

Dan-yang JIANG, Hong-hui CHEN, 2019. Cohort-based personalized query auto-completion. *Frontiers of Information Technology & Electronic Engineering*, 20(9):1246-1258. <https://doi.org/10.1631/FITEE.1800010>

Cohort-based personalized query auto-completion

Key words: Query auto-completion; Cohort-based retrieval; Topic models

Corresponding author: Dan-yang JIANG

E-mail: danyangjiang@nudt.edu.cn

 ORCID: <http://orcid.org/0000-0001-6012-5313>

Motivation

Non-personalized QAC models are far from optimal because the one-size-fits-all approach fails to take users' context information, such as submitted queries and click-through data, into consideration while such information often influences users' intended queries.

Personalized QAC models are effective only when there is a large amount of data available about individuals. Unfortunately, users' context is often very sparse and insufficient to identify their interest and intentions.

Main idea

We build an individual's interest profile by learning his/her topic preferences through topic models and then identify users who share similar profiles.

We use cohorts' contextual information together with query frequency to rank completions.

Method

1. Cohort topic models

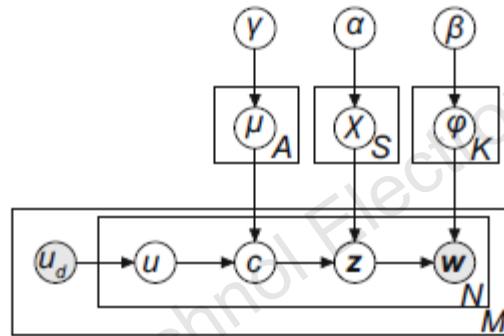


Fig. 5 Graphical representation for cohort topic model 1 (CTM1)

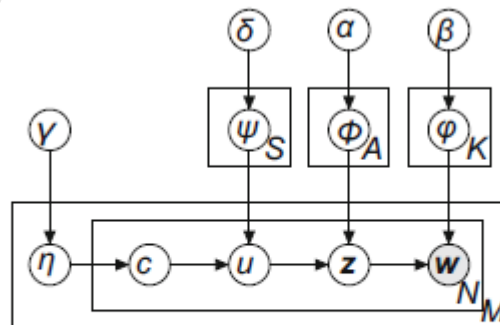


Fig. 6 Graphical representation of cohort topic model 2 (CTM2)

2. Ranking models

$$\text{Score}(q) = \lambda \text{FreqScore}(q) + (1 - \lambda) \text{CoScore}(q),$$

$$\begin{cases} \text{FreqScore}(q) = \frac{f(q)}{\max_{q_c \in C(p)} f(q_c)}, \\ \text{CoScore}(q) = \text{norm}(\omega_j) \cdot \text{sim}(q, q_j), \end{cases}$$

$$\text{sim}(a, y) = \begin{cases} 1/D(a, y), & \text{LDA or ATM,} \\ \mu_{a,s} \cdot \mu_{y,s}, & \text{CTM1,} \\ \psi_{s,a} \cdot \psi_{s,y}, & \text{CTM2,} \end{cases}$$

3. Experimental setup

- (1) Dataset: AOL query log;
- (2) Evaluation metric: mean reciprocal rank (MRR);
- (3) Baselines: MPC, P-QAC, C-QAC, and N-QAC.

Major results

Table 3 Mean reciprocal rank (MRR) results of different QAC models for prefix p consisting of 1–5 characters

p	MRR							
	MPC	P-QAC	C-QAC	N-QAC	LDA-QAC	ATM-QAC	CTM1-QAC	CTM2-QAC
1	0.0981	<u>0.3535</u>	0.3287	0.3365	0.3600	0.3534	Δ 0.3670 [▲]	Δ 0.3684 [▲]
2	0.1851	<u>0.4434</u>	0.4365	0.4413	0.4497 ^Δ	0.4448	Δ 0.4524 [▲]	Δ 0.4553 [▲]
3	0.3165	<u>0.5246</u>	0.5183	0.5207	0.5280 ^Δ	0.5258	Δ 0.5313 ^Δ	Δ 0.5348 [▲]
4	0.4249	0.5925	0.5947	<u>0.5963</u>	0.5920	0.5931	Δ 0.5969 ^Δ	Δ 0.5997 ^Δ
5	0.4921	0.6355	0.6364	<u>0.6381</u>	0.6318 [▽]	0.6356	Δ 0.6385 ^Δ	Δ 0.6407 ^Δ
Overall	0.2991	<u>0.5071</u>	0.4992	0.5035	0.5096	0.5077	Δ 0.5145 ^Δ	Δ 0.5170 [▲]

The best results of baselines and of all models in each row are underlined and boldfaced, respectively. The statistical significance of pairwise differences is determined by the t -test [▲] for $p < 0.01$ and ^{Δ/▽} for $p < 0.05$

- Cohort-based personalized QAC models had advantages over competitive baselines on relevance ranking.
- Models based on CTMs produced better ranking results than models using the conventional topic models and the K -means algorithm in identifying cohorts.

