

Explainable AI for Health Care based Retrieval System

Lakshmi Reghu¹, Gayathri Ashok² and Remya RK Menon³

¹⁻³Department of Computer Science and Applications, Amrita School of Computing, Amrita Vishwa Vidyapeetham
Amritapuri
India

Email: am.en.p2mca22023@am. students.amrita.edu, am.en.p2mca22060@am.students.amrita.edu, ramya@am.amrita.edu

Abstract—Explainable Artificial Intelligence or (XAI) be-coming more popular in machine learning field because it allows the users to interpret most clear result and also provide public trust and facilitate their development and deployment of AI systems. Encouraging AI models to be more transparent and understandable is the aim of the burgeoning field of explainable AI. Here, we cover a varied scope of XAI for a healthcare-based retrieval system. When a user queries a healthcare-related document, XAI may explain why it was retrieved. Patients who make use of the XAI tools will be better able to comprehend the advice and be more inclined to believe it. The article also covers several additional uses for XAI models, including manufacturing and banking. Within the healthcare industry, medical diagnostic systems' predictions may also be defined by using explainable AI models. This is crucial for applications in the healthcare industry since it's required to comprehend the reasons behind predictions made by AI.

Index Terms— XAI, LIME, SHAP, XIR, BEIR, Knowledge Graph.

I. INTRODUCTION

In recent years, Explainable Artificial Intelligence (XAI) has attracted a lot of interest and its goal is to improve transparency and comprehensibility of AI models. This is significant because the usage of AI models in high-stakes fields like criminal justice, healthcare, and finance is growing. It is challenging to comprehend how these models decide and whether they are impartial and fair in the absence of XAI.

Healthcare could undergo a transformation thanks to artificial intelligence (AI), which could lead to more dependable and precise diagnosis and treatments. However, a number of significant ethical and societal issues are brought up by the application of AI in healthcare. The requirement for AI systems to be dependable and explicable ranks among the most crucial issues [13]. Additionally, they ought to be accountable and open so that patients and healthcare professionals may comprehend the decision-making process. Explainable AI systems consist of those that can explain their selections in detail. This is important as it helps medical professionals to make well-informed judgments about patient care and to have faith in the AI system's output [12]. It aims to provide transparency and human-understandability of machine learning models. The goal of XAI approaches is to shed light on these models' decision-making processes so that users may comprehend the logic underlying the models' predictions and results. Building trust in AI systems and guaranteeing their responsible and ethical use depend heavily on this transparency. XAI can be applied to enhance healthcare recommendation systems based on AI's explainability. A higher level of explainability may result in more patients having faith in these systems. Better healthcare output can result from increased patient trust.

While global models explain a machine learning model's overall behavior, local models explain the predictions of a machine learning model for a single data point. Medical diagnostic systems' decisions can be explained using XAI models in the healthcare industry. XAI models can be used in finance to explain loan application choices. XAI models can be used in manufacturing to explain the choices made by quality control systems. The usage of AI models in high-stake fields like criminal justice, banking, and healthcare is growing. It is challenging to comprehend how these models decide and whether they are impartial and fair in the absence of XAI.

Ranking in Information Retrieval(Fig.2): A increasing amount of characteristics are being merged by web search engines with learning-to-rank algorithms. Still, there are a number of useful issues about how ranking education could be implemented. For example, a sufficiently large sample of documents is employed, such that the learnt model reranks the sample, placing the relevant documents at the front. Information retrieval (IR) is paying more and more attention to learning to rank (Liu 2009), where machine learning techniques are applied to learn an optimal combination of characteristics into an efficient ranking model [9]. Using methods from machine learning to create a learnt model that combines several document attributes for an information retrieval system is known as "learning to rank" [8]. Web search engines frequently utilise learning to rank approaches in order to integrate different content weighting models with other query-independent information. A taught neural network or a sequence of regression trees may be represented by the model created using a particular learning to rank approach. For others, it may simply be a vector of weights for linearly combining each feature. The following are the typical stages for applying a learning to rank approach to create a learnt model, regardless of the used technique:

1. Pooling: A pooling process is used to identify documents for each query in a training set for which human relevance evaluations are acquired. To achieve a high-quality pool, a variety of retrieval mechanisms for each query might be combined.
2. Top K Retrieval: Create a sample of documents using an initial retrieval strategy for a set of training queries.
3. Feature extraction: Extract a vector of feature values for every document in the sample. A feature is a number or binary indication that shows how well a document matches the query or how good it is.
4. Learning: Use a learning to rank strategy to acquire knowledge about a model. Learning a rank problem in information retrieval involves two stages [11]. In the first stage, known as the training phase, a set of annotated training data is used to teach us a ranking function. The second part is the test phase, in which a collection of texts is ranked using the learnt ranking function. A search engine retrieves these documents in response to a user query. A learning algorithm is given a set of queries and the retrieved documents that correspond to them during the training phase. The documents are then given labels indicating the degree of relevance that the documents have for the related question. Human annotators assign the relevance findings. An annotator may annotate a collection of documents by giving each one a rating score based on how relevant it is to the query. A document that has a higher ranking score is deemed more relevant to the user query and should be ranked highest. Building a ranking model, such as a ranking function, with the goal of achieving the best agreement with the ranking produced by the scores awarded by the human annotators is the aim of learning.

II. LITERATURE REVIEW

A. Explainable online health information truthfulness in Consumer Health Search

These days, more and more people rely on internet health resources to help them make decisions that might affect their emotional and physical health. The majority of current literature solutions, according to the author, tackle the problem as a binary classification task, differentiating between accurate and incorrect information [14]. These techniques use machine learning or knowledge-based approaches. These solutions provide a number of issues with user decision-making, including: Users should take for granted that there are just two predetermined options available to them regarding the veracity of the information in the binary classification task. (ii) The methods used to obtain the results are frequently opaque, and the results themselves require little to no interpretation. We tackle the problem as an ad-hoc retrieval job instead of a classification task to overcome these problems, specifically referring to the Consumer Health Search task. To achieve this, a ranked list of both honest and topically-relevant texts is obtained by applying a previously suggested Information Retrieval model, which takes information truthfulness into account as a dimension of relevance. This research is unusual in that it extends a model of that kind and provides an explanation for the outcomes by using a knowledge base of scientific data in the form of medical journal articles. This study examines the

”explained” ranked list of documents by looking at it both subjectively via a user study and quantitatively using a typical classification task.

B. Explainable Information Retrieval: A Survey

The aim of explainable information retrieval (XIR) is to Enhancing the usability and transparency of information retrieval (IR) systems. This is noteworthy since IR systems are increasingly being used in high-stakes industries like law and medicine. Users are able to assess the validity of the documents they have gotten and discover the reason behind the retrieval of a particular document using XIR. This study examines many techniques for developing transparent and reliable information retrieval systems [1] Below is the summary of the many approaches examined in the paper:

- * Feature Attribution Methods: These methods aim to pinpoint the specific features (terms or characteristics) in a document that contribute to its relevance for a particular query. LIME (Local Interpretable Model-agnostic Explanations) is used for this. LIME generates explanations by creating a simpler, interpretable model around a specific document and analyzing its weights to understand feature importance.

- * Rule-based Explanation Methods: These methods rely on pre-defined rules that encode knowledge about how information retrieval systems work. They translate these rules into explanations for users.

- * Contrastive Explanation Methods: This approach focuses on explaining the difference in ranking between two documents. It highlights the features that make one document more relevant than the other for the given query.

- * Model-Agnostic Meta-learning Explanation Methods: This category leverages meta-learning techniques to learn explanation models from various information retrieval systems. This allows for explanations to be adapted to different retrieval models.

- * Shapley Additive exPlanations (SHAP): Based on each feature’s influence on the model’s prediction (document ranking), SHAP provides a contribution score to it. This makes it easier to comprehend how each feature affects a document’s rating.

- * Gradient-based Explanation Methods: These methods analyze the gradients of the ranking model to identify features with the most significant influence on the ranking of a particular document.

