

Explainable Large Language Model for Islamic and Humanities Sciences Based on Knowledge Graphs

Ali Mirarab 

Assistant Professor, Information Dissemination and Knowledge Exchange Department, Islamic Science and Culture Research Institute, Qom, Iran (**Corresponding author**). alimirarab@isca.ac.ir

Fatemeh Darestani Farahani 

B.A., Department of Knowledge and Information Science, University of Qom, Qom, Iran.
darestanyfatemeh@gmail.com

Abstract

Purpose: This research explores the challenges and opportunities associated with integrating knowledge graphs and large language models in the fields of Islamic studies and the humanities. The primary motivation is the complexity and diversity of concepts and categories within these fields. Religious and humanistic concepts inherently possess multiple dimensions and layers, demanding innovative and efficient analytical approaches. Knowledge graphs are powerful tools for organizing information and representing relationships between concepts. They clarify and visualize intricate connections, enabling researchers to easily understand the interrelations among various concepts. The core objective of this study is to present a framework that integrates these two concepts within Islamic and humanities studies, thereby enabling a more precise and efficient analysis of Islamic and humanistic concepts.

Method: This research is a conceptual study that employs an analytical approach. Data were collected through a systematic literature review, which identified key and fundamental concepts. Existing resources in both the knowledge graph and large language model domains were purposefully selected for analysis.

Findings: The primary finding of this research is the development of an innovative framework that integrates knowledge graphs and large language models to analyze and infer concepts within Islamic and humanities studies. Knowledge graphs serve as tools for organizing information and representing relationships between concepts, helping to clarify and visualize intricate connections among them. Furthermore, large language models can process natural language and extract meaningful information from texts. The integration of these two approaches can foster the creation of innovative frameworks for data analysis and inference within the field of religious studies.

Conclusion: The integration of knowledge graphs and large language models represents a novel and efficient approach to analyzing and inferring complex concepts in Islamic and humanities studies. The proposed framework can serve as a model for future research and the development of intelligent systems in these fields, thereby enhancing the research and analysis processes in Islamic and humanities studies.

Keywords: Knowledge graph, Large language models, Explainable artificial intelligence, Islamic and humanities studies, Logic, Conceptual inference.

Cite this article: Mirarab, A. & Darestani Farahani, F. (2024). Explainable Large Language Model for Islamic and Humanities Sciences Based on Knowledge Graphs. *Sciences and Techniques of Information Management*, 10(4): 7-38. <https://doi.org/10.22091/stim.2024.11352.2162>

Received: 2024-06-20 ; **Revised:** 2024-07-24 ; **Accepted:** 2024-10-01 ; **Published online:** 2024-12-26

© The Author(s).

Article type: Research Article

Published by: University of Qom.





ارائه الگویی برای مدل زبانی بزرگ توضیح‌پذیر علوم اسلامی - انسانی مبتنی بر گراف دانش

علی میرعرب^{id}

استادیار، گروه اشاعه اطلاعات و تبادل دانش، پژوهشگاه علوم و فرهنگ اسلامی، قم، ایران (نویسنده مسئول).
alimirab@isca.ac.ir

فاطمه دارستانی فراهانی^{id}

کارشناسی، گروه علم اطلاعات و دانش‌شناسی، دانشگاه قم، قم، ایران. darestanyfatemeh@gmail.com

چکیده

هدف: پژوهش حاضر به بررسی چالش‌ها و فرصت‌های ترکیب گراف دانش^۲ و مدل زبانی بزرگ^۳ در حوزه علوم اسلامی - انسانی می‌پردازد. یکی از دلایل اصلی این مسئله، پیچیدگی و تنوع مفاهیم و مقولات در این حوزه است. مفاهیم دینی و انسانی به‌طور طبیعی دارای ابعاد و لایه‌های متعددی هستند که تحلیل آن‌ها نیازمند رویکردهایی نوین و کارآمد است. گراف دانش به‌عنوان ابزاری قدرتمند برای سازمان‌دهی اطلاعات و نمایش روابط بین مفاهیم شناخته می‌شود. این ابزار می‌تواند به شفاف‌سازی و تجسم روابط پیچیده میان مفاهیم کمک کند و به پژوهشگران امکان می‌دهد تا به‌راحتی روابط میان مفاهیم مختلف را درک کنند. هدف اصلی پژوهش حاضر، ارائه الگویی جهت ترکیب این دو مفهوم در علوم اسلامی - انسانی برای تحلیل دقیق‌تر و کارآمدتر مفاهیم اسلامی و انسانی است.

روش: پژوهش پیش‌رو از لحاظ هدف توسعه‌ای بوده و با رویکرد کیفی انجام شده است. روش گردآوری داده‌ها کتابخانه‌ای و برای مرور نظام‌مند متون بوده که به شناسایی مفاهیم کلیدی و اساسی می‌پردازد. منابع موجود در دو حوزه گراف دانش و مدل‌های بزرگ زبانی با روش هدفمند جهت تحلیل، انتخاب و مورد بررسی قرار گرفته است.

یافته‌ها: یافته اصلی پژوهش حاضر، طراحی یک الگوی جدید مبتنی بر ترکیب گراف دانش و مدل‌های زبانی بزرگ برای تحلیل و استنتاج مفاهیم علوم اسلامی - انسانی است. گراف دانش به‌عنوان ابزاری برای سازمان‌دهی اطلاعات و نمایش روابط میان مفاهیم می‌تواند به شفاف‌سازی و تجسم روابط پیچیده میان مفاهیم کمک کند. همچنین، مدل‌های زبانی بزرگ قادر به پردازش زبان طبیعی^۴ و استخراج اطلاعات معنادار از متون هستند. ترکیب این دو رویکرد می‌تواند به توسعه چارچوب‌های نوینی برای تحلیل و استنتاج داده‌ها در حوزه‌های علوم دینی کمک کند.

استناد به این مقاله: میرعرب، علی؛ دارستانی فراهانی، فاطمه (۱۴۰۳). ارائه الگویی برای مدل زبانی بزرگ توضیح‌پذیر علوم اسلامی - انسانی مبتنی

بر گراف دانش. *علوم و فنون مدیریت اطلاعات*. ۱۰(۴): ۷-۳۸. <https://doi.org/10.22091/stim.2024.11352.2162>

تاریخ دریافت: ۱۴۰۳/۰۳/۳۱؛ تاریخ اصلاح: ۱۴۰۳/۰۵/۳۰؛ تاریخ پذیرش: ۱۴۰۳/۰۷/۱۰؛ تاریخ انتشار آنلاین: ۱۴۰۳/۱۰/۰۶

ناشر: دانشگاه قم

نوع مقاله: پژوهشی

© نویسندگان.

2. Knowledge graph

3. LLM

4. Natural language processing (NLP)



نتیجه‌گیری: تلفیق گراف دانش و مدل‌های زبانی بزرگ به‌عنوان رویکردی نوین و کارآمد در تحلیل و استنتاج مفاهیم پیچیده در علوم اسلامی- انسانی مطرح است. این الگوی پیشنهادی می‌تواند به‌عنوان الگویی برای تحقیقات آتی و توسعه سیستم‌های هوشمند در این حوزه‌ها مورد بهره‌برداری قرار گیرد و به بهبود فرایندهای تحقیق و تحلیل در علوم اسلامی- انسانی کمک کند.

کلیدواژه‌ها: گراف دانش، مدل‌های زبانی بزرگ، هوش مصنوعی توضیح‌پذیر^۱، علوم اسلامی- انسانی، منطق، استنتاج مفهومی.

۱. مقدمه

با توجه به اینکه حرکت علم داده و تحلیل داده، به سمت استفاده از ساختار گراف گونه برای ذخیره، بازیابی و توسعه دانش بوده و به زودی، جای پایگاه داده‌های رابطه‌ای را خواهد گرفت، طراحی و ایجاد چنین بستری می‌تواند ایران را به عنوان یکی از قطب‌های مبحث وب معنایی تبدیل کند. لازم به ذکر است در حال حاضر بسیاری از انواع استنتاج‌های غیرتوصیفی (مانند استنتاج در فضای منطق پیش فرض، استنتاج در فضای منطق پویا و غیره) به طور صنعتی و در مقیاس بالا در لبه مرز دانش نیز پیاده‌سازی نشده است. امروزه، برای مدل‌سازی دقیق‌تر و کاراتر دانش، به جای بهره‌گیری از جداول و پایگاه داده‌های قدیمی‌تر، از گراف‌ها استفاده می‌شود. این بهره‌گیری را به طور نمونه می‌توان در عملکرد گوگل، به عنوان یکی از پیشتازهای مدیریت و بازیابی دانش، مشاهده کرد که از سال ۲۰۱۲، قدم‌های اولیه برای ایجاد گراف دانش را برداشته و تاکنون توجه قابل ملاحظه‌ای به توسعه آن داشته است (سالیوان، ۲۰۲۰). در همین راستا در داخل کشور نیز به همت آزمایشگاه داده‌کاوی دانشگاه علم و صنعت، گراف دانشی با عنوان «فارس بیس» طراحی و ساخته شده است. کارکردهای مختلفی برای گراف دانش متصور است که به طور نمونه می‌توان به پاسخ‌گویی خودکار، سیستم‌های توصیه‌کننده، بازیابی اطلاعات، تحلیل اطلاعات در یک دامنه خاص، دسته‌بندی و جز اینها اشاره کرد.

از مهم‌ترین مسائلی که در زمینه گراف دانش مطرح بوده، آشکارسازی اطلاعات نهانی است که استدلال‌گرها می‌توانند آن‌ها را هویدا کنند. اخیراً استنتاج روی گراف دانش و کسب دانش جدید بر پایه دانش قبل، یکی از موضوعات مطرح و قابل توجه مجامع علمی شده است. این در حالی است که استنتاج، پیچیدگی‌های زیادی دارد و همین پیچیدگی‌ها سبب شده که منطق‌های مختلفی برای نیازهای متنوع شکل گیرد. یکی از مسائل مهم در تجزیه و تحلیل گراف دانش، بحث اعتبارسنجی واقعیاتی هستند که در گراف دانش موجود است. اینکه مستند هر واقعیت ارائه شده در گراف چیست، به شدت در فرایند استنتاج مؤثر است. به طور مثال در مواردی که برای دامنه‌ای، هستان‌نگاری طراحی می‌شود، ممکن است نتایج استنتاج ماشینی با هم در تعارض باشند. اگر چارچوبی برای اعتبارسنجی گراف دانش وجود داشته باشد، می‌توان، فرض‌های مختلف ناسازگاری را به صورت کامل مدل‌سازی کرد و دیگر در مواجهه با تناقض و تعارض، فرایند استنتاج ماشینی از حرکت نایستد و با توجه به قوانین اعتبارسنجی، فرایند استنتاج در جهتی دیگر به حرکت خود ادامه دهد. برای این منظور در مواردی لازم است گراف‌های دانش از جهاتی توسعه یابند.

پیش‌بینی می‌شود مطالعات حوزه گراف دانش در داخل کشور با توجه به پیشرفت‌ها و مزایای آن، به سمت مدل‌سازی دانش و اطلاعات رود و دیگر در زمینه مدیریت دانش از پایگاه داده‌های متداول

استفاده چندانی صورت نگیرد. در این راستا، وجود ابزار مناسبی برای اعتبارسنجی گراف دانش از جهت تحلیل دادگان مدل شده در گراف دانش ضروری خواهد بود. به عنوان نمونه می توان به کمک هستان نگارها در ارتباط با اعتبارسنجی گراف دانش موارد ذیل را متصور بود:

● می توان نوع موجودیت های احتمالی الفاظی که در یک متن به کار رفته اند را به طور خودکار پیدا نمود.

● الفاظ مشترک را ابهام زدایی نمود.

● از نظر منطقی گزاره ها را تحلیل کرد و لازمه های فاسد یک گزاره، تناقض ها و نتایج استدلال های مورد غفلت واقع شده را یافت.

● پاسخ برخی سوالاتی که بالقوه معلوم، ولی در اطلاعات موجود نهفته هستند را رصد نمود.

● استدلال های ناقص را به صورت پیشنهادی تکمیل کرد.

● جست وجوهای معنایی و غیر وابسته به لفظ انجام داد.

