

Fast Set Intersection and Two Patterns Matching

Hagai Cohen and Ely Porat

Department of Computer Science, Bar-Ilan University, 52900 Ramat-Gan, Israel
{cohenh5,porately}@cs.biu.ac.il

Abstract. In this paper we present a new problem, the *fast set intersection* problem, which is to preprocess a collection of sets in order to efficiently report the intersection of any two sets in the collection. In addition we suggest new solutions for the *two-dimensional substring indexing* problem and the *document listing* problem for two patterns by reduction to the *fast set intersection* problem.

1 Introduction and Related Work

The intersection of large sets is a common problem in the context of retrieval algorithms, search engines, evaluation of relational queries and more. Relational databases use indices to decrease query time, but when a query involves two different indices, each one returning a different set of results, we have to intersect these two sets to get the final answer. The running time of this task depends on the size of each set, which can be large and make the query evaluation take longer even if the number of results is small. In information retrieval there is a great use of inverted index as a major indexing structure for mapping a word to the set of documents that contain that word. Given a word, it is easy to get from the inverted index the set of all the documents that contain that word. Nevertheless, if we would like to search for two words to get all documents that contain both, the inverted index doesn't help us that much. We have to calculate the occurrences set for each word and intersect these two sets. The problem of intersecting sets finds its motivation also in web search engines where the dataset is very large.

Various algorithms to improve the problem of intersecting sets have been introduced in the literature. Demaine et al. [1] proposed a method for computing the intersection of k sorted sets using an adaptive algorithm. Barbay et al. [2] showed that using interpolation search improves the performance of adaptive intersection algorithms. Baeza et al. [3] proposed an algorithm to improve the multiple searching problem which is related directly to computing the intersection of two sets. They introduced an intersection algorithm for two sorted sequences that is fast on average.

In this paper we present a new problem, the *fast set intersection* problem. This problem is to preprocess a databases of size N consisting of a collection of m sets to answer queries in which we are given two set indices $i, j \leq m$, and wish

to find their intersection. This problem has lots of applications where there is a need to intersect two sets in a lot of different fields like Information Retrieval, Web Searching, Document Indexing, Databases etc. An optimal solution for this problem will bring better solutions to various applications. As far as we know there is no previous work on this concrete problem.

We solve this problem using minimal space and still decrease the query time by using a preprocessing part. Our solution is the first non-trivial algorithm for this problem. We give a solution that requires linear space with worst case query time bounded by $O(\sqrt{N \text{output}}) + \text{output}$ where *output* is the intersection size.

In addition, we present a solution for the *two-dimensional substring indexing* problem, introduced by Muthukrishnan et al. [4]. In this problem we preprocess a database D of size N . So when given a string pair (σ_1, σ_2) , we wish to return all the database string pairs $\alpha_i \in D$ such that σ_1 is a substring of $\alpha_{i,1}$ and σ_2 is a substring of $\alpha_{i,2}$. Muthukrishnan et al. suggested a tunable solution for this problem which uses $O(N^{2-y})$ space for a positive fraction y and query time of $O(N^y + \text{output})$ where *output* is the number of such string pairs. We present a solution for this problem, based on solving the *fast set intersection* problem, that uses $O(N \log N)$ space with $O((\sqrt{N \log N \text{output}}) + \text{output}) \log^2 N$ query time.

In the *document listing* problem which was presented by Muthukrishnan et al. [5], we are given a collection of size N of text documents which may be preprocessed so when given a pattern p we want to return the set of all the documents that contain that pattern. Muthukrishnan et al. suggested an optimal solution for this problem which requires $O(N)$ space with $O(|p| + \text{output})$ query time where *output* is the number of documents that contain the pattern. However, there is no optimal solution when given a query consists of two patterns p, q to return the set of all the documents that contain them both. The only known solution for this problem is of Muthukrishnan et al. [5] which suggested a solution that uses $O(N\sqrt{N})$ space which supports queries in time $O(|p| + |q| + \sqrt{N} + \text{output})$. We present a solution for the *document listing* problem when the query consists of two patterns. Our solution uses $O(N \log N)$ space with $O(|p| + |q| + (\sqrt{N \log N \text{output}}) + \text{output}) \log^2 N$ query time.

