوند. هدف هر عامل به حداکثر رساندن مطلوبیت شخصی خود است. مطلوبیت آنها از دو بخش تشکیل شده است: لذت حاصل از داشتن پول و زحمت ناشی از کار کردن (که از مطلوبیت کم می‌کند). این تابع مطلوبیت برای مدل‌سازی مطلوبیت نهایی نزولی پول (Diminishing Marginal Utility of Money) به کار می‌رود. بدین معنا که ارزش هر واحد پول اضافی، با افزایش ثروت کاهش می‌یابد. در چهارچوب یادگیری تقویتی، پاداش لحظه‌ای هر عامل، تغییر در مطلوبیت اوست؛ بنابراین، عامل‌ها یاد می‌گیرند اقداماتی را انجام دهند که بیشترین افزایش مثبت را در مطلوبیت‌شان ایجاد کند.

$$1) u_i = (x_{i,t}, l_{i,t}) = crra(x_{i,t}^c) - l_{i,t}, \quad crra(z) = \frac{z^{1-\eta}-1}{1-\eta}, \quad \eta > 0$$

$x_{i,t}$ ، کل پول انباشته شده عامل i تا زمان t ، $l_{i,t}$ مجموع کل کار یا تلاشی که عامل i از ابتدای بازی تا زمان t انجام داده است و z درآمدی است که در یک بازه زمانی مشخص به دست می‌آید. عبارت $crra$ بیانگر تابع مطلوبیت با ریسک‌گریزی نسبی ثابت (Constant Relative Risk Aversion) است. وجود بخش $crra$ تضمین می‌کند که عاملان به صورت بی‌نهایت و بدون فکر کار نکنند. در یک نقطه‌ای، زحمت و هزینه کار برای ساختن یک خانه جدید، بیشتر از مطلوبیت ناچیزی است که از سکه‌های اضافی آن به دست می‌آید. η

پارامتر ریسک‌گریزی است و مشخص می‌کند که با افزایش ثروت یک عامل، ارزش هر سکه اضافی برای او چقدر سریع کاهش می‌یابد. هر چه η بزرگتر باشد، عامل ریسک‌گریزتر می‌شود و مطلوبیت پول سریع‌تر کاهش می‌یابد؛ در مقابل، هر چه η کوچکتر باشد، عامل ریسک‌پذیری بیشتری دارد و مطلوبیت پول برای او دیرتر کاهش پیدا می‌کند. در مطالعه استفان ژنگ و همکاران، مقدار این پارامتر ۰.۲۳ در نظر گرفته شده است. این انتخاب، که بیانگر سطح ریسک‌گریزی نسبتاً پایینی برای عاملان اقتصادی است، منجر به خلق رفتار اقتصادی واقع‌گرایانه می‌شود. به‌طور عاملان از کسب درآمد بیشتر لذت می‌برند اما با کاهش ملایم مطلوبیت نهایی برای هر واحد اضافی درآمد مواجه‌اند. از آنجا که این کاهش مطلوبیت نهایی شدید نیست، انگیزه برای کار و تلاش جهت کسب درآمد بیشتر حفظ می‌شود. این رویکرد، تعادل رفتاری منطقی بین یک عامل کاملاً حریص (که بی‌نهایت کار می‌کند) و یک عامل بسیار محتاط (که با درآمد اندک اشیاع می‌شود) برقرار کرده و تضمین می‌کند که عاملان برای تحلیل بده‌بستان میان بهره‌وری و برابری، رفتاری قابل پیش‌بینی و منطقی از خود نشان دهند.

مرحله ۳: طراحی یک دولت هوش مصنوعی (AI Planner)

یک هوش مصنوعی دیگر در جایگاه برنامه‌ریز اجتماعی یا دولت عمل می‌کند. این هوش مصنوعی نمی‌تواند عاملان را مستقیماً کنترل کند. تنها ابزار دولت تعیین نرخ‌های مالیات در پلکان‌های درآمدی مختلف است. هدف دولت به حداکثر رساندن رفاه اجتماعی است.