مدل های زبانی بزرگ مانند GPT که توسط OpenAI توسعه یافته، این توانایی را دارند که متنی شبیه انسان تولید کنند. حضور این فناوری، حوزه تولید و پردازش زبان طبیعی را متحول کرده است. دانش قابل توجه و خلاقیت ظاهری آن ها باعث شده که افراد بیشتری از طریق این ربات های قدرتمند هوش مصنوعی برای پاسخ به فوری ترین سؤالات خود استفاده کنند؛ ولی با این حال، مانند هر مدل یادگیری ماشینی، محدودیت های خود را نیز دارد. یکی از محدودیت های مدل های زبانی، عدم درک صریح و منطق محور از بافت و دانش پس زمینه متنی است که او تولید می کند. به عنوان مثال، اگر از آن خواسته شود در مورد یک موضوع خاص بنویسد، ممکن است متنی تولید کند که از نظر نحو و دستور زبان صحیح باشد؛ اما فاقد عمق و ظرافت یک متخصص در آن بافت باشد. محدودیت دیگر، ناتوانی آن در استدلال و ایجاد ارتباط منطقی بین مفاهیم مختلف است. در حالی که می تواند متنی تولید کند که ظاهر منطقی مناسبی دارد، اما توانایی استنتاج و نتیجه گیری براساس اطلاعات ارائه شده را حداقل تا نسخه موجود ندارد. یکی از راه های غلبه بر محدودیت های اشاره شده، ترکیب مدل های زبانی با گراف دانش است. علاوه بر این، گراف دانش می تواند اطلاعات دقیقی را در اختیار مدل زبانی قرار دهد که آن اطلاعات را در تولید متن خود بگنجاند. این رویکرد می تواند اطلاعات خاص و استواری را در مورد هر موضوعی ارائه دهد و متن تولید شده را آموزنده تر و مفیدتر کند. در نتیجه، یک مدل زبانی قدرتمند می تواند با بهره گیری از داده های گراف های دانش و استدلال ورزی های صریح روی آن، بر محدودیت های منطقی خود غلبه کرده و در نتیجه یک سیستم تولید متن هوشمندتر و آموزنده تر ایجاد کند.

با توجه به مطالب ارائه شده، مشکلاتی در زمینه توضیح پذیری وجود دارد که ضرورت تشکیل

- یک مدل زبانی بزرگ علوم اسلامی - انسانی را مشخص می‌کند؛ مشکلات شامل موارد زیر است:
- مدل‌های زبانی بزرگ موجود، خروجی در سطح روابط و قواعد منطقی - استنتاجی ندارند،
 - مدل‌های زبانی بزرگ در خدمت پنج زبان پرکاربرد دنیا،
 - عدم دسترس‌پذیری داده‌های نیمه‌ساخت‌یافته‌ای که مدل‌های زبانی بزرگ مبتنی بر آموزش دیده‌اند،

• جهت‌دار بودن خروجی‌های مدل‌های زبانی بزرگ.

اهداف اصلی ایجاد مدل زبانی بزرگ علوم اسلامی - انسانی متناسب با بافت پژوهش حاضر مشتمل بر موارد ذیل است:

ارتقاء درک زبانی: با استفاده از مدل زبانی بزرگ، درک زبانی در حوزه علوم اسلامی - انسانی ارتقاء و بهبود پیدا خواهد کرد. این مدل قادر است متون علوم اسلامی - انسانی را تجزیه و تحلیل کند و درک معنا و مفهوم را بهبود دهد.

تولید متن: با استفاده از مدل زبانی بزرگ، قابلیت تولید متنی باکیفیت و طبیعی در حوزه علوم اسلامی - انسانی بهبود پیدا خواهد کرد. این مدل قادر است ساختار و معنای جملات مربوط به علوم اسلامی - انسانی را درک کرده و به صورت خودکار متون جدید تولید کند.

استفاده در تحقیقات دینی: مدل‌های زبانی بزرگ را می‌توان برای استفاده در تحقیقات دینی و مطالعات علوم اسلامی - انسانی آموزش داد. این کار به توسعه فناوری زبانی در حوزه تحقیق و تولید متن در علوم اسلامی - انسانی کمک خواهد کرد.

پشتیبانی از آموزش و تعلیم: مدل زبانی بزرگ پیشنهاد شده قابلیت انتقال یادگیری به برنامه‌های کاربردی مرتبط با آموزش و تعلیم در حوزه علوم اسلامی - انسانی را داراست. از این مدل می‌توان در زمینه‌هایی مانند تفسیر و تحلیل متون دینی، پاسخگویی به سؤالات دینی، ترجمه متون مذهبی و سیستم‌های هوشمند در زمینه علوم اسلامی - انسانی استفاده کرد.

با توجه به مطالب پیش‌گفته و کاربردهای متصور از ترکیب مدل‌های زبانی بزرگ و گراف دانش، هدف پژوهش حاضر، طراحی و توسعه یک مدل زبانی در حوزه علوم اسلامی - انسانی با استفاده از گراف دانش علوم اسلامی - انسانی در سطح روابط و قواعد منطقی - استنتاجی است.

۲. آثار مرتبط

لو و اسپسیا^۱ (۲۰۲۴) در پژوهشی با هدف بررسی روش‌های توضیح‌پذیری برای مدل‌های زبانی

بزرگ، نشان دادند که با توجه به توانایی‌های قابل توجه مدل‌های زبانی بزرگ در پردازش زبان طبیعی، اما پیچیدگی داخلی آن‌ها همچنان ناشناخته است و عدم شفافیت می‌تواند خطراتی در برنامه‌های کاربردی مبتنی بر آنها ایجاد کند. بنابراین، فهم و توضیح این مدل‌ها برای تبیین رفتار، محدودیت‌ها و تأثیرات اجتماعی آن‌ها ضروری است. تکنیک‌های توضیح‌پذیری براساس پارادایم‌های آموزشی را می‌توان به دودسته تقسیم کرد: پارادایم سنتی مبتنی بر تنظیم دقیق و پارادایم مبتنی بر پرامپت.^۱ در پارادایم سنتی، مدل‌ها ابتدا با یک مجموعه داده بزرگ و بدون برچسب آموزش می‌بینند و سپس با داده‌های برچسب‌دار خاص تنظیم می‌شوند. در پارادایم مبتنی بر پرامپت، مدل‌ها از طریق پرامپت و بدون نیاز به داده‌های آموزشی اضافی یاد می‌گیرند.

ماورپیس^۲ و همکاران (۲۰۲۴) رویکردی نوآورانه در زمینه توضیح‌پذیری انسان-محور هوش مصنوعی^۳ ارائه کردند. پژوهشگران یک مدل زبانی بزرگ مبتنی بر GPT به نام «x- [p]lAIIn» را توسعه داده‌اند که می‌تواند خروجی‌های پیچیده الگوریتم‌های XAI را به توضیحات متنی ساده و قابل فهم برای انواع گروه‌های مخاطب از جمله متخصصان تجاری و دانشگاهی، تبدیل کند. روش‌شناسی تحقیق شامل ادغام روش‌های XAI نظیر SHAP، LIME و GradCam در زیرساخت LLM سفارشی با استفاده از GPT-Builder است. ویژگی‌های کلیدی مدل پیشنهادی عبارتند از: توانایی تولید خلاصه‌های روان و قابل درک از روش‌های مختلف XAI، سفارشی شده برای سطوح مختلف دانش و علایق گروه‌های مختلف مخاطب. این امر درگیری و درک کاربران را در بخش‌های مختلف افزایش می‌دهد. همچنین مستقل از روش‌های XAI خاص که کاربرد آن را در طیف گسترده‌ای از روش‌های XAI و حوزه‌های دانش گسترش می‌دهد؛ فرایند تصمیم‌گیری برای کاربران نهایی با ارائه توضیحات واضح، به موقع و متناسب با زمینه را تسهیل می‌کند. اعتبارسنجی تجربی از طریق مطالعات موردی نشان داده که این مدل می‌تواند توضیحات متناسب با گروه مخاطب ارائه کند که قابل فهم و مرتبط هستند. نویسندگان بر این باورند که این رویکرد می‌تواند دسترسی به مفاهیم پیشرفته XAI را برای غیرمتخصصان آسان کند و فاصله بین فناوری‌های پیچیده هوش مصنوعی و کاربردهای عملی آنها را پر کند. در نهایت، یافته‌ها نشان می‌دهد که این رویکرد مسیری امیدبخش برای کاربرد مدل‌های زبان بزرگ در قابل فهم کردن مفاهیم پیشرفته هوش مصنوعی برای طیف گسترده‌ای از کاربران است.

<http://stun.gom.ac.ir>

1. Prompt
2. Mavrepisa
3. Human-Centric XAI

عباس کوهی و همکاران (۲۰۲۴) در پژوهشی به بررسی کارایی مدل‌های زبانی بزرگ‌در زبان فارسی پرداختند. اگرچه مدل‌هایی مانند چت جی‌بی‌تی در زبان انگلیسی عملکرد چشمگیری دارند، اما کارایی آن‌ها در زبان‌های کم‌منبع همچنان مورد سوال و بررسی است. تمرکز اصلی این پژوهش بر روی GPT-3.5-turbo بوده، اما شامل مدل‌های GPT-4 و OpenChat-3.5 نیز است، تا ارزیابی کامل‌تری ارائه دهد. ارزیابی شامل مجموعه‌ای متنوع از وظایف در دسته‌های کلاسیک، استدلال و مبتنی بر دانش است. به منظور مقایسه جامع، این مدل‌ها با مدل‌های موجود که مخصوص هر وظیفه بهینه‌سازی شده‌اند، مورد ارزیابی قرار گرفتند. یافته‌ها نشان داد که مدل‌های زبانی بزرگ، به ویژه GPT-4، در وظایفی که نیاز به توانایی استدلال و درک گسترده‌ای از دانش عمومی دارند، عملکرد خوبی داشته، اما در بسیاری از موارد، مدل‌های کوچک‌تر که بهینه‌سازی شده‌اند، عملکرد بهتری نشان دادند. این پژوهش همچنین بهبود عملکرد را زمانی که مجموعه‌های آزمایشی قبل از ورودی به GPT-3.5 به انگلیسی ترجمه می‌شوند، نشان داد. این نتایج نشان‌دهنده پتانسیل قابل توجهی برای بهبود عملکرد مدل‌های زبانی بزرگ در زبان فارسی است. این امر به ویژه به دلیل ویژگی‌های منحصر به فرد زبان فارسی، از جمله الفبای متفاوت و سبک‌های نوشتاری متنوع، قابل توجه است. در این مطالعه، دو معیار جدید برای وظایف استدلالی با توجه به کمبود مجموعه داده‌های فارسی معرفی شده است: یکی براساس سوالات ریاضی دبستان و دیگری مشتق شده از آزمون‌های ورودی برای پایه‌های هفتم و دهم. طبق نتایج، بهترین عملکرد GPT-4 در تحلیل احساسات با امتیاز F1 برابر ۰/۹۰۶ در حالت سه‌شات با دستورالعمل انگلیسی است. در ترجمه از انگلیسی به فارسی، GPT-4 بالاترین امتیاز BLEU برابر ۸/۷ را در حالت سه‌شات کسب کرده است. عملکرد GPT-4 در پاسخ‌های چندگزینه‌ای (ریاضی و منطق) به ۷۲/۵٪ دقت رسیده است.

هوانگ و چانگ^۱ (۲۰۲۳) در پژوهشی به بررسی تکنیک‌های بهبود و برانگیختن استدلال در مدل‌های زبان بزرگ پرداختند. نتایج نشان داد که در سال‌های اخیر، مدل‌های زبانی بزرگ پیشرفت‌های قابل توجهی در پردازش زبان طبیعی داشته‌اند و این مدل‌ها ممکن است با رسیدن به اندازه کافی بزرگ، توانایی‌های استدلالی از خود نشان دهند. استدلال به عنوان بخش اساسی از هوش انسانی نقش مهمی در حل مسئله، تصمیم‌گیری و تفکر انتقادی ایفا می‌کند. در حالی که مدل‌های زبانی بزرگ در برخی از وظایف استدلالی عملکرد خوبی دارند، اما هنوز مشخص نیست که این مدل‌ها واقعاً به چه میزان قادر به استدلال هستند.

1. Huang & Chang

ژنگ^۱ و همکاران (۲۰۲۳) با هدف نشان دادن چگونگی مدل‌های زبانی بزرگ در پیش‌بینی خواص مولکولی و ارائه توضیحات قابل درک به این یافته رسیدند که مدل‌های زبانی بزرگ می‌توانند برای استنتاج و توضیح در علوم طبیعی استفاده شوند؛ LLMها می‌توانند دانش علمی را از طریق ترکیب اطلاعات از متون علمی و استنتاج از داده‌های علمی به دست آورند و توضیحاتی برای نتایج پیش‌بینی‌های خود ارائه دهند. روش پیشنهادی آن‌ها، LLM4SD، توانسته است در پیش‌بینی خواص مولکول‌ها عملکرد بهتری نسبت به روش‌های موجود داشته باشد. این سیستم با استفاده از داده‌ها و ادبیات علمی، قوانینی برای پیش‌بینی ویژگی‌های مولکولی استخراج کرده و مدل‌های قابل تفسیر ایجاد می‌کند. نتایج به دست آمده نشان می‌دهد که LLM4SD در ۵۸ وظیفه در چهار حوزه علمی مختلف، شامل فیزیولوژی^۲، بیوفیزیک^۳، شیمی فیزیک و مکانیک کوانتومی، بهبود قابل توجهی در دقت پیش‌بینی‌ها داشته است. به طور خاص، این سیستم در حوزه مکانیک کوانتومی ۴۸/۲٪ بهبود و در شیمی فیزیک ۱۸/۵٪ بهبود در شاخص‌های ارزیابی عملکرد نسبت به بهترین مدل‌های موجود نشان داده است.

