

UMass at TREC 2004: Novelty and HARD

Nasreen Abdul-Jaleel, James Allan, W. Bruce Croft, Fernando Diaz, Leah Larkey, Xiaoyan Li,
Mark D. Smucker, Courtney Wade

Center for Intelligent Information Retrieval
Department of Computer Science
University of Massachusetts
Amherst, MA 01003

Abstract

For the TREC 2004 Novelty track, UMass participated in all four tasks. Although finding relevant sentences was harder this year than last, we continue to show marked improvements over the baseline of calling all sentences relevant, with a variant of tfidf being the most successful approach. We achieve 5–9% improvements over the baseline in locating novel sentences, primarily by looking at the similarity of a sentence to earlier sentences and focusing on named entities.

For the High Accuracy Retrieval from Documents (HARD) track, we investigated the use of clarification forms, fixed- and variable-length passage retrieval, and the use of metadata. Clarification form results indicate that passage level feedback can provide improvements comparable to user supplied related-text for document evaluation and outperforms related-text for passage evaluation. Document retrieval methods without a query expansion component show the most gains from related-text. We also found that displaying the top passages for feedback outperformed displaying centroid passages. Named entity feedback resulted in mixed performance. Our primary findings for passage retrieval are that document retrieval methods performed better than passage retrieval methods on the passage evaluation metric of binary preference at 12,000 characters, and that clarification forms improved passage retrieval for every retrieval method explored. We found no benefit to using variable-length passages over fixed-length passages for this corpus. Our use of geography and genre metadata resulted in no significant changes in retrieval performance.

1 Introduction

The University of Massachusetts Amherst participated in three tracks this year. This report discusses work done on the Novelty and High Accuracy Retrieval from Documents (HARD) tracks. Work on the Terabyte track is reported elsewhere [24].

2 Novelty

2.1 Overview of Our Approaches for the Four Tasks

There are four tasks in this year's novelty track and we participated in all of them. For the 50 topics in the 2004 track, each of them has 25 relevant documents, and zero or more non-relevant documents. Task 1 was to identify all relevant and novel sentences, given the full set of documents for the 50 topics. Task 2 was to identify all novel sentences, given the full set of relevant sentences in all documents. Task 3 was to find the relevant and novel sentences in the remaining documents, given the relevant and novel sentences in the first 5 documents only. Task 4 was to find the novel sentences in the remaining documents, given all relevant sentences from all documents and the novel sentences from the first 5 documents.

We compared the statistics of the 2004 track with both 2002 and 2003 tracks, and have found that the statistics of the 2004 track is closer to the 2003 track. The comparison of the statistics of the 2003 and 2004 novelty track data is shown in 1. Therefore we decided to train our system with the 2003 data when no training from this year's track was available for Task 1 and Task 2, and used the training data from this year's track as it was available for Task 3 and task4. We have already developed an answer-updating approach to novelty detection [1], which gave better performance in terms precision at low recall on both the 2002 and the 2003 novelty track data than the baseline approaches reported in that work. However, we could not use the answer-updating approach directly in the tasks of this year's novelty track because the evaluation measure used in novelty track was the F measure, which is the harmonic mean of precision and recall. Therefore, we used TFIDF techniques with selective feedback for finding relevant sentences and considered the maximum similarity of a sentence to its previous sentences and new named entities to identify novel sentences. The detail descriptions about our approaches are elaborated in the following subsections. Only the main approach for each task will be

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE 2004		2. REPORT TYPE		3. DATES COVERED 00-00-2004 to 00-00-2004	
4. TITLE AND SUBTITLE UMass at TREC 2004: Novelty and HARD				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Center for Intelligent Information Retrieval, Department of Computer Science, University of Massachusetts Amherst, Amherst, MA, 01003-9264				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited					
13. SUPPLEMENTARY NOTES					
14. ABSTRACT					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT	18. NUMBER OF PAGES 13	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

Feature	Track 2003	Track 2004
Num. of Event Topics	28	25
Num. of Opinion Topics	22	25
Num. of Relevant Documents/Topic	25	25
Num. of Non-relevant Documents/Topic	0	11.16
Avg. Num. Sentences/Topic	797.4	1048.8

Table 1: Statistics comparison of 2003 and 2004 track data

Approaches	F-score (2003)	F-score (2004)
0. The Original full set of sentences	0.5398	0.303
1. TFIDF models with pseudo feedback	0.6429	0.393 (CIIRT1R1)
2. TFIDF models with selective pseudo feedback	0.6593	0.395 (CIIRT1R2)

Table 2: Performance of finding relevant sentences in Task 1 on 2003 and 2004 data

Approaches	F-score(2003)	F-score(2004)
0. The Original full set of sentences	0.5271	0.306
1. TFIDF models with relevance feedback	0.6229	0.405 (CIIRT3R2)
2. TFIDF models with selective relevance feedback	0.6554	0.406(CIIR31R1)

Table 3: Performance of finding relevant sentences in Task 3 on 2003 and 2004 data

reported in this paper even though multiple runs for each topic were submitted to TREC from us.

2.2 Relevant Sentence Retrieval

For relevant sentence retrieval, our system treated sentences as documents and used the words in the title fields of the topics as queries. TFIDF techniques with pseudo feedback or selective pseudo feedback were used for finding relevant sentences for Task 1 and TFIDF techniques with relevance feedback or selective relevance feedback were used for Task3. Selective pseudo feedback means pseudo feedback was performed on some queries but not on other queries based on an automatic analysis on query words across different topics. Basically, a query with more focused query words that rarely appear in relevant documents related to other queries is likely to have a better performance without pseudo feedback. Selective relevance feedback means whether to performance relevance feedback on a query was determined by the comparison between the performance with and without relevance feedback in the top five documents for this query because the judgment of the top five documents was given for Task 3. Short sentences, non-informative sentences as well as non-normal sentences were removed in the final results. Non-informative sentences are the sentences that have less than n non-stopwords, where the best value of n is 3 (which was learned from the 2003 data). Sentences that have less than m terms are short sentences, where the best value of m is 7 from the 2003 data. Non-normal sentences refer to some special formats for some purposes other than offering the information about the

story discussed in a news story. In addition to short sentences, non-informative sentences and non-normal sentences, sentences similar to given non-relevant sentences were also removed for Task 3 when partial judgment was available. Basically if the maximum similarity between a sentence and given non-relevant sentences is greater than a preset threshold (which was trained with the 2003 data), the sentence was treated as non-relevant sentence and thus removed from the result list.

