

دریافت مقاله: ۱۴۰۳/۰۹/۲۰

پذیرش مقاله: ۱۴۰۳/۱۰/۲۹

نوع مقاله: پژوهشی

کاربرد هوش مصنوعی قابل توضیح برای تحلیل نیازمندی‌های نرم‌افزار

بهناز جاماسب

دانشجوی دکتری، دانشکده مهندسی کامپیوتر و فناوری اطلاعات، دانشگاه صنعتی شیراز

پست الکترونیکی: b.jamasb@sutech.ac.ir

سید رئوف خیامی*

دانشیار، دانشکده مهندسی کامپیوتر و فناوری اطلاعات، دانشگاه صنعتی شیراز، شیراز، ایران

پست الکترونیکی: khayami@sutech.ac.ir

رضا اکبری

دانشیار، دانشکده مهندسی کامپیوتر و فناوری اطلاعات، دانشگاه صنعتی شیراز، شیراز، ایران

پست الکترونیکی: akbari@sutech.ac.ir

چکیده

نیازمندی‌های تولید شده توسط طبقه‌بند می‌پردازد. نتایج نشان می‌دهد که XAI می‌تواند به تحلیل‌گر کمک کند تا بهتر درک کند که چرا یک نیازمندی خاص به‌عنوان کارکردی، غیرکارکردی یا زیرشاخه‌ای از غیرکارکردی طبقه‌بندی شده است. در طی این فرآیند، کلمات کلیدی مهم برای چنین تصمیماتی شناسایی شده و به‌طور دقیق تحلیل می‌شوند. کلمات کلیدی: هوش مصنوعی قابل توضیح، تحلیل نیازمندی نرم‌افزار، طبقه‌بندی نیازمندی‌ها

۱. مقدمه

مهندسی نیازمندی‌ها^۱ یکی از مراحل اصلی در چرخه حیات توسعه نرم‌افزار است و شامل فعالیت‌های مختلفی است که توسط تحلیل‌گران نرم‌افزار انجام می‌شود. استخراج نیازمندی‌ها، تحلیل، اعتبارسنجی و مدیریت معمولاً به عنوان فعالیت‌های عمومی در مهندسی نیازمندی‌ها شناخته می‌شوند [۱]. فعالیت استخراج نیازمندی‌ها به شناسایی

الگوریتم‌های یادگیری ماشین در مسایل مختلف مانند طبقه‌بندی نیازمندی‌ها استفاده شده‌اند. با وجود این که این الگوریتم‌ها عملکرد خوبی دارند، اما نمی‌توانند نحوه انجام تصمیم‌گیریشان را توضیح دهند. توضیح نحوه تصمیم‌گیری کمک بزرگی برای متخصصین این حوزه خواهد بود. هوش مصنوعی قابل توضیح (XAI) به دنبال توضیح پیش‌بینی‌های الگوریتم‌های یادگیری ماشین است. با توجه به کاربرد موفق هوش مصنوعی قابل توضیح در حوزه پردازش زبان طبیعی و عدم وجود مطالعات کافی در زمینه به‌کارگیری آنها در تحلیل نیازمندی‌های نرم‌افزار، در این مقاله، کاربرد XAI برای طبقه‌بندی نیازمندی‌های نرم‌افزار بررسی شده است. روش پیشنهادی روی مجموعه داده نیازمندی‌های نرم‌افزار PROMISE پیاده‌سازی شده است. الگوریتم ماشین بردار پشتیبان (SVM) به‌عنوان الگوریتم طبقه‌بندی استفاده شده و ابزار LIME به تفسیر دسته‌بندی

1-Requirements Engineering

* نویسنده مسئول

خدماتی که نرم افزار باید ارائه دهد، همراه با محدودیت‌ها و معیارهای کیفیتی که باید توسط نرم افزار برآورده شود، می‌پردازد. این فعالیت شامل مراحل مختلفی از جمله کشف نیازمندی‌ها، طبقه‌بندی نیازمندی‌ها، اولویت‌بندی نیازمندی‌ها و مشخصات نیازمندی‌ها است.

فعالیت استخراج نیازمندی‌ها تأثیر مهمی بر کیفیت نهایی نرم افزار دارد. هرگونه نقص در این فعالیت مانند نیازمندی‌های از دست رفته، نیازمندی‌های غیر قابل پیگیری، نیازمندی‌های متناقض، ابهام در نیازمندی‌ها و غیره، شانس توسعه نرم افزار با کیفیت بالا را کاهش می‌دهد. کیفیت نرم افزار به نیازمندی‌های صحیح وابسته است و چون استخراج نیازمندی‌ها یک کار زمان‌بر است، استفاده از روش‌ها و راه‌حل‌های خودکار جهت استخراج نیازمندی‌ها، پیشنهاد می‌شود. به عنوان مثال، روش‌های یادگیری ماشین مانند طبقه‌بندی کنندگان متنی می‌توانند برای دسته‌بندی نیازمندی‌های کارکردی^۲ و نیازمندی‌های غیرکارکردی^۳ مورد استفاده قرار گیرند.

معمولاً نیازمندی‌های نرم افزار به دو دسته کارکردی و غیرکارکردی دسته‌بندی می‌شوند. نیازمندی‌های کارکردی مسئولیت‌ها یا وظایف نرم افزار را نشان می‌دهند و نیازمندی‌های غیرکارکردی نمایانگر کیفیت یا محدودیت‌هایی هستند که نرم افزار در حین انجام وظایف خود باید آن‌ها را برآورده کند. بیشتر نیازمندی‌های غیرکارکردی معیارهای کیفیت نرم افزار مانند امنیت، عملکرد، قابلیت اطمینان، قابلیت استفاده و غیره را نشان می‌دهند [۱۸].

در سال‌های اخیر، تأثیر هوش مصنوعی^۴ بر حوزه‌های مختلف افزایش یافته است. بسیاری از راه‌حل‌های هوش مصنوعی توسط پژوهشگران ارائه شده‌اند تا مشکلات پیچیده را حل کنند و به افراد و متخصصان حوزه‌های مختلف کمک کنند تا کارهای خود را بهتر انجام دهند. یادگیری ماشین یک زیرمجموعه از هوش مصنوعی است که شامل دسته‌های

مختلف روش‌های نظارت شده^۵، نظارت نشده^۶ و تقویتی^۷ است.

هوش مصنوعی قابل توضیح (XAI)^۸، یک زیرمجموعه از هوش مصنوعی است که کمک می‌کند سیستم‌های هوش مصنوعی شفاف و قابل درک برای انسان‌ها توسعه داده شوند. هوش مصنوعی قابل توضیح نقش مهمی در حوزه‌های مختلف دارد و به انسان‌ها این امکان را می‌دهد تا استدلال پشت نتایج تولید شده توسط هوش مصنوعی را درک کنند. به این ترتیب، XAI اشکال‌زدایی و عیب‌یابی را تسهیل می‌کند، انطباق با قوانین را تضمین می‌کند و در نهایت باعث توسعه سیستم‌های هوش مصنوعی قابل اعتماد و قابل اطمینان می‌شود.

روش‌های یادگیری نظارت‌شده که شامل طبقه‌بندی و رگرسیون هستند، به‌طور موفقیت‌آمیزی برای حل مشکلات مهندسی نرم افزار به کار رفته‌اند. با این حال، در بیشتر موارد این روش‌ها به عنوان یک جعبه سیاه استفاده می‌شوند و راه‌حل‌های آن‌ها قابل تفسیر نیستند. به دلیل اهمیت تحلیل نیازمندی‌ها، قابلیت اطمینان به مدل جعبه سیاه بسیار مهم است. XAI می‌تواند کمک کند تا درجه اطمینان روش‌های نظارت‌شده در طبقه‌بندی نیازمندی‌ها را تعیین کنیم. با وجود به کارگیری موفق XAI در حوزه‌های مختلف پردازش زبان طبیعی و تحلیل متن، اما تحقیقات بسیار محدودی در زمینه استفاده از XAI جهت تفسیر نتایج الگوریتم‌های یادگیری ماشین که برای طبقه‌بندی نیازمندی‌های نرم افزار استفاده می‌شوند، گزارش شده است. با توجه به این که پردازش نیازمندی‌های نرم افزار را می‌توان زیرمجموعه‌ای از مسایل مورد بحث در حوزه پردازش زبان طبیعی در نظر گرفت، اما ویژگی‌های منحصر بفردی در این مسئله وجود دارد که در سایر حوزه‌های پردازش زبان طبیعی وجود ندارد. برای مثال کلماتی که اولویت نیازمندی‌ها را تعیین می‌کنند، یا واژه‌هایی که

5-Supervised Learning

6-Unsupervised Learning

7-Reinforcement

8-Explainable Artificial Intelligence

2-Functional Requirements

3-Non-functional Requirements

4-Artificial Intelligence

یک نیازمندی به عنوان کارکردی یا غیرکارکردی طبقه‌بندی شده است. نوآوری‌های اصلی این مقاله عبارتند از:

- برای اولین بار از XAI برای توضیح طبقه‌بندی‌های نیازمندی‌های نرم‌افزار استفاده شده است.
- یک چارچوب برای توضیح الگوریتم طبقه‌بندی نیازمندی‌های نرم‌افزار پیشنهاد شده است.
- یک مدل توضیح‌دهنده پیشرفته برای توضیح کیفیت طبقه‌بندی‌کننده نیازمندی‌های نرم‌افزار استفاده شده است. ساختار مقاله به این صورت است: بخش ۲ به کارهای پیشین می‌پردازد. معرفی XAI و الگوریتم LIME در بخش ۳ ارائه می‌شود. بخش ۴ روش پیشنهادی را ارائه می‌کند. بخش ۵ به کاربرد مدل و نتایج تجربی اختصاص دارد. در نهایت، در بخش ۶ نتیجه‌گیری این کار ارائه می‌شود.