مرحله ۴: فرآیند یادگیری دو سطحی (Two-Level Learning)

عاملان و دولت به‌صورت همزمان یاد می‌گیرند:

۱. حلقه داخلی (Inner Loop): عامل‌های اقتصادی که خودشان AI هستند با سطوح مهارت متفاوت از طریق آزمون و خطا یاد می‌گیرند که چگونه رفتار خود (کار کردن، تجارت، استراحت) را برای به حداکثر رساندن مطلوبیت خود بهینه کنند. آنها با جمع‌آوری منابع (چوب و سنگ) و ساخت خانه، سکه (درآمد) کسب می‌کنند. هر فعالیتی هزینه نیروی کار دارد. آنها به‌صورت پویا به سیاست‌های مالیاتی واکنش نشان می‌دهند.

۲. حلقه بیرونی (Outer Loop): یک هوش مصنوعی دیگر در نقش برنامه‌ریز اجتماعی یا دولت با مشاهده رفتار تطبیقی عامل‌ها یاد می‌گیرد که نرخ‌های مالیاتی را به گونه‌ای تنظیم کند که یک تابع رفاه اجتماعی به حداکثر برسد.

در این پژوهش، نقش برنامه‌ریز هوشمند به صورت دستی ایفا می‌شود. به این صورت که با آزمایش سه سناریوی مختلف مالیاتی از پیش تعریف شده به دنبال یافتن سیاستی هستیم که رفاه اجتماعی را بهینه کند. این فرایند (برای عاملان) میلیون‌ها بار تکرار می‌شود تا عاملان به بهترین راهبرد شخصی خود دست یابند.

مرحله ۵: اجرای سیاست‌های مالیاتی

پس از اتمام فرایند یادگیری عاملان، سه سناریوی مختلف سیاست مالیاتی برای آزمایش تعریف می‌شود: سناریو A (بازار آزاد) که در آن نرخ همه طبقات مالیاتی صفر است ($\tau = 0$); سناریو B (مالیات سنگین) مبتنی بر یک سیستم مالیات بر درآمد تصاعدی با نرخ‌های بالا برای طبقات مختلف درآمدی ($\tau = \{0.20, 0.40, 0.60, 0.80\}$) و سناریو C (مالیات متعادل) مبتنی بر یک سیستم مالیات بر درآمد تصاعدی و منطقی برای طبقات درآمدی ($\tau = \{0.10, 0.20, 0.30, 0.40\}$).

مالیات در پایان هر دوره مالیاتی بر درآمد کسب شده در همان دوره اعمال می‌شود. بازه‌های درآمدی از پیش تعریف شده وجود دارد و برنامه‌ریز اجتماعی (هوش مصنوعی دولت) برای هر بازه، یک نرخ مالیات نهایی (τ) تعیین می‌کند. این نرخ‌ها می‌توانند در هر دوره مالیاتی جدید، براساس وضعیت اقتصاد، تغییر کنند. کل مالیات پرداختی یک عامل، مجموع مالیات پرداختی در هر پلکان است.

$$T(z) = \sum_{b=0}^{B-1} \tau_b ((m_{b+1} - m_b) 1[z > m_{b+1}] + (z - m_b) 1[m_b < z \leq m_{b+1}])$$

معادله (۲)، میزان کل مالیات $T(z)$ را برای یک درآمد z در یک دوره مالیاتی براساس یک سیستم مالیاتی پلکانی (Bracketed) محاسبه می‌کند. b شمارنده یا اندیس پلکان مالیاتی، B تعداد کل پلکان‌های مالیاتی، m_b و m_{b+1} آستانه‌های درآمدی که مرزهای پلکان b را مشخص می‌کنند، τ_b نرخ مالیات نهایی برای پلکان b ، درصدی از مالیات است که فقط به آن بخش از درآمد که در این پلکان قرار می‌گیرد اعمال می‌شود. $1[\dots]$ تابع نشانگر (Indicator Function) اگر شرط داخل کروشه درست باشد، مقدار آن ۱ و اگر شرط نادرست باشد، مقدار آن ۰ است.