Major results

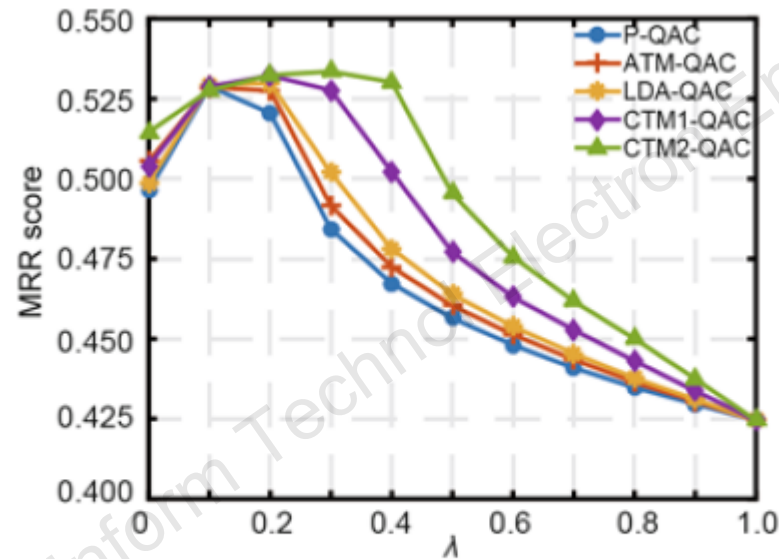


Fig. 7 Performance of the five personalized QAC rankers when varying the contribution weight λ from 0 to 1

The performance of QAC models depending only on query similarity was much better than that of models considering frequency alone.

Major results

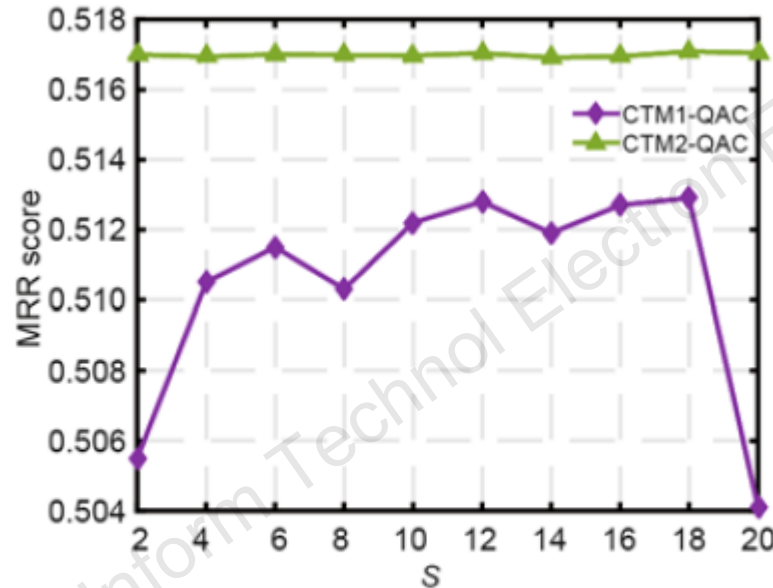


Fig. 8 Performance of CTM1-QAC and CTM2-QAC at various cohort numbers in cohort discovery

- MRR score of CTM1-QAC evidently fluctuated with the change of S .
- CTM2-QAC demonstrated insensitivity over different S values.

Major results

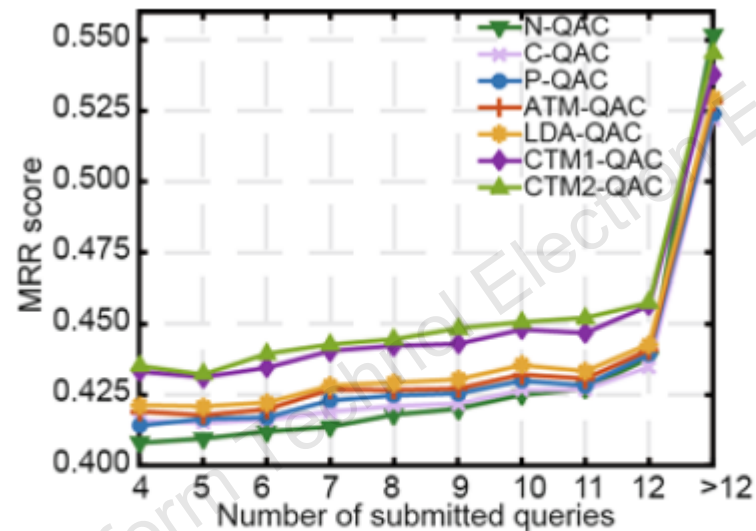


Fig. 9 Performance of seven personalized QAC rankers for users who submit different numbers of queries

On the whole, the greater the query volume, the better the ranking accuracy of the personalized QAC models.

Conclusions

Cohort-based personalized QAC models greatly improve ranking effectiveness with the mean reciprocal rank (MRR) score increased by 1.5% compared with the personalized QAC without cohorts.

The personalized rankers based on CTMs can appropriately tackle the context sparseness problem while remaining robust.