۲. مطالعات پیشین

مسئله طبقه‌بندی نیازمندی‌های نرم‌افزار را می‌توان به عنوان یکی از مسایل موجود در حوزه پردازش زبان طبیعی در نظر گرفت. زیرا نیازمندی‌های نرم‌افزاری معمولاً به زبان طبیعی بیان می‌شوند، چه به صورت مستندات رسمی (مانند مشخصات نیازمندی‌های نرم‌افزار) و چه به شکل‌های غیررسمی (مانند داستان‌های کاربر یا درخواست‌های ویژگی). این وظیفه شامل پردازش و درک این متون زبان طبیعی برای دسته‌بندی آن‌ها به گروه‌های معناداری مانند نیازمندی‌های عملکردی و نیازمندی‌های غیرعملکردی است. لذا برخی از مطالعات انجام شده در این حوزه با تمرکز بر هوش مصنوعی قابل توضیح ارائه می‌شود.

اخیراً کاربردهای هوش مصنوعی قابل توضیح در زمینه پردازش داده‌های متنی و پردازش زبان طبیعی مورد مطالعه قرار گرفته است. برای مثال، دیوالی و همکاران مروری را بر روی روش‌های تحلیل احساسات با تمرکز بر استفاده از هوش مصنوعی قابل توضیح ارائه داده‌اند [۱۹]. در مطالعه دیگری، ماتئوس کاربردهای هوش مصنوعی

نشان‌دهنده نوع نیازمندی‌های عملکردی و غیرعملکردی هستند نشان‌دهنده تفاوت‌هایی است که به طور خاص در پردازش نیازمندی‌ها دارای اهمیت هستند. بنابراین می‌توان این مورد را به عنوان یکی از موضوعات باز در زمینه پردازش نیازمندی‌های نرم‌افزار که به زبان طبیعی بیان می‌شوند در نظر گرفت. این مسئله نویسندگان مقاله را بر این داشته است تا یک مدل توضیح‌دهنده شناخته شده برای توضیح الگوریتم‌های طبقه‌بندی نیازمندی نرم‌افزار ارائه دهند.

در این پژوهش روشی برای توضیح نتایج تولید شده توسط یک الگوریتم طبقه‌بندی نیازمندی ارائه می‌گردد. الگوریتم‌های طبقه‌بندی مختلفی وجود دارند که بخشی از آن‌ها مانند درخت تصمیم^۹ و رگرسیون خطی^{۱۰} قابل توضیح هستند، در حالی که برخی دیگر مانند شبکه‌های عصبی عمیق^{۱۱}، ماشین‌های بردار پشتیبان^{۱۲} و جنگل تصادفی^{۱۳} به دلیل پیچیدگی ذاتی یا عدم شفافیت، قابل توضیح نیستند. معمولاً این طبقه‌بندی‌های پیچیده در مقایسه با مدل‌های ساده و قابل توضیح، نتایج بهتری ارائه می‌دهند و می‌توانند الگوهای پیچیده و تعاملات آن‌ها را درک کنند. بنابراین، استفاده از مدل‌های پیچیده به دلیل عملکردشان خوب است و XAI می‌تواند کمبود ذاتی تفسیرپذیری آن‌ها را برطرف کند.

در اینجا از طبقه‌بند SVM استفاده شده است [۲]. با این حال، امکان جایگزینی SVM با سایر الگوریتم‌های طبقه‌بندی پیچیده وجود دارد. همچنین از توضیح‌دهنده LIME به همراه دانش متخصصان این حوزه برای تفسیر دسته‌بندی نیازمندی‌هایی که توسط الگوریتم طبقه‌بندی تولید شده، استفاده شده است. برای این منظور، ما روی کلمات کلیدی در نیازمندی‌ها که از راه حل طبقه‌بند مورد نظر حمایت می‌کنند یا با آن مخالفت دارند، تمرکز می‌کنیم. در نهایت به متخصصان حوزه کمک می‌شود تا بدانند چرا

9-Decision Tree
10-Linear Regression
11-Deep Neural Networks
12-Support Vector Machines
13-Random Forest

قابل توضیح را در حوزه‌های پردازش زبان طبیعی، زیست پزشکی و تشخیص بدافزار بررسی کرده است [۲۰]. ارزیابی وی بر روی پردازش زبان طبیعی برای طبقه‌بندی داده‌های توییت انجام می‌شود، طبقه‌بندی سیگنال‌های زیست‌پزشکی برای دسته‌بندی سرطان و طبقه‌بندی بدافزار رایانه شخصی است. زهور و زکریا از رهیافتی مبتنی بر هوش مصنوعی قابل توضیح برای دسته‌بندی نظرات کاربران در مورد برنامه کاربردی موبایل استفاده کرده‌اند [۲۱]. ژانگ و آوگاساوارا از یک روش مبتنی بر هوش مصنوعی قابل توضیح برای دسته‌بندی متون پزشکی استفاده کرده‌اند. آنها نشان داده‌اند الگوریتم موردنظر از چه کلماتی برای دسته‌بندی استفاده می‌کند [۲۲]. در پژوهش دیگری، از هوش مصنوعی قابل توضیح در پردازش زبان طبیعی استفاده شده است. نویسندگان کارایی دو روش توضیح مدل‌محور بر کاربرد را بر روی مجموعه داده جملات محاوره‌ای اعمال کرده‌اند و نتیجه آنها را گزارش داده‌اند [۲۳]. این مطالعات نشان می‌دهند هوش مصنوعی قابل توضیح به‌طور موفق برای حل مسایل موجود در حوزه پردازش زبان طبیعی قابل استفاده هستند و لذا می‌توان از آنها برای تحلیل نیازمندی‌های نرم‌افزار به همراه الگوریتم‌های طبقه‌بندی یادگیری ماشین استفاده کرد.

در سال‌های اخیر از الگوریتم‌های مختلف نظارت‌شده توسط پژوهشگران برای طبقه‌بندی نیازمندی‌ها استفاده شده است. در ادامه به کارهای انجام شده در زمینه طبقه‌بندی نیازمندی‌ها و کاربرد هوش مصنوعی قابل توضیح در حوزه مهندسی نیازمندی‌ها پرداخته می‌شود. بر اساس تحقیقات و اطلاعات نویسندگان مقاله حاضر، مطالعات بسیار محدودی در زمینه کاربرد هوش مصنوعی قابل توضیح برای تحلیل نیازمندی‌های نرم‌افزار رایج شده است. در ابتدا برخی از مطالعات مربوط به الگوریتم‌های طبقه‌بندی یادگیری ماشین جهت طبقه‌بندی نیازمندی‌ها رایج می‌شود و سپس مطالعات مرتبط به هوش مصنوعی قابل توضیح در ارتباط با مهندسی نیازمندی‌ها مطرح می‌گردد.

بین‌خونین و ژائو مروری نظام‌مند بر روش‌های یادگیری ماشین برای شناسایی و طبقه‌بندی نیازمندی‌های غیرکارکردی ارائه کردند [۳]. مطالعات آنها نشان داد که ۱۶ روش مختلف یادگیری ماشین توسط پژوهشگران استفاده شده و الگوریتم‌های نظارت‌شده محبوب‌ترین آنها هستند. آنها ذکر کردند که در حالی که روش‌های یادگیری ماشین در تحلیل نیازمندی‌ها موفق بوده‌اند، با چالش‌های باز در این حوزه مواجه هستیم که می‌توان آنها را با همکاری بین پژوهشگران مهندسی نرم‌افزار و یادگیری ماشین حل کرد. خیاشی و همکاران از تکنیک‌های مختلف یادگیری عمیق برای طبقه‌بندی نیازمندی‌های نرم‌افزار به عنوان کارکردی یا غیرکارکردی استفاده کردند [۴]. آنها پنج الگوریتم یادگیری عمیق همراه با دو روش رای‌گیری^{۱۴} را روی مجموعه داده PURE اعمال کردند. نتایج نشان داد که اعمال روش‌های رای‌گیری بر روی الگوریتم‌های یادگیری عمیق نتایج بهتری تولید می‌کند. در پژوهشی دیگر، کانیدو و مندس به بررسی ترکیب سه روش استخراج ویژگی به نام‌های بسته کلمات (BoW)^{۱۵}، فراوانی اصطلاح-معکوس فراوانی متن (TF-IDF)^{۱۶} و کای‌دو (CHI²)^{۱۷} با چهار الگوریتم طبقه‌بندی، ماشین بردار پشتیبان (SVM)^{۱۸}، بیز ساده چندجمله‌ای^{۱۹}، نزدیک‌ترین همسایه^{۲۰} و رگرسیون لجستیک^{۲۱} پرداختند [۵]. آنها دریافتند که ترکیب TF-IDF و رگرسیون لجستیک بهترین عملکرد را در طبقه‌بندی دوتایی دارد.

طبقه‌بندی نیازمندی‌های کارکردی توسط رحیمی و همکاران بررسی شد [۶]. آنها یک روش تجمعی با ترکیب پنج الگوریتم طبقه‌بندی پیشنهاد کردند. در کار دیگری، آنها الگوریتم‌های مختلف یادگیری عمیق را برای طبقه‌بندی نیازمندی‌ها در روش‌های یک‌مرحله‌ای و دو مرحله‌ای اعمال کردند. ابتدا نیازمندی‌ها را به عنوان کارکردی و غیرکارکردی طبقه‌بندی کردند. سپس

14-Voting Method

15-Bag of Words

16-Term Frequency-Inverse Document Frequency

17-Chi-Squared

18-Support Vector Machine

19-Multinomial Naive Bayes

20-k-Nearest Neighbors

21-Logistic Regression

تحت تاثیر قرار می‌گیرد و فاقد بینش‌های سراسری است، زیرا تنها بر روی رفتارهای محلی تمرکز می‌کند. کیفیت توضیحات ارایه شده توسط این الگوریتم به راهبرد تغییرات بستگی دارد، بنابراین ممکن است اثرات منفی روی توضیحات ارایه شده برای داده‌های متنی باشد. استفاده از مدل‌های خطی به عنوان مدل‌های جایگزین، می‌تواند مرزهای تصمیم‌گیری پیچیده را ساده‌سازی کند و تعادل بین تفسیرپذیری و دقت معمولاً منجر به توضیحاتی می‌شود که یا بیش از حد ساده یا بسیار پیچیده هستند. LIME از نظر محاسباتی پرهزینه است، به مقیاس‌بندی ویژگی‌ها حساس است و در برابر نمونه‌های خصمانه آسیب‌پذیر است، که می‌تواند منجر به توضیحات گمراه‌کننده شود. علاوه بر این، هیچ معیاری برای کمی‌سازی اعتبار توضیحات آن وجود ندارد. رسیدگی به این مشکلات از طریق بهبود ثبات، روش‌های قوی‌تر برای ایجاد تغییرات، مدل‌های جایگزین پیشرفته و معیارهای بهتر می‌تواند قابلیت اطمینان و مفید بودن LIME را افزایش دهد.