The performance of finding relevance sentences using our approaches on the 2003 and 2004 data for Task1 and Task 3 are given in Table 2 and Table 3 respectively. There are three conclusions that can be drawn from the results. First, the F scores of the original full set of sentences show that how difficult the task is on different data set. It is clear to us that the task of finding relevant sentences on the 2004 data is more difficult than that on the 2003 data. Second, TFIDF techniques work well for relevant sentences retrieval on both the 2003 and 2004 data sets. Third, selective feedback gives better performance than applying feedback on all queries on the two data sets.

2.3 Identifying Novel Sentences

Similarities of a sentence to its previous sentences and the occurrence of new named entities in the sentence are two main factors considered in our approach to identifying novel sentences. New named entities have been used successfully in our answer-updating approach in novelty detection [21].

For Task 1 and Task3, our system started with the list of sentences returned from the relevant sentences retrieval,

Approaches	Starting set of sentences	Identify novel sentences
F-scoreTask 1 (Ch%)	0.195	0.211(+8.2%)
F-scoreTask 2 (Ch%)	0.577	0.610(+5.7%)
F-scoreTask 3 (Ch%)	0.194	0.210(+8.2%)
F-scoreTask 4 (Ch%)	0.541	0.577(+6.7%)

Table 4: Performance of identifying novel sentences for Tasks 1-4

which unavoidably contains many non-relevant sentences in addition to relevant sentences. For Task 2 and Task 4, our system started with the set of given relevant sentences only. In either case, the cosine similarity between a sentence and each its previous sentence was calculated. The maximum similarity of a sentence to its previous sentences was used to eliminate redundant sentences. Sentences with a maximum similarity value greater than a preset threshold may be treated as redundant sentences. The value of the same threshold for all topics was tuned with the TREC 2003 track data when no training data from this year’s was available. The value of the threshold for each topic was trained with the training data when the judgment of the top five documents was given for Task 3 and Task 4. In addition to the maximum similarity between a sentence and its previous sentences, new named entities were also considered in identifying novel sentences. A person’s name or an organization in a sentence that did not appear in the previous sentences may give new information about who was related to an event or an opinion [21]. Therefore, a sentence with previously unseen named entities was treated as novel sentences. About 20 types of named entities were considered in our system, which included PERSON, LOCATION, ORGANIZATION, DATE and MONEY, etc. BBN’s IdentiFinder [2] and our approach [21] were used for identifying named entities.

The performance of identifying novel sentences for Task 1, Task 2, Task 3 and Task 4 on the 2004 novelty track data are given in Table 4. The F-score on the starting set of sentences (as described above) for each task establishes a bottom line for performance of a novelty detection algorithm. The F-scores were evaluated when we simply assumed all the sentences were novel (without any novelty detection). Any successful novelty detection approach should beat the F-score bottom-line for each task. Table 4 shows that the F-scores of our approaches have significant increases from the bottom lines for all the four tasks.

3 HARD

UMass explored four different sub-tasks in the course of HARD 2004: fixed-length passage retrieval, variable-length passage retrieval, metadata, and clarification form feedback.

First, we generate a clarification form and receive user feedback. Using the response, the first clarification form module constructs a new, possibly modified query representation. Depending on the retrieval element, the query representation is passed to either a passage retrieval module or a document retrieval module. Both of these modules return a ranked list of items (passages or documents). These items are then re-ranked based upon the satisfaction of topic metadata value. As a post-processing step, the ranked list is further altered by feedback elicited from the clarification form.

3.1 Methods and Materials

3.1.1 COLLECTION PROCESSING

We processed the HARD collection differently for retrieval and metadata classification. For both retrieval and classification, only text between the <TITLE> and <TEXT> tags were handled.

For retrieval, tokenization was based on non-alphanumeric characters. If a token was not in a list of Acrophile [18] acronyms, then it was down-cased. If a down-cased token was in the libbow stopword list [23], then it was ignored. The Krovetz stemmer [15] packaged with Lemur [1] was used to stem all remaining down-cased words. The topics and related text metadata were processed in the same manner with the additional processing step that `http://` URLs were automatically stripped from the related text.

For metadata classification, contiguous digits were replaced by a token representing a number. The paragraph tag, <P>, was retained as a token. Quotation marks, “ ” and “ ’ ”, were converted to the double quote mark, “ ”. Contractions were pulled off and became their own tokens (n’t, ’s, ’d, ’m, ’ll, ’ve, and ’re). All punctuation was treated as separate tokens. All remaining text was down-cased and broken at whitespace boundaries.

3.1.2 TRAINING TOPICS

The LDC supplied training data consists of 21 topics. For each topic, the LDC judged the top 100 documents returned by their search system. We augmented the training topics with additional judgments by obtaining in-house judgments on an additional 100 documents for each topic. This expanded set of judgments was used for parameter tuning.

3.1.3 QUERY REPRESENTATION

A query model refers to a probability distribution over words representing the user’s information need. In the simplest case, we have the maximum likelihood query model based on the the user’s title and description fields. Here, the we would process the text according to section 3.1.1 and then form a maximum likelihood language model using remaining terms as evidence.

3.1.4 RETRIEVAL USING LANGUAGE MODELS

A description of retrieval using language models is beyond the scope of this document. We refer readers to the several papers on the subject [4]. We used a modified version of the Lemur language modeling toolkit to perform retrieval [1].

It has been shown that query likelihood and divergence ranking using a maximum likelihood query model are equivalent [17]. Therefore, without loss of generality, we confine our description to divergence-based retrieval. In this approach, we take a query model, $P(w|Q)$, and rank all documents in the collection according to the Kullback-Leibler divergence with $P(w|Q)$,

$$score(D, Q) = \sum_w P(w|Q) \log \frac{P(w|Q)}{P(w|D)} \quad (1)$$

Here, the document language model, $P(w|D)$, may be estimated using a number of different techniques [33]; smoothing parameters used will be described whenever language model retrieval is used.