According to the paper, one key aspect of Explainable Information Retrieval EIR is explaining why a document is relevant to a specific query [10]. The authors discuss an explainable search system called EXS (Explainable Search System) that addresses this very concern. EXS provides explanations to users through a technique called feature attribution. Feature attribution helps pinpoint the factors that contribute to a document’s relevance for a given query. Essentially, it breaks down the decision-making process of the search engine and highlights which features (terms, phrases, or other characteristics) in the document were most influential in ranking it as relevant. So, if you use EXS and it retrieves a document about ”birdwatching in California,” the explanation might reveal that the document was ranked highly because it contained phrases like ”best birdwatching locations,” ”California birdwatching guide,” and ”types of birds in California.” This way, you understand not only that the document is relevant, but also why the search engine considers it relevant based on specific content within the document.

C. Learning to Rank for Information Retrieval

The goal of learning to rank for information retrieval (IR) is to automatically build a ranking model with training data that can be used to rank new items based on their relevance, significance, or preference. By their very nature, many IR challenges are ranking problems, and learning-to-rank approaches have the potential to improve many IR systems. This paper’s goal is to provide an overview of this line of inquiry. To be more precise, the current learning-to-rank algorithms are examined and divided into three groups: listwise, pairwise, and pointwise. Each approach’s benefits and drawbacks are examined, and the connections between these techniques’ loss functions and IR assessment metrics are explored. Subsequently, a statistical ranking theory is presented that can be utilised to analyse the query level generalisation abilities of various learning-to-rank algorithms and to explain them. Numerous learning-to-rank algorithms exist. They categorise the algorithms into three approaches—the pointwise approach, the pairwise approach, and the listwise method—in accordance with the four ML pillars in order to better comprehend them [7].

D. The survey of Explainable AI techniques in Healthcare

In addition to providing guidance for creating better deep learning model interpretations utilising XAI principles in medical picture and text analysis, the study surveys explainable AI (XAI) approaches utilised in healthcare and focuses on difficult XAI difficulties in medical applications. See [2]. The use of both computer- and human-centered assessment techniques is mentioned in the study, which emphasises how crucial it is to assess the calibre of explanations offered by XAI models. By exposing the decision-making processes behind deep learning

black-box models, XAI seeks to replicate human judgement and interpretation abilities in AI algorithms. The lack of transparency in deep learning-based artificial intelligence (AI) results in a "black-box" gap between AI models and human comprehension. In order to attain credibility, causality, transferability, confidence, fairness, accessibility, and interaction, XAI models offer insights into the prediction process.

E. Explainable AI (XAI): Core Ideas, Techniques, and Solutions

In addition to providing an overview of cutting-edge programming methods and methodologies for explainable artificial intelligence (XAI), the article addresses the significance of XAI in fostering trust and implementing AI systems [5]. It compares methodologies, software toolkits, and programming frameworks to assist stakeholders in choosing the best options for creating XAI-enabled apps using a certain taxonomy.

In order to put these XAI ideas into practice, the article also provides software toolkits and programming frameworks.

III. OTHER METHODS AND ALGORITHMS FOR EXPLAINABILITY

A. Local Interpretable Model-agnostic Explanations (LIME)

LIME (Local Interpretable Model-agnostic Explanations) is a novel explanation technique that learns an interpretable model that rotates locally around the forecast in order to authentically and interpretably explain any classifier's prediction. LIME offers local explanation by locally replacing complex models with simpler models or by predicting a Convolutional Neural Network (CNN) with a linear model. Modifications to the provided data result in modifications to the complex model's output. LIME uses feature engineering to make a prediction more understandable in order to ensure every person can compare and improve an inaccurate model. The ideal model explanation is comprehensible, model-agnostic, local-fidelity, and global-perspective. LIME illustrates how different features impact the model's forecast for the chosen data point [11]. Values for each attribute indicate how much it contributed to the forecast; the characteristics may be listed in a prioritised list. LIME can be used, for instance, to make a decision tree from the original model human-understandable. This tree helps demonstrate the sequence of decisions the model made to reach its prediction based on the specific features of the data point. LIME makes use of the simpler model to learn the relationship between the perturbed input data and the output change. A weight depending on the degree of similarity between the perturbed and original inputs is used to ensure that explanations from simple models with significantly perturbed data have less of an effect on the final explanation. LIME has been used for the evaluation of medical pictures in several research. It uses LIME, for example, to determine whether areas of gastric endoscopic images were bloody. The LIME classifier with the best real-world performance may be selected by non-experts. Moreover, they obtain a much improved version of an unstable classifier trained on 20 newsgroups by using LIME for feature engineering. Any explanation that is comprehensible must have a human-friendly representation, no matter what components the model uses. For categorization of text, an example of a possibly interpretable format is a binary vector that indicates whether a word is present or absent, even if word embeddings and other advanced features may be used by the classifier.

