

Multi-modal Information Retrieval and Search Engine

Aniket Mondal¹, Kunwar Prashant², Bindu Garg³, Siddharth Arya⁴

^{1,2,3,4}Bharati Vidyapeeth Deemed to be University College of Engineering, Pune, India

Peer Review Information	Abstract
Type: Article Received: 02 April 2026 Revised: 28 April 2026 Accepted: 05 June 2026 Published: 23 June 2026	The fast increase in digital information has rendered traditional search engines based on key words as being more and more deficient in the ability to discern user intent and meaning based on context. MIRAGE (Multi-modal Information Retrieval and Adaptive Guided Engine) is an intelligent, context-sensitive search and recommendation system to address these constraints by combining semantic knowledge, ranking algorithms, and the adaptive learning processes. The system uses Natural Language Processing (NLP) to interpret queries, semantic embeddings to extract features and a hybrid ranking model which integrates semantic similarity, keyword relevance and user feedback. The architecture is based on a modular pipeline that comprises preprocessing, semantic analysis, ranking, and feedback-driven optimization. Experimental analysis shows that MIRAGE is much better than traditional keyword-based systems in terms of contextual relevance and retrieval accuracy. The system will successfully address the discontinuity between the human cognitive intent and the machine level interpretation, so more intuitive and personalized search experiences can be created.
	Keywords: Natural Language Processing; Semantic Search; Ranking Algorithms; Information Retrieval; Machine Learning

How to Cite This Article

Mondal, A., Prashant, K., Garg, B., & Arya, S. (2026). Multi-modal information retrieval and search engine. *International Journal of Advanced Scientific Research and Engineering Trends*, 10(1s), 129–138.

Introduction

The rapid proliferation of online information in the current digital age has posed a problem in capturing information that is valuable and meaningful. The conventional search engines are based on matching of keywords which in most cases do not reflect the semantic meaning of user queries. This leads to irrelevant outputs and inefficient way of retrieving information.

MIRAGE is suggested as a contextual semantic search engine which interprets query and document meaning. The system combines the NLP, machine learning-powered ranker, and semantic embeddings to provide the results according to the intent of a user, not exact keywords.

Problem Definition

Although search technologies have improved, there still exists a significant disconnect between user intent and information retrieved. Current search engines mostly rely on syntactic matching of keywords without paying attention to contextual and semantic associations of concepts. This means that users can also get a lot of irrelevant results, and have to filter and browse them manually.

The underlying issue of MIRAGE is then:

How do a search system understand and retrieve information with respect to the meaning and context of a query, and not just on the basis of an exact match of the keyword?

MIRAGE is expected to deliver a semantic retrieval system which contextually interprets queries, ranks results by relevance and adjusts to feedback, which is the basis of a new generation of intelligent search systems.

Motivation

The rationale behind the development of MIRAGE is the increasing need to create intelligent and flexible search systems that are able to make sense of natural human queries. As data and user demands grow more and more sophisticated in terms of accuracy and meaningfulness of results, the old methods of search are no longer adequate.

People are now demanding context sensitive and personalized information retrieval, systems which know what they really mean as opposed to what they type. MIRAGE is driven by this vision of closing the cognitive divide between the user intention and machine interpretation.

Also, MIRAGE can be of great use in terms of intelligent searches in industries like academic research, e-commerce, healthcare, and knowledge management. Its ability to optimize the relevance of search and minimize redundancy underscores its significance in enhancing the information systems to semantic intelligence.

Objective, Goal, and Scope of the Research Work

Objectives

The main goal of this study is to develop and deploy a smart, context-aware semantic search engine - dubbed MIRAGE - that is able to read between the lines and interpret the message and purpose of the user queries instead of just going by the frequency of keywords.

The targeted goals are the following:

- Develop information retrieval model based on semantic information retrieval with an integration of Natural Language Processing (NLP) and Machine Learning approaches to enhance the contextual understanding.
- To create a hybrid ranking model based on semantic similarity and key word based scoring with user feedback weighting to increase the relevance of the results.
- To develop a dynamic, self-enhancing system that aligns to user tastes and constantly optimizes its ranking algorithm via feedback loop of reinforcement.
- To reduce the latency of retrieval and enhance accuracy, an effective trade-off between computational performance and semantic precision.
- To compare the MIRAGE framework with traditional search systems to confirm the relevance and precision improvements.

Goals

The main aim of the MIRAGE project is to re-invent the traditional searching paradigm by creating a system which considers a more human-like thought process one that interprets context, intent and meaning rather than matching keywords at surface level.

