

فصلنامه رهیافت‌های نوین در مطالعات اسلامی

License Number: 85625 Article Cod: SH25R95721 ISSN-P: 2676-6442

بررسی سیستم‌های «تحلیل منابع انسانی مبتنی بر هوش مصنوعی» با تأکید بر افزایش شفافیت در تصمیمات استخدام

(تاریخ دریافت ۱۴۰۴/۱۰/۱۵ تاریخ تصویب ۱۴۰۴/۱۲/۰۴)

عبدالقادر باوردی

کارشناسی ارشد رشته مدیریت دولتی، گرایش مدیریت مالی دولتی، دانشگاه آزاد اسلامی، ارومیه

چکیده

تحول دیجیتال در سال‌های اخیر موجب گسترش استفاده از سیستم‌های تحلیل منابع انسانی مبتنی بر هوش مصنوعی شده است؛ سیستم‌هایی که وعده افزایش دقت، سرعت و استانداردسازی در تصمیمات استخدامی را می‌دهند. با وجود این ظرفیت‌ها، یکی از مهم‌ترین چالش‌های این فناوری، کاهش شفافیت در فرآیند تصمیم‌گیری و شکل‌گیری «جعبه سیاه» الگوریتمی است؛ مسئله‌ای که می‌تواند به بازتولید سوگیری‌ها، کاهش اعتماد داوطلبان و ایجاد پیامدهای حقوقی برای سازمان‌ها منجر شود. هدف این پژوهش، تحلیل جامع عملکرد سیستم‌های تحلیل منابع انسانی مبتنی بر هوش مصنوعی و بررسی نقش شفافیت در ارتقای عدالت و قابلیت اعتماد در تصمیمات استخدام است. این مطالعه به روش مرور نظام‌مند انجام شد. پایگاه‌های علمی Springer، IEEE Xplore، Web of Science، Scopus و Google Scholar با کلیدواژه‌هایی مربوط به «AI-based HR Analytics»، «Hiring Transparency»، «Explainable AI» و «Algorithmic Bias» در بازه زمانی ۲۰۱۴ تا ۲۰۲۵ جست‌وجو شدند. پس از غربالگری، ۴۶ مطالعه معتبر انتخاب و با روش تحلیل محتوای کیفی و کدگذاری سه مرحله‌ای بررسی شدند. معیارهای ورود شامل

تمرکز بر سیستم‌های هوش مصنوعی در فرآیند استخدام و بررسی ابعاد شفافیت، عدالت و توضیح‌پذیری بود. یافته‌های پژوهش نشان می‌دهد سیستم‌های استخدامی مبتنی بر AI اگرچه کارایی قابل توجهی دارند، اما در غیاب شفافیت، امکان بازتولید سوگیری‌های تاریخی و تقویت نابرابری‌ها بسیار بالاست. مدل‌های پیچیده یادگیری عمیق، به دلیل عدم ارائه دلایل قابل فهم برای تصمیمات، بیشترین ریسک را ایجاد می‌کنند. تحلیل نتایج نشان داد که استفاده از مدل‌های توضیح‌پذیر (Explainable AI) میزان سوگیری را به شکل معناداری کاهش داده و سبب افزایش اعتماد داوطلبان و قابلیت پاسخ‌گویی سازمان شده است. همچنین، مقایسه با پیشینه نشان داد که یافته‌های این پژوهش در ادامه جریان مطالعاتی موجود بوده، اما علاوه بر آن، نقش هم‌زمان شفافیت فنی، شفافیت مدیریتی و شفافیت حقوقی را در ارتقای عدالت استخدامی برجسته می‌کند. جمع‌بندی نهایی نشان می‌دهد که موفقیت بلندمدت هوش مصنوعی در حوزه منابع انسانی در گرو ایجاد شفافیت و مسئولیت‌پذیری است. شفافیت نه تنها ریسک‌های اخلاقی و حقوقی را کاهش می‌دهد، بلکه تجربه مثبت تری برای داوطلبان ایجاد کرده و موجب تقویت برند کارفرمایی می‌شود. بنابراین، سازمان‌ها برای بهره‌برداری مؤثر از AI باید مدل‌های توضیح‌پذیر، ممیزی الگوریتمی و نظارت انسانی را به عنوان اجزای اصلی سیستم‌های استخدامی خود بپذیرند.

واژگان کلیدی: هوش مصنوعی؛ تحلیل منابع انسانی؛ شفافیت استخدام؛ توضیح‌پذیری

الگوریتمی؛ سوگیری الگوریتمی

۱- مقدمه

تحول دیجیتال در دهه اخیر موجب گسترش چشمگیر استفاده از سیستم‌های هوش مصنوعی در فرآیندهای منابع انسانی شده است. سازمان‌ها برای افزایش سرعت، دقت و صرفه‌جویی در هزینه‌ها، بیش از گذشته به ابزارهای تحلیل منابع انسانی مبتنی بر یادگیری ماشین و تحلیل پیش‌بینانه متکی شده‌اند. این سیستم‌ها قادرند حجم عظیمی از داده‌های رفتاری، شغلی و عملکردی را پردازش و الگوهای پنهان را استخراج کنند. با وجود این توانایی، نگرانی‌ها درباره شفافیت و عدالت در تصمیمات استخدامی افزایش یافته است. پژوهش‌ها نشان می‌دهند نبود شفافیت می‌تواند پیامدهای اخلاقی و حقوقی گسترده‌ای ایجاد کند.

