

طراحی چارچوب حسابرسی توضیح‌پذیر برای ارزیابی ریسک صورت‌های مالی مبتنی بر هوش مصنوعی قابل تبیین

نوع مقاله: پژوهشی

رحیم ملائی^۱

علی جعفری^۲

عسگر پاکمرام^۳

نادر رضائی^۴

تاریخ پذیرش: ۱۴۰۵/۳/۴

تاریخ دریافت: ۱۴۰۴/۱۱/۱۰

چکیده

هدف این پژوهش ارائه چارچوب حسابرسی توضیح‌پذیر برای ارزیابی ریسک صورت‌های مالی مبتنی بر هوش مصنوعی قابل تبیین است. این پژوهش از لحاظ هدف، اکتشافی و از لحاظ نوع استدلال، استقرایی و از نظر ماهیت داده‌ها، پژوهش کیفی است. جامعه آماری این پژوهش شامل ۳۹ مقاله نهایی می‌باشد که از طریق جستجوی نظام‌مند در پایگاه‌های علمی (تمرکز بر منابع خارجی از سال ۲۰۲۰ به بعد و منابع داخلی از سال ۱۴۰۰ به بعد) استخراج شدند. تحلیل داده‌ها با استفاده از تکنیک تحلیل محتوای اسنادی انجام شد و برای اطمینان از قابلیت اعتماد تحلیل‌ها، کیفیت آن با محاسبه ضریب کاپای کوهن (۰/۷۴۲) تأیید گردید. یافته‌های محوری پژوهش به چهار کد فراگیر دسته‌بندی شدند: مفهوم حسابرسی توضیح‌پذیر مبتنی بر هوش مصنوعی قابل تبیین، پیشایندها، پیامدها و عوامل تعدیل‌گر. نتایج نشان داد که هوش مصنوعی قابل تبیین با تقویت ارتباط تبیین‌ها با ریسک تحریف با اهمیت و بهبود قضاوت انسانی، کیفیت حسابرسی و پاسخگویی حرفه‌ای را ارتقا می‌دهد. دستیابی به این پیامدهای مثبت، مستلزم فراهم بودن پیشایندهایی نظیر آمادگی فنی و زیرساختی، استانداردهای حرفه‌ای و مهارت‌های ترکیبی حساب‌رسان است. همچنین، پیامدهای منفی بالقوه‌ای مانند چالش‌های فنی و شناختی، حرفه‌ای و اخلاقی شناسایی شدند. در نهایت، اثرگذاری حسابرسی توضیح‌پذیر تابعی از عوامل تعدیل‌گر همچون نگرش مدیریت، فشارهای محیطی و بلوغ اکوسیستم فناوری می‌باشد.

واژه‌های کلیدی: حسابرسی توضیح‌پذیر، ریسک صورت‌های مالی، هوش مصنوعی قابل تبیین.

طبقه‌بندی JEL: M41, M42, O33, C88

۱. گروه حسابداری، واحد بناب، دانشگاه آزاد اسلامی، بناب، ایران
Rahim.mollaei@iau.ac.ir

۲. گروه حسابداری، واحد بناب، دانشگاه آزاد اسلامی، بناب، ایران (نویسنده مسئول)
alijafari1355@iau.ac.ir

۳. گروه حسابداری، واحد بناب، دانشگاه آزاد اسلامی، بناب، ایران
pakmaram@iau.ac.ir

۴. گروه حسابداری، واحد بناب، دانشگاه آزاد اسلامی، بناب، ایران
naderrezaei@iau.ac.ir

۱. گروه حسابداری، واحد بناب، دانشگاه آزاد اسلامی، بناب، ایران

۲. گروه حسابداری، واحد بناب، دانشگاه آزاد اسلامی، بناب، ایران

۳. گروه حسابداری، واحد بناب، دانشگاه آزاد اسلامی، بناب، ایران

۴. گروه حسابداری، واحد بناب، دانشگاه آزاد اسلامی، بناب، ایران

مقدمه

حسابرسي به عنوان سازوكاري حياتي، نه تنها زمينه دسترسي سرمايه گذاران به اطلاعات معتبر را فراهم مي آورد، بلکه با ارتقاي شفافيت، از منافع كلي بازار و اقتصاد حمايت مي کند (آذر سعيد و رستمي، ۱۴۰۲). پژوهشگران حوزه تصميم گيري و قضاوت در حسابرسي دريافته اند که عليرغم اينکه محتوای اطلاعات حسابداری از معاملات اقتصادی نشأت می گیرد، لیکن آماده سازی و البته تفسیر این اطلاعات به شدت تحت تاثیر ویژگیهای پردازش اطلاعات انسانی حسابداران و دیگر کاربران اطلاعات حسابداری مانند مدیران، سرمايه گذاران و ... قرار گرفته است (صحت بر مچ و همکاران، ۱۴۰۴). در سالهاي اخير، حرفه حسابرسي تحت تأثير عوامل و رخدادهاي گوناگوني قرار گرفته است (منصوري، کرمي و يزداني، ۱۴۰۳). این حرفه، به طور سنتي، به شدت به فرايندهاي دستي و مهارت هاي انساني متكي بوده است. با رشد فناوري اطلاعات، تبديلي بنيادين در روشهاي حسابرسي ايجاد شده است (الحاتمي^۱، ۲۰۲۲). ابزارهای تحلیل داده ها، يادگيري ماشين و سيستم هاي خبره براي شناسايي تقلب، افزايش دقت و بهبود کارايي حسابرسي مورد استفاده قرار گرفته اند (مينکینن و همکاران^۲، ۲۰۲۲). به طور كلي، حسابرسي در بستر جامعه و سازمان ها قرار مي گيرد، لذا نمي توان آن را از تغييرات و هرآنچه در آن است، جدا در نظر گرفت. حسابرسان براي باقي ماندن به عنوان يك خدمت ارزشمند و مرتبط براي سرمايه گذاران، اعتباردهندگان و ساير ذينفعان و ارائه خدمات با كيفيت، نمي توانند به تبديلاتي كه در محيط حسابرسي رخ مي دهد، رويكرد منفعل داشته باشند و بايد با شناسايي روندهاي آتي و واكنش مناسب به آنها، خود را براي آينده آماده سازند (فدراسيون حسابداران اروپا^۳، ۲۰۱۶) و پيوسته در حال تكامل باشند (صديقي و همكاران، ۱۴۰۲). در صورت عدم واكنش به تغييرات، ممكن است مناسب و مفيد بودن نقش حسابرسي براي ذينفعان از بين برود و منابع اطلاعاتي جديدي جايگزين خدمات حسابرسي شود (گريستون^۴، ۲۰۲۲)؛ هرچند كه نمي توان جهان را بدون حسابرسي تصور كرد.

در سالهای اخیر هوش مصنوعی^۵ و یادگیری ماشين^۶ با توجه به کاربردهای بالقوه شان در حسابرسي، توجه فزاينده ای را به خود جلب کرده اند. يکي از چالش های اصلي پذيرش آنها در حسابرسي، عدم قابليت توضيح نتايج آنهاست. با بلوغ هوش مصنوعي و يادگيري ماشين، تکنیک‌هایی که می‌توانند قابليت تفسير هوش مصنوعي يا به اختصار هوش مصنوعي قابل

1 Al-Hattami

2 Minkinen , Laine & Mäntymäki.

3 Federation of European Accountants (FEE)

4 Grayston

5 Artificial Intelligence

6 Machine Learning

تبیین^۱ را افزایش دهند، تکامل می‌یابند (ژانگ و همکاران^۲، ۲۰۲۲). هوش مصنوعی به توانایی یک رایانه یا ربات برای انجام وظایف اشاره دارد که در بیشتر موارد به هوش انسانی نیاز دارند (زمانکوا^۳، ۲۰۱۹). یکی از مزایای استفاده از هوش مصنوعی، کاهش خطای انسانی در ورود خودکار داده‌ها و تشخیص داده‌های ساختگی و تقلبی است که از نیاز به مداخله انسانی در برخی فرایندها می‌کاهد. این فناوری همچنین با خودکارسازی وظایف تخصصی، دقت و سرعت ارزیابی سیستم را افزایش می‌دهد (یون^۴، ۲۰۲۰). از طرف دیگر، یادگیری ماشین یک روش محاسباتی است که می‌تواند الگوهای پنهان را از داده‌ها شناسایی کرده و پیش‌بینی یا طبقه‌بندی کند (ژانگ و همکاران^۵، ۲۰۲۲).

