

From Vector Space to Neural Ranking: A Comparative Study of Modern Information Retrieval Models

Prapitha Gopi K

Research Scholar, United College of Arts and Science, Periyanaickenpalayam, Coimbatore.

prapithakishore1@gmail.com • ORCID: 0000-0000-0000-0000 (optional)

Abstract

Information retrieval has changed dramatically over the past decades. Early systems relied on simple keyword matching, but modern search engines must understand meaning, context, and user intent. This paper examines three major families of retrieval models that have shaped this evolution: vector space models, probabilistic retrieval, and neural retrieval. Vector space models represent documents and queries as weighted term vectors and rank them by similarity, providing a simple yet effective way to handle partial matches. Probabilistic models, such as BM25, treat relevance as a probability and rank documents according to how likely they are to satisfy a query, offering a stronger theoretical foundation for ranking. Neural retrieval goes further by learning dense semantic representations that can capture meaning beyond exact word overlap, enabling more accurate matching and reranking. We review key works including Salton et al.'s foundational vector space model, Robertson and Zaragoza's probabilistic relevance framework, and recent neural approaches such as Dense Passage Retrieval and large language model-based retrieval surveys. The discussion shows that modern search systems rarely rely on a single model. Instead, they combine fast lexical retrieval with powerful neural rerankers to balance speed and accuracy. This hybrid approach reflects the current state of the field and points toward future research directions.

Keywords: information retrieval, vector space model, probabilistic retrieval, BM25, neural retrieval, reranking

1. Introduction

Information retrieval (IR) is the field that deals with finding relevant information from large collections of documents, web pages, or other records. The core challenge in IR is that user queries are often short, vague, and worded very differently from the documents that actually answer them. Early search systems mainly relied on simple keyword matching, but this approach struggled to capture meaning, synonyms, and variations in language. To overcome these limitations, researchers developed more advanced ranking models that use mathematical, probabilistic, and learned representations to improve search quality.

The evolution of IR models can be understood in three broad stages. First, the vector space model introduced a geometric way of representing text, which made it possible to rank documents by similarity instead of just matching keywords. Second, probabilistic retrieval treated relevance as a probability problem and led to highly effective ranking functions such as BM25. Third, neural retrieval uses deep learning and transformer-based models to learn dense semantic representations, allowing systems to rerank candidate documents with a better understanding of meaning. This paper compares these three approaches and explains why modern search systems increasingly combine them into hybrid retrieval architectures that balance speed and accuracy.



International Journal of Technology and Emerging Research (IJTER)

Volume 2 | ICITT-2026 | Article ID: 3677-7020-4045 | <https://doi.org/10.64823/ijter.2621011>

IORO Publications · Texas, USA · www.ioro.org

Unlike many existing surveys that focus only on classical models or only on modern neural methods, this paper offers a concise, comparative overview that brings together vector space, probabilistic, and neural retrieval in a single narrative. Our contribution has three main parts: (i) a structured review of how IR models have evolved over time, (ii) a critical comparison of their strengths, limitations, and practical roles, and (iii) a discussion of hybrid architectures and possible directions for future research. We hope this perspective helps researchers and practitioners see how different retrieval approaches can work together in modern search systems.

2. Literature Review

2.1. Early Information Retrieval Models

The first computer-based retrieval systems emerged in the late 1940s and 1950s, drawing inspiration from the manual cataloging and indexing methods long used in libraries. These early systems mainly relied on Boolean keyword searches, retrieving documents only if they matched exact query terms combined with logical operators like AND, OR, and NOT. Although straightforward and easy to interpret, Boolean retrieval had serious limitations: it provided no ranking of results, demanded highly precise queries, and could not deal with synonyms or partial matches.

During the 1960s and 1970s, pioneering work by researchers such as Gerard Salton and the SMART project at Cornell University laid the groundwork for modern IR. They introduced key ideas like the inverted index and early term-weighting schemes, which made it possible to move beyond simple exact-match retrieval. This era marked the beginning of a shift toward ranked retrieval, where systems could order documents by their estimated relevance instead of returning an unranked list of matches.

2.2. Vector Space and Algebraic Models

The vector space model, introduced in the 1970s, was a major breakthrough because it treated documents and queries as weighted vectors and ranked them by similarity. The influential paper “A Vector Space Model for Automatic Indexing” by G. Salton, A. Wong, and C. S. Yang (1975), published in *Communications of the ACM*, is widely regarded as the first formal presentation of this model in information retrieval. This algebraic approach made partial matching possible and introduced term-weighting schemes such as TF-IDF, which quickly became a standard tool in IR.

Salton’s work sparked a wave of research on term weighting, relevance feedback, and evaluation methods. Later, the textbook “Introduction to Information Retrieval” by C. D. Manning, P. Raghavan, and H. Schütze (2008), published by Cambridge University Press, offered a clear and comprehensive treatment of vector space models, TF-IDF, and cosine similarity, helping to make these ideas core material in IR courses. Comparative studies between vector space and probabilistic models have shown that vector space methods are simple and intuitive, but they do not provide a direct probabilistic interpretation of relevance. Even so, they remain important as baseline models and as building blocks in more complex systems, especially when combined with efficient indexing structures.

2.3. Probabilistic Retrieval and Language Models

The vector space model, introduced in the 1970s, was a major breakthrough because it treated documents and queries as weighted vectors and ranked them by similarity. The influential paper “A Vector Space Model for Automatic Indexing” by G. Salton, A. Wong, and C. S. Yang (1975), published in *Communications of the ACM*, is widely regarded as the first formal presentation of this model in information retrieval. This algebraic

approach made partial matching possible and introduced term-weighting schemes such as TF-IDF, which quickly became a standard tool in IR.

Salton's work sparked a wave of research on term weighting, relevance feedback, and evaluation methods. Later, the textbook "Introduction to Information Retrieval" by C. D. Manning, P. Raghavan, and H. Schütze (2008), published by Cambridge University Press, offered a clear and comprehensive treatment of vector space models, TF-IDF, and cosine similarity, helping to make these ideas core material in IR courses. Comparative studies between vector space and probabilistic models have shown that vector space methods are simple and intuitive, but they do not provide a direct probabilistic interpretation of relevance. Even so, they remain important as baseline models and as building blocks in more complex systems, especially when combined with efficient indexing structures.

2.4. Probabilistic Retrieval and Language Models

Probabilistic information retrieval emerged in the 1970s and 1980s with the aim of placing ranking on a firmer probabilistic foundation. The **Probability Ranking Principle (PRP)**, formalized by **S. E. Robertson (1977)** and later refined in subsequent work, states that retrieval works best when documents are ranked in decreasing order of their estimated probability of relevance. This idea led to models such as the Binary Independence Model and, eventually, to **BM25**, which is still widely used in modern search engines.

The article "**The Probabilistic Relevance Framework: BM25 and Beyond**" by **S. E. Robertson and H. Zaragoza (2009)**, published in *Foundations and Trends in Information Retrieval*, offers a detailed survey of this framework and the development of BM25. BM25 builds on earlier probabilistic models but adds term frequency and document length normalization, making it both effective and practical for real-world systems. Language modeling approaches further extended probabilistic IR by treating retrieval as the problem of estimating the probability that a document would generate a given query. These models provided a more principled way to handle term distributions and document length, and they later influenced work on neural ranking and language model-based retrieval.

2.5. Neural Information Retrieval

The rapid advancement of deep learning has significantly transformed the field of Information Retrieval (IR), making neural information retrieval one of the most active areas of research. Unlike traditional retrieval methods that rely mainly on exact keyword matching, neural IR models are capable of understanding the semantic meaning of queries and documents. These models are generally categorized into three types: sparse, dense, and hybrid approaches. Sparse neural models improve traditional term-based representations by learning the importance of individual terms, whereas dense models use transformer-based encoders to convert queries and documents into continuous vector representations. This enables the system to retrieve relevant information based on semantic similarity rather than simple word overlap. As a result, dense retrieval methods have become widely adopted in applications such as question answering, semantic search, and retrieval-augmented generation (RAG).

A major breakthrough in this field came with the paper "**Dense Passage Retrieval for Open-Domain Question Answering**" by **V. Karpukhin, B. Oguz, S. Min, P. Lewis, L. Wu, S. Edunov, D. Chen, and W.-T. Yih**, presented at **EMNLP 2020**. The authors introduced a dual-encoder architecture that learns dense vector representations for queries and documents, demonstrating better retrieval performance than the widely used BM25 model in open-domain question answering tasks. This work established dense retrieval as a practical and effective alternative to traditional lexical retrieval methods. Since then, researchers have developed several neural ranking architectures, including bi-encoders, which provide efficient large-scale

retrieval, and cross-encoders, which achieve higher ranking accuracy by jointly processing queries and documents.

Recent survey papers provide a broader understanding of the progress in neural information retrieval. The survey “**Large Language Models for Information Retrieval: A Survey**” by Y. Gao and colleagues (2023) explores the growing role of large language models (LLMs) in modern IR systems. It explains how LLMs can function as query rewriters, retrievers, rerankers, and readers, enabling more intelligent and context-aware search. Similarly, “**A Survey of Model Architectures in Information Retrieval**” by A. Petrov and co-authors (2025) reviews the evolution of IR models, beginning with traditional keyword-based approaches and progressing to transformer-based and LLM-driven architectures. The survey emphasizes the development of backbone models and end-to-end retrieval systems that improve both efficiency and retrieval quality. Together, these studies highlight that neural approaches consistently achieve superior performance in tasks requiring semantic understanding, while traditional methods such as BM25 continue to remain highly effective for large-scale web and document retrieval because of their speed, simplicity, and robustness.

2.6. Evaluation and Large-Scale Experiments

The Text Retrieval Conference (TREC), established in 1992, has played a vital role in advancing Information Retrieval (IR) research by providing standardized evaluation frameworks and large-scale benchmark datasets. These benchmarks have enabled researchers to systematically compare the performance of different retrieval models, including vector space, probabilistic, and more recently, neural approaches, under consistent and controlled experimental settings. As a result, TREC has become one of the most influential platforms for measuring progress in IR and validating new retrieval techniques. Evaluations conducted using TREC and similar benchmark collections have consistently shown that neural retrieval models excel in tasks requiring semantic understanding and contextual relevance. At the same time, traditional methods such as BM25 continue to demonstrate strong performance in large-scale web and document retrieval, making them a reliable and widely used baseline in modern information retrieval systems.

Overall, the existing body of research highlights the steady evolution of Information Retrieval (IR) systems. The field has progressed from early Boolean and keyword-based retrieval methods to probabilistic ranking techniques and, more recently, to neural and hybrid models that combine computational efficiency with a deeper understanding of semantic relationships. This evolution reflects the continuous effort to improve both the accuracy and relevance of information retrieval in increasingly complex and diverse search environments.

3. Vector Space Model

The **Vector Space Model (VSM)** represents both documents and user queries as vectors in a multi-dimensional term space, where each dimension corresponds to a unique term. Each term is assigned a weight, commonly calculated using techniques such as **Term Frequency-Inverse Document Frequency (TF-IDF)**, and the relevance of a document is determined by measuring its similarity to the query, typically using cosine similarity. Unlike earlier retrieval models that provide only a binary decision, the Vector Space Model ranks documents based on their degree of relevance, making it more flexible and effective for information retrieval.

One of the key strengths of the Vector Space Model is its ability to perform **partial matching**. A document does not need to contain every query term to be considered relevant; instead, it can still receive a high ranking if its overall pattern of weighted terms closely matches the query. This allows the model to retrieve a

broader set of potentially useful documents and often improves the user experience. However, the model has important limitations. Since it treats documents as a **bag of words**, it ignores the order in which words appear and does not capture the context or relationships between them. As a result, it has limited semantic understanding and often struggles to recognize that different words or phrases, such as synonyms and paraphrases, may convey the same meaning.

4. Probabilistic Retrieval

Probabilistic retrieval models approach information retrieval by estimating the likelihood that a document is relevant to a user's query. Instead of simply measuring the similarity between a query and a document, these models assign a probability of relevance and rank documents accordingly. This approach is based on the **Probability Ranking Principle (PRP)**, which states that the most effective retrieval results can be achieved by presenting documents in decreasing order of their estimated probability of relevance. This principle provides a strong theoretical foundation and has made probabilistic retrieval one of the most influential approaches in the field of Information Retrieval (IR).

Among probabilistic retrieval methods, **BM25** is the most widely used and successful ranking algorithm. It calculates document relevance by considering factors such as query terms, term frequency, and document length normalization. Because of its effectiveness, efficiency, and simplicity, BM25 has become the standard ranking method in many modern search engines and retrieval systems. Compared with the Vector Space Model, BM25 generally produces more accurate rankings because it is based on probabilistic relevance estimation rather than similarity alone. However, despite its strong performance, BM25 remains a lexical retrieval method that primarily depends on matching query terms with document terms. As a result, it has limited ability to understand the semantic meaning of text and often struggles when the same concept is expressed using different words or phrases.

5. Neural Retrieval

Neural retrieval has emerged as a significant advancement in Information Retrieval (IR) by leveraging deep learning techniques to better understand the meaning of queries and documents. Instead of relying solely on exact keyword matching, neural retrieval models represent queries and documents as dense vector embeddings or directly estimate the relevance of a query-document pair. One of the most influential approaches is **Dense Passage Retrieval (DPR)**, which uses dual-encoder models to generate semantic embeddings for both queries and passages. These embeddings enable the retrieval system to identify relevant documents based on meaning rather than exact word overlap, making DPR particularly effective for large-scale retrieval tasks. After the initial retrieval stage, neural reranking models, such as cross-encoders and transformer-based architectures, are often employed to further improve the ranking of the retrieved documents.

One of the greatest strengths of neural retrieval is its ability to capture **semantic relationships** between queries and documents. Unlike traditional lexical retrieval methods, neural models can identify relevant information even when different words or phrases are used to express the same idea. This capability makes them highly effective for applications such as question answering, recommendation systems, semantic search, and retrieval-augmented generation (RAG). However, these advantages come with increased computational requirements. Training and deploying neural retrieval models require considerably more processing power and memory than classical methods. For this reason, many real-world retrieval systems adopt a hybrid strategy, using a fast-lexical model such as **BM25** for the initial retrieval stage and then

applying neural models to rerank the most promising candidate documents, thereby achieving a balance between efficiency and retrieval accuracy.

6. Review Methodology

This article uses a narrative review approach to trace how information retrieval models have evolved from vector space methods to neural ranking. We selected works based on three main criteria:

- **Foundational influence:** Papers that introduced key concepts or frameworks, such as the vector space model, the Probability Ranking Principle, BM25, and dense retrieval.
- **Practical impact:** Models and methods that are widely used in real-world search systems or standard benchmarks, for example BM25 in web search and transformer-based rerankers in question answering.
- **Recency and coverage:** Recent surveys and empirical studies that capture current trends in neural and large language model-based retrieval.

Rather than aiming for an exhaustive systematic review, our goal is to provide a coherent overview that connects classical and modern retrieval paradigms and highlights their practical implications.

7. Comparison of Models

Model	Core idea	Strengths	Limitations
Vector space model	Represent text as weighted term vectors and rank by similarity	Simple, intuitive, supports partial matching	Weak semantic understanding, ignores order
Probabilistic model	Rank by estimated probability of relevance	Strong theoretical basis, effective ranking, BM25 is practical	Still mostly lexical, uses simplifying assumptions
Neural retrieval	Learn dense semantic representations or pairwise relevance scores	Captures meaning, handles paraphrases, strong reranking	Expensive, data-hungry, less interpretable

The comparison of these retrieval models highlights that they differ primarily in the way they interpret and measure document relevance. The Vector Space Model determines relevance based on the geometric similarity between query and document vectors, whereas probabilistic models estimate the likelihood that a document is relevant to a given query. In contrast, neural retrieval models rely on learned semantic representations, allowing them to capture the underlying meaning of text rather than depending only on exact keyword matches. Recognizing that each approach has its own strengths and limitations, many modern information retrieval systems integrate these methods within a retrieve-then-rerank framework. In such systems, a fast lexical model is first used to retrieve a set of candidate documents, after which a neural model reranks them to improve retrieval accuracy while maintaining computational efficiency.

7.1. Critical Insights on Reviewed Models

Each family of retrieval models has clear strengths, but also important limitations that affect where and how well it can be used.

Vector space models are simple, easy to interpret, and computationally efficient, which makes them attractive for baseline systems and teaching. However, they rely on bag-of-words representations, ignore word order, and struggle with synonymy and polysemy. Their lack of a probabilistic interpretation also makes it harder to integrate them with more advanced ranking frameworks.

Probabilistic models, especially BM25, have a stronger theoretical foundation and have shown robust performance across many different collections. They handle term frequency and document length more effectively than basic vector space models. Still, they remain largely lexical and do not capture deeper semantic relationships. Their performance can drop on queries that require understanding of intent, context, or paraphrase.

Neural retrieval models are strong at capturing semantic similarity and can significantly improve ranking quality on complex tasks such as question answering and conversational search. However, they need large amounts of training data, substantial computational resources, and careful engineering to scale to large corpora. Their reduced interpretability also raises concerns about explainability and trust, particularly in sensitive domains such as healthcare and law.

These trade-offs help explain why modern systems increasingly use hybrid architectures that combine efficient lexical retrieval with more powerful neural reranking.

7.2. Quantitative Comparison of Retrieval Models

Although it is hard to compare models directly because studies use different datasets and evaluation metrics, several works give useful indications. On standard ad-hoc retrieval benchmarks such as TREC and MS MARCO, BM25 usually delivers strong baseline performance, with Mean Average Precision (MAP) and nDCG scores that are competitive with more complex models on purely lexical tasks.

Neural models, especially dense retrievers and transformer-based rerankers, often outperform BM25 on tasks that require semantic understanding. For instance, Dense Passage Retrieval (DPR) has been shown to improve top-k retrieval accuracy on open-domain question answering datasets compared to BM25, particularly when questions and answers use different wording. Recent surveys of neural ranking models report consistent improvements in nDCG@10 and MRR when neural rerankers are added on top of lexical first-stage retrieval.

Overall, these findings suggest that BM25 remains a strong and efficient baseline, while neural models can deliver extra gains on semantically complex queries, at the cost of higher computational requirements.

8. Discussion

The evolution of Information Retrieval (IR) models reflects the field's gradual transition from simple keyword-based matching to a deeper understanding of the meaning and context of information. Early models, such as the **Vector Space Model**, laid the foundation for ranked retrieval by representing documents and queries mathematically and comparing them using similarity measures. This marked an important shift from binary retrieval to relevance-based ranking. Later, **probabilistic models** provided a stronger theoretical basis for ranking by estimating the likelihood that a document would satisfy a user's information need. More recently, **neural retrieval models**, particularly dense retrievers and transformer-based rerankers, have become the focus of research because of their ability to capture contextual and semantic relationships between queries and documents more effectively than traditional methods.

Despite their impressive performance, neural retrieval models are computationally expensive and are therefore not commonly used as the only retrieval method in large-scale search systems. Instead, most modern information retrieval systems adopt a hybrid approach. In this architecture, a fast lexical retrieval method, such as **BM25**, is first used to retrieve a manageable set of candidate documents. These candidates are then passed to a neural reranker, which performs a more detailed semantic analysis to produce a more

accurate final ranking. This **retrieve-then-rerank** strategy has become the standard architecture for contemporary search engines and question-answering systems, as it effectively balances computational efficiency with high retrieval accuracy.

8.1. Future Research Directions

Several emerging trends and challenges are likely to shape the next generation of information retrieval systems.

- **Large language models (LLMs) for retrieval:** LLMs are increasingly used not only as rerankers but also as query rewriters, document generators, and even end-to-end retrievers. Bringing LLMs into retrieval pipelines raises important questions about cost, latency, hallucination, and how to evaluate such systems reliably.
- **Efficiency and scalability:** Neural models are computationally expensive, especially at web scale. Current research on efficient transformers, model distillation, and approximate nearest neighbor search aims to reduce latency and resource usage while keeping retrieval effectiveness high.
- **Explainability and fairness:** Because retrieval systems influence what information users see, issues of bias, fairness, and explainability are becoming more critical. Future work needs to explore how to make ranking decisions more transparent and how to reduce biases present in training data and models.
- **Multimodal and conversational retrieval:** Users are increasingly interacting with systems through natural language, voice, and images. Extending retrieval models to handle multimodal inputs and multi-turn conversations remains an active and important research area.
- **Robustness and security:** Neural retrieval models can be vulnerable to adversarial attacks and shifts in data distribution. Ensuring that these models remain robust under changing conditions and potentially hostile inputs is an important open challenge.

Addressing these challenges will require close collaboration between retrieval researchers, system engineers, and domain experts to build IR systems that are not only effective, but also efficient, fair, and trustworthy.

9. Conclusion

The field of Information Retrieval (IR) has evolved significantly over the years, moving from vector-based similarity models to probabilistic ranking techniques and, more recently, to neural approaches that focus on semantic understanding. The Vector Space Model continues to hold an important place in IR, not only because it introduced the concept of ranked retrieval but also because it remains a valuable baseline for evaluating new retrieval methods. Probabilistic models, particularly BM25, have maintained their relevance due to their balance of effectiveness, efficiency, and scalability, making them widely adopted in both research and real-world search systems. In recent years, neural retrieval models have further advanced the field by capturing contextual meaning and semantic relationships, leading to more accurate retrieval and reranking performance. As a result, the current trend in Information Retrieval research is toward hybrid retrieval systems that combine the speed and efficiency of lexical methods with the semantic understanding of neural models, offering an effective balance between retrieval quality and computational cost. A promising direction for future research is to keep combining lexical and neural methods into hybrid retrieval pipelines, while also working on efficiency, explainability, and the use of LLM-based components in retrieval systems.

References

- [1] G. Salton, A. Wong, and C. S. Yang, "A Vector Space Model for Automatic Indexing," *Communications of the ACM*, vol. 18, no. 11, pp. 613–620, Nov. 1975.
- [2] K. Sparck Jones, "A Statistical Interpretation of Term Specificity and Its Application in Retrieval," *Journal of Documentation*, vol. 28, no. 1, pp. 11–21, 1972.
- [3] S. E. Robertson and K. Sparck Jones, "Relevance Weighting of Search Terms," *Journal of the American Society for Information Science*, vol. 27, no. 3, pp. 129–146, 1976.
- [4] S. E. Robertson and S. Walker, "Some Simple Effective Approximations to the 2-Poisson Model for Probabilistic Weighted Retrieval," in *Proceedings of the 17th Annual International ACM SIGIR Conference on Research and Development in Information Retrieval (SIGIR)*, 1994, pp. 232–241.
- [5] C. D. Manning, P. Raghavan, and H. Schütze, *Introduction to Information Retrieval*. Cambridge, UK: Cambridge University Press, 2008.
- [6] V. Karpukhin, B. Oguz, S. Min, P. Lewis, L. Wu, S. Edunov, D. Chen, and W.-T. Yih, "Dense Passage Retrieval for Open-Domain Question Answering," in *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 2020, pp. 6769–6781.
- [7] L. Xiong, C. Xiong, Y. Li, K.-F. Tang, J. Liu, P. N. Bennett, J. Ahmed, and A. Overwijk, "Approximate Nearest Neighbor Negative Contrastive Learning for Dense Text Retrieval," in *International Conference on Learning Representations (ICLR)*, 2021.
- [8] O. Khattab and M. Zaharia, "ColBERT: Efficient and Effective Passage Search via Contextualized Late Interaction over BERT," in *Proceedings of the 43rd International ACM SIGIR Conference on Research and Development in Information Retrieval (SIGIR)*, 2020, pp. 39–48.
- [9] Y. Gao, X. Xiong, X. Gao, K. Jia, J. Pan, Y. Bi, Y. Dai, J. Sun, and H. Wang, "Large Language Models for Information Retrieval: A Survey," *arXiv preprint arXiv:2308.07107*, 2023.
- [10] A. Petrov *et al.*, "A Survey of Model Architectures in Information Retrieval," *arXiv preprint*, 2025.
- [11] E. M. Voorhees and D. K. Harman (Eds.), *TREC: Experiment and Evaluation in Information Retrieval*. Cambridge, MA, USA: MIT Press, 2005.
- [12] D. Metzler and W. B. Croft, "A Markov Random Field Model for Term Dependencies," in *Proceedings of the 28th Annual International ACM SIGIR Conference on Research and Development in Information Retrieval (SIGIR)*, 2005, pp. 472–479.
- [13] J. Lin, R. Nogueira, and A. Yates, "Pretrained Transformers for Text Ranking: BERT and Beyond," *Synthesis Lectures on Human Language Technologies*, vol. 14, no. 4, pp. 1–325, 2021.