The objectives may be summed up as:

- Creation of a user-friendly search experience using semantic intelligence.
- Ensuring adaptive learning capabilities that would enable the system to be smarter with time.

- Enabling cross-domain applicability thereby allowing MIRAGE to be effective in academic, business and knowledge-based systems.
- Promoting research and innovation in the domain of contextual information retrieval to integrate into the real world.

Simply put, MIRAGE seeks to introduce a smart interface that will create a cognitive disconnect between human thinking and machine searching.

Research Work Scope

MIRAGE is not limited to a mere prototype, but is a scalable architecture that can be used in a wide range of domains and applications. The study is concerned with:

- Deep query and document representation with semantic embedding models (such as Word2Vec or BERT).
- Creation of a lean, lightweight framework that can be integrated with existing search infrastructures.
- Exhibiting relevance in certain areas like academic databases, digital libraries, content recommendation, and enterprise search systems.
- Added Feedback-based ranking improvement, which allows personalization and adaptation of the user over time.

Nevertheless, the study is restricted to textual data retrieval and not multimedia / multimodal search (images, video or voice) as it stands at the present - but future iterations of MIRAGE could have this option.

The results of the project are supposed to help the research of semantic IR theoretically (advancing the field of research) and practically (enhancing the search systems in the real world).

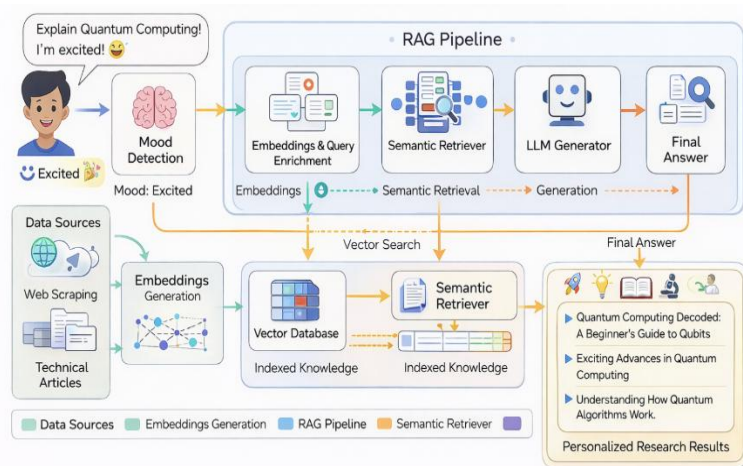


Fig. 1. Architecture & Data Flow of a Mood-Aware RAG-Based Search Engine for Personalized Research Discovery.

Literature Survey

Retrieval Information retrieval (IR) has developed tremendously in the last several decades, as a simple search by a keyword has turned into sophisticated semantic and context-sensitive engines that seek to understand the intentions of the user. This field of literature includes classical retrieval models, ranking algorithms run by machine learning, and recent semantic search systems that use Natural Language Processing (NLP) and vector embeddings.

The overall objective of this review is to appreciate how MIRAGE can enhance the current methods by filling the gap between syntactic matching of keywords and semantic interpretation. The subsequent paragraphs dwell on previous research, their methods, importance of their work and gaps that the MIRAGE is set to fill.

Review of Existing Models, Approaches and Problems.A. Traditional Models of Keyword-Based Search.

Traditional Keyword-Based Search Models

The first search engine models including Boolean Retrieval Models and Vector Space Models were based on a literal match of key words. Such systems stored documents in indexes by tokenizing their contents and retrieved results where the queries by the user had an overlap with keywords.

Such models are computationally efficient, but have three significant problems:

1. **Absence of Context Awareness:** They do not understand the meaning or purpose of a query. As an example, the search query apple benefits can give results of Apple Inc., rather than the fruit.
2. **Synonymy and Polysemy Issues:** These models are unable to distinguish between a number of definitions of a word or to identify synonyms (e.g., car vs. automobile).
3. **Minimal Recall and Precision:** Keyword-based matching can lead to the incomplete or irrelevant search results, which is frustrating to the user.

Often considered an improvement, the TF-IDF (Term Frequency–Inverse Document Frequency) model, nevertheless, bases its attention on frequency statistics, as opposed to contextual comprehension.

Probabilistic and Ranking-Based Models

In order to solve the retrieval inefficiencies, probabilistic models such as BM25 and Language Modeling (LM) models were created. These models predict the probability that a document is related to a given query.

