

Information retrieval: a view from the Chinese IR community

Zhumin CHEN¹, Xueqi CHENG², Shoubin DONG³, Zhicheng DOU⁴, Jiafeng GUO (✉)²,
Xuanjing HUANG⁵, Yanyan LAN (✉)², Chenliang LI⁶, Ru LI⁷, Tie-Yan LIU⁸,
Yiqun LIU (✉)⁹, Jun MA¹, Bing QIN¹⁰, Mingwen WANG¹¹, Jirong WEN⁴, Jun XU⁴,
Min ZHANG⁹, Peng ZHANG¹², Qi ZHANG⁵

¹ School of Computer Science and Technology, Shandong University, Jinan 250100, China

² Institute of Computing Technology, Chinese Academy of Sciences, Beijing 100190, China

³ School of Computer Science and Engineering, South China University of Technology, Guangzhou 510006, China

⁴ School of Information, Renmin University of China, Beijing 100872, China

⁵ School of Computer Science, Fudan University, Shanghai 200433, China

⁶ School of Cyber Science and Engineering, Wuhan University, Wuhan 430072, China

⁷ School of Big Data, Shanxi University, Taiyuan 200433, China

⁸ Microsoft Research Asia, Beijing 100080, China

⁹ Department of Computer Science and Technology, Tsinghua University, Beijing 100084, China

¹⁰ School of Computer Science and Technology, Harbin Institute of Technology, Harbin 150001, China

¹¹ School of Computer Information and Engineering, Jiangxi Normal University, Nanchang 330022, China

¹² School of Computer Science and Technology, Tianjin University, Tianjin 300072, China

© Higher Education Press 2020

Abstract During a two-day strategic workshop in February 2018, 22 information retrieval researchers met to discuss the future challenges and opportunities within the field. The outcome is a list of potential research directions, project ideas, and challenges. This report describes the major conclusions we have obtained during the workshop. A key result is that we need to open our mind to embrace a broader IR field by rethink the definition of information, retrieval, user, system, and evaluation of IR. By providing detailed discussions on these topics, this report is expected to inspire our IR researchers in both academia and industry, and help the future growth of the IR research community.

Keywords information retrieval, redefinition, information, scope of retrieval, retrieval models, users, system architecture, evaluation

1 Introduction

The early idea of using computers to search for relevant pieces of information was popularized in the article “As We May Think” by Bush in 1945 [1]. During the six decades that followed, we have witnessed the booming of information retrieval (IR) in both industry and academia. Especially after Web search engines were invented, the arising need for advanced IR technologies led to a huge wave of IR researches in our community. This was reflected by the growing numbers of IR related conferences, workshops, and contests, as well as the huge volume

of ideas generated in those events.

However, during the past few years, many members in our community have raised their concerns that IR researches seem to be shrinking [2]. For example, the numbers of submissions to major IR conferences (e.g., SIGIR) are keeping stable if not declining, while those to the sibling conferences (e.g., KDD and ACL) have significantly increased. Many IR related research topics, such as recommender systems, multimedia retrieval, and human computer interactions, have faded out and found their new home (e.g., RecSys, ICMR, CHI). Some closely related research areas, like natural language processing and data mining, have been driven by the new engine of Big Data and Artificial Intelligence, while the IR community seems to remain in its traditional pace in contrast.

Given all of this, it seems necessary for us to have a reflection on our current situation, and figure out the major challenges and opportunities that the community is facing. With such a motivation, the Chinese IR community organized a strategy workshop to discuss future challenges and opportunities within the broad IR field on February 2nd and 3rd, 2018. The goal is to open our mind by rethinking the definition of information, retrieval, user, system, and evaluation. The expected output is a list of exciting and challenging future research directions that we should devote our time and energy to. In this paper, we report the major conclusions we have obtained during this strategy workshop.

2 Redefining information retrieval

2.1 Motivations

Over decades, IR researches and applications have achieved great success. Especially after the computer was invented, many

Received March 12, 2019; accepted July 21, 2019

E-mail: {guojiafeng, lanyanyan}@ict.ac.cn; yiqunliu@tsinghua.edu.cn

solid IR technologies have emerged. For instance, the inverted index [3], the vector space models [4], the probabilistic retrieval models [5], the language models [6, 7], and the Cranfield evaluation methodology, etc. Driven by the invention of the World Wide Web in late 1990s, Web search became one of the main research areas in IR. Correspondingly, link analysis (e.g., PageRank [8] and HITS [9]), query logs based ranking signals, and the learning-to-rank techniques have been developed, which enabled us to leverage the interconnectivity of billions of web pages, the behaviors of millions of users, and the combination of thousands of signals to make IR systems stronger than ever.

However, at the same time, the gap between the accessible Web data in industry and academia is getting wider. This limits the healthy evolution of the IR research community, especially in the Big Data and Artificial Intelligence (BigData+AI) era.

As we have noticed, in the BigData+AI era, many research fields have been greatly boosted, e.g., machine learning, natural language processing, computer vision, etc. However, our IR community is a little bit losing its track comparatively. It would be timely and helpful for us to perform a serious reflection on the current situation of IR research whether the current components of IR, e.g., system architecture, retrieval models, user modeling, evaluation methodologies need to be refined or even re-defined. With this motivation, we have organized a workshop to collect the wisdom of crowds. We noted down some of our discussions during the workshop in the following sections, with the goal of pushing forward the Renaissance of IR research.

2.2 Dimensions of redefinition

During the discussions of this workshop, we consider the redefinition of seven dimensions of Information Retrieval. Here, we summarize some basic ideas, whereas details of some of them will be discussed in the next sections.

- **Redefine information**

The terminology information does not only refer to documents or webpages, but also involves a variety of information. The richness of information could be characterized by its openness of information (e.g., private or public), its formats (e.g., webpages, microblogs, WeChat dialogue or APPs), and its structure (e.g., free texts, tables or knowledge graphs).

- **Redefine scope of retrieval**

The scope of retrieval should not be restricted to searching indexed documents/web pages. It should also cover the searching of user generated information (e.g., generated microblog posts/CQA contents), the combination of different information resources, and the reasoning from knowledges. Even further, we should not restrict ourselves to retrieval, and should think about recommendation, text analytics, question answering, text summarization, chatbot, as well.

- **Redefine retrieval models**

Traditional retrieval models, which have been used for decades, should also be refined. For example, the target should be much border than a listwise ranking result. Besides, more advanced technologies should be leveraged to enhance the ability of retrieval models, e.g. deep learn-

ing, reinforcement learning, and adversarial learning. The general principle is to synchronize the design of retrieval models with the technical frontier of other fields, especially machine learning and artificial intelligence.

- **Redefine users**

Previously users are regarded as customer of IR systems, however, today they are also contributors. The interactions between users and IR systems should be emphasized and investigated with a new paradigm. Game theory can be utilized to model such user-system interplay, and corresponding learning algorithms, e.g., generative adversarial networks, reinforcement learning, and game-theoretic learning should be leveraged.

- **Redefine system architecture**

The architecture of IR systems should be redefined or re-designed in order to index different types of data (structured or unstructured), integrate the cloud computing technologies, have clearer interfaces, and build end-to-end systems.

- **Redefine evaluation**

We should emphasize the online and interactive nature of IR evaluation. A well-defined simulation system might be necessary to bridge the gap between offline and online evaluations.

- **Others**

Except the aforementioned topics, we also need to consider other important factors, such as new IR theory, interpretation of retrieval models, and so on.

The relations of different dimensions is illustrated in Fig. 1. Based on the above outline, we have conducted extensive discussions on each aspect of redefinition as well as their related topics (except the IR theory, which we believe requires more in-depth study in a dedicated paper). We summarize these detailed discussions as potential exciting and challenging IR research topics in the next a few years.

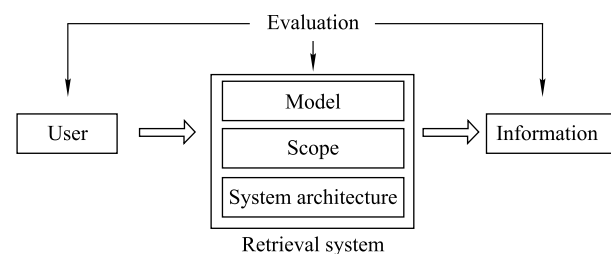


Fig. 1 The relations of different dimensions

3 New definition and representation of information

3.1 Motivation

Traditional informational retrieval systems are mainly based on text retrieval. However, with the rapid development of Internet, more and more data are generated by different ways. The data are large-scale, multi-source, heterogeneous, cross-domain, cross-media, cross-language, and dynamically evolving [10]. All these characteristics of big data bring new research topics and massive challenges to information retrieval.

3.2 Proposed research topics

Fusion of big data from heterogeneous sources

Data indexed by current information retrieval systems mainly come from news websites, video websites, image websites etc. However, with the development of Internet of things and mobile Internet, the data sources should also include wearable devices, mobile devices, smart home sensors, APP stores, etc. Based on these different kinds of data resources, an effective sharing mechanisms for data should be established for efficient retrieval.

Representation of indexed data to be retrieved and understood effectively

To index and retrieve the big and complex data accurately, effective data representation is a necessary step. The traditional data representation in current retrieval systems is one-hot representation, i.e., bag-of-words. For a given Web page with different terms, it is represented as a vector where each dimension is 0/1 or a weight computed by TF-IDF corresponding to a term. One-hot representation assumes that all dimensions are independent. As a result, the similarity relations between words are not inaccurate. With the development of representation learning, distributed representation denotes an object as a dense, real-valued and low-dimensional vector. The object maybe a letter, a word, a sentence, a document, an APP, a user, an item, a query etc. A document can be represented by fusing different levels of granularity (character, word, sentence, passage, document). After objects and queries are denoted into a unified space, it is easy to compute the similarity between them accurately.

Representation of extracted knowledge to answer users queries directly

After knowledge is extracted, it should be represented within a common semantic space in order to be used for information retrieval systems. At present, knowledge graph is the most popular form of knowledge which is represented as a series of triples, i.e., head entity, tail entity, relation between head and tail. Entities and relations are also denoted as the dense, real-valued and low-dimensional vectors with distributed representation learning. These vectors form a unified semantic space in which knowledge can be transferred across sources, domains and models. Finally, the next generation retrieval systems can understand users queries better, retrieve relevant objects more accurately with the help of knowledge. Furthermore, for some difficult queries that cannot be matched from the Web directly, knowledge graph can be used to infer the correct answers.

3.3 Research challenges

Some challenges will be faced by the next generation retrieval systems with the development of Internet and new technologies.

(1) Cross source/domain/mode/language data indexing. Data to be indexed will come from many new sources such as wearable devices, tensors. How to design a new approach to aggregate and index data from multiple sources, domains, modes, languages and views in order to retrieve the diversity results effectively and accurately is a challenge.

(2) Cross source/domain/mode/language data representation. How to learn distributed representations for different information objects or users information needs from multiple sources,

domains, modes, languages and views considering their contents, relations, structures and faces in order to compute the similarity accurately between them is another challenge.

(3) Knowledge graph representation and indexing. Knowledge graph is especially useful for open domain information retrieval tasks. How to best construct, represent, index and utilize knowledge graphs so that they are maximally useful for the next generation retrieval system is also a challenge.

(4) Data sharing mechanism. How to build an open mechanism to share deep processed and high value data (private or public) among different research institutions and enterprises under the protection of privacy in order to satisfy different users needs is a challenge, too.

4 Enlarged scope of retrieval

4.1 Motivation

Conventional information retrieval [11] has, to some extent, the underlying assumption that users' information needs, represented with the issued keyword queries, can be satisfied with a list of ranked documents. The documents are retrieved from a static document collection and ranked according to the relevance of the document to the query. Based on the assumption, different relevant ranking models, also called retrieval models, have been proposed and successfully applied in search engines.

Recently, however, this assumption has been challenged by the development of Web and the emergence of variant approaches to accessing information. Researchers realize that information retrieval is not equal to document ranking or retrieving a list of documents. The limitations of conventional IR models and systems include:

(1) Users must know in advance what information they need, and then try to pull the information from the static document set.

(2) Information retrieval systems are generally unidirectionally query-based. They are only able to respond to specific user requests. They can generally neither proactively generate information for users, nor even respond to queries in a user-specific fashion.

(3) The documents indexed in IR systems are relatively fixed and the provided information is directly selected from the index. Any indirect knowledge available through analysis of current information, or implicit knowledge inherent in the patterns of information retrieval, cannot be exploited to enable push of user-specific content or to enhance semantic representations of content.

(4) Listing a list of documents ranked with relevance is far from the optimal way for representing the retrieval results. To satisfy users' information requirement, an answer directly summarized from the relevant documents is a better output, which is beyond the current ranking based optimization target.

Overall, for better satisfying users information needs and improving their search experiences, IR systems are expected to go beyond the simple relevant document ranking. There are more and more demands on moving from existing information retrieval paradigm to general information access.

4.2 Proposed research topics

From active information retrieval to passive information retrieval

At present, IR systems assume that the users know exactly in advance what information they need. They also have the ability of summarizing their information needs as keyword queries. This is usually called active information retrieval as the users actively issue queries and wait the systems to answer the queries with document lists. In active information retrieval, the burden on search users is high as the users need to represent their search intent with keyword queries accurately and instantly. In many cases, this is hard or even impossible for search users.

In passive information retrieval, on the other hand, the keyword queries are no longer necessary for pushing the information. The user is kept to be updated with new information after some initial configurations. For example, the users may want to be notified if there are some new publications in a particular research topic or there are some new citations to some particular papers. In some extreme cases, the retrieved resources are not merely stored statically but are reported more or less promptly to those who are interested or are assumed to be interested. For example, the information resources can be periodic reports generated by some information systems.

Going beyond retrieving a closed set of documents

Current IR systems are designed to retrieve a static set of documents. Thus, once been deployed, the indexed documents are remains relatively fixed and can be updated periodically with newly retrieved documents. Users can get the most relevant documents, but of course, they are fetched by the IR systems from a closed document set.

Ideally, an IR system can not only retrieve information from the closed document set, it can also provide users the generated and modified new information, for better answering the user queries.

Going from search engines to analytic engines

Existing search engines are designed for finding the documents that may satisfy users information need, by issuing simple keywords as queries. Users need to browse and read the results and summarize the information contained in these documents by themselves. This usually costs much user efforts when users have complex information needs, such as doing a survey on a research topic or learning about the latest progress on an event. It would be great if search engines can directly extract and aggregate relevant information nuggets (such as topics, events, person names, locations, organizations, etc) from search results, and provide a kind of multi-dimensional interactive analytics to users. Users can click on a dimension item, and drill down into the information they are interested in. This is a kind of text analytical service like OLAP [12] in the database field. This could significantly reduce user efforts on reading and summarizing documents by themselves, and narrow the gap between real user information need and the information returned by IR systems.

New functions for information retrieval

With the development of internet services, there are much more diverse information needs that can be categorized into the area of information access, which can be an important research problems in the broader IR area. In addition to a traditional IR system, users may need other kinds of services, like recommen-

dation, text analytics, question answering, text summarization, chatbot, etc. Different from the existing IR systems (such as search engines) which focus on retrieving relevant documents that already exist on the web, these services may provide direct answers or higher-order knowledge to users, or move from the traditional ten blue links to conversational IR [13].

The comparison of traditional and new methods in the scope of retrieval is illustrated in Table 1.

Table 1 The comparison of traditional and new methods in the scope of retrieval

Traditional methods	New methods
Active information retrieval	Passive information retrieval
A closed set of documents	Generated and modified new information
Search engines	Analytic engines
Retrieval function	Other services

4.3 Research challenges

There are several challenges when we go beyond the existing IR paradigm.

(1) Currently, search engine is a successful IR application and it significantly impacts the IR research community. New killer applications are needed to demonstrate the usefulness and effectiveness of these new IR approaches.

(2) How to effectively evaluate the new IR approaches is a challenge. It is usually hard to improve the quality of algorithms without a clear evaluation criterion. The complexity of the new IR functions we proposed above is much higher than retrieving a simple document list, and we need carefully design corresponding evaluation metrics for these tasks.

5 AI-Enhanced retrieval models

5.1 Motivation

IR models lie in the heart of the information retrieval research field. Different techniques have been proposed and applied in IR models, from traditional heuristic methods, probabilistic models, to new machine learning to rank techniques. Recently, with the advance of new AI technology, such as deep learning and reinforcement learning, a lot of research areas have been pushed forward, including speech recognition [14], computer vision [15], and natural language processing [16]. This has led to expectations that these novel AI techniques are likely to demonstrate similar scale of breakthroughs on IR tasks.

During the past few years, we have witnessed the growing body of work in applying deep neural networks in IR models. There have been related workshops, such as SIGIR Neu-IR workshops [17, 18], encouraging the discussions and development of new IR models with neural networks. However, we are still in the early days in leveraging AI techniques for IR models. Unlike in computer vision or natural language processing, few positive results have been reported in IR with new AI techniques. We are still waiting for exciting new breakthroughs in this field. Meanwhile, there are broad AI methods, beyond deep neural networks, to be explored to enhance IR models. These new techniques, like generative adversarial networks [19] and deep reinforcement learning [20], could lead to new IR models as well as the definition of new IR tasks.

To incorporate novel AI techniques to enhance IR models,

we need to clear up, to what IR applications could AI techniques be applied, to what extent could AI techniques bring in retrieval performance, what will be the challenges in this enhancement, what are the basic influences in IR models, and the deficiency as well as the most overlooked places.

5.2 Proposed research

The proposed research can be divided into the following four major areas. The comparison of traditional and new retrieval models is illustrated in Table 2.

Table 2 The comparison of traditional and new retrieval models

Traditional models	New models
BM25	Neural representations
Language model for IR	Deep neural networks
Learning to rank	Reinforcement learning
	Adversarial methods

IR models with neural representations

By encoding texts or images into real-valued vector representations, neural representations have shown their ability in capturing the semantic meanings of the objects, and demonstrated incredibly power in natural language processing and image recognition. These neural representation techniques could also bring significant changes to IR by altering the fundamental representations of queries, documents and so on. Example research questions include:

- How to leverage neural representations to enhance the modeling of different IR objects, such as queries, documents, images, questions and answers? What are the major benefits?
- How similarity/relevance estimation can be enhanced via neural representations?
- How to improve the efficiency of learning good representations for queries and documents?
- How can neural representations be efficiently indexed for online access?
- How to leverage neural representations to enhance cross-modal search?

IR models using deep neural networks

Deep neural networks have shown promising results in modeling complex tasks with their ability of approximating arbitrary functions. It is natural to apply these powerful models for IR, but it is limited by simply using them to pushing up the retrieval performances. There are still many important questions we need to tackle when applying deep neural models for IR, such as:

- What are the right architectures of neural networks for different IR tasks?
- How can deep architectures give us new insights about core IR problems?
- What are the appropriate training data, test data and toolkits for neural models for IR?
- How can we interpret the learned deep neural models for IR?

- What are the relationships between neural models and traditional models for IR?

IR models using reinforcement learning

Reinforcement learning is to learn how to interact with environment by maximizing a future reward. The recent progress in combining deep learning with reinforcement learning has achieved a lot of momentum especially in computer games [20, 21]. Since IR in essential is about the interaction between users and information, or users and search systems, it is natural to utilize reinforcement learning methods for IR tasks. We list a few potential research directions where reinforcement learning could be applied to enhance IR models:

- Online learning to rank in order to optimize an IR model dynamically over time
- Session search where users aim to complete a complex IR task with multiple steps in a session
- Conversational search where multi-turn interactions between users and systems are involved to obtain the target information
- Some complex ranking tasks, e.g., diversified search, where independence between documents no longer holds and a list needs to be optimized according to multiple criteria
- User modeling where a dynamic user profile needs to be obtained in order to achieve personalized search or recommendation

IR models using adversarial methods

Adversarial methods is among the important progresses in recent AI researches, which is useful to produce robust models by introducing adversarial samples or signals. The same idea has also been applied in generative models, i.e., GAN [19], where the generator tries to generate adversarial instances that can cheat the discriminator, while the discriminator tries to be robust to the adversarial instances. It would be promising to borrow these ideas into IR and we have already witnessed some success in this direction, like IRGAN [22], but there is still much space to be explored:

- How to construct a robust ranking model by involving adversarial instances?
- What are the adversarial samples in different IR tasks and how to generate them automatically?
- Is it possible to leverage GAN to automatically generate search results, user queries, answers and so on?
- Is it possible to use adversarial idea to model the learning and evaluation processes?

5.3 Research challenges

Reproducibility: Many AI enhanced IR models (e.g., neural IR models) often consist of a large number of free parameters to be learned, thus require large quantities of labeled training examples. Due to the lack of large scale public datasets, many recently published models have been trained and evaluated on private industrial datasets. Their

results cannot be fully recovered as their datasets are often not available to academia. Moreover, since these models often consist of many hyper-parameters as well as some detailed tuning tricks which are often missing in published papers, these problems make the issue of reproducibility more serious.

Generalizability: The generalization ability of modern complex AI models is still unclear. Most of them involve large number of hyper-parameters, making them vulnerable in transferring from one dataset to another. This is largely different from many traditional IR models, which are usually simple but can achieve robust performance out of the box over different datasets. Therefore, we require better understanding on the key design principles of AI enhanced IR models to guide us to choose the proper architectures for different IR tasks.

Interpretability: AI enhanced IR models, e.g., some deep neural IR models, may behave like a black box and hard to interpret. This is mainly due to many nested non-linearity functions and the end-to-end training fashion. However, when IR models are applied to accomplish many practical tasks, the interpretability of the results as well as the model itself becomes critical.

Evaluation: With AI enhanced IR models, many traditional benchmark datasets as well as evaluation metrics may no longer fit well. Meanwhile, many complex tasks, such as conversational search and proactive search, could be taken into consideration. Therefore, the development of stable and discriminative metrics for these novel scenarios and new models become urgent. A combined progress on metrics and models may serve as a new breakthrough on this direction.

6 Expanded role of users

6.1 Motivation

User is an important part of an IR system. Although historical user behaviors have been widely used in IR systems to improve the ranking quality [23, 24], users are mainly treated as consumers who use IR systems to find the information they need. Users are usually separated from the IR system, and are not considered as a part of the information process workflow. From another perspective, users are assumed to be consuming data, but not “generating” data.

In recent years, with the rapid growth of social networks [25], users are no longer simple consumers, but can produce data, proactively or passively. These data can be consumed by other users or be used for boosting quality of the retrieval service. Users play a more important role in an IR system. We need to carefully rethink the definition of users when designing a new IR system, taking users as an important part of the system.

6.2 Proposed research topics

Including users in the IR circle

The fundamental problem is to redefine the role of users in IR systems. More system functions need to be designed and studied for proactively encouraging users to generate data, without additional overload. Efforts need to be spent on going from existing document-centric information architecture, to

user-centric one. More information retrieval activities should be redefined to include users in the information producing and consuming circle.

User profiling and personalization

In existing IR systems, personalization is not fully demonstrated to be effective [26]. Sometimes, improper personalization may harm the relevance of results and affect user experience. In most search engines, personalization is mainly about customizing search results based on users’ location and languages. The full advantage of personalization is not exploited due to the complexity of user interest and personalization algorithms. How to control the quality of personalization and implement an effective search result personalization system is still a grand challenge.

Personalization and diversification

In recommendation systems, it becomes a concern to generate over-personalized results. Users usually need diversified results that are relevant to their information need. It is an interesting problem that how to balance personalization and diversification, to improve overall user satisfaction.

User-centric evaluation for IR systems

The goal of IR system is to fulfill users’ information need and make the user satisfied with the overall search experience. Traditional system-centric evaluation paradigms (e.g., the Cranfield-style evaluation) are based on a set of (over-) simplified assumptions about users and therefore cannot perfectly measure the actual user satisfaction and preference. By re-considering and emphasizing user’s role in the search process, we can develop new evaluation paradigms and methods that are more aligned to users experience for IR systems.

6.3 Research challenges

Privacy protection

There is always a trade-off between utilizing user data to improve service quality, and protecting user privacy. The user privacy protection problem raises more and more public attention in recent years.

Supporting complex search tasks

When completing a complex search task, the user usually needs to submit multiple queries to search engines to tackle each subtask separately or to gradually learn about a topic in the multi-query session. Such a complex search task is still considered as challenging for the user and the success of this task may heavily depend on users’ iterative querying behavior. Therefore, it is important for the IR system to build a task-level or session-level user model, and to provide necessary supports (e.g., query suggestions and task-aware ranking) in this scenario.

Understanding users in heterogeneous environments

In recent years, users use IR systems in different environments such as mobile search, image search [27], product search, job and talent search. Users’ search intents, strategies, and behavior may change dramatically in different search environments. Therefore, it is crucial to analyze and understand users using different IR systems in different environments.

Correct and accurate interpretation of user behavior

Although historical user behavior is a valuable information source for improving the quality of retrieval service, it is also known as noisy and biased [28]. For example, previous studies have shown that user clicks are heavily influenced by the position bias, and proposed a series of click models to calibrate the biased CTR signals and extract unbiased relevance feedback. We need further careful investigation on the intrinsic bias in user behavior if we want to exploit it to improve the search system.

7 New system architecture

7.1 Motivation

In recent years, the rapid development of mobile information services and large-scale expansion of personalized and specialized information resources provide unprecedented potential and new opportunities for information growth.

With the development of the next-generation Internet, information resources tend to be largely distributed. Every Web user would publish messages and share information in social circles, which makes Web information substantially heterogeneous, highly aggregated, and socially shared. The social aggregation feature and distributed divergence trend present grand challenges to the existing information retrieval system architecture.

Though the HTTP-based Web encourages a high degree of centralization, emerging technologies and protocols such as Block Chain [29] and IPFS (inter planetary file system) [30] are trying to establish a trust-driven decentralized Web service. In addition, information resources will probably be content-based addressed, instead of domain-based addressed. Information sharing is limited to certain range and requires credit as well as permission. It is foreseeable that new technologies will be of profound impact to information retrieval.

According to the re-definition of IR, users will be core part of the IR framework, not simply participants. To adapt to the rapid accumulation of information resources and the development of new technologies, revolutionary changes for the IR architecture are necessary.

7.2 Proposed research topics

Information retrieval framework for fully decentralized network

Current Web information retrieval mainly relies on super large search engines such as Baidu, Google, and Bing, etc. In a fully decentralized network, each user may build and maintain a search engine and provide search services in a distributed cooperation environment. The capability of users to obtain and retrieve information depends on their credit and authority.

Existing distributed information retrieval (DIR) methods may not be suitable for distributed retrieval in the decentralized network because most of the existing DIR methods are inherently centralized. Therefore, it is necessary to re-design a new distributed information retrieval framework for a fully decentralized Web, in which the retrieval requests can be quickly and accurately responded wherever the search request comes from, as long as there is enough credit and permission.

Information collection on complex network structure

With the development of next generation network technologies, the network structure will inevitably be more complex and divergent. The network structure significantly influences the efficiency and effectiveness of information collection. Another factor that affects the collection of information is the way information is provided and shared. The development of social network will essentially change the way of traditional information crawling. It is necessary to study the complex network structure of future Web, and deeply harness the relationship between network structure and distribution characteristics of information resources, to enable the efficient access to high-quality information collection.

Distributed indexing and exchanging mechanism

Future information resources will probably be content-based addressed, each file (content) has the uniqueness of existence. When a file is added to Web, a unique encrypted hash value is assigned to the content based on calculation. Such mechanism will change the way that using domain names to access web content, and present challenges to the existing inverted indexing mechanism. It is necessary to design efficient distributed indexing mechanism that facilitates various content representations, and meanwhile efficiently supports the sharing and exchange of distributed information with permissions under a well-established credit system.

Search results fusion in distributed heterogeneous environment

The fusion mechanism of retrieval results has always been a critical issue in distributed IR systems. On a decentralized network, the existence of massive retrieval units and numerous distributed heterogeneous resources will result in more critical challenges to traditional solutions. How to select and integrate the most useful information from massive search engine units is a crucial problem. The credit of information sources will play a key role. It is necessary to study credit rewarding mechanism and develop credit based fusion algorithms to reduce the cost and enhance the efficiency of result fusion.

7.3 Research challenges

Due to the rapid development of the distributed Web, and the potential changes caused by new technologies, the IR framework will face a wide range of challenges.

(1) More and more data structures are appearing in the Web. How to define a unified data structure is a critical problem for content-based addressing. Moreover, information will be historic versioning, allowing multiple nodes to manipulate different versions of the content. It is a great challenge to design content-based addressing to support unified information representation of different data formats and support multiple versions' retrieval.

(2) New retrieval framework must support consensus mechanism in the fully distributed and decentralized settings (such as blockchain and its service of transaction). However, current consensus mechanism of blockchain is highly dependent on computing resources and energy consumption [31, 32], and it is challenging to develop new consensus mechanism which can support fast and efficient retrieval.

(3) In a distributed environment, the cooperative ways of distributed nodes are very important. How to establish a well-

defined credit system and apply credit to reliable and trustworthy exchange of original information or indexing is a major challenge.

(4) Due to the numerous search engine nodes and a huge number of documents, the storage architecture of indexing should be carefully re-designed. It is a significant challenge to design an efficient distributed indexing architecture for content representation, and exchange protocols of index files among distributed search engines under privacy protection for highly efficient and secure retrieval.

(5) A huge number of search engines in decentralized distributed Web will exist. How to quickly locate related search engines based on content-based addressing and efficiently respond to users with precise fusion results will be a major challenge.

8 New evaluation methodology

8.1 Motivation

Evaluation has a long history in the area of IR [33]. The basic IR evaluation methods are based on the test collections shared by researchers, which contain a corpus, queries, and relevance assessments. In recent decades, researchers have explored various strategies to evaluate search performance [34].

Previous evaluation methods can be roughly divided into three types: (1) The basic evaluation methods usually try to directly measure returned relevant information objects of the system. (2) Large-scale log studies have also been proposed and used for this task at search engine companies. With the collected retrieval logs, researchers can observe interests and information needs of searchers. (3) Researchers defined and prescribed a small number of topics. Then, users are asked to find information of these topics. They may also be required to provide feedback via questionnaires.

Although previous methods have been successfully used for evaluating IR tasks, there are several issues which we need to pay attention to. Firstly, the information needs of users in existing evaluation models are less considered. Although questionnaires-based methods can provide some information about the actions of users, they cannot be used for large scale evaluation. Secondly, the dynamic nature of Web may highly impact the retrieval objective over time. Thirdly, there are many complex IR tasks which do not have stable and definite end points. Hence, the aim of this proposed research area is to combine the advantages of these two kinds of methods to provide new effective methods for IR evaluation.

8.2 Proposed research

The proposed research can be divided into four major areas: (1) evaluation of spatial and temporal aware multi-recourse retrieval, (2) modeling the users, (3) crowdsourcing evaluation methods, and (4) dynamic datasets for complexity tasks.

Evaluation of spatial and temporal aware retrieval systems

The first research area is related to how to evaluate the performance of spatial and temporal aware retrieval. Due to the continuous increase of mobile search, the objects of users may dynamically change in different location at different time. The retrieval needs are complex and cannot be clearly defined. It becomes a challenging task to evaluate whether the retrieved

documents, images, point-of-interests or even different kinds of sensors satisfy these needs by combining traditional notions of relevance and spatial and temporal information.

Understand the general Web search users

The second research area is related to how to understand the general Web search users [35]. Existing evaluation methods are usually focused on the experienced users with clearly defined tasks. However, there are many different kinds of Web search users. Modeling users can be achieved from the characters of users and the aims of users on the topic, using different kinds of resources. Hence, how to model the differences of users and how to integrate the users' modelling into the evaluation metric is an important and challenge research problem.

Crowdsourcing based evaluation

The third research area is how to use crowdsourcing platform to do the evaluation. Traditional interactive evaluation method needs laboratory for users to evaluate the systems. Hence, it is a time consuming and expensive task, and cannot be used for large-scale data sets. Crowdsourcing platforms provide an opportunity to achieve the problem. A large number of users may participate in the evaluation. However, the quality of evaluation from crowdsourcing platforms is usually lower than on site evaluation. How to design and provide description information for ordinary users is a challenge problem.

Evaluation dataset construction

The fourth research area is related to how to construct datasets used for evaluation. The datasets should contain documents, information seeking requirements, and golden standards. Previous methods usually constructed static dataset to facilitate evaluation. Different types of documents and different kinds of queries should be incorporated into the dataset. Moreover, a novel direction is to construct datasets that can be dynamically changed for realistic evaluation.

8.3 Research challenges

Evaluating IR systems need participation of different kinds of users. However, users are difficult to measure. Peoples with different backgrounds, cultures, languages may provide different feedback for the same thing. Even a single person under different circumstance may give different result. Hence, how to model user is an important task and is also one of the most important challenges in this task. This kind of information can be represented by user models.

How to obtain large-scale resources for academic community is another challenge for this task. Potential datasets should contain millions of searching records by thousands of users. From laboratory, the interviews of users and videos recordings of screens may also provide valuable information for this task. However, these datasets are expensive to construct and the privacy protection is also an important issue for publicly sharing these dataset.

9 Other research topics

9.1 Explainable information retrieval

9.1.1 Motivation

For a long time period IR systems mostly focus on finding rel-

evant results as efficiently and effectively as possible. However, the explainability of IR systems were largely neglected [36,37]. The lack of explainability mainly exists in terms of two perspectives, 1) the outputs of the IR systems (i.e., search or recommendation results) are presented to end users without explanations, and 2) the inner mechanisms of the IR systems (i.e., search or recommendation algorithms) are gray or black boxes to system designers.

This lack of explainability for IR systems leads to major problems in practice [38]. Without making the users aware of why certain results are provided, IR systems may lose its reliability and become less effective in making the users trust the results. More importantly, many IR systems nowadays are not only useful for information seeking, but also useful for complicated decision making by providing supportive information and evidences. For example, medical workers need retrieve comprehensive healthcare documents to make medical diagnosis [39]. In these critical decision-making tasks, explainability of the IR systems are vital so that users can understand why a particular result is provided and how to take advantage of the result for taking actions.

Recently, deep neural models have been widely used in IR systems [40–44]. Though researchers have achieved notable success in neural IR systems, the complexity and the lack of explainability of neural models further emphasize the importance of the research of explainable IR. To bring explainability to IR systems, there is a wide range of research topics for the community to address in the coming years.

9.1.2 Proposed research topics

Leveraging heterogeneous information for explainable information retrieval

Modern IR systems not only deal with textual documents, but also a lot of heterogeneous multi-modal information sources [45]. For example, Web search engines have access to documents, images, videos, audios as candidate results for queries; e-commerce recommendation system works on user numerical ratings, textual reviews, product images, demographic information, etc., for user personalization and recommendation; and social networks leverage user social relations and contextual information such as time and location for search and recommendation.

Current systems mostly leverage heterogeneous information sources to improve search and recommendation performance. A lot of research efforts are needed regarding how to jointly leverage heterogeneous information sources for explainable IR, including research tasks such as multi-modal explanation based on aligning two or more different information sources, transfer learning over heterogeneous information sources for explainable IR, cross-domain explanation in IR systems, and so on.

Personalization in explainable information retrieval

To improve the persuasiveness, trustworthiness, transparency and effectiveness of explainable information retrieval systems, explanations should be personalized for different users. Currently, most of explanations used in search and recommendation are generated based on data mining techniques such as frequent pattern mining and association rule mining. For example, in e-commerce, the most commonly used explanation

is “a certain percentage of users who bought this also bought that” [46], and in social networks the explanation “a certain percentage of your friends also viewed” is generated based on graph mining algorithms [47]. These explanations are not closely coupled with the user’s personalized preferences, and they are also not necessarily related to the IR models that generate the search/recommendation results. As a result, more research efforts are needed to explore personalized explainable IR algorithms and systems.

Fusion of explanations from different models

Different explainable IR models may generate different explanations. We usually have to design different explainable models to generate different explanations for different purposes, and the explanations may not be logically consistent. When the system generates a lot of candidate explanations for a search or recommendation result, a great challenge is how to select the best combined explanations to display in a limited space, and how to fuse different explanations into a logically consistent unified format. Solving this problem requires extensive efforts to integrate statistical and logical machine learning approaches, and to bring in a certain ability of logical inference to explain the results.

Evaluation of explainable information retrieval systems

Evaluation of explainable IR systems remains an important problem. Explainable IR systems can be readily evaluated with traditional IR measures to test its search/recommendation performance. To evaluate the explanation performance, a currently reliable protocol is to test explainable vs. non-explainable IR models based on real-world user study, such as A/B testing in practical systems or evaluation with online workers in M-Turk [48]. However, there is still a lack of offline measures to evaluate the explanation performance. Evaluation of explanations is related to multiple perspectives of information systems, including persuasiveness, effectiveness, efficiency, transparency, trustworthiness, user satisfaction, etc. Developing reliable and easily usable evaluation measures will save a lot of efforts for offline evaluation of explainable IR systems.

9.1.3 Research challenges

Due to the increasing demand of explainability to support comprehensive decision-making tasks in information systems, there are a lot of challenges and opportunities ahead.

(1) Due to the heterogeneous nature of the available data in many information systems, how to integrate information sources with various forms for explainable retrieval is a great challenge.

(2) To support comprehensive decision-making tasks, explainable retrieval models should have the ability to organize different explanations in a logically consistent manner to better help decision makers.

(3) Similar to personalized search and recommendation tasks, the explanations should also be personalized and carefully tailored to different users to improve the effectiveness.

(4) Reliable and general applicable offline evaluation measures and protocols will help evaluate explainable retrieval systems more efficiently.

9.2 Ubiquitous information retrieval

9.2.1 Motivation

With the rapid development of IoT (Internet of Things), users are able to address their information needs anywhere at anytime for **anything** [49]. Traditional search engines based on desktop computers are therefore facing both challenges and opportunities. Imagine that one day you wake up at home, your personal assistant (like Amazon echo) will tell you the most important news today. Then you may drive to work, and the car will automatically find the most convenient route to your office. The very powerful search engine will greatly boost your confidence at work. These are not just imagination, but what is going on in our daily life. Almost every application is constructed based on search. We believe that in the next few years, ubiquitous search will be one of the most promising directions for IR community.

9.2.2 Proposed research topics

Search with ubiquitous devices

According to a report from Google in 2015, the search traffic from mobile devices has exceeded that from conventional desktop devices. This massive shift in search scenario forces both industry and academia to redesign existing technologies in the context of mobile search.

Mobile search is different from desktop search in many aspects: 1) Users are searching for different things using mobile device, e.g., more for entertainment and image, but less for business. 2) Most mobile devices are equipped with a touchable screen, which enables a completely different way of interaction. Meanwhile, mobile devices present much less content at a time due to the limited screen size. Thus, users have to incur a higher interaction cost in order to access the same amount of information. 3) Mobile devices also provide much more space for search. For example, it is much easier for search engines to optimize search results with the geographic location. Considering the differences between mobile and desktop search, it is necessary to calibrate existing techniques to provide better search. Some important research issues include ranking, presentation, personalization, and evaluation.

Search should not be limited to mobile phones or desktop computers. It is possible to search with various intelligent devices. Users' interactions may be performed in different ways with different devices. For instance, with a smart speaker, users may ask some questions. The results might be retrieved from a knowledge base and will be presented with voice. The response should be as accurate as possible since users can hardly select a relevant result from multiple ones as the current desktop search engines. When applying IR application to a new device, the technique should be carefully designed and evaluated adaptively. It is also possible that the future search on various devices will be unified by something like Apple Siri or Microsoft Cortana.

Search for ubiquitous data

Traditional search engines originate from library search and mainly focus on textual search (e.g., Web search). Modern search engines can automatically identify the specific kind of resource that a query may reflect (like maroon 5 for music) and integrate these results into search result pages, which are referred to as vertical results. Although existing researches have

already invested a lot to support search among multi-modal and multi-source data, we believe that the ability of searching for ubiquitous data is still far below our expectation.

In the near future, query is not restricted to textual content. Some intents can be described by language while some others are not. For example, it is difficult to describe the shape of maple leaf, or the smell of grass. Search for ubiquitous data is facing the following important problems:

1) **Intent understanding** Intent understanding is always the bottleneck of search techniques. Search engines need to identify the specific kind of data (such as images, video, and etc.) that users may want to find.

2) **Cross-modal representation** To support cross-modal search, data resource, as well as users' queries need to be represented in a unified semantic space.

3) **Whole-page optimization** By far most of existing ranking models are explicitly or implicitly based on PRP framework, i.e., ranking the documents by their relevance probabilities. With more and more vertical results embedded in the search result page (SERP), it is a necessity to optimize the utility of the entire SERP, rather than the utility of each individual result.

Search in ubiquitous scenarios

We believe with explosive growth of data, IR needs to provide support in more and more scenarios. As we stated before, smart search engines will be incorporated into more mobile devices, such as cellphone, car, even television and refrigerator at home. Privacy is another important issue in ubiquitous IR research. On the one hand, everything generated by users (including queries, viewed content, etc.) may contain users' privacy. On the other hand, search in various scenarios needs to handle a lot of users' private data, for example, email, photo, notes. These private data may be stored locally, and search engines can only apply their retrieval algorithms without knowing the exact content.

9.2.3 Research challenges

We believe that in ubiquitous search, there are a few important research challenges as follows:

- For various search devices, it is important to extend existing "keyword - results" paradigm to various forms. For example, users should be able to search with voice, images, or videos.
- For heterogeneous data, the central problem is to build a cross-modal representation to support search tasks. Though existing studies have already made the first steps, there are still a lot of problems in this field.
- For ubiquitous scenarios, the privacy preserved search may raise research challenges in a number of domains, such as information retrieval and network security.

The overview of ubiquitous information retrieval is illustrated in Fig. 2. To summarize, we think that in the next few years, the edge between different modals will dissolve. It is promising to develop ubiquitous information retrieval techniques to help people collect information effectively and efficiently.

10 Further suggestions

10.1 Education of IR

IR is an important course for students in the majors of cyberspace security, computer science, information science, etc. The research scope of IR includes large scale content crawling, analyzing, organizing and accessing. It also includes the use of natural language processing (NLP), machine learning (ML) and data mining techniques for content processing. The recent advances and trends for IR focus on the combination of other technologies. On one hand, through the learning of IR course, students are able to understand the fundamental theories, models and algorithms of IR, which can be the basis for future research. On the other hand, through the practice, reading classic and advanced research papers, the research ability of students can be trained for future work on intelligent information processing, big data analyzing and processing or practical work and further enrich the professionals in IR.

According to the recent progresses of IR, the content of IR course should also include: the fundamental concepts and knowledge of IR, the combination of NLP, ML and IR, multimedia retrieval and the applications of IR.

For teaching skills, IR course emphasizes the fundamental and advanced knowledge. Teachers are required to master the basic knowledge and classic methods as well as introduce advanced development of IR. They should also leverage diverse teaching methods, such as discussion in class, recommending further reading of textbook and conference proceedings, training the presentation skills of students and synchronizing with the international level. IR is a course that has the aspects of practice and application. Experiments are especially important, and we call for more collaborations with industrial people in teaching for better effect.

10.2 Open source communities

Open source platforms often play a crucial role in the development of new technologies. For example, the open source platforms for deep learning (e.g., Tensorflow [50], Caffe [51], Pytorch [52], etc.) simplify the building of complex deep learning models, which makes deep learning accessible to everyone. IR needs such platforms too. Lucene [53] is an excellent project for IR application. However, it no longer meets the needs of recent IR researches. Especially, when we want to develop the search engines under new IR models, e.g., the model for interactive information retrieval.

The new platforms should have some essential highlights to

attract IR researchers and developers. First, the new platforms should be flexible to support new IR models or patterns. Second, the new platforms should take the recent advances in IR (e.g., deep learning-based IR especially) into account. Third, the platform should be scalable, reliable and upgradable to meet the demands of distributed computational environments.

Clearly, the unification and standardization for developing new IR platforms should be discussed in depth first. To build such a platform, we first need to make a thorough investigation of actual application and research requirements. The investigation reports should include detailed requirement analysis about what the researchers need to do their experiments and what the developers need to build their systems. Then, we can define flexible architectures and self-contained modules based on the investigation results. Especially, some common used parts should be provided independently. For example, the soft modules for Web crawler, document indexing, ranking and re-ranking, training module, etc. After that, we can provide well-defined APIs so that the IR community can bring more and more new models to the platform based on the APIs, overcoming the limits of the existing infrastructure. Most importantly, we should setup some incentive policies to attract the volunteers who can share the costs of producing such a powerful open source search system.

Usually the search engines needed by the middle and small-size enterprises are vertical search engines or the search engines with special characteristics. If we want to set up a search engine platform, which can help the middle and small-size enterprises to develop their search engines, the different soft modules of the platform should be developed as independently as possible, and the platform should be tailorable. Besides, it is better to provide some helpful tools and platforms for the volunteers to contribute their APIs and the forum for exchange their experience and expertise.

Currently one of the major challenges in the research of IR is the lack of data, especially the user behavior data. The IR community should collect data and release them for the open research purpose. The data can be either contributed by the community members or collected by the open IR platform.

11 Conclusions

Information retrieval remains a fundamental way for users to explore the big data and to access information, facts, and knowledge. It is also evolving very fast, bringing in the new

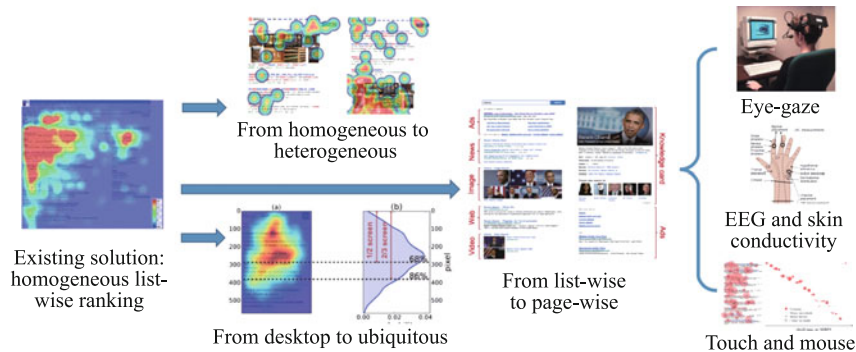


Fig. 2 Ubiquitous information retrieval

possibilities in redefining users, information, retrieval, and evaluation. In this report, we have taken a serious look at these possibilities, and suggested many important research themes for future study. We hope that this strategic report could inspire our IR researchers in both academia and industry, and could help the future growth of the IR research community.

Acknowledgements We would like to thank Chinese Information Processing Society of China, CAS Key Laboratory of Network Data Science and Technology, Institute of Computing Technology, Chinese Academy of Sciences, and ACM SIGIR Beijing Chapter for supporting the strategic workshop. Thank Professor Bo Zhang (Tsinghua University) and Ming Zhou (Microsoft Research Asia) for contributing the keynotes and valuable discussions in the workshop.

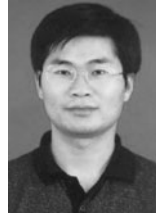
References

- Bush V. As we may think. *The Atlantic Monthly*, 1945, 176(1): 101–108
- Clarke C. From the chair... *ACM SIGIR Forum*, 2016, 50(1): 1
- Zobel J, Moffat A. Inverted files for text search engines. *ACM Computing Surveys (CSUR)*, 2006, 38(2): 6
- Salton G, Wong A, Yang C S. A vector space model for automatic indexing. *Communications of the ACM*, 1975, 18(11): 613–620
- Robertson S, Zaragoza H. The probabilistic relevance framework: Bm25 and beyond. *Foundations and Trends® in Information Retrieval*, 2009, 3(4): 333–389
- Lv Y, Zhai C. Positional language models for information retrieval. In: *Proceedings of the 32nd International ACM SIGIR Conference on Research and Development in Information Retrieval*. 2009, 299–306
- Zhai C, Lafferty J. A study of smoothing methods for language models applied to ad hoc information retrieval. *ACM SIGIR Forum*, 2017, 51(2): 268–276
- Page L, Brin S, Motwani R, Winograd T. The pagerank citation ranking: bringing order to the web. Technical Report, Stanford InfoLab, 1999
- Kleinberg J M. Authoritative sources in a hyperlinked environment. *Journal of the ACM*, 1999, 46(5): 604–632
- Chen C P, Zhang C Y. Data-intensive applications, challenges, techniques and technologies: a survey on big data. *Information Sciences*, 2014, 275: 314–347
- Sanderson M, Croft W B. The history of information retrieval research. *Proceedings of the IEEE*, 2012, 100 (Special Centennial Issue): 1444–1451
- Chaudhuri S, Dayal U. An overview of data warehousing and olap technology. *ACM Sigmod Record*, 1997, 26(1): 65–74
- Borlund P. The IIR evaluation model: a framework for evaluation of interactive information retrieval systems. *Information Research*, 2003, 8(3): 289–291
- Hinton G, Deng L, Yu D, Dahl G, Mohamed A R, Jaitly N, Senior A, Vanhoucke V, Nguyen P, Kingsbury B. Deep neural networks for acoustic modeling in speech recognition. *IEEE Signal Processing Magazine*, 2012, 29(6): 82–97
- LeCun Y, Bengio Y. Convolutional networks for images, speech, and time series. *The Handbook of Brain Theory and Neural Networks*, 1995, 3361(10): 1995
- Socher R, Huang E H, Pennin J, Manning C D, Ng A Y. Dynamic pooling and unfolding recursive autoencoders for paraphrase detection. In: *Proceedings of Advances in Neural Information Processing Systems*. 2011, 801–809
- Craswell N, Croft W B, Guo J, Mitra B, de Rijke M. Neu-IR: the SIGIR 2016 workshop on neural information retrieval. In: *Proceedings of the 39th International ACM SIGIR Conference on Research and Development in Information Retrieval*. 2016, 1245–1246
- Craswell N, Croft W B, de Rijke M, Guo J, Mitra B. SIGIR 2017 workshop on neural information retrieval (Neu-IR'17). In: *Proceedings of the 40th International ACM SIGIR Conference on Research and Development in Information Retrieval*. 2017, 1431–1432
- Goodfellow I, Pouget-Abadie J, Mirza M, Xu B, Warde-Farley D, Ozair S, Courville A, Bengio Y. Generative adversarial nets. In: *Proceedings of Advances in Neural Information Processing Systems*. 2014, 2672–2680
- Mnih V, Kavukcuoglu K, Silver D, Rusu A A, Veness J, Bellemare M G, Graves A, Riedmiller M, Fidjeland A K, Ostrovski G, Petersen S, Beattie C, Sadik A, Antonoglou I, King H, Kumaran D, Wierstra D, Legg S, Hassabis D. Human-level control through deep reinforcement learning. *Nature*, 2015, 518(7540): 529–533
- Silver D, Schrittwieser J, Simonyan K, Antonoglou I, Huang A, Guez A, Hubert T, Baker L, Lai M, Bolton A, Chen Y, Lillicrap T, Hui F, Sifre L, Driessche G V D, Graepel T, Hassabis D. Mastering the game of go without human knowledge. *Nature*, 2017, 550(7676): 354
- Wang J, Yu L, Zhang W, Gong Y, Xu Y, Wang B, Zhang P, Zhang D. Irgan: a minimax game for unifying generative and discriminative information retrieval models. In: *Proceedings of the 40th International ACM SIGIR Conference on Research and Development in Information Retrieval*. 2017, 515–524
- Agichtein E, Brill E, Dumais S. Improving web search ranking by incorporating user behavior information. In: *Proceedings of the 29th Annual International ACM SIGIR Conference on Research and Development in Information Retrieval*. 2006, 19–26
- Granka L A, Joachims T, Gay G. Eye-tracking analysis of user behavior in www search. In: *Proceedings of the 27th Annual International ACM SIGIR Conference on Research and Development in Information Retrieval*. 2004, 478–479
- Morris M R, Teevan J, Panovich K. What do people ask their social networks, and why?: a survey study of status message q&a behavior. In: *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*. 2010, 1739–1748
- Croft W B, Cronen-Townsend S, Lavrenko V. Relevance feedback and personalization: a language modeling perspective. In: *Proceedings of the 2nd DELOS Network of Excellence Workshop on Personalisation and Recommender Systems in Digital Libraries*. 2001
- Thomee B, Lew M S. Interactive search in image retrieval: a survey. *International Journal of Multimedia Information Retrieval*, 2012, 1(2): 71–86
- Said A, Jain B J, Narr S, Plumbaum T. Users and noise: the magic barrier of recommender systems. In: *Proceedings of International Conference on User Modeling, Adaptation, and Personalization*. 2012, 237–248
- Swan M. *Blockchain: Blueprint for a New Economy*. O'Reilly Media, Inc., 2015
- Akyildiz I F, Akan Ö B, Chen C, Fang J, Su W. Interplanetary internet: state-of-the-art and research challenges. *Computer Networks*, 2003, 43(2): 75–112
- Lavanya B M. Blockchain technology beyond bitcoin: an overview. *International Journal of Computer Science and Mobile Applications*, 2018, 6(1): 76–80
- Seebacher S, Schüritz R. Blockchain technology as an enabler of service systems: a structured literature review. In: *Proceedings of International Conference on Exploring Services Science*. 2017, 12–23
- Croft W B, Metzler D, Strohman T. *Search Engines: Information Retrieval in Practice*. Addison-Wesley Reading, 2010
- Voorhees E M, Harman D K. *TREC: Experiment and Evaluation in Information Retrieval*. Cambridge: MIT Press, 2005
- Kelly D. Methods for evaluating interactive information retrieval systems with users. *Foundations and Trends® in Information Retrieval*, 2009, 3(1–2): 1–224
- Ellis D. Theory and explanation in information retrieval research. *Journal of Information Science*, 1984, 8(1): 25–38

37. Vakkari P, Järvelin K. Explanation in information seeking and retrieval. *New Directions in Cognitive Information Retrieval*, 2006, 19: 113–138
38. Singh J, Anand A. EXS: explainable search using local model agnostic interpretability. In: *Proceedings of the 12th ACM International Conference on Web Search and Data Mining*. 2019, 770–773
39. Luo G, Tang C, Yang H, Wei X. Medsearch: a specialized search engine for medical information retrieval. In: *Proceedings of the 17th ACM Conference on Information and Knowledge Management*. 2008, 143–152
40. Huang P S, He X, Gao J, Deng L, Acero A, Heck L. Learning deep structured semantic models for Web search using clickthrough data. In: *Proceedings of the 22nd ACM International Conference on Information & Knowledge Management*. 2013, 2333–2338
41. Guo J, Fan Y, Ai Q, Croft W B. A deep relevance matching model for ad-hoc retrieval. In: *Proceedings of the 25th ACM International Conference on Information and Knowledge Management*. 2016, 55–64
42. Zhang Y, Rahman M M, Braylan A, Dang B, Chang H L, Kim H, McNamara Q, Angert A, Banner E, Khetan V, McDonnell T, Nguyen A T, Xu D, Wallace B C, Leasey M. Neural information retrieval: a literature review. 2016, arXiv preprint arXiv:1611.06792
43. Mitra B, Craswell N. Neural models for information retrieval. 2017, arXiv preprint arXiv:1705.01509
44. Guo J, Fan Y, Pang L, Yang L, Ai Q, Zamani H, Wu C, Croft W B, Cheng X. A deep look into neural ranking models for information retrieval. 2019, arXiv preprint arXiv:1903.06902
45. Sharma D, Kumar S, Kholia C. Multi-modal information retrieval system. US Patent 7,054,818, 2006
46. Lee D, Park J, Ahn J H. On the explanation of factors affecting e-commerce adoption. In: *Proceedings of the International Conference on Information Systems*. 2001, 109–120
47. Jamali M, Ester M. A matrix factorization technique with trust propagation for recommendation in social networks. In: *Proceedings of the 4th ACM Conference on Recommender Systems*. 2010, 135–142
48. Callison-Burch C. Fast, cheap, and creative: evaluating translation quality using amazon's mechanical turk. In: *Proceedings of the 2009 Conference on Empirical Methods in Natural Language Processing*. 2009, 286–295
49. Gubbi J, Buyya R, Marusic S, Palaniswami M. Internet of Things (IoT): a vision, architectural elements, and future directions. *Future Generation Computer Systems*, 2013, 29(7): 1645–1660
50. Abadi M, Barham P, Chen J, Chen Z, Davis A, Dean J, Devin M, Ghemawat S, Irving G, Isard M, Kudlur M, Levenberg J, Monga R, Moore S, Murray D G, Steiner B, Tucker P, Vasudevan V, Warden P, Wicke M, Yu Y, Zheng X. Tensorflow: a system for large-scale machine learning. In: *Proceedings of the 12th USENIX Symposium on Operating Systems Design and Implementation*. 2016, 265–283
51. Jia Y, Shelhamer E, Donahue J, Karayev S, Long J, Girshick R, Guadarrama S, Darrell T. Caffe: convolutional architecture for fast feature embedding. In: *Proceedings of the 22nd ACM International Conference on Multimedia*. 2014, 675–678
52. Paszke A, Gross S, Chintala S, Chanan G. Pytorch: tensors and dynamic neural networks in python with strong GPU acceleration. 2017
53. McCandless M, Hatcher E, Gospodnetic O. Lucene in Action: Covers Apache Lucene 3.0. Greenwich, CT: Manning Publications Co., 2010



Zhumin Chen is an associate professor in School of Computer Science and Technology, Shandong University, China. His research interests include information retrieval and natural language processing. His research is supported by the Natural Science Fund of China, Key Science and Technology Innovation Project of Shandong Province, etc.



Xueqi Cheng is a full professor and vice director of the Institute of Computing Technology, Chinese Academy of Sciences (CAS), and the director of the CAS Key Laboratory of Network Data Science and Technology, China. His research areas include Web search and data mining, data science, and social media analytics. He is the general secretary of CCF Committee on Big Data, the vice-chair of CIPS Committee on Information Retrieval, the general co-chair of SIGIR'20 and WSDM'15. He is the principal investigator of more than 10 major research projects, funded by NSFC and MOST. He was awarded the NSFC Distinguished Youth Scientist (2014), the National Prize for Progress in Science and Technology (2012), the China Youth Science and Technology Award (2011).



Shoubin Dong received the PhD degree in electronic engineering from the University of Science and Technology of China (USTC), China in 1994. She was a visiting scholar at the School of Computer Science and Language Technology Institute, Carnegie Mellon University (CMU), Pittsburgh, USA from 2001 to 2002. She is now a professor with the School of Computer Science and Engineering, South China University of Technology (SCUT), China. Her research interests include information retrieval, natural language processing, high-performance computing, etc.



Zhicheng Dou is currently a professor at School of Information, Renmin University of China, China. He received his PhD and BS degrees in computer science and technology from the Nankai University, China in 2008 and 2003, respectively. He worked at Microsoft Research Asia from July 2008 to September 2014. His current research interests are information retrieval, natural language processing, and big data analytics.



Jiafeng Guo is a professor in Institute of Computing Technology, Chinese Academy of Sciences, and University of Chinese Academy of Sciences, China. He has worked on a number of topics related to web search and data mining. His current research is focused on representation learning and neural models for information retrieval and filtering. He has published more than 80 papers in several top conferences/journals. His work on IR has received the Best Paper Award in ACM CIKM (2011), Best Student Paper Award in ACM SIGIR (2012) and Best Full Paper Runner-up Award in ACM CIKM (2017). Moreover, he has served as the PC member for the prestigious conferences including SIGIR, WWW, KDD, WSDM, and ACL, and the associate editor of TOIS.



Xuanjing Huang is a professor of the School of Computer Science, Fudan University, China. Her research interest includes natural language processing, information retrieval, artificial intelligence, deep learning and data intensive computing. She has published more than 100 papers in major conferences including ACL, SIGIR, IJCAI, AAAI, NIPS, ICML, CIKM, EMNLP, WSDM, and COLING. In the research community, she served as the PC Co-

Chair of CCL 2019, NLPCC 2017, CCL 2016, SMP 2015, and SMP 2014, the organizer of WSDM 2015, competition chair of CIKM 2014, tutorial chair of COLING 2010, SPC or PC member of past WSDM, SIGIR, WWW, CIKM, ACL, IJCAI, KDD, EMNLP, COLING, and many other conferences.



Yanyan Lan is a professor in Institute of Computing Technology, Chinese Academy of Sciences, China. She leads a research group working on Big Data and Machine Learning. Her current research interests include machine learning, information retrieval and natural language processing. From April 2018 to March 2019, she acted as a visiting scholar in the department of statistics, UC

Berkeley. She has published over 70 papers on top conferences including ICML, NIPS, SIGIR, WWW, etc. Her paper entitled “Top-k Learning to Rank: Labeling, Ranking, and Evaluation” has won the Best Student Paper Award of SIGIR 2012, and the paper entitled “Learning Visual Features from Snapshots for Web Search” has won the Best Paper RunnerUp Award of CIKM2017.



Chenliang Li received PhD from Nanyang Technological University, Singapore in 2013. Currently, he is an associate professor at School of Cyber Science and Engineering, Wuhan University, China. His research interests include information retrieval, text/web mining, data mining and natural language processing. He is a co-recipient of Best Student Paper Award Honorable Mention in ACM

SIGIR 2016, and serves as an editorial board member of JASIST and IPM.



Ru Li, Professor, PhD Supervisor. She is the Vice-dean of the School of Computer and Information Technology, and the School of Big Data of Shanxi University, the standing council member of Chinese Information Processing Society (CIPS), the committee member of Computational Linguistics, Information Retrieval, and Language and Knowledge Computing of CIPS. Her research interests

include Chinese information processing and information retrieval. She has published more than 70 papers in both international and national important academic journals and conferences, including in the IEEE Transactions on Knowledge and Data Engineering, the Annual Meeting of the Association for Computational Linguistics, and the International Conference on Computational Linguistics, and so on. She has won three Second Prize for scientific and technological progress in Shanxi.



Tie-Yan Liu, assistant managing director of Microsoft Research Asia, fellow of the IEEE, and distinguished member of the ACM. He is an adjunct professor at Carnegie Mellon University (CMU) and Tsinghua University. His research interest includes learning to rank, deep learning, reinforcement learning, and distributed learning. He has

served as general chair, program committee chair, local chair, or area chair for a dozen of international conferences including WWW/WebConf, SIGIR, KDD, ICML, NIPS, IJCAI, AAAI, ACL, ICTIR, as well as associate editor of ACM Transactions on Information Systems, ACM Transactions on the Web, and Neurocomputing.



Yiqun Liu is professor and Department co-Chair at the Department of Computer Science and Technology in Tsinghua University, China. His major research interests are in Web search, user behavior analysis, and natural language processing. He also works as a visiting research professor of National University of Singapore and a visiting professor of National Institute of Informatics (NII) of Japan, as well as a member of Tiangong AI Research Center which is founded by Tsinghua and Sogou Inc.



Jun Ma received his BSc, MSc, and PhD degrees from Shandong University in China, Ibaraki University and Kyushu University in Japan, respectively. He is a full professor at Shandong University. He was a senior researcher in the Department of Computer Science at Ibaraki University in 1994 and German GMD and Fraunhofer from the year 1999 to 2003. His research interests include information retrieval, Web data mining, recommendation systems, and machine learning. He has published more than 150 papers in the international journals and conferences, including SIGIR, WWW, MM, TOIS, and TKDE. He is a member of the ACM and IEEE.



Bing Qin, a professor and doctoral supervisor of the School of Computer Science and Technology, at Harbin Institute of Technology, China. Her main research directions are natural language processing, information extraction, text mining, sentiment analysis, etc. She has published more than 80 papers in the several international top journals and conferences such as ACL, COLING, EMNLP, IEEE TKDE, IEEE TASLP, etc. She has led over several the National Natural Science Foundations of China, as well as the key research and development projects of the National Science and Technology Ministry. She was awarded the first prize of Qian Weichang Chinese Information Processing Science and Technology Award by the Chinese Information Processing Society and the second prize of Heilongjiang Province Technical Invention.



Mingwen Wang is currently a professor of Jiangxi Normal University, China. He received the BS (1985) and MS (1988) degrees in mathematics from Jiangxi Normal University, China, and the PhD (2000) degree in computer science from Shanghai Jiaotong University, China. His research interests include information retrieval, natural language processing, and machine learning.



Jirong Wen is a professor and the dean of School of Information, Renmin University of China (RUC), China. He is also the Executive Dean of Gaoling School of Artificial Intelligence, and Director of Beijing Key Laboratory of Big Data Research. He received his PhD degree in 1999 from the Institute of Computing Technology, the Chinese Academy of Science, China. His main research interests include information retrieval, data mining and machine learning. He worked at Microsoft Research Asia (MSRA) for 14 years and once was the group manager of the Web Search and Mining Group.



Jun Xu is a professor with the School of Information, Renmin University of China, China. His research interests include learning to rank and semantic matching in search. He has published more than 50 papers in international conferences (e.g., SIGIR, WSDM) and journals (e.g., ACM TOIS, IEEE TKDE). He has won the Best Paper Award in AIRS (2010), Best Paper Runner-up in CIKM (2017), and Test of Time Award Honorable Mention in SIGIR (2019). He is serving as SPC for SIGIR, WWW, AAAI, and ACML, editorial board member for JASIST, and associate editor for ACM TIST.



Min Zhang is a tenured associate professor in the Department of Computer Science & Technology, Tsinghua University, China, specializes in Web search and recommendation, and user modeling. She is the vice director of Intelligent Technology & Systems lab at CS Dept., and vice director of Intelligent Information Acquisition, AI Institute, Tsinghua University. She also serves as ACM SIGIR Executive Committee Member, Associate Editor for the ACM Transaction of Information Systems (TOIS), Web Mining and Content Analysis Track Chair of the WebConf 2020, Short Paper Chair of SIGIR 2018, Program Chair of WSDM 2017, etc. She also owns 12 patents.



Peng Zhang is an associate professor and Vice Dean of School of Computer Science and Technology, College of Computing and Intelligence, Tianjin University, China. He obtained his PhD at Robert Gordon University, United Kingdom in 2013. His research is focused on the tensor space language models, explainable neural network design and quantum theory inspired language models. He has published papers on refereed journals such as IEEE TNN, IEEE TKDE, ACM TIST, ACM TALIP, JASIST, IP&M, and on top conferences such as NeurIPS, SIGIR, ACL, AAAI, IJCAI, CIKM, WWW, and EMNLP. He won ECIR 2011 Best Poster Award and SIGIR 2017 Best Paper Award Honorable Mention.



Qi Zhang received the PhD degree in computer science from Fudan University, China. He is a professor of computer science at Fudan University, China. His research interests include natural language processing and information retrieval.