

LAPI @ 2016 Retrieving Diverse Social Images Task: A Pseudo-Relevance Feedback Diversification Perspective

Bogdan Boteanu, Mihai Gabriel Constantin, Bogdan Ionescu
LAPI, University “Politehnica” of Bucharest, Romania
{bboteanu, mgconstantin, bionescu}@alpha.imag.pub.ro

ABSTRACT

In this paper we present the results achieved during the 2016 MediaEval Retrieving Diverse Social Images Task, using an approach based on pseudo-relevance feedback, in which human feedback is replaced by an automatic selection of images. The proposed approach is designed to have in priority the diversification of the results, in contrast to most of the existing techniques that address only the relevance. Diversification is achieved by exploiting a hierarchical clustering scheme followed by a diversification strategy. Methods are tested on the benchmarking data and results are analyzed. Insights for future work conclude the paper.

1. INTRODUCTION

An efficient information retrieval system should be able to provide search results which are in the same time *relevant* for the query and cover different aspects of it, i.e., *diverse*. The 2016 Retrieving Diverse Social Images Task [1] addresses this issue in the context of a general ad-hoc image retrieval system, which provides the user with diverse representations of the queries. The system should be able to tackle complex and general-purpose multi-concept queries. Given a ranked list of photos retrieved from Flickr¹, participating systems are expected to refine the results by providing up to 50 images that are in the same time relevant and provide a diversified summary of the query. The process is based on the social metadata associated with the images and/or on the visual characteristics. A complete overview of the task is presented in [1].

Despite the current advances of machine intelligence techniques used in the area of information retrieval and multimedia, in search for achieving high performance and adapting to user needs, more and more research is turning now towards the concept of “*human in the loop*” [2]. The idea is to bring the human expertise in the processing chain, thus combining the accuracy of human judgements with the computational power of machines.

In this work we use an adapted version of the work in [4] that exploits the concept of pseudo-relevance feedback (RF). RF techniques attempt to introduce the user in the loop by harvesting feedback about the relevance of the search results. This information is used as ground truth for re-computing a better representation of the data needed. Relevance feedback proved efficient in improving the precision of the results [7], but its potential was not fully exploited to diversification. The main contribution of our approach is in proposing a pseudo-relevance feedback technique which sub-

¹<http://flickr.com/>.

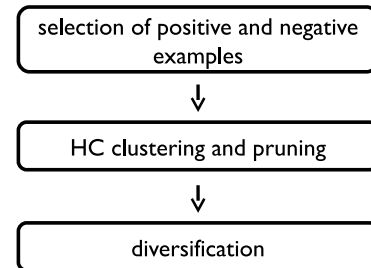


Figure 1: General scheme of the proposed approach

stitutes the user needed in traditional RF and in proposing several diversity-adapted relevance feedback schemes.

2. PROPOSED APPROACH

In traditional RF techniques, recording actual user feedback is inefficient in terms of time and human resources. The proposed approach, denoted in the following *HC-RF*, attempts to replace user input with machine generated ground truth. It exploits the concept of pseudo-relevance feedback. The concept is based on the assumption that top k ranked documents are relevant and the feedback is learned as in traditional RF under this assumption [5]. A general diagram of the approach is depicted in Figure 1.

This year we decided not to use the pre-processing step, i.e., the use of filters to remove un-relevant images. The motivation relies on the specificity of the dataset proposed for this year [1], in particular the use of multi-topic queries in the development and evaluation sets. An image containing people or depicting a location or a place which is geographically far away from the query, can be considered relevant as long as it is a common photo representation of the query topics (all at once). Also, we noticed in an extensive study [6] that the blur filter does not improve significantly the overall performance, thus we decided to remove it to reduce complexity. The algorithm is as follows.

First, we employ a pseudo-relevance feedback scheme based on the selection of the images assessed in an automated manner. We consider that most of the first returned results are relevant (i.e., positive examples). For instance, on *devset* [1], in average, 35 out of 50 returned images are relevant which supports our assumption. In contrast, the very last of the results are more likely non-relevant and considered accordingly (i.e., negative examples). The positive and negative examples are fed to an Hierarchical Clustering² scheme

²<http://www.mathworks.com/help/stats/hierarchical-clustering.html>

Table 1: Best pseudo-relevance feedback results for each modality or combination of modalities on devset (best results are depicted in bold).

metric/run	HC-RF visual	HC-RF text	HC-RF vis-text	HC-RF cred.	HC-RF All	Flickr init. res.
$P@20$	0.6514	0.7479	0.7479	0.645	0.7479	0.6979
$CR@20$	0.446	0.4572	0.4572	0.4274	0.4572	0.3717
$F1@20$	0.5202	0.5468	0.5468	0.5014	0.5468	0.4674

which yields a dendrogram of classes. For a certain cutting point (i.e., number of classes), a class is declared non-relevant if contains only negative examples or the number of negative examples is higher than the positive ones. The final step is the actual diversification scheme, which is a round robin approach. We select from each of the relevant classes one image which has the highest rank according to the initial ranking of the system. Then, we remove the selected images from the clusters and proceed by selecting the remaining ones in the same manner. The process is repeated until a maximum number of images is reached. The resulting images represent the output of the proposed system.

3. EXPERIMENTAL RESULTS

This section presents the experimental results achieved on *devset* which consists of 70 multi-topic queries and 20,757 images and *testset*, respectively, which consists in 64 multi-topic queries and 19,017 images. We optimized the parameters of the proposed approaches on *devset* to obtain best precision and diversity. Ground truth was also provided with the data for this set for preliminary validation of the approaches. The final benchmarking is conducted however on *testset*.

In our approaches, images are represented with the content descriptors that were provided with the task data, i.e., visual (e.g., convolutional neural network based descriptors), text (e.g., term frequency - inverse document frequency representations of metadata) and user annotation credibility (e.g., face proportions, upload frequency) information. Detailed information about provided content descriptors is available in [1]. Performance is assessed with Precision at X images ($P@X$), Cluster Recall at X ($CR@X$) and F1-measure at X ($F1@X$).

3.1 Results on devset

Several tests were performed with different descriptor combinations and various cutoff points. Descriptors are combined with an early fusion approach. We varied the number of initial images considered as positive examples, from 100 to 280 with a step of 20 images, the number of last images considered as negative examples, from 0 to 20 with a step of 10, and the inconsistency coefficient threshold for which HC divides the data into well-separated clusters, from 0.5 to 1.3 with a step of 0.2. We select the combinations yielding the highest $F1@20$, which is the official metric.

While experimenting, we observed that, by increasing the number of analyzed images, precision tends to slightly decrease as the probability of obtaining un-relevant images increases; in the same time, diversity increases as having more images is more likely to get more diverse representations. For brevity reasons, in the following we focus on presenting only the results at a cutoff of 20 images which is the official cutoff point. These results are presented in Table 1. To serve as baseline for the evaluation, we present also the

Table 2: Results for the official runs on testset (best results are depicted in bold).

metric/run	<i>Run1</i>	<i>Run2</i>	<i>Run3</i>	<i>Run4</i>	<i>Run5</i>
$P@20$	0.5437	0.5758	0.5219	0.5484	0.5539
$CR@20$	0.3455	0.3881	0.3868	0.4374	0.3732
$F1@20$	0.4037	0.447	0.4208	0.4638	0.4223

Flickr initial retrieval results. From the modality point of view, text descriptors (TF, DF, TFIDF), separately and combined with other descriptors (all visual-all textual, all visual-all textual-all credibility), lead to the highest results ($F1@20=0.5468$), followed closely by visual (CNN-ad) descriptors ($F1@20=0.5202$) and then by all credibility information ($F1@20=0.5014$).

3.2 Official results on testset

Following the previous experiments, the final runs were determined for the best modality/parameter combinations obtained on *devset* (see Table 1). We submitted five runs, computed as following: *Run1* - automated using visual information only: HC-RF visual CNN-ad; *Run2* - automated using text information only: HC-RF all text; *Run3* - automated using visual-text information: HC-RF all visual-all text; *Run4* - everything allowed: HC-RF all cred.; and *Run5* - everything allowed: HC-RF all visual-all text-all cred. Results are presented in Table 2.

What is interesting to observe is the fact that the highest precision is achieved using textual information, (*Run2* - $P@20 = 0.5758$), whereas maximum diversification is achieved using credibility information (*Run4* - $CR@20 = 0.4374$), with more than 2% over other types of descriptors, although this type of information had the lowest performance on devset in terms of diversification. Another interesting observation is that relevance was also preserved in this case, compared to other modalities, which leads to the conclusion that credibility information was useful in the context of overall diversification. Credibility information gives an automatic estimation of the quality of tag-image content relationships, telling which users are most likely to share relevant images in Flickr. Best diversification is achieved, $CR@20 = 0.4374$, due to the high probability that different relevant images belong to different users with a good credibility score. In terms of $F1$ metric score, the use of credibility information, *Run4* - $F1@20 = 0.4638$, allows for better performance over all text descriptors by almost 2% and by 6% over visual descriptor (CNN-ad).

4. CONCLUSIONS

We approached the image search result diversification issue from the perspective of relevance feedback techniques, when user feedback is substituted with an automatic pseudo-relevance feedback approach. Results show that in general, the automatic techniques improve the precision and diversification, which proves the real potential of relevance feedback to the diversification. Future developments will mainly address different efficient exploitations of re-ranking approaches, e.g., relevance-score estimation techniques, to improve the relevance and consequently the overall diversification. Another perspective is to also exploit the advantages of deep neural networks and use them in the context of automatic relevance-feedback-based diversification scenarios, by classifying the selected positive and negative examples using unsupervised deep-learning-based classifiers.

5. REFERENCES

- [1] B. Ionescu, A.L. Gînscă, M. Zaharieva, B. Boteanu, M. Lupu, H. Müller, “*Retrieving Diverse Social Images at MediaEval 2016: Challenge, Dataset and Evaluation*”, MediaEval 2016 Workshop, October 20-21, Hilversum, Netherlands, 2016.
- [2] B. Emond, “*Multimedia and Human-in-the-loop: Interaction as Content Enrichment*”, ACM Int. Workshop on Human-Centered Multimedia, pp. 77–84, 2007.
- [3] B. Boteanu, I. Mironică, B. Ionescu, “*A Relevance Feedback Perspective to Image Search Result Diversification*”, IEEE ICCP, September 4-6, Cluj-Napoca, Romania, 2014.
- [4] B. Boteanu, I. Mironică, B. Ionescu, “*LAPI @ 2016 Retrieving Diverse Social Images Task: A Pseudo-Relevance Feedback Diversification Perspective*”, MediaEval 2015 Workshop, September 15-16, Wurzen, Germany, 2016.
- [5] B. Boteanu, I. Mironică, B. Ionescu, “*Hierarchical Clustering Pseudo-Relevance Feedback for Social Image Search Result Diversification*”, IEEE CBMI, June 10-12, Prague, Czech Republic, 2015.
- [6] B. Boteanu, I. Mironică, B. Ionescu, “*Pseudo-Relevance Feedback Diversification of Social Image Retrieval Results*”, Multimedia Tools and Applications, pp. 1–28, Springer 2016.
- [7] J. Li, N.M. Allinson, “*Relevance Feedback in Content-Based Image Retrieval: A Survey*”, Handbook on Neural Information Processing, 49, pp. 433–469, Springer 2013.