رشد توضیح‌پذیری مدل‌های زبانی بزرگ پژوهشی دیگر بود که توسط بینز^۴ و شولز^۵ (۲۰۲۳) انجام شد. مدل‌های زبانی بزرگ پس از تنظیم دقیق می‌توانند رفتار انسانی را بهتر از مدل‌های شناختی سنتی در دو حوزه تصمیم‌گیری توصیف و پیش‌بینی کنند. نتایج نشان داد که تنظیم دقیق این مدل‌ها به آنها اجازه می‌دهد رفتار افراد را در سطح فردی نیز مدل‌سازی کنند. همچنین، مدل‌هایی که بر روی چندین وظیفه تنظیم شده‌اند، قادر به پیش‌بینی رفتار انسانی در وظایف جدید و نادیده هستند. این پژوهش از مدل LLaMA با ۶۵ میلیارد پارامتر استفاده کرده و نشان داده که پس از تنظیم، مدل در کارهای مرتبط با تصمیم‌گیری از توصیف‌ها و تجربه‌ها عملکرد بهتری از مدل‌های مرتبط با یک خاص حوزه دارد. برای مثال، در مجموعه داده choices13k، مدل CENTaUR بهبود قابل توجهی در میزان احتمال منفی (۴۸۰۰۲/۳) در مقایسه با مدل BEAST (49448.1) داشت. در وظیفه horizon، CENTaUR با میزان احتمال منفی ۲۵۹۶۸/۶ بهتر از مدل هیبریدی^۶ (۲۹۰۴۲/۵) بود. عملکرد این نتایج نشان می‌دهد که مدل‌های زبانی بزرگ می‌توانند به عنوان مدل‌های شناختی عمومی

1. Zheng
2. Physiology
3. Biophysics
4. Binz
5. Schulz
6. Hybrid model

مورد استفاده قرار گیرند.

ژائو و همکاران (۲۰۲۴) در پژوهشی به بررسی زمینه رو به رشد توضیح‌پذیری مدل‌های زبانی بزرگ پرداختند که جنبه‌ای حیاتی و در عین حال چالش‌برانگیز از پردازش زبان طبیعی است. با توجه به نقش مهم LLMها در کاربردهای مختلف، ماهیت جعبه سیاه آن‌ها نگرانی‌هایی درباره شفافیت و استفاده اخلاقی به وجود می‌آورد. این پژوهش بر ضرورت بهبود توضیح‌پذیری در LLMها تأکید دارد و به دو جنبه اعتماد عمومی و نیاز جامعه فنی به درک عمیق‌تر این مدل‌ها می‌پردازد. این پژوهش به بررسی روش‌های موجود توضیح‌پذیری پرداخته و کاربردهای آنها در بهبود شفافیت و قابلیت اطمینان مدل‌ها را مورد بحث قرار می‌دهد. همچنین روش‌های ارزیابی نماینده را بررسی کرده و نقاط قوت و محدودیت‌های آنها را برجسته می‌کند.

هدف این بررسی پر کردن شکاف بین درک نظری و کاربرد عملی و ارائه بینش‌هایی برای تحقیقات و توسعه‌های آینده در زمینه توضیح‌پذیری LLMها است. این پژوهش تمرکز خود را بر مدل‌های بزرگ مبتنی بر ترنسفورمر^۱ مانند LLaMA گذاشته که چالش‌های خاصی از نظر تفسیرپذیری به دلیل مقیاس و پیچیدگی خود دارند. روش‌های توضیح‌پذیری به دو دسته تحلیل محلی و تحلیل جهانی تقسیم می‌شوند. تحلیل محلی به بررسی ویژگی‌های خاص و تحلیل بلوک‌های ترنسفورمر می‌پردازد، در حالی که تحلیل جهانی شامل روش‌های مبتنی بر پروب^۲ و تفسیر مکانیکی است. این پژوهش همچنین به کاربردهای بهبود توانایی‌های LLMها از جمله ویرایش مدل، بهبود قابلیت‌ها و تولید کنترل شده پرداخته است. در نهایت، این پژوهش بر اهمیت استفاده از توضیح‌پذیری به عنوان ابزاری برای رفع اشکال و بهبود مدل‌ها تأکید دارد.

هانگ^۳ و همکاران (2023b) نیز در پژوهشی به توانایی مدل‌های زبانی بزرگ مانند GPT در ارائه توضیحات خودکار برای پیش‌بینی‌هایشان پرداخته‌اند. هدف اصلی این است که مشخص شود این توضیحات تا چه حد دقیق و قابل اعتماد هستند، به ویژه در بررسی تحلیل احساسات و تخصیص ویژگی‌ها. این پژوهش روش‌های مختلفی را برای استخراج این توضیحات بررسی کرده و آن‌ها را با روش‌های سنتی مانند نقشه‌های برجستگی LIME و انسداد مقایسه می‌کند. نتایج نشان داد که توضیحات خودکار تولید شده توسط GPT عملکردی مشابه با روش‌های سنتی دارند؛ اما از نظر معیارهای مختلف توافقی، تفاوت‌های زیادی وجود دارد. این روش‌ها مقرون به صرفه‌تر نیز هستند؛

1. Transformer

2. Probing-based

3. Huang

زیرا همراه با پیش‌بینی تولید می‌شوند. در آزمایش‌ها، دو رویکرد برای تولید توضیحات مورد بررسی قرار گرفته است: ابتدا توضیح و سپس پیش‌بینی (E-P) و ابتدا پیش‌بینی و سپس توضیح (P-E). نتایج نشان می‌دهد که هر دو رویکرد، دقت بالایی دارند (به ترتیب ۸۵٪ و ۸۸٪)، اما هنوز کمتر از مدل‌هایی هستند که نیازی به تولید توضیحات ندارند (۹۲٪). این یافته‌ها نشان می‌دهند که ممکن است نیاز به بازنگری در روش‌های تفسیری فعلی برای این مدل‌ها در عصر مدل‌های زبانی بزرگ مانند GPT وجود داشته باشد.

چن^۱، سینگ^۲، و سراً^۳ (۲۰۲۳) در پژوهشی به چگونگی بهبود قابلیت توضیح‌دهندگی مدل‌های زبانی بزرگ پرداخته‌اند. این مدل‌ها، با وجود عملکرد بسیار قدرتمند در پردازش زبان طبیعی، همچنان به دلیل ساختار پیچیده و غیرشفاف خود، در ارائه توضیحات قابل فهم برای انسان‌ها، با چالش‌هایی مواجه‌اند. روش‌های اخیر عمدتاً بر استفاده از وزن‌های توجه برای توضیح پیش‌بینی‌های مدل‌های زبانی تمرکز کرده‌اند، اما این روش‌ها به‌تنهایی قادر به پوشش پیچیدگی‌های فزاینده این مدل‌ها نیستند. برای حل این مشکل، محققان LMExplainer را پیشنهاد می‌دهند که با استفاده از گراف دانش و شبکه عصبی^۴ توجه گراف، توضیحات قابل فهم‌تری ارائه می‌دهد. نتایج آزمایش‌ها نشان می‌دهد که LMExplainer در مقایسه با روش‌های موجود در زمینه پرسش پاسخ، عملکرد بهتری داشته و توضیحات جامع‌تر و واضح‌تری ارائه می‌کند. این روش با استفاده از گراف‌های دانش و شبکه‌های عصبی توجه، توانسته است مدل را بهبود بخشیده و فرایند استدلال مدل را به زبان طبیعی توضیح دهد. این امر نشان‌دهنده پتانسیل بالای LMExplainer در بهبود عملکرد و ارائه توضیحات مدل‌های زبانی است.

پاتل، کین، و پاتل^۵ (۲۰۲۳) بیان می‌کنند که مدل‌های زبانی بزرگ در حل مسائل متنوع زبان طبیعی عملکرد چشمگیری از خود نشان داده‌اند. با این حال، در زمینه دین اسلام، نمایش دقیق و واقعی از اعتقادات و آموزه‌های آن براساس قرآن و سنت، امری حیاتی است. این پژوهش به چالش ساخت LLM‌های مخصوص یک حوزه که معطوف به جهان‌بینی اسلامی باشند، می‌پردازد. چالش‌های استفاده از LLM‌های موجود عبارتند از: تمایل به تولید اطلاعات نادرست، سوگیری علیه مسلمانان و عدم توانایی در مدیریت محتوای مضر. هرچند روش‌هایی مانند A-INLP برای کاهش

<http://stun.gom.ac.ir>

1. Chen
2. Singh
3. Sra
4. Neural network
5. Patel, Kane & Patel

سوگیری‌ها ارائه شده‌اند؛ اما همچنان تلاش‌های بسیاری باید انجام شود.

از این‌رو توسعه مدل‌های زبانی مخصوص یک حوزه، برای افزایش دقت آن‌ها ضروری است. در این پژوهش از مجموعه داده IslamQA، معیارهای ارزیابی ترانسفورمر مانند BERTScore و فاصله بردار تعبیه شده استفاده شده است. همچنین برای مهندسی پرامپت از روش‌های zero-shot، few-shot و instruction-based روی GPT-3.5، GPT-4 و LLAMA استفاده شده است. به‌منظور تقویت بازیابی براساس مجموعه داده نیز از سیستم RAG با مجموعه داده احادیث برای بازیابی و استفاده از متون اسلامی مرتبط استفاده شده است. ریزتنظیم GPT-3.5 نیز بر روی داده‌های احادیث، سؤالات اسلامی و ترکیبی از هر دو انجام شده است. همچنین برای محافظت و جلوگیری از سوگیری و محتوای مضر از بسته Guardrails AI و زبان RAIL برای جلوگیری از پاسخ‌های خسونت‌آمیز یا توهین‌آمیز استفاده شده است. این پژوهش یک چارچوب جامع برای توسعه LLM‌های تخصصی، با تمرکز بر جهان‌بینی اسلامی ارائه داده و بیشتر می‌تواند برای بهبود ماژول‌ها، داده‌ها و محدودیت‌های فنی به کار برده شود و طراحی سیستم جدیدی نیست.

راستی میمندی و همکاران (۲۰۲۴) در پژوهشی به بررسی استفاده از تکنیک‌های پیشرفته پردازش زبان طبیعی و هوش مصنوعی برای تحلیل و تفسیر ادبیات فارسی، با تمرکز بر شعرهای فروغ فرخزاد پرداخته‌اند. با استفاده از روش‌های محاسباتی، هدف این است که الگوهای تماتیک، سبکی و زبانی در شعرهای فارسی شناسایی شوند. به‌طور خاص، این تحقیق از مدل‌های هوش مصنوعی مانند مدل‌های زبانی مبتنی بر ترانسفورمر برای خوشه‌بندی شعرها در یک چارچوب بدون نظارت بهره می‌برد. این پژوهش بر پتانسیل هوش مصنوعی در ارتقای فهم ما از میراث ادبی فارسی تأکید دارد و آثار فروغ فرخزاد به عنوان یک مطالعه موردی جامع ارائه می‌شود. این رویکرد نه تنها به حوزه علوم انسانی دیجیتال فارسی کمک می‌کند، بلکه معیاری برای پژوهش‌های آینده در مطالعات ادبیات فارسی با استفاده از تکنیک‌های محاسباتی فراهم می‌سازد.

در بخش روش‌شناسی، پنج مجموعه شعر فروغ فرخزاد تحلیل شده و از تکنیک‌هایی مانند تحلیل فراوانی، خوشه‌بندی و مدل‌سازی موضوعی با استفاده از LDA و ParsBERT استفاده شده است. نتایج نشان می‌دهد که مدل‌های پیشرفته می‌توانند به درک عمیق‌تری از اشعار کمک کنند و الگوهای معنایی پیچیده‌تری را آشکار سازند. یافته‌ها بر اهمیت استفاده از هوش مصنوعی و NLP در مطالعات ادبی به‌ویژه برای زبان‌ها و ادبیاتی که در پژوهش‌های علوم انسانی دیجیتال کمتر نمایان شده‌اند، تأکید می‌کند. این تکنیک‌ها می‌توانند بینش‌های جدیدی درباره متون ادبی ارائه دهند و به درک بهتر مضامین و ویژگی‌های سبکی کمک نمایند.