The paper is structured as follows: In Sect. 2 we describe the *fast set intersection* problem. In Sect. 3 we describe our solution for this problem. In Sect. 4 we present similar problems with their solutions. In Sect. 5 we present our solution for the *two dimensional substring indexing* problem and the *document listing* problem for two patterns. In Sect. 6 we present some concluding remarks.

2 Fast Set Intersection Problem

We formally define the fast set intersection (FSI) problem.

Definition 1. Let D be a database of size N consisting of a collection of m sets. Each set has elements drawn from $1 \dots c$. We want to preprocess D so that given a query of two indices $i, j \leq m$, we will be able to calculate the intersection between sets i, j efficiently.

A naive solution for this problem is to save the sets sorted. Given a query of two sets i, j , go over the smaller set and check for each element if it exists in the second set. This costs $O(\min(|i|, |j|) \log \max(|i|, |j|))$. This solution can be further improved using hash tables. A static hash table [6] can store n elements using $O(n)$ space and build time, with $O(1)$ query time. For each set we can build a hash table to check in $O(1)$ time if an element is in the set or not. This way the query time is reduced to $O(\min(|i|, |j|))$ using linear space. The disadvantage of using this solution is that on the worst case we go over a lot of elements even if the intersection is small. A better query time can be gained by using more space for saving the intersection between every two sets. Using $O(m^2c)$ space we get an optimal query time of $O(\text{output})$ where output is the size of the intersection. Nevertheless, this solution uses extremely more space. In the next section we present our solution for the *fast set intersection* problem which bounds the query time on the worst case.

3 Fast Set Intersection Solution

Here we present our algorithm for solving the FSI problem. To the output of the algorithm, i.e. the intersection of the two sets, we call the *result set*. By *output* we denote the size of the result set.

3.1 Preprocessing

For each set in D we save a hash table to know in $O(1)$ time if an element is in that set or not. In addition, we save the inverse structure, i.e. for each element we save a hash table to know in $O(1)$ time if it belongs to a given set or not.

Our main data structure consists of a binary tree. Starting from the root node at level 0, each node in that tree handles number of subsets of the original sets from D . The cost of a node in that tree is the sum of the sizes of all the sets it handles. The root node handles all the m sets in D , therefore, it costs N .

Definition 2. Let d be a node which costs n . A large set in d is a set which has more than $\sqrt{n}$ elements.

Lemma 1. By definition, a node d which costs n , can handle at most $\sqrt{n}$ large sets.

A *set intersection matrix* is a matrix that saves for each set if it has an intersection with any other set. For m sets this matrix costs $O(m^2)$ bits space with $O(1)$ query time for answering if set i and set j have a non-empty intersection.

For each node we construct a set intersection matrix for the large sets in that node. By lemma 1, saving the set intersection matrix only for the large sets in a node that costs n space will cost only another n space.

Only large sets in a node will be propagated down to its two sons, we call them the *propagated group*. We divide the propagated group between the two sons by the elements each son will handle. Let d be the father node which costs

n , therefore, the propagated group which consists of only the large sets in d costs at most n as well. We want each son to cost at most $\frac{n}{2}$. Nevertheless, it is a hard problem to partition the propagated group into two groups of the same size, if even possible at all. Therefore, we shall do it in the following way: In the left son we add elements until adding another element will make it cost over $\frac{n}{2}$. This one element will not be handled in any son. It will only be remarked in d to check, during query time, if this element belongs to both sets. The other elements will be added to the right son. This way each son costs at most $\frac{n}{2}$.

A leaf in this binary tree is a node which is in constant size. Because each node in the tree costs half the space of its father then this tree has $\log N$ levels.

Theorem 1. *The space needed for this data structure is $O(N)$ space.*

Proof. The hash tables for all the sets cost $O(N)$ space. As well the inverse hash tables for all the elements cost $O(N)$ space.

The binary tree structure space cost is as follows: The root costs $O(N)$ bits for saving the set intersection matrix. In each level we save only another $O(N)$ bits because every two sons don't cost more than their father. Hence, the total cost of this tree structure is $O(N \log N)$ bits which is $O(N)$ space in term of words. $\square$

3.2 Query Answering

Given sets i, j (without loss of generality we assume $|i| \leq |j|$), we start traversing the tree from the root node. If i is not a large set in the root we check each element from it in the hash table of j . As there can be at most $\sqrt{N}$ elements in i because it is not a large set, this will cost $O(\sqrt{N})$. If both i, j are large sets we do as follows: We check in the set intersection matrix of the root whether there is a non-empty intersection between i and j . If there isn't we add nothing to the result set. If there is an intersection we check the element which is remarked in that node and check in its hash table if it belongs to i and j and add that element to the intersection if it belongs to both. Next we go down to the sons of the root and continue the traversing recursively.