In addition to the maximum likelihood query model presented in section 3.1.3, we also used *relevance models* for query representation [19]. Relevance models are a form of massive query expansion through blind feedback. Constructing a relevance model entails first ranking the collection according to the maximum likelihood query model. Some set of documents at the top of this ranking become evidence for the relevance model, $P(w|\mathcal{R})$. If we call this set $\mathcal{R}$, then the relevance model is estimated according to,

$$P(w|\mathcal{R}) = \sum_{D \in \mathcal{R}} P(w|D) \frac{P(Q|D)}{\sum_{D' \in \mathcal{R}} P(Q|D')} \quad (2)$$

where the query likelihood score, $P(Q|D)$, can be easily computed from the divergence measure [27]. The relevance model replaces the maximum likelihood query model in a second round of document ranking.

Ideally, we would include the entire collection in the set $\mathcal{R}$ and, therefore, $P(w|\mathcal{R})$ would have no terms with zero probability. However, computational limitations force us to let $|\mathcal{R}|$ be fixed; that is, we only consider the top N documents. Furthermore, we also *truncate* and *normalize* the relevance model to include only the M terms with highest probability. The first parameter, N , does not affect the

estimation of the relevance model since we are normalizing the query likelihoods. The second parameter, M , requires a little explanation. First, we compute the relevance model as in Equation 2. Second, we order the terms in $P(w|\mathcal{R})$ in decreasing order of probability. Third, we select the top M terms from this ordering. Finally, we normalize these term weights to sum to one.

A relevance model captures behavior of the returned documents but throws away the original query. In order to maintain information in the original query model, we linearly interpolate the relevance model with the original query model: $P'(w|\mathcal{R}) = \lambda P(w|\mathcal{R}) + (1 - \lambda)P(w|Q)$. For our runs where we do this, we specify λ . In our experiments the relevance model is truncated prior to interpolation with the query. Depending on the module, a second truncation and normalization process is performed in a similar manner.

3.1.5 RETRIEVAL USING SUPPORT VECTOR MACHINES

Of the runs that UMass submitted, several runs involved the use of support vector machines for passage or document retrieval [26]. This technique applies discriminative models to information retrieval. Previous work has demonstrated that the performance of support vector machines on the document retrieval task is on par with that of language models. Our document retrieval and passage retrieval experiments on HARD 2003 test queries and HARD 2004 training queries showed that using SVMs gave better results than traditional language models.

Support vector machines are a class of discriminative supervised learning models. SVMs used for classification create a hyperplane that maximizes the margin from the training examples. The discriminant function used to separate the two classes is given by: $g(R|D, Q) = \mathbf{w} \bullet \phi(\mathbf{f}(D, Q)) + b$, where R denotes the relevant class, D is a document, Q is a query, $\mathbf{f}(D, Q)$ is the vector of features, $\mathbf{w}$ is the weight vector in kernel space that is learned by the SVM from the training examples, $\bullet$ denotes inner product, b is a constant and ϕ is the mapping from input space to kernel space. The value of this discriminant function is proportional to the distance between the document D and the separating hyperplane in the kernel space.

The features are term-based statistics commonly used in information retrieval systems such as *tf*, *idf* and their combinations as shown in Table 5. Each of the six features is a sum over the query terms.

In order to provide a finer-grained weighting of query terms, we incorporated the query models described in Section 3.1.3 into our features. These hybrid features are presented in Table 5. Unless otherwise noted, all SVMs were trained using the regular features. In all cases, retrieval was performed using the hybrid features.

The corpora, queries and relevance judgments for

	Features	Hybrid Features
1	$\sum_{q_i \in Q \cap D} \log(c(q_i, D))$	$\sum_{w \in V} P(w Q) \log(c(w, D))$
2	$\sum_{i=1}^n \log(1 + \frac{c(q_i, D)}{ D })$	$\sum_{w \in V} P(w Q) \log(1 + \frac{c(w, D)}{ D })$
3	$\sum_{q_i \in Q \cap D} \log(idf(q_i))$	$\sum_{w \in V} P(w Q) \log(idf(w))$
4	$\sum_{q_i \in Q \cap D} (\log(\frac{ C }{c(q_i, C)}))$	$\sum_{w \in V} P(w Q) (\log(\frac{ C }{c(w, C)}))$
5	$\sum_{i=1}^n \log(1 + \frac{c(q_i, D)}{ D } idf(q_i))$	$\sum_{w \in V} P(w Q) \log(1 + \frac{c(w, D)}{ D } idf(w))$
6	$\sum_{i=1}^n \log(1 + \frac{c(q_i, D)}{ D } \frac{ C }{c(q_i, C)})$	$\sum_{w \in V} P(w Q) \log(1 + \frac{c(w, D)}{ D } \frac{ C }{c(w, C)})$

Table 5: Features in the discriminative models: $c(w, D)$ represents the raw count of word w in document D , C represents the collection, n is the number of terms in the query, $|\cdot|$ is the *size-of* function and $idf(\cdot)$ is the inverse document frequency. In the case of the hybrid features, $P(w|Q)$ refers to a query model as described in Section 3.1.3. We define $\log(0) = 0$.

TREC 1 and TREC 2 provided training data. All the documents marked relevant for a query were used as positive training instances. An equal number of negative instances were obtained by random sampling of the remaining documents. These training instances are represented in terms of their transformed feature vectors in the kernel space. The support vector machine then learns the hyperplane that separates the positive and negative training instances with the highest margin. For our runs, we used a linear kernel. Hence the hyperplane is drawn in the original feature space. The equation of this hyperplane provides the discriminant function $g(R|D, Q)$ that is subsequently used for scoring documents (or fixed length passages).

The indexed elements (documents or passages) are treated as instances in the feature space. For a test topic, Q , an instance D is scored based on the value of the discriminant function $g(R|D, Q)$. The instances are then ranked based on this score.

3.1.6 BOOTSTRAPPING SVMs

Previous work has balanced classes by random sampling from the negative training instances [26]. We propose another technique for instance sampling, which we refer to as *bootstrapping*. This method differs from the random sampling technique in the selection of negative training instances. All positive instances are used for training as in the previously described sampling method. In bootstrapping, negative instances are selected in the following way. First, an initial SVM is created using the technique described in section 3.1.5. Then, negative training instances are selected from only the negative examples *misclassified* by the initial SVM created in step 1. As many nega-

tive instances are selected as there are positive instances. This training set is used to create an SVM boundary as described in section 3.1.5.

Sampling from the set of misclassified negative documents, as opposed to sampling from all the negatives, will produce a set of negative training instances that are closer to the positive instances in the feature space. The intuition is that this will result in a boundary that is still good for ranking but has fewer misclassified instances on the positive side.

3.2 Clarification Form Feedback

This year’s HARD track again permitted sites to request one round of feedback from the topic creator. UMass studied four methods for eliciting user feedback. Different manifestations of these methods appeared on our submitted clarification forms.

3.2.1 CLARIFICATION FORM SUBSECTIONS

Passages Although the three minute time limit constrained our ability to request true document-level relevance judgments, we assumed that the presentation the most relevant *passages* retrieved would serve as an acceptable surrogate. Specifically, we performed SVM-based retrieval on a passage index comprised of 150-word overlapping passages. We used a linear model trained on TREC collections 1 and 2. We then split the top 15 document-unique passages into 25-word passages and selected the passage which the SVM scored the highest. These 15 25-word passages were then presented to the user with the document title and time stamp for feedback.

In addition to selecting the top 15 document-unique

150-word passages, we also experimented with using agglomerative clustering to remove redundancy from the passages presented. We used group-average, agglomerative clustering [22]. Term vectors were weighted according to a tf.idf scheme, $weight(x_i) = x_i / (\log((|C| + 1) / (0.5 + df_i)))$. Using these vectors, a cosine measure was used to compute the similarity matrix. We clustered clustered 200 150-word passages until a threshold similarity of 0.6 was reached. At that point, the largest 15 clusters were selected. The 15 150-word centroid passages of these clusters were then split into 25-word passages to be handled as above.

Query Reformulation Because the title and description subsections of the topic do not often serve as a good representation of a realistic user query, we allowed users to modify the stopped and stemmed version of their title and description query using a free entry text box.

Extracted Entities Previous work has shown that user feedback of term lists tends to have little (and sometimes negative) impact on retrieval performance [29]. We were interested in exploring the potential advantage of using different types of words as feedback candidates [16]. In particular, we were interested in the use of proper names rather than arbitrary terms as feedback sources. To accomplish this, we gathered the top 200 150-word passages after an initial retrieval and ran BBN's Identifinder across this set of passages [2]. We extracted the person, place, and organization names from this run and normalized the names by down-casing and removing punctuation and spaces. After removing names such as "New York Times", "AFP", and other source tags, we presented the user with the 15 most frequently occurring people, places, and organizations.

For each of these types of named entities, the user was also presented with a text box in which to enter named entities not in the top 15 for that type.

Temporal Feedback Previous work has shown that some topics demonstrate strong temporal structure [7, 20]. In order to elicit temporal biases in the information need, we asked the user for relevant months in the year spanned by the collection.

3.2.2 OFFICIAL CLARIFICATION FORMS

CF1 Our first clarification form included a list of 15 25-word passages derived from the top 15 150-word passages, a query reformulation text box, a free-text named entity text box, and a temporal feedback interface.

CF2 Our second clarification form included a list of 15 25-word passages derived from clustering, a query reformulation text box, a free-text named entity text box, and a temporal feedback interface.

CF3 Our third clarification form included the list of 15 people, 15 places, and 15 organizations with free-entry for each entity type, a temporal feedback interface, and a query reformulation text box.

3.2.3 INCORPORATION OF RESPONSES

Passages Passage feedback was used in two ways. First, we performed query expansion based upon the relevant passages. A query model was constructed by uniformly combining the language models of the relevant documents. We selected the top 200 terms from this distribution and renormalized the weights. This was our final query model for relevant passages. Secondly, passage feedback was used in order to re-rank documents at the end of the treatment. Specifically, we multiplied all final scores by 1 if they were from a document marked relevant, 0 if from a document marked non-relevant, and 0.5 otherwise.

Query Reformulation Whenever the user reformulated a query, we discarded the original query and constructed a query model from the new query strings.

Extracted Entities All relevant named entities or named entities entered in the free text box were combined to construct a named entity query model.

Temporal Feedback Temporal feedback was used in order to re-rank documents at the end of the treatment. Specifically, we multiplied all final scores by 1 if they were from a month marked relevant, 0 if from a month marked non-relevant, and 0.5 otherwise.

3.3 Fixed-Length Passage Retrieval

Passage retrieval was one of the issues that were studied as part of the HARD track. The central goal of the track was to perform high accuracy retrieval. Retrieving passages instead of whole documents could potentially return less non-relevant text at the top of the ranked list, thereby increasing the accuracy of the search.

Experiments on HARD 2003 test queries indicated that retrieval using 100 word half overlapping passages gave the best results. This was the passage size that was used for all the fixed-length passage experiments.

We explored various approaches to fixed-length passage retrieval. We studied the performance of passage retrieval systems that used query likelihood, relevance models and support vector machines. Passage retrieval using SVMs, described in 3.1.5 performed better than the other systems. We also explored the comparative utility of retrieving the best passages from top ranked documents versus indexing overlapping passages and scoring each of these independent of the document that the passage came from. The latter method gave higher precision

on our training data. Therefore, we scored pre-indexed passages for our final run.

3.4 Variable-Length Passage Retrieval

One of the questions UMass explored through the passage retrieval portion of the HARD track was whether retrieving passages of different lengths could improve our ability to return only the relevant portions of documents. In order to keep our text index relatively small and maintain the theoretical possibility that any passage of any document could be retrieved by the system, we chose to extract passages from highly ranked documents at the time of retrieval, rather than indexing particular passages in advance.

Previous work has been inconclusive as to whether there is benefit to retrieving passages of different lengths [13, 3, 25]. However, most past studies have only evaluated passage retrieval by its ability to retrieve relevant documents, due in part to the unavailability of passage-level relevance judgments. Now that the HARD track has provided passage judgments and the evaluation is based on more fine-grained retrieval, we decided to revisit this question.

3.4.1 EXTRACTING RELEVANT PASSAGES

Our method of extracting relevant passages from documents is inspired by work by de Kretser and Moffat [5], that assigned a relevance score to every word in a document. They used term frequency within the query and inverse term frequency in the corpus to determine the score of each word, and used several different functions to determine how much query terms contributed to the scores of surrounding words.

Our approach to selecting relevant passages is similar, in that each term from an expanded query representation is assigned a score which affects the scores of proximal words. However, the scores we use are derived from language models, and the task is somewhat different.

This process of extracting passages for a topic starts with the top-ranked documents from some document run and a language model representing the topic. Of the different topic models we tried, the best-performing one was a mixture model between the maximum likelihood representation of the original query and the top 50 terms from the relevance model for the query, as described in section 3.1.4.

We refer to the range of word positions in a document that a particular query word affects as its *region of influence*. The *spread* of a query term is the number of words before it and after it that that query term influences. Thus, the size of the region of influence is equal to $(2 \times \text{spread}) + 1$. This method takes as parameters the minimum spread and the maximum spread that any particular query term can have. The weights of the topic model are then linearly

scaled to fall between these minimum and maximum values. For all of our submitted runs that used passages of varying lengths, the minimum spread was 1 and the maximum spread was 25.

We extract any group of words that falls within the region of influence of any query term as a passage, discarding passages with fewer than 400 characters. Next we score the remaining passages as described in the following section.

3.4.2 SCORING PASSAGES

We experimented with several passage-scoring methods that fall into two basic classes. The first group used SVMs to score passages, as described in section 3.1.5. The second assigned scores equal to the negative relative entropy between the topic and passage language models, but differed in how the passage was modeled.

The SVM models did not perform as well as the relative entropy-based methods on the training data, regardless of which topic representation we used. For the class of relative-entropy-based measures, we tried three different topic models. The first used Dirichlet smoothing of the maximum-likelihood passage model with the collection model as the background model. The second used Dirichlet smoothing of the maximum-likelihood passage model with the document model as the background. Neither of these methods performed well on the training data.

The third and best-performing passage representation, used in UMassVPMM and UMassCVC, was a mixture of the collection, document, and passage models.

$$p(w|\Theta_{PSG}) = \lambda_c p(w|\Theta_{ML_c}) + \lambda_d p(w|\Theta_{ML_d}) + \lambda_p p(w|\Theta_{ML_p}) \quad (3)$$

Θ_{ML_c} , Θ_{ML_d} , and Θ_{ML_p} are the maximum likelihood collection, document, and passage models respectively. The three lambdas sum to 1. In our submitted runs, λ_c was 0.8, and the other two parameters were 0.1.

Future work will investigate the possibility of using two different topic models for the passage extraction and passage scoring stages of this technique.

3.5 Metadata

For metadata our approach was to take a ranked list of documents and rerank the list based on the topic’s metadata values. For the genre and geography metadata values we trained classifiers to determine to what degree a document satisfies the metadata value. Documents that better satisfy the metadata values are moved up in the ranked list compared to those that do not satisfy the metadata values.

3.5.1 DATA COLLECTION FOR CLASSIFIERS

We used several human annotators to obtain metadata judgments on documents from the collection. The majority of the judgments came from one of the authors. Table

Metadata	Pos.	Neg.	Total
Genre news-report	848	491	1339
Genre opinion-editorial	147	1192	1339
Genre other	344	995	1339
Geography US	590	758	1348

Table 6: Counts of human judgments collected for the genre and geography metadata broken down by positive and negative judgments.

Metadata	Pos.	Neg.	Total
Genre news-report	2603	2280	4883
Genre opinion-editorial	1633	3250	4883
Genre other	647	4236	4883
Geography US	1470	1451	2921

Table 7: Counts of judgments obtained by using the <KEYWORD> element of the documents to automatically guess a document’s genre and geography.

6 shows the breakdown of judgments obtained by humans for each metadata category.

To boost performance, we automatically extracted training data from the corpus using the corpus’ existing metadata. The AP wire, New York Times, and LA Times either contained explicit metadata in the <KEYWORD> element or was discernible in some other manner. The number of judgments collected in this mainly automatic fashion are shown in Table 7.

While we knew that this process would lead to mistakes, we did spot check the extracted documents, and we felt the gain from the additional training data exceeded the cost in misclassified examples. Also, we had counter balanced this automatically extracted data with over 1000 human judgments covering all subcollections.

3.5.2 CLASSIFIER TECHNOLOGY

We used linear support vector machines (SVMs) as our classifiers because of their success at text classification [32, 11, 8] and their ability to produce a ranking rather than merely a class prediction. The linear SVM learns a hyperplane in the feature space of the training examples that separates positive from negative examples. A document’s distance from the hyperplane determines the degree to which the SVM predicts the document is a positive or negative example of the learned class. We used *SVM^{light}* with its default settings compiled for Windows to perform all classification [10].

3.5.3 CLASSIFIER FEATURES

We used the same set of features for each of our classifiers. Our selection of features was guided by the choices others have used for the classification of text genre [12, 14, 30, 6, 9]. We used the 10K most frequently

Metadata	Avg. Prec.	Accuracy	F1
Genre news	0.99	0.96	0.96
Genre op-ed	0.97	0.95	0.91
Genre other	0.82	0.92	0.76
Geo. US	0.96	0.92	0.91

Table 8: This table describes the performance of SVM classifiers on the labeled data. All performance measures are averages from 3-fold cross validation. The class examples are oversampled so that positive examples comprise 50% of the training examples.

occurring tokens in the corpus. If a document contained one of these tokens, the corresponding feature value was 1 otherwise it was 0. We also used the out of vocabulary probability mass. The 10K most frequently occurring tokens constituted our vocabulary. We made eight binary features, one for each subcollection in the HARD collection: AFE, APE, CNE, LAT, NYT, SLN, UME, and XIE. Finally, we constructed a set of features focused on various length measures of a document: number of tokens, average token length, average sentence length, average paragraph length, variance in paragraph lengths, average corpus frequency of tokens, and four features that measured the number of words $\leq X$ characters long where X was one of 6,7,8, and 9. We normalized each of these measures to vary between 0 and 1. We first took the log of the sentence, paragraph, and document length features before normalizing them.

3.5.4 CLASSIFIER TRAINING

To deal with imbalances in the number of positive examples per class, we randomly oversampled from either the positive or negative examples, whichever was in the minority until 50% of the examples were positive [31]. No other special techniques were used.

The performance of the classifiers on the final datasets is shown in Table 8. We aimed to improve average precision, which measures ranking ability, while keeping an eye on the other measures. One could obtain a high average precision while doing poorly on accuracy.

While these metrics are certainly indicative of the classifiers’ power, some caveats must be stated. The HARD corpus contains many articles that are posted to the newswires multiple times in order to add more information or make small corrections. Our automatically judged articles may in fact contain several near copies of the same document. In addition, we included many examples from the same columnists. It is likely that a columnist’s pieces are more similar to each other than a selection of opinion pieces written by different authors. These duplicates can thus straddle the train and test sets of the 3-fold cross validation and artificially inflate the performance metrics.