در طول سال‌ها، تمرکز تحقیقات و کاربرد مدل‌های هوش مصنوعی در حسابداری و حسابرسی بر بهبود عملکرد پیش‌بینی‌کننده آنها بوده است. با این حال، با افزایش عملکرد پیش‌بینی‌کننده یک مدل هوش مصنوعی، قابلیت توضیح مدل به طور کلی کاهش می‌یابد. برای مؤثرتر کردن یک مدل، اغلب به متغیرها و محاسبات چندلایه زیادی نیاز است که مدل را مانند یک «جعبه سیاه» مبهم می‌کند. عدم قابلیت توضیح مدل‌های پیچیده هوش مصنوعی یکی از چالش‌های اصلی پذیرش هوش مصنوعی در حسابداری و حسابرسی است. برای رفع نیاز جهانی به تفسیر بهتر فرایندها و خروجی‌های هوش مصنوعی، دانشمندان کامپیوتر جریانی از تحقیقات را به هوش مصنوعی قابل تبیین اختصاص داده‌اند. هوش مصنوعی قابل تبیین مجموعه‌ای از تکنیک‌هایی است که یک الگوریتم یادگیری ماشین «جعبه سیاه» را توضیح می‌دهد (میلر^۵، ۲۰۱۹). اگرچه ابزارهای نوین محاسباتی مانند یادگیری ماشینی پتانسیل بالایی برای بهبود کارایی دارند، اما ماهیت «جعبه سیاه» آنها، به‌ویژه در زمینه‌های حساس مانند قضاوت حسابرسی و ارزیابی ریسک تحریف بااهمیت، پذیرش گسترده و اتکای حرفه‌ای را محدود ساخته است. این چالش، شکاف تحقیقاتی اصلی را در حوزه حسابرسی نوین شکل می‌دهد: چگونه می‌توان از مزایای نوآورانه هوش مصنوعی در ارزیابی ریسک صورت‌های مالی به طور مؤثر و قابل اطمینان استفاده کرد، در حالی که شفافیت، توضیح‌پذیری و مسئولیت‌پذیری انسانی حفظ شود؟

هوش مصنوعی قابل تبیین به عنوان یک پاسخ نوین و متعهد به این چالش، در ادبیات علمی و صنعتی ظهور کرده است. اگرچه تحقیقات علوم کامپیوتر بسیاری از تکنیک‌های هوش مصنوعی قابل تبیین را توسعه داده است، اما پژوهش‌های محدودی در معرفی این تکنیک‌ها به متخصصان حسابداری و حسابرسی انجام شده است. بیشتر پژوهش‌های موجود در این حوزه، بر پیش‌بینی مدل‌های هوش مصنوعی تمرکز دارند، یا به مباحث کلی و تئوری مربوط

1 Explainable Artificial Intelligence (XAI)

2 Zhang et al

3 Zemankova

4 Yoon

5 Miller

به توضیح‌پذیری پرداخته‌اند، بنابراین طراحی یک چارچوب مفهومی برای ادغام هوش مصنوعی قابل تبیین در فرایند حسابرسی، به ویژه در کشور ایران، ضروری بوده و نیازمند پژوهش است. از این رو، هدف اصلی این پژوهش، طراحی یک چارچوب مفهومی حسابرسی توضیح‌پذیر مبتنی بر هوش مصنوعی قابل تبیین است تا اطمینان حاصل شود که توضیحات ارائه شده توسط مدل‌ها، ضمن افزایش اعتماد کاربر، الزامات استاندارد حسابرسی و ریسک تحریف بااهمیت را پوشش می‌دهند.

۱. مبانی نظری و پیشینه پژوهش

حسابرسی به عنوان فرایندی نظام‌مند و مستقل شناخته می‌شود که طی آن شواهد مرتبط با ادعاهای مالی جمع‌آوری و ارزیابی شده و میزان انطباق این ادعاها با معیارهای از پیش تعیین شده بررسی می‌شود. هدف اصلی این فرایند، ارائه گزارش‌هایی است که قابلیت اتکا به اطلاعات مالی را برای ذینفعان افزایش دهد (پیری و آشتاب، ۱۴۰۴). اصطلاح «هوش مصنوعی قابل تبیین» اولین بار توسط ون لنت، فیشر و مانکوسو (۲۰۰۴) برای توصیف توانایی یک سیستم آموزشی توسعه‌یافته برای ارتش ایالات متحده در توضیح تصمیمات مبتنی بر هوش مصنوعی آن ابداع شد. در سال ۲۰۱۷، سازمان پروژه‌های پژوهشی پیشرفته دفاعی^۱ ایالات متحده آمریکا برنامه هوش مصنوعی قابل تبیین خود را برای توسعه تکنیک‌هایی که می‌توانند سیستم‌های هوشمند را توضیح دهند، راه‌اندازی کرد. این سازمان، هوش مصنوعی قابل تبیین را به عنوان مجموعه‌ای از تکنیک‌هایی تعریف می‌کند که «مدل‌های قابل توضیحی تولید می‌کنند که وقتی با تکنیک‌های توضیح مؤثر ترکیب می‌شوند، کاربران نهایی را قادر می‌سازند تا نسل نوظهور سیستم‌های هوش مصنوعی را درک کنند، به طور مناسب به آنها اعتماد کنند و به طور مؤثر مدیریت کنند». این مفهوم از اواسط دهه ۱۹۷۰ وجود داشته است، زمانی که سیستم‌های خبره برای هوش مصنوعی رایج بودند (میلر، ۲۰۱۹). سیستم‌های خبره به عنوان «موج اول هوش مصنوعی» در نظر گرفته می‌شوند، که سیستم‌های مبتنی بر قانون هستند که بر اساس دانش تخصصی ساخته شده‌اند (لانچبری، ۲۰۱۷).

با پیشرفت‌های «موج دوم هوش مصنوعی»، که با روش‌های یادگیری آماری مانند یادگیری ماشین نشان داده می‌شود، مفهوم هوش مصنوعی قابل تبیین کند شد، زیرا تمرکز تحقیقات هوش مصنوعی به سمت پیاده‌سازی مدل و بهبود قدرت پیش‌بینی تغییر یافت (آدادی و برادا، ۲۰۱۸). در نتیجه، پیشرفت‌ها در یادگیری ماشین سیستم‌های خودمختاری را ایجاد کرده است که می‌توانند به طور مستقل درک کنند، یاد بگیرند، تصمیم بگیرند و عمل کنند (سازمان پروژه‌های پژوهشی پیشرفته دفاعی، ۲۰۱۶). با این حال، سال‌های اخیر شاهد تجدید حیات در بحث‌ها و تحقیقات هوش مصنوعی قابل تبیین بوده‌ایم، بیشتر به این دلیل که بسیاری از برنامه‌های مدرن

هوش مصنوعی فاقد شفافیت و قابلیت تفسیر بوده‌اند و منجر به کاربردهای محدود و نگرانی‌هایی در مورد اخلاق و اعتماد شده‌اند. هوش مصنوعی قابل تبیین به عنوان بخشی از «موج سوم هوش مصنوعی» در نظر گرفته می‌شود که هدف آن تولید الگوریتم‌هایی است که می‌توانند خود را توضیح دهند (لانچیری، ۲۰۱۷). آدادی و برادا (۲۰۱۸) نیاز به هوش مصنوعی قابل تبیین را در چهار جنبه خلاصه کردند: توضیح برای توجیه، توضیح برای کنترل، توضیح برای بهبود و توضیح برای کشف. هوش مصنوعی قابل تبیین با مفهوم «انسان در حلقه» سازگار است، زیرا هدف آن فراهم کردن درک بهتر از سیستم‌های هوش مصنوعی مبهم است تا انسان‌ها بتوانند از چنین ابزارهایی برای در کار خود بهتر استفاده کنند (ژو و همکاران، ۲۰۱۸).

هوش مصنوعی و یادگیری ماشینی با توجه به کاربردهای بالقوه‌شان در حسابرسی، توجه فزاینده‌ای را به خود جلب می‌کنند. یکی از چالش‌های اصلی پذیرش آنها در حسابرسی، عدم قابلیت توضیح نتایج آنهاست. وقتی هوش مصنوعی و یادگیری ماشینی در رویه‌های حسابرسی استفاده می‌شود، مشمول استانداردهای حسابرسی، مانند شواهد حسابرسی (استاندارد ۱۱۰۵) و مستندات حسابرسی (استاندارد ۱۲۱۵) هیئت نظارت بر حسابداری شرکت‌های سهامی عام^۱ است. استاندارد شواهد حسابرسی، حساب‌برسان را ملزم می‌کند که شواهد حسابرسی کافی و مناسب را برای ارائه مبنایی معقول برای اظهار نظر خود هنگام انجام رویه‌های حسابرسی به دست آورند. کفایت، کمیت شواهد حسابرسی را اندازه‌گیری می‌کند و مناسب بودن، کیفیت شواهد حسابرسی را نشان می‌دهد که می‌تواند به مربوط بودن و قابل اعتماد بودن تقسیم شود. وقتی هوش مصنوعی و یادگیری ماشینی در رویه‌های حسابرسی استفاده می‌شود، کفایت شواهد حسابرسی می‌تواند افزایش یابد، زیرا هوش مصنوعی و یادگیری ماشینی می‌تواند رویه‌های حسابرسی خاصی را خودکار کند، مانند خودکارسازی آزمون کل جامعه (نو و همکاران، ۲۰۱۹). وقتی هوش مصنوعی و یادگیری ماشینی با تشخیص ناهنجاری‌ها یا انجام پیش‌بینی‌ها به رویه‌های حسابرسی کمک می‌کند، حساب‌برسان باید در نظر بگیرند که کدام ادعاها بیشترین ارتباط را با خروجی‌های هوش مصنوعی و یادگیری ماشینی دارند و آیا نتایج هوش مصنوعی و یادگیری ماشینی قابل اعتماد هستند یا خیر. برای مثال، وقتی یک الگوریتم یادگیری ماشینی پیش‌بینی می‌کند که یک تراکنش جعلی است، حساب‌برس باید قبل از هرگونه اقدامی، تصمیم بگیرد که کدام ادعا تحت تاثیر قرار گرفته است (به عنوان مثال، ارزش‌گذاری یا وجود) و الگوریتم یادگیری ماشینی چقدر این پیش‌بینی را قابل اعتماد می‌کند.

مستندات حسابرسی باید با جزئیات کافی باشد تا هدف، منبع و نتایج به دست آمده به وضوح درک شوند (استاندارد ۱۲۱۵). هنگامی که هوش مصنوعی و یادگیری ماشینی به قضاوت حساب‌برسان در مورد هر گونه ادعای صورت‌های مالی کمک می‌کند، حساب‌برسان باید شواهد

حسابرسی (استاندارد ۱۱۰۵) به دست آمده از ابزارهای هوش مصنوعی و یادگیری ماشین و ماهیت (یعنی سازوکار، قابلیت اطمینان) چنین ابزارهایی را مستند کنند. مستندات حسابرسی ضروری است، زیرا مبنای بررسی کیفیت کار است و برنامه‌ریزی، اجرا و نظارت بر کار را تسهیل می‌کند (استاندارد ۱۲۱۵).

گولوبوفسکی^۱ (۲۰۲۵) در پژوهشی با عنوان «قابلیت تبیین و شفافیت هوش مصنوعی در حسابرسی» به بررسی میزان قابلیت تبیین و شفافیت سیستم‌های هوش مصنوعی مورد استفاده در حسابرسی پرداخت. وی با استفاده از یک رویکرد کیفی، مصاحبه‌های نیمه-ساختاریافته‌ای با نمایندگان شرکت‌های حسابرسی از شرکت‌هایی مانند دیلویت^۲، ارنست اند یانگ^۳ و رابوبانک^۴ انجام داد. یافته‌های پژوهش نشان داد که اگرچه سیستم‌های هوش مصنوعی کمک ارزشمندی در فرآیند حسابرسی ارائه می‌دهند، اما قابلیت توضیح آنها اغلب به بینش‌های سطحی محدود شده است. طراحی مدل برای تعاملات پیچیده یا پرخطر، فاقد عمق بوده است. به دلیل دسترسی محدود به داده‌های آموزشی و منطق مدل، شفافیت پایین بوده است. علاوه بر این، سیستم‌های هوش مصنوعی تا اندازه‌ای شفاف و قابل توضیح هستند و از نظر فنی به اندازه کافی برای ارائه ارزش مستقل بدون دخالت انسان کافی نیستند.

احمد صالح و همکاران (۲۰۲۵) در پژوهشی با عنوان «نگاهی به روش‌های هوش مصنوعی قابل تبیین: SHAP و LIME» به بررسی دو روش پرکاربرد هوش مصنوعی قابل تبیین برای افزایش شفافیت مدل‌های یادگیری ماشین پرداختند. نتایج این پژوهش، مبتنی بر یک مطالعه موردی زیست‌پزشکی، نشان داد که خروجی‌های SHAP و LIME به شدت وابسته به نوع مدل یادگیری ماشین بوده و در شرایط وجود همخطی بین متغیرها با محدودیت‌های تفسیری مواجه بوده است، لذا استفاده از این روش‌ها نیازمند احتیاط بوده است.

فرانسیسکو هررا (۲۰۲۵) در پژوهشی با عنوان «تأملاتی بر هوش مصنوعی قابل تبیین» به بررسی ضرورت هوش مصنوعی قابل تبیین در عصر غلبه مدل‌های جعبه‌سیاه پرداخته و بر نقش توضیح‌پذیری در تعامل انسان-هوش مصنوعی تأکید کرده است. وی هوش مصنوعی قابل تبیین را ابزاری برای افزایش درک انسانی، بهبود تصمیم‌گیری مشترک و تقویت اعتماد اجتماعی به هوش مصنوعی دانسته و بر لزوم تکامل آن در پاسخ به چرایی و هدف توضیحات تأکید کرده است.

ژونگ و گوئل^۵ (۲۰۲۴) در پژوهشی با عنوان «هوش مصنوعی شفاف در حسابرسی از طریق هوش مصنوعی قابل تبیین» پیاده‌سازی هوش مصنوعی قابل تبیین را از طریق یک

1 Golubovskij

2 Deloitte

3 Ernst & Young (EY)

4 Rabobank

5 Zhong & Goel

مورد استفاده برای تشخیص تقلب بررسی کرد و نشان داد که چگونه ادغام یک لایه قابلیت توضیح با استفاده از هوش مصنوعی قابل تبیین می‌تواند قابلیت تفسیر مدل‌های هوش مصنوعی را بهبود بخشد و ذینفعان را قادر سازد تا فرآیند تصمیم‌گیری مدل‌ها را درک کنند.

صحت برمجه و همکاران (۱۴۰۵) در پژوهشی با عنوان «ارائه الگوی قضاوت حسابرسی مبتنی بر هوش مصنوعی» ابعاد مؤثر بر قضاوت حرفه‌ای حساب‌رسان در بستر هوش مصنوعی را شناسایی کرده‌اند. یافته‌ها نشان داد عواملی مانند جمع‌آوری داده‌ها، ارزیابی مخاطرات، گزارشگری ریسک و حسابرسی مبتنی بر فناوری اطلاعات تأثیر مثبت و معناداری بر قضاوت حساب‌رسان داشته است، در حالی که برخی عوامل کنترلی و ساختاری از نظر آماری معنادار نبوده‌اند.

ریاحی‌نیا و همکاران (۱۴۰۴) در پژوهشی با عنوان «هوش مصنوعی در علم اطلاعات و دانش‌شناسی» با استفاده از روش‌های علم‌سنجی، ساختار موضوعی و روندهای پژوهشی این حوزه را تحلیل کردند. نتایج حاکی از گسترش کاربردهای هوش مصنوعی در مدیریت داده‌ها، تعامل کاربران و افزایش توجه به شفافیت الگوریتمی و ملاحظات اخلاقی در سال‌های اخیر بوده است.

رحمانی و همکاران (۱۴۰۴) در پژوهشی با عنوان «ادغام هوش مصنوعی در حسابرسی؛ چالش‌ها و مزایا» نشان دادند که هوش مصنوعی می‌تواند کارایی و دقت فرایندهای حسابرسی را بهبود بخشد، اما چالش‌هایی نظیر مسائل اخلاقی، کمبود مهارت‌های تخصصی و ریسک اتکای بیش از حد به سیستم‌های هوشمند، پیاده‌سازی آن را با محدودیت مواجه می‌کند.

۲. روش‌شناسی پژوهش

این پژوهش از لحاظ هدف، اکتشافی و از لحاظ نوع استدلال، استقرایی و از نظر ماهیت داده‌ها، پژوهش کیفی است و با به‌کارگیری روش تحلیل محتوای اسنادی به دنبال طراحی یک چارچوب مفهومی برای حسابرسی توضیح‌پذیر مبتنی بر هوش مصنوعی قابل تبیین در ارزیابی ریسک صورت‌های مالی می‌باشد. جامعه آماری این پژوهش شامل ۳۹ مقاله علمی-تخصصی است که از طریق جستجوی نظام‌مند در پایگاه‌های علمی (تمرکز بر منابع خارجی از سال ۲۰۲۰ به بعد و منابع داخلی از سال ۱۴۰۰ به بعد) استخراج و با ارزیابی کیفی دقیق، انتخاب شدند. گردآوری داده‌ها از طریق استخراج مفاهیم و مضامین اصلی (کدها) از این متون انجام شده است. تحلیل داده‌ها بر مبنای رویکرد تحلیل محتوای اسنادی استوار بود، که در آن مضامین استخراجی به صورت سلسله مراتبی دسته‌بندی و در نهایت در چهار مقوله اصلی ترکیب شدند. کیفیت این تحلیل مفهومی با اجرای گام کنترل کیفیت از طریق محاسبه ضریب کاپای کوهن برای سنجش پایایی بین تحلیل‌گران، تأیید و اعتبار روش‌شناسی پژوهش مستند گردید.

۳. یافته‌های پژوهش

گام نخست: تدوین پرسش‌های پژوهش

پرسش‌های ماهوی پژوهش براساس ابعاد گوناگون نظیر چه چیز، چه کسانی، چه زمانی و چگونه است. از پارامترهای مختلفی برای تنظیم پرسش‌های پژوهش استفاده شده است که عبارت‌اند از:

جدول ۱: سؤال‌های پژوهش

مؤلفه‌ها	سؤال‌های پژوهش
(What چه چیزی)	مروری بر مدل‌ها، چارچوب‌ها و تکنیک‌های اصلی هوش مصنوعی قابل تبیین که برای ارزیابی ریسک صورت‌های مالی در ادبیات علمی استفاده شده‌اند و همچنین شناسایی چارچوب‌های تحلیل محتوای اسنادی مرتبط.
(Who جامعه مورد مطالعه)	پایگاه‌های داده علمی قابل استناد
پایگاه‌های داده علمی قابل استناد	Scopus, Web of Science (WoS), Google Scholar, AIS Electronic Library (AisEL), IEEE Xplore.
(When حدودیت زمانی)	پژوهش‌های خارجی منتشر شده در سال‌های ۲۰۲۰ میلادی و ۱۴۰۰ شمسی به بعد
(How چگونه روش)	داده‌ها با استفاده از روش تحلیل اسناد تحلیل شدند. (در چارچوب تحلیل محتوای اسنادی، این شامل استخراج، کدگذاری، و سنتز کیفی/کمی داده‌ها از مقالات منتخب است).

منبع: نتایج تحقیق

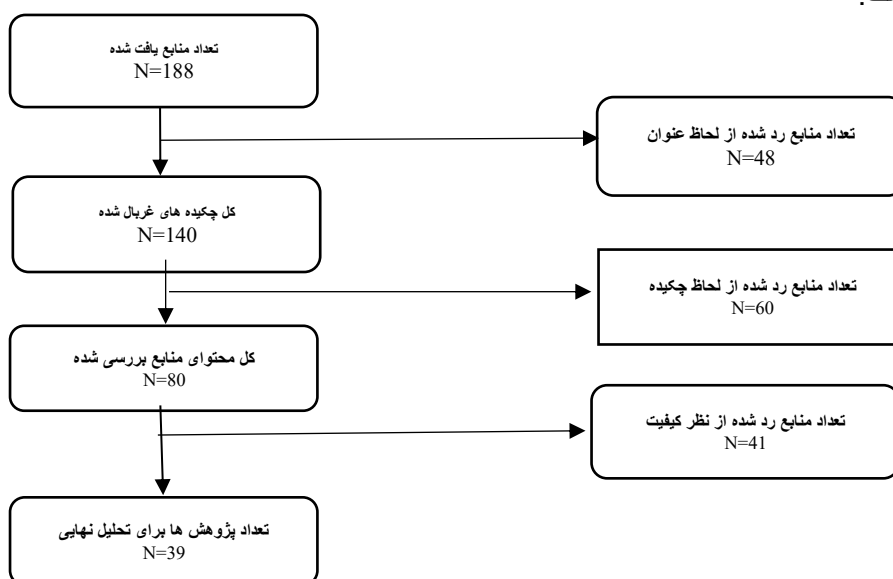
گام دوم: جستجوی سیستماتیک منابع پژوهش

پس از مشخص شدن پرسش‌های پژوهش، منابع و مستندات مرتبط با موضوع پژوهش بررسی شدند. در پژوهش حاضر، پایگاه‌های داده، مجله‌ها، مستندات و موتورهای جستجوی مختلف از سال ۲۰۲۰ به بعد برای مقالات لاتین و از سال ۱۴۰۰ به بعد برای مقالات ایرانی بررسی شدند. در این راستا، جستجوی منابع از پایگاه‌های معتبر نظیر ریسچ‌گیت، ساینس دایرکت، پروکوئست، امرالد و موتورهای جستجوی گوگل و گوگل اسکولار در دستور کار قرار گرفت. علت جستجوی پایگاه‌های خارجی این بود که تاکنون مقاله‌ای درخصوص حسابرسی توضیح‌پذیر برای ارزیابی ریسک صورت‌های مالی مبتنی بر هوش مصنوعی قابل تبیین در فصلنامه‌های ایران چاپ نشده است.

گام سوم: انتخاب مقالات و منابع مرتبط با پژوهش

در این مرحله، پژوهشگر کیفیت منابع استخراج شده را ارزیابی نمود. هدف این مرحله در فرآیند تحلیل محتوای اسنادی، حذف منابع با اعتبار کم و افزایش دقت و کیفیت تحلیل است. برای ارزیابی کیفی مطالعات اولیه از ابزار مفیدی به نام برنامه مهارت‌های ارزیابی حیاتی

استفاده می‌شود. این ابزار با مطرح کردن ده سؤال کلیدی، در ارزیابی دقت، اعتبار و اهمیت مطالعات کیفی پژوهش به پژوهشگر کمک می‌کند. پرسش‌های برنامه مهارت‌های ارزیابی حیاتی شامل هدف‌های پژوهش، منطق پژوهش، طرح پژوهش، روش نمونه‌برداری، جمع‌آوری داده‌ها، انعکاس‌پذیری، ملاحظات اخلاقی، دقت تجزیه و تحلیل داده‌ها، بیان واضح و روشن یافته‌ها و در نهایت، ارزش پژوهش است. بر همین مبنای ابتدا تعداد ۱۸۸ منبع مرتبط با موضوع پژوهش گردآوری شد. پس از بررسی عنوان و ارتباط موضوعی، ۴۸ منبع حذف شدند و تعداد منابع به ۱۴۰ رسید. در گام بعد، چکیده‌های ۱۴۰ منبع مورد ارزیابی قرار گرفت که ۶۰ منبع به دلیل عدم تطابق با موضوع پژوهش و حسابرسی توزیع پذیر کنار گذاشته شدند. سپس محتوای ۸۰ مقاله به طور کامل بررسی و بازبینی شد و بر اساس شاخص‌های کیفیتی، ۴۱ منبع دیگر حذف شدند. در نهایت، با توجه به معیارهای کیفیت و مرتبط بودن با موضوع، تعداد ۳۹ مقاله نهایی جهت تحلیل و استفاده در تحلیل محتوای اسنادی انتخاب شد. در جدول ۳، مشخصات این ۳۹ منبع انتخاب‌شده برای استخراج چارچوب نظری حسابرسی توضیح‌پذیر برای ارزیابی ریسک صورت‌های مالی مبتنی بر هوش مصنوعی قابل تبیین مطرح شدند؛ به گونه‌ای که مشخصه‌های آنها از قبیل نام نویسنده، سال انتشار، عنوان نشریه یا انتشارات و در نهایت، عنوان پژوهش ارائه شده‌اند؛ همان‌طور که ملاحظه می‌شود بازه زمانی منابع انتخاب‌شده از سال ۲۰۲۰ به بعد برای مقالات لاتین و از سال ۱۴۰۰ برای مقالات ایرانی است.



شکل ۱: نحوه انتخاب منابع نهایی تجزیه و تحلیل

منبع: یافته‌های تحقیق

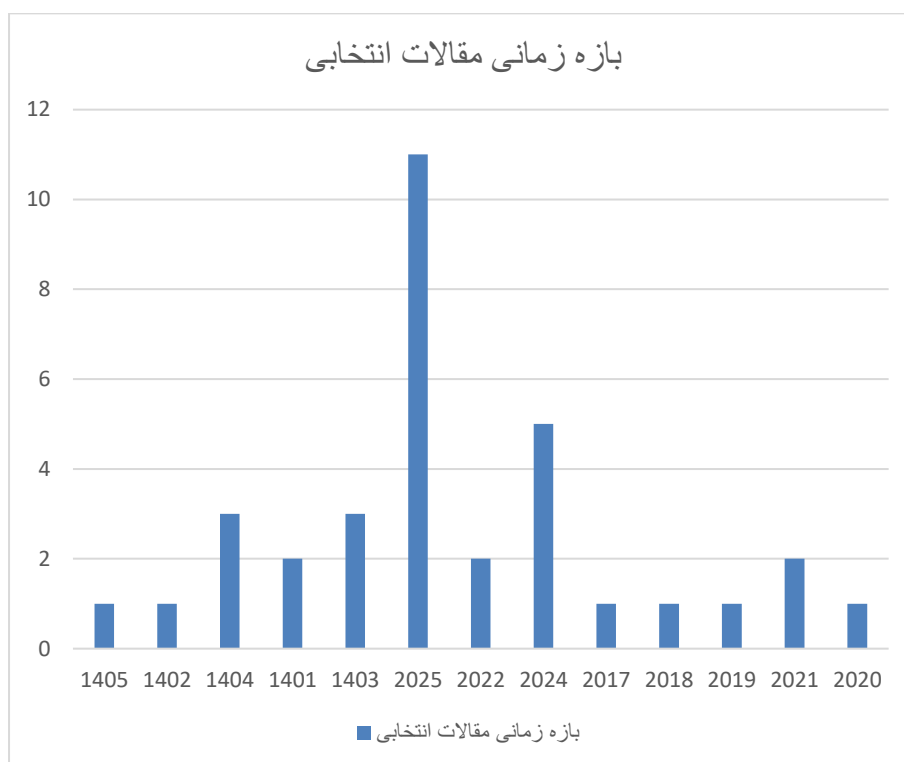
نحوه انتخاب منابع نهایی تجزیه و تحلیل در پیوست مقاله ارائه شده است. جدول ۲ و نمودار ۱ درصد فراوانی و توزیع فراوانی مقالات منتخب ۵ سال اخیر را نشان می‌دهد.

جدول ۲: درصد فراوانی مقالات منتخب

سال	تعداد فراوانی	درصد فراوانی (%)
1405	1	2.56%
1402	1	2.56%
1404	3	7.69%
1401	2	5.13%
1403	3	7.69%
2025	11	28.21%
2022	2	5.13%
2024	5	12.82%
2017	1	2.56%
2018	1	2.56%
2019	1	2.56%
2021	2	5.13%
2020	1	2.56%
مجموع	39	100.00%

منبع: نتایج تحقیق

-۱



نمودار ۱: بازه زمانی مقالات انتخاب شده (خروجی نرم افزار)

جدول ۲ نشان می‌دهد که بیشترین سهم مقالات منتخب مربوط به سال ۲۰۲۵ میلادی با ۲۸.۲۱ درصد است، پس از آن سال ۲۰۲۴ با ۱۲.۸۲ درصد و سال‌های ۱۴۰۴ و ۱۴۰۳ شمسی هر کدام با ۷.۶۹ درصد در رتبه‌های بعدی قرار دارند. سال‌های ۱۴۰۱، ۲۰۲۲ و ۲۰۲۱ هر کدام حدود ۵.۱۳ درصد سهم دارند و سایر سال‌ها هر کدام حدود ۲.۵۶ درصد را به خود اختصاص داده‌اند.

گام چهارم: استخراج مفاهیم و کدهای مرتبط با پژوهش

در این مرحله، اطلاعات مقاله‌ها و مستندات با توجه به پرسش‌های پژوهش درخصوص چابستی، چگونگی، محدوده زمانی و مکانی، طبقه‌بندی شد. سپس با استفاده از روش کدگذاری محوری مبادرت به انتخاب کدهای منحصر به فرد مرتبط با حسابرسی توضیح‌پذیر برای ارزیابی ریسک صورت‌های مالی مبتنی بر هوش مصنوعی قابل تبیین شد. در نهایت، ۷۳ کد استخراج شد. جدول شناسایی کدهای منحصر به فرد مرتبط و استخراج مفاهیم در پیوست مقاله ارائه شده

است. به طور کلی، در جدول ۳ شاخص‌هایی که از پژوهشگران استخراج گردیده، دسته‌بندی شده است.

جدول ۳: دسته‌بندی شاخص‌های استخراجی از متون مقالات براساس محققین

نام محققین	سال	شاخص‌های استخراج شده
صحت برمچه، اعظم، شماخی، حمیدرضا و یلفانی، امیر	1405	ثبت گام‌به‌گام مراحل حسابرسی، ارتباط با ریسک تحریف با اهمیت، افزایش عمق تحلیل، شفافیت منطق تصمیم
مرشد زاده، مهناز و سمیعی، ندا و مشایخ، فاطمه	1402	شواهد پشتیبان خودکار، کاهش خطر حسابرسی
عدنان حمود، محمد، پیری، پرویز و آشتاب، علی	1404	قابلیت بازبینی مسیر تحلیل، پشتیبانی تصمیم‌گیری حسابرسان
بیگ پناه، بهزاد، اثنی عشری، حمیده، هشی، عباس و اسدی، غلامحسین	1401	کاهش پیچیدگی، کاهش اتکای مکانیکی به سیستم
منصوری، محمد جواد، کرمی، غلامرضا و یزدانی، حمیدرضا	1403	ارائه توضیح متناسب با سطح کاربر، قابلیت دفاع‌پذیری تصمیمات
نورالله‌زاده، نوروز و صادقی، مرتضی	1403	استفاده از زبان حرفه‌ای حسابرسی، کاهش شکاف انتظارات
واعظ، محمد و اکبری، نعمت‌الله	1401	تمرکز بر ریسک‌های ذاتی و کنترلی، قابلیت پیگیری تصمیمات
عبدی براق‌تابی، عبدالله، صمدی، فاطمه، شاهوردیانی، شادی و حاجیها، زهره	1404	تقویت تصمیم‌گیری انسانی، کاهش زمان و هزینه حسابرسی
ریاحی نیا، نصرت، دانیالی، سمیرا، و حاصلی، داود	1403	مستندسازی منطق قضاوت، بهینه‌سازی تخصیص منابع
بزرگ‌اصل، موسی، پورمهدی، پارسا، و کریمی قره‌عمر، ادريس	1403	درخت تصمیم، انتقال دانش در مؤسسه
Anisykurlillah, I., Mukhibad, H., Jati, K. W., Rizkyana, F. W., Nurkhin, A., & Hapsoro, B. B.	2025	مدل‌های قاعده‌محور، تقویت قابلیت‌های فناوریانه

نام محققین	سال	شاخص‌های استخراج شده
Ariany, V.	2025	SHAP و LIME، پیچیدگی فزاینده
Bagwe, C.	2025	تبیین محلی و جهانی مدل، توهم درک و تفسیر نادرست
Bracci, E., Tallaki, M., & Otia, J. E.	2025	تمرکز بر اقلام پرریسک، انتقال یا ابهام مسئولیت
Carlioni, G., Berti, A., & Colantonio, S.	2025	بهینه‌سازی برنامه حسابرسی، تشدید یا عادی‌سازی سوگیری
Dayeh, J.	2025	امکان پرسش‌وپاسخ حسابرس با سیستم، تنوع فعالیت‌ها
Edward, E.	2025	تسهیل ارتباطات در تیم حسابرسی، سطح برآوردهای حسابداری
Eulerich, M., Pawlowski, J., Waddoups, N. J., & Wood, D. A.	2022	کاهش ابهام، پذیرش فناوری‌های نوین
Golubovskij, V.	2024	پاسخگویی الگوریتمی، همکاری در ارائه اطلاعات
Kalyanathaya, K. P., & Prasad, K. K.	2022	شناسایی سوگیری داده، تشویق نوآوری و آزمایش
Kassar, M., & Jizi, M.	2025	همسویی با اصول اخلاق حرفه‌ای، یادگیری از اشتباهات
Kesari, A., Sele, D., Ash, E., & Bechtold, S.	2024	داده‌های تمیز و ساختاریافته، برنامه‌های آموزشی منظم
Kokina, J., & Davenport, T. H.	2017	قابلیت ردیابی داده‌ها، مهارت‌های حسابرسان، سرمایه‌گذاری در فناوری، استانداردها و آیین‌نامه‌ها
Langer, M., Baum, K., Hartmann, K., Hessel, S., Speith, T., & Wahl, J.	2024	ملاحظات اخلاقی و حریم خصوصی، رویکرد نهادهای ناظر
Liu, B., Liu, P., Lu, W., & Olofsson, T.	2025	وجود تکنیک‌های معتبر تبیین، فشار رقابتی بازار حسابرسی

نام محققین	سال	شاخص‌های استخراج شده
Lyu, Q., & Wu, S. M.	2025	معیارهای اعتبارسنجی تبیین، انتظارات فزاینده ذی‌نفعان
Miller, T.	2019	تبیین محلی و جهانی مدل، دسترسی به زیرساخت ابری
Mumuni, F., & Mumuni, A.	2025	حسابرسان آشنا با مفاهیم داده، وجود شرکت‌های فناور معتبر
Muñoz-Ordóñez, J., Cobos, C., Vidal-Rojas, J. C., & Herrera, F.	2025	متخصصان داده آشنا با اصول حسابرسی، وجود نیروی متخصص
Ramos, M., & Esther, D.	2024	فرهنگ یادگیری و نوآوری، دوره‌های آموزشی کاربردی
Razali, F. M., Sulaiman, N., Abdul Manan, D. I., & Said, J.	2025	حمایت مدیریت ارشد، قابلیت ردیابی و فراداده غنی
Rudin, C., Chen, C., Chen, Z., Huang, H., Semenova, L., & Zhong, C.	2021	به‌روزرسانی استانداردهای حسابرسی
Salih, A. M., Raisi-Estabragh, Z., Boscolo Galazzo, I., Radeva, P., Petersen, S. E., Lekadir, K., & Menegaz, G.	2024	راهنمای عملی از سوی مؤسسات حرفه‌ای
Schumann, G., & Marx Gómez, J.	2021	وضوح مسئولیت نهایی حسابرسان
Taofoek, A.	2025	به‌رسمیت شناختن تبیین‌ها توسط ناظران
Thành, L. C., Duyen, V. T. M., & Hang, N. T. T.	2025	بلوغ سیستم‌های اطلاعاتی مالی
Whitaker, J. A., Brooks, D. M., Price, S. R., Reynolds, O. K., & Paul, C.	2020	اثر بخشی کنترل‌های داخلی
Zavitsanos, E., Spyropoulou, E., Giannakopoulos, G., & Paliouras, G.	2025	دسترسی به داده‌ها
Zhong, C., & Goel, S.	2024	درک مشترک از اهداف حسابرسی

منبع: نتایج تحقیق

گام پنجم: تجزیه و تحلیل و ترکیب یافته‌های کیفی

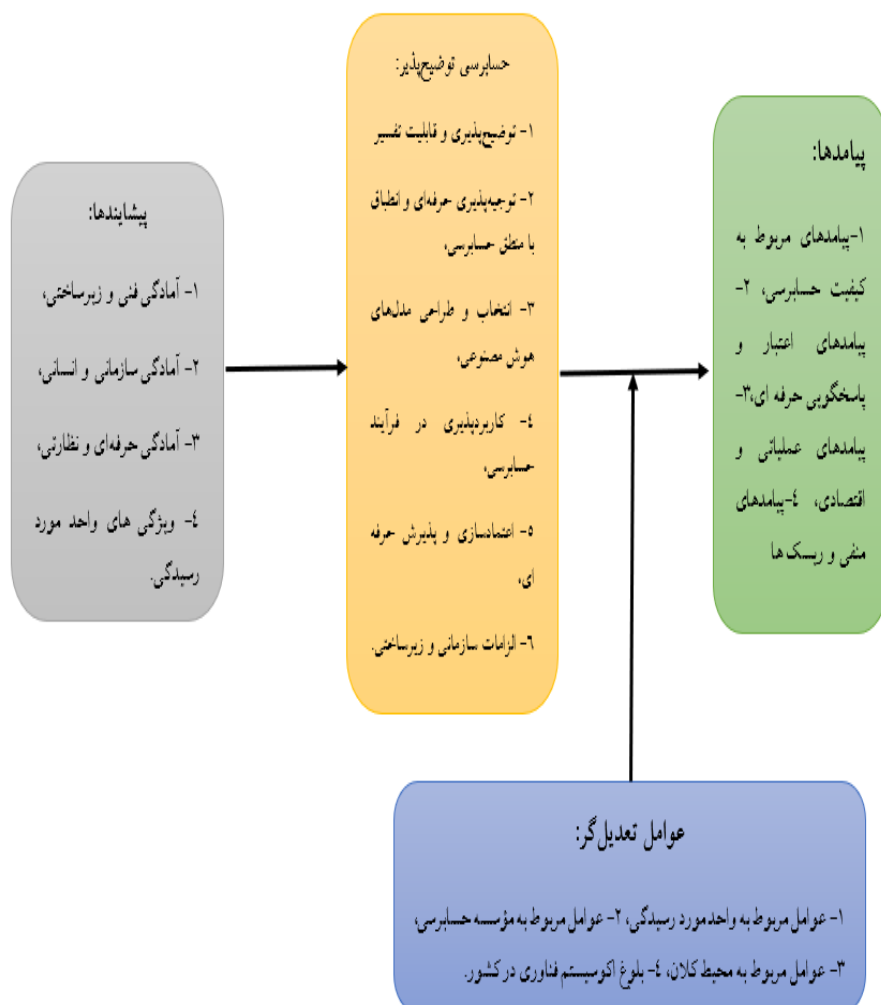
در این مرحله، محتوای مقالات به طور دقیق بررسی شد و کدهایی که با واژه‌های کلیدی پژوهش، ارتباط داشتند، انتخاب و بر مبنای آنها مفاهیم و مقوله‌های پژوهش شکل گرفتند. درواقع، ابتدا همه عوامل استخراج شده از پژوهش‌های انتخاب شده به عنوان کد، در نظر گرفته و سپس با در نظر گرفتن مفهوم هر یک از این کدهای استخراج شده، در یک مفهوم مشابه طبقه‌بندی شده‌اند؛ در نتیجه، مقوله‌های پژوهش استخراج شدند. همان‌طور که مشاهده می‌شود ۷۳ کد استخراج شده است؛ کدهای استخراج شده بدلیل شباهت معنایی و مفهومی به ۳۶ کد اولیه دسته بندی شدند و در نهایت کدهای اولیه نیز به دلیل شباهت معنایی و مفهومی به ۱۸ کد سازمان دهنده (توضیح‌پذیری و قابلیت تفسیر نتایج؛ توجیه‌پذیری حرفه‌ای و انطباق با منطق حسابرسی؛ انتخاب و طراحی مدل‌های هوش مصنوعی؛ کاربردپذیری در فرآیند حسابرسی؛ اعتمادسازی و پذیرش حرفه‌ای؛ الزامات سازمانی و زیرساختی؛ آمادگی فنی و زیرساختی؛ آمادگی سازمانی و انسانی؛ آمادگی حرفه‌ای و نظارتی؛ ویژگی‌های واحد مورد رسیدگی؛ پیامدهای مربوط به کیفیت حسابرسی؛ پیامدهای اعتبار و پاسخگویی حرفه‌ای؛ پیامدهای عملیاتی و اقتصادی؛ پیامدهای منفی و ریسک‌ها؛ عوامل مرتبط با واحد مورد رسیدگی؛ عوامل مرتبط با مؤسسه حسابرسی؛ عوامل مرتبط با محیط کلان؛ بلوغ اکوسیستم فناوری کشور) دسته بندی شدند که به ۴ کد فراگیر حسابرسی توضیح‌پذیر مبتنی بر هوش مصنوعی قابل تبیین؛ پیشایندها؛ پیامدها؛ عوامل تعدیل‌گر جای داده شدند که در جدول ۴ نشان داده شده است.

جدول ۴: ابعاد، مقوله‌ها و شاخص‌های نهایی

مفاهیم استخراجی	کدهای اولیه	کدهای سازمان دهنده	کدهای فراگیر
ثبت گام‌به‌گام مراحل حسابرسی	ردیابی‌پذیری فرآیند حسابرسی	توضیح‌پذیری و قابلیت تفسیر نتایج	حسابرسی توضیح‌پذیر مبتنی بر هوش مصنوعی قابل تبیین
شواهد پشتیبان خودکار			
قابلیت بازبینی مسیر تحلیل			
کاهش پیچیدگی			
ارائه توضیح متناسب با سطح کاربر	صورته‌بندی قابل فهم توضیحات		
استفاده از زبان حرفه‌ای حسابرسی			
ارتباط با ریسک تحریف بااهمیت	پیوند تبیین‌ها با ریسک‌های حسابرسی	توجیه‌پذیری حرفه‌ای و انطباق با منطق حسابرسی	
تمرکز بر ریسک‌های ذاتی و کنترلی			
تقویت تصمیم‌گیری انسانی	پشتیبانی از قضاوت حرفه‌ای حسابرس		
مستندسازی منطق قضاوت			
درخت تصمیم	مدل‌های ذاتاً قابل تبیین	انتخاب و طراحی مدل‌های هوش مصنوعی	
مدل‌های قاعده‌محور			
LIME و SHAP	مدل‌های پیچیده با تبیین پس‌پردازشی		
تبیین محلی و جهانی مدل	هدایت و اولویت‌بندی آزمون‌های حسابرسی	کاربردپذیری در فرآیند حسابرسی	
تمرکز بر اقلام پرریسک			
بهینه‌سازی برنامه حسابرسی	تعامل انسان-سیستم		
امکان پرسش و پاسخ حسابرس با سیستم			
تسهیل ارتباطات در تیم حسابرسی	اعتبار ادراک‌شده سیستم	اعتمادسازی و پذیرش حرفه‌ای	
کاهش ابهام			
پاسخگویی الگوریتمی	انصاف و کنترل سوگیری		
شناسایی سوگیری داده			
همسویی با اصول اخلاق حرفه‌ای	کیفیت و حاکمیت داده	الزامات سازمانی و زیرساختی	پیشایندها
داده‌های تمیز و ساختاریافته			
قابلیت ردیابی داده‌ها			
مهارت‌های حسابرسان	توانمندی سازمانی و اخلاقی		
ملاحظات اخلاقی و حریم خصوصی			
وجود تکنیک‌های معتبر تبیین	قابلیت تبیین‌پذیری مدل‌های هوش مصنوعی	آمادگی فنی و زیرساختی	
معیارهای اعتبارسنجی تبیین			
داده‌های تمیز و ساختاریافته	زیرساخت داده‌ای یکپارچه و باکیفیت		
قابلیت ردیابی و فراداده غنی			

کدهای فراگیر	کدهای سازمان دهنده	کدهای اولیه	مفاهیم استخراجی
	آمادگی سازمانی و انسانی	مهارت‌های ترکیبی تیم‌ها	حسابرسان آشنا با مفاهیم داده
			متخصصان داده آشنا با اصول حسابرسی
		فرهنگ سازمانی حمایت‌گر	فرهنگ یادگیری و نوآوری
			حمایت مدیریت ارشد
	آمادگی حرفه‌ای و نظارتی	راهنمایی‌ها و استانداردهای حرفه‌ای	بمروزرسانی استانداردهای حسابرسی
			راهنمای عملی از سوی مؤسسات حرفه‌ای
		چارچوب قانونی و نظارتی پذیرا	وضوح مسئولیت نهایی حسابرس
			بمسئمت شناختن تبیین‌ها توسط ناظران
	ویژگی‌های واحد مورد رسیدگی	بلوغ سیستم‌های اطلاعاتی مشتری	بلوغ سیستم‌های اطلاعاتی مالی
			اثر بخشی کنترل‌های داخلی
		همکاری مدیریت واحد مورد رسیدگی	دسترسی به داده‌ها
			درک مشترک از اهداف حسابرسی
پیامدها	پیامدهای مربوط به کیفیت حسابرسی	ارتقای تحلیل و ارزیابی ریسک	افزایش عمق تحلیل
			کاهش خطر حسابرسی
		تقویت قضاوت حرفه‌ای	پشتیبانی تصمیم‌گیری حسابرس
			کاهش اتکای مکانیکی به سیستم
	پیامدهای اعتبار و پاسخگویی حرفه‌ای	افزایش شفافیت و اعتماد عمومی	قابلیت دفاع‌پذیری تصمیمات
			کاهش شکاف انتظارات
		پاسخگویی و مسئولیت‌پذیری	شفافیت منطق تصمیم
			قابلیت پیگیری تصمیمات
	پیامدهای عملیاتی و اقتصادی	افزایش کارایی حسابرسی	کاهش زمان و هزینه حسابرسی
			بهینه‌سازی تخصیص منابع
		یادگیری و توسعه سازمانی	انتقال دانش در مؤسسه
			تقویت قابلیت‌های فناورانه

کدهای فراگیر	کدهای سازمان دهنده	کدهای اولیه	مفاهیم استخراجی
عوامل تعدیل‌گر	پیامدهای منفی و ریسک‌ها	چالش‌های فنی و شناختی	پیچیدگی فزاینده
			توهم درک و تفسیر نادرست
		چالش‌های حرفه‌ای و اخلاقی	انتقال یا ابهام مسئولیت
			تشدید یا عادی‌سازی سوگیری
	عوامل مرتبط با واحد مورد رسیدگی	پیچیدگی عملیات و ساختار مالی	تنوع فعالیت‌ها
			سطح برآوردهای حسابداری
		نگرش و همکاری مدیریت	پذیرش فناوری‌های نوین
			همکاری در ارائه اطلاعات
	عوامل مرتبط با مؤسسه حسابرسی	فرهنگ سازمانی	تشویق نوآوری و آزمایش
			یادگیری از اشتباهات
		مهارت‌ها و منابع	برنامه‌های آموزشی منظم
			سرمایه‌گذاری در فناوری
	عوامل مرتبط با محیط کلان	چارچوب نظارتی و حرفه‌ای	استانداردها و آیین‌نامه‌ها
			رویکرد نهادهای ناظر
		فشارهای محیطی	فشار رقابتی بازار حسابرسی
			انتظارات فزاینده ذی‌نفعان
	بلوغ اکوسیستم فناوری کشور	زیرساخت و خدمات فنی	دسترسی به زیرساخت ابری
			وجود شرکت‌های فناوری معتبر
		سرمایه انسانی و دانش فنی	وجود نیروی متخصص
			دوره‌های آموزشی کاربردی



۱- شکل ۲: چارچوب حسابرسی توضیح‌پذیر مبتنی بر هوش مصنوعی قابل تبیین (خروجی نرم افزار)

گام ششم: کنترل کیفیت

با وجود اینکه در مرحله تجزیه و تحلیل منابع انتخاب‌شده و استخراج کدهای پژوهش، حداکثر دقت لازم به عمل آمد و پژوهشگر، بازبینی و کدگذاری مجدد داده‌های استخراج‌شده را در دستورکار مستمر قرار داد و تمام اقدامات لازم برای تضمین کیفیت یافته‌های پژوهش انجام شده، برای آزمون پایایی درونی و کیفیت از ضریب کاپای کوهن استفاده شده است. بر همین مبنا، نتایج حاصل از کدهای استخراج‌شده و مقوله‌های منحصربه‌فرد حسابرسی توضیح‌پذیر برای ارزیابی ریسک صورت‌های مالی مبتنی بر هوش مصنوعی قابل تبیین برای دو نفر از نخبگان ارسال شد. پس از دریافت نظرات نخبگان و گردآوری آنها، ضریب کاپای کوهن بر مبنای توافق یا عدم توافق آنها با کدها و مقوله‌ها محاسبه شد. در نهایت، پس از محاسبه داده‌های حاصل از نظرات دو نفر از خبرگان، ضریب کاپای کوهن برابر با 0.742 با سطح معناداری 0.000 به دست آمد. این رقم نشان‌دهنده توافق معتبر و مناسب است؛ زیرا بالاتر از 0.6 است.

گام هفتم: ارائه یافته‌ها

در گام هفتم پژوهش، یافته‌های حاصل از فراتحلیل به صورت ساختاریافته در قالب چهار کد فراگیر ارائه می‌شود. کد فراگیر اول، «حسابرسی توضیح‌پذیر مبتنی بر هوش مصنوعی قابل تبیین»، بر جنبه‌های کلیدی شامل لزوم ارائه توضیحات شفاف و متناسب با مخاطب، توجیه‌پذیری منطقی تبیین‌ها با ریسک تحریف بااهمیت، و انتخاب استراتژیک مدل‌های هوش مصنوعی (ذاتاً تبیین‌پذیر یا مجهز به تکنیک‌های پس‌پردازشی نظیر SHAP/LIME) تأکید دارد؛ هدف نهایی، کاربردپذیری این سیستم‌ها در بهینه‌سازی برنامه‌های حسابرسی و تقویت تعامل انسان-سیستم است. کد فراگیر دوم، «پیشایندها»، شرایط لازم برای تحقق این هدف را در حوزه‌های فنی (مانند اعتبارسنجی تبیین)، سازمانی (مانند مهارت‌های ترکیبی تیم‌ها و حمایت مدیریتی)، حرفه‌ای (مانند بهروزرسانی استانداردها)، و ویژگی‌های واحد مورد رسیدگی (مانند بلوغ کنترل‌های داخلی) دسته‌بندی می‌کند.

کدهای فراگیر سوم و چهارم به ترتیب به پیامدها و عوامل تعدیل‌گر می‌پردازند. «پیامدها» نشان می‌دهند که اجرای موفق حسابرسی توضیح‌پذیر، کیفیت حسابرسی (کاهش خطر و تقویت قضاوت) و پاسخگویی حرفه‌ای (افزایش شفافیت و قابلیت دفاع) را بهبود بخشیده و در عین حال هزینه‌ها را کاهش می‌دهد. در نهایت، «عوامل تعدیل‌گر» شامل پیچیدگی عملیاتی واحد مورد رسیدگی، فرهنگ سازمانی مؤسسه حسابرسی، و بلوغ اکوسیستم فناوری کلان کشور هستند که اثربخشی نهایی پیامدها را تحت تأثیر قرار می‌دهند. در مجموع، یافته‌ها چارچوبی نظری و عملی برای ارتقای کیفیت و کارایی حسابرسی از طریق ادغام هوش مصنوعی قابل تبیین فراهم می‌سازد.

۴. بحث و نتیجه‌گیری

در بررسی پیامدهای مستقیم و مثبت استقرار هوش مصنوعی قابل تبیین، یافته‌ها به روشنی نشان می‌دهند که استفاده از این رویکرد منجر به ارتقاء چشمگیر کیفیت فرآیند حسابرسی می‌شود. این ارتقاء بیشتر از طریق بهبود عمق تحلیل‌های انجام شده و کاهش خطاهای سیستمی ناشی از خطای نمونه‌گیری و قضاوت انسانی غیرمدعم حاصل می‌گردد. این دستاوردها به طور مستقیم، میزان پاسخگویی حسابرس را ارتقا می‌دهند، زیرا مستندسازی و توجیه تصمیمات گرفته شده مبتنی بر هوش مصنوعی، شفاف‌تر و قابل دفاع‌تر خواهد بود. با این حال، پژوهش تأکید می‌کند که کسب این مزایا به صورت خودکار نیست و نیازمند توجه به پیشایندها است. موفقیت در گرو تأمین دو شرط اساسی است: نخست، آمادگی نیروی انسانی از طریق کسب مهارت‌های ترکیبی برای تفسیر دقیق خروجی‌های فنی و درک ریسک‌های مرتبط؛ و دوم، آمادگی نظارتی از طریق بازنگری و به‌روزرسانی استانداردهای حسابرسی متناسب با فناوری‌های نوین. علاوه بر این، نتایج به طور قاطع بر پیامدهای منفی بالقوه تمرکز دارند، به‌ویژه اتکای بیش از حد مکانیکی به خروجی‌های الگوریتم که می‌تواند منجر به نادیده گرفتن شواهد متناقض شود، و همچنین ابهام در تخصیص مسئولیت نهایی در صورت بروز خطا در فرآیند تصمیم‌گیری. این عوامل، ضرورت طراحی مکانیسم‌های داخلی قوی را برای حفظ و تقویت قضاوت حرفه‌ای حسابرس گوشزد می‌کنند. در نهایت، نقش عوامل تعدیل‌گر نظیر بلوغ فنی اکوسیستم فناوری سازمان و سطح پیچیدگی عملیاتی واحد مورد رسیدگی، مشخص می‌سازد که اثرگذاری XAI دارای یکنواختی نیست و اجرای آن نیازمند انطباق استراتژیک با محیط عملیاتی خاص است.

نتایج حاصل از این پژوهش، تائیدی بر روند جهانی و در عین حال، ارائه بینش‌های تکمیل‌کننده در بافتار خاص مورد مطالعه است. یافته‌های پژوهش نشان می‌دهند که موفقیت در شفاف‌سازی مدل‌ها، مستلزم رعایت دقیق اصل ارائه توضیح متناسب با سطح کاربر و برقراری ارتباط مستقیم میان تبیین‌های مدل و ریسک تحریف با اهمیت است؛ این رویکرد، به‌طور مستقیم دیدگاه محققانی چون هررا (۲۰۲۵) را که بر ضرورت هوش مصنوعی قابل تبیین برای تقویت درک انسانی و تعامل مؤثر انسان-هوش مصنوعی تأکید داشتند، تأیید می‌نماید. در حوزه ابزارهای فنی، استفاده از تکنیک‌هایی مانند SHAP/LIME به عنوان ابزار محوری برای شفاف‌سازی مدل‌های پیچیده، با نتایج اثبات‌شده توسط احمد صالح و همکاران (۲۰۲۵) همخوانی کامل دارد. همچنین، تأکید پژوهش حاضر بر ضرورت آمادگی‌های سازمانی و نظارتی (شامل مهارت و استانداردسازی) به مثابه پیش‌نیاز، مستقیماً با چالش‌های کمبود مهارت و مسائل اخلاقی که توسط رحمانی و همکاران (۱۴۰۴) مطرح شده بود، ارتباط پیدا می‌کند و لزوم توجه به سرمایه‌گذاری انسانی را تقویت می‌کند. برخلاف این مزایا، نتایج پژوهش ریسک‌های نهفته را برجسته می‌سازد؛ به‌ویژه خطر اتکای مکانیکی و ابهام در مسئولیت‌پذیری که بیشتر

توسط نورالزهاده و صادقی (۱۴۰۳) مورد توجه قرار گرفته بود، مورد تأکید مجدد قرار می‌گیرد و لزوم پوشش این ابعاد در کدهای فراگیر منفی را پررنگ می‌سازد. نهادهای ناظر باید بر تدوین دستورالعمل‌های عملیاتی متمرکز شوند تا سازوکارهای اعتبارسنجی خروجی‌های XAI تعریف شده و شفافیت منطق تصمیم به عنوان یک معیار بازرسی کیفیت حسابرسی لحاظ گردد. برای مؤسسات حسابرسی، سرمایه‌گذاری بر آموزش‌های تخصصی در زمینه تفسیر خروجی‌های مدل‌ها و پیوند دادن آن‌ها به برآورد ریسک تحریف بااهمیت، یک الزام است تا از وقوع اتکای صرف به نتایج ماشینی ممانعت شود. از منظر زیرساختی، کیفیت داده‌ها، شامل ردیابی‌پذیری و پاک‌سازی داده‌ها، باید به عنوان مهم‌ترین پیشایندها استقرار XAI در اولویت قرار گیرد. همچنین، لازم است تا با توجه به نگرانی‌های موجود در خصوص مسئولیت‌پذیری، سازوکارهای داخلی دقیقی برای حفظ وضوح مسئولیت نهایی حسابرس انسانی در هرگونه تصمیم‌گیری نهایی مبتنی بر هوش مصنوعی تعریف و پیاده‌سازی شود.

منابع

۱. بزرگ‌اصل، موسی، پورمهدی، پارسا و کریمی قره‌عمر، ادريس. (۱۴۰۳). هوش مصنوعی در فرایندهای حسابرسي، *حسابدار رسمي*، ۲۱(۶۷)، ۲۷.
۲. بیگ پناه، بهزاد، اثنی عشري، حمیده، هشی، عباس و اسدی، غلامحسین. (۱۴۰۱). پاسخ‌گویی مؤسسه‌های حسابرسي: رویکرد تحلیل محتوا. *بررسی‌های حسابداری و حسابرسي*، ۲۹(۲)، ۲۴۱-۲۱۳.
۳. صحت برمچه، اعظم، شماخی، حمیدرضا و یلفانی، امیر. (۱۴۰۴). تاثیر عوامل مبتنی بر مؤلفه‌های هوش مصنوعی (هوش، طراحی، انتخاب) بر قضاوت حسابرسان. *دانش حسابداری و حسابرسي مدیریت*، ۱۶(۶۱)، ۲۶۵-۲۷۸.
۴. صحت برمچه، اعظم، شماخی، حمیدرضا و یلفانی، امیر. (۱۴۰۵). ارائه الگوی قضاوت حسابرسي مبتنی بر هوش مصنوعی. *دانش حسابداری و حسابرسي مدیریت*، ۱۵(۶۰)، ۲۱۰-۱۹۱.
۵. عبدی برآفتابی، عبدالله، صمدی، فاطمه، شاهوردیانی، شادی و حاجیها، زهره. (۱۴۰۴). نقش چرخه حسابرسي داخلی مبتنی بر فنون یادگیرنده و هوش مصنوعی و تأثیرات آن در بهبود فرآیندهای حسابرسي *دانش حسابداری و حسابرسي مدیریت*، ۱۶(۶۲)، ۴۴۳-۴۵۷.
۶. عدنان حمود، محمد، پیری، پرویز و آشتاب، علی. (۱۴۰۴). امکان‌سنجی بهره‌گیری از فناوری‌های نوین هوش مصنوعی در بهبود فرایندهای حسابرسي در کشور. *بررسی‌های حسابداری و حسابرسي*، ۳۲(۳)، ۵۳۵-۵۵۹.
۷. مرشدزاده، مهناز، سمیعی، ندا و مشایخ، فاطمه. (۱۴۰۲). ارزیابی ریسک مالی شرکت‌ها بر اساس اطلاعات صورت‌های مالی، اولین کنفرانس ملی توسعه و آینده‌نگری در صنایع با رویکرد فرایندهای داخلی، تهران، <https://civilica.com/doc/2035672>.
۸. منصوری، محمدجواد، کرمی، غلامرضا و یزدانی، حمیدرضا. (۱۴۰۳). شناسایی پیشران‌های مؤثر بر آینده حسابرسي: رویکرد فراترکیب، *بررسی‌های حسابداری و حسابرسي*، ۳۱(۲)، ۳۹۰-۴۲۷.

9. Anisykurlillah, I., Mukhibad, H., Jati, K. W., Rizkyana, F. W., Nurkhin, A., & Hapsoro, B. B. (2025). Audit committee attributes and financial statement fraud risk: evidence from Indonesian banks. *COGENT BUSINESS & MANAGEMENT*, 12(1), 2566443. <https://doi.org/10.1080/23311975.2025.2566443>
10. Ariany, V. (2025). The Impact of Artificial Intelligence on Audit Quality and Auditor Judgement: A Multi Country Analysis. *ACCOUNTING STUDIES AND TAX JOURNAL (COUNT)*, 2(3), 557–569. <https://doi.org/10.62207/qpy82263>
11. Bagwe, C. (2025). Explainable AI (XAI) in Compliance Audits: Bridging the Gap Between AI and Regulatory Transparency. *INTERNATIONAL JOURNAL OF SCIENTIFIC ENGINEERING AND SCIENCE*, 9(3), 137–144.
12. Bracci, E., Tallaki, M., & Otia, J. E. (2025). Propensity factors of artificial intelligence technology adoption by public sector auditors. *JOURNAL OF PUBLIC BUDGETING, ACCOUNTING & FINANCIAL MANAGEMENT*. DOI: 10.1108/JPBAFM-09-2024-0189
13. Carloni, G., Berti, A., & Colantonio, S. (2025). The Role of Causality in Explainable Artificial Intelligence. *WILEY INTERDISCIPLINARY REVIEWS: DATA MINING AND KNOWLEDGE DISCOVERY*.
14. Dayeh, J. (2025). *AUDIT AND ARTIFICIAL INTELLIGENCE: AUDIT DATA ANALYTICS AND AUDITING AI*. Copenhagen Business School [Phd]. PhD Series No. 27.2025. <https://doi.org/10.22439/phd.27.2025>
15. Edward, E. (2025, September). *EXPLAINABLE AI FRAMEWORKS FOR ENHANCING AUDITOR INTERPRETABILITY AND REGULATORY COMPLIANCE*
16. Eulerich, M., Pawlowski, J., Waddoups, N. J., & Wood, D. A. (2022). A framework for using robotic process automation for audit tasks. *CONTEMPORARY ACCOUNTING RESEARCH*, 39(1), 691–720. <https://doi.org/10.1111/1911-3846.12723>
17. Golubovskij, V. (2025). Explainability and Transparency of AI in Auditing (Bachelor's thesis, University of Twente).
18. Kalyanathaya, K. P., & Prasad, K. K. (2022). A literature review and research agenda on explainable artificial intelligence (XAI). *INTERNATIONAL JOURNAL OF APPLIED ENGINEERING AND MANAGEMENT LETTERS (IAEML)*, 6(1), 43–59. <https://doi.org/10.5281/zenodo.5998488>

19. Kassar, M., & Jizi, M. (2025). Artificial intelligence and robotic process automation in auditing and accounting: a systematic literature review. *JOURNAL OF APPLIED ACCOUNTING RESEARCH*. DOI 10.1108/JAAR-05-2024-0175
20. Kesari, A., Sele, D., Ash, E., & Bechtold, S. (2024). A LEGAL FRAMEWORK FOR EXPLAINABLE ARTIFICIAL INTELLIGENCE (Center for Law & Economics Working Paper Series 09/2024). ETH Zurich. <https://doi.org/10.3929/ethz-b-000699762>
21. Kokina, J., & Davenport, T. H. (2017). The emergence of artificial intelligence: How automation is changing auditing. *JOURNAL OF EMERGING TECHNOLOGIES IN ACCOUNTING*, 14(1), 115–122. <https://doi.org/10.2308/jeta-51730>
22. Kokina, J., Blanchette, S., Davenport, T. H., & Pachamanova, D. (n.d.). Challenges and opportunities for artificial intelligence in auditing: Evidence from the field. *INTERNATIONAL JOURNAL OF ACCOUNTING INFORMATION SYSTEMS*.
23. Langer, M., Baum, K., Hartmann, K., Hessel, S., Speith, T., & Wahl, J. (2024). EXPLAINABILITY AUDITING FOR INTELLIGENT SYSTEMS: A RATIONALE FOR MULTI-DISCIPLINARY PERSPECTIVES.
24. Liu, B., Liu, P., Lu, W., & Olofsson, T. (2025). Explainable Artificial Intelligence (XAI) for Material Design and Engineering Applications: A Quantitative Computational Framework. *INTERNATIONAL JOURNAL OF MECHANICAL SYSTEM DYNAMICS*, 5, 236–265. <https://doi.org/10.1002/msd2.70017>
25. Lyu, Q., & Wu, S. M. (2025). Explainable Artificial Intelligence for Business and Economics: Methods, Applications and Challenges. *EXPERT SYSTEMS*, 42, e70017. <https://doi.org/10.1111/exsy.70017>
26. Miller, T. (2019). Explanation in artificial intelligence: Insights from the social sciences. *ARTIFICIAL INTELLIGENCE*, 267, 1–38.
27. Mumuni, F., & Mumuni, A. (2025). EXPLAINABLE ARTIFICIAL INTELLIGENCE (XAI): FROM INHERENT EXPLAINABILITY TO LARGE LANGUAGE MODELS. arXiv:2501.09967. Retrieved from [URL یا آدرس arXiv]

28. Muñoz-Ordóñez, J., Cobos, C., Vidal-Rojas, J. C., & Herrera, F. (2025). A maturity model for practical explainability in artificial intelligence-based applications: Integrating analysis and evaluation (MM4XAI-AE) models. *INTERNATIONAL JOURNAL OF INTELLIGENT SYSTEMS*, XX(Y), pp-pp. <https://doi.org/10.1155/int/4934696>
29. Ramos, M., & Esther, D. (2024). *DEVELOPING EXPLAINABLE AI FRAMEWORKS FOR REGULATED INDUSTRIES*. ResearchGate. Retrieved from <https://www.researchgate.net/publication/390175104>
30. Razali, F. M., Sulaiman, N., Abdul Manan, D. I., & Said, J. (2025). Sustainability of Audit Profession in Digital Technology Era: The Role of Competencies and Digital Technology Capabilities to Detect Fraud Risk. *SAGE OPEN*, 15(1), 21582440241304974. <https://doi.org/10.1177/21582440241304974>
31. Rudin, C., Chen, C., Chen, Z., Huang, H., Semenova, L., & Zhong, C. (2021). Interpretable machine learning: Fundamental principles and 10 grand challenges. *DUKE UNIVERSITY AND UNIVERSITY OF MAINE*. <https://www.duke.edu/~rudin/papers/interpretable-ml.pdf>
32. Salih, A. M., Raisi-Estabragh, Z., Boscolo Galazzo, I., Radeva, P., Petersen, S. E., Lekadir, K., & Menegaz, G. (2024). A perspective on explainable artificial intelligence methods: SHAP and LIME. *ADVANCED INTELLIGENT SYSTEMS*, 7(1), 2400304. <https://doi.org/10.1002/aisy.202400304>
33. Schumann, G., & Marx Gómez, J. (2021). *NATURAL LANGUAGE PROCESSING IN INTERNAL AUDITING – A STRUCTURED LITERATURE REVIEW*. In *AMCIS 2021 PROCEEDINGS AMERICAS CONFERENCE ON INFORMATION SYSTEMS (AMCIS)*. Association for Information Systems (AIS Electronic Library (AISeL)). https://aisel.aisnet.org/amcis2021/sig_acctinfosystem_asys/sig_acctinfosystem_asys/1
34. Taofeek, A. (2025, September). *HUMAN-IN-THE-LOOP APPROACHES FOR COMBINING AI RECOMMENDATIONS WITH AUDITOR EXPERTISE*.
35. Thành, L. C., Duyen, V. T. M., & Hang, N. T. T. (2025). Explainable artificial intelligence in auditing: Factors influencing auditors' acceptance in Vietnam. *EDELWEISS APPLIED SCIENCE AND TECHNOLOGY*, 9(12), 81–92. <https://doi.org/10.55214/2576-8484.v9i12.11281>

36. Whitaker, J. A., Brooks, D. M., Price, S. R., Reynolds, O. K., & Paul, C. (2020). DEEP LEARNING APPROACHES FOR ANOMALY DETECTION IN FINANCIAL TRANSACTIONS.
37. Zavitsanos, E., Spyropoulou, E., Giannakopoulos, G., & Paliouras, G. (2025). Machine learning for identifying risk in financial statements: A survey. *ACM COMPUTING SURVEYS*, 57(9), Article 220. <https://doi.org/10.1145/3723157>
38. Zhong, C., & Goel, S. (2024). Transparent AI in auditing through explainable AI. *CURRENT ISSUES IN AUDITING*, 18(2), A1–A14. American Accounting Association. <https://doi.org/10.2308/CIIA-2023-009>
39. Zhang, C. A., Cho, S., & Vasarhelyi, M. (2022). Explainable artificial intelligence (XAI) in auditing. *International Journal of Accounting Information Systems*, 46, 100572.