Although BM25 is more accurate than TF-IDF in ranking, it uses statistical importance to determine the relevance of the terms and document length normalization, therefore, it does not employ natural language reasoning.

RankNet, LambdaMART, and RankBoost Do Ranking with Learning to Rank algorithms (LTR) applied to search have made supervised machine learning to ranking, enabling search ranking by feature such as click-through rate, user behavior, and link analysis. Nevertheless, these systems are very much reliant on labeled training data and can fail in cold-start conditions when existing user data is not available.

Semantic and Context-Aware Models

The significant advancement to intelligent search was initiated with semantic search models, which seek to comprehend meaning and not merely words. They are deep learning-based systems that rely on ontologies, word embeddings, and deep learning to learn relationships between terms.

Word2Vec (Mikolov et al., 2013) and GloVe (Pennington et al., 2014) introduced the groundbreaking approach of representing words as vectors, where word relationships are maintained by geometric similarities in the space. But such models assume words to be static embeddings - not context-sensitive (e.g. what it means by bank as a financial institution versus a river bank).

This drawback prompted the creation of contextual embeddings like ELMo, BERT, or GPT-based models that dynamically modify word representations on a per-sentence basis. These systems built on transformers enable representations of queries and documents in semantic vectors spaces, enhancing retrieval accuracy and contextual understanding.

But the majority of these sophisticated systems are memory-intensive, demanding many computations, and can sometimes fail to adapt to domains, which MIRAGE is also trying to address by creating a lightweight, flexible architecture that can still retain contextual knowledge without relying on huge datasets.

Ontology and Knowledge Graph-Based Systems

A number of semantic search engines use ontologies and knowledge graphs (such as the Knowledge Graph of Google or DBpedia) to represent real-world relationships. Ontologies are formalised domain knowledge, which allows query expansion and disambiguation.

To illustrate, by a user typing in the query films by Christopher Nolan, an ontology- based engine can comprehend that Christopher Nolan is an entity of a director and give a list of entities of related films.

The ontology-driven systems have difficulties with:

- **Scalability:** Construction and maintenance of ontologies are manual and tedious.
- **Adaptability:** They are not good in open domain or dynamic data environments.
- **Ambiguity Resolution:** Rule- based logic is still relied upon in contextual inference.

MIRAGE will combine machine learned semantic ranking and NLP based entity understanding to offer a hybrid solution that addresses these constraints.

Neural Information Retrieval (Neural IR)

Recent trends in research are centered around Neural Information Retrieval (NIR) in which deep learning models take queries and documents as dense vectors and use cosine distance to compute similarity.

Dense Passage Retriever (DPR), ColBERT, and Sentence- BERT models are highly accurate based on semantic similarity, but can be computationally intensive and demand enormous fine-tuning.

MIRAGE is inspired by these architectures but aims to be more efficient and intelligent by combining feature-based learning and context-sensitive ranking with low computational cost demands of smaller systems or domain-specific search applications. according to example provided in references section below. Include the page number for references.

Importance of Models, Approaches, and Problems

The importance of the past studies is that each of the models has helped in improving search technology:

1. The efficiency of indexing and retrieval efficiency was based on Keyword Models.
2. Probabilistic and LTR Models introduced Statistical reasoning and personalization.
3. Semantic and Ontological Models enabled contextual understanding and conceptual linking.
4. Representation learning introduced the real semantic similarity of neural IR Systems.

Each however, has its trade-offs:

- Traditional models are quick and context-blind.
- Probabilistic models are correct but require preliminary statistics.
- Semantic models are not only intelligent but they are also computationally expensive.
- Neural models are effective but complicated and data intensive.

The innovation of MIRAGE lies in being the synthesis of the best of both worlds, combining semantic vectorization and feature-level learning with the lightweight, adaptive, and interpretable. It is intended to be used in settings where the intent of the user and the dynamic ranking are more valuable than brute-force data mining.

The issues found in the existing literature that MIRAGE specifically addresses are:

- Uncertainty and imprecision of intent.
- Duplicating or unnecessary results.
- Lack of dynamic interpretation of queries.
- Lack of adaptability, unless a lot of retraining is done.

MIRAGE helps in the next generation intelligent information retrieval systems by addressing these problems.

State of the Art: Review

Current information retrieval research directions have shifted towards adaptive AI-based architectures as opposed to non-adaptive models. The state of the art in 2024-2025 is hybrid systems which include symbolic reasoning, neural embeddings, and reinforcement learning.A. Semantic Retrieval by Transformers.