شو^۱ و همکاران (۲۰۲۴) روشی نوآورانه به نام کی جی-ال.ال.ام.^۲ با هدف ارائه یک روش جدید برای پیش‌بینی چندگانه لینک‌ها در گراف دانش ارائه داده‌اند که براساس مدل‌های بزرگ زبانی است و قابلیت استدلال و تعمیم‌پذیری برای پیش‌بینی ارتباطات چندگانه در گراف‌های دانش پیشنهاد می‌کند. این روش از الگوهای پردازش زبان طبیعی و روش‌های تعبیه شده در گراف‌های دانش استفاده کرده است. در این پژوهش داده‌های گراف دانش، ساختاری را به زبان طبیعی تبدیل کرده و از این طریق مدل‌های زبانی بزرگ را آموزش داده تا در پیش‌بینی ارتباطات، گراف‌های دانش را بهبود ببخشند. این چارچوب با تبدیل گراف دانش به سوالات زبان طبیعی، برای تشخیص و یادگیری نماینده‌های نهفته از موجودیت‌ها و ارتباطات آن‌ها طراحی شده است. با استفاده از این چارچوب، سه مدل زبانی بزرگ را در چارچوب آزمایشی KG-LLM آموزش داده و ارزیابی کاملی انجام شده است. همچنین، توانایی این فریم‌ورک برای ارائه قابلیت‌های صفر به مدل‌های زبان بزرگ، برای کار با سوالاتی که قبلاً دیده نشده‌اند، مورد بررسی قرار گرفته است. نتایج آزمایش‌ها نشان می‌دهد که عملکرد این روش می‌تواند ظرفیت عمومی مدل‌ها را به طور قابل توجهی افزایش دهد و اطمینان حاصل کند که پیش‌بینی‌ها در سناریوهای ناآشنا دقیق‌تر انجام می‌شود. در مجموع، این مطالعه روش جدیدی برای تبدیل داده‌های ساختاریافته گراف دانش به زبان طبیعی و استفاده از مدل‌های زبان بزرگ برای وظایف مرتبط با گراف دانش ارائه می‌کند.

کانفالونیری^۳ و همکاران (۲۰۲۰) در پژوهشی به بررسی تاریخی هوش مصنوعی توضیح‌پذیر پرداخته و به اهمیت توضیح‌پذیری در سیستم‌های هوش مصنوعی اشاره می‌کنند، به‌ویژه در زمینه‌هایی مانند رانندگی خودکار، تشخیص پزشکی و خدمات مالی. این پژوهش به چگونگی تعریف و درک توضیح‌پذیری در طول تاریخ پرداخته و نشان می‌دهد که چگونه این مفهوم از سیستم‌های خبره به یادگیری ماشین و سیستم‌های توصیه‌گر تکامل یافته است. همچنین به رویکردهای نوین مانند یادگیری و استدلال عصبی-نمادین و اهمیت توضیح‌پذیری پرداخته شده است.

همان‌طور که در تحقیقات مذکور مشاهده می‌شود، اقدامات مختلفی بر روی هستان‌نگاری، گراف دانش و مدل زبانی بزرگ انجام شده است. توضیح‌پذیری مدل‌های زبانی بزرگ با بهره‌گیری از گراف دانش، از موضوعات قابل توجه در هوش مصنوعی است. با بررسی پژوهش‌های موجود می‌توان ادعان کرد که تاکنون مطالعات کافی در این زمینه به‌ویژه در زبان فارسی انجام نشده است. برخلاف

1. Shu

2. Knowledge Graph Large Language Model (KG-LLm)

3. Confalonieri

زبان انگلیسی که در آن تلاش‌هایی برای توسعه مدل‌های توضیح‌پذیر صورت گرفته، زبان فارسی در این حوزه کار نشده است. در این راستا، هدف اصلی پژوهش حاضر تمرکز بر زبان فارسی با استفاده از گراف دانش در جهت توضیح‌پذیری مدل‌های زبانی بزرگ است. کاربست الگوی پیشنهادی در پژوهش حاضر می‌تواند تحولی در این زمینه ایجاد کند.

۳. روش پژوهش

پژوهش حاضر از لحاظ هدف یک پژوهش توسعه‌ای با رویکرد کیفی است. روش گردآوری داده‌ها کتابخانه‌ای بوده و به منظور مرور نظام‌مند متون انجام شده که با هدف توسعه یک الگوی پیشنهادی در جهت ترکیب گراف دانش و مدل‌های زبانی بزرگ ارائه شده است. در مطالعات کتابخانه‌ای، مقالات، کتاب‌ها و منابع معتبر مرتبط با گراف دانش، مدل‌های زبانی بزرگ و ساختارهای مرتبط با علوم اسلامی- انسانی که در چهار سال اخیر در مجلات معتبر به زبان انگلیسی منتشر شده‌اند، جمع‌آوری و بررسی شده‌اند. این مطالعات شامل مروری بر آخرین پژوهش‌ها و پیشرفت‌های علمی در زمینه‌های گراف دانش و مدل‌های زبانی است. پس از جمع‌آوری داده‌ها و منابع مرتبط، از روش تحلیل مفهومی برای شناسایی و استخراج اصول، مفاهیم و عناصر کلیدی مرتبط با موضوع تحقیق استفاده شده است. این تحلیل به استخراج ساختارهای بنیادین گراف دانش و نحوه پیوند آن با مدل‌های زبانی بزرگ انجامیده است. روش این پژوهش براساس ارائه یک الگو و چارچوب نظری برای ترکیب گراف دانش با مدل‌های زبانی بزرگ در علوم اسلامی- انسانی طراحی شده است. مراحل اصلی روش پژوهش به شرح زیر است:

● **شناسایی و بررسی ادبیات موضوع:** این مرحله شامل مطالعه گسترده پژوهش‌های پیشین در حوزه‌های مرتبط با گراف دانش، مدل‌های زبانی و کاربردهای آن‌ها در علوم اسلامی- انسانی است. هدف از این مرحله، شناسایی شکاف‌های موجود در دانش و ارائه الگویی برای پر کردن این شکاف‌ها است.

● **ارائه الگوی پیشنهادی:** در این مرحله، با تکیه بر اصول و مفاهیم به‌دست آمده از تحلیل‌های مفهومی، یک الگوی پیشنهادی برای تلفیق گراف دانش و مدل‌های زبانی بزرگ تدوین شده است. این الگو به شکلی طراحی شده است که بتواند به طور نظری و عملی، نیازهای تحلیل مفاهیم علوم اسلامی- انسانی را پوشش دهد.

● **ساختاردهی به مفاهیم در قالب گراف دانش:** الگوی پیشنهادی براساس ایجاد پیوندهای مفهومی بین داده‌های مختلف طراحی شده است. این ساختاردهی به کمک گراف دانش امکان‌پذیر است، به‌طوری که داده‌ها به‌صورت سلسله‌مراتبی و مفهومی سازمان‌دهی می‌شوند و روابط بین مفاهیم

به وضوح مشخص می شود.

● **ارائه چارچوب نظری:** در نهایت، چارچوب نظری این تحقیق شامل اصول و قواعد مورد استفاده برای طراحی الگوی پیشنهادی و نحوه تلفیق گراف دانش با مدل های زبانی بزرگ است. این چارچوب به گونه ای تنظیم شده که امکان توسعه آن در آینده برای تحقیقات عملیاتی وجود داشته باشد.

۴. الگوی پیشنهادی

در این بخش، الگوی پیشنهادی پژوهش برای ترکیب مدل های زبانی بزرگ با گراف دانش در زمینه علوم اسلامی- انسانی ارائه می شود. این راهکار به منظور ایجاد یک چارچوب کارآمد برای استنتاج مفهومی و سازمان دهی دانش در حوزه های متنوع علمی، خصوصاً در علوم اسلامی طراحی شده است. در طراحی این الگو، تلاش شده تا از ساختار گراف دانش برای ارتباط دهی مفاهیم پیچیده استفاده شود و به کمک مدل های زبانی بزرگ، امکان تولید و استنتاج خودکار متون فراهم گردد.

الگوی پیشنهادی شامل دو بخش کلیدی است:

● **ساختاردهی مفاهیم در قالب گراف دانش:** در این بخش، روابط و مفاهیم کلیدی در قالب یک ساختار گرافی سازماندهی می شوند تا به صورت سلسله مراتبی و مفهومی بین عناصر علمی و دینی پیوند برقرار شود.

● **تلفیق مدل زبانی بزرگ با گراف دانش:** این بخش به بررسی چگونگی ادغام مدل های زبانی با گراف دانش می پردازد، تا به کمک داده های متنی، قابلیت تولید و تفسیر خودکار متون و استنتاج های معنایی بهبود یابد.

این الگوی پیشنهادی به گونه ای طراحی شده است که بتواند زمینه ساز توسعه ابزارهای تحلیلی پیشرفته تر در حوزه های مختلف علمی و اسلامی شود و بهبود دقت و جامعیت در تحلیل و استنتاج را ممکن سازد.

<http://stn.gom.ac.ir>

۴-۱. ساختاردهی مفاهیم در قالب گراف دانش

در بخش اول الگوی پیشنهادی، به ساختاردهی مفاهیم در قالب گراف دانش پرداخته می شود. گراف دانش به عنوان یک ساختار منظم برای نمایش و سازمان دهی روابط میان مفاهیم عمل می کند. در این روش، مفاهیم به صورت گره هایی در یک شبکه گرافی تعریف می شوند و این گره ها از طریق روابط معنایی و مفهومی به یکدیگر متصل می گردند. چنین ساختاری به درک بهتر مفاهیم پیچیده و پیوندهای آن ها کمک می کند. در حوزه علوم اسلامی، مفاهیم دینی به شکل عمیقی به یکدیگر مرتبط بوده و استفاده از این ساختار می تواند روابط بین آیات، احادیث و مفاهیم فقهی را به صورت بصری و

ساختارمند نمایش دهد و دسترسی به این اطلاعات را تسهیل کند. در جدول (۱) مراحل ساختاردهی مفاهیم در قالب گراف دانش بیان شده است.

جدول ۱- مراحل ساختاردهی مفاهیم در قالب گراف دانش

مرحله اول	مرحله دوم	مرحله سوم	مرحله چهار	مرحله پنج
تجمع معنایی متون	تجمع ادله ^۱	تجمع کلام فقها ^۲		
نهایی سازی استنتاج‌های اولیه منطق توصیفی در دامنه‌ها- منطق احتمالاتی در زمان و مکان	پیاده‌سازی انواع مبانی رجالی قاعده‌مند مانند توثیقات عام و...- استنتاج احتمالاتی مذهب	بازسازی اصول اربعمائه براساس قواعد استنتاجی	مفروض‌یابی و ارزش‌گذاری پژوهشی ^۳	تحریر محل نزاع ^۴
استدلال ماشینی	به‌روزرسانی فتاوا با تغییر مبنا	بررسی پایبندی به اصول در فقه و کشف سایر تناقضات ^۵		
بازیابی هوشمند اطلاعات و نمایش دانش	نهایی سازی بازیابی اطلاعات حول شخص و کتاب (نمودار، نقشه و...)	توسعه بازیابی در موجودیت‌ها در زمینه‌های دیگر	طراحی میزهای پژوهشی اختصاصی	

۱. تجمع ادله یعنی بازیابی معنایی روایات که همان ادله مربوط به موضوعی هستند که فقیه قصد دارد جستجو شود.
۲. تجمع کلام فقها یعنی بازیابی معنایی گزاره‌های فقهی- اصولی که همان فحوص کلام فقهاست. وقتی یک فقه‌پژوه قصد دارد به دنبال یک مطلب در بین کلام فقهای گوناگون بگردد.
۳. مفروض‌یابی و ارزش‌گذاری پژوهشی یعنی ماشین بتواند درخت استدلال هر یک از گزاره‌هایی که یک فقیه اثبات کرده را تبیین کند و عملاً به یک درخت بزرگ چگال (پُر یال) از روابط استدلالی برسد. وقتی درخت استدلال ایجاد شد، می‌توان فهمید کدام گزاره‌ها کلیدی‌تر بوده و در قضایای مهم‌تر و بیشتری استفاده شده‌اند.
۴. تحریر محل نزاع یعنی ماشین بتواند درخت استدلال کلام دو نفر که با هم اختلاف نظر فقهی/اصولی دارند، استخراج کند، هم درخت استدلال اثبات هر دو نفر نسبت به دو قضیه متفاوت، هم درخت استدلال نفی دو نفر نسبت به اثبات طرف دیگر. سپس ماشین با توجه به ادراک منطقی- استنتاجی که دارد، بتواند مبتنی بر نزدیک‌ترین گزاره‌هایی که طرفین قبول دارند، اولین و مهم‌ترین گزاره‌ای که موجب این اختلاف شده است را تبیین منطقی کند (منظور از نزدیک‌ترین، نزدیک‌ترین گروه در درخت استدلال نسبت به قضایای اثبات شده است).
۵. کشف فروعی که یک فقیه قائل است ولی التفات ندارد یعنی ماشین به‌طور هوشمند فتاوی جدیدی را مبتنی بر مبانی یک فقیه تولید کند.
۶. بررسی پایبندی به اصول در فقه و کشف سایر تناقضات یعنی تشخیص هوشمند ناسازگاری بین نظرات فقیه و مبانی او (که در درخت استدلال قبل‌تر/بالتر از قضایای اثبات‌شده‌ی او هستند).