لونبرگ و لی الگوریتم SHAP را ارایه کردند [۱۰]. مشابه LIME، الگوریتم SHAP می‌تواند توضیح دهد که هر ویژگی یک الگوریتم طبقه‌بندی یا رگرسیون چگونه به یک پیش‌بینی کمک می‌کند. SHAP یک چارچوب توضیحی است که از رویکرد نظریه‌بازی‌ها استفاده می‌کند و به ما کمک می‌کند بفهمیم چرا یک مدل یادگیری ماشین خروجی خاصی را تولید می‌کند. الگوریتم SHAP، در حالی که برای ارایه توضیحات قابل تفسیر در یادگیری ماشین قدرتمند است، با چندین چالش روبه‌رو است. مشابه با الگوریتم LIME، این الگوریتم نیز از نظر محاسباتی پرهزینه است، به‌ویژه برای مدل‌هایی با ویژگی‌های زیاد یا داده‌های بزرگ، زیرا نیاز به محاسبه مقادیر شاپلی برای هر ترکیب ممکن از ویژگی‌ها دارد. SHAP همچنین در فضاهای با ابعاد بالا که تعداد ترکیب‌های ممکن ویژگی‌ها به طور نمایی رشد می‌کند، با مقیاس‌پذیری مشکل دارد. وابستگی SHAP به نظریه بازی ممکن است منجر به ناسازگاری در

نیازمندی‌های غیرکارکردی را زیرشاخه‌بندی کردند. همچنین قوبا و همکاران به بررسی عملکرد ترکیب BOW با SVM و KNN برای طبقه‌بندی نیازمندی‌های نرم‌افزار پرداختند [۷]. مشخص شد که ترکیب BOW و SVM نتایج بهتری تولید می‌کند.

مطالعات قبلی نشان دادند که روش‌های یادگیری ماشین در طبقه‌بندی نیازمندی‌ها عملکرد خوبی دارند. با وجود کاربرد موفقیت‌آمیز یادگیری ماشین، این روش‌ها به عنوان راه‌حل‌های جعبه سیاه استفاده می‌شوند. این روش‌ها به دلیل عدم شفافیت برای کاربران، جعبه سیاه نامیده می‌شوند. فرض کنید که یک روش یادگیری نظارت‌شده پیچیده داشته باشیم، به دلیل پیچیدگی آن، نمی‌توانیم درک کنیم که الگوریتم چگونه نتیجه خاصی را تولید می‌کند. در برخی موارد، یک متخصص حوزه نیاز دارد بداند الگوریتم چگونه به چنین نتیجه‌ای رسیده است تا بتواند به آن اعتماد کند.

در اینجا، XAI وارد عمل می‌شود. XAI روشی است که از روش‌ها و تکنیک‌های مختلف برای توضیح مدل‌های هوش مصنوعی استفاده می‌کند. در سال‌های اخیر، روش‌های XAI مختلفی پیشنهاد شده‌اند. توضیحات مدل‌های محلی و مستقل از الگوریتم (LIME) [۲۲] و توضیحات افزودنی شاپلی (SHAP) [۲۳] از جمله متداول‌ترین توضیح‌دهنده‌ها هستند [۸]. هر دوی آن‌ها به عنوان مدل‌های جانشین شناخته می‌شوند و از مدل‌های یادگیری ماشین جعبه سیاه استفاده می‌کنند.

LIME توسط ریبیرو و همکاران ارایه شد [۹]. هدف LIME توضیح نتایج تولیدشده توسط هر مدل طبقه‌بندی یا رگرسیون است. LIME یک پیش‌بینی را با تقریب زدن آن به صورت محلی توضیح می‌دهد. البته الگوریتم LIME که به طور گسترده‌ای برای توضیح پیش‌بینی‌های مدل‌های یادگیری ماشین استفاده می‌شود، با چندین چالش روبه‌رو است که روی قابلیت‌های آن تاثیرگذار است. کارایی این الگوریتم به دلیل تصادفی بودن در نمونه‌برداری تغییرات،

22-Local Interpretable Model-Agnostic Explanations

23-SHapley Additive exPlanation

کیفیت توضیحات شود، به‌ویژه برای مدل‌های پیچیده و غیرخطی. علاوه بر این، مقادیر شاپلی به نقطه مرجع یا پایه انتخاب شده حساس هستند و انتخاب پایه مناسب می‌تواند تأثیر زیادی بر تفسیر داشته باشد. در نهایت، در حالی که SHAP بینش‌های کلی را ارائه می‌دهد، تفسیر خروجی آن برای مدل‌های بسیار پیچیده همچنان می‌تواند برای کاربران غیرمتخصص به دلیل پیچیدگی توضیحات تولید شده چالش برانگیز باشد.

از زمان معرفی XAI، این حوزه توجه بسیاری از پژوهشگران را به خود جلب کرده است. مدل‌های توضیح‌دهنده مختلفی پیشنهاد شده‌اند و این توضیح‌دهنده‌ها با موفقیت به مسایل عملی مختلف اعمال شده‌اند. با این که حوزه‌های کاربردی مختلفی توسط XAI پوشش داده شده‌اند، تعداد کمی از کارها به‌طور هم‌زمان به XAI و مهندسی نرم‌افزار یا تحلیل نیازمندی‌های نرم‌افزاری پرداخته‌اند.

حبیبیه و همکاران بحثی درباره هم‌افزایی بین XAI و مهندسی نیازمندی‌ها ارائه کردند [۱۱]. آن‌ها چالش‌ها را بیان کرده و چارچوبی برای مقابله با آن‌ها پیشنهاد کردند. چالش‌های اصلی شامل نبود نقش واسطه، عدم وجود تعریف منسجم از توضیح‌پذیری، کمبود روش‌های متمرکز بر سودبران، و عدم وجود واژگان مشترک است. چارچوب پیشنهادی برای مقابله با این چالش‌ها شامل پنج مرحله است که با شناسایی سودبران آغاز شده و با طبقه‌بندی نیازمندی‌ها پایان می‌یابد.

چارت و همکاران در [۲۴] ذکر کرده‌اند که توضیح‌پذیری در حوزه نیازمندی‌های نرم‌افزار مورد مطالعه قرار نگرفته است و محدودیت‌هایی برای اعمال توضیح‌پذیری در این زمینه وجود دارد. آن‌ها راهکاری را جهت رفع موانع و محدودیت‌های موجود با تمرکز بر توضیح‌پذیری، ارائه یک مدل مفهومی، کاتالوگ دانش و یک مدل مرجع برای سیستم‌های توضیح‌پذیر ارائه داده‌اند. در پژوهش دیگری، بهره‌گیری از هوش مصنوعی قابل توضیح در مرحله

مهندسی نیازمندی‌ها مورد مطالعه قرار گرفته است [۲۵]. مگباله در این پژوهش بیان کرده است هوش مصنوعی قابل توضیح هنوز به‌طور گسترده در مرحله مهندسی نیازمندی‌ها استفاده نشده است، زیرا هیچ رویکرد استاندارد یا پیشنهاد نشده است. هیچ یک از مقالات بررسی شده توسط ایشان به‌طور واقعی از هوش مصنوعی و زیرتکنیک‌های مرتبط با آن برای کار با جنبه‌گرایی استفاده نکرده‌اند و همچنین توجه کمتری به ادبیات مربوط به استفاده از تکنیک‌های هوش مصنوعی در مرحله اعتبارسنجی نیازمندی‌ها شده است، در حالی که این مرحله یکی از مراحل بسیار مهم در کل فرآیند مهندسی نیازمندی‌ها محسوب می‌شود و استفاده از هوش مصنوعی بیشتر بر استخراج نیازمندی‌ها و مدیریت نیازمندی‌ها متمرکز بوده است. بای و همکاران از ابزار تحلیل احساسات، مدل‌سازی موضوعی ساختاری، ابزار تحلیل وب، یادگیری ماشین و تکنیک‌های هوش مصنوعی قابل توضیح برای تحلیل عمیق نیازهای کاربران در سه برنامه کاربردی ویدئو کنفرانس استفاده کرده‌اند. آن‌ها از هوش مصنوعی قابل توضیح برای اولویت‌بندی و رتبه‌بندی نیازمندی‌های نرم‌افزار استفاده کرده‌اند [۲۶].

کلمان و همکاران مروری بر هم‌راستایی XAI با فرآیند توسعه نرم‌افزار ارائه کردند [۱۲]. آن‌ها مراحل اصلی فرآیند توسعه نرم‌افزار مانند مهندسی نیازمندی‌ها، تحلیل و طراحی را در نظر گرفتند و روش‌ها و ابزارهای XAI را با این مراحل توسعه هم‌راستا کردند. هدف آن‌ها ارائه یک نقشه راه برای توسعه برنامه‌های هوش مصنوعی با استفاده از توضیح‌پذیری بود که چرخه عمر کلاسیک توسعه سیستم‌های نرم‌افزاری را در نظر می‌گرفت. آن‌ها معتقدند که چالش چنین روش‌هایی نحوه انتخاب تنظیمات بهینه، مانند تعداد نمونه‌های اولیه است و چالش اصلی این است که توضیحات برای انسان طراحی شده است، در نتیجه ارزیابی چنین توضیحاتی را تا حدی ذهنی می‌کند. دریافت‌کنندگان توضیحات می‌توانند با توجه به انتظارات، تخصص یادگیری ماشین و دانش زمینه‌ای بسیار متفاوت

را حل‌های پیچیده هوش مصنوعی فراهم کند. این مفهوم به دنبال حل چالش‌های مربوط به درک و اعتماد کاربران به مدل‌های هوش مصنوعی، به‌ویژه در سیستم‌های پیچیده، است. اهمیت هوش مصنوعی قابل توضیح در کاربردهایی که تصمیمات مدل تأثیر مستقیم بر زندگی انسان‌ها دارند، مانند حوزه‌های پزشکی، مالی و حقوقی، به‌طور چشمگیری افزایش یافته است. یک مدل قابل توضیح باید نه تنها برای کاربران انسانی قابل فهم باشد، بلکه باید اعتماد کاربران را نیز جلب کند. همچنین، توجیه‌پذیری نتایج مدل‌ها و ایجاد شفافیت در فرآیند تصمیم‌گیری از الزامات کلیدی هوش مصنوعی قابل توضیح هستند. شفافیت، به معنای فراهم کردن بینش در مورد چگونگی عملکرد مدل‌ها و دلایل اتخاذ یک تصمیم خاص، به کاربران این امکان را می‌دهد که مدل‌ها را بهتر درک کنند و در صورت لزوم عملکرد آن‌ها را بهبود بخشند. هوش مصنوعی قابل توضیح دو هدف را دنبال می‌کند [۱۴]:

ایجاد راه‌حل‌های دقیق و کارآمد: بهبود دقت و عملکرد مدل‌های هوش مصنوعی بدون قربانی کردن قابلیت توضیح‌پذیری.