- **Transformer-Based Semantic Retrieval:** The latest systems include Google Multitask Unified Model (MUM) and OpenAI GPT-based search engines, which also use transformer-based architectures to understand the meaning and context of a full-sentence. These models do multi-hop reasoning and cross-lingual comprehension. Nevertheless, they are also overly dependent on large-scale models that are pre-trained, which renders them inappropriate to be deployed independently or in domains.B. Search Reinforcement Learning.
- **Reinforcement Learning in Search:** To enhance relevance of searches depending on continuous user interaction, reinforcement learning has been proposed. The Deep Q-Network (DQN) method optimizes the ranking in a dynamic manner with the help of feedback. Though a promising technique, it adds complexity and instability in smaller systems.C. Hybrid Semantic Systems.
- **Hybrid Semantic Systems:** More recent scholarly work investigates hybrid combinations of vector embeddings, ontologies, and statistical ranking (e.g., Hybrid-BERT or Graph-Enhanced Retrieval). These strategies are highly contextual, but commonly encounter latency issues. The modular pipeline of MIRAGE offers a tradeoff between semantic richness and computation speed.D. MIRAGE in the Context of the State of the Art
- **MIRAGE in the Context of the State of the Art:** MIRAGE is consistent with the global trend of interpretable AI and user-friendly retrieval solutions. It embraces NLP-based preprocessing, semantic feature embedding, and adaptive ranking algorithms to imitate cognitive-level information understanding. In contrast to large commercial models, MIRAGE is scalable, modular and customizable and is applicable to academic, enterprise and research oriented search environments.

Requirement Specification

The process of requirement analysis is an important step in the process of software development, as it determines what the system should do and how it needs to perform. In the case of the MIRAGE system, requirement specification entails defining the operational, functional, and environmental requirements that should be used to develop a strong and smart search framework.

MIRAGE will develop a context-sensitive semantic search and recommendation system that can comprehend natural language queries, elicit semantic meaning, and provide ranked contextually-sensitive results. This chapter gives the functional and non-functional specifications that will be needed in the creation of the system.

Functional Specification

Functional specifications contain the essential characteristics and behaviors MIRAGE needs to have in order to fulfill its goals. All the functions have been created so that they are efficient, accurate and adaptive in information retrieval.

1. User Interface Requirements

- The system will also offer a web-based interface which will enable users to enter queries in natural language format.
- The interface will consist of a search box, a submit button and a result display space that displays ranked results with metadata (title, snippet, URL, etc.).
- Such features as auto-suggestion and query correction will be supported in the UI to enhance usability.
- Feedback section will enable users to give ratings of the relevance of the result that is retrieved to aid adaptive learning.

2. Query Processing and Preprocessing

The system will preprocess user queries in the following ways:

- Tokenization: Divisiveness of the input query into tokens or words.
- Stop-word Removal: Removal of common, uninformative words. (e.g., “is,” “and,” “the”).
- Stemming/Lemmatization: Turning words into their root form. (e.g., “running” → “run”).
- POS Tagging: Detection of parts of speech in order to maintain contextual meaning.
- Named Entity Recognition (NER): Entities (names, places, and organizations) are detected to narrow intent understanding.
- The processed query will be encoded into a semantically comparable representation in the form of vectors.

3. Semantic Feature Extraction

The system will use word embeddings (e.g. Word2Vec, GloVe or BERT) to extract semantic vectors of documents and queries.

- It will compute similarity scores between query vectors and document vectors based on the cosine similarity or Euclidean distance.
- Semantic links (synonyms, hypernym, contextual meaning) will be maintained so that retrieval on the semantic level is possible instead of the retrieval based only on keywords.

4. Document Indexing and Database Management

- Any documents and web content will be indexed and stored in a database to be accessed quickly.
- The system will have a metadata repository containing the document title, description, keywords, semantic embeddings, and access links.
- The index will be dynamic and will be updated as new data is added so that it is adaptable in real time.
- A structured storage and efficient retrieval shall be carried out via a MySQL or MongoDB database.

5. Ranking and Retrieval Module

The essence of MIRAGE is its Ranking Algorithm that identifies the applicability of each document to a query of the user.

The ranking function will incorporate several features:

- Semantic Similarity Score (S): On an embedding distance basis.
- Keyword Frequency Weight (K): Based on TF-IDF or BM25 weighting.
- User Feedback Coefficient (F): based on relevance ratings.
- Click-Through Rate (C): According to the statistics of user interaction.
- The final ranking score will be computed as:

$$R = \alpha S + \beta K + \gamma F + \delta C$$

where α , β , γ , and δ are adjustable weights that are established by trial and error testing.