3.5.5 METADATA RERANKING

We reranked the results based on a linear combination of the normalized outputs of both the retrieval and classifier outputs. We normalize each classifier’s output across the whole corpus. For each topic, the document scores were normalized with the rank 1 document score set to 1 and rank 1000 document score set to 0. We rerank passages as though they were documents.

We tuned the linear combination with a simple parameter sweep using the LDC hard-relevance training data augmented with additional UMass judgments. The best coefficients found weighted the original IR results at 0.4, geography at 0.3, and genre at 0.3.

3.5.6 USE OF RELATED TEXT

To utilize the related text metadata, we created a maximum likelihood model of the related text provided with the topic and linearly mixed this model with a model created for the title and description. This mixture model was used as the query. A parameter sweep was used to find the best mixture ratio on the training topics. The title and description model had a weight of 0.4 and the related text model had a weight of 0.6. We did not differentiate between on-topic and relevant related text and used both together.

3.6 HARD Runs

We submitted three baseline runs (UMassBaseQL, UMassBaseRM3, UMassBaseSVM) that did not use any of the metadata, clarification form, or passage techniques described earlier. Our other ten runs aimed to investigate the use of these techniques.

UMassBaseQL This run uses the maximum likelihood query model as described in section 3.1.4. It used both the title and the description. Smoothing was performed using the Dirichlet prior with its parameter set to 1000.

UMassBaseRM3 For this run, we used the title and description and the relevance modeling approach described in section 3.1.4. We used the first 50 documents retrieved to build the relevance model. The model was truncated to include only the 200 words of highest probability with a minimum probability of 0.001. The foreground model (the title and description) received a weight of 0.6 when mixed with the relevance model. Smoothing was performed using the Dirichlet prior with its parameter set to 1000.

UMassBaseSVM This run used a support vector machine built from the normal features in Table 5 to retrieve documents using a hybrid representation.

UMassMerge This run merged three different rankings. The first ranking used CF1 and all associated feedback.

This run used the passage feedback and reformulation for building a query model. A hybrid SVM was used for an initial retrieval. This ranked list was reranked using temporal and document feedback. List two is UMassF. List three was identical to UMassRGG for the document topics. For the passage topics, passages were reranked using the genre and geography metadata as described in section 3.5.5. The source of the passages came from the fixed length SVM passage retrieval used by run UMassCFMC, which used a query model produced by CF1 and related text. These passages were reranked prior to removal of overlap as opposed to the passages in UMassCFMC which were reranked after overlap in the passages had been removed. The three lists were each normalized and merged by summing the scores of identical documents or passages and ranked according to this sum. Overlap in passages were removed and the lists were trimmed to the top 1000 results.

UMassCFMC This run was a pipeline of the CF1 clarification form, bootstrapped SVM retrieval, and genre and geography metadata reranking. The linear bootstrapped model used for UMassF was used with the query generated from the responses to CF1 as well the related text. Ranked lists were generated for document and passage topics in the same manner as for UMassF. The results were then normalized and reranked using the genre and geography metadata as per section 3.5.5. We performed temporal and document feedback to provide a final ranking.

UMassCFC The linear bootstrapped model used for UMassF was used with the query generated from the responses to the clarification form, CF1. Ranked lists were generated for document and passage topics in the same manner as for UMassF. We performed temporal and document feedback to provide a final ranking.

UMassCMC The initial retrieval was performed using a query model built from CF1. These results were then reranked using topic metadata values. We utilized the geography, and genre metadata to rerank the results from the clarification form. We performed temporal and document feedback to provide a final ranking.

UMassCVC UMassCVC used variable-length passage techniques described in section 3.4, starting from the baseline document run UMassBaseSVM and the top 50 terms from the query model generated from the response to clarification form CF1. After variable-length passage retrieval, we post-processed the results as described in 3.2.3. For the 25 topics where the retrieval element was documents, the results we submitted were identical to the results from our baseline run UMassBaseSVM.

UMassF For the 25 document topics, query models were generated using the top 10 results of a preliminary ranked list as described in section 3.1.3. This preliminary list was obtained by retrieving 100 word passages using query likelihood. The title and description was used as the query for each topic. A linear bootstrapped model was used for retrieval. The top 1000 documents were returned for each of the 25 document topics. The same process as above was repeated for the 25 passage topics, except that a passage index was used for retrieval. The top 1000 non-overlapping passages were returned for each of these topics.

UMassRGG This run utilized the related text, geography, and genre metadata. Documents were returned for all topics. The metadata was utilized as described in sections 3.5.6 and 3.5.5. Retrieval was via query likelihood with Dirichlet smoothing. The smoothing parameter was set to 1000.

UMassVPMM UMassVPMM was a baseline passage run of sorts; it does not use any metadata or clarification form feedback for retrieval. It used variable length passage retrieval as described in section 3.4. We used the interpolated relevance model query model described in 3.1.4. We used the baseline run UMassBaseSVM as the starting ranked document list. Because we found in training that boosting the scores of passages from the top 25 documents improved results, we added a constant to the score of each of these passages, large enough to ensure that they would be ranked above all other passages. For the 25 topics where the retrieval element was documents, the results we submitted were identical to the results from our baseline run UMassBaseSVM.

UMassC2 This run used the passage feedback and reformulation for building a query model. A hybrid SVM was used for an initial retrieval. This ranked list was reranked using temporal and document feedback.

UMassC3 This run used the named entity feedback and reformulation for building a query model. A hybrid SVM was used for an initial retrieval. This ranked list was reranked using temporal feedback.

3.7 Results and Discussion

In order to allow an initial analysis of our various techniques, we generated several new runs based on different combinations of feedback, metadata handling, and retrieval granularity. These runs were evaluated using relevance judgments for the HARD 2004 topics. Results are presented in Tables 9 and 10.

3.7.1 CLARIFICATION FORMS

Our initial experiments allow us to investigate broad issues in ranking alternatives and named entity performance.

Passage Ranking Comparing the baseline, CF1, and CF2 rows in Tables 9, and 10, we observe that, in general passage feedback tends to improve performance. This result is not surprising given previous work in relevance feedback. What is a little more surprising is that clustering the results did not provide any advantage over the standard ranking. In fact, clustering often resulted in worse performance. One explanation for this behavior is the strictly positive nature of our feedback. Query models were built from positive documents. Negative information was essentially discarded. Therefore, to maximize the amount of information it receives, a system should get feedback from the documents which it is *most confident* about. By definition, these documents (or passages) will be the ones at the top of the ranked list. This intuition is confirmed by the number of passages marked relevant in the CF1 and CF2 clarification forms. On average, CF1 garnered more positive responses from users.

This result motivates two questions. First, how do we incorporate negative feedback into our existing framework? Research in retrieval by language models has ignored the question of negative feedback. If interaction and relevance feedback is to be considered an important aspect of HARD, it seems necessary to develop models for negative feedback. Second, how do we improve clustering so that removing redundancy does not result in detrimental loss of information in feedback? This question assumes both that the feedback in the likelihood ranking approach is redundant and that the feedback in the clustered approach is inferior. These assumptions need to be confirmed. Moreover, a similar question presents itself in novelty and subtopic retrieval and models from work in that field could improve future passage-based feedback forms.

Named Entities The results for runs using named entity information seem to confirm the difficulty of handling term-based feedback. The impact of named entity expansion is inconclusive. Training experiments demonstrate that, given the proper weighting of named entities, retrieval can be improved to the level of document feedback. That is, if we can detect that a person name is more important than a geographic name *for a particular query*, then we can match document feedback performance. However, the models we constructed used a uniform weight for all queries; person names always weighed the same as geographic and organizational names. Future experiments will attempt predict the relative import of entity types based on the query and corpus statistics.

	QL_{doc}	RM_{doc}	SVM_{doc}^{trec12}	$VPMM$	$SVMH_{psq}^{trec12}$	SVM_{psq}^{boot}
baseline	0.222	0.218	0.223	0.211	0.174	0.185
GG	0.226	0.223	0.233	0.171	0.186	0.201
RT	0.272	0.262	0.214	0.196	0.231	0.214
GG+RT	0.276	0.251	0.232	0.191	0.238	0.226
CF1	0.335	0.331	0.307	0.308	0.298	0.294
CF2	0.263	0.262	0.252	0.306	0.253	0.275
CF3	0.228	0.230	0.246	0.191	0.192	0.207
GG+CF1	0.312	0.306	0.295	0.304	0.298	0.292

Table 9: Binary Preference at 12,000 characters for passage and document runs. QL refers to query-likelihood retrieval, RM to relevance model retrieval, SVM to retrieval using a support vector machine *trained* using normal features, and $SVMH$ to retrieval using a support vector machine *trained* using hybrid features. Both SVM and $SVMH$ used hybrid feature vectors for *retrieval*. Subscripts indicate whether documents or passages were presented in the ranked list. Superscripts indicate the data which the SVM was built from. GG runs used genre and geography metadata, RT used related text, and CF* used clarification form query models and re-ranking.

	QL_{doc}		RM_{doc}		SVM_{doc}^{trec12}		$SVMH_{doc}^{trec12}$		SVM_{doc}^{boot}	
	hard	soft	hard	soft	hard	soft	hard	soft	hard	soft
baseline	0.327	0.320	0.343	0.339	0.322	0.314	0.286	0.300	0.287	0.318
GG	0.316	0.320	0.319	0.330	0.315	0.325	0.314	0.303	0.292	0.312
RT	0.359	0.373	0.348	0.355	0.365	0.361	0.332	0.347	0.342	0.354
GG+RT	0.363	0.369	0.349	0.353	0.350	0.354	0.358	0.351	0.321	0.348
CF1	0.341	0.361	0.339	0.362	0.357	0.374	0.335	0.358	0.347	0.363
CF2	0.316	0.330	0.315	0.329	0.323	0.345	0.320	0.338	0.336	0.372
CF3	0.333	0.298	0.333	0.303	0.332	0.329	0.319	0.323	0.297	0.326
GG+CF1	0.334	0.349	0.338	0.352	0.344	0.364	0.327	0.349	0.322	0.355

Table 10: Document R-Precision for hard and soft relevance. Labels are described in the caption to Table 9.

3.7.2 METADATA

The results of using the related-text (RT), reranking results by genre and geography (GG), and the combination of RT and GG can be seen in Tables 9 and 10.

For the submitted runs, our implementation of genre and geography reranking techniques was incorrect. Following the TREC conference, we fixed the mistake. The *notebook* version of this paper reports incorrect results for runs utilizing the genre and geography metadata.

For document retrieval, our use of related-text resulted in results as good as the use of the clarification form. The related-text significantly improves the results of retrieval methods that do not perform query expansion. When compared to the relevance models retrieval (RM_{doc}), which effectively performs query expansion, the related-text is on par or only slightly better. For passage retrieval, clarification forms performed better than related-text. Related-text may not provide feedback as precise as that collected with clarification forms.

There is no evidence we were able to leverage genre and geography. The results for genre and geography reranking differ little from our previously reported incorrect results. As such, we believe our technique for genre and geography reranking was akin to merely adding noise to the ranks of documents. We have since developed a new technique for metadata reranking that shows promise.

Topics may in fact disambiguate themselves with respect to metadata such that the majority of on-topic documents already satisfy the metadata. We expected topics to be ambiguous with respect to their metadata, but many were not. Eleven of the 45 topics were completely unambiguous, i.e. all on-topic documents satisfied the metadata. Looking at the fraction of on-topic documents that were relevant, across topics the median fraction was 0.83. The training topics were similarly unambiguous with respect to metadata. The more a topic is unambiguous with respect to the metadata, the less power metadata has for improving retrieval quality.

Another factor that may limit the power of genre and geography metadata is that searchers may be unable to express their metadata needs correctly. On an initial exploratory analysis of the retrieval results, we discovered many documents judged relevant that clearly fall outside the requested metadata. Searchers know a relevant document when they see one, but a priori they don't fully know what metadata is required of a relevant document. Successful techniques for using metadata will need to take this user error into consideration.

3.7.3 PASSAGE RETRIEVAL

Table 9 reveals two major findings in passage retrieval. First, document runs (shown in the first three columns) generally tend to do better than passage runs (columns 4-6) at passage retrieval, when a high-precision character-

level metric such as binary preference at 12,000 characters is used for evaluation. Second, CF1 seems to provide big improvements over the baseline for every retrieval method.

As for the question of whether variable-length passages improve high-accuracy passage retrieval, the results in table 9 are somewhat misleading. Although VPMM did better than both the bootstrap SVM and the hybrid SVM as a baseline, experiments performed after the TREC submission deadline showed that the gain there comes from the difference in retrieval method, not in passage length. Preliminary experiments using the mixture model of VPMM on fixed-length passages provide a better baseline than any of the document or passage runs presented here.

The bootstrap SVM method provides a small gain over the hybrid SVM method for all combinations of clarification forms and metadata, except for those involving related text. Interestingly, it seems that within the group of passage runs, the lower the baseline score, the bigger the boost from related text. In fact, VPMM is even hurt by the use of related text.

The major question raised by our findings for passage retrieval is whether passage retrieval is worthwhile, given that document retrieval almost always does better than passage retrieval for this evaluation metric. Or are we simply using the wrong evaluation metric for what we are really trying to measure? The official TREC 2003 HARD track metric of passage R-precision got at the notion that systems should be rewarded for returning text from many different documents. The character level measures correct a flaw in passage-level R-precision that favored very short passages, but remove this notion that there is some inherent good in returning text from a variety of documents. The problem of how to evaluate passage retrieval has clearly not been solved yet.

4 Acknowledgments

This work was supported in part by the Center for Intelligent Information Retrieval and in part by SPAWARSYSCEN-SD grant number N66001-02-1-8903. Any opinions, findings and conclusions or recommendations expressed in this material are the author(s)' and do not necessarily reflect those of the sponsor.

References

- [1] J. Allan, J. Callan, K. Collins-Thompson, B. Croft, F. Feng, D. Fisher, J. Lafferty, L. Larkey, T. N. Truong, P. Ogilvie, L. Si, T. Strohman, H. Turtle, and C. Zhai. The lemur toolkit for language modeling and information retrieval. <http://www.cs.cmu.edu/~lemur/>, 2003.
- [2] D. M. Bikel, R. L. Schwartz, and R. M. Weischedel. An algorithm that learns what's in a name. *Machine Learning*, 34(1-3):211–231, 1999.

- [3] J. P. Callan. Passage-level evidence in document retrieval. In *SIGIR 1994*, pp. 302–310.
- [4] W. B. Croft and J. Lafferty. *Language Modeling for Information Retrieval*. Kluwer, 2003.
- [5] O. de Kretser and A. Moffat. Effective document presentation with a locality-based similarity heuristic. In *SIGIR 1999*, pp. 113–120.
- [6] N. Dewdney, C. VanEss-Dykema, and R. MacMillan. The form is the substance: Classification of genres in text. In *Workshop on Human Language Technology and Knowledge Management, ACL 2001*.
- [7] F. Diaz and R. Jones. Using temporal profiles of queries for precision prediction. In *SIGIR 2004*, pp. 18–24.
- [8] S. Dumais, J. Platt, D. Heckerman, and M. Sahami. Inductive learning algorithms and representations for text categorization. In *CIKM 98*, pp. 148–155.
- [9] A. Finn and N. Kushmerick. Learning to classify documents according to genre. In *IJCAI-03 Workshop on Computational Approaches to Style Analysis and Synthesis*.
- [10] T. Joachims. Making large-scale SVM learning practical. In B. Schölkopf, C. Burges, and A. Smola, editors, *Advances in Kernel Methods - Support Vector Learning*. MIT Press, 1999.
- [11] T. Joachims. A statistical learning model of text classification for support vector machines. In *SIGIR 2001*, pp. 128–136.
- [12] J. Karlgren and D. Cutting. Recognizing text genres with simple metrics using discriminant analysis. In *15th conference on Computational linguistics*, pp. 1071–1075. ACL, 1994.
- [13] M. Kaszkiel and J. Zobel. Passage retrieval revisited. In *SIGIR 1997*, pp. 178–185.
- [14] B. Kessler and G. N. H. Schütze. Automatic detection of text genre. In *ACL/EACL 1997*, pp. 32–38.
- [15] R. Krovetz. Viewing morphology as an inference process. In *SIGIR 1993*, pp. 191–202.
- [16] G. Kumaran and J. Allan. Text classification and named entities for new event detection. In *SIGIR 2004*, pp. 297–304.
- [17] J. Lafferty and C. Zhai. Document language models, query models, and risk minimization for information retrieval. In *SIGIR 2001*, pp. 111–119.
- [18] L. S. Larkey, P. Ogilvie, M. A. Price, and B. Tamilio. Acrophile: an automated acronym extractor and server. In *5th Conf. Digital libraries*, pp. 205–214. ACM, 2000.
- [19] V. Lavrenko. *A Generative Theory of Relevance*. PhD thesis, University of Massachusetts, Amherst, 2004.
- [20] X. Li and W. B. Croft. Time-based language models. In *CIKM 2003*, pp. 469–475.
- [21] X. Li and W. B. Croft. An answer-updating approach to novelty detection. TR IR-359, CIIR, UMass, 2004.
- [22] C. D. Manning and H. Schütze. *Foundations of Statistical Natural Language Processing*. MIT Press, 1999.
- [23] A. K. McCallum. Bow: A toolkit for statistical language modeling, text retrieval, classification and clustering. <http://www.cs.cmu.edu/~mccallum/bow>, 1996.
- [24] D. Metzler, T. Strohman, H. Turtle, W. B. Croft. Indri at TREC 2004: Terabyte Track. In *TREC 13*, 2004.
- [25] E. Mittendorf and P. Schäuble. Document and passage retrieval based on hidden Markov models. In *SIGIR 1994*, pp. 318–327.
- [26] R. Nallapati. Discriminative models for information retrieval. In *SIGIR 2004*, pp. 64–71.
- [27] P. Ogilvie and J. Callan. Experiments using the lemur toolkit. In *TREC-10*, pp. 103–108, 2001.
- [28] J. M. Ponte and W. B. Croft. A language modeling approach to information retrieval. In *SIGIR 1998*, pp. 275–281.
- [29] I. Ruthven. Re-examining the potential effectiveness of interactive query expansion. In *SIGIR 2003*, pp. 213–220.
- [30] E. Stamatos, N. Fakotakis, and G. Kokkinakis. Text genre detection using common word frequencies. In *18th Int. Conf. on Comp, Linguistics*, 2000.
- [31] R. Yan, Y. Liu, R. Jin, and A. Hauptmann. On predicting rare classes with SVM ensembles in scene classification. In *ICASSP*, 2003.
- [32] Y. Yang and X. Liu. A re-examination of text categorization methods. In *SIGIR 1999*, pp. 42–49.
- [33] C. Zhai and J. Lafferty. A study of smoothing methods for language models applied to information retrieval. *ACM Trans. Inf. Syst.*, 22(2):179–214, 2004.