ارایه توضیحات شفاف و واضح: تضمین این که نتایج پیش‌بینی شده برای کاربران قابل فهم باشند و بتوان به آن‌ها اعتماد کرد.

برای دستیابی به این اهداف روش‌های مختلف هوش مصنوعی قابل توضیح ارایه شده‌اند که به سه دسته قابل تقسیم هستند:

روش‌های تفسیرپذیری: تکنیک‌های تفسیر مدل که به دنبال درک عملکرد داخلی مدل‌های هوش مصنوعی و شناسایی ویژگی‌هایی هستند که در پیش‌بینی‌ها بیشترین اهمیت را دارند.

روش‌های تولید توضیح: این روش‌ها برای توضیح واضح پیش‌بینی مدل‌های هوش مصنوعی استفاده می‌شوند.

روش‌های تحلیل حساسیت: این روش‌ها نشان می‌دهند

باشند. بنابراین، یک توضیح یکسان می‌تواند برای یک کاربر راضی‌کننده باشد اما برای کاربر دیگر کاملاً طاقت‌فرسا یا غیرقابل درک باشد.

یک نظرسنجی کیفی برای بررسی درک کارشناسان از اهداف مدل‌های پیش‌بینی نقض و تکنیک‌های مدل‌محور مستقل که برای توصیف بصری مدل‌های پیش‌بینی نقض استفاده می‌شوند، در [۱۳] ارایه شده است. این کار نشان داد که تکنیک‌های مدل‌محور مستقل مانند LIME به کارشناسان کمک می‌کنند تا مدل‌های پیش‌بینی نقض و پیش‌بینی‌های آن‌ها را درک کنند.

با وجود استفاده موثر از الگوریتم‌های SHAP و LIME در مسایل پردازش زبان طبیعی این الگوریتم‌ها با چالش‌هایی در این زمینه روبه‌رو هستند. به دلیل ابعاد بالای داده‌های متنی، وابستگی‌های ناشی از ماهیت زبان مورد پردازش و پیچیدگی مدل‌های پردازش زبان طبیعی این چالش‌ها به‌وجود می‌آیند. هر دو الگوریتم در ایجاد تغییرات معنادار در متن مشکل دارند، زیرا تغییر کلمات می‌تواند به طور چشمگیری معنای جمله را تغییر دهد و درک روابط وابسته به زمینه دشوار است. تعداد زیاد ویژگی‌ها (کلمات) در مدل‌های پردازش زبان طبیعی بار محاسباتی را برای هر دو الگوریتم افزایش می‌دهد و آن‌ها را از نظر محاسباتی پرهزینه می‌کند. علاوه بر این، این الگوریتم‌ها ممکن است توضیحاتی برای مدل‌های پیچیده مانند ترنسفورمرها ایجاد کنند که درک آنها ساده نباشد. هر دو روش همچنین با حساسیت به حملات خصمانه، ناهماهنگی‌های پیش‌پردازش و مقیاس‌پذیری روبه‌رو هستند که می‌تواند باعث ایجاد اشکالاتی در کارایی آنها شود.

۳- هوش مصنوعی قابل توضیح

هوش مصنوعی قابل توضیح یک مفهوم نوظهور در حوزه هوش مصنوعی است که به طور خاص بر افزایش قابلیت توضیح‌پذیری در راه‌حل‌های هوش مصنوعی تمرکز دارد. یک مدل قابل توضیح باید برای کاربران انسانی قابل فهم، قابل اعتماد و توجیه‌پذیر باشد و شفافیت را در

که چگونه پیش‌بینی مدل تحت تأثیر تغییرات کوچک در داده‌های ورودی قرار می‌گیرد.

در سال‌های اخیر مدل‌های توضیحی مختلفی توسط پژوهشگران پیشنهاد شده است که SHAP و LIME از محبوب‌ترین آن‌ها هستند [۱۵]. به دلیل تعداد زیاد مدل‌های توضیح‌دهنده، انتخاب یک توضیح‌دهنده مناسب ممکن است یک کار چالش‌برانگیز باشد. ژیزوس و همکاران یک مدل آزمایشی پیشنهاد دادند که به کاربر در انتخاب مدل مناسب کمک می‌کند [۱۵].

۳-۱- توضیح الگوریتم LIME

در این مقاله از روش LIME استفاده شده است که یکی از پرکاربردترین و محبوب‌ترین تکنیک‌های توضیح‌پذیری مدل‌های یادگیری ماشین در سال‌های اخیر است. این روش با هدف ارائه توضیحات محلی برای پیش‌بینی‌های مدل‌های یادگیری ماشین طراحی شده و قابلیت اعمال بر روی انواع مختلف مدل‌ها را دارد. از آنجایی که LIME مستقل از مدل است، می‌تواند برای توضیح پیش‌بینی‌های الگوریتم‌های گوناگون از جمله طبقه‌بندی و رگرسیون مورد استفاده قرار گیرد. روش کار LIME به این صورت است که ابتدا داده‌های ورودی را تغییر داده و تأثیر این تغییرات را بر پیش‌بینی‌های مدل تحلیل می‌کند. برای یک نمونه خاص، LIME مجموعه‌ای از نمونه‌های تصادفی را در اطراف داده ورودی اصلی تولید می‌کند و مدل اصلی را برای پیش‌بینی این نمونه‌ها مورد استفاده قرار می‌دهد. سپس، با استفاده از پیش‌بینی‌های مدل برای این نمونه‌ها، یک مدل ساده و قابل فهم (مانند یک مدل خطی) ایجاد می‌کند که رفتار مدل اصلی را در اطراف داده ورودی توضیح می‌دهد. این مدل ساده به صورت محلی عمل می‌کند، به این معنا که تنها پیش‌بینی مربوط به نمونه ورودی اصلی را توضیح می‌دهد و نه کل مجموعه داده یا ساختار کلی مدل.

یکی از ویژگی‌های برجسته LIME، توانایی آن در ارائه توضیحات محلی دقیق است. این ویژگی به خصوص برای

مدل‌های پیچیده و جعبه سیاه، مانند شبکه‌های عصبی عمیق یا مدل‌های تقویتی، بسیار مفید است. چنین مدل‌هایی به دلیل ساختار داخلی پیچیده، معمولاً غیرشفاف هستند و کاربران نمی‌توانند به سادگی علت پیش‌بینی‌های آن‌ها را درک کنند. LIME این چالش را با تولید توضیحات قابل فهم برای پیش‌بینی‌های خاص برطرف می‌کند و اعتماد کاربران را به نتایج مدل افزایش می‌دهد. علاوه بر این، LIME از نظر پیاده‌سازی نیز انعطاف‌پذیری بالایی دارد. این روش می‌تواند بر روی طیف وسیعی از مدل‌های یادگیری ماشین اعمال شود و برای کاربران امکان تحلیل و درک نتایج پیش‌بینی را فراهم می‌کند. ابزارهایی مانند LIME به پژوهشگران و توسعه‌دهندگان این امکان را می‌دهند که علاوه بر بهبود دقت مدل‌های خود، شفافیت و توضیح‌پذیری آن‌ها را نیز ارتقا دهند. این امر به ویژه در زمینه‌هایی مانند پزشکی، مالی و امنیت سایبری که اعتماد به مدل‌های هوش مصنوعی ضروری است، اهمیت بسیاری دارد. به طور کلی، LIME به عنوان یک ابزار قدرتمند و مستقل از مدل، روشی قابل اعتماد برای درک پیش‌بینی‌های مدل‌های یادگیری ماشین ارائه می‌دهد و نقشی کلیدی در ترویج استفاده مسئولانه و شفاف از هوش مصنوعی ایفا می‌کند.

الگوریتم LIME می‌تواند به صورت زیر مدل شود. برای سادگی، جزئیات در اینجا نادیده گرفته شده است. فرض کنید که یک نمونه ورودی x داریم و می‌خواهیم پیش‌بینی ارائه شده توسط طبقه‌بند C را توضیح دهیم. پیش‌بینی انجام شده توسط مدل $y = C(x)$ است. الگوریتم LIME نمونه‌های n تایی $x' = \{x'_1, x'_2, \dots, x'_n\}$ را با تغییر دادن نمونه اصلی تولید می‌کند.

خروجی $y'_i = C(x'_i)$ برای نمونه x'_i توسط الگوریتم جعبه سیاه تولید و نمونه برچسب‌گذاری می‌شود. این نمونه‌ها به عنوان یک مجموعه داده برای آموزش مدل قابل توضیح استفاده می‌شوند. سپس، الگوریتم LIME نمونه‌ها را بر اساس نزدیکی به نمونه اصلی x وزن‌دهی می‌کند. به این ترتیب، نمونه‌های نزدیک‌تر به x وزن بیشتری دارند.