Elements are added to the result set when we get to a node which in that node i is not a large set. In this case, we stop traversing down the tree from that node. Instead we step over all the elements of i in that node checking for each one of them if it belongs to j . We call such a node a *stopper node*.

Theorem 2. *The query time is bounded by $O(\sqrt{N \text{output}} + \text{output})$.*

Proof. The query computation consists of two parts. The tree traversal part and the time we spend on stopper nodes.

There are *output* elements in the result set, therefore, there can be at most $O(\text{output})$ stopper nodes. Because the tree height is $\log N$, for each stopper node we visit at most $\log N$ nodes for the tree traversal until we get to it. Therefore, the tree traversal part adds at most $O(\text{output} \log N)$ to the query time. But this is more than what we actually pay for the tree traversal because some stopper

nodes share their path from the root. This can be bounded better. Because the tree is a binary tree if we fully traverse the tree till $\log \text{output}$ height it will cost $O(\text{output})$ time. Now, from this height if we continue traverse the tree we visit for each stopper node at most $\log N - \log \text{output}$ nodes because we are already at $\log \text{output}$ height. Thus, the tree traversal part is bounded by $O(\text{output} + \text{output}(\log N - \log \text{output}))$. By log rules this equals to $O(\text{output} + \text{output} \log \frac{N}{\text{output}})$.

Now, we calculate how much time we spent on all the stopper nodes. A stopper node is a node which during the tree traversal we have to go over all elements of a non-large set in that node. The size of a non-large set in a stopper at level l is $\sqrt{\frac{N}{2^l}}$. Consider there are x stopper nodes. We denote by l_i the level for stopper node i . For all stopper nodes we pay at most:

$$\sum_{i=1}^x \sqrt{\frac{N}{2^{l_i}}} = \sqrt{N} \sum_{i=1}^x 2^{-\frac{1}{2}l_i} = \sqrt{N} \sum_{i=1}^x 1 \cdot 2^{-\frac{1}{2}l_i}$$

The Cauchy Schwarz inequality is that $(\sum_{i=1}^n x_i y_i)^2 \leq (\sum_{i=1}^n x_i^2)(\sum_{i=1}^n y_i^2)$. We use it in our case to get:

$$\begin{aligned} &\leq \sqrt{N} \sqrt{\sum_{i=1}^x 1^2} \sqrt{\sum_{i=1}^x (2^{-\frac{1}{2}l_i})^2} \\ &= \sqrt{N} \sqrt{x} \sqrt{\sum_{i=1}^x 2^{-l_i}} \end{aligned}$$

Kraft inequality from Information Theory states that for any binary tree:

$$\sum_{l \in \text{leaves}} 2^{-\text{depth}(l)} \leq 1$$

Because we never visit a subtree rooted by a stopper node, then in our case each stopper node can be viewed as a leaf in the binary tree. Therefore, we can transform Kraft inequality for all the stopper nodes instead of all tree leaves to get that $\sum_{i=1}^x 2^{-l_i} \leq 1$. Using this inequality gives us that:

$$\leq \sqrt{N} \sqrt{x} = \sqrt{Nx} \leq \sqrt{N \text{output}} = \text{output} \sqrt{\frac{N}{\text{output}}}$$

Thus, we pay $O(\text{output} \sqrt{\frac{N}{\text{output}}})$, for the time we spend in the stopper nodes.

Therefore, the tree traversal part and the time we spend on all stopper nodes is $O(\text{output} + \text{output} \log \frac{N}{\text{output}} + \text{output} \sqrt{\frac{N}{\text{output}}})$. Hence, the final query time is bounded by $O(\sqrt{N \text{output}} + \text{output})$. $\square$

Corollary 1. *The fast set intersection problem can be solved in linear space with worst case query time of $O(\sqrt{N \text{output}} + \text{output})$.*

4 Intersection Empty Query and Intersection Size Query

In the FSI problem given a query we want to return the result set, i.e. the intersection between two sets. What if we only want to know if there is any intersection between two sets? We call that the *intersection empty query* problem. Moreover, sometimes we would like only to know the size of the intersection without calculating the actual result set. We define these problems as follows:

Definition 3. Let D be a database of size N consisting of a collection of m sets. Each set has elements drawn from $1 \dots c$. The intersection empty query problem is to preprocess D so that given a query of two indices $i, j \leq m$, we want to calculate if sets i, j have any intersection. In the intersection size query problem when given a query we want to calculate the size of the result set.