ساختاردهی به مفاهیم در چند بخش و در سطح منطق و استدلال قواعد مورد بحث قرار می‌گیرد. در ادامه چگونگی ساختاردهی به تقضیل بیان شده است.

۴-۱-۱. ساختار منطق

منطق، علم استدلال است و در هر دانش و معرفتی که سخن از استدلال باشد، ردّ پای منطق به چشم می‌خورد. مبنای مبحث گراف دانش نیز در چارچوب زبان‌های هستان‌نگاری نظیر اودیلول، فرایند استنتاج از مبادی معرفتی و پایگاه‌های دانشی به کوثری‌ها است که مبتنی بر روش‌های محاسباتی الگوریتمیک منطقی کارا و تصمیم‌پذیر مانند الگوریتم‌های تابلو^۱ و رزولوشن^۲ در منطق توصیفی و توسعه‌های نیمه‌کلاسیک آن (مانند زمان، مکان، پویا، معرفتی، بایایی و...) و غیرکلاسیک آن (مانند فازی، احتمالاتی، ربطی، شهودی، آزاد و غیره) است. تا شناخت درستی نسبت به این الگوریتم‌ها و قواعد آن‌ها و مبادی نظریه برهانی و نظریه مدل الگوریتم تابلو حاصل نشود؛ اولاً فرایند استنتاج ماشینی و خودکار (که مبنای نسل سوم وب، یعنی وب معنایی است) و رفتار استدلال‌گر در هستان‌نگارهای گوناگون قابل فهم نخواهد بود و ثانیاً نوآوری و ابداع در حوزه مباحث هوش مصنوعی مرتبط با این حوزه ممکن نخواهد بود.

۴-۱-۲. توسعه استدلال‌گر در منطق پیش‌فرض

هدف استفاده از منطق پیش‌فرض این است که بتوان با وجود برخی فرضیات، پیش‌فرض عمل استدلال را انجام داد. قابل ذکر است که تمرکز این پژوهش بر روی منطق پیش‌فرض اولویت‌دار است؛ یعنی اینکه اگر بین دو قانون تناقض وجود داشته باشد، قانونی را استفاده می‌کند که اولویت بیشتری از سمت کاربر به آن اختصاص داده شده است.

۴-۱-۳. توسعه استدلال‌گر در منطق پویا

بسیاری از نیازمندی‌های مدل‌سازی، ساختاری پویا دارند؛ از این‌رو برخی از مدل‌سازی‌ها با

<http://stim.gom.ac.ir>

۱. الگوریتم تابلو یا روش تابلو یکی از روش‌های گرافیکی قدرتمند در منطق است که برای بررسی صحت فرمول‌های منطقی و استنتاج نتایج از آن‌ها به‌کار می‌رود. این روش براساس ساختن یک درخت منطقی عمل می‌کند که هر شاخه آن نشان‌دهنده یک حالت ممکن برای فرمول است. با بررسی این شاخه‌ها می‌توان به نتیجه‌گیری در مورد صحت یا عدم صحت فرمول رسید.

۲. الگوریتم رزولوشن یکی دیگر از روش‌های مهم و پرکاربرد در منطق ریاضی و هوش مصنوعی است که برای استنتاج نتایج از مجموعه‌ای از گزاره‌ها به‌کار می‌رود. این الگوریتم براساس مفهوم تضاد کار می‌کند و به دنبال یافتن یک تناقض در مجموعه‌ای از گزاره‌ها است.

توجه به ماهیت پویا و وجود جهان‌های گوناگون، دارای قواعد پویا هستند. منطق پویا که در رده منطق‌های نیمه‌کلاسیک قرار دارد، این امکان را فراهم می‌کند که بتوان درک خوبی از توالی جهان‌های ایجاد شده توسط کنشگرهای تغییردهنده حالت داشته و استدلالی با علم بر پویایی رفتار محیط ارائه داد. استدلال‌گر پویایی که توسعه داده خواهد شد، علاوه بر امکان استنتاج پویا، قادر به امکان‌سنجی و تحقق‌پذیری قواعد پویای تعریف شده و نیز تصمیم‌گیری و ارائه مسیری برای رسیدن به مطلوب وارد شده توسط کاربر می‌باشد. مجموعه‌ای از امکانات افزوده شده در کنار یکدیگر، قابلیت استدلال‌ورزی پویا را مهیا خواهد کرد. در فرایند تصمیم‌گیری و برنامه‌ریزی جهت ارائه مسیر رسیدن به مطلوب، ممکن است مطلوب خواسته شده و نیز شرط‌های رسیدن به آن مطلوب، به صورت مستقیم در میان قواعد یافت نشوند؛ در این شرایط لحاظ کردن T-BOX می‌تواند کمک‌کننده باشد.

۴-۱-۴. طراحی و ساخت استدلال‌گر احتمالاتی

بخش قابل‌توجهی از اطلاعات در دنیای واقعی، براساس احتمالات و عدم قطعیت بیان می‌شود. از این‌رو، نیاز به استدلال‌گرهایی است که براساس این احتمالات عمل کنند و استنتاج‌هایی با تفسیرهای متناسب با پایگاه دانش با دامنه احتمالاتی در اختیار متخصصین مربوطه قرار دهند. استدلال‌گر احتمالاتی در قالب افزونه‌ای در وبگاه بر پایه ایده‌ای که از استدلال‌گرهای احتمالاتی بر پایه (TRILL) و (TRILL^P) گرفته شد، پیاده‌سازی خواهد شد. استدلال‌گر احتمالاتی، مانند استدلال‌گرهای TRILL، بر فرض مستقل بودن مفروضات اولیه هستان‌نگار استوار است و تنها در این صورت پاسخ صحیح را خواهد داشت. اطلاعات می‌تواند به دو شکلی قطعی (بدون بیان احتمال) و یا احتمالاتی (به صورت عددی بین صفر تا یک) بیان شوند. نتایج استنتاج شده به سه زبان طبیعی کنترل شده (فارسی، عربی، انگلیسی) و دو شکل احتمالاتی کمی و کیفی قابل نمایش برای کاربران می‌باشد. برای بیان کیفی نتایج، کاربر می‌تواند بازه‌هایی برای بیان قطعیت قوی تا احتمال ضعیف تعیین کند که بر مبنای این اعداد، جملاتی کیفی متناسب با زبان طبیعی انتخاب شده به کاربر نمایش داده می‌شود. همچنین این ساختار پس از توسعه، می‌تواند برای بیان احتمالات در فرایند تجمیع قرائن طرفینی استفاده گردد، که با توجه به بررسی انجام شده، تاکنون این فرایند در فضای دانشگاهی پیاده‌سازی نشده و امید می‌رود که نتایج پژوهشی و عملیاتی مناسبی حاصل شود.

۴-۱-۵. طراحی چارچوب ساخت و تغییر قواعد به کمک قواعد

استنتاج، به عنوان فرایند نتیجه‌گیری از داده‌های موجود است که در علوم مختلف از جمله فقه کاربرد فراوانی دارد و به طور نمونه برای استنتاج احکام شرعی، گاهی به مدل‌سازی مفاهیم معقول

ثانی فلسفی^۱ و روابط بین آن‌ها نیاز است. این کار به شکل ساده و معمول در منطق توصیفی قابل پیاده‌سازی نیست؛ زیرا اساساً منطق توصیفی تنها یک صورت تصمیم‌پذیر از منطق مرتبه اول است و این کار به منطق مرتبه دوم نیاز دارد. همچنین بعضی از روش‌ها (استفاده وسیع از کلاس‌های مرکب) منجر به پیچیدگی بیش از حد فرایند استنتاج می‌شود و بعضی دیگر (استفاده از یال‌های غیردخیل در استنتاج) نتایج استنتاج شده را به مقدار زیادی کاهش می‌دهد.

از طرفی یکی از مهم‌ترین چالش‌های پیش‌روی هر دانشمندی و به طور نمونه هر فقیهی در فرایند اجتهاد، حل تعارض بین ادله است. بعضی از تعارضات بین ادله با توجه به قرائن خاصی قابل رفع است؛ به طور نمونه تعارض بین عام و خاص و حاکم و محکوم در فضای فقه و اصول از این دست می‌باشد.

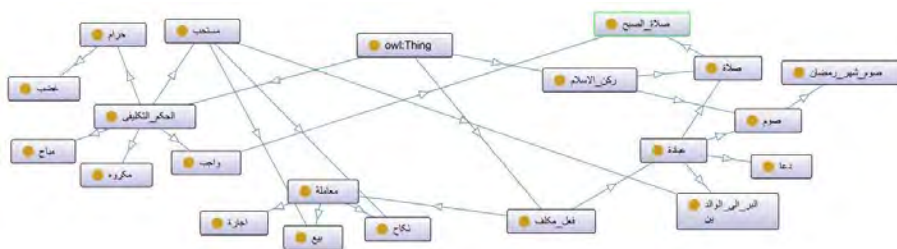
باتوجه به اینکه سازوکار رفع تناقض (تعارض) بر پایه اعتبار دو دلیل و تقدم یکی بر دیگری است، یک هستان‌نگار برای ادله نیز نیاز است که در دستور ساخت قرار گیرد. در این هستان‌نگار مفاهیمی که برای مدل‌سازی سازوکار رفع تعارض بین ادله نیاز است، مدل شده است؛ مفاهیمی همچون لفظ، حجت، روایت، اجماع و... در مجموع، هستان‌نگار نیمه مرتبه دو برای مدیریت مفاهیم معقول ثانی، کمک به حل تعارض بین ادله، داشتن کنترل بیشتر روی قواعد و ساخت قواعد جدید، ایجاد می‌شود. به طور کلی، روش کار در این قسمت به این صورت است که تمامی مفاهیم و روابط موجود در هستان‌نگار مرتبه اول، به منفرد تبدیل و با همان iri در هستان‌نگار نیمه مرتبه دوم قرار می‌گیرند. با این کار، می‌توان روی کلاس‌های موجود در هستان‌نگار مرتبه اول، سور قرار داد و بین آن‌ها روابط مختلفی ایجاد کرد.

در مرحله بعد، در هستان‌نگار مورد بحث، قوانین و قواعد مختلف جهت ساخت قوانین مختلف ایجاد شده و در نهایت، به کمک هستان‌نگار ادله، قواعدی جهت رفع بعضی از تعارضات، به هستان‌نگار اضافه شد. نمونه‌ای از هستان‌نگار در حوزه قواعد اصول فقه به شرح ذیل می‌باشد:

- هر واجبی مادامی که مزاحمی نداشته باشد و در مورد فعلیت آن اطلاعی نداشته باشیم، فعلی است.
- دو واجب متضاد مادامی که هیچ کدام بر دیگری رجحان نداشته باشند، عدل دیگری هستند.
- در فرض تراحم دو واجب انشائی، واجبی که بدل دارد، مرجوح بوده و آن بدل فعلی است.
- مادامی که حکم تکلیفی یک فعل خارجی، مشخص نباشد، انجام آن جایز است.

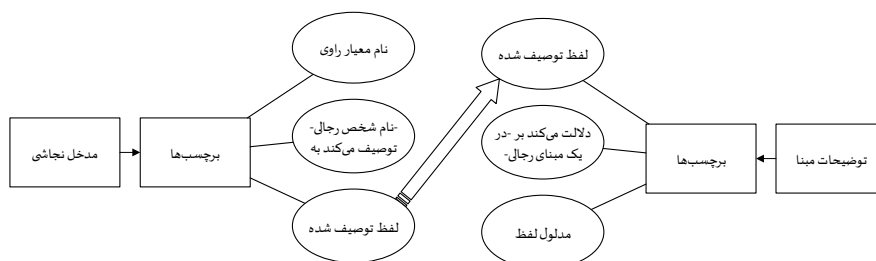
۱. معقول ثانی فلسفی مفاهیمی هستند که ذهن ما از ترکیب و پردازش مفاهیم ساده‌تر (معقولات اولی) می‌سازد. این مفاهیم به روابط بین اشیاء اشاره دارند و در دنیای واقعی به شکل مستقل وجود ندارند.

- اگر دو فعل و قیودشان با هم تشابه داشته باشند، بین آن‌ها خصوصیت تشابه و مثلث برقرار است.
- اگر ذات و قیود مأمور به و ذات و قیود ماتی‌به، به‌طور کامل تطابق داشتند و ماتی‌به، به‌وجود آمده بود، امثال محقق می‌شود.
- در فرض تراحم دو واجب انشائی، واجبی که مرجوح است، از فعلیت می‌افتد.
- امکان ندارد که یک شخص، در آن واحد، علم به ثبوت و علم به عدم ثبوت یک شیء داشته باشد.
- اگر یکی از دو عدل واجب تخییری، امثال شده باشد (موجود شده باشد)، عدل دیگر از فعلیت ساقط است.
- مکلف شرعی، مکلف بالغ عاقل است.
- مکلف شرعی است که مخاطب (پذیرنده) بازداشتن یا طلب کردن شارع است.
- عامل انجام فعل اثبات کردن، اخبار و انشاء، فقط انسان یا شخص حقوقی است.
- فعلی که مورد الزام شارع قرار گرفته، نمی‌تواند رخصت در ترک کردن از طرف شارع داشته باشد.
- هر تکلیف کردنی، یک حکم کردن است و حتماً مخاطب (پذیرنده) تکلیف کردن، جوهر قابل حیات است.
- هر تکلیف کردن، یا بازداشتن است، یا طلب کردن.
- هر الزام کردن یک تکلیف کردن است.
- رخصت دادن یک نوع حکم کردن است.
- هر فعل که مقید به قصد قربت است، فعل مقررانه است.
- مباح به معنای اعم، فعلی است که حرام نباشد و می‌تواند مستحب، واجب، مکروه یا مباح به معنای اخص باشد.
- مباح به معنای اخص و حرام و واجب و مستحب و مکروه قابل جمع شدن در یک مصداق نیستند.
- مباح به معنای اعمی که مورد رخصت شارع باشد، مباح به معنای اخص است.
- فعلی که مورد طلب شارع باشد و در عین حال مورد الزام شارع نباشد، مستحب است.
- فعلی که مورد طلب شارع باشد و در عین حال مورد الزام شارع هم باشد، واجب شرعی است.
- فعلی که مورد بازداشتن شارع باشد و در عین حال مورد الزام شارع نباشد، مکروه است.
- فعلی که مورد بازداشتن شارع باشد و در عین حال مورد الزام شارع هم باشد، حرام شرعی است.
- واجب یا غیری است یا نفسی، یا فعلی است یا غیرفعلی، یا تعینی است یا تخییری، یا عینی است یا کفایی، یا شرعی است یا عقلی، یا تعبدی است یا توصیلی.
- واجبی که بدل اختیاری دارد، واجب تخییری است.



۱) نجاشی آن شخص را به «عین» توصیف کرده است. ۲) «عین» به معنای مورد اعتماد است. این در حالی است که این دو گزاره، کاملاً دو واقعیت جدای از هم هستند. مهم‌تر آنکه گزاره اول جنبه حسّی پررنگ‌تر و گزاره دوم جنبه حدسی و اجتهادی جدی‌تری را دارا بوده و با تفکیک این دو

گزاره از یکدیگر، مدیریت دانش خصوصاً در زمان تغییر مبنا راحت تر و دقیق تر محقق می شود. به عبارت دیگر، اگر فردی دیگر خواست مبانی خود را در معانی الفاظ وارد کند، دیگر لازم نیست که تعبیری که هر راوی به آن ها متّصف شده را مجدداً دسته بندی و برچسب گذاری کند. برای این منظور، طرح زیر برای ثبت دانش موجود در کتاب رجال نجاشی، به جهت تفکیک لایه حسی توصیفات از لایه حدسی برداشت از توصیفات طراحی شد:



نمودار ۲- ساختار پایگاه داده در مدل سازی تفکیک لفظ از معنا

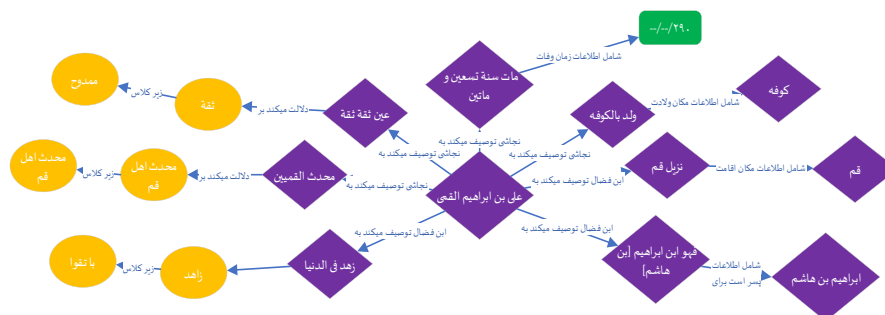
به طور مثال، نجاشی در مدخل «ابراهیم بن سلیمان بن ابی داحه»، او را به عبارت «وجه اصحابنا البصریین فی الفقه و الکلام و الادب و الشعر» توصیف کرده است. در نتیجه به عنوان یک واقعیت، این گزارشی که نجاشی بدون استناد به فرد دیگری مطرح کرده، از طرف خودش ذخیره می شود. از طرفی ممکن است این عبارت، مدالیل مختلفی را به همراه داشته باشد: مانند اینکه او از اصحاب امامیه بوده است؛ بصری بوده است؛ چهره شاخصی در فقه بوده است؛ چهره شاخصی در کلام بوده است و... هر کدام از این دلالت ها به صورت فوقانی برای مدلول اخص این لفظ (یعنی «چهره شاخص شیعیان بصره در فقه و کلام و ادب و شعر») ثبت می شوند. این تفکیک در ثبت دانش، علاوه بر اینکه لایه حسی و حدسی ثبت گزارش را از یکدیگر جدا می کند، انعطاف زیادی در ثبت مدالیل چندگانه برای الفاظ، تغییر مبنا و کشف جنبه های مختلف توصیفات خواهد داد. همچنین باتوجه به ثبت شخص توصیف کننده در قسمت ثبت گزارش اولیه، مواردی که نجاشی به استناد افراد دیگری همچون ابن فضال و... به توصیف راویان پرداخته، از مواردی که خودش صریحاً آنها را توصیف کرده، جدا می شود. بنابراین، با این سبک از مدل سازی اطلاعات، بسیاری از مشکلات ساختارهای متداول برطرف شده و زمینه انجام محاسبات و تحلیل های آماری پدید می آید.

به طور کلی و خلاصه می توان گفت، موجودیت ها (که شامل افراد، آثار و الفاظ توصیفی، نسخه، حدیث، سند، قبیله و مکان)، توسط افرادی (که خود این افراد داخل در موجودیت افراد هستند) به الفاظی (که آن هم موجودیت است)، توصیف می شوند. این الفاظ در مبانی مختلف و یا

۱) یا لفظ به مدلول اخص (معنا) ای دلالت می‌کند که این مدلول در شبکه‌ای از معانی تعریف شده و سبب می‌شود شخص توصیف‌شده به تمامی معانی فوقانی مدلول اخص توصیف شود (مانند مثال ذکر شده که سبب می‌شود ابراهیم بن سلیمان بن ابی داحه، هم در مذهب، شیعه شود و به تبع آن مسلمان، هم چهره شاخص گردد و به تبع آن چهره اجتماعی و هم...).

۲) یا دارای اطلاعات یک نسبت خاص میان موجودیت توصیف شده و یک موجودیت دیگر هستند؛ به این معنا که به واسطه این لفظ مکان و زمان ولادت، وفات و... و همچنین لقب، کنیه و نسبت‌های فامیلی و استاد-شاگردی را بین شخص توصیف شده با موارد مختلف برقرار می‌کند.

به طور نمونه می‌توان گراف دانش نمادینی برای علی بن ابراهیم القمی را که شامل هر دو نوع داده بوده، اینچنین به تصویر کشید:



نمودار ۳- گراف دانش موجودیت از نوع شخص (علی بن ابراهیم القمی)

تحلیل اطلاعات و استنباط داده جدید، یک فرایند است که در آن از داده‌ها و اطلاعات موجود استفاده می‌شود، تا به نتایج و اطلاعات جدیدی برسیم. در ادامه به برخی از جوانب مهم این فرایند اشاره می‌شود:

● متمایزسازی اطلاعات استنباط شده با اطلاعات پرداخت شده از متن: در این مرحله، اطلاعات استنباط شده از داده‌ها با اطلاعاتی که از متن یا منابع دیگر به دست آمده، متمایزسازی می‌شود. به عبارت دیگر، سعی گردد بین اطلاعاتی که از داده‌ها به دست آمده و اطلاعاتی که از متن خوانده شده، تمایز قائل شود.

● **استنباط به صورت همزمان و غیرهمزمان با ثبت دادگان جدید:** در تحلیل داده، استنباط می‌تواند به صورت همزمان یا غیرهمزمان انجام شود. در موارد ساده، استنباط همزمان ممکن است که در طول

فرایند تجزیه و تحلیل داده انجام شود، در حالی که در موارد پیچیده، استنباط غیرهمزمان ممکن است نیاز به ثبت و ذخیره داده‌های جدید داشته باشد.

● **استنباط مبنامحور است و تنها براساس مبانی انتخاب شده کاربر تحقق می‌یابد:** در استنباط داده، از مبانی و قوانینی که توسط کاربر تعیین می‌شود، استفاده می‌شود. با انتخاب مبانی مناسب و قوانین درست، استنباط داده جدید و مفید امکان‌پذیر است.

● **تناقض‌یابی میان دادگان:** در تحلیل داده، ممکن است با تناقض‌هایی در داده‌ها مواجه شد. این تناقض‌ها می‌توانند به صورت همزمان یا غیرهمزمان در داده‌ها ظاهر شوند. بررسی و تحلیل این تناقض‌ها می‌تواند به درک بهتر از داده‌ها و استنباط داده جدید کمک کند.

۴-۲. تلفیق مدل زبانی بزرگ مبتنی با گراف دانش

در بخش دوم الگوی پیشنهادی، تلفیق گراف دانش با ابزارهای پردازش زبان طبیعی بررسی می‌شود. هدف این تلفیق، استفاده از قدرت گراف دانش در سازمان‌دهی روابط بین مفاهیم و استفاده از داده‌های متنی برای بهبود فرایند تولید و استنتاج اطلاعات است. این روش می‌تواند به صورت خودکار به تحلیل و استنتاج متون پرداخته و به تولید متون جدید براساس داده‌های موجود کمک کند. در حوزه علوم اسلامی، این امکان فراهم می‌شود که مفاهیم و روابط پیچیده بین آیات قرآن، احادیث و منابع فقهی به‌طور دقیق تحلیل و استنتاج شود. این تلفیق می‌تواند به بهبود دقت و جامعیت در تحلیل داده‌های علمی و دینی منجر گردد.

جدول ۲- گام‌های کلی طراحی و توسعه مدل زبانی بزرگ علوم اسلامی- انسانی مبتنی بر گراف دانش

تکمیل کلان‌داده	جمع‌آوری متون و منابع علمی در حوزه علوم اسلامی- انسانی، مانند قرآن، حدیث، فقه و عرفان و...
	تبدیل هوشمند کلان‌داده صوتی به متن
	تبدیل کلان‌داده تصویر نسخ خطی به متن
	ذخیره SPINها در بانک کلان‌داده براساس ساختار استاندارد
پیش‌پردازش و استخراج ویژگی‌ها	تمیز کردن و پاک‌سازی متون از نقص‌ها و اشتباهات
	استخراج ویژگی‌های زبانی و معنایی از متون
	ایجاد نماینده‌های مناسب برای مفاهیم و مفاهیم
	استخراج SPINهای جدید از SPINهای موجود ^۱

۱. براساس SPINهای ایجاد شده در گراف دانش فقه، اصول و حدیث، SPINهای جدید از متون دیگر موجود در کلان‌داده استخراج می‌شود.

آموزش مدل	طراحی معماری مدل زبانی بزرگ مبتنی بر سه‌تایی‌های دانشی در حوزه علوم اسلامی- انسانی
	آموزش مدل با استفاده از داده‌های متنی و معنوی جمع‌آوری شده
	تنظیم پارامترها و وزن‌ها برای بهینه‌سازی عملکرد مدل
ارزیابی و بهینه‌سازی مدل	ارزیابی دقیق عملکرد مدل در حوزه علوم اسلامی- انسانی با استفاده از معیارهای مناسب
	بهبود عملکرد مدل با تنظیم پارامترها و وزن‌ها
	اعمال روش‌های بهبودی در مدل براساس نتایج ارزیابی
استفاده از مدل زبانی بزرگ برای تکمیل گراف دانش علوم اسلامی- انسانی	ارزیابی دقیق عملکرد مدل در حوزه علوم اسلامی- انسانی با استفاده از معیارهای مناسب
	استفاده از مدل برای تولید خروجی‌های متنی مرتبط با حوزه علوم اسلامی- انسانی
	پاسخ به سؤالات و مشکلات مربوط به حوزه علوم اسلامی- انسانی

۴-۲-۱. تکمیل کلان‌داده و پیش‌پردازش داده‌ها^۱

در این مرحله، ابتدا بایستی متون مرتبط با علوم اسلامی- انسانی جمع‌آوری شوند. این متون می‌توانند شامل قرآن، فلسفه، منطق، فقه و حقوق و سایر مراجع باشند. تلاش برای جمع‌آوری متون باکیفیت و متنوع از منابع معتبر و خبره در زمینه علوم اسلامی- انسانی نیز بسیار مهم است. پس از جمع‌آوری داده‌ها، بایستی فرایند پیش‌پردازش بر روی داده‌ها انجام شود. این امر شامل مراحل مانند تمیز کردن داده‌ها، حذف علائم نامربوط و توضیحات غیرضروری، تقسیم‌بندی به جملات و کلمات و استخراج ویژگی‌های زبانی است. همچنین در این مرحله از SPIN‌های تولید شده در گراف دانش علوم اسلامی- انسانی در سطح روابط و قواعد منطقی- استنتاجی استفاده می‌شود، تا به ارتباطات معنایی عمیق‌تر بین کلمات و عبارات، تبدیل داده‌ها و متون به شکلی منظم و قابل‌فهم‌تر و استفاده از قواعد استنتاجی و منطقی کمک کند. این امکانات باعث بهبود عملکرد و دقت مدل و درک بهتر متون در حوزه علوم اسلامی- انسانی می‌شوند.

۴-۲-۲. آموزش مدل

پس از پیش‌پردازش داده‌ها، بایستی مدل زبانی بزرگ آموزش داده شود. برای آموزش مدل، از روش‌های یادگیری عمیق مانند شبکه‌های عصبی بازگشتی^۲، شبکه‌های عصبی ترنسفرمر و یا معماری‌های پیشرفته‌تری مانند GPT می‌توان استفاده کرد. در این مرحله، مراقبت از بزرگ‌نمایی و تعمیم‌پذیری مدل در حوزه علوم اسلامی- انسانی بسیار حائز اهمیت است.

1. ASR, OCR, SPINs

2. Recurrent Neural Networks (RNN)

۴-۲-۳. ارزیابی و بهبود مدل

پس از آموزش مدل، بایستی عملکرد آن ارزیابی شود. از معیارهایی مانند درستی نتایج، پوشش متنوع موضوعات و نگرش‌ها، پرسش و پاسخ دقیق و منسجم و تولید متن طبیعی می‌توان برای ارزیابی استفاده کرد. در صورت نیاز، مدل بایستی بهبود داده شده و مجدداً آموزش داده شود.

۴-۲-۴. استفاده از مدل زبانی بزرگ برای تکمیل گراف دانش علوم اسلامی - انسانی و انتقال

یادگیری

پس از ارزیابی و بهبود، مدل زبانی بزرگ می‌تواند در برنامه‌ها و سامانه‌های مختلفی در حوزه علوم اسلامی - انسانی استفاده شود. این امر شامل سامانه‌های تفسیر و تحلیل متون دینی، پاسخگویی به سؤالات دینی، ترجمه متون مذهبی، سیستم‌های هوشمند برای پژوهش و آموزش در علوم اسلامی - انسانی و سایر برنامه‌های مرتبط است.

۵. بحث و بررسی

روند استفاده از گراف‌های دانش برای ذخیره، بازیابی و توسعه اطلاعات به سرعت در حال پیشی گرفتن از روش‌های مبتنی بر پایگاه‌های داده رابطه‌ای است. این تغییرات می‌تواند ایران را به یکی از مراکز پیشرو در حوزه وب معنایی تبدیل کند. در مقایسه با جداول و پایگاه‌های داده سنتی، گراف‌های دانش ابزار قدرتمندتری برای مدل‌سازی و مدیریت دانش فراهم می‌کنند. به عنوان نمونه، شرکت‌های بزرگ نظیر گوگل، از سال ۲۰۱۲ به توسعه گراف دانش پرداخته‌اند. گراف‌های دانش دارای کاربردهای متنوعی هستند؛ از جمله آن‌ها می‌توان به پاسخ‌گویی خودکار، سیستم‌های توصیه‌گر و تحلیل داده‌های تخصصی اشاره کرد. یکی از چالش‌های کلیدی در این زمینه، کشف اطلاعات پنهان و اعتبارسنجی داده‌ها در گراف دانش است. به دلیل پیچیدگی‌های متعدد در فرایند استنتاج از گراف دانش، نیاز به روش‌های منطقی مختلف برای مدیریت تضادها و تناقضات موجود در داده‌ها وجود دارد. این فرایند اعتبارسنجی می‌تواند از توقف فرایند استنتاج ماشینی در مواجهه با تضادها جلوگیری کند.

ترکیب گراف‌های دانش در سطح روابط و قواعد منطقی - استنتاجی با مدل‌های زبانی بزرگ می‌تواند به سیستم‌های پردازش زبان طبیعی قدرتمندتر و مؤثرتر منجر شود. در ادامه چند نمونه از تعامل و ترکیب گراف‌های دانش با مدل‌های زبانی بزرگ پیشنهاد می‌شود:

۱. **پیش‌آموزش با داده‌های گراف دانش:** تزریق دانش از یک گراف دانش به مرحله پیش‌آموزش یک مدل زبانی بزرگ می‌تواند به مدل کمک کند تا روابط بین موجودیت‌ها و ویژگی‌های آنها را بیاموزد. این کار را می‌توان با تبدیل سه گانه‌های گراف دانش به جملات زبان طبیعی و افزودن آنها به مجموعه

آموزشی انجام داد.

۲. تنظیم دقیق با اهداف مبتنی بر گراف: پس از پیش آموزش، تنظیم دقیق مدل برای وظایف خاص با اهداف مبتنی بر گراف می‌تواند به آن کمک کند تا با استفاده از دانش کدگذاری شده در گراف، استدلال و استنتاج را بیاموزد. به عنوان مثال، مدل را می‌توان برای پیش‌بینی موجودیت‌ها یا روابط گمشده در گراف تنظیم کرد که می‌تواند درک آن را از ساختار زیربنایی بهبود بخشد.

۳. پرس‌وجو از گراف‌های دانش استنتاج شده: زمانی که مدل با یک سؤال یا وظیفه‌ای مواجه می‌شود که به دانش دقیق و ساختاریافته‌ای نیاز دارد، می‌تواند برای به دست آوردن اطلاعات مرتبط، از گراف دانش پرس‌وجو کند. این کار را می‌توان با تبدیل و ترجمه پرس‌وجوی زبان طبیعی به یک پرس‌وجو مبتنی بر گراف (به عنوان مثال، با استفاده از SPARQL) و سپس استفاده از اطلاعات بازیابی شده برای پاسخ به سؤال انجام داد.

۴. تولید دانش مبتنی بر راهنمایی گراف: با ترکیب قابلیت‌های تولید مدل‌های زبانی بزرگ با اطلاعات ساختاریافته در گراف‌های دانش، می‌توان متون دقیق‌تر و مرتبط‌تری را تولید کرد. به عنوان مثال، هنگام ایجاد خلاصه یا توصیفی از یک موجودیت، مدل می‌تواند اطلاعات برگرفته از گراف دانش را اولویت‌بندی کند تا از صحت واقعی اطلاعات بدست آمده اطمینان حاصل شود.

۵. مدل‌های ترکیبی: توسعه مدل‌های ترکیبی که هر دو مؤلفه مبتنی بر گراف و مبتنی بر شبکه عصبی را دربرمی‌گیرد، می‌تواند از نقاط قوت هر دو رویکرد استفاده کند. به عنوان مثال، شبکه‌های عصبی گراف (GNN) می‌توانند برای یادگیری بازنمون‌های پنهان داده‌های ساختاریافته گراف مورد استفاده قرار بگیرند، تا با مدل‌های زبانی بزرگ برای وظایف پردازش زبان طبیعی ادغام شوند.

بهبود مدل‌های زبانی بزرگ برای سازگاری بهتر با گراف‌های دانش در سطح روابط و قواعد منطقی - استنتاجی می‌تواند منجر به بهبود عملکرد در کارهایی شود که نیاز به دانش ساختاریافته و استدلال دارند. در اینجا چند راهبرد برای سازگاری بیشتر مدل‌های زبانی بزرگ با گراف‌های دانش ذکر می‌شود:

۱. ترکیب داده‌های ساختاریافته در طول آموزش: همانطور که گذشت، بهتر است دانش برگرفته از گراف‌های دانش به مرحله پیش‌آموزش مدل زبانی بزرگ منتقل شود. این کار را می‌توان با تبدیل سه‌گانه گراف دانش به جملات زبان طبیعی یا تولید متن ترکیبی که حاوی اطلاعات گراف دانش است، انجام داد.

۲. طراحی معماری‌های مبتنی بر گراف: پیشنهاد می‌شود معماری مدل زبانی با تمرکز بر آگاهی بیشتر از ساختارهای گراف اصلاح شوند. این کار می‌تواند شامل یکپارچه‌سازی شبکه‌های عصبی

گراف یا دیگر لایه‌های مبتنی بر گراف برای پردازش اطلاعات ساختاریافته باشد.

۳. **تنظیم دقیق با وظایف مبتنی بر گراف:** باید مدل زبانی با وظایفی که مستلزم استدلال با گراف‌های دانش است، دقیق تنظیم شود. به عنوان مثال، مدل برای پیش‌بینی روابط یا موجودیت‌های گمشده در گراف، یا انجام وظایف پیش‌بینی پیوند و تفکیک موجودیت آموزش ببیند.

۴. **بهبود پیوند موجودیت و ابهام‌زدایی:** افزایش توانایی مدل برای شناسایی و پیوند موجودیت‌ها در متن به گره‌های مربوطه آن‌ها در گراف دانش، باید مورد توجه قرار گیرد. این هدف با ترکیب کردن تکنیک‌های پیشرفته پیوند موجودیت و ابهام‌زدایی در طول آموزش یا تنظیم دقیق، به دست خواهد آمد.

۵. **طراحی روش‌هایی با هدف استدلال مبتنی بر گراف:** الگوریتم‌ها و تکنیک‌هایی طراحی می‌شوند که مدل‌های زبانی بزرگ را قادر می‌سازد تا استدلال و استنتاج را با استفاده از ساختارهای گراف دانش ایجاد کنند. این کار می‌تواند شامل ترجمه پرس‌وجوهای زبان طبیعی به پرس‌وجوهای مبتنی بر گراف و نیز توسعه رویکردهایی برای ترکیب استدلال مبتنی بر گراف با استدلال مبتنی بر متن باشد. ترجمه متون خام زبان طبیعی به پرس‌وجوهای مبتنی بر زبان استنتاج وب معنایی و استفاده از آن به عنوان ورودی مدل زبانی بزرگ، مدل‌های زبانی بزرگ را قادر می‌سازد تا استدلال و استنتاج را با استفاده از ساختارهای گراف دانش ایجاد کنند.

۶. **ادغام تعبیه‌های گراف:** گنجانیدن تعبیه‌های موجودیت‌ها و روابط گراف دانش در مدل زبانی بزرگ ضروری است. این تعبیه‌ها می‌توانند با استفاده از شبکه‌های عصبی گراف یا سایر روش‌های یادگیری بازنمون گراف آموزش ببینند و برای غنی‌سازی درک مدل از موجودیت‌ها و روابط آنها مورد استفاده قرار گیرند.

۷. **تشویق تبیین‌پذیری و تفسیرپذیری:** پس از آنکه پیش آموزش مدل مبتنی بر کلان‌داده متنی و کلان‌داده استنتاجی انجام شد، آموزش مدل مبتنی بر گفتگوهای منطقی - استنتاجی با خبرگان علوم اسلامی - انسانی و دریافت بازخورد اصلاحی در خصوص استنتاج‌های نادرست، به مدل زبانی بزرگ اجازه می‌دهد تا بعد از اتمام فرایند آموزش، توضیحاتی منطقی را برای استدلال و پیش‌بینی‌های خود براساس وب معنایی ارائه کند. این هدف می‌تواند به کاربران کمک کند تا بفهمند مدل چگونه از اطلاعات ساختاریافته گراف دانش استفاده می‌کند و همچنین می‌تواند به اشکال‌زدایی و اصلاح مدل کمک نماید. ایجاد یک مدل زبانی بزرگ مبتنی بر گراف دانش که به‌طور خاص برای حوزه علوم اسلامی - انسانی طراحی شده باشد، امکان کنترل کامل بر معماری، پارامترها و اجزای مدل را فراهم می‌کند. این کنترل جامع به پژوهشگران اجازه می‌دهد تا مدل را به‌طور دقیق براساس نیازهای علوم اسلامی و انسانی طراحی و بهینه‌سازی کنند. چنین مدلی می‌تواند به‌طور خاص به پیچیدگی‌های

مفاهیم این علوم پاسخ دهد و از قابلیت‌های گراف دانش برای بهبود درک مفاهیم استفاده نماید. علاوه بر این، ایجاد یک مدل زبانی اختصاصی، انعطاف‌پذیری بیشتری را در طراحی ساختار و اجزای مدل به همراه خواهد داشت. محققان می‌توانند لایه‌ها و معماری مدل را مطابق با نیازهای خاص این علوم تنظیم کنند و مدل را به شکلی پویا بهبود بخشند. این انعطاف به بهبود مداوم عملکرد مدل و تطبیق آن با داده‌های جدید کمک می‌کند. همچنین، با ایجاد این مدل، امکان استفاده از داده‌های خاص و منحصر به فرد حوزه علوم اسلامی و انسانی نیز فراهم می‌شود. این ویژگی به مدل اجازه می‌دهد تا با داده‌هایی که به‌طور ویژه با مفاهیم دینی و علمی این حوزه در ارتباط هستند، آموزش ببیند و به سطح دقت و عملکرد بهتری دست یابد.

یکی دیگر از مزایای این مدل، قابلیت توسعه و بهبودهای بعدی است. با ایجاد این مدل، می‌توان به‌مرور زمان قابلیت‌های جدیدی به آن اضافه کرد و براساس نیازهای جدید آن را ارتقاء داد. این فرآیند منجر به تکامل و تقویت مدل خواهد شد، به‌گونه‌ای که در طول زمان بتواند به بهترین نحو پاسخگوی نیازهای پژوهشی و تحلیلی باشد. از سوی دیگر، ساخت یک مدل زبانی بزرگ مختص این حوزه باعث کاهش وابستگی به مدل‌های پیشین می‌شود. در برخی موارد، مدل‌های موجود ممکن است در پردازش مفاهیم خاص علوم اسلامی و انسانی دچار نقص باشند و نتوانند به‌درستی آن‌ها را تفسیر و تحلیل کنند. با ایجاد یک مدل اختصاصی، این مشکل رفع می‌شود و پژوهشگران می‌توانند مدلی داشته باشند که به‌طور کامل نیازهای مفهومی و معنایی علوم اسلامی و انسانی را پوشش دهد. در نهایت، باید به این نکته اشاره کرد که مدل‌های پیشین اغلب ناسازگاری‌هایی با مفاهیم و اصطلاحات خاص علوم اسلامی و انسانی دارند و همین امر می‌تواند منجر به کاهش دقت و کارایی آن‌ها در این حوزه شود. همچنین، انتقال یادگیری ناکافی از این مدل‌های پیشین نیز می‌تواند مانع از بهبود قابل توجه عملکرد مدل‌ها در حوزه علوم اسلامی و انسانی شود. از این‌رو، ایجاد یک مدل اختصاصی و بومی، بهترین راه‌حل برای مواجهه با این چالش‌هاست و می‌تواند به بهبود دقت، جامعیت و صحت پاسخ‌های مربوط به علوم اسلامی و انسانی منجر شود.

<http://stjm.gom.ac.ir>

۶. نتیجه‌گیری

در این پژوهش، به ارائه یک الگوی پیشنهادی برای تلفیق گراف دانش و مدل‌های زبانی بزرگ در حوزه‌های علوم اسلامی - انسانی پرداخته شد. با توجه به پیچیدگی مفاهیم علمی و دینی و نیاز به سازمان‌دهی معنادار این داده‌ها، استفاده از گراف دانش به‌عنوان ابزاری برای نمایش سلسله‌مراتبی و ارتباط‌دهی مفاهیم، در کنار مدل‌های زبانی بزرگ که قابلیت تولید متون و استنتاج خودکار را دارند، ضروری به نظر می‌رسد. بررسی منابع و پیشینه تحقیقاتی نشان داد که علی‌رغم پیشرفت‌های گسترده

در هر یک از این حوزه‌ها به‌طور جداگانه، شکاف‌هایی در ترکیب این دو رویکرد در علوم اسلامی- انسانی وجود دارد که به‌طور مؤثر پاسخ داده نشده‌اند. الگوی پیشنهادی پژوهش حاضر، چارچوبی را فراهم می‌کند که از گراف دانش برای ساختاردهی و سازمان‌دهی مفاهیم استفاده کرده و از مدل‌های زبانی بزرگ برای پاسخ‌دهی خودکار به سؤالات و تولید متون مرتبط بهره می‌گیرد. این الگو، با تمرکز بر استنتاج مفهومی و تحلیل سلسله‌مراتبی داده‌ها، قادر است به‌عنوان یک ابزار نظری برای توسعه سیستم‌های هوشمند در حوزه علوم اسلامی- انسانی عمل کند. همچنین این الگو می‌تواند به پژوهشگران و توسعه‌دهندگان کمک کند تا از یک رویکرد کارآمد و ساختاریافته برای مدیریت و تحلیل داده‌های مفهومی در این حوزه‌ها استفاده کنند. مزیت رقابتی مدل زبانی بزرگ پیشنهادی شامل ترکیب گراف‌های دانش در سطح روابط و قواعد منطقی- استنتاجی با مدل‌های زبانی بزرگ و بهبود سیستم‌های گراف دانش با مدل‌های زبانی بزرگ است.

پژوهش حاضر با ارائه این الگو، گام مهمی در جهت تلفیق مدل‌های زبانی و گراف دانش در حوزه‌های علوم اسلامی- انسانی برداشته است؛ اما همچنان مسیر برای تحقیقات و توسعه‌های بیشتر باز است. برای پژوهش‌های آینده، توصیه می‌شود که این الگو به‌صورت عملیاتی پیاده‌سازی شده و کارایی آن در محیط‌های واقعی مورد ارزیابی قرار گیرد. علاوه بر این، بررسی چگونگی به‌کارگیری این الگو در سایر حوزه‌های علمی و فرهنگی می‌تواند به تعمیق و بسط این رویکرد کمک کند. همچنین امکان توسعه روش‌های استنتاج دقیق‌تر و ارتقای قابلیت پاسخ‌دهی خودکار به به‌کارگیری تکنیک‌های جدید در گراف دانش و مدل‌های زبانی بزرگ می‌تواند از محورهای پژوهشی آتی باشد. در نهایت، این الگو می‌تواند به‌عنوان یک مبنای نظری برای توسعه سیستم‌های پیشرفته‌تر در حوزه‌های اسلامی- انسانی به‌کار گرفته شود و به افزایش دقت، کارایی و جامعیت در تحلیل و استنتاج‌های مرتبط کمک شایانی نماید.

<http://stun.gom.ac.ir>

باتوجه به الگوی ارائه شده در این پژوهش و پیشرفت‌های اخیر در حوزه گراف دانش و مدل‌های زبانی بزرگ، فرصت‌های متعددی برای ادامه تحقیقات در این زمینه وجود دارد. در ادامه چند پیشنهاد برای پژوهش‌های آتی ارائه می‌شود:

● **پیاده‌سازی عملی الگوی پیشنهادی:** یکی از گام‌های مهم پس از ارائه الگو، ارزیابی عملی آن در محیط‌های واقعی است. پیشنهاد می‌شود پژوهشگران در آینده این الگو را در سیستم‌های هوشمند مدیریت دانش، به‌ویژه در حوزه‌های علوم اسلامی- انسانی، به‌صورت عملی پیاده‌سازی کرده و کارایی آن را در کاربردهای واقعی بررسی کنند. این ارزیابی می‌تواند شامل تحلیل دقت استنتاج‌ها و قابلیت پاسخ‌دهی خودکار به سؤالات علمی و دینی باشد.

● **توسعه گراف دانش در حوزه‌های تخصصی:** درحالی‌که این پژوهش بیشتر بر روی علوم اسلامی- انسانی تمرکز دارد، می‌توان این الگو را به سایر حوزه‌های تخصصی مانند فلسفه، ادبیات و تاریخ نیز گسترش داد. پیشنهاد می‌شود پژوهش‌های آینده به توسعه گراف‌های دانش در این حوزه‌ها بپردازند و از این الگو برای پیوند دادن بین مفاهیم تخصصی استفاده کنند.

● **بهبود روش‌های استنتاج:** در الگوی پیشنهادی، از گراف دانش برای ایجاد روابط معنایی بین مفاهیم استفاده می‌شود. با این حال، امکان بهبود روش‌های استنتاج و ارتقای دقت نتایج از طریق ترکیب تکنیک‌های پیشرفته‌تر مانند استدلال مبتنی بر منطق‌های غیرکلاسیک و الگوریتم‌های استنتاج معنایی وجود دارد. پژوهش‌های آینده می‌توانند به بررسی و به‌کارگیری این روش‌ها در چارچوب‌های ارائه شده بپردازند.

● **بررسی تعامل مدل‌های زبانی بزرگ و گراف دانش در زبان‌های مختلف:** با توجه به اهمیت تحلیل‌های چندزبانه، پیشنهاد می‌شود پژوهش‌های آتی به بررسی تعامل بین گراف دانش و مدل‌های زبانی بزرگ در زبان‌های مختلف، به‌ویژه زبان‌های کمتر پرداخته شده مانند فارسی و عربی بپردازند. این بررسی می‌تواند به گسترش الگو در زمینه‌های جدید و افزایش قابلیت‌های آن در حوزه‌های زبانی مختلف منجر شود.

● **استفاده از تکنیک‌های جدید در پردازش زبان طبیعی:** با پیشرفت سریع در تکنیک‌های پردازش زبان طبیعی، می‌توان از این تکنیک‌ها برای بهبود فرایندهای استنتاج و تولید متن در گراف دانش استفاده کرد. پژوهش‌های آینده می‌توانند به بررسی استفاده از مدل‌های زبانی جدیدتر و پیشرفته‌تر برای ارتقای کارایی الگوی پیشنهادی بپردازند.

● **کاربرد الگو در سیستم‌های آموزشی و یادگیری:** استفاده از گراف دانش و مدل‌های زبانی بزرگ در سیستم‌های آموزشی می‌تواند امکان یادگیری هوشمند و فراهم‌سازی محتوای آموزشی خودکار را تسهیل کند. پیشنهاد می‌شود پژوهشگران به بررسی کاربرد این الگو در سیستم‌های آموزش الکترونیکی و پلتفرم‌های یادگیری تطبیقی بپردازند.

References

- Abaskohi, A., Baruni, S., Masoudi, M., Abbasi, N., Babalou, Mh., Edalat, A., Kamahi, S. & et al. (2024). Benchmarking Large Language Models for Persian: A Preliminary Study Focusing on ChatGPT. *arXiv*. <https://doi.org/10.48550/arXiv.2404.02403>
- Binz, M. & Schulz, E. (2023). Turning large language models into cognitive models. *arXiv*. <https://doi.org/10.48550/arXiv.2306.03917>
- Chen, Z., Singh, A.K. & Sra, M. (2023). LMExplainer: a Knowledge-Enhanced Explainer for Language Models. *arXiv*. <https://doi.org/10.48550/arXiv.2303.16537>
- Confalonieri, R., Coba, L., Wagner, B. & Besold, T.R. (2020). A historical perspective of explainable Artificial Intelligence. *WIREs Data Mining Know Discov*, no. 11. <https://doi.org/10.1002/widm.1391>
- Huang, J. & Chang, K.C.-C. (2023a). Towards Reasoning in Large Language Models: A Survey. *arXiv*. <https://doi.org/10.48550/arXiv.2212.10403>
- Huang, S., Mamidanna, S., Jangam, S., Zhou, Y. & Gilpain, L.H. (2023b). Can Large Language Models Explain Themselves? A Study of LLM-Generated Self-Explanations. *arXiv*. <https://doi.org/10.48550/arXiv.2310.11207>
- Luo, H. & Specia, L. (2024). From Understanding to Utilization: A Survey on Explainability for Large Language Models. *arXiv*. <https://doi.org/10.48550/arXiv.2401.12874>
- Mavrepis, P., Makridis, G., Fatouros, G., Koukos, V., Separdani, M.M. & Kyriazis, D. (2024). XAI for All: Can Large Language Models Simplify Explainable AI? *arXiv preprint arXiv: 2401.11310*.
- Patel, S., Kane, H. & Patel, R. (2023). Building Domain-Specific LLMs Faithful To The Islamic Worldview: Mirage or Technical Possibility?. *arXiv*. <https://doi.org/10.48550/arXiv.2312.06652>
- Rasti Meymandi, A., Hosseini, Z., Davari, S., Moshiri, A., Rahimi Golkhandan, S., Namdar, K., Feizi, N., Tavakoli-Targhi, M. & Khalvati, F. (2024). Opportunities for Persian Digital Humanities Research with Artificial Intelligence Language Models; Case Study: Forough Farokhzadeh. *arXiv*. <https://doi.org/10.48550/arXiv.2405.06760>
- Shu, D., Chen, T., Jin, M., Zhang, Y., Zhang, C., Du, M. & Zhang, Y. (2024). Knowledge Graph Large Language Model (KG-LLM) for Link Prediction. *arXiv*. <https://doi.org/10.48550/arXiv.2403.07311>
- Sullivan, D. (2020). *A reintroduction to our Knowledge Graph and knowledge panels*. URL= <https://blog.google/products/search/about-knowledge-graph-and-knowledge-panels>
- Zhao, H., Chen, H., Yang, F., Liu, N., Deng, H., Cia, H., Wang, S. & Yin, D. (2024). Explainability for Large Language Models: A Survey. *ACM Transactions on Intelligent Systems and Technology*, 15(2). <https://doi.org/10.1145/3639372>
- Zheng, Y., Koh, H.Y., Ju, J., Nguyen, A.T.N., May, L.T., Webb, G.L. & Pan, S. (2023). Large Language Models for Scientific Synthesis, Inference and Explanation. *arXiv*. <https://doi.org/10.48550/arXiv.2310.07984>