B. SHapley Additive exPlanations (SHAP)

Any model based on machine learning may be used with SHAP, a post-hoc method. Game theory, which establishes the amount that each player contributes to the payoff, forms its basis. SHAP offers details on how certain attributes influence forecasts. Whereas a "low" Shapley value shows the relevant information has a comparatively smaller impact on the prediction, a "high" Shapley value indicates a higher influence. It is essential to bear in mind that the exact interpretation of "low" and "high" might change based on the dataset, context, and issue under study. The relative value of certain features may vary depending on the domain and particular machine learning technique being used. The understanding of model functioning and the discovery of important details are facilitated by SHAP Values.

The size of feature attributions is the basis for the formulation of SHAP [6]. Traditional techniques, such as permutation, on the other hand, rely on the decline in model performance. Furthermore, while useful, the traditional feature significance plot only provides information relevant to the feature. Additionally, because of the axiomatic assumptions of SHAP, SHAP explanations can yield more trustworthy findings than other conventional metrics like the Gini index. Deep explainers are required due to the growing trend in ML models' complexity. The SHAP, a cooperative gaming approach, offers a thorough analysis of each input's contribution. Information to produce the finished product. It can accurately explain to decision-makers and end users how the models forecast data and how best to use it for the landslide susceptibility application. The SHAP

plot analyzed individual plots to emphasize the function for predictor (input datasets) and their accessibility, and it demonstrated how predictors and their internal relationship supported the output predictions [4]. The SHAP summary map shows both positive and negative interactions between the predictors and the objective (target variable), which is a key distinction from a standard feature importance plot. Since the summary graphic is made up of data points, it seems dotted. An alternative to conventional feature importance (such as permutation, Gini index, etc.) is SHAP feature importance. These importance metrics differ significantly in that the SHAP is determined by the magnitude of feature attributions. Traditional techniques, such as permutation, on the other hand, rely on the decline in model performance. Furthermore, although useful, the traditional feature significance plot just shows the value of the feature. Additionally, because of SHAP's inherent properties, worldwide SHAP explanations can yield more trustworthy conclusions than other traditional metrics like the Gini index. axiomatic presumptions. Three main advantages using SHAP over the conventional feature set are briefly mentioned below:

First: The SHAP values allow for the display of the relative contributions of every single predictor, both positively and negatively.

Second: The SHAP values are unique to each observation. This significantly increases its interpretability locally (transparency). On the other hand, the conventional variable significance methodologies only show the outcomes throughout the population, not on each unique or specific instance.

Third: While other systems use logistic regression or linear regression algorithms as substitute models, SHAP offers a quick implementation and allows values to be found for any tree-based model.

IV. PROPOSED METHODOLOGY

A. Architectural Diagram for an XAI-based model

A document retrieval system's block diagram is shown in Fig.1. An application that assists users in locating certain documents or information within a large document collection is known as a document retrieval system. The user query is initially parsed by the system before it functions. This indicates that the user's query is divided into smaller segments by the system, such keywords and phrases. After that, the system searches an index of documents (beir index) in the collection using these terms and phrases. A data structure called an index aids the system in finding documents related to the user's query rapidly. BEIR isn't just a test, it's a toolbox for IR researchers and developers. It focuses on evaluating how well new information retrieval systems that use neural networks (NIR models) perform on various real-world tasks. BEIR is different from older tests because it covers a wide range of tasks, giving a more complete picture of how well an NIR model works. BEIR also makes it easy to assess these models by providing a user-friendly system. This system lets researchers see how their NIR models perform on different tasks within BEIR, allowing for a fair comparison of different models and approaches to find the best solution for a specific problem. A list of terms or phrases that are used in the documents as well as a list of the papers that include each term or phrase are usually included in the index. After identifying a collection of papers that may be relevant to the user's query, the system assigns a ranking to the documents according to how relevant they are likely to be. Finally, a collection of the most relevant documents is returned to the user with explainability feature.

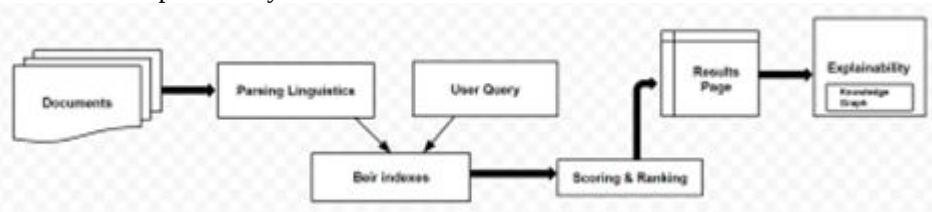


Fig. 1. Architectural Diagram for an XAI-based model

B. Knowledge Graph Creation

An organised representation of information that shows the connections between various components is called a knowledge graph. In essence, it's a graph database in which nodes represent entities and edges reflect the interactions between them. Knowledge graphs are a useful tool for simulating human comprehension through the organisation and connection of data, which facilitates the interpretation and extraction of insights from complicated data. *Entities: Entities are concepts, things, or occurrences that exist in the actual world. They might be anything, including individuals, groups, locations, occasions, and abstract notions like concepts or categories. *Relationship: Relationships are the relationships that exist between different things. They give the

data context and specify the relationships between the items. Diverse forms and trajectories of relationships enable complex depictions of linkages.

- *Properties: Properties define the traits or qualities of things and connections. They give more details on the relationships or entities in the graph.

- * Graph Structure: Knowledge graphs are shown as graphs with nodes (entities) and edges (relationships) as their constituent parts. The graph structure makes it possible to query and traverse the data efficiently, which makes it easier to explore related information.

- *Semantic Meaning: Semantic meaning is typically incorporated into the representations of knowledge graphs. In order to enable robots to comprehend and make sense of the data, entities, relationships, and characteristics must be precisely described with semantics.

- *Knowledge graphs possess the potential to handle substantial amounts of data and may be expanded to incorporate additional entities, connections, and characteristics as the knowledge base expands.

- *Use Cases: Knowledge graphs are applicable in a number of fields, such as but not restricted to.

- *Semantic search: Enhancing search engines to deliver more contextual and relevant search results.

- *Recommendation systems: Producing tailored suggestions according to user choices and item associations.

- * Natural Language understanding:Using semantic link-ages, natural language understanding helps robots comprehend and interpret human discourse.

- * Data Integration :Integrating and harmonising data from several sources to present a cohesive picture of the information is known as data integration.

- *Decision Support: Offering organised analysis and suggestions to help in decision-making processes.

In the first step we read a dataset from a fixed-width file into a Pandas DataFrame using a function. Then we take the query results as input and extract entities from it using spaCy's natural language processing capabilities. Then we create a list of entity pairs and extract relations from it. Finally, we visualize a directed graph based on these relations.

Recently, there has been a lot of interest in representation learning approaches based on deep learning. They rapidly compute the semantic links between relationships and objects in low-dimensional domains, effectively resolving the problem of data sparseness and significantly improving overall accuracy and efficiency. Expert systems developed in the late 1960s gave rise to knowledge graphs, which are today recognised for an intelligent system capable of large-scale data and knowledge integration. Knowledge graphs are a kind of semantic network, which is the primary difference between them and regular graph data. Reasoning within the constraints of the graph network topology will result in a large loss of crucial information and fail to achieve the intended completeness effect. To improve the KGC model's learning capacity and ultimately attain higher prediction and completion accuracy, more extensive characteristics might be retrieved. This would include considering the semantic information concealed inside the semantic network, such as the many kinds and characteristics of entities and their associations, the semantic linkages between entities, the knowledge about domain rules, and other multi-source data. Knowledge graph reasoning has been extensively investigated and applied in several knowledge-driven tasks since KG's invention, which greatly enhances their performance. With its special benefits of expressing and managing vast knowledge, Knowledge Graph (KG) offers high-quality structured knowledge for a variety of downstream knowledge-aware applications (including recommendation and intelligent question answering). The efficacy of the downstream tasks is mostly determined on the completeness and quality of KGs. Knowledge Graphs (KGs) offer a more efficient way to store, organise, and comprehend the vast amount of knowledge available in the world by describing concepts, entities, and their relationships in an organised triple form. In many knowledge-aware jobs, KG has become more and more significant in recent years. It particularly enhances intelligent question responding, information extraction, and other artificial intelligence tasks [3].

V. RESULT ANALYSIS

Natural language processing, or NLP, is a fascinating area of artificial intelligence (AI) that studies how language

is understood and used by computers. NLP utilizes a combination of techniques from various fields like computer science, linguistics, and machine learning. NLP acts as a bridge between unstructured text data and a structured knowledge graph. NLP acts like a detective sifting through textual evidence to uncover the who, what, and how of a story. For that first we should recognise the entities. Once the entities are identified, NLP goes a step further to understand how these entities are connected. Techniques like dependency parsing analyze the

grammatical structure of the sentence to figure out the relationships between the entities. Then we build the Knowledge Graph.

Rank	Document ID	Title	Body
0	0	539 Hypoglycemia increases the risk of dementia.	NaN
1	1	13262296 Hypoglycemic episodes and risk of dementia in ...	CONTEXT Although acute hypoglycemia may be ass...
2	2	11953232 Impaired Awareness of Hypoglycemia in a Popula...	OBJECTIVE To determine the prevalence and clin...
3	3	21472388 Hypoglycemia: incidence and clinical predictor...	OBJECTIVE To determine the frequency of modera...
4	4	24921368 Impaired awareness of hypoglycaemia: a review.	Impaired awareness of hypoglycaemia (IAH) is a...
5	5	7011850 Unawareness of hypoglycaemia and inadequate hy...	OBJECTIVE To examine the traditional view that...
6	6	20697217 Symptoms of hypoglycemia in children with IDDM.	OBJECTIVE To examine the symptoms of hypoglyc...
7	7	32463364 Antihypertensive classes, cognitive decline an...	OBJECTIVES Prevention of cognitive decline and...
8	8	4866637 Diabetes, metabolic syndrome, and breast cance...	Incidences of breast cancer, type 2 diabetes, ...
9	9	11616424 Copyright © 1996, American Society for Microbi...	Hypoglycemia is among the most injurious metab...
10	10	8668863 Interactions between sleep, circadian function...	Sleep disturbances, including sleep insufficie...

Fig. 2. Ranking

	source	target	edge
0	Hypoglycemia	dementia	increases
1		type 2 diabetes mellitus	episodes
2	Impaired Awareness	Population Based Type	Diabetes
3		based IDDM	Hypoglycemia
4		Impaired hypoglycaemia	awareness
5		autonomic neuropathy	Unawareness
6		IDDM	Symptoms
7		cognitive dementia	classes
8		current breast evidence	Diabetes
9		Glucose Influx	Copyright
10		circadian glucose diabetes	Interactions

Fig. 3. Entity Relation Mapping

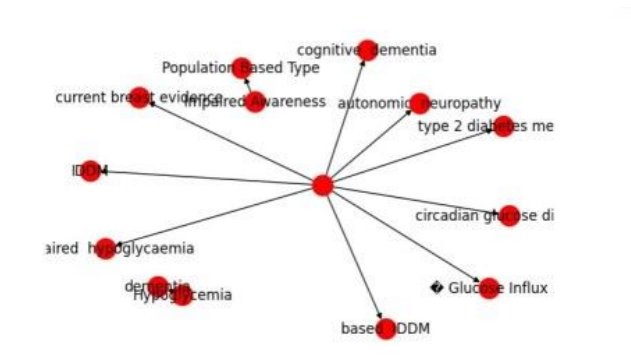


Fig. 4. Knowledge Graph Creation for Explainability

The result for Knowledge graph creation for explain-ability was shown in Fig.4. Using the findings from entity recognition and relationship extraction, NLP creates a network where entities are nodes and the relationships are the links between them. This allows us to explore the data in a structured way.

VI. CONCLUSION

Our main objective in this study is to apply XAI to support the anticipated results. Arguments for choosing the best output are made using measures like as precision, recall, and others to get the intended outcomes. This will make things simpler on the user and assist them in selecting the correct response, hence avoiding

misunderstanding. User decisions will improve because of the transparent and effective explainability feature. In conclusion, this paper offers details on XAI and how it may be used to give predictability. XAI is transforming the world of healthcare information retrieval. Because XAI delivers transparency, efficiency, and confidence, it has the potential to dramatically transform the manner in which humans identify, treat, and manage illnesses. Going forward, ethical and responsible use of XAI will be necessary to ensure that this potent innovation serves the greatest good.

REFERENCES

- [1] ANAND, A., LYU, L., IDAHL, M., WANG, Y., WALLAT, J., AND ZHANG, Z. Explainable information retrieval: A survey.
- [2] CHADDAD, A., PENG, J., XU, J., AND BOURIDANE, A. Survey of explainable ai techniques in healthcare. vol. 23, MDPI, p. 634.
- [3] CHEN, Z., WANG, Y., ZHAO, B., CHENG, J., ZHAO, X., AND DUAN, Z. Knowledge graph completion: A review. vol. 8, IEEE, pp. 192435–192456.
- [4] CHOUDARY MOVVA, B. V. T. G. R. G. P. G. G. Predicting covid-19 positive cases and analysis on the relevance of features using shap (shapley additive explanation). In 2021 Second International Conference on Electronics and Sustainable Communication Systems (ICESC) (2021), IEEE, pp. 1892–1896.
- [5] DWIVEDI, R., DAVE, D., NAIK, H., SINGHAL, S., OMER, R., PATEL, P., QIAN, B., WEN, Z., SHAH, T., MORGAN, G., ET AL. Explainable ai (xai): Core ideas, techniques, and solutions. vol. 55, ACM New York, NY, pp. 1–33.
- [6] GANESHKUMAR, M, R. V. S. V. G. E. S. K. Explainable deep learning-based approach for multilabel classification of electrocardiogram. IEEE.
- [7] LIU, T.-Y., ET AL. Learning to rank for information retrieval. vol. 3, Now Publishers, Inc., pp. 225–331.
- [8] MACDONALD, C., SANTOS, R. L., AND OUNIS, I. The whens and hows of learning to rank for web search. vol. 16, Springer, pp. 584– 628.
- [9] MENON, REMYA, K. A. Improving ranking in document based search systems. In 2020 4th International Conference on Trends in Electronics and Informatics (ICOEI)(48184) (2020), IEEE, pp. 914– 921.
- [10] MENON, R. R., GAYATHRI, S., AND AMINA, A. Representation of documents using minimal dictionary of embeddings. In 14th International Conference on Advances in Computing, Control, and Telecommunication Technologies, ACT 2023 (2023), Grenze Scientific Society, p. 1897 – 1903.
- [11] MENON, R. R., RAHUL, R., AND BHADRAKRISHNAN, V. Graph auto encoders for content-based document recommendation system. In 2023 4th IEEE Global Conference for Advancement in Technology (GCAT) (2023), IEEE, pp. 1–8.
- [12] RANI, N. S., NACHAPPA, C., KRISHNA, A. S., AND NAIR, B. B. Multi disease diagnosis model for chest x-ray images with explain-able ai–grad-cam feature map visualization. In 2022 International Conference on Futuristic Technologies (INCOFT) (2022), IEEE, pp. 1–5.
- [13] SARASWAT, D., BHATTACHARYA, P., VERMA, A., PRASAD, V. K., TANWAR, S., SHARMA, G., BOKORO, P. N., AND SHARMA, R. Explainable ai for healthcare 5.0: opportunities and challenges. vol. 10, IEEE, pp. 84486–84517.
- [14] UPADHYAY R., KNOTH, P., PASI, G., AND VIVIANI, M. Ex-plainable online health information truthfulness in consumer health search. vol. 6, Frontiers, p. 1184851.