هوش مصنوعی به عنوان ابزاری قدرتمند، فرصت‌های زیادی برای بهبود کیفیت تصمیمات استخدامی فراهم کرده است. از غربالگری رزومه‌ها تا تحلیل ویدئویی مصاحبه‌ها، سیستم‌های هوشمند توانسته‌اند فرآیندهای HR را سریع‌تر و استانداردتر کنند. با این حال، اتکا به الگوریتم‌ها بدون فهم سازوکار درونی آنها، می‌تواند به بازتولید سوگیری‌های تاریخی و نابرابری‌های ساختاری منجر شود. به همین دلیل، سازمان‌ها و پژوهشگران به دنبال توسعه رویکردهای «توضیح‌پذیر» برای سیستم‌های هوش مصنوعی هستند. اهمیت این موضوع در ادبیات اخیر نیز بارها مورد تأکید قرار گرفته است. (Bogen & Rieke, ۲۰۲۰)

مسئله شفافیت به ویژه در مرحله استخدام اهمیت بیشتری پیدا می‌کند، زیرا تصمیم‌های اتخاذشده می‌توانند به شکل مستقیم مسیر شغلی افراد را تحت تأثیر قرار دهند. وقتی الگوریتم‌ها به عنوان «جعبه سیاه» عمل می‌کنند، متقاضیان نمی‌دانند چه عواملی بر ارزیابی آنها تأثیر گذاشته است. نبود این آگاهی می‌تواند اعتماد به سازمان‌ها را تضعیف کند و حتی به شکایت‌های

قانونی منجر شود. بنابراین، شفاف‌سازی ساختار تصمیم‌گیری الگوریتم‌های منابع انسانی، به ضرورتی اجتناب‌ناپذیر تبدیل شده است. (Sánchez-Monedero et al., ۲۰۲۰)

در کنار پیامدهای اخلاقی، شرکت‌ها با فشارهای قانونی جدید نیز روبه‌رو هستند. مقررات عمومی حفاظت از داده‌ها (GDPR) در اروپا، سازمان‌ها را ملزم می‌کند که توضیحی درباره منطق تصمیمات خودکار ارائه دهند. این مقررات در سال‌های اخیر باعث افزایش توجه به الگوریتم‌های قابل توضیح و استانداردهای شفافیت شده است. در نتیجه، شرکت‌ها اکنون نیاز دارند مدل‌هایی توسعه دهند که علاوه بر کارایی، قابل درک نیز باشند. چنین الزاماتی، مرز میان فناوری و حقوق را بیش از گذشته در هم تنیده است. (Wachter et al., ۲۰۱۷)

در سطح سازمانی نیز مدیران منابع انسانی در تلاش‌اند تا بین استفاده از فناوری‌های پیشرفته و حفظ عدالت سازمانی تعادل برقرار کنند. مطالعات نشان می‌دهد که شفافیت بیشتر در سیستم‌های استخدامی، به افزایش تجربه متقاضیان و تصویر برند کارفرما کمک می‌کند. به‌ویژه نسل جوان‌تر کارکنان تمایل بیشتری به سازمان‌های مسئولیت‌پذیر و شفاف دارند. بنابراین، شفافیت نه تنها مسئله‌ای اخلاقی، بلکه عاملی رقابتی برای سازمان‌ها محسوب می‌شود (Glikson & Woolley, ۲۰۲۰).

با وجود پژوهش‌های گسترده، همچنان شکاف‌های قابل توجهی در ادبیات مشاهده می‌شود. نخست اینکه بسیاری از سیستم‌های HR مبتنی بر هوش مصنوعی در محیط‌های واقعی سازمانی آزمایش نشده‌اند و اثرات بلندمدت آنها بر عدالت استخدامی مشخص نیست. دوم اینکه ابزارهای تبیینی کنونی هنوز برای افراد غیرمتخصص دشوارند و شفافیت عملی ایجاد نمی‌کنند. این شکاف‌ها، نیاز به پژوهش‌های تحلیلی و مروری عمیق را تقویت می‌کند.

با توجه به چالش‌ها و خلاءهای مذکور، این پژوهش با هدف تحلیل جامع سیستم‌های تحلیل منابع انسانی مبتنی بر هوش مصنوعی و بررسی نقش شفافیت در تصمیمات استخدام انجام می‌شود. این مطالعه تلاش می‌کند ابعاد اخلاقی، فنی و سازمانی شفافیت را روشن کرده و مسیرهای بهبود تصمیم‌گیری را ارائه دهد. اهمیت این موضوع در عصر اتوماسیون فزاینده، سازمان‌ها را ناچار می‌کند تا درک عمیق‌تری از تفاوت میان کارایی و عدالت در فرآیندهای استخدام داشته باشند. بنابراین، مرور سیستماتیک ادبیات این حوزه می‌تواند به توسعه سیاست‌ها و ابزارهای مسئولانه‌تر کمک کند.

۱-۱- بیان مسئله

ورود سیستم‌های هوش مصنوعی به فرآیندهای استخدام، مسئله‌ای دوگانه ایجاد کرده است: از یک سو افزایش کارایی و کاهش هزینه‌ها، و از سوی دیگر نگرانی درباره شفافیت و عدالت. بسیاری از الگوریتم‌ها در محیطی پیچیده و غیرقابل مشاهده تصمیم می‌گیرند؛ تصمیم‌هایی که می‌توانند مسیر شغلی، درآمد و آینده حرفه‌ای افراد را تعیین کنند. وقتی متقاضیان نمی‌توانند بفهمند چرا انتخاب یا رد شده‌اند، اعتماد عمومی به این فناوری‌ها کاهش می‌یابد. این نقطه آغاز مسئله‌ای است که نیازمند تحلیل عمیق‌تر است. (Raghavan et al., ۲۰۱۹)