براساس معادله (۳) در پایان هر دوره مالیاتی، کل مالیات جمع‌آوری شده از تمام چهار عامل در یک صندوق مشترک واریز شده و به‌طور مساوی بین هر چهار عامل تقسیم و به

دارایی آنها اضافه می‌شود. این الگو باعث می‌شود عاملان کم‌درآمد، دریافت‌کننده خالص (Net Receiver) و عاملان پردرآمد، پرداخت‌کننده خالص (Net Payer) باشند. در نتیجه عامل‌هایی که درآمد بالایی داشته و مالیات زیادی پرداخته‌اند، مبلغی کمتر از مالیات پرداختی خود را پس می‌گیرند (پرداخت‌کننده خالص). عامل‌هایی که درآمد کمی داشته و مالیات کمی پرداخته‌اند، مبلغی بیشتر از مالیات پرداختی خود دریافت می‌کنند (دریافت‌کننده خالص یا دریافت‌کننده یارانه). این مکانیزم، ابزار اصلی دولت برای کاهش نابرابری است.

$$3) \tilde{z}_i^p = z_i^p - T(z_i^p) + \frac{1}{N} \sum_{j=1}^N T(z_j^p)$$

$\tilde{z}_i^p$ درآمد خالص (پس از کسر مالیات و دریافت یارانه) عامل i در دوره مالیاتی p ، z_i^p درآمد ناخالص (قبل از کسر مالیات و دریافت یارانه) عامل i در دوره مالیاتی p ، $T(z_j^p)$ میزان مالیاتی که عامل i براساس درآمد ناخالص z_i^p خود در دوره p پرداخت می‌کند و N تعداد کل عاملان در اقتصاد است.

مرحله ۶: ارزیابی عملکرد

برای ارزیابی عملکرد هر سیاست مالیاتی از سه شاخص اصلی استفاده شده است:

۱. بهره‌وری (Productivity): مجموع کل سکه‌هایی که توسط تمام عامل‌ها در پایان شبیه‌سازی به دست آمده است. این شاخص معادل اندازه کل اقتصاد یا GDP است.

$$4) prod(x^c) = \sum_{i=1}^N x_i^c$$

۲. برابری (Equality): شاخص جینی، نابرابری توزیع درآمد را اندازه‌گیری می‌کند (۰ برابری کامل، ۱ نابرابری مطلق). شاخص برابری به صورت معادله (۵) تعریف می‌شود. مقدار ۱ به معنای برابری کامل و مقدار ۰ به معنای نابرابری مطلق است.

$$5) eq(x^c) = 1 - gini(x^c) \frac{N}{N-1}, \quad 0 \leq eq(x^c) < 1$$

$$6) gini(x^c) = \frac{\sum_{i=1}^N \sum_{j=1}^N |x_i^c - x_j^c|}{2N \sum_{i=1}^N x_i^c}, \quad 0 \leq gini(x^c) \leq \frac{N-1}{N}$$

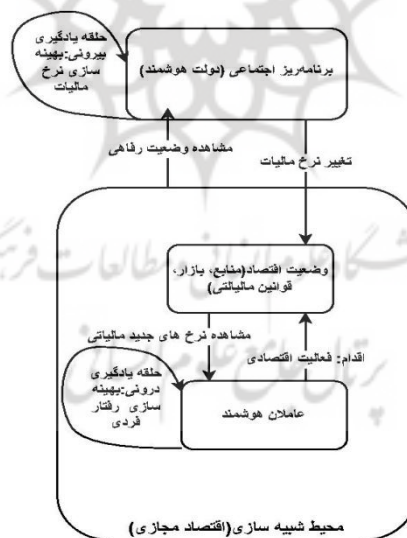
تعداد کل عاملان اقتصادی، i و z شمارنده‌هایی که برای پیمایش در میان عاملان اقتصادی از ۱ تا N به کار می‌روند، x_t^c موجودی سکه عامل i ام و x_t^z موجودی سکه عامل z ام است.

معادله (۶) میانگین اختلاف ثروت بین تمام زوج‌های ممکن از افراد جامعه را محاسبه و آن را بر کل درآمد جامعه نرمال‌سازی می‌کند. نتیجه عددی بین ۰ و $(N-1)/N$ است؛ سپس معادله (۵) این ضریب جینی را برای تطابق با بازه ۰ تا ۱ اصلاح می‌کند.

۳. تابع رفاه اجتماعی (Social Welfare Function) (SWF): این شاخص نهایی برای قضاوت است. این تابع به سیاستی پاداش می‌دهد که همزمان هم اقتصاد بزرگی ایجاد کند و هم آن را به شکل عادلانه‌ای توزیع نماید. هدف اصلی برنامه‌ریز اجتماعی هوش مصنوعی، یافتن سیاست مالیاتی‌ای است که این شاخص را به حداکثر برساند.

$$7) swf_i(x_t^c) = eq_t(x_t^c).prod_t(x_t^c)$$

نمودار (۱): معماری یادگیری هوش مصنوعی دو سطحی



منبع: یافته‌های پژوهش

نمودار (۱)، معماری یادگیری دو سطحی مدل را به تصویر می‌کشد. این نمودار چگونگی تعامل غیرمستقیم و پویا میان برنامه‌ریز اجتماعی (دولت هوشمند) و عاملان هوشمند را از طریق یک واسطه مرکزی، یعنی محیط شبیه‌سازی، نمایش می‌دهد. در «حلقه بیرونی»، برنامه‌ریز اجتماعی با مشاهده شاخص وضعیت رفاهی که خروجی کل محیط است به ارزیابی عملکرد اقتصاد می‌پردازد؛ سپس براساس فرایند یادگیری تقویتی خود برای بهینه‌سازی نرخ مالیات، اقدام به تغییر نرخ مالیات می‌کند. این سیاست مستقیماً بر قوانین حاکم بر وضعیت اقتصاد در دل محیط شبیه‌سازی اثر می‌گذارد. در «حلقه درونی»، عاملان هوشمند با مشاهده نرخ‌های جدید مالیاتی و سایر شرایط از بلوک وضعیت اقتصاد، اطلاعات لازم برای تصمیم‌گیری را کسب می‌کنند؛ سپس براساس فرایند یادگیری خود برای بهینه‌سازی رفتار فردی، دست به فعالیت اقتصادی می‌زنند. این اقدامات جمعی به نوبه خود، «وضعیت اقتصاد» را برای گام بعدی شبیه‌سازی تغییر می‌دهند. این دو حلقه به‌صورت همزمان عمل می‌کنند و فرایندی هم‌تکاملی را شکل می‌دهند. به عبارت دیگر، سیاست‌های تعیین‌شده توسط برنامه‌ریز، محیطی را که عاملان در آن عمل می‌کنند تغییر می‌دهد و اقدامات عاملان نیز نتایج کلانی را ایجاد می‌کند که مبنای تصمیم‌گیری‌های آتی برنامه‌ریز قرار می‌گیرد.

۲-۳. روش‌شناسی پژوهش

در این پژوهش برای ارزیابی تأثیر سیاست‌های مالیاتی مختلف بر رفاه اجتماعی از یک مدل شبیه‌سازی مبتنی بر عامل استفاده شد. فرایند شبیه‌سازی در محیط اویوز و با تعریف پارامترهای اولیه مدل، شامل موجودی اولیه منابع (سنگ و چوب) و همچنین تخصیص سطح مهارت متفاوت به هر یک از چهار عامل اقتصادی آغاز شد. هسته مرکزی مدل، یک حلقه تکرار (Do-Loop) است که به مدت ۱۰۰۰ گام زمانی (time-step) اجرا می‌شود و هر گام، نماینده یک دوره فعالیت در اقتصاد شبیه‌سازی شده است. در هر گام از این حلقه، مجموعه‌ای از دستورات برای هر یک از چهار عامل به‌صورت مستقل و براساس یک منطق تصمیم‌گیری از پیش تعریف‌شده عمل می‌کند که شامل مراحل زیر است:

۱. ارزیابی منابع: عامل ابتدا بررسی می‌کند که آیا منابع لازم برای تولید (یک واحد سنگ و یک واحد چوب) را در اختیار دارد یا خیر.

۲. تحلیل هزینه-فایده: در صورت وجود منابع، عامل، سود خالص حاصل از ساخت و

فرصت‌ها و چالش‌های هوش مصنوعی در سیاست‌گذاری عادلانه ... / صالحی رزوه ۷۷

فروش یک خانه را براساس هزینه‌های تولید و قیمت فروش محاسبه می‌کند.

۳. تصمیم‌گیری برای تولید: عامل فقط در صورتی اقدام به تولید می‌کند که سود خالص محاسبه‌شده از یک آستانه سودآوری مشخص فراتر رود. این مکانیزم، رفتار عقلایی اقتصادی را شبیه‌سازی می‌کند.

۴. اجرای تولید و تراکنش: پس از تصمیم مثبت، منابع مورد نیاز از موجودی عامل کسر شده است و درآمد حاصل از فروش (پس از کسر مالیات متعلقه) به موجودی سکه او افزوده می‌شود.

۵. بازسازی منابع محیطی: به منظور تضمین پایداری بلندمدت اقتصاد و جلوگیری از اتمام منابع، مکانیزمی برای بازتولید تدریجی سنگ و چوب در محیط شبیه‌سازی تعبیه شده است.

این فرایند شبیه‌سازی به‌طور کامل و مستقل، سه بار برای هر یک از سناریوهای مالیاتی (صفر، سنگین و متعادل) تکرار شد. در هر اجرا، تنها پارامتر متغیر، نرخ مالیات بود. پس از اتمام ۱۰۰۰ گام زمانی برای هر سناریو، شاخص‌های بهره‌وری کل و برابری براساس عملکرد تجمیعی عاملان محاسبه و ثبت شد.

۳-۳. نتایج

پس از اجرای شبیه‌سازی برای هر سه سناریو، نتایج کمی به‌دست‌آمده در جدول (۶) خلاصه شده است:

جدول (۶): نتایج شبیه‌سازی

سناریو	بهره‌وری	برابری	رفاه اجتماعی	تحلیل
مالیات صفر	۱۴۵۰	۰۰۰۳۴	۴۹۰۵۵	سناریوی حداکثر کارایی: در غیاب مالیات، عامل‌های هوشمند هیچ مانعی برای تولید نمی‌بینند و اقتصاد با تمام پتانسیل خود کار می‌کند؛ اما این رشد به قیمت نابرابری شدید تمام می‌شود و ثروت به شکل ناعادلانه‌ای توزیع می‌شود.
مالیات سنگین	۱۲۳۰	۰۰۰۳۶۸	۴۵۰۲۹	سناریوی تله رکود: مالیات سنگین، برابری را افزایش می‌دهد، اما انگیزه تولید به‌شدت سرکوب می‌شود،

بهره‌وری کاهش می‌یابد و در نهایت، رفاه کل جامعه حتی از حالت بازار آزاد نیز کمتر می‌شود.				
سناریوی بهینه: این سیاست، نقطه تعادل طلایی را پیدا کرده است. با اعمال یک مالیات منطقی، بخش کوچکی از بهره‌وری فدای افزایش چشمگیر برابری می‌شود. ترکیب این دو عامل، منجر به بالاترین رفاه اجتماعی ممکن در میان گزینه‌های آزمایشی می‌شود.	۴۹.۹۴	۰۰.۳۶۱	۱۳۸۰	مالیات متعادل

منبع: یافته‌های پژوهش

نتایج به دست آمده به روشنی فرضیه پژوهش را تأیید می‌کند:

۱. تحلیل بهره‌وری و برابری: یک رابطه معکوس بین نرخ مالیات و بهره‌وری و یک رابطه مستقیم بین نرخ مالیات و برابری مشاهده می‌شود. سناریوی مالیات صفر با ارائه حداکثر انگیزه به عاملان به بیشترین بهره‌وری (۱۴۵۰) دست یافت، اما کمترین برابری (۰۰.۳۴) را به همراه داشت. سناریوی مالیات سنگین با بازتوزیع گسترده ثروت، بیشترین برابری (۰۰.۳۶۸) را محقق کرد، اما این کار به قیمت سرکوب شدید انگیزه تولید و کسب کمترین بهره‌وری (۱۲۳۰) تمام شد.

۲. تحلیل رفاه اجتماعی و شناسایی سیاست بهینه: سیاست مالیات سنگین، علیرغم دستیابی به بالاترین سطح برابری به دلیل آسیب شدید به بهره‌وری، ناکارآمدترین سیاست از منظر رفاه کل جامعه بود و کمترین امتیاز (۴۵.۲۹) را کسب کرد. سیاست مالیات صفر با وجود کارایی اقتصادی بالا به دلیل نابرابری شدید، نتوانست به بالاترین رفاه اجتماعی برسد (۴۹.۵۵). سیاست مالیات متعادل به مثابه سیاست بهینه در این شبیه‌سازی، منجر به کسب بالاترین امتیاز رفاه اجتماعی (۴۹.۹۴) شده است.

این پژوهش نشان داد که چگونه می‌توان از شبیه‌سازی به عنوان یک آزمایشگاه سیاستی مبتنی بر هوش مصنوعی استفاده کرد. این ابزار به ما اجازه داد تا سیاست‌های مختلف را به صورت مجازی و بدون ریسک آزمایش کرده و براساس داده‌های کمی، بهترین گزینه را انتخاب کنیم. این آزمایش چشم‌اندازی برای آینده سیاست‌گذاری اقتصادی ترسیم می‌کند که در آن از AI برای طراحی سیاست‌های اقتصادی هوشمندانه‌تر و مبتنی بر شواهد استفاده می‌شود.

۴. نتیجه‌گیری و پیشنهادهای سیاستی

این پژوهش به بررسی فرصت‌ها و چالش‌های هوش مصنوعی در سیاست‌گذاری عادلانه توزیع می‌پردازد. نتایج نشان می‌دهد که هوش مصنوعی پتانسیل بالایی برای بهبود عدالت توزیعی دارد، اما کاربرد آن بدون توجه به چالش‌های بنیادی همچون سوگیری الگوریتمی، نقض حریم خصوصی و عدم شفافیت مدل‌ها می‌تواند نابرابری‌ها را تشدید کند. تحقق عدالت در هوش مصنوعی، مستلزم گذار از معیارهای صرفاً فنی به رویکردی جامع‌نگر است که با تلفیق نظریه‌های عدالت، الزامات دو وجه توزیعی و بازنمایی را در بستر پیچیده اجتماعی و فرهنگی یکپارچه کند. در مطالعه موردی به این پرسش محوری پاسخ داده شد که هوش مصنوعی چگونه می‌تواند بده‌بستان میان بهره‌وری اقتصادی و برابری درآمد را از طریق سیاست‌گذاری مالیاتی بهینه سازد. نتایج نشان داد که سیاست مالیاتی متعادل به بالاترین سطح رفاه اجتماعی دست یافت. این سیاست در مقایسه با بازار آزاد (مالیات صفر) با پذیرش کاهش ۴.۸۳ درصدی در بهره‌وری، برابری را ۶.۱۸ درصد بهبود بخشید. همزمان در مقایسه با سیاست مالیات سنگین با فدا کردن تنها ۱.۹۰ درصد از برابری، بهره‌وری را به میزان چشمگیر ۱۲.۲۰ درصد افزایش داد. این بده‌بستان، رفاه اجتماعی را ۰.۷۹ درصد نسبت به بازار آزاد و ۱۰.۲۷ درصد نسبت به مالیات سنگین افزایش داد. این یافته بر اهمیت هوش مصنوعی به عنوان آزمایشگاه سیاستی برای طراحی و آزمایش سیاست‌های اقتصادی هوشمندانه‌تر و مبتنی بر شواهد تأکید می‌کند، مشروط بر آنکه حفاظت از حقوق افراد، شفافیت و مسئولیت‌پذیری در فرایند توسعه و اجرای سیستم‌های هوشمند رعایت شود.

بر این اساس برای بهره‌برداری مؤثر از هوش مصنوعی در بهبود توزیع درآمد و کاهش نابرابری، سیاست‌گذاران باید چهارچوبی جامع و چندبعدی را اتخاذ کنند. این چهارچوب شامل مشارکت فعال ذینفعان (از جمله جامعه‌شناسان، سیاست‌گذاران و نمایندگان گروه‌های آسیب‌پذیر)، راه‌اندازی سازوکارهای شفاف‌سازی و نظارت انسانی بر تصمیم‌های الگوریتمی از طریق رویکرد انسان در حلقه و تعریف نقاط مداخله انسانی در فرایند تصمیم‌گیری است. حفاظت از حریم خصوصی با به کارگیری تکنیک‌های پیشرفته و تدوین قوانین جامع حفاظت از داده‌های شخصی با نظارت مستقل، همراه با استانداردهای داده‌های آموزشی برای کاهش سوگیری‌ها و ارزیابی مستمر الگوریتم‌ها توسط نهادهای مستقل از دیگر الزامات است؛ همچنین سرمایه‌گذاری در زیرساخت‌های دیجیتال و دسترسی برابر به فناوری و طراحی پلتفرم‌های آموزشی و بهداشتی هوشمند با محتوای شخصی‌سازی شده برای کاهش

شکاف‌های دسترسی، می‌تواند به بهبود توزیع عادلانه خدمات کمک کند.

در راستای تکمیل مطالعه حاضر، پیشنهادها زیر ارائه می‌شود:

تحلیل مقایسه‌ای برای ارزیابی عملکرد مدل هوش مصنوعی: سیاست مالیاتی بهینه‌ای که توسط این مدل کشف شد با مدل‌های شناخته‌شده دیگر مانند مدل نظری سائر مقایسه شود.

اعتبارسنجی با انسان: طراحی یک مطالعه تجربی با شرکت‌کنندگان انسانی (برای مثال از طریق آزمایشگاه‌های اقتصاد تجربی) برای اعتبارسنجی نتایج شبیه‌سازی. در این آزمایش، انسان‌ها به جای عاملان هوش مصنوعی بازی می‌کنند و سیاست مالیاتی کشف‌شده روی آنها اعمال می‌شود تا کارایی آن در محیطی واقعی‌تر سنجیده شود.

توسعه مدل برای تحلیل‌های عمیق‌تر: افزودن پیچیدگی‌های بیشتر به مدل مانند چند کالا، بودجه دولت، انتقالات اجتماعی، تأثیر تغییرات تقاضا در اقتصاد شبیه‌سازی‌شده و نیز اجرای حلقه‌های بازخورد دینامیک بین دولت و عاملان (برای مثال تغییرات در سیاست مالیاتی به مرور زمان براساس پاسخ‌های جمعی عاملان).

سیاهه منابع

الف - منابع فارسی:

عربی، سید هادی. *نظریات عدالت توزیعی*. چاپ دوم. قم: پژوهشگاه حوزه و دانشگاه، ۱۳۹۹.

ب - منابع لاتین:

- Capraro, Valerio, Austin Lentsch, Daron Acemoglu, Selin Akgun, Aisel Akhmedova, Ennio Bilancini, Jean-François Bonnefon, Pablo Brañas-Garza, Luigi Butera, Karen M Douglas, Jim A C Everett, Gerd Gigerenzer, Christine Greenhow, Daniel A Hashimoto, Julianne Holt-Lunstad, Jolanda Jetten, Simon Johnson, Werner H Kunz, Chiara Longoni, Pete Lunn, Simone Natale, Stefanie Paluch, Iyad Rahwan, Neil Selwyn, Vivek Singh, Siddharth Suri, Jennifer Sutcliffe, Joe Tomlinson, Sander van der Linden, Paul A M Van Lange, Friederike Wall, Jay J Van Bavel, and Riccardo Viale. "The impact of generative artificial intelligence on socioeconomic inequalities and policy making." *PNAS Nexus* 3, no. 6 (June 2024): 1-18. <https://doi.org/10.1093/pnasnexus/pgae191>
- Chancel, Lucas., T. Piketty, E. Saez, and G. Zucman. "World Inequality Report 2022." *World Inequality Lab* (2022): 1-19. <https://wir2022.wid.world>
- Goodfellow, Ian, Yoshua Bengio, and Aaron Courville. *Deep Learning*. Cambridge: The MIT Press, 2016.
- Hamdan, Rusnita Isnin, Azuraliza Abu Bakar, and Nur Samsiah Sani. "Does Artificial Intelligence Prevail in Poverty Measurement." *Journal of Physics: Conference Series* 1529, no. 4 (2020): 042082. 10.1088/1742-6596/1529/4/042082
- Jean, Neal, Marshall Burke, Michael Xie, W. Matthew Davis, David B. Lobell, and Stefano Ermon. "Combining Satellite Imagery and Machine Learning to Predict Poverty." *Science* 353, no. 6301 (August 18, 2016): 790-794. 10.1126/science.aaf7894
- Jordan, M. I., and T. M. Mitchell. "Machine Learning: Trends, Perspectives, and Prospects." *Science* 349, no. 6245 (2015): 255-260. <https://doi.org/10.1126/science.aaa8415>.
- Kasy, Maximilian, and Rediet Abebe. "Fairness, Equality, and Power in Algorithmic Decision-Making." In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency (FAccT '21)*, 576-586. New York: Association for Computing Machinery, 2021. <https://doi.org/10.1145/3442188.3445919>.
- Koster, Raphael, Jan Balaguer, Andrea Tacchetti, Ari Weinstein, Tina Zhu, Oliver Hauser, Duncan Williams, Lucy Campbell-Gillingham, Phoebe Thacker, Matthew Botvinick, Christopher Summerfield. "Human-centred mechanism design with Democratic AI." *Nature Human Behavior* 6 (2022): 1398-1407. <https://doi.org/10.1038/s41562-022-01383-x>