A naive solution for the intersection empty query problem is to build a matrix saving if there is any intersection between every two sets. This solution uses $O(m^2)$ bits space with query time of $O(1)$. For the intersection size query problem we save the intersection size for every two sets by using slightly more space, $O(m^2)$ space, with query time of $O(1)$.

We can use part of our FSI solution method to solve the intersection empty query problem using $O(N)$ space with $O(\sqrt{N})$ query time. Instead of the whole tree structure we save only the root node with its set intersection matrix using $O(N)$ space. Given sets i, j (without loss of generality let's assume $|i| \leq |j|$), if i is not large set in the root we check each element from it in the hash table of j . Because i is not large set, this will cost at most $O(\sqrt{N})$ time. If i is a large set then we check in the set intersection matrix of the root to see if there is any intersection in $O(1)$ time. Hence, we can solve the intersection empty query problem in $O(\sqrt{N})$ time using $O(N)$ space.

With the same method we can solve the intersection size query problem by saving the size of the intersection instead of saving if there is any intersection in the set intersection matrix. This way we can solve the intersection size query problem in $O(\sqrt{N})$ time using $O(N)$ space.

5 Two Dimensional Substring Indexing Solution

In this section, we show how to solve the *two-dimensional substring indexing* problem and the *document listing* problem for two patterns using our FSI solution. The *two-dimensional substring indexing* problem was showed by Muthukrishnan et al. [4]. It is defined as follows:

Definition 4. Let D be a database consisting of a collection of string pairs $\alpha_i = (\alpha_{i,1}, \alpha_{i,2}), 1 \leq i \leq c$, which may be preprocessed. Given a query string pair (σ_1, σ_2) , the 2-d substring indexing problem is to identify all string pairs $\alpha_i \in D$, such that σ_1 is a substring of $\alpha_{i,1}$ and σ_2 is a substring of $\alpha_{i,2}$.

Muthukrishnan et al. [4] reduced the *two-dimensional substring indexing* problem to the *common colors query* problem which is defined as follows:

Definition 5. We are given an array $A[1 \dots N]$ of colors drawn from $1 \dots C$. We want to preprocess this array so that the following query can be answered efficiently: Given two non-overlapping intervals I_1, I_2 in $[1, N]$, list the distinct colors that occur in both intervals I_1 and I_2 .

The common colors query (CCQ) problem is another intersection problem where we have to intersect two intervals on the same array. We now show how to solve the CCQ problem by solving the FSI problem. By that we solve the two-dimensional substring indexing problem as well.

Given array A of size N , we build a data structure consisting of $\log N$ levels over this array. In the top level we partition A into two sets of size at most $\frac{N}{2}$, the first set containing colors, i.e. elements, of A in range $A[1 \dots \frac{N}{2}]$ and the second set containing colors in range $A[\frac{N}{2} + 1 \dots N]$. As well, each level i is partitioned into 2^i sets, each respectively, containing a successive set of $\frac{N}{2^i}$ colors from A . The bottom level, in similar fashion, is therefore partitioned into N sets each containing one different color from array A . The size of all the sets in each level is $O(N)$. Therefore, the size needed for all the sets in all levels is $O(N \log N)$.

Lemma 2. An interval I on A can be covered by at most $2 \log N$ sets.

Proof. Assume, by contradiction, that there exists an interval for which at least $m > 2 \log N$ sets are needed. This implies that there is some level that at least 3 (consecutive) sets are selected. However, for every 2 consecutive sets there have to be a set in the upper level that contains them both, so we can take it instead, and cover the same interval with only $m - 1$ sets, in contradiction to the assumption that at least m sets are required for the cover. $\square$

Theorem 3. The CCQ problem can be solved using $O(N \log N)$ space with $O((\sqrt{N \log N \text{output}} + \text{output}) \log^2 N)$ query time where output is the number of distinct colors that occur in both I_1 and I_2 .

Proof. Given two intervals I_1, I_2 we want to calculate their intersection, By lemma 2, I_1, I_2 are each covered by a group of $2 \log n$ sets at the most. To get the intersection of I_1, I_2 we will take each set from the first group and intersect it with each set from the second group using our FSI solution. Hence, we have to solve the FSI problem $O(\log^2 N)$ times. Our FSI solution takes $O(\sqrt{N \text{output}} + \text{output})$ time and $O(N)$ space for dataset which costs $O(N)$ space. Here the dataset costs $O(N \log N)$ space, therefore, we can solve the common colors query problem in $O((\sqrt{N \log N \text{output}} + \text{output}) \log^2 N)$ time using $O(N \log N)$ space. $\square$

As showed in [4] to solve the two dimensional substring problem we can solve a CCQ problem. As a result, the two dimensional substring problem can be solved in $O((\sqrt{N \log N \text{output}} + \text{output}) \log^2 N)$ time using $O(N \log N)$ space.

5.1 Document Listing Solution For Two Patterns

The document listing problem was presented by Muthukrishnan et al. [5]. In this problem we are given a collection D of text documents $d_1, \dots, d_c$, with

$\sum_i |d_i| = N$, which may be preprocessed, so when given a query comprising of a pattern p our goal is to return the set of all documents that contain one or more copies of p . Muthukrishnan et al. presented an optimal solution for this problem by building a suffix tree for D , searching the suffix tree for p and getting an interval I on an array with all the occurrences of p in D . Then they solve the colored range query problem on I to get each document only once. This solution requires $O(N)$ space with optimal query time of $O(|p| + output)$ where $output$ is the number of documents that contain p .

We are interested in solving this problem for a two patterns query. Given two patterns p, q , our goal is to return the set of all documents that contain both p and q . In [5] there is a solution that uses $O(N\sqrt{N})$ space with $O(|p| + |q| + \sqrt{N} + output)$ query time. Their solution is based on searching a suffix tree of all the documents for the two patterns p, q in $O(|p| + |q|)$ time. From this they get two intervals: I_1 with p occurrences and I_2 with q occurrences.. On this intervals they solve a CCQ problem to get the intersection between I_1 and I_2 for all the documents that contain both p and q .

We suggest a new solution based on solving the FSI problem. We use the same method as Muthukrishnan et al. [5] until we get the two intervals: I_1 with p occurrences and I_2 with q occurrences. Now, we have to solve a CCQ problem which can be solved as shown above in theorem 3. Therefore, the document listing problem for two patterns can be solved in $O(|p| + |q| + (\sqrt{N \log N output} + output) \log^2 N)$ time using $O(N \log N)$ space where $output$ is the number of documents that contain both p and q .

6 Conclusions

In this paper we developed a method to improve algorithms which intersects sets as a common task. We solved the fast set intersection problem using $O(N)$ space with query time bounded by $O(\sqrt{N output} + output)$. We showed how to improve some other problems, the two dimensional substring indexing problem and the document listing problem for two patterns, using the fast set intersection problem.

There is still a lot of research to be done in regards to the fast set intersection problem. It is open if the query time can be bounded better. Moreover, we showed only two applications for the fast set intersection problem. We are sure that the fast set intersection problem can be useful in other fields as well.

References

1. Demaine, E.D., Lo'pez-Ortiz, A., Munro, J.I.: Adaptive set intersections, unions, and differences. In: In Proceedings of the 11th Annual ACM-SIAM Symposium on Discrete Algorithms (SODA). (2000) 743–752
2. Barbay, J., Lo'pez-ortiz, R., Lu, T.: Faster adaptive set intersections for text searching. In: Experimental Algorithms: 5th International Workshop, WEA 2006, Cala Galdana, Menorca, Springer (2006) 146–157

3. Baeza-Yates, R.A.: A fast set intersection algorithm for sorted sequences. In Sahinalp, S.C., Muthukrishnan, S., Dogrusöz, U., eds.: CPM. Volume 3109 of Lecture Notes in Computer Science., Springer (2004) 400–408
4. Ferragina, P., Koudas, N., Srivastava, D., Muthukrishnan, S.: Two-dimensional substring indexing. In: PODS '01: Proceedings of the twentieth ACM SIGMOD-SIGACT-SIGART symposium on Principles of database systems, New York, NY, USA, ACM Press (2001) 282–288
5. Muthukrishnan, S.: Efficient algorithms for document retrieval problems. In: SODA '02: Proceedings of the thirteenth annual ACM-SIAM symposium on Discrete algorithms, Philadelphia, PA, USