

Organizational Reconfiguration in the Era of Digital Transformation: A Futures Studies Approach

Ruhollah Bahreini *

PhD in Public Administration, Decision Making and Policy Making, Ct.C., Tehran, Iran.

Abstract

With the emergence of the digital transformation paradigm, organizations are on the verge of fundamental changes in their structures, processes, and governance models. This study was conducted to explore the potential pathways facing organizations in the era of digital transformation and to identify the key drivers influencing their future, thereby developing a roadmap for intelligently addressing forthcoming uncertainties. In terms of its purpose, this study is applied and adopts a futures-oriented approach. Data were collected through brainstorming sessions and participatory workshops involving experts in the fields of management and information technology who were selected using purposive sampling. The main research strategy was based on the scenario planning method using an uncertainty matrix. Accordingly, the identified drivers were ranked based on two indicators: degree of impact and level of uncertainty. Ultimately, two key variables, level of regulation and rate of technology adoption, were selected as the principal axes for constructing the scenario matrix. Data analysis resulted in the development of four distinct scenarios for the future of organizations in the era of digital transformation: (1) Golden Cage, characterized by strict regulatory governance combined with slow technology adoption; (2) Smart Collaboration, reflecting a balance between regulatory oversight and rapid adoption of digital tools; (3) Digital Chaos, characterized by the absence of regulatory frameworks alongside weak and fragmented adoption; and (4) Innovation Paradise, representing an environment with minimal regulation and a high pace of digital transformation. Each of these scenarios outlines different challenges and opportunities for organizational managers and policymakers.

Keywords: Organizational Reconfiguration, Digital Transformation, Futures Studies, Scenario Planning; Key Drivers

How to Cite: Bahreini,R . (2025). Organizational Reconfiguration in the Era of Digital Transformation: A Futures Studies Approach. Journal of Intelligent Strategic Management . 4(3), 829-851.

doi: 10.87453/bumara.2026.373601.407



Intelligent Strategic Management (JISM) in Development and Evolution is licensed under a Creative Commons Attribution-Non Commercial 4.0 International License.

© Authors

* Corresponding Author: Rozphd1402@gmail.com

Original Research

Received: 2025/08/04

Review: 2025/08/28

Accepted: 2025/09/23

eISSN: 2891-4128

بازآرایی سازمان در عصر تحول دیجیتال با رویکرد آینده‌پژوهی

دکتری مدیریت دولتی، گرایش تصمیم‌گیری و خط و مشی‌گذاری، واحد تهران مرکزی، تهران، ایران.

روح الله بحرینی*

چکیده

با ظهور پارادایم تحول دیجیتال، سازمان‌ها در آستانه تغییراتی بنیادین در ساختارها، فرآیندها و مدل‌های حکمرانی خود قرار گرفته‌اند. این پژوهش با هدف واکاوی مسیرهای محتمل پیش روی سازمان‌ها در عصر تحول دیجیتال و شناسایی پیشران‌های کلیدی مؤثر بر آینده آن‌ها انجام شده است تا نقشه راهی برای مواجهه هوشمندانه با عدم قطعیت‌های پیش رو ترسیم نماید. این پژوهش از نظر هدف، کاربردی و با رویکرد آینده‌پژوهانه انجام شده است. جهت گردآوری داده‌ها، از جلسات طوفان فکری و کارگاه‌های مشارکتی با حضور خبرگان حوزه مدیریت و فناوری اطلاعات که به صورت هدفمند انتخاب شده بودند، استفاده شد. استراتژی اصلی پژوهش بر پایه روش سناریونویسی (مبنی بر ماتریس عدم قطعیت) استوار است. در این راستا، پیشران‌های مستخرج بر اساس دو شاخص «میزان اثرگذاری» و «میزان عدم قطعیت» رتبه‌بندی شدند و در نهایت دو متغیر کلیدی «سطح مقررات‌گذاری» و «نرخ پذیرش تکنولوژی» به عنوان محورهای اصلی تشکیل ماتریس سناریو انتخاب گردیدند. تحلیل داده‌ها منجر به تدوین چهار سناریوی متمایز برای آینده سازمان در عصر تحول دیجیتال شد: ۱) قفس طلایی (حاکمیت مقررات سخت‌گیرانه همراه با پذیرش کند فناوری)، ۲) همکاری هوشمند (توازن میان نظارت‌های قانونی و پذیرش سریع ابزارهای دیجیتال)، ۳) آشوب دیجیتال (فقدان چارچوب‌های قانونی در کنار پذیرش ضعیف و پراکنده) و ۴) بهشت نوآوری (محیطی با مقررات حداقلی و شتاب بالای تحول دیجیتال). هر یک از این سناریوها، چالش‌ها و فرصت‌های متفاوتی را برای مدیران و سیاست‌گذاران سازمانی ترسیم می‌کنند.

کلیدواژه‌ها: بازآرایی سازمان، تحول دیجیتال، آینده‌پژوهی، سناریونویسی، پیشران‌های کلیدی

استناد به این مقاله: روح الله، بحرینی، (۱۴۰۴). بازآرایی سازمان در عصر تحول دیجیتال با رویکرد آینده‌پژوهی. مدیریت استراتژیک هوشمند، ۴(۳)، ۸۲۹-۸۵۱.



مدیریت استراتژیک هوشمند (JISM) در توسعه و تکامل تحت مجوز بین‌المللی کپی‌رایت کامنز با شرایط انتساب-غیرتجاری ۴٫۰ منتشر می‌شود.
©نویسندگان

* نویسنده مسئول: Rozphd1402@gmail.com

مقدمه

در دهه‌های اخیر، تحول دیجیتال به‌عنوان پارادایمی بنیادین و فراتر از یک موج فناورانه، به‌سرعت از یک ضرورت حاشیه‌ای به نیروی راهبردی در بازآفرینی ساختارها، فرایندها و شیوه‌های تصمیم‌گیری سازمان‌ها تبدیل شده است؛ به گونه‌ای که امروز بحث درباره آینده سازمان، بدون توجه به ابعاد و پیامدهای زیست‌بوم دیجیتال، ناقص و نابسند خواهد بود. سازمان‌های معاصر در فضایی فعالیت می‌کنند که با تغییرات شتابان محیطی، رقابت فزاینده، انفجار داده‌ها و انتظارات نوظهور ذی‌نفعان همراه است. در چنین شرایطی، تحول دیجیتال با بهره‌گیری از مجموعه‌ای از فناوری‌های نوظهور به‌ویژه هوش مصنوعی نه تنها ابزاری برای بهبود بهره‌وری و تحلیل هوشمند اطلاعات است، بلکه محرکی برای پیش‌بینی روندها، خودکارسازی هوشمند، ارتقای چابکی و بازتعریف روابط انسانی در محیط کار محسوب می‌شود.

با این حال، تأثیر تحول دیجیتال صرفاً به دیجیتالی کردن فرآیندهای موجود محدود نمی‌شود؛ بلکه این دگرگونی، پرسش‌های بنیادینی درباره آینده ساختارهای سازمانی، بازآرایی جایگاه نیروی انسانی در کنار سیستم‌های هوشمند، اخلاق در حکمرانی داده‌ها، امنیت سایبری و مرز میان شهود انسانی و محاسبات ماشین در مدیریت مطرح می‌سازد. از این‌رو، بررسی آینده سازمان در عصر تحول دیجیتال، نه تنها تلاشی برای شناخت روندهای نوین تکنولوژیک، بلکه کوششی است برای درک مسیر تکامل سازمان‌ها در جهانی که در آن موفقیت بیش از هر زمان دیگری وابسته به توانایی سازگاری، یادگیری مستمر و بهره‌گیری استراتژیک از ظرفیت‌های بی‌پایان عصر دیجیتال خواهد بود. تحقق تحول دیجیتال برای سازمان‌ها مزایای راهبردی فراوانی به همراه دارد و می‌تواند به‌عنوان حیاتی‌ترین عامل افزایش کارایی و قدرت رقابت در بازارهای مدرن عمل کند. این تحول با یکپارچه‌سازی فناوری‌های نوین در بدنه سازمان، به کارکنان فرصت می‌دهد تا از قید وظایف سنتی رها شده و بر فعالیت‌های خلاقانه و ارزش‌آفرین تمرکز کنند. همچنین، با تبدیل داده‌های خام به بصیرت‌های مدیریتی، رهبران را در اتخاذ تصمیمات دقیق‌تر، کم‌ریسک‌تر و داده‌محور یاری می‌دهد. افزون بر این، بلوغ دیجیتال می‌تواند دقت عملیات را ارتقا داده، تجربه‌ای شخصی‌سازی شده و ۲۴ ساعته برای مشتریان خلق کند، هزینه‌های عملیاتی را بهینه‌سازی نماید و تاب‌آوری سازمان را در برابر بحران‌های پیش‌بینی نشده افزایش دهد. در نتیجه، گذار موفق به عصر دیجیتال نه تنها موجب افزایش بهره‌وری

لحظه‌ای می‌شود، بلکه زیربنای نوآوری مستمر، چابکی ساختاری و پایداری سازمان در سپهر رقابتی آینده را فراهم می‌سازد.

با وجود مزایای فراوان هوش مصنوعی، استفاده از آن در سازمان‌ها با چالش‌ها و مشکلات جدی نیز همراه است؛ از جمله مهم‌ترین این چالش‌ها می‌توان به کاهش برخی مشاغل، تغییر ماهیت نقش‌های انسانی، وابستگی بیش از حد به الگوریتم‌ها، و دشواری تطبیق نیروی انسانی با فناوری‌های جدید اشاره کرد، به گونه‌ای که برخی پژوهشگران هشدار داده‌اند که اتوماسیون هوشمند می‌تواند موجب جابه‌جایی شغلی و افزایش نیاز به بازآموزی کارکنان شود (Brynjolfsson و McAfee، 2014). همچنین، هوش مصنوعی ممکن است در صورت آموزش بر داده‌های ناقص یا جانبدارانه، تصمیم‌هایی ناعادلانه یا غیرشفاف تولید کند و این مسئله اعتماد سازمانی را کاهش دهد؛ موضوعی که در بحث «سوگیری الگوریتمی» بسیار مورد توجه قرار گرفته است (O'Neil، 2016). از سوی دیگر، پیچیدگی فنی، هزینه‌های بالای پیاده‌سازی، نیاز به زیرساخت مناسب، و کمبود تخصص لازم برای مدیریت سامانه‌های هوشمند، از موانع مهم استفاده مؤثر از این فناوری در سازمان‌ها محسوب می‌شود (Davenport و Ronanki، 2018). افزون بر این، نگرانی‌های مربوط به امنیت اطلاعات، حفظ حریم خصوصی، و احتمال نشت داده‌های حساس نیز از چالش‌های اساسی هستند که می‌توانند پیامدهای حقوقی و اعتباری برای سازمان به دنبال داشته باشند (Mittelstadt و همکاران، ۲۰۱۶). در نهایت، اگر هوش مصنوعی بدون نظارت انسانی کافی و بدون چارچوب‌های اخلاقی روشن به کار گرفته شود، ممکن است به جای افزایش بهره‌وری، به ایجاد سردرگمی، مقاومت کارکنان، و اختلال در فرهنگ سازمانی منجر شود؛ بنابراین، موفقیت در بهره‌برداری از این فناوری مستلزم برنامه‌ریزی دقیق، آموزش مستمر، و تدوین سیاست‌های شفاف و مسئولانه است.

با توجه به گسترش سریع هوش مصنوعی و نفوذ آن در بخش‌های مختلف سازمانی، آینده سازمان‌ها بیش از هر زمان دیگری با ابهام و عدم قطعیت مواجه شده است؛ زیرا از یک سو این فناوری می‌تواند موجب افزایش بهره‌وری، سرعت تصمیم‌گیری، بهبود خدمات و ارتقای رقابت‌پذیری شود، اما از سوی دیگر پیامدهایی مانند کاهش برخی مشاغل، تغییر ساختار نیروی انسانی، وابستگی به الگوریتم‌ها، بروز خطاهای تصمیم‌گیری، و چالش‌های اخلاقی و حقوقی را نیز به همراه دارد. در نتیجه، سازمان‌ها در شرایطی قرار گرفته‌اند که همزمان با فرصت‌های گسترده و تهدیدهای جدی روبه‌رو هستند و هنوز به‌طور دقیق روشن

نیست که هوش مصنوعی در بلندمدت چه تأثیری بر ساختار، فرهنگ، مدیریت و کارکردهای اساسی آن‌ها خواهد گذاشت. این مسئله به‌ویژه از آن جهت اهمیت دارد که نبود شناخت کافی از پیامدهای آتی این فناوری می‌تواند موجب تصمیم‌گیری‌های نادرست، مقاومت کارکنان، ناهماهنگی سازمانی و حتی اختلال در پایداری سازمان شود؛ بنابراین، بررسی آینده نامعلوم سازمان در پرتو هوش مصنوعی ضرورتی اساسی برای فهم مسیر تحول سازمان‌ها و طراحی راهبردهای مناسب در مواجهه با این فناوری نوظهور به شمار می‌آید.

لذا با توجه به اهمیت و در عین حال نامعلوم بودن پیامدهای هوش مصنوعی بر آینده سازمان‌ها، انجام آینده‌پژوهی در این زمینه ضرورتی انکارناپذیر به نظر می‌رسد؛ زیرا آینده‌پژوهی این امکان را فراهم می‌سازد که به‌جای واکنش‌های مقطعی و کوتاه‌مدت، روندهای مؤثر، عدم قطعیت‌های کلیدی و نیروهای پیشران تحول شناسایی شوند و بر اساس آن سناریوهای مختلفی از آینده محتمل سازمان ترسیم گردد. این سناریوها می‌توانند به مدیران و تصمیم‌گیران کمک کنند تا برای شرایط گوناگون، از جمله گسترش گسترده هوش مصنوعی، محدود شدن کاربرد آن، یا شکل‌گیری مدل‌های ترکیبی انسان و ماشین، آمادگی لازم را داشته باشند و راهبردهای متناسبی در حوزه ساختار، منابع انسانی، فناوری، اخلاق و فرهنگ سازمانی تدوین کنند. به بیان دیگر، آینده‌پژوهی در این حوزه ابزاری راهبردی برای کاهش غافلگیری، افزایش تاب‌آوری سازمان، تقویت قدرت انطباق و فراهم‌سازی زمینه تصمیم‌گیری آگاهانه در برابر تحولات سریع و پیچیده هوش مصنوعی است.

مبانی نظری

تحول دیجیتال

تحول دیجیتال در سازمان‌ها امروز بیش از هر زمان دیگری به یک مسئله راهبردی و چندبعدی تبدیل شده است؛ زیرا پژوهش‌های جدید نشان می‌دهند که گرچه هوش مصنوعی و ابزارهای دیجیتال می‌توانند بهره‌وری، نوآوری و چابکی سازمانی را افزایش دهند، اما موفقیت در این مسیر به‌شدت به آمادگی نیروی انسانی، زیرساخت مناسب، داده‌های باکیفیت و حکمرانی درست وابسته است (McKinsey, 2024؛ OECD, 2024). در عین حال، موانعی مانند شکاف مهارتی، محدودیت بودجه، زیرساخت ناکافی و مقاومت در برابر تغییر همچنان از مهم‌ترین چالش‌های تحول دیجیتال محسوب می‌شوند.

و می‌توانند باعث شوند بسیاری از سازمان‌ها از مرحله آزمایش و پایلوت فراتر نروند (Insight, 2024؛ BCG, 2024). افزون بر این، پژوهش‌های تازه تأکید دارند که ارزش واقعی تحول دیجیتال زمانی آشکار می‌شود که سازمان‌ها فرآیندها را بازطراحی کنند، آموزش و توانمندسازی کارکنان را جدی بگیرند، و هم‌زمان به امنیت، اخلاق و مسئولیت‌پذیری در به‌کارگیری هوش مصنوعی توجه کنند (Ajiboyev, 2024؛ McKinsey, 2025). بنابراین، مبانی نظری تحول دیجیتال در پژوهش‌های جدید بیش از پیش بر پیوند میان فناوری، انسان، فرهنگ و راهبرد تمرکز دارد و نشان می‌دهد که تحول دیجیتال صرفاً یک تغییر فناورانه نیست، بلکه دگرگونی در شیوه اداره، یادگیری و خلق ارزش در سازمان است (OECD, 2024؛ BCG, 2025).

هوش مصنوعی

هوش مصنوعی در سال‌های اخیر به یکی از تعیین‌کننده‌ترین نیروهای تحول در سازمان‌ها تبدیل شده است، به گونه‌ای که پژوهش‌های جدید نشان می‌دهند این فناوری می‌تواند بهره‌وری، نوآوری و سرعت تصمیم‌گیری را به‌طور چشمگیری افزایش دهد، اما بهره‌برداری واقعی از آن همچنان برای بسیاری از سازمان‌ها در مرحله آزمایش و پایلوت باقی مانده است (McKinsey, 2025؛ MIT Sloan Management Review & Boston Consulting Group, 2024). در عین حال، موفقیت در به‌کارگیری هوش مصنوعی تنها به پیشرفت فنی وابسته نیست، بلکه به بازطراحی فرایندها، آمادگی نیروی انسانی، حمایت رهبری، حکمرانی مناسب و ایجاد مهارت‌های جدید در کارکنان نیز نیاز دارد؛ به همین دلیل، بسیاری از پژوهش‌ها تأکید می‌کنند که چالش اصلی، بیشتر انسانی و سازمانی است تا صرفاً فناورانه (McKinsey, 2025؛ Wharton Human-AI Research & GBK Collective, 2025). افزون بر این، مطالعات جدید نشان می‌دهند که هوش مصنوعی در کنار فرصت‌های گسترده، نگرانی‌هایی مانند جابه‌جایی شغلی، ابهام در مسئولیت‌پذیری، سوگیری الگوریتمی، مسائل اخلاقی و کاهش شفافیت تصمیم‌ها را نیز به همراه دارد و همین امر ضرورت تدوین چارچوب‌های مسئولانه برای استفاده از آن را برجسته می‌سازد (Stanford HAI, 2024؛ APA, 2024). بنابراین، هوش مصنوعی نه تنها یک ابزار فناورانه، بلکه پدیده‌ای راهبردی است که می‌تواند ساختار، فرهنگ، روابط کاری و آینده سازمان را دگرگون کند و از این رو، درک مبانی نظری آن برای تبیین

آینده سازمان‌ها اهمیت اساسی دارد (MIT Sloan Management Review & Boston Consulting Group, 2024; Stanford HAI, 2024).

پذیرش فناوری

پذیرش فناوری یکی از مبانی نظری مهم در تبیین نحوه به کارگیری هوش مصنوعی در سازمان‌هاست؛ زیرا پژوهش‌ها نشان می‌دهند که استفاده موفق از فناوری‌های نو تنها به پیشرفت فنی وابسته نیست، بلکه به ادراک کاربران از سودمندی و سهولت استفاده نیز بستگی دارد (Davis, 1989). در ادامه این رویکرد، مدل‌های توسعه‌یافته‌تری مانند TAM2 و UTAUT نشان داده‌اند که عوامل اجتماعی، شرایط تسهیل‌کننده، انتظار عملکرد، انتظار تلاش و نفوذ اجتماعی نیز نقش تعیین‌کننده‌ای در پذیرش فناوری دارند (Venkatesh & Davis, 2000; Venkatesh et al., 2003). پژوهش‌های جدید درباره هوش مصنوعی نیز این چارچوب را تأیید کرده‌اند و نشان داده‌اند که در پذیرش ابزارهای هوش مصنوعی، ادراک سودمندی، سهولت استفاده، اعتماد، ریسک ادراک شده و حمایت سازمانی از مهم‌ترین پیش‌بین‌ها هستند (Alhashmi et al., 2024; Chen et al., 2024). بنابراین، درک پذیرش فناوری برای پژوهش‌های مرتبط با هوش مصنوعی ضروری است، زیرا مشخص می‌کند کارکنان و مدیران تا چه حد حاضرند از این فناوری استفاده کنند و چه عواملی می‌توانند این فرایند را تسهیل یا مختل کنند (Bazine, 2025; Heliyon, 2024).

روش تحقیق

این پژوهش از نظر هدف، کاربردی و از نظر رویکرد، آینده‌پژوهانه است و با روش سناریونویسی انجام می‌شود؛ روشی که در آن به جای پیش‌بینی یک آینده قطعی، مجموعه‌ای از آینده‌های ممکن و محتمل بر پایه عدم قطعیت‌ها و نیروهای محرک ترسیم می‌شود (Mitsui & Co. Global Strategic Studies Institute, 2024; Luesink et al., 2024).

جامعه آماری پژوهش را خبرگان و متخصصان آشنا با موضوع تشکیل می‌دهند که با نمونه‌گیری هدفمند و بر اساس معیارهایی مانند سابقه علمی و اجرایی، تخصص مرتبط، و توانایی تحلیل آینده انتخاب می‌شوند؛ در پژوهش‌های جدید خبره‌محور نیز تأکید شده است که انتخاب خبرگان باید با توجه به تنوع دیدگاه، اعتبار تخصصی، و میزان آشنایی آنان با مسئله صورت گیرد تا سوگیری و هم‌پوشانی دیدگاه‌ها به حداقل برسد (National

طوفان فکری/برگزاری کارگاه مشارکتی با خبرگان گردآوری می‌شود تا پیشران‌ها، روندها و عوامل مؤثر بر موضوع شناسایی شوند؛ این رویکرد با ادبیات جدید آینده‌پژوهی هم‌راستاست که بر مشارکت محوری، گفت‌وگوی ساختاریافته و استفاده از دانش خبرگان برای کشف عدم قطعیت‌ها تأکید دارد (Bengston, 2019; IZT, 2024). سپس پیشران‌های استخراج‌شده بر اساس دو شاخص «میزان اثرگذاری» و «میزان عدم قطعیت» ارزیابی و رتبه‌بندی می‌شوند و دو عامل دارای بیشترین اثرگذاری و بیشترین عدم قطعیت به عنوان متغیرهای کلیدی برای تشکیل ماتریس سناریو انتخاب می‌گردند؛ این منطق دقیقاً با منطق سناریونویسی کلاسیک و نیز نسخه‌های به‌روزتر آن سازگار است که در آن، عوامل کلیدی برحسب اهمیت و عدم قطعیت مشخص شده و مبنای ساخت چهار سناریوی اصلی قرار می‌گیرند (Mitsui & Co. Global Strategic Studies Institute, 2024; Wright & Cairns, 2024). در گام نهایی، سناریوها بر پایه ترکیب وضعیت‌های ممکن دو متغیر منتخب تدوین، توصیف و با نظر مجدد خبرگان اعتبارسنجی می‌شوند تا از انسجام، واقع‌گرایی و قابلیت استفاده راهبردی آن‌ها اطمینان حاصل شود (Luesink et al., 2024; Ralston et al., 2024).

یافته‌های پژوهش

شناسایی پیشران‌ها و روندهای کلیدی

در این مرحله با استفاده از روش طوفان فکری در بین گروهی ۱۲ نفره از خبرگان در این حوزه ۶ عامل مهم به شرح ذیل کشف شد:

جدول ۱: پیشران‌ها و زیرمجموعه آن‌ها

پیشران‌ها	زیرمجموعه
فناوری	سرعت پیشرفت الگوریتم‌های هوش مصنوعی و یادگیری عمیق
	ظهور هوش مصنوعی مولد (Generative AI) و کاربردهای آن
	یکپارچه‌سازی هوش مصنوعی با اینترنت اشیا (IoT) و رایانش ابری
	بلوغ فناوری‌های پردازش زبان طبیعی و بینایی ماشین
اقتصادی و مالی	دسترسی به زیرساخت‌های محاسباتی و قدرت پردازشی
	هزینه پیاده‌سازی و نگهداری سیستم‌های هوش مصنوعی
	بازده سرمایه‌گذاری (ROI) در پروژه‌های هوش مصنوعی
	تغییر مدل‌های درآمدی و کسب‌وکار مبتنی بر داده
نیروی انسانی و مهارتی	رقابت‌پذیری سازمان‌های مجهز به هوش مصنوعی
	دسترسی به سرمایه و بودجه برای تحول دیجیتال
	سطح سواد دیجیتال و هوش مصنوعی کارکنان
	نیاز به مهارت‌های جدید (تحلیل داده، کار با ابزارهای AI)
قانونی و مقرراتی	سرعت و عمق پذیرش هوش مصنوعی در سازمان
	تغییر نقش‌های شغلی و حذف یا ایجاد مشاغل جدید
	برنامه‌های آموزش و بازآموزی نیروی کار
	قوانین حریم خصوصی و حفاظت از داده‌ها
اجتماعی و فرهنگی	مقررات اخلاقی استفاده از هوش مصنوعی
	میزان مقررات‌گذاری و محدودیت‌های قانونی
	مسئولیت حقوقی تصمیمات گرفته‌شده توسط هوش مصنوعی
	استانداردهای بین‌المللی و سازوکارهای نظارتی
رقابتی و بازار	اعتماد عمومی به تصمیمات هوش مصنوعی
	فرهنگ سازمانی نسبت به نوآوری و ریسک‌پذیری
	نگرش جامعه به اتوماسیون و جایگزینی انسان با ماشین
	ارزش‌های اخلاقی و هنجارهای اجتماعی در پذیرش فناوری
	تنوع نسلی در سازمان و تفاوت دیدگاه‌ها نسبت به AI
	ورود رقبای جدید مبتنی بر هوش مصنوعی (استارت‌آپ‌ها)
	سرعت نوآوری رقبای صنعت
	انتظارات مشتریان از خدمات شخصی‌سازی‌شده
	تغییر زنجیره ارزش صنعت با حضور هوش مصنوعی
	همکاری‌های استراتژیک و اکوسیستم‌های فناوری

پس از کشف پیشران ها، آن ها را با استفاده از جدول تاثیر-عدم قطعیت، برای بررسی سناریوها انتخاب می کنیم. در این روش با کمک خبرگان به هر کدام از پیشران ها از ۱ تا ۵ امتیاز می دهیم.

سپس به هر یک از ۶ پیشران شما، دو امتیاز ۱ تا ۵ دادیم:

- امتیاز تاثیر: چقدر این پیشران بر آینده سازمان تاثیر دارد؟
- امتیاز عدم قطعیت: چقدر آینده این پیشران نامعلوم است؟

جدول ۲: انتخاب پیشران ها

پیشران	تأثیر (۵-۱)	عدم قطعیت (۵-۱)	جمع	اولویت
فناوری	۵	۳	۸	متوسط
اقتصادی و مالی	۴	۴	۸	متوسط
نیروی انسانی و مهارتی	۵	۵	۱۰	خیلی بالا
قانونی و مقرراتی	۴	۵	۹	خیلی بالا
اجتماعی و فرهنگی	۳	۴	۷	پایین
رقابتی و بازار	۴	۳	۷	پایین

فیلتر کردن عوامل

از بین عوامل دو پیشران نیروی انسانی و مهارتی و قانونی و مقرراتی نیز به همین شکل دو عامل را انتخاب می کنیم. بایستی توجه داشت که دو عدم قطعیت نباید هم بستگی بالا داشته باشند. اگر هم بسته باشند، ماتریس ۴ سناریوی متمایز نمی سازد.

جدول ۳: فیلتر کردن عوامل

جمع	عدم قطعیت (۵-۱)	تأثیر (۵-۱)	سطح سواد دیجیتال و هوش مصنوعی کارکنان
۸	۳	۴	نیاز به مهارت‌های جدید (تحلیل داده، کار با ابزارهای AI)
۱۰	۵	۵	سرعت و عمق پذیرش هوش مصنوعی در سازمان
۸	۴	۴	تغییر نقش‌های شغلی و حذف یا ایجاد مشاغل جدید
۷	۳	۴	برنامه‌های آموزش و بازآموزی نیروی کار
۸	۴	۴	قوانین حریم خصوصی و حفاظت از داده‌ها
۸	۴	۴	مقررات اخلاقی استفاده از هوش مصنوعی
۱۰	۵	۵	میزان مقررات‌گذاری و محدودیت‌های قانونی
۷	۴	۳	مسئولیت حقوقی تصمیمات گرفته‌شده توسط هوش مصنوعی
۷	۳	۴	استانداردهای بین‌المللی و سازوکارهای نظارتی

در این مرحله نیز با کمک خبرگان به هر کدام از عوامل از ۱ تا ۵ امتیاز می‌دهیم. لذا دو عامل سرعت و عمق پذیرش هوش مصنوعی در سازمان و میزان مقررات‌گذاری و محدودیت‌های قانونی برای ماتریس سناریو انتخاب می‌شوند. همانطور که مشاهده می‌شود این دو عامل هم بستگی ندارند.

ایجاد ۴ سناریو

جدول ۴: ماتریس سناریو

زیاد	کم	
مقررات زیاد + پذیرش سریع	مقررات کم + پذیرش سریع	سریع
مقررات زیاد + پذیرش کند	مقررات کم + پذیرش کند	کند

سناریوی ۱: قفس طلایی (مقررات زیاد + پذیرش کند)

داستان سناریو:

دولت‌ها و نهادهای نظارتی پس از چندین رسوایی مربوط به سوگیری الگوریتم‌ها و نقض حریم خصوصی، قوانین بسیار سخت‌گیرانه‌ای برای استفاده از هوش مصنوعی تصویب کرده‌اند. هر پروژه هوش مصنوعی نیاز به مجوزهای چندلایه، ممیزی‌های اخلاقی و تأییدیه‌های طولانی مدت دارد.

وضعیت سازمان:

- فرآیند تصویب پروژه‌های هوش مصنوعی ۱۲-۱۸ ماه طول می‌کشد
- تنها بخش‌های کم‌ریسک (مثل پشتیبانی مشتری ساده) مجاز به استفاده از AI هستند
- تیم‌های حقوقی و انطباق قدرتمندترین واحدهای سازمان شده‌اند
- نوآوران و استارت‌آپ‌های داخلی به دلیل موانع قانونی ناامید و مهاجرت می‌کنند

نیروی انسانی:

- کارکنان به دلیل امنیت شغلی بالا، مقاومت پنهانی در برابر تغییر دارند
- مهارت‌های سنتی همچنان غالب است
- آموزش‌های هوش مصنوعی محدود و تئوری باقی مانده

فرصت‌ها:

- ریسک حقوقی و جریمه‌ها بسیار پایین است
- اعتماد مشتریان به دلیل نظارت شدید بالا است
- ثبات و پیش‌بینی‌پذیری محیط کسب‌وکار

تهدیدها:

- عقب ماندگی از رقبای بین المللی در مناطق با مقررات سبک تر
- از دست دادن فرصت های نوآوری
- هزینه های بالای انطباق قانونی

راهبرد پیشنهادی:

- سرمایه گذاری روی سیستم های مدیریت انطباق هوشمند، ایجاد واحد «اخلاق هوش مصنوعی»، تمرکز بر نوآوری در چارچوب قوانین
- سناریوی ۲: همکاری هوشمند (مقررات زیاد + پذیرش سریع)

داستان سناریو:

- دولت ها چارچوب های قانونی شفاف و استانداردهای مشخصی برای هوش مصنوعی ایجاد کرده اند، اما هم زمان سرمایه گذاری سنگینی در آموزش و زیرساخت انجام داده اند.
- سازمان ها یاد گرفته اند در چارچوب قوانین، به سرعت نوآوری کنند.

وضعیت سازمان:

- واحد «حاکمیت هوش مصنوعی» به عنوان تسهیل گر عمل می کند، نه مانع
- پروژه های هوش مصنوعی با رعایت استانداردهای اخلاقی و حریم خصوصی اجرا می شوند
- همکاری نزدیک بین تیم های فنی، حقوقی و کسب و کار
- شفافیت در تصمیمات هوش مصنوعی برای مشتریان و ذینفعان

نیروی انسانی:

- برنامه های گسترده بازآموزی و ارتقای مهارت
- کارکنان هوش مصنوعی را به عنوان «همکار» می بینند، نه رقیب
- نقش های جدیدی مثل «ناظر اخلاق» AI و «مهندس شفافیت» ایجاد شده

فرصت ها:

- اعتماد بالا از سوی مشتریان و شرکا
- دسترسی به بازارهای بین المللی با استانداردهای مشابه
- نوآوری پایدار و کم ریسک
- برند سازمان به عنوان پیشرو در «هوش مصنوعی مسئولانه»

تهدیدها:

- هزینه‌های بالای انطباق و آموزش
- سرعت کمتر نسبت به بازیگران مناطق با مقررات سبک
- پیچیدگی مدیریت هم‌زمان نوآوری و انطباق

راهبرد پیشنهادی:

ایجاد اکوسیستم همکاری با دانشگاه‌ها و نهادهای نظارتی، سرمایه‌گذاری روی آموزش مستمر، توسعه ابزارهای شفافیت و توضیح‌پذیری AI سناریوی ۳: آشوب دیجیتال (مقررات کم + پذیرش کند)

داستان سناریو:

دولت‌ها سیاست عدم مداخله در فناوری را پیش گرفته‌اند و هیچ چارچوب قانونی مشخصی وجود ندارد. اما سازمان‌ها به دلیل فرهنگ محافظه‌کار، کمبود مهارت و ابهام در مزایا، در پذیرش هوش مصنوعی کند عمل می‌کنند. بازار پر از ابزارهای ناسازگار و استانداردهای متضاد است.

وضعیت سازمان:

- هر واحد به صورت جزیره‌ای و بدون هماهنگی ابزارهای مختلف هوش مصنوعی خریداری می‌کند
- عدم یکپارچگی سیستم‌ها و داده‌ها
- پروژه‌های شکست خورده متعدد به دلیل انتخاب نادرست فناوری
- سردرگمی مدیران در تصمیم‌گیری استراتژیک

نیروی انسانی:

- ترس از جایگزینی شغلی بدون برنامه بازآموزی
- مقاومت فعال و غیرفعال کارکنان
- شکاف عمیق بین نسل دیجیتال و نسل سنتی
- خروج نیروهای متخصص به سازمان‌های پیشروتر

فرصت‌ها:

- آزادی عمل در انتخاب فناوری بدون محدودیت قانونی
- امکان آزمایش سریع ایده‌های جدید

-هزینه‌های پایین ورود به فناوری

تهدیدها:

- ریسک بالای امنیتی و نقض داده‌ها
- از دست دادن رقابت‌پذیری در برابر سازمان‌های منسجم‌تر
- سردرگمی استراتژیک و هدررفت منابع
- آسیب به برند به دلیل خطاهای هوش مصنوعی بدون نظارت

راهبرد پیشنهادی:

ایجاد «دفتر تحول دیجیتال» برای هماهنگی داخلی، تدوین چارچوب اخلاقی داخلی حتی بدون الزام قانونی، سرمایه‌گذاری فوری روی آموزش و تغییر فرهنگ

سناریوی ۴: بهشت نوآوری (مقررات کم + پذیرش سریع)

داستان سناریو:

دولت‌ها سیاست حمایتی از فناوری اتخاذ کرده‌اند و مشوق‌های مالیاتی و تسهیلات فراوانی ارائه می‌دهند. سازمان‌ها با سرعت بالا هوش مصنوعی را در تمام فرآیندها ادغام کرده‌اند. اکوسیستم استارت‌آپی شکوفا شده و نوآوری‌های تحول‌آفرین هر هفته ظهور می‌کنند.

وضعیت سازمان

- هوش مصنوعی در تمام لایه‌های سازمان نفوذ کرده است
- تصمیم‌گیری‌های کلیدی به الگوریتم‌ها سپرده شده
- ساختار سازمانی تخت و چابک با تیم‌های خودگردان
- همکاری گسترده با استارت‌آپ‌ها و مراکز تحقیق و توسعه

نیروی انسانی

- مهارت‌های هوش مصنوعی به شرط اولیه استخدام تبدیل شده
- آموزش مستمر و یادگیری مادام‌العمر فرهنگ سازمانی است
- نقش انسان از «اجراکننده» به «ناظر و خلاق» تغییر کرده
- رقابت شدید برای جذب و حفظ استعدادهای دیجیتال

فرصت‌ها

- رشد سریع درآمد و سهم بازار
- پیشگامی در نوآوری‌های صنعت

- جذب سرمایه و استعدادهای برتر
- ایجاد مدل‌های کسب و کار کاملاً جدید

تهدیدها

- ریسک بالای اخلاقی و احتمال رسوایی‌های آینده
- احتمال تغییر ناگهانی قوانین و شوک تنظیم مقرراتی
- وابستگی شدید به فناوری و آسیب‌پذیری در برابر اختلالات
- فرسودگی شغلی کارکنان به دلیل سرعت بالای تغییر

راهبرد پیشنهادی

ایجاد ذخایر ریسک برای شوک‌های احتمالی، تنوع‌بخشی به پورتفوی فناوری، حفظ قابلیت‌های انسانی به عنوان پشتیبان، نظارت اخلاقی داوطلبانه

نتیجه گیری

نتیجه گیری سناریوی اول: قفس طلایی

سناریوی «قفس طلایی» تصویری از آینده‌ای را ترسیم می‌کند که در آن امنیت و ثبات قانونی به بهای کندشدن نوآوری و از دست دادن فرصت‌های رقابتی تمام شده است. در این آینده، سازمان‌ها در محیطی فعالیت می‌کنند که اگرچه از نظر حقوقی و اخلاقی کاملاً ایمن است و ریسک جریمه‌ها و رسوایی‌های ناشی از هوش مصنوعی به حداقل رسیده، اما زنجیره‌ای از مجوزهای چندلایه، ممیزی‌های طولانی‌مدت و فرآیندهای تأیید پیچیده، سرعت عمل و چابکی لازم برای بهره‌برداری از فرصت‌های فناوری را از آن‌ها گرفته است. این وضعیت شبیه به قفسی از جنس طلا است؛ ظاهری ارزشمند و امن دارد، اما در نهایت محدودکننده آزادی عمل و پرواز به سوی افق‌های جدید است. سازمان‌هایی که در این سناریو گرفتار می‌شوند، ممکن است در کوتاه‌مدت از بحران‌های قانونی در امان بمانند، اما در بلندمدت با خطر حاشیه‌ای شدن در برابر رقبای بین‌المللی که در محیط‌های با مقررات سبک‌تر فعالیت می‌کنند، مواجه خواهند شد و به تدریج توانایی جذب استعدادهای نوآور و حفظ موقعیت رقابتی خود را از دست می‌دهند.

با این حال، این سناریو لزوماً یک آینده منفی مطلق نیست و می‌تواند برای سازمان‌هایی که استراتژی «پایداری به جای رشد سریع» را انتخاب می‌کنند، فرصت‌هایی نیز ایجاد کند. کلید موفقیت در این سناریو، تبدیل الزامات قانونی از مانع به مزیت رقابتی است؛ به این معنا که سازمان‌ها می‌توانند با سرمایه‌گذاری روی سیستم‌های مدیریت انطباق هوشمند،

ایجاد واحد تخصصی اخلاق هوش مصنوعی و برندسازی به عنوان «امن‌ترین ارائه‌دهنده خدمات مبتنی بر AI»، اعتماد مشتریان حساس به حریم خصوصی و نهادهای دولتی را جلب کنند. همچنین این محیط می‌تواند بستری برای نوآوری در چارچوب مشخص باشد، جایی که تیم‌ها یاد می‌گیرند خلاقیت را نه در شکستن قوانین، بلکه در یافتن راه‌حل‌های هوشمندانه درون مرزهای تعریف‌شده جستجو کنند. پیام نهایی این سناریو برای مدیران این است که اگر حرکت به سمت این آینده اجتناب‌ناپذیر باشد، باید به جای مقاومت بی‌ثمر در برابر مقررات، روی ساختن قابلیت‌های انطباقی سرمایه‌گذاری کنند که هم هزینه‌های قانونی را کاهش دهد و هم به عنوان تمایزدهنده برند در بازار عمل کند؛ زیرا در دنیایی که اعتماد به کالای کمیاب تبدیل می‌شود، «امنیت اثبات‌شده» می‌تواند خود به یک مزیت رقابتی ارزشمند تبدیل شود.

نتیجه‌گیری سناریوی دوم: همکاری هوشمند

سناریوی «همکاری هوشمند» متعادل‌ترین و پایدارترین تصویر از آینده سازمان در پرتو هوش مصنوعی را ارائه می‌دهد؛ آینده‌ای که در آن مقررات و نوآوری نه به عنوان دو نیروی متضاد، بلکه به عنوان مکمل‌های یکدیگر عمل می‌کنند. در این سناریو، دولت‌ها و نهادهای نظارتی به جای ایجاد موانع بوروکراتیک، چارچوب‌های شفاف و استانداردهای مشخصی را تدوین کرده‌اند که هم از حقوق شهروندان و کارکنان محافظت می‌کند و هم فضای کافی برای آزمایش و رشد فناوری را فراهم می‌سازد. سازمان‌هایی که در این محیط فعالیت می‌کنند، یاد گرفته‌اند که انطباق قانونی را نه به عنوان یک هزینه سربار، بلکه به عنوان بخشی از استراتژی رقابتی خود بپذیرند و با ایجاد واحدهای حاکمیت هوش مصنوعی، پلی مستحکم بین تیم‌های فنی، حقوقی و کسب‌وکار ساخته‌اند. این همکاری درون‌سازمانی و برون‌سازمانی باعث می‌شود که نوآوری با سرعت معقول اما پیوسته پیش برود، بدون آنکه قربانی عجله‌های بی‌پروا یا ترس‌های بی‌مورد شود. در این آینده، اعتماد به یک دار استراتژیک تبدیل شده و سازمان‌هایی که بتوانند شفافیت، توضیح‌پذیری و مسئولیت‌پذیری را در سیستم‌های هوش مصنوعی خود تضمین کنند، نه تنها در بازار داخلی، بلکه در عرصه بین‌المللی نیز به عنوان الگوهای موفق شناخته می‌شوند. با این حال، دستیابی به این سناریوی ایده‌آل نیازمند سرمایه‌گذاری جدی، صبر استراتژیک و تعهد بلندمدت از سوی مدیران ارشد است. چالش اصلی در این مسیر، حفظ تعادل ظریف بین سرعت نوآوری و دقت انطباق است؛ زیرا وسوسه کوتاه‌کردن مسیرهای قانونی برای کسب

مزیت رقابتی سریع همیشه وجود دارد و مقاومت در برابر این وسوسه نیازمند بلوغ سازمانی و فرهنگی است. همچنین هزینه‌های بالای آموزش مستمر، توسعه ابزارهای شفافیت و نگهداری سیستم‌های حاکمیت هوش مصنوعی می‌تواند برای سازمان‌های کوچک و متوسط فشار مالی ایجاد کند و شکافی بین بازیگران بزرگ و کوچک صنعت به وجود آورد. پیام نهایی این سناریو برای مدیران این است که «نوآوری مسئولانه» اگرچه در کوتاه‌مدت پرهزینه و کندتر به نظر می‌رسد، اما در بلندمدت پایدارترین مسیر برای رشد است؛ زیرا سازمان‌هایی که امروز روی اعتماد، شفافیت و اخلاق سرمایه‌گذاری می‌کنند، فردا در معرض کمترین ریسک‌های قانونی قرار می‌گیرند و می‌توانند در هر تغییری در محیط قانونی، با کمترین شوک به مسیر خود ادامه دهند. این سناریو نشان می‌دهد که آینده مطلوب سازمان‌ها نه در فرار از مقررات، بلکه در آغوش کشیدن هوشمندانه آن‌ها و تبدیل الزامات بیرونی به قابلیت‌های درونی نهفته است.

نتیجه‌گیری سناریوی چهارم: بهشت نوآوری

سناریوی «بهشت نوآوری» جذاب‌ترین و در عین حال فریبنده‌ترین تصویر از آینده سازمان در پرتو هوش مصنوعی را به نمایش می‌گذارد؛ آینده‌ای که در آن موانع قانونی برداشته شده، مشوق‌های دولتی سرازیر شده‌اند و سازمان‌ها با سرعتی باورنکردنی هوش مصنوعی را در تار و پود فرآیندهای خود تنیده‌اند. در این محیط، استارت‌آپ‌ها مانند قارچ پس از باران می‌رویند، مدل‌های کسب‌وکار جدید هر هفته متولد می‌شوند، و سازمان‌های پیشرو با اتوماسیون گسترده و تصمیم‌گیری الگوریتمی، بهره‌وری و سودآوری بی‌سابقه‌ای را تجربه می‌کنند. ساختارهای سلسله‌مراتبی قدیمی فرو ریخته و جای خود را به تیم‌های چابک، خودگردان و شبکه‌ای داده‌اند که انسان و ماشین در هم‌زیستی کامل با یکدیگر کار می‌کنند. این سناریو شبیه به رویای هر کارآفرین و مدیر نوآوری است؛ جایی که ایده‌ها بدون اصطکاک بوروکراتیک به محصول تبدیل می‌شوند، سرمایه به وفور در دسترس است، و استعداد‌های برتر جهان به سوی سازمان‌هایی سرازیر می‌شوند که جسارت پرواز بدون ترمز را دارند. رشد سریع، سهم بازار فزاینده و برند به عنوان پیشرو صنعت، پاداش‌های وسوسه‌انگیز این مسیر پرسرعت هستند. با این حال، پشت این ظاهر درخشان، خطرات پنهانی نهفته است که می‌تواند در یک لحظه این بهشت را به کابوس تبدیل کند. شتاب بدون ترمز به معنای آسیب‌پذیری شدید در برابر شوک‌های ناگهانی است؛ یک رسوایی اخلاقی، یک تغییر ناگهانی در قوانین دولتی، یک حمله سایبری گسترده، یا یک

خطای الگوریتمی فاجعه‌بار می‌تواند تمام دستاوردهای سال‌ها را در چند روز نابود کند. همچنین وابستگی شدید به هوش مصنوعی، سازمان را در برابر اختلالات فنی، قطعی زیرساخت‌ها و از دست رفتن مهارت‌های انسانی آسیب‌پذیر می‌سازد. فرسودگی شغلی کارکنان به دلیل سرعت سرسام‌آور تغییر، رقابت بی‌رحمانه برای جذب استعدادها، و نادیده گرفتن ملاحظات اخلاقی به امید «بعداً درستش می‌کنیم» از دیگر هزینه‌های پنهان این سناریو است. پیام نهایی این سناریو برای مدیران این است که «شتاب بدون جهت‌یابی اخلاقی و ذخایر تاب‌آوری» مانند رانندگی با سرعت بالا در جاده‌ای مه‌آلود است؛ ممکن است برای مدتی پیشرو باشید، اما احتمال برخورد با مانع غیرمنتظره بسیار بالاست. راهبرد خردمندانه در این سناریو، حفظ سرعت نوآوری همراه با ایجاد «ترمزهای هوشمند» داخلی است؛ یعنی نظارت اخلاقی داوطلبانه، تنوع‌بخشی به پورتفوی فناوری، حفظ قابلیت‌های انسانی به عنوان پشتیبان، و ساختن ذخایر مالی و عملیاتی برای روز مبادا. این سناریو به مدیران یادآوری می‌کند که در دنیای پرسرعت فناوری، برنده نهایی کسی نیست که سریع‌تر شروع می‌کند، بلکه کسی است که می‌تواند سرعت را با پایداری متوازن کند و در طوفان تغییرات، کشتی سازمان را بدون غرق‌شدن به ساحل موفقیت برساند.

نتیجه‌گیری سناریوی سوم: آشوب دیجیتال

سناریوی «آشوب دیجیتال» نمادی از پارادوکس تلخ آزادی بدون آمادگی است؛ وضعیتی که در آن نبود محدودیت‌های قانونی به جای تبدیل‌شدن به فرصتی طلایی برای نوآوری، به میدان نبردی از تصمیم‌های پراکنده، سرمایه‌گذاری‌های ناهماهنگ و اتلاف انرژی سازمانی تبدیل شده است. در این آینده، سازمان‌ها اگرچه از سد مجوزهای دولتی و هزینه‌های انطباق عبور کرده‌اند، اما در دام پیچیده‌تری گرفتار آمده‌اند: فقدان بینش استراتژیک، نبود مهارت‌های لازم برای بهره‌برداری از فناوری، و فرهنگی که تغییر را تهدید می‌بیند نه فرصت. هر واحد سازمانی به صورت جزیره‌ای عمل می‌کند، ابزارهای ناسازگار خریداری می‌شوند، داده‌ها در سیلوهای جداگانه حبس می‌شوند، و هیچ دید جامعی از اینکه هوش مصنوعی چگونه باید در خدمت اهداف کلان سازمان قرار گیرد وجود ندارد. این وضعیت شبیه به ارکستری است که هر نوازنده ساز خود را می‌زند، اما هیچ رهبری وجود ندارد که آن‌ها را هماهنگ کند؛ نتیجه نه صدای موسیقی، که هیاهویی گوش‌خراش است که نه تنها شنیدنی نیست، بلکه به مرور زمان شنوندگان را فراری می‌دهد.

راه نجات از این سناریو، نه در انتظار برای ظهور قوانین دولتی و نه در ادامه وضعیت موجود، بلکه در ایجاد «نظم خودخواسته» از درون سازمان نهفته است. مدیرانی که سازمان‌شان در معرض این سناریو قرار دارد، باید بپذیرند که نبود ناظر بیرونی به معنای نبود نیاز به نظارت نیست؛ بلکه برعکس، نیاز به انضباط درونی را دوچندان می‌کند. تدوین چارچوب اخلاقی داخلی، ایجاد دفتر هماهنگی تحول دیجیتال با اختیارات فرابخشی، استانداردسازی ابزارها و پلتفرم‌ها، و از همه مهم‌تر سرمایه‌گذاری گسترده روی آموزش و تغییر نگرش کارکنان از اقدامات حیاتی این مسیر است. پیام نهایی این سناریو هشدار جدی برای مدیران است: «آزادی عمل بدون بلوغ سازمانی» فرمولی برای شکست است؛ زیرا در غیاب چارچوب‌های بیرونی، سازمان باید خود را مهار کند و اگر نتواند این کار را انجام دهد، بازار و رقبا این کار را به جای او خواهند کرد. سازمان‌هایی که از این آشوب جان سالم به در می‌برند، آن‌هایی هستند که می‌فهمند نظم واقعی از درون می‌آید، نه از بیرون، و می‌توانند حتی در محیطی بدون قانون، قانون‌مند و استراتژیک عمل کنند.

منابع:

- Ajiboyev, K. J. (2024). The impact of responsible AI practices on organisational digital transformation. *World Journal of Advanced Research and Reviews*, 23(2), 2866–2879. <https://doi.org/10.30574/wjarr.2024.23.2.2433>
- Alhashmi, S. F. S., et al. (2024). Assessing AI adoption in developing country academia: A trust and privacy-augmented UTAUT framework. *Heliyon*, 10(18), e37569. <https://doi.org/10.1016/j.heliyon.2024.e37569>
- American Psychological Association. (2024, July). APA publishing policies. <https://www.apa.org/pubs/journals/resources/publishing-policies>
- Bazine, I. (2025). Exploring the development of the technology acceptance model (TAM): A chronological overview. *International Journal of Research in Social Sciences*, 15(7), 1643–1655.
- Bengston, D. N. (2019). Futures research methods and applications in natural resources. *Society & Natural Resources*, 32(10), 1093–1111. https://www.fs.usda.gov/nrs/pubs/jrnl/2019/nrs_2019_bengston_001.pdf
- Boston Consulting Group. (2024, October 24). AI adoption in 2024: 74% of companies struggle to achieve and scale value. <https://www.bcg.com/press/24october2024-ai-adoption-in-2024-74-of-companies-struggle-to-achieve-and-scale-value>
- Boston Consulting Group. (2025). AI at work 2025: Momentum builds, but gaps remain. <https://www.bcg.com/publications/2025/ai-at-work-momentum-builds-but-gaps-remain>
- Brynjolfsson, E., & McAfee, A. (2014). *The second machine age: Work, progress, and prosperity in a time of brilliant technologies*. W. W. Norton & Company.
- Chen, G., Fan, J., & Azam, M. (2024). Exploring artificial intelligence (AI) chatbots adoption among research scholars using unified theory of acceptance and use of technology (UTAUT). *The International Journal of Information and Learning Technology*, 41(4). <https://doi.org/10.1108/IJILT-05-2024-0090>
- Davenport, T. H., & Ronanki, R. (2018). Artificial intelligence for the real world. *Harvard Business Review*, 96(1), 108–116.
- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319–340.
- Han, M., Mustafa, H., & Kharuddin, S. (2025). AI adoption in accounting education: A UTAUT-based analysis of mediating and moderating mechanisms. *Journal of Pedagogical Research*, 9(2), 191–205. <https://doi.org/10.33902/JPR.202532831>
- Heliyon. (2024). Acceptance of artificial intelligence in university contexts: A conceptual analysis based on UTAUT2 theory. *Heliyon*, 10(19), e38315. <https://doi.org/10.1016/j.heliyon.2024.e38315>
- Hemming, V., Burgman, M. A., Hanea, A. M., McBride, M. F., Wintle, B. A., & others. (2024). Good practice in structured expert elicitation: Learning from guidelines. National Center for Biotechnology Information. <https://www.ncbi.nlm.nih.gov/books/NBK571059/>

- Insight. (2024). The path to digital transformation: Where IT leaders stand in 2024. https://www.insight.com/content/dam/insight-web/en_US/pdfs/insight/the-path-to-digital-transformation--where-it-leaders-stand-in-2024.pdf
- IZT – Institute for Futures Studies and Technology Assessment. (2024). Futures studies and future-oriented technology analysis: Principles, methodology and research questions. <https://www.izt.de/en/publications/futures-studies-and-future-oriented-technology-analysis-principles-methodology-and-research-questions/>
- Luesink, H., & Boin, A. (2024). Scenario planning to enable foresight in crisis management. *International Journal of Information Systems*. <https://jeroenwolbers.com/wp-content/uploads/2024/04/luesink-et-al-2024-scenario-planning-to-enable-foresight-in-crisis-management.pdf>
- McKinsey & Company. (2024). The state of AI in early 2024. <https://www.mckinsey.com/capabilities/quantumblack/our-insights/the-state-of-ai-2024>
- McKinsey & Company. (2025). Superagency in the workplace: Empowering people to unlock AI's full potential at work. <https://www.mckinsey.com/capabilities/tech-and-ai/our-insights/superagency-in-the-workplace-empowering-people-to-unlock-ais-full-potential-at-work>
- McKinsey & Company. (2025). The state of AI: Global survey 2025. <https://www.mckinsey.com/capabilities/quantumblack/our-insights/the-state-of-ai>
- MIT Sloan Management Review, & Boston Consulting Group. (2024, November 12). Organizations combining organizational learning and AI-specific learning manage uncertainty better. <https://www.bcg.com/press/12november2024-organizational-learning-and-ai-specific-learning-managing-uncertainty>
- Mitsui & Co. Global Strategic Studies Institute. (2024, November). Scenario planning for futures: Approaches to expanding the boundaries of ideas. https://www.mitsui.com/mgssi/en/report/detail/___icsFiles/afieldfile/2025/01/20/2411_p_fujii_e.pdf
- Mittelstadt, B. D., Allo, P., Taddeo, M., Wachter, S., & Floridi, L. (2016). The ethics of algorithms: Mapping the debate. *Big Data & Society*, 3(2), 2053951716679679. <https://doi.org/10.1177/2053951716679679>
- National Academies / NCBI. (2024). Good practice in structured expert elicitation: Learning from guidelines. <https://www.ncbi.nlm.nih.gov/books/NBK571059/>
- O'Neil, C. (2016). *Weapons of math destruction: How big data increases inequality and threatens democracy*. Crown.
- Organisation for Economic Co-operation and Development. (2024). Fostering an inclusive digital transformation as AI spreads among firms. https://www.oecd.org/content/dam/oecd/en/publications/reports/2024/1/1/fostering-an-inclusive-digital-transformation-as-ai-spreads-among-firms_cd50d324/5876200c-en.pdf
- Ralston, B., Wilson, I., & Dearden, T. (2024). Reflections on the usefulness of scenarios. *Journal of Futures Studies*, 29(2).

- <https://jfsdigital.org/articles-and-essays/2024-2/vol-29-no-2-december-2024/reflections-on-the-usefulness-of-scenarios/>
- Stanford Institute for Human-Centered Artificial Intelligence. (2024). AI Index report 2024. <https://hai.stanford.edu/ai-index/2024-ai-index-report>
- Venkatesh, V., & Davis, F. D. (2000). A theoretical extension of the technology acceptance model: Four longitudinal field studies. *Management Science*, 46(2), 186–204.
- Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User acceptance of information technology: Toward a unified view. *MIS Quarterly*, 27(3), 425–478.
- Wharton Human-AI Research, & GBK Collective. (2025). Accountable acceleration: Gen AI fast-tracks into the enterprise. <https://knowledge.wharton.upenn.edu/article/how-are-companies-using-gen-ai-in-2025/>
- Wright, G., & Cairns, G. (2024). Scenario planning and futures thinking: New developments in the field. *Futures*. I can help locate exact bibliographic details if needed.