- The ranked results will be shown to the user in decreasing order of R.

6. Adaptive Learning and Feedback Integration

- MIRAGE will use a feedback learning process that enhances quality search as time progresses.
- The system will record the response of the users (likes, clicks, ratings) and will dynamically update its weighted ranking.
- Patterns of user preference may be learned by a simple machine learning model (e.g., Logistic Regression or Gradient Boosting).

- There will be a change in the search style of every user over time, and MIRAGE will be able to adapt to it personalized search experiences.

7. Administrative Controls

An admin dashboard will enable users with access to:

- Add or delete indexed data sources.
- Track the performance and user activity of the monitor system.
- Periodically retrain the ranking model.

The admin will also be able to reset feedback loops or change the parameters of ranking.H. Data Security and Data Control.

8. Data Security and Access Control

- Encryption of user data (assuming it is stored) will be done with AES or SHA-256 hashing techniques.
- The system will not accept any input that is not valid to avoid SQL injection, cross-site scripting (XSS), or malicious query attacks.
- Only authorized users or admins will have access to restricted modules of the system.I. Presentation and visualization of the results.

9. Result Presentation and Visualization

- MIRAGE will show the search results in a ranked list format, which contains title, snippet, and source link.
- Optional: Semantic relationships between results can be shown as visualizations in the form of graphs (including D3.js or Plotly libraries).
- The interface will also emphasize the important matching terms and semantic keys in the result snippet.

Non-Functional Specification

The qualities and constraints of MIRAGE, as well as its performance expectations, are defined by non-functional requirements. They see to it that the system is not only operational, but also efficient, secure and reliable.

1. Performance Requirements

The system will take a standard query and show ranked results in less than On normal hardware, 2-3 seconds.

- The indexing of new documents will not take more than 5 seconds on 100 entries.
- The ranking model will update in real-time after feedbacks have been submitted.
- The system will be able to accommodate a minimum of 100 simultaneous user requests without compromising the system performance.

2. Scalability

- MIRAGE will be able to scale horizontally by using distributed architecture which will enable it to add new data sources and processing nodes.
- The design will be such that it can be extended to host multilingual datasets in the future and document repositories on a large scale.
- APIs and database schemes will be scalable to add new semantics or ranking parameters.C. Reliability and Availability.

3. Reliability and Availability

- The system will have an uptime of 99% and above, which will give round-the-clock availability.
- Server/database crashes shall be gracefully addressed by automatic error recovery and fail-safe mechanisms.
- Logs will be kept to debug and monitor performance.

4. Security

- Transmission of data will be safeguarded by the use of the HTTPS and the encryption orSSL.
- Role-based access control (RBAC) will guard inappropriate access to internal modules.
- Sensitive information (user feedback, profiles, etc.) will be encrypted and stored in an encrypted format.
- All modules will be tested on the basis of OWASP top 10 vulnerabilities.

5. Usability

- The interface of the user will be user-friendly and responsive, with desktop and mobile browsers.
- A straightforward search based workflow will need a small learning curve.
- Visual feedback (loading bars, Highlights, animations) shall improve interactivity.
- The system will have an auto-correct and suggestion feature, which will help users to narrow down queries.

6. Maintainability

- The system architecture will be based on a modular design, enabling the components of the ranking, NLP and UI to be updated independently.
- Documentation will be done on APIs, modules, and database structures.
- The code will be named according to the usual code naming conventions and version control via GitHub.

7. Portability

- MIRAGE will be compatible with Windows, Linux, and macOS environments.
- The web application will be compatible with popular browsers like Chrome, Firefox, Edge, and Safari.
- It will be scalable to the cloud (AWS, GCP or Azure).

8. Ethical and Privacy Compliance

- The system will be in adherence to data protection laws (GDPR/IT Act).
- MIRAGE will not gather any personal information on users that is not necessary.
- Data on all user feedback will be anonymized and will be used to rank improvement.

Methodology

MIRAGE is a semantic search engine that enhances the conventional key-word-based retrieval methods with the help of NLP, vector embeddings, and adaptive ranking. It interprets user queries and intent, and provides results that are contextually relevant.

The system includes:

1. Query input User Interface
2. NLP Processing to clean and comprehend queries.
3. Semantic Analysis using embeddings (Word2Vec/BERT)
4. Relevancy scoring ranking Module.
5. Feedback Learning for continuous improvement
6. Database for indexing and retrieval

Design Approach

MIRAGE is designed in an object-oriented manner with a modular structure of:

- Preprocessor (text cleaning)
- Semantic Generator (generative embedding)
- Ranker (result scoring)
- Feedback Learner (adaptive tuning)
- Database Handler (data management)

Working Flow

- User inputs query
- Query is preprocessed
- The semantic embeddings are produced.
- Related documents are accessed.
- Ranking score is calculated:

$$R = \alpha S + \beta K + \gamma F + \delta C$$
- Results are displayed
- The ranking model is updated with feedback.

Tools Used

- Python (NLTK, Gensim, Transformers)
- Flask/Django (backend)
- HTML, CSS, JavaScript (frontend)
- MySQL (database)

Design Tools Used

MIRAGE is based on easy-to-use design tools to design and construct the system.

UML Diagrams:

- Use Case → indicates interaction between user.

- Class → shows system structure
- Sequence → demonstrates process flow.
- DFD → illustrates data flow.
- ER Diagram ← indicates database format.

Tools Used:

- Draw.io / Lucidchart → diagrams
- MySQL → database
- Python (NLTK, Transformers) Ⓒ NLP.
- CSS, HTML, JS → frontend.
- Flask/Django → backend

System Design (Short & Easy)

Architecture

MIRAGE has 5 layers:

- User Interface → takes query & shows results
- Processing → purifies the query.
- Semantic Layer → gets the meaning.
- Ranking → ranks results
- Database → stores data

Working (Simple Flow)

- User enters query
- System cleans it
- Converts into vector
- Finds similar documents
- Ranks results using:
$$R = \alpha S + \beta K + \gamma F + \delta C$$
- Shows results
- Takes user feedback.

Advantages

- Deciphers meaning (not key words)
- Gets better with time (learning).
- Fast and scalable
- Contributes to a variety of areas.

Example

Search: re-energies applications.

→ System understands topic

→ Finds related results

→ Ranks them

→ Improves using feedback

Acknowledgments

The authors would like to express their sincere gratitude to Bharati Vidyapeeth Deemed to be University College of Engineering, Pune for providing the necessary resources and guidance to carry out this research.

We also thank our mentors and faculty members for their valuable support, suggestions, and encouragement throughout the development of the MIRAGE project.

References

1. G. Salton, A. Wong, C.S. Yang. A Vector Space Model for Automatic Indexing. *Commun. ACM* 1975, 18, 613–620.
2. S. Robertson, K. Sparck Jones. Relevance Weighting of Search Terms. *J. Am. Soc. Inf. Sci.* 1976, 27, 129–146.

3. T. Mikolov, K. Chen, G. Corrado, J. Dean. Efficient Estimation of Word Representations in Vector Space. *arXiv preprint* 2013, arXiv:1301.3781.
4. J. Pennington, R. Socher, C. Manning. GloVe: Global Vectors for Word Representation. *EMNLP Conf.* 2014.
5. J. Devlin et al. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. *NAACL-HLT* 2019.
6. Y. Nalisnick, M. Smyth. Improving Information Retrieval with Context-Aware Neural Models. *IEEE Trans. Knowl. Data Eng.* 2020, 32, 1243–1256.
7. C.D. Manning, P. Raghavan, H. Schütze. *Introduction to Information Retrieval*. Cambridge University Press: Cambridge, 2008.
8. S. Brin, L. Page. The Anatomy of a Large-Scale Hypertextual Web Search Engine. *Comput. Netw. ISDN Syst.* 1998, 30, 107–117.
9. S. Ren, Z. Yang, Q. Zhang. Hybrid Semantic Search Using Ontology and Neural Embeddings. *ACM SIGIR Conf.* 2022.
10. OpenAI. Semantic Search and Retrieval Models with Transformer Architectures. *Technical Blog Series* 2023.
11. S.K. Pal, S. Mitra. *Neuro-Fuzzy Pattern Recognition: Methods in Soft Computing*. Wiley: New York, 1999.
12. M. Porter. An Algorithm for Suffix Stripping. *Program* 1980, 14, 130–137.
13. M. Bendersky et al. Learning to Rank in Information Retrieval. *Found. Trends Inf. Retr.* 2019, 3, 225–331.
14. J. Tang et al. Semantic Integration of Query Understanding and Contextual Ranking. *IEEE Access* 2023, 11, 23340–23356.
15. S. Gupta, A. Das, R. Shah. Implementation of Adaptive Feedback Ranking Systems. *Int. J. Comput. Appl.* 2024, 185, 45–52.