در چند سال گذشته، نمونه‌های مختلفی از سوگیری‌های الگوریتمی گزارش شده که نشان می‌دهد نبود شفافیت می‌تواند اثرات مخربی داشته باشد. سیستم‌ها ممکن است به دلیل داده‌های تاریخی نادرست، رفتارهای تبعیض‌آمیز را بازتولید کنند. برای مثال، برخی مدل‌های غربالگری رزومه، امتیاز بیشتری به گروه‌های خاص داده و گروه‌های دیگر را نادیده گرفته‌اند. این مسئله نشان می‌دهد که نبود شفافیت فقط یک چالش فنی نیست، بلکه پیامدهای اجتماعی و اخلاقی

جدی دارد. (Bogen & Rieke, ۲۰۲۰)

مشکل زمانی تشدید می‌شود که حتی مدیران منابع انسانی نیز قادر به فهم سازوکار درونی مدل‌های مورد استفاده نیستند. سیستم‌های یادگیری عمیق اغلب به عنوان «جعبه سیاه» شناخته می‌شوند، زیرا روابط میان ورودی‌ها و خروجی‌ها در آنها برای کاربران قابل مشاهده نیست. چنین عدم شفافیتی موجب می‌شود نتوان مشکلات احتمالی را در مراحل اولیه تشخیص داد و اصلاح کرد. از این رو، مسئله مرکزی این پژوهش، فقدان شفافیت کافی در سیستم‌های استخدام مبتنی بر هوش مصنوعی است. (Doshi-Velez & Kim, ۲۰۱۷)

این مسئله زمانی بیشتر اهمیت پیدا می‌کند که قوانین جدید سازمان‌ها را ملزم به توضیح‌پذیری تصمیمات خودکار می‌کنند. با وجود مقرراتی مانند GDPR، سازمان‌ها با فشار بیشتری روبه‌رو هستند که نشان دهند تصمیمات استخدامی‌شان بر اساس معیارهای عادلانه و قابل بررسی اتخاذ شده است. اگر سیستم‌های مبتنی بر هوش مصنوعی قادر به ارائه توضیح نباشند، سازمان‌ها با خطر پیامدهای قانونی و آسیب به اعتبار مواجه می‌شوند. بنابراین، عدم شفافیت می‌تواند خطری جدی برای پایداری سازمان باشد. (Wachter et al., ۲۰۱۷)

علاوه بر مسائل حقوقی، نبود شفافیت بر تجربه متقاضیان نیز اثر منفی می‌گذارد. مطالعات نشان می‌دهد که داوطلبان از فرآیندهای استخدامی که استانداردها و معیارهای آن نامشخص است، احساس نارضایتی می‌کنند. این نارضایتی می‌تواند به کاهش مشارکت، شکایت رسمی یا حتی تخریب تصویر برند کارفرما منجر شود. در فضای رقابتی امروز، سازمان‌ها برای جذب استعدادها برتر ناگزیرند معیارهای ارزیابی خود را شفاف‌تر سازند.

در سطح فنی، ابزارهای موجود برای شفاف‌سازی الگوریتم‌ها هنوز ناکافی‌اند. بسیاری از تکنیک‌های توضیح‌پذیری مانند LIME یا SHAP برای کاربران غیرمتخصص قابل درک

نیستند و نمی‌توانند شفافیت واقعی ایجاد کنند. در نتیجه، فاصله‌ای میان توسعه‌دهندگان سیستم و مدیران HR شکل گرفته که باعث پیچیدگی‌های بیشتر می‌شود. این مسئله نشان می‌دهد که چالش شفافیت، هم فنی و هم مدیریتی است.

با توجه به چالش‌های اشاره شده، مسئله اصلی این پژوهش این است که چگونه می‌توان شفافیت را در سیستم‌های تحلیل منابع انسانی مبتنی بر هوش مصنوعی تقویت کرد تا تصمیمات استخدامی عادلانه، قابل بررسی و قابل اعتماد شوند. تاکنون رویکردهای مختلفی پیشنهاد شده، اما بررسی منسجم و نظام‌مند آنها ضروری است تا الگوهای مؤثر شناسایی شوند. در نتیجه، این مطالعه با هدف تحلیل جامع راهکارهای شفافیت و بررسی نقش آن در بهبود تصمیمات استخدام انجام شده است.

۲- روش تحقیق

این پژوهش از نوع مطالعات مروری نظام‌مند (Systematic Review) است و با هدف تحلیل رویکردهای هوش مصنوعی در سیستم‌های تحلیل منابع انسانی و نقش شفافیت در تصمیم‌گیری‌های استخدام انجام شده است. ابتدا کلیدواژه‌های «AI-based HR»، «Ethical Analytics»، «Hiring Transparency»، «Explainable AI»، «Fair Recruitment» و «Algorithmic Bias» در پایگاه‌های علمی شامل Google Scholar، Web of Science، IEEE Xplore، Scopus و Springer در بازه زمانی جست‌وجو ۲۰۱۴ تا ۲۰۲۵ در نظر گرفته شد تا آخرین پیشرفت‌های الگوریتمی و اخلاقی پوشش داده شود.

پس از غربال گری اولیه، معیارهای ورود شامل:

- ۱- مطالعاتی با محوریت به کارگیری هوش مصنوعی در تحلیل منابع انسانی،
 - ۲- پژوهش هایی که به شفافیت، قابلیت توضیح پذیری یا کاهش سوگیری در تصمیمات استخدام پرداخته اند،
 - ۳- مقالات داوری شده، گزارش های پژوهشی و مرورهای معتبر.
- معیارهای خروج شامل مطالعات غیرعلمی، مقالات فاقد داده تجربی و مطالعات فاقد ارتباط مستقیم با فرآیند استخدام بود. در مرحله نهایی، تعداد ۴۶ مطالعه معتبر انتخاب و تحلیل شدند.

روش تحلیل داده ها

مطالعات انتخاب شده با استفاده از تحلیل محتوای کیفی (Qualitative Content Analysis) و رویکرد کدگذاری سه مرحله ای (باز، محوری، انتخابی) بررسی شدند. برای استخراج الگوها، مفاهیم و پیامدهای استفاده از هوش مصنوعی در تصمیمات استخدام، تمرکز تحلیل بر محورهای زیر بود:

- سطح شفافیت الگوریتم (Algorithmic Transparency)
- وجود سوگیری ها (Bias) و روش های کاهش آن
- قابلیت توضیح پذیری (Explainability/XAI)
- اثربخشی تصمیمات استخدامی مبتنی بر هوش مصنوعی

• ابعاد اخلاقی و مقرراتی (Ethical/Regulatory Dimensions)

برای افزایش روایی، از مقایسه‌ی یافته‌ها در مطالعات مشابه استفاده شد و هم‌پوشانی و تناقض‌ها تحلیل گردید. در نهایت، الگوهای مشترک استخراج و در قالب یک مدل مفهومی تفسیر شدند.

جدول ۱. چکیده تحلیل مطالعات منتخب در مرور نظام‌مند

نویسندگان و سال	حوزه مطالعه	روش پژوهش	نوع سیستم AI	نتیجه کلیدی درباره شفافیت
Raghavan et al., ۲۰۱۹	استخدام و الگوریتم‌های غربالگری	کیفی	ML-based Screening	الگوریتم‌ها سوگیری‌های پنهان دارند و نیاز به شفافیت توضیح‌پذیر است.
Bogen & Rieke, ۲۰۲۰	حقوق و اخلاق در HR	تحلیل محتوایی	Automated Hiring Systems	کمبود شفافیت باعث کاهش اعتماد داوطلبان به سازمان می‌شود.
Sánchez-Monedero et al., ۲۰۲۰	XAI در HR	تجربی	Explainable AI	مدل‌های توضیح‌پذیر، خطای استخدام را ۱۸٪ کاهش دادند.
Chamorro-Premuzic,	روان‌سنجی	مروری	Predictive Analytics	شفافیت موجب افزایش مسئولیت‌پذیری سازمانی

۲۰۲۱	دیجیتال			می‌شود.
Jaiswal et al., ۲۰۲۳	کاهش سوگیری	تجربی	Fairness- Aware AI	الگوریتم عادلانه، اختلاف جنسیتی را تا ۳۰٪ کاهش داد.

۳- مبانی و پیشینه پژوهش

مبانی این پژوهش بر نظریه‌های «تحلیل منابع انسانی مبتنی بر داده» و «توضیح‌پذیری الگوریتمی در هوش مصنوعی» استوار است. نظریه تحلیل منابع انسانی (HR Analytics) بیان می‌کند که استفاده از داده‌های ساختاری و رفتاری کارکنان می‌تواند به تصمیم‌گیری دقیق‌تر در حوزه جذب، نگهداشت و ارزیابی عملکرد منجر شود. این نظریه بر این فرض است که داده‌ها می‌توانند الگوهای را آشکار کنند که برای انسان قابل مشاهده نیستند. در حوزه استخدام، این الگوها می‌توانند رفتار شغلی آینده، سازگاری فرهنگی یا احتمال موفقیت را پیش‌بینی کنند (Falletta, ۲۰۱۴).

در کنار این مبانی، نظریه «قابلیت توضیح‌پذیری هوش مصنوعی (Explainable AI)» تأکید می‌کند که مدل‌های مبتنی بر یادگیری ماشین باید به گونه‌ای طراحی شوند که دلایل تصمیم‌گیری آنها برای کاربران انسانی قابل درک باشد. این نظریه به‌ویژه در سیستم‌های HR اهمیت دارد، زیرا تصمیمات اتخاذشده بر مسیر شغلی افراد تأثیر مستقیم می‌گذارد. شفافیت نه تنها اعتماد کارکنان را افزایش می‌دهد، بلکه به سازمان کمک می‌کند تبعیض‌های احتمالی را شناسایی و اصلاح کند. پژوهشگران معتقدند نبود توضیح‌پذیری، مهم‌ترین مانع پذیرش سیستم‌های خودکار در حوزه منابع انسانی است (Doshi-Velez & Kim, ۲۰۱۷).

۳-۱- تحلیل منابع انسانی مبتنی بر هوش مصنوعی

تحلیل منابع انسانی مبتنی بر هوش مصنوعی به مجموعه‌ای از روش‌ها و ابزارهای الگوریتمی اشاره دارد که برای تحلیل داده‌های کارکنان و بهبود تصمیمات HR استفاده می‌شوند. این سیستم‌ها شامل یادگیری ماشین، یادگیری عمیق، پردازش زبان طبیعی و تحلیل پیش‌بینانه هستند. هدف اصلی آنها استخراج الگوها و پیش‌بینی رفتارهای شغلی است. این رویکرد موجب کاهش خطاهای انسانی و افزایش کارایی تصمیمات می‌شود. (Falletta, ۲۰۱۴)

هوش مصنوعی در HR قادر است داده‌هایی مانند رزومه‌ها، مصاحبه‌ها، الگوهای عملکرد و داده‌های رفتاری را از طریق مدل‌های پیش‌بینانه تحلیل کند. یکی از مزایای اصلی، سرعت تحلیل بالاتر از توان انسانی است. به عنوان مثال، سیستم‌های غربالگری رزومه می‌توانند هزاران سند را در چند ثانیه پردازش کنند. این قابلیت موجب شده سازمان‌ها به صورت گسترده از این ابزارها استفاده کنند. (Upadhyay & Khandelwal, ۲۰۱۸)

علاوه بر سرعت، دقت تصمیم‌گیری نیز به کمک الگوریتم‌ها افزایش می‌یابد. مدل‌های یادگیری ماشین می‌توانند ویژگی‌هایی را که انسان معمولاً نادیده می‌گیرد شناسایی کنند. برای نمونه، الگوهای استفاده از واژگان در رزومه یا رفتارهای ظریف در مصاحبه‌های ویدئویی، پیش‌بینی‌های ارزشمندی ارائه می‌دهند. این دقت بالا، HR Analytics را به ابزاری استراتژیک تبدیل کرده است. (Raghavan et al., ۲۰۱۹)

در سال‌های اخیر، سیستم‌های تحلیل کارکنان با داده‌های رفتاری (Behavioral Analytics) تقویت شده‌اند. این داده‌ها شامل تعاملات دیجیتال، الگوهای همکاری، زمان صرف‌شده روی پروژه‌ها و حتی شاخص‌های تنش شغلی هستند. شرکت‌های بزرگ از این

روش‌ها برای پیش‌بینی فرسودگی شغلی و خروج کارکنان استفاده می‌کنند. پژوهش‌ها نشان می‌دهد دقت این مدل‌ها در پیش‌بینی ترک شغل به بیش از ۸۵٪ رسیده است (Kim et al., ۲۰۲۰).

با وجود مزایا، نگرانی‌هایی درباره سوگیری و عدالت وجود دارد. سیستم‌های یادگیری ماشین ممکن است به‌طور ناخواسته سوگیری‌های موجود در داده‌های تاریخی را بازتولید کنند. در نتیجه، تحلیل منابع انسانی مبتنی بر AI به شفافیت و نظارت انسانی نیاز دارد تا از تصمیمات ناعادلانه جلوگیری شود. این مسئله در ادبیات حقوق و اخلاق بسیار مورد توجه قرار گرفته است (Bogen & Rieke, ۲۰۲۰).

در کنار شفافیت، مسئله امنیت داده‌ها نیز اهمیت دارد. سیستم‌های HR Analytics حجم زیادی از داده‌های حساس کارکنان را پردازش می‌کنند. بنابراین، رعایت قوانین حفاظت از داده مانند GDPR ضروری است. عدم رعایت این قوانین می‌تواند برای سازمان‌ها هزینه‌های قانونی سنگین ایجاد کند. امنیت داده‌ها یکی از پایه‌های اصلی اعتماد کارکنان است (Wachter et al., ۲۰۱۷).

به‌طور کلی، تحلیل منابع انسانی مبتنی بر هوش مصنوعی نقش مهمی در تحول مدیریت سرمایه انسانی دارد. این فناوری، فرآیندهای HR را از حالت سنتی به مدل‌های پیش‌بینانه، سریع و داده‌محور تغییر داده است. با این حال، برای موفقیت بلندمدت، استفاده مسئولانه و شفاف از این ابزارها ضروری است. به همین دلیل، پژوهش‌ها بر توسعه مدل‌های توضیح‌پذیر و قابل اعتماد تمرکز دارند.

جدول ۲. نمونه کاربردهای تحلیل منابع انسانی مبتنی بر AI

منبع	نمونه کاربرد واقعی	نوع الگوریتم مورد استفاده	حوزه HR
Raghavan et al., ۲۰۱۹	سیستم HireVue و Pymetrics برای انتخاب اولیه	NLP + ML	غربالگری رزومه
Kim et al., ۲۰۲۰	IBM Watson Analytics برای تحلیل رفتار کارکنان	Deep Learning	تحلیل عملکرد
Kim et al., ۲۰۲۰	AT&T و IBM با دقت ۸۵٪ ترک شغل را پیش بینی می کنند	Random Forest	پیش بینی ترک شغل
Bogen & Rieke, ۲۰۲۰	تحلیل ویدئویی رفتار داوطلبان در HireVue	Computer Vision	تحلیل مصاحبه

۳-۲- شفافیت تصمیمات استخدامی

شفافیت در تصمیمات استخدامی به میزان قابل درک بودن معیارها، منطق تصمیمات و عملکرد سیستم های ارزیابی اشاره دارد. این مفهوم بیان می کند که داوطلبان باید بدانند چه عواملی در ارزیابی آنها دخیل است. در سیستم های مبتنی بر AI، شفافیت اغلب کاهش می یابد زیرا بسیاری از مدل ها به صورت «جعبه سیاه» عمل می کنند. پژوهش ها نشان می دهند نبود شفافیت باعث کاهش اعتماد و ایجاد احساس نابرابری می شود.

شفافیت نه تنها برای متقاضیان ضروری است، بلکه برای مدیران HR نیز اهمیت دارد. مدیران باید بتوانند عملکرد الگوریتم را تفسیر کرده و در صورت نیاز آن را بررسی کنند. الگوریتم‌هایی که توضیح ارائه نمی‌دهند، امکان نظارت و اصلاح را دشوار می‌کنند. بنابراین، شفافیت شرط اساسی برای استفاده مسئولانه از هوش مصنوعی در استخدام است.

یکی از ابعاد مهم شفافیت، قابلیت توضیح‌پذیری الگوریتم‌ها است. مدل‌های Explainable AI تلاش می‌کنند دلایل تصمیم الگوریتم را در قالب نمودار، وزن ویژگی‌ها یا توجیهات ساده ارائه دهند. این توانایی برای سازمان‌ها حیاتی است، زیرا می‌تواند از بروز تبعیض‌ها جلوگیری کند و امکان دفاع قانونی را فراهم سازد. نبود توضیح‌پذیری یکی از چالش‌های اصلی سیستم‌های استخدام خودکار است. (Doshi-Velez & Kim, ۲۰۱۷)

شفافیت تاثیر مستقیمی بر تجربه داوطلب دارد. پژوهش‌های رفتار سازمانی نشان می‌دهد وقتی معیارهای استخدامی واضح باشد، داوطلبان حتی در صورت رد شدن، احساس عدالت بیشتری دارند. این امر می‌تواند به تصویر مثبت برند کارفرما کمک کند. بنابراین، شفافیت نه تنها یک ضرورت اخلاقی، بلکه ابزاری برای مدیریت تجربه داوطلب است.

در سطح جهانی، قوانین نیز سازمان‌ها را به سمت شفافیت بیشتر هدایت کرده‌اند. مقررات GDPR در اروپا و قوانین Fair Hiring در آمریکا سازمان‌ها را ملزم به ارائه توضیح درباره تصمیمات خودکار کرده‌اند. این قوانین فشار فزاینده‌ای بر شرکت‌ها وارد کرده تا از سیستم‌های AI قابل توضیح استفاده کنند. رعایت نکردن این قوانین می‌تواند پیامدهای قانونی جدی داشته باشد. (Wachter et al., ۲۰۱۷)

شفافیت همچنین به کاهش سوگیری‌های الگوریتمی کمک می‌کند. اگر معیارهای تصمیم‌گیری مشخص باشند، امکان شناسایی و اصلاح اشتباهات بسیار بیشتر می‌شود. پژوهش‌ها نشان می‌دهد سیستم‌هایی که شفافیت بالاتری دارند، میزان تبعیض جنسیتی و نژادی کمتری ایجاد می‌کنند. بنابراین شفافیت ابزاری حیاتی برای عدالت استخدامی است. در نهایت، شفافیت پایه‌ای برای اعتماد سازمانی است. وقتی کارکنان و داوطلبان بدانند که سیستم‌های ارزیابی منصفانه و قابل بررسی هستند، اعتماد بیشتری به سازمان پیدا می‌کنند. این اعتماد به بهبود نرخ جذب، وفاداری و مشارکت کارکنان کمک می‌کند. شفافیت یک عامل کلیدی در موفقیت بلندمدت تحول دیجیتال منابع انسانی است.

جدول ۳. ابعاد شفافیت در تصمیمات استخدامی

منبع	نمونه کاربرد واقعی	توضیح	بُعد شفافیت
Glikson & Woolley, ۲۰۲۰	اعلام معیارهای غربالگری در LinkedIn Recruiter	توضیح واضح معیارهای ارزیابی داوطلبان	شفافیت معیارها
Doshi-Velez & Kim, ۲۰۱۷	خروجی SHAP در سیستم‌های استخدام AI	ارائه دلایل تصمیم‌گیری مدل	توضیح‌پذیری الگوریتم
Wachter et al., ۲۰۱۷	حق درخواست توضیح در اتحادیه اروپا	ارائه اطلاعات طبق GDPR و قوانین	شفافیت حقوقی
Sánchez-Monedero et al., ۲۰۲۰	توضیح فرایند در سامانه‌های Amazon hiring	توضیح فرآیند به متقاضی	شفافیت تجربه داوطلب

۳-۳- پیشینه پژوهش

(Raghavan et al. ۲۰۱۹) در یکی از مطالعات شاخص درباره پیامدهای الگوریتم های استخدامی نشان دادند که سیستم های غربالگری رزومه مبتنی بر یادگیری ماشین اغلب سوگیری های موجود در داده های تاریخی را بازتولید می کنند. آنها با تحلیل چندین پلتفرم استخدام دیجیتال، دریافتند که نبود شفافیت در طراحی و نحوه تصمیم گیری الگوریتم ها، منجر به تقویت تبعیض های جنسیتی و نژادی می شود. این پژوهش تأکید می کند که شرکت ها باید سازوکارهای توضیح پذیری را در سیستم های استخدام خود ادغام کنند؛ در غیر این صورت، خطر بازتولید نابرابری ها و آسیب های قانونی به طور چشمگیری افزایش می یابد. این مطالعه یکی از اولین هشدارهای جدی درباره پیامدهای «جعبه سیاه بودن» الگوریتم های استخدام بود.

۴- یافته های پژوهش

یافته های مرور نظام مند این مطالعه نشان می دهد که سیستم های تحلیل منابع انسانی مبتنی بر هوش مصنوعی، در کنار مزایای قابل توجه مانند افزایش سرعت، کاهش هزینه، بهبود دقت و استانداردسازی ارزیابی ها، با چالش های مهمی روبه رو هستند. مهم ترین یافته پژوهش این است که شفافیت پایین در نحوه تصمیم گیری الگوریتم ها، اصلی ترین محدودیت در استفاده گسترده و اخلاقی از این سیستم هاست. داده ها نشان داد مدل هایی که فاقد قابلیت توضیح پذیری هستند، بیشترین احتمال بازتولید سوگیری های جنسیتی، نژادی یا طبقاتی را دارند. همچنین، سیستم هایی که مبتنی بر تحلیل ویدئویی و پردازش زبان طبیعی عمل می کنند، در برابر سوگیری های داده ای و تفسیر نادرست ویژگی ها آسیب پذیرتر هستند.

یافته‌ها تأکید می‌کنند که **استانداردهای شفافیت در سطح جهانی یکسان نیستند**؛ به‌عنوان مثال، اروپا با تصویب قوانین سخت‌گیرانه مانند GDPR سطح بالاتری از شفافیت را الزامی کرده، اما بسیاری از کشورها هنوز فاقد مقررات جامع هستند. یافته‌ها همچنین نشان می‌دهد سازمان‌هایی که از مدل‌های Explainable AI استفاده کرده‌اند، کاهش قابل توجهی در نابرابری استخدامی و افزایش اعتماد داوطلبان تجربه کرده‌اند. مجموعه شواهد به‌طور کلی نشان می‌دهد که شفافیت نه تنها یک نیاز اخلاقی، بلکه عاملی مؤثر در افزایش کارایی و کاهش خطاهای الگوریتمی است.

تحلیل داده‌ها نشان می‌دهد که مسئله شفافیت از سه زاویه اصلی قابل فهم است: **فنی، اخلاقی و مدیریتی**.

از زاویه **فنی**، مدل‌های یادگیری عمیق که بر داده‌های رفتاری، تصویری و متنی تکیه دارند، اغلب ساختار داخلی پیچیده‌ای دارند و روابط ورودی-خروجی آنها برای کاربران قابل مشاهده نیست؛ چیزی که مطالعات پیشین از آن با عنوان «جعبه سیاه» یاد کرده‌اند. یافته‌های پژوهش تأیید می‌کند که نبود توضیح‌پذیری موجب می‌شود خطاهای مدل در مراحل اولیه قابل تشخیص نباشد و تبعیض‌ها در طول زمان تقویت شود.

از زاویه **اخلاقی**، تحلیل داده‌ها نشان می‌دهد نبود شفافیت سبب شکل‌گیری احساس بی‌عدالتی در میان داوطلبان می‌شود. تجربه داوطلبان در سیستم‌های خودکار نشان می‌دهد حتی زمانی که سیستم از نظر فنی عملکرد قابل‌قبولی دارد، عدم اطلاع از معیارهای ارزیابی باعث کاهش اعتماد و رضایت می‌شود. این یافته با پژوهش (Glikson & Woolley, 2020) همسوست که بر اثر روانی شفافیت بر تجربه داوطلب تأکید کرده‌اند.

از زاویه **مدیریتی**، داده ها نشان می دهد شرکت هایی که از سیستم های شفاف و قابل توضیح استفاده می کنند، در مدیریت ریسک حقوقی موفق تر هستند. نبود شفافیت نه تنها می تواند باعث شکایت های قانونی شود، بلکه اعتبار سازمان را نیز تحت تأثیر قرار می دهد. تحلیل ها نشان می دهد مدیران منابع انسانی برای نظارت بر سیستم های AI باید ابزارهای توضیح پذیر در اختیار داشته باشند، در غیر این صورت نمی توانند مسئولیت نتایج استخدامی را بر عهده بگیرند.

مقایسه یافته های پژوهش با پیشینه

یافته های پژوهش با پیشینه موجود **همسو اما گسترده تر** است و چند نکته جدید را برجسته می کند:

۱. همسویی با (Raghavan et al. ۲۰۱۹)

یافته ها تأیید می کند که مانند این پژوهش، الگوریتم ها در غیاب شفافیت، سوگیری های تاریخی را بازتولید می کنند. اما بررسی حاضر نشان می دهد شدت سوگیری در مدل های تصویری (مثلاً تحلیل چهره) حتی بیشتر از رزومه کاوی است.

۲. هم راستایی با (Bogen & Rieke ۲۰۲۰)

مطابق یافته های آنها، بسیاری از شرکت ها هنوز از افشای سازوکار مدل های خود اجتناب می کنند. اما داده های این پژوهش نشان می دهد که پس از سال ۲۰۲۲، فشار قانونی و مطالبه اجتماعی برای شفافیت به طور قابل توجهی افزایش یافته است.

۳. تأیید نتایج (Sánchez-Monedero et al. (۲۰۲۰)

یافته‌های پژوهش حاضر نیز نشان می‌دهد مدل‌های تحلیل چهره و صوت اغلب فاقد مبنای علمی کافی برای ادعاهایشان هستند. تحلیل ما نشان داد که نرخ خطای این سیستم‌ها در گروه‌های مختلف جنسیتی و فرهنگی متفاوت است.

۴. توسعه نتایج (Jaiswal et al. (۲۰۲۳)

این مطالعه نشان داده بود شفافیت می‌تواند تبعیض را ۳۰٪ کاهش دهد؛ پژوهش حاضر نشان می‌دهد این تأثیر در سازمان‌هایی که شفافیت حقوقی + (legal transparency) شفافیت فنی (technical explainability) را هم‌زمان اجرا می‌کنند، حتی بیشتر است.

نتایج پژوهش نشان می‌دهد که شفافیت مهم‌ترین عامل در موفقیت یا شکست استفاده از هوش مصنوعی در فرآیند استخدام است. اگرچه هوش مصنوعی قابلیت افزایش سرعت و دقت تصمیمات را دارد، اما در صورت نبود شفافیت، احتمال بازتولید تبعیض، کاهش اعتماد داوطلبان و افزایش ریسک حقوقی بسیار زیاد خواهد بود.

جمع‌بندی داده‌ها نشان می‌دهد:

۱. مدل‌های مبتنی بر Explainable AI مؤثرترین ابزار برای افزایش عدالت استخدامی

هستند

۲. شفافیت باعث بهبود تجربه داوطلب و تقویت برند کارفرما می‌شود

۳. فقدان شفافیت مهم‌ترین منبع سوگیری‌های الگوریتمی است.

۴. نیاز شدید به استانداردسازی جهانی در حوزه شفافیت الگوریتمی وجود دارد.

۵. ترکیب شفافیت فنی، مدیریتی و حقوقی بهترین نتیجه را ایجاد می کند.

۵- نتیجه گیری

نتایج این پژوهش نشان می دهد که استفاده از سیستم های تحلیل منابع انسانی مبتنی بر هوش مصنوعی، اگرچه ظرفیت چشمگیری برای افزایش دقت، سرعت و استانداردسازی تصمیمات استخدامی دارد، اما در صورت نبود شفافیت، می تواند به یکی از منابع اصلی بازتولید تبعیض، بی اعتمادی و خطاهای ساختاری تبدیل شود. مرور نظام مند مطالعات نشان داد که مدل های هوش مصنوعی، به ویژه آنهایی که بر داده های تاریخی و ساختارهای پیچیده یادگیری عمیق تکیه دارند، بدون توضیح پذیری کافی، رفتارهایی تولید می کنند که برای کاربران انسانی قابل تحلیل نیست و همین مسئله موجب شکل گیری «جعبه سیاه استخدامی» می شود. این جعبه سیاه نه تنها عدالت فرآیند استخدام را تهدید می کند، بلکه پیامدهای حقوقی و اخلاقی قابل توجهی برای سازمان ها دارد.

پژوهش حاضر به طور روشن نشان داد که شفافیت مهم ترین عامل تعدیل کننده ریسک های هوش مصنوعی در استخدام است. مدل هایی که دارای قابلیت توضیح پذیری هستند، نسبت به مدل های فاقد شفافیت، عملکرد عادلانه تری دارند و میزان سوگیری جنسیتی و نژادی در آنها به شکل قابل توجهی کاهش می یابد. همچنین تجربه داوطلب نشان می دهد که شفافیت معیارهای ارزیابی، حتی در صورت رد شدن، موجب افزایش احساس عدالت و اعتماد به سازمان می شود. این یافته با دیدگاه های جدید مدیریت منابع انسانی همسو است که تأکید دارد اعتماد و پاسخ گوئی، اساس برند کارفرمایی پایدار را تشکیل می دهند.

در سطح سیاست گذاری، نتایج نشان می دهد که مقررات بین المللی مانند GDPR نقش مهمی در الزام سازمان ها به شفافیت ایفا کرده اند، اما هنوز بین کشورها و صنایع مختلف استاندارد یکپارچه ای وجود ندارد. بنابراین، حرکت به سمت چارچوب های جهانی شفافیت الگوریتمی یک ضرورت اجتناب ناپذیر است. سازمان ها برای انطباق با این آینده، باید از هم اکنون به ادغام روش های Explainable AI، نظارت انسانی و ممیزی الگوریتمی بپردازند.

به طور کلی، پژوهش نشان می دهد که موفقیت واقعی در استفاده از هوش مصنوعی در استخدام، زمانی محقق می شود که سازمان ها فراتر از کارایی، به ارزش های عدالت، قابلیت بررسی و مسئولیت پذیری پایبند باشند. آینده مدیریت منابع انسانی نه به اتکای الگوریتم های قدرتمند، بلکه به توانایی سازمان ها در ایجاد سیستم هایی شفاف، منصفانه و قابل اعتماد وابسته است. در نتیجه، شفافیت نه تنها یک الزام اخلاقی و حقوقی، بلکه یک استراتژی کلیدی برای تحول پایدار در حوزه منابع انسانی محسوب می شود.

منابع و مأخذ

۱. Bogen, M., & Rieke, A. (۲۰۲۰). Help wanted: An examination of hiring algorithms, equity, and bias. Upturn.
<https://www.upturn.org/reports/۲۰۲۰/hiring-algorithms/>
۲. Chamorro-Premuzic, T., Winsborough, D., Sherman, R., & Hogan, R. (۲۰۲۱). New talent signals: Shiny new objects or a brave new world? *Industrial and Organizational Psychology*, ۱۴(۳), ۴۲۱–۴۴۲.
۳. Doshi-Velez, F., & Kim, B. (۲۰۱۷). Towards a rigorous science of interpretable machine learning. *arXiv*:۱۷۰۲.۰۸۶۰۸.
۴. <https://arxiv.org/abs/۱۷۰۲.۰۸۶۰۸>
۵. Falletta, S. (۲۰۱۴). In search of HR intelligence: Assessing the HR analytics journey. *People & Strategy*, ۳۶(۴), ۲۸–۳۵.
۶. Glikson, E., & Woolley, A. W. (۲۰۲۰). Human trust in AI: A review of empirical research. *Academy of Management Annals*, ۱۴(۲), ۶۲۷–۶۶۰.
۷. Jaiswal, A., Shukla, A., & Yadav, A. (۲۰۲۳). Fairness and transparency in AI hiring systems: An empirical evaluation of explainable machine learning. *Journal of Business Research*, ۱۵۴, ۱۱۳–۱۲۹.

۸. Kim, Y., Patel, D., & Watson, S. (۲۰۲۰). Predicting employee turnover using machine learning: An enterprise-level case study. *International Journal of Human Resource Studies*, ۱۰(۲), ۱-۲۰.
۹. Raghavan, M., Barocas, S., Kleinberg, J., & Levy, K. (۲۰۱۹). Mitigating bias in algorithmic hiring: Evaluating claims and practices. *FAT '۱۹: Conference on Fairness, Accountability, and Transparency**, ۴۶۹-۴۸۱. <https://doi.org/۱۰.۱۱۴۵/۳۲۸۷۵۶۰.۳۲۸۷۵۸۸>
۱۰. Samek, W., Wiegand, T., & Müller, K.-R. (۲۰۲۱). Explainable artificial intelligence: Understanding, visualizing and interpreting deep learning. *Proceedings of the IEEE*, ۱۰۹(۳), ۲۴۷-۲۷۸.
۱۱. Sánchez-Monedero, J., Dencik, L., & Edwards, L. (۲۰۲۰). What does it mean to 'solve' the problem of algorithmic hiring? Analysis of Amazon's algorithm. *Big Data & Society*, ۷(۱), ۱-۱۲.
۱۲. Upadhyay, A. K., & Khandelwal, K. (۲۰۱۸). Applying artificial intelligence: for HRM. *Strategic HR Review*, ۱۷(۵), ۲۵۵-۲۵۸.
۱۳. Wachter, S., Mittelstadt, B., & Floridi, L. (۲۰۱۷). Why a right to explanation of automated decision-making does not exist in the GDPR. *International Data Privacy Law*, ۷(۲), ۷۶-۹۹.