

Un silenzio assordante

Original

Un silenzio assordante / Morriello, R.. - In: BIBLIOTECHE OGGI. - ISSN 0392-8586. - STAMPA. - 5(2006), pp. 73-73.

Availability:

This version is available at: 11583/2743098 since: 2019-07-22T14:04:54Z

Publisher:

Editrice Bibliografica

Published

DOI:

Terms of use:

This article is made available under terms and conditions as specified in the corresponding bibliographic description in the repository

Publisher copyright

(Article begins on next page)



Overview of the PIR Track at FIRE 2024: Evaluation of Personalised Information Retrieval

Pranav Kasela
University of Milano-Bicocca
Milan, Italy
pranav.kasela@unimib.it

Gian Carlo Milanese
University of Milano-Bicocca
Milan, Italy
giancarlo.milanese@unimib.it

Alessandro Raganato
University of Milano-Bicocca
Milan, Italy
alessandro.raganato@unimib.it

Marco Braga
University of Milano-Bicocca
Milan, Italy
m.braga@campus.unimib.it

Georgios Peikos
University of Milano-Bicocca
Milan, Italy
georgios.peikos@unimib.it

Marco Viviani
University of Milano-Bicocca
Milan, Italy
marco.viviani@unimib.it

Effrosyni Sokli
University of Milano-Bicocca
Milan, Italy
effrosyni.sokli@unimib.it

Sandip Modha
University of Milano-Bicocca
Milan, Italy
sandip.modha@unimib.it

Gabriella Pasi
University of Milano-Bicocca
Milan, Italy
gabriella.pasi@unimib.it

Abstract

This abstract provides a short overview of the first edition of the shared task on Personalised Information Retrieval (PIR) organized at the 16th Forum for Information Retrieval Evaluation (FIRE 2024). A more detailed discussion of the approaches used by the participating teams is available in the track overview paper. PIR 2024 consisted of two sub-tasks. The first sub-task aims to explore the personalisation in cQA based on user profiles, following the standard IR pipeline. The second one, instead, aims to investigate the personalisation in cQA based on user profiles using recent LLMs and prompt engineering. Although the tasks saw an enthusiastic response in registrations, with 10 teams requesting the dataset, only 1 team finally submitted the runs, and 2 of them submitted the working notes.

CCS Concepts

• **Information systems** → **Information retrieval**; **Personalization**.

Keywords

Information Retrieval, Large Language Model, Personalization, Question Answering

ACM Reference Format:

Pranav Kasela, Marco Braga, Effrosyni Sokli, Gian Carlo Milanese, Georgios Peikos, Sandip Modha, Alessandro Raganato, Marco Viviani, and Gabriella

Pasi. 2024. Overview of the PIR Track at FIRE 2024: Evaluation of Personalised Information Retrieval. In *Proceedings of the 16th Annual Meeting of the Forum for Information Retrieval Evaluation (FIRE '24)*, December 12–15, 2024, Gandhinagar, India. ACM, New York, NY, USA, 3 pages. <https://doi.org/10.1145/3734947.3735665>

1 Introduction

The aim of the proposed shared tasks in FIRE 2024 is to elaborate on activities in the evaluation of Personalised Information Retrieval (PIR). Personalisation in Information Retrieval (IR) remains an important topic both in research and the development of practical applications for Information Retrieval. However, suitable means of evaluation of PIR remain an underexplored topic. PIR-FIRE aims to bring together researchers with a common interest in developing and evaluating novel methods for personalised operation of novel information retrieval applications.

Personalisation [3, 4, 6] and other adaptation of the search experience [7] to the user and search context is an important topic in IR, studied by the research community for a long time. Personalised search aims to tailor the search outcome to a specific user (or group of users) based on the knowledge of their interests and online behaviour. Unfortunately, text collections shared by initiatives on search evaluation do not usually provide the information needed for evaluating personalisation, i.e., information about specific users and their preferences. At the same time, the field has witnessed a transformation with the recent availability of pre-trained Large Language Models (LLMs). Despite recent research that has emphasized the advantages and drawbacks of personalizing LLMs [9], the development and evaluation of LLMs specifically designed to generate personalized responses remains inadequately explored.

For this reason, we propose running the PIR-FIRE laboratory shared task with the goal of organizing and introducing a new benchmark dataset coupled with user information allowing the development of personalised approaches, putting forward two sub-tasks: (i) Personalised Information Retrieval, (ii) Personalised Prompt-based Information Retrieval.



This work is licensed under a Creative Commons Attribution International 4.0 License.

FIRE '24, Gandhinagar, India

© 2024 Copyright held by the owner/author(s).

ACM ISBN 979-8-4007-1318-7/24/12

<https://doi.org/10.1145/3734947.3735665>

2 Task Definition

The first edition of PIR consisted of the following two sub-tasks:

2.1 Task 1: Standard IR

The cQA task will be tackled as a standard ad-hoc IR task, where the questions are going to be considered as the queries, and the collection, from which the answers will be retrieved, is composed by all the answers available in the dataset. In this case, personalization can be tackled using any standard or novel technique to create a user profile and inject it in the retrieval model. We plan to provide multiple baselines that utilize, as first stage retrievers, both classical approaches such as BM25, and neural approaches based on BERT-like models [5]. As a second stage, we plan to provide re-rankers, using cross-encoders, like Mono-T5 [8], for non-personalized baselines, and for personalized baselines, using of a mix of tags and historical documents related to the users and weighted according to their importance for the current question.

2.2 Task 2: Prompt-based IR

Differently from the second stage of the standard IR task, the proposed prompt-based baselines personalise the results by using models like Phi [1] and GPT [2] with prompts similar to the following one: “To which degree between 0 and 1 does the document [DOCUMENT] answer the question [QUESTION], and is relevant to a user with the following profile [USER PROFILE]”, where the [USER PROFILE] is a series of user interests that are inferred from their activities and ordered according to their timestamp (most recent first) and importance.

More details about how this dataset was created can be found in the original resource paper [6].

3 Dataset and Evaluation

In this section, we discuss the datasets for each sub-task and the evaluation metrics used for each of them.

The PIR-FIRE will use data from StackExchange, a very popular community Question Answering (cQA) platform. The data is publicly available¹ under a cc-by-sa 4.0 license. The dataset is composed of questions, and their answers, collected from fifty communities, which can be categorized under the large umbrella of humanistic communities. In Table 1 we report the basic statistics for the dataset. Specifically: *document length*, measured in the number of words, *document score*, which is the difference between the number of up- and down-votes assigned by the community; *answers' count*, the number of answers given to a question; *comments' count*, the number of user comments to a given question or answer; *favorite count*, that indicates the number of users that flagged the question as their favorite, showing their interest in that topic; *tags count*, the number of tags associated to the question by the asking user.

The dump is curated and merged to tackle the cQA task as a retrieval task. The PIR-FIRE test collections provide the traditional components used in IR experiments, i.e. access to a document collection, search topics, and corresponding relevance judgments. Regarding the judgments, we only consider relevant the single answer

Table 1: Basic feature statistics for question and answers.

	Type	mean	std	median	25%	75%	99%
Questions #docs = 1,125,407	Length	125.69	112.90	94	60	153	553
	Score	5.13	10.73	2	1	6	45
	Answers Count	1.93	1.92	1	1	2	9
	Comments Count	2.78	3.37	2	0	4	15
	Favorite Count	2.07	4.88	1	1	2	16
	Tags Count	2.45	1.21	2	1	3	5
Answers #docs = 2,173,139	Length	178.15	210.09	117	61	218	1000
	Score	5.13	12.43	2	1	5	51
	Comments Count	1.62	2.62	1	0	2	12

that is explicitly labelled as the best answer by the user who submitted the question. In addition, our PIR evaluation test collections [6] are accompanied by user-related information for modelling and introducing profiles in evaluation experiments. The user-related information includes the text and the number of views of documents they have generated, and in many cases also the tags associated with these documents, the date since they are registered on the website, the badges they obtained, their reputation score, and some times also their autobiography. The information collected as previously explained can be used for personalising and adapting the search process to the current user, e.g. by creating and exploiting personal user profiles.

3.0.1 Evaluation setup. We will provide the participants with several baselines, including keyword and dense-based representations of user profiles (anonymised information gathered about individual users) as part of our data collection. For this shared task, we will use traditional evaluation metrics in the IR literature that can be applied also to personalized search, Precision (P), Recall (R), Mean Average Precision (MAP), Mean Reciprocal Rank (MRR), and (normalized) Discounted Cumulative Gain (nDCG).

4 Participation

A total of 10 teams registered across both sub-tasks. Only one team submitted runs for task 1. For both tasks, 2 teams submitted the working notes.

5 Conclusion

The Personalised Information Retrieval (PIR) track at FIRE'24 focus on the evaluation of Personalised Information Retrieval (PIR), which remains an important topic both in research and the development of practical applications.

In future efforts, we plan to address the challenges encountered by making datasets more manageable, enhancing promotion, and broadening the scope of personalization to include diverse tasks in Information Retrieval (IR) and possibly Natural Language Processing (NLP). By providing resources such as pre-trained models, smaller datasets, and novel tasks, we aim to encourage a stronger focus on personalization across varied domains. These steps will help attract a broader range of participants and methodologies, driving greater engagement and impact.

References

- [1] Marah Abdin, Sam Ade Jacobs, Ammar Ahmad Awan, Jyoti Aneja, Ahmed Awadallah, Hany Awadalla, Nguyen Bach, Amit Bahree, Arash Bakhtiari, Harkirat Behl,

¹<https://zenodo.org/records/10679181>

- et al. 2024. Phi-3 technical report: A highly capable language model locally on your phone. *arXiv preprint arXiv:2404.14219* (2024).
- [2] Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, et al. 2023. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774* (2023).
- [3] Marco Braga. 2024. Personalized large language models through parameter efficient fine-tuning techniques. In *Proceedings of the 47th International ACM SIGIR Conference on Research and Development in Information Retrieval*. 3076–3076.
- [4] Marco Braga, Pranav Kasela, Alessandro Raganato, and Gabriella Pasi. 2024. Synthetic Data Generation with Large Language Models for Personalized Community Question Answering. *arXiv preprint arXiv:2410.22182* (2024).
- [5] Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2018. Bert: Pre-training of deep bidirectional transformers for language understanding. *arXiv preprint arXiv:1810.04805* (2018).
- [6] Pranav Kasela, Marco Braga, Gabriella Pasi, and Raffaele Perego. 2024. SE-PQA: Personalized Community Question Answering. In *Companion Proceedings of the ACM on Web Conference 2024 (Singapore) (WWW '24)*. Association for Computing Machinery, New York, NY, USA, 1095–1098. <https://doi.org/10.1145/3589335.3651445>
- [7] Pranav Kasela, Gabriella Pasi, Raffaele Perego, and Nicola Tonellotto. 2024. DESIRE-ME: Domain-Enhanced Supervised Information Retrieval Using Mixture-of-Experts. In *European Conference on Information Retrieval*. Springer, 111–125.
- [8] Rodrigo Nogueira, Zhiying Jiang, Ronak Pradeep, and Jimmy Lin. 2020. Document Ranking with a Pretrained Sequence-to-Sequence Model. In *Findings of the Association for Computational Linguistics: EMNLP 2020*, Trevor Cohn, Yulan He, and Yang Liu (Eds.). Association for Computational Linguistics, Online, 708–718. <https://doi.org/10.18653/v1/2020.findings-emnlp.63>
- [9] Alireza Salemi, Sheshera Mysore, Michael Bendersky, and Hamed Zamani. 2023. LaMP: When Large Language Models Meet Personalization. *arXiv:2304.11406* [cs.CL]