- Lipton, Zachary C. "The Mythos of Model Interpretability: In Machine Learning, the Concept of Interpretability Is Both Important and Slippery." *Queue* 16, no. 3 (2018): 31-57. <https://doi.org/10.1145/3236386.3241340>.
- Lundgard, Alan. "Measuring Justice in Machine Learning." *arXiv.2009.10050* (2020): 1-35. <https://doi.org/10.1145/3351095.3372838>
- Mehrabi, Ninareh, Fred Morstatter, Nripsuta Saxena, Kristina Lerman, and Aram Galstyan. "A Survey on Bias and Fairness in Machine Learning." *ACM Computing Surveys* 54, no. 6 (2021): 1-35. <https://doi.org/10.1145/3457607>.
- Monica, Andini, Ciani Emanuele, de Blasio Guido, D'Ignazio Alessio, and Salvestrini Viola. "Targeting with machine learning: An application to a tax rebate program in Italy." *Journal of Economic Behavior & Organization* 156, no. 1 (2018): 86-102. <https://doi.org/10.1016/J.JEBO.2018.09.010>
- Peter, Sekiswa, and Namuyanga Rebecca. "Data-Driven Approaches to Addressing Income Inequality: A Case Study of Nansana Municipality, Uganda." *Metropolitan Journal of Business & Economics (MJB)* 3, no. 10 (2024): 21-27.
- Ramakrishnan, Karthik. "AI-Driven Ethical Pricing: Balancing Profitability and Social Responsibility." *International Journal of Scientific Research in Computer Science, Engineering and Information Technology* 11, no. 1 (2025): 185-194. <https://doi.org/10.32628/cseit25111220>.
- Raventós, Daniel. *Basic income: The material conditions of freedom*. London: Pluto Press, 2007.
- Rockall, Emma J., Marina Mendes Tavares, and Carlo Pizzinelli. "AI Adoption and Inequality." *IMF Working Papers*, No. 2025/068 (2025). <https://doi.org/10.5089/9798229006828.001>
- Russell, Stuart J., and Peter Norvig. *Artificial Intelligence: A Modern Approach*. 3rd edition. Hoboken, NJ: Pearson, 2009.
- Violante, Giovanni L. "Skill-Biased Technical Change." In *The New Palgrave Dictionary of Economics*, edited by Steven N. Durlauf and Lawrence E. Blume, 1-6. London: Palgrave Macmillan, 2008. https://doi.org/10.1057/978-1-349-95121-5_2388-1.
- Wang, Xiaomeng, Yishi Zhang, and Ruilin Zhu. "A Brief Review on Algorithmic Fairness." *MSE* 1, no. 7 (2022): 1-13. <https://doi.org/10.1007/s44176-022-00006-z>.
- Zanzotto, Fabio Massimo. "Viewpoint: Human-in-the-Loop Artificial Intelligence." *Journal of Artificial Intelligence Research* 64 (2019): 243-252. <https://doi.org/10.1613/JAIR.1.11345>.
- Zheng, Stephan, Alexander Trott, Sunil Srinivasa, Nikhil Naik, Melvin Gruesbeck, David C. Parkes, and Richard Socher. "The AI Economist: Improving Equality and Productivity with AI-Driven Tax Policies." *arXiv: General Economics* (2020): 1-46. <https://arxiv.org/abs/2004.13332>.