نرم‌افزار مشخص شود. بدین‌منظور یک روش که شامل یک طبقه‌بند یادگیری ماشین و توضیح‌دهنده LIME است، برای تحلیل نیازمندی‌های نرم‌افزاری استفاده شده است. ساختار روش پیشنهادی در شکل ۱ ارایه شده است. این شکل نشان می‌دهد که ساختار دارای دو بخش اصلی است. بخش اول به طبقه‌بند اختصاص دارد و بخش دوم متعلق به توضیح‌دهنده است. برای استفاده از مدل XAI همراه با یک طبقه‌بند، ابتدا طبقه‌بند طراحی می‌شود. سپس، مدل XAI می‌تواند برای توضیح نمونه‌های مورد نیاز استفاده شود. این روش یک مجموعه نیازمندی به عنوان ورودی دریافت می‌کند. هر رکورد در مجموعه داده دارای دو ستون است که ستون اول شامل یک جمله یا پاراگراف از یک نیازمندی و ستون دوم برچسب کلاس است. نیازمندی‌ها به زبان طبیعی هستند. بنابراین، داده‌ها باید پیش‌پردازش شوند. مرحله پیش‌پردازش شامل برخی مراحل شناخته شده است که به طور عمومی در برنامه‌های NLP^{۲۶} استفاده می‌شوند. وظایف مهم شامل توکن‌سازی، حذف اعداد و کلمات بی‌فایده و ریشه‌یابی است.

علاوه بر این مراحل، یک مرحله جدید اضافه می‌شود که به کارشناس حوزه این امکان را می‌دهد که پردازش خاص برنامه‌کاربردی را روی خروجی مرحله ریشه‌یابی، اعمال کند. دلیل اضافه کردن چنین مرحله‌ای در بخش بعدی توضیح داده می‌شود. به عنوان مثال، کلمات کیفی مانند «باید^{۲۷}»، «می‌بایست^{۲۸}» و «لازم است^{۲۹}» معمولاً برای تعیین اولویت نیازمندی‌ها استفاده می‌شوند و به نظر می‌رسد که ممکن است در طبقه‌بندی نیازمندی‌ها نادیده گرفته شوند. بنابراین، این کلمات می‌توانند از نیازمندی‌ها حذف شوند. همچنین، نام نرم‌افزار ممکن است اطلاعات اضافی برای طبقه‌بندی نیازمندی‌ها ارایه نکند.

خروجی مرحله پیش‌پردازش یک مجموعه بردار است که در آن هر بردار شامل ریشه‌های کلمات مهم در یک

وزن هر نمونه x'_i با استفاده از هسته گاوسی^{۲۴} به صورت زیر است:

$$w(x'_i) = \exp\left(-\frac{d(x, x'_i)^2}{2\sigma^2}\right)$$

که در آن $d(x, x'_i)$ فاصله اقلیدسی بین x و x' است و σ میزان گسترش هسته را کنترل می‌کند. پس از آن، میانگین وزنی پیش‌بینی‌های y'_i محاسبه می‌شود. سپس، پیش‌بینی‌های وزنی برای تطبیق یک مدل قابل توضیح g استفاده می‌شوند. LIME از مدل‌های ذاتی قابل توضیح مانند رگرسیون خطی، درخت تصمیم‌گیری و مدل‌های ذهنی قاعده‌محور^{۲۵} برای توضیح نتایج استفاده می‌کند. فرض کنید که از رگرسیون خطی استفاده می‌شود. مدل قابل توضیح g برای نمونه‌های تغییر یافته x'_i به صورت زیر برازش می‌شود:

$$g(x'_i) = \theta_0 + \sum_{i=1}^k \theta_i x'_i$$

که در آن θ_i ضریب ویژگی i ام در رگرسیون خطی است. در نهایت، مدل تولید شده g برای توضیح پیش‌بینی انجام‌شده توسط مدل C برای ورودی اصلی x استفاده می‌شود. توضیح الگوریتم LIME مجموعه‌ای از امتیازهای اهمیت ویژگی است. هر وزن نشان می‌دهد که یک ویژگی در x چقدر به پیش‌بینی کمک می‌کند. برای رگرسیون خطی، ضرایب θ به عنوان وزن‌هایی در نظر گرفته می‌شوند که توسط الگوریتم LIME گزارش می‌شوند.

۴- روش پیشنهادی

با توجه به کاربردهای موفق XAI در زمینه‌های مختلف، در این مقاله تلاش شده تا توضیح‌دهنده LIME به مشکل طبقه‌بندی نیازمندی‌ها اعمال شود. هدف این مطالعه این است که نحوه استفاده از XAI در مرحله توسعه نرم‌افزار برای اعضای تیم توسعه توضیح داده شود، همچنین نتایج اعمال XAI در مرحله مهندسی نیازمندی‌ها برای تیم توسعه

26-Natural Language Processing
27-Should
28-Should
29-Must

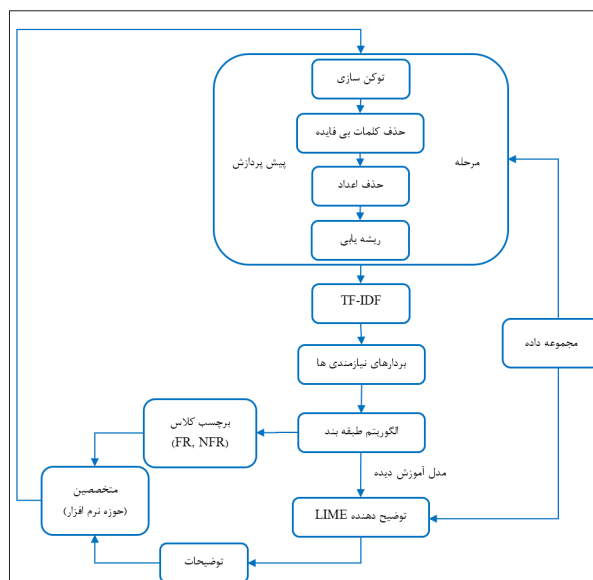
24-Gaussian kernel
25-Rule-Based Heuristic (RBH)

توضیح و احتمال پیش‌بینی کلاس نیازمندی تولید می‌کند. در نهایت، خروجی مدل XAI همراه با کلاس پیش‌بینی شده توسط طبقه‌بند می‌تواند توسط کارشناسان دامنه مانند تحلیل‌گران، مشخص‌کنندگان نیازمندی یا معماران، به‌عنوان بخشی از کارهایشان استفاده شود. سپس کارشناس دامنه برچسب‌های کلاس پیش‌بینی شده را همراه با توضیحات تولید شده توسط الگوریتم LIME تحلیل می‌کند. نتیجه این تحلیل می‌تواند به عنوان بازخورد برای مرحله پیش‌پردازش استفاده شود.

متد پیشنهادی به کارشناسان دامنه در فرآیند توسعه نرم‌افزار کمک می‌کند تا ویژگی‌های مؤثر در پیش‌بینی کلاس یک نیازمندی را مشاهده کنند. بنابراین، کارشناسان دامنه می‌توانند به خروجی مدل‌های یادگیری ماشین در تحلیل نیازمندی‌ها اعتماد کنند.

۵- مطالعه تجربی

برای بررسی قابلیت کاربرد روش پیشنهادی که شامل ماژول XAI برای تحلیل نیازمندی‌های نرم‌افزاری است، از مجموعه داده PROMISE استفاده شده است. این مجموعه داده شامل نیازمندی‌های کارکردی و غیرکارکردی است. نیازمندی‌های کارکردی با برچسب (F) و نیازمندی‌های غیرکارکردی به یازده دسته تقسیم می‌شوند که برچسب‌های زیر را دارند: دسترس‌پذیری (A)^{۳۲}، تحمل خطا (FT)^{۳۳}، قانونی (L)^{۳۴}، ظاهر و احساس (LF)^{۳۵}، قابلیت نگهداری (MN)^{۳۶}، عملیاتی (O)^{۳۷}، عملکرد (PE)^{۳۸}، قابلیت حمل (PO)^{۳۹}، مقیاس‌پذیری (SC)^{۴۰}، امنیت (SE)^{۴۱} و کاربرپذیری (US)^{۴۲}. در این کار، طبقه‌بندی نیازمندی‌ها به عنوان یک مسئله طبقه‌بندی دوتایی در نظر گرفته شده است. بنابراین، تمام زیرمجموعه‌های



شکل ۱. ساختار روش پیشنهادی برای تحلیل نیازمندی‌ها با استفاده از هوش مصنوعی قابل توضیح و طبقه‌بند SVM

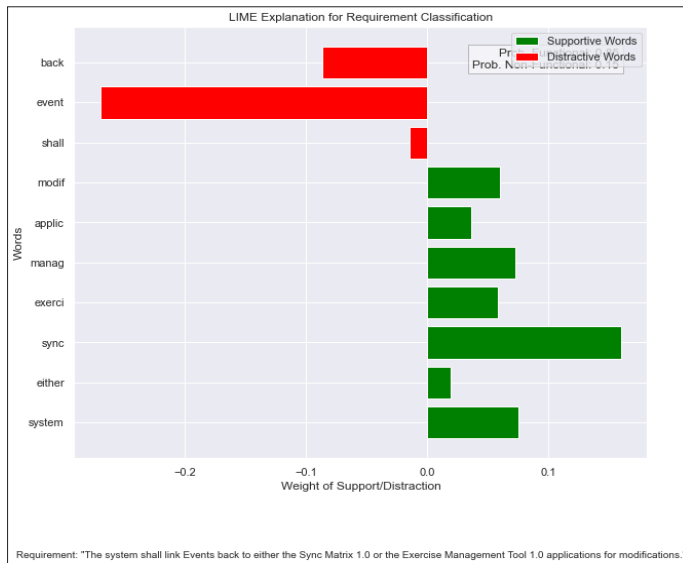
نیازمندی همراه با برچسب کلاس آن است. این بردارها به عنوان ورودی برای مرحله بعدی طبقه‌بندی استفاده می‌شوند.

بیشتر طبقه‌بندها با داده‌های متنی به عنوان ورودی کار نمی‌کنند. بنابراین، داده‌های متنی معمولاً به بردارهای عددی با استفاده از روش‌های مختلف تبدیل می‌شوند. در اینجا، بر اساس تحقیقات قبلی [۱۶]، [۱۷] از روش TF-IDF برای تبدیل نیازمندی‌ها به بردارهای عددی استفاده کردیم. پس از اعمال روش TF-IDF، از طبقه‌بند برای طبقه‌بندی نیازمندی‌ها استفاده می‌شود. با آموزش طبقه‌بند، مدل تولید شده می‌تواند توسط مدل XAI برای توضیح این که چرا یک نیازمندی به عنوان کارکردی یا غیرکارکردی طبقه‌بندی می‌شود، استفاده شود.

در اینجا از مدل توضیح‌دهنده LIME استفاده شده است. نیازمندی که می‌خواهیم آن را توضیح دهیم و طبقه‌بندی تولید شده را به عنوان ورودی به LIME داده و توضیح مربوط به آن تولید می‌شود. برای داده‌های متنی مانند نیازمندی‌های کاربران، مدل LIME مجموعه‌ای از کلمات حمایتی^{۳۰} و مختل‌کننده^{۳۱} را همراه با وزن‌های آن‌ها در

32-Availability
33-Fault Tolerance
34-Legal
35-Look & Feel
36-Maintainability
37-Operational
38-Performance
39-Portability
40-Scalability
41-Security
42-Usability

30-Supportive
31-Distractive



شکل ۲. خروجی توضیح‌دهنده LIME بر روی یک نیازمندی انتخاب‌شده به صورت تصادفی در مجموعه داده PROMISE

The system shall link Events back to either the Sync Matrix 1.0 or the Exercise Management Tool 1.0 applications for modifications.

شکل ۳. نیازمندی مربوط به توضیح شکل ۲

می‌دهد. کلمات حمایتی و مختل‌کننده به ترتیب با رنگ‌های سبز و قرمز نمایش داده می‌شوند. محور افقی وزن حمایت یا اختلال کلمات را نشان می‌دهد. همچنین، LIME احتمال در نظر گرفتن یک نیازمندی به عنوان کارکردی یا غیرکارکردی را ارایه می‌دهد.

نیازمندی مربوطه در شکل ۳ ارایه شده است. متن نیازمندی پس از پیش‌پردازش در اینجا نشان داده شده است. این یک نیازمندی غیرکارکردی است و احتمال غیرکارکردی بودن آن توسط LIME بیشتر از ۰/۹ است. این شکل نشان می‌دهد که برخی از کلمات مانند «لینک»، «ابزار» و «ماتریس» برای طبقه‌بندی این نیازمندی در نظر گرفته نشده‌اند.

۵-۲- قابلیت توضیح مدل

این آزمایش به منظور تحلیل توضیح پیش‌بینی‌های

جدول ۱. عملکرد ماشین بردار پشتیبان در مجموعه آزمایشی

مقدار	معیار
۸۹٪	F1-score
۹۰٪	دقت (Accuracy)
۸۸٪	دقت (Precision)
۹۱٪	یادآوری (Recall)

نیازمندی‌های غیرکارکردی به عنوان غیرکارکردی در نظر گرفته می‌شوند. به طور کلی، ۴۰ درصد نیازمندی‌ها کارکردی و ۶۰ درصد غیرکارکردی هستند. مجموعه داده به یک مجموعه آموزش و یک مجموعه تست تقسیم شده است که ۸۰ درصد برای آموزش و باقیمانده برای آزمایش استفاده می‌شود. نمونه‌ها در مجموعه‌های آموزشی و آزمایشی به صورت تصادفی انتخاب شده‌اند. این روش با استفاده از زبان برنامه‌نویسی پایتون پیاده‌سازی شد. از بسته‌های Seaborn و NumPy، NLTK، Keras، Lime، Pandas استفاده شده است.

الگوریتم LIME می‌تواند با هر روش طبقه‌بندی کار کند. بنابراین، هر الگوریتم یادگیری ماشین یا یادگیری عمیق می‌تواند به عنوان ماژول طبقه‌بندی در شکل ۱ قرار گیرد. در اینجا، از SVM به عنوان الگوریتم طبقه‌بندی استفاده شده است. عملکرد طبقه‌بند در داده‌های آزمایشی در جدول ۱ ارایه شده است. نتایج میانگین پس از ۳۰ بار اجرای این روش در این جدول گزارش شده است. نتایج نشان می‌دهد که ماشین‌بردار پشتیبان (Support Vector Machine) عملکرد قابل قبولی در داده‌های آزمایشی دارد.

۵-۱- خروجی توضیح‌دهنده

قبل از تحلیل قابلیت کاربرد الگوریتم LIME برای تحلیل نیازمندی‌ها، خروجی LIME در اینجا ارایه شده است. راه‌حلی که توسط LIME تولید می‌شود شامل مجموعه‌ای از کلمات است که تصمیم را پشتیبانی یا مختل می‌کند. یک مثال از خروجی LIME در شکل ۲ آورده شده است. محور عمودی لیست کلمات حمایتی و کلمات مختل‌کننده را نشان

جدول ۲. تعریف کلمات حمایتی و مختل کننده

معیار	تعریف
کلمات حمایتی (S)	مجموعه‌ای از تمام کلمات با اثر مثبت در طبقه‌بندی یک نیازمندی به‌عنوان نیازمندی غیرکارکردی.
کلمات مختل کننده (D)	مجموعه‌ای از تمام کلمات با اثر منفی در طبقه‌بندی یک نیازمندی به‌عنوان نیازمندی غیرکارکردی.

این کلمات را نشان می‌دهد. درصد وقوع به‌صورت زیر تعریف می‌شود:

$$pc_i = \frac{n_i}{N} \times 100$$

در این فرمول n_i تعداد وقوع کلمه i است و N تعداد کل کلمات می‌باشد. الگوریتم LIME ده‌ها کلمه حمایتی و مختل‌کننده را برای توضیح تمام نمونه‌های آزمایشی گزارش کرده است. نمایش تمام آن‌ها ممکن نیست. به همین دلیل، تنها ۳۰ کلمه برتر با بالاترین درصد وقوع در اینجا نمایش داده شده است. ریشه‌های کلمات در این شکل‌ها ارایه شده است.

همان‌طور که در شکل ۴ نشان داده شده است، «محصول»، «کاربر»، «استفاده»، «ثانیه» و «دقیقه» پنج کلمه حمایتی با بالاترین تکرار هستند. اگرچه این کلمات در پیش‌بینی یک نیازمندی به‌عنوان غیرکارکردی بسیار تکرار شده‌اند، اما در میان کلمات مهم‌ترین نیستند. معمولاً یک کارشناس انسانی نمی‌تواند تنها با استفاده از این کلمات، یک نیازمندی غیرکارکردی را تعیین کند و برای طبقه‌بندی یک نیازمندی به‌عنوان غیرکارکردی، به کلمات دقیق‌تری نیاز دارد. احتمالاً در مدل یادگیری ماشین نیز همین وضعیت وجود دارد. به‌نظر می‌رسد که کلمات خاص غیرکارکردی توسط روش‌های یادگیری ماشین برای طبقه‌بندی استفاده می‌شود.

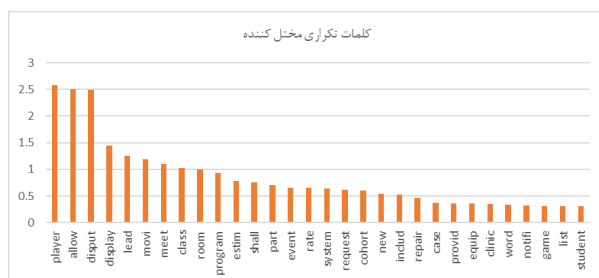
همچنین کلمات «بازیکن»، «اجازه»، «اختلاف»، «هدایت» و «حرکت» پنج کلمه مختل‌کننده با بالاترین تکرار هستند که در شکل ۵ نشان داده شده‌اند. برخی از کلمات با تکرار بالا مانند «کلمه»، هم در کلمات حمایتی و هم در کلمات مختل‌کننده مشاهده می‌شود.

برای تفکیک و تحلیل بهتر کلمات، تمام کلمات استفاده

نیازمندی‌های غیرکارکردی در مجموعه داده با استفاده از الگوریتم LIME طراحی شده است. به‌نظر می‌رسد خروجی توضیح‌دهنده که در شکل ۲ نشان داده شده است، بینش ارزشمندی برای ما فراهم می‌کند تا قابلیت کاربرد الگوریتم LIME را تحلیل کنیم. به‌عنوان مثال، کلمه «استفاده» که یک کلمه حمایتی قوی است، برای در نظر گرفتن این نیازمندی به‌عنوان غیرکارکردی استفاده می‌شود. نتایجی که در اینجا و در بخش بعدی گزارش شده است، تنها مجموعه‌های آزمایشی را در نظر می‌گیرد.

برای تحلیل قابلیت کاربرد LIME برای تحلیل نیازمندی‌ها، از دو مجموعه کلمه که توسط الگوریتم LIME گزارش شده است، برای تعیین نیازمندی‌ها به‌عنوان غیرکارکردی استفاده می‌شود. در اینجا، ما فقط توضیح یک نیازمندی به‌عنوان غیرکارکردی را در نظر می‌گیریم. سناریوی مشابهی می‌تواند برای نیازمندی‌های کارکردی استفاده شود. این کلمات به‌عنوان کلمات حمایتی و مختل‌کننده شناخته می‌شوند. تعاریف کلمات حمایتی و مختل‌کننده در جدول ۲ ارایه شده است.

سه آزمایش برای نشان دادن قابلیت توضیح الگوریتم LIME برای نیازمندی‌های نرم‌افزاری انجام شده است. در آزمایش اول، تمام کلمات حمایتی و مختل‌کننده در نیازمندی‌ها در نظر گرفته می‌شوند. برای این منظور، داده‌های آزمایشی به‌صورت نمونه به الگوریتم LIME داده می‌شود و کلمات حمایتی و مختل‌کننده جمع‌آوری می‌شوند. تمام این کلمات تجمیع شده و رخداد آنها محاسبه می‌شود. شکل‌های ۴ و ۵ کلمات حمایتی و مختل‌کننده را برای توضیح طبقه‌بندی یک نیازمندی به‌عنوان غیرکارکردی نشان می‌دهند. محور افقی نشان‌دهنده کلماتی است که در غیرکارکردی‌ها ظاهر شده‌اند و محور عمودی درصد وقوع



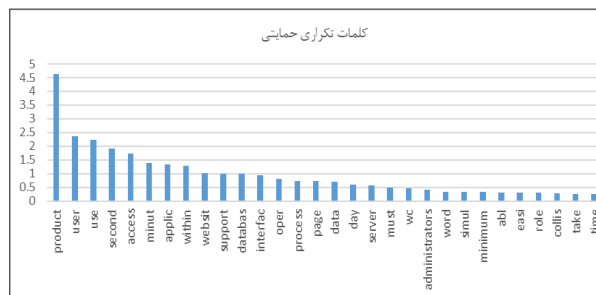
شکل ۵. کلمات تکراری مختل کننده در طبقه‌بندی یک نیازمندی به‌عنوان غیرکارکردی در مجموعه داده قبل از حذف کلمات کیفی

و یک کارشناس انسانی نمی‌تواند تنها به کلمات مشترک برای طبقه‌بندی یک نیازمندی نرم‌افزاری تکیه کند.

مجموعه A شامل کلمات مختل کننده است، کلماتی مانند «نمایش»، «تغییر»، «کلاس»، «حذف» و غیره. معمولاً، این کلمات برای نشان دادن مسئولیت یا خدماتی که توسط نرم‌افزار ارائه می‌شود، استفاده می‌شوند. این کلمات می‌توانند به‌عنوان نامزدهایی برای تعیین یک نیازمندی به‌عنوان غیرکارکردی در نظر گرفته شوند. بنابراین، به‌نظر می‌رسد که توضیح‌دهنده این کلمات را به‌عنوان کلمات مختل کننده شناسایی کرده است. باید توجه داشت که به دلیل خطاهای طبقه‌بندی، دسته‌بندی کلمات در این مجموعه‌ها ممکن است شامل خطا باشد.

مجموعه B که شامل کلمات حمایتی است، شامل کلماتی مانند «ثانیه»، «دقیقه»، «نصب»، «مجوز»، «تهدید» و غیره می‌شود که برای طبقه‌بندی نیازمندی‌ها به‌عنوان غیرکارکردی معقول هستند. این کلمات جالب توجه هستند و توسط طبقه‌بند برای طبقه‌بندی یک نیازمندی به‌عنوان غیرکارکردی استفاده می‌شوند. به‌عنوان مثال، «ثانیه» یک واحد زمان است و کاربران از این کلمه برای بیان معیارهای نگهداری یا عملکرد مانند «زمان پاسخ» یا «زمان پاسخ اصلاح» استفاده می‌کنند. کلمه «نصب» نشان می‌دهد که برای عملیات یا نگهداری نرم‌افزار استفاده می‌شود و انتظار داریم با چنین معیاری مواجه شویم.

به‌عنوان یک مثال، دو نیازمندی غیرکارکردی که شامل کلمات «ثانیه» و «نصب» در مجموعه داده هستند در شکل ۶ نشان داده شده‌اند. معمولاً در دنیای واقعی، یک کارشناس



شکل ۴. لیست تکراری ترین کلمات حمایتی در طبقه‌بندی یک نیازمندی به‌عنوان غیرکارکردی در مجموعه داده قبل از حذف کلمات کیفی

شده توسط توضیح‌دهنده برای توضیح طبقه‌بند SVM را به سه مجموعه تقسیم کردیم. این مجموعه‌ها A، B و C نامیده می‌شوند و به شرح زیر تعریف شده‌اند:

$A = D - C$ مجموعه کلمات مختل کننده منهای کلمات مشترک.

$B = S - C$ مجموعه کلمات حمایتی منهای کلمات مشترک
 $C = S \cap D$ مجموعه کلمات مشترک که تقاطع مجموعه‌های S و D است.

بعضی از اعضای هر مجموعه عبارتند از:

$A = \{player, allow, dispute., diplay, meet, ... \}$

$B = \{support, second, minute., server, ... \}$

$C = \{word, shall, system, user, ... \}$

مجموعه کلمات حمایتی (S) و مجموعه کلمات مختل کننده (D) در جدول ۲ تعریف شده‌اند.

این دسته‌بندی کلمات نشان می‌دهد که توضیح‌دهنده به‌طور معقولی از آن‌ها برای پیش‌بینی یک نیازمندی به‌عنوان کارکردی یا غیرکارکردی استفاده می‌کند. مجموعه‌های A و B از مجموعه C مهم‌تر هستند. با این حال، کلمات مشترک نیاز به توجه بیشتری دارند.

مجموعه C از کلمات مشترک، کلماتی مانند «باید»، «بایستی»، «کاربر» تشکیل شده است. اگرچه این کلمات به‌طور مکرر تکرار شده‌اند، به‌نظر می‌رسد که چنین کلماتی نمی‌توانند به‌طور قابل توجهی توسط یک طبقه‌بند برای طبقه‌بندی یک نیازمندی به‌عنوان کارکردی یا غیرکارکردی استفاده شوند. همین وضعیت در دنیای واقعی وجود دارد

Req 1: When purchasing a streaming movie or pre-paid card via credit card the processing time should have a maximum response time of 15 seconds

Req 2: The product will function along-side server software on any operating system where the Java runtime can be installed

.....

نیازمندی یک: هنگام خرید یک فیلم استریم شده یا کارت پیش پرداخت از طریق کارت اعتباری، زمان پردازش باید حداکثر زمان پاسخ ۱۵ ثانیه باشد.

نیازمندی دو: این محصول در کنار نرم افزار کارساز در هر سیستم عاملی که جاوا قابل نصب باشد، کار خواهد کرد.

شکل ۶. نمونه نیازمندی هایی که به درستی توسط طبقه بند به عنوان نیازمندی غیرکارکردی طبقه بندی شده اند.

جدول ۳. زیرمجموعه ای از کلمات استفاده شده توسط روش پیشنهادی برای پیش بینی نیازمندی ها به عنوان نیازمندی های غیرکارکردی

کلمه	کلاس نیازمندی غیرکارکردی
موجود، شکست	در دسترس بودن ^۱
عملیات	تحمل خطا ^۲
قوانین	قانونی ^۳
ناوبری، حل	ظاهر و احساس ^۴
بودجه	قابلیت نگهداری ^۵
نصب، مجوز	عملیاتی ^۶
ثانیه، حداکثر، بار	عملکرد ^۷
عملیاتی	قابل حمل ^۸
پشتیبانی، هم زمان، افزایش	مقیاس پذیری ^۹
ورود به سیستم، تهدید	امنیت ^{۱۰}
درک، صفحه، راحتی	قابلیت استفاده ^{۱۱}

1-Availability

2-Fault Tolerance

3-Legal

4-Look & Feel

5-Maintainability

6-Operational

7-Performance

8-Portability

9-Scalability

10-Security

11-Usability

انسانی از این کلمات کلیدی برای دسته بندی یک نیازمندی به عنوان نیازمندی غیرکارکردی استفاده می کند. به نظر می رسد که مدل این کلمات را به عنوان بخشی از فرآیند تصمیم گیری خود برای پیش بینی یک نیازمندی به عنوان نیازمندی غیرکارکردی در نظر می گیرد. در مجموعه داده، بیشتر وقوع های این کلمه متعلق به نیازمندی های غیرکارکردی هستند.

فهرست کلمات حمایتی تولید شده توسط الگوریتم LIME توسط کارشناسان انسانی مورد تجزیه و تحلیل قرار گرفت. کارشناسان قصد داشتند مجموعه ای از کلمات را استخراج کنند که معمولاً در دنیای واقعی برای دسته بندی یک نیازمندی به عنوان نیازمندی غیرکارکردی استفاده می شوند. جدول ۳ بخشی از این کلمات را که توسط مدل پیشنهادی برای پیش بینی یک نیازمندی به عنوان نیازمندی غیرکارکردی استفاده می شود، نشان می دهد.

برخی از کلمات در بین بیش از یک زیرکلاس از نیازمندی های غیرکارکردی مشترک هستند. با این حال، هر کلمه تنها در یک زیرکلاس از نیازمندی غیرکارکردی نشان داده می شود که در آن بالاترین احتمال وقوع را دارد. این آزمایش نشان می دهد که طبقه بند تقریباً مشابه یک کارشناس انسانی در تعیین یک نیازمندی به عنوان نیازمندی غیرکارکردی عمل می کند. در مقایسه با کلمات عمومی مانند «محصول» یا «کاربر»، این کلمات خاص دارای فراوانی کمتری هستند. اما، آن ها به طور موفقیت آمیزی توسط مدل برای پیش بینی یک نیازمندی به عنوان نیازمندی غیرکارکردی استفاده می شوند.

۵-۳- تحلیل قابلیت کاربرد

با در نظر گرفتن نتایج به دست آمده از اعمال الگوریتم LIME روی مجموعه داده، می توان مشاهده کرد که طبقه بند تقریباً از ویژگی های مرتبط (مانند کلمات) برای طبقه بندی نیازمندی استفاده می کند. با این حال، برخی از کلمات عمومی مانند «باید»، «کاربر»، «سیستم» و «محصولات» به عنوان

روی نیازمندی‌های متنی اعمال می‌شوند، دارد. مهم‌ترین کلماتی که توسط طبقه‌بندها در دسته‌بندی نیازمندی‌ها استفاده می‌شوند، شناسایی شدند. بررسی و تحلیل این کلمات توسط کارشناسان حوزه (مانند تحلیلگران، معماران و مشخص‌کنندگان نیاز) نشان داد که کلمات پیشنهادی معنادار هستند. در این کار، توضیح روش‌های یادگیری ماشین برای طبقه‌بندی نیازمندی مورد توجه قرار گرفت. اخیراً، روش‌های مختلف یادگیری عمیق برای طبقه‌بندی نیازمندی‌ها استفاده شده است. توضیح این روش‌ها در حوزه‌های مهندسی نیازمندی ممکن است نتایج جالبی را به دلیل عملکرد خوب این روش‌ها در طبقه‌بندی نیازمندی‌ها نشان دهد.

در سال‌های اخیر انواع مختلفی از توضیح‌دهنده‌ها توسط محققان طراحی شده‌اند که می‌توانند در این حوزه استفاده شوند. در این کار، فقط قابلیت کاربرد LIME به عنوان یک توضیح‌دهنده مورد توجه قرار گرفت. انجام یک مطالعه مقایسه‌ای بین این توضیح‌دهنده‌ها توصیه می‌شود. همچنین، توصیه می‌شود که XAI به عنوان ابزاری برای انتخاب ویژگی‌ها مورد استفاده قرار گیرد. همچنین یکی از چالش‌های الگوریتم LIME در زمینه پردازش نیازمندی‌های نرم‌افزار این است که توضیحات محلی را با نمونه‌برداری از یک همسایگی اطراف نمونه ورودی ایجاد می‌کند. لذا ایجاد یک همسایگی معنادار و واقع‌گرایانه برای یک نیازمندی نرم‌افزاری می‌تواند تا حدودی دشوار باشد، به‌ویژه برای نیازمندی‌های خاصی که ممکن است تفسیرهای مختلفی داشته باشند. برای کاهش اثر این محدودیت پیشنهاد می‌شود از تکنیک‌های شباهت معنایی برای ایجاد تغییراتی استفاده شود که از نظر معنایی به نیازمندی اصلی نزدیک‌تر باشند. با ایجاد تغییراتی که با زمینه معنایی مطابقت داشته باشند، می‌توان اطمینان حاصل کرد که همسایگی مرتبط‌تر و سازگارتر با ماهیت واقعی نیازمندی‌های نرم‌افزاری است.

هم حمایتی و هم مختل‌کننده در طبقه‌بندی یک نیازمندی به عنوان غیرکارکردی عمل می‌کنند. به نظر می‌رسد که این نوع کلمات در ترکیب با سایر کلمات حمایتی، یک نیازمندی را به عنوان غیرکارکردی طبقه‌بندی می‌کنند. همچنین، نتایج نشان‌دهنده برخی افعالی هستند که معمولاً برای مشخص کردن قابلیت (یا وظیفه) یک سیستم استفاده می‌شوند. این افعال به مجموعه مختل‌کننده تعلق دارند و به نظر منطقی می‌رسند.

به طور کلی، آزمایش‌ها نشان می‌دهند که XAI برای تحلیل نیازمندی و به‌ویژه مشکل طبقه‌بندی نیازمندی، کاربردی است. یک روش جدید برای اعمال XAI روی مشکل طبقه‌بندی نیازمندی ارائه شد. با این حال، در این کار فقط تحلیل پایه‌ای در نظر گرفته شده است و امکان انجام تحلیل‌های دقیق‌تر و پیچیده‌تر وجود دارد. اعضای تیم که مسئولیت‌هایی در مهندسی نیازمندی دارند، مانند مشخص‌کنندگان نیاز، تحلیلگران و معماران می‌توانند از توضیح‌دهنده‌های XAI به‌طور هم‌زمان با طبقه‌بندها استفاده کنند.

XAI بینش‌هایی برای اعضای تیم درباره پیش‌بینی‌های تولیدشده توسط مدل‌های یادگیری ماشین فراهم می‌کند. خروجی‌های XAI به آن‌ها کمک می‌کند تا بدانند چگونه مدل‌های یادگیری ماشین راه‌حل‌ها را تولید می‌کنند. در مورد طبقه‌بندی نیازمندی، XAI به اعضای تیم کمک می‌کند تا کلمات حمایتی، مختل‌کننده و عمومی را در طبقه‌بندی یک نیازمندی به عنوان کارکردی، نیازمندی غیرکارکردی یا یک زیرکلاس از نیازمندی غیرکارکردی مشاهده کنند.

۶- نتیجه‌گیری

در این مقاله، قابلیت کاربرد XAI در توضیح رفتار الگوریتم‌های طبقه‌بندی نیازمندی مورد مطالعه قرار گرفت. الگوریتم LIME نتایج تولیدشده توسط طبقه‌بند را روی مجموعه داده PROMISE تحلیل کرد. این مطالعه نشان داد که XAI نقش مثبتی در درک بهتر رفتار طبقه‌بندها زمانی که

۷- مراجع

- [14] Arrieta, A. B., Díaz-Rodríguez, N., Del Ser, J., Bennetot, A., Tabik, S., Barbado, A., ... & Herrera, F. (2020). Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI. *Information fusion*, 58, 82-115.
- [15] Jesus, S., Belém, C., Balayan, V., Bento, J., Saleiro, P., Bizarro, P., & Gama, J. (2021, March). How can I choose an explainer? An application-grounded evaluation of post-hoc explanations. In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency* (pp. 805-815).
- [16] Kim, D., Seo, D., Cho, S., & Kang, P. (2019). Multi-co-training for document classification using various document representations: TF-IDF, LDA, and Doc2Vec. *Information Sciences*, 477, 15-29.
- [17] Jabri, S., Dahbi, A., Gadi, T., & Bassir, A. (2018, April). Ranking of text documents using TF-IDF weighting and association rules mining. In *2018 4th international conference on optimization and applications (ICOA)* (pp. 1-6). IEEE.
- [18] Ramadhani, D. A., Rochimah, S., & Yuhana, U. L. (2015). Classification of non-functional requirements using semantic-FSKNN based ISO/IEC 9126. *TELKOMNIKA (Telecommunication Computing Electronics and Control)*, 13(4), 1456-1465.
- [19] Diwali, A., Saeedi, K., Dashtipour, K., Gogate, M., Cambria, E., & Hussain, A. (2023). Sentiment analysis meets explainable artificial intelligence: A survey on explainable sentiment analysis. *IEEE Transactions on Affective Computing*.
- [20] Mathews, S. M. (2019). Explainable artificial intelligence applications in NLP, biomedical, and malware classification: a literature review. In *Intelligent Computing: Proceedings of the 2019 Computing Conference, Volume 2* (pp. 1269-1292). Springer International Publishing.
- [21] Zahoor, K., & Bawany, N. Z. (2024). Explainable artificial intelligence approach towards classifying educational android app reviews using deep learning. *Interactive Learning Environments*, 32(9), 5227-5252.
- [22] Zhang, H., & Ogasawara, K. (2023). Grad-CAM-based explainable artificial intelligence related to medical text processing. *Bioengineering*, 10(9), 1070.
- [23] Pospelova, N., Tarasova, A., Subbotina, N., Koroleva, N., Raimova, N., & Lydia, E. L. (2024). Explainable Artificial Intelligence and Natural Language Processing for Unraveling Deceptive Contents. *Fusion: Practice & Applications*, 14(2).
- [24] Chazette, L., Brunotte, W., & Speith, T. (2022). Explainable software systems: from requirements analysis to system evaluation. *Requirements Engineering*, 27(4), 457-487.
- [25] Magableh, A. A. (2023, August). Towards leveraging explainable artificial intelligent (XAI) in requirements engineering (RE) to identify aspect (crosscutting concern): a systematic literature review (SLR) and bibliometric analysis. In *2023 International Conference on Information Technology (ICIT)* (pp. 319-326). IEEE.
- [26] Bai, S., Shi, S., Han, C., Yang, M., Gupta, B. B., & Arya, V. (2024). Prioritizing user requirements for digital products using explainable artificial intelligence: A data-driven analysis on video conferencing apps. *Future Generation Computer Systems*, 158, 167-182.
- [1] Sommerville, I. (2015). *Software engineering*. 10th. Book Software Engineering. 10th, Series Software Engineering.
- [2] Zhang, Y., Xu, F., Zou, J., Petrosian, O. L., & Krinkin, K. V. (2021, June). XAI evaluation: evaluating black-box model explanations for prediction. In *2021 II International Conference on Neural Networks and Neurotechnologies (NeuroNT)* (pp. 13-16). IEEE.
- [3] Binkhonain, M., & Zhao, L. (2019). A review of machine learning algorithms for identification and classification of non-functional requirements. *Expert Systems with Applications: X*, 1, 100001.
- [4] Khayashi, F., Jamasb, B., Akbari, R., & Shamsinejadbakaki, P. (2022). Deep Learning Methods for Software Requirement Classification: A Performance Study on the PURE dataset. *arXiv preprint arXiv:2211.05286*.
- [5] Dias Canedo, E., & Cordeiro Mendes, B. (2020). Software requirements classification using machine learning algorithms. *Entropy*, 22(9), 1057.
- [6] Rahimi, N., Eassa, F., & Elrefaei, L. (2021). One-and two-phase software requirement classification using ensemble deep learning. *Entropy*, 23(10), 1264.
- [7] Quba, G. Y., Al Qaisi, H., Althunibat, A., & AlZu'bi, S. (2021, July). Software requirements classification using machine learning algorithm's. In *2021 International Conference on Information Technology (ICIT)* (pp. 685-690). IEEE.
- [8] Jesus, S., Belém, C., Balayan, V., Bento, J., Saleiro, P., Bizarro, P., & Gama, J. (2021, March). How can I choose an explainer? An application-grounded evaluation of post-hoc explanations. In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency* (pp. 805-815).
- [9] Ribeiro, M. T., Singh, S., & Guestrin, C. (2016, August). "Why should i trust you?" Explaining the predictions of any classifier. In *Proceedings of the 22nd ACM SIGKDD international conference on knowledge discovery and data mining* (pp. 1135-1144).
- [10] Lundberg, S. M., & Lee, S. I. (2017, December). A unified approach to interpreting model predictions. In *Proceedings of the 31st International Conference on Neural Information Processing Systems* (pp. 4768-4777).
- [11] Habiba, U. E., Bogner, J., & Wagner, S. (2022, August). Can Requirements Engineering Support Explainable Artificial Intelligence? Towards a User-Centric Approach for Explainability Requirements. In *2022 IEEE 30th International Requirements Engineering Conference Workshops (REW)* (pp. 162-165). IEEE.
- [12] Clement, T., Kemmerzell, N., Abdelaal, M., & Amberg, M. (2023). XAIR: A Systematic Metareview of Explainable AI (XAI) Aligned to the Software Development Process. *Machine Learning and Knowledge Extraction*, 5(1), 78-108.
- [13] Jiarpakdee, J., Tantithamthavorn, C. K., & Grundy, J. (2021, May). Practitioners' perceptions of the goals and visual explanations of defect prediction models. In *2021 IEEE/ACM 18th International Conference on Mining Software Repositories (MSR)* (pp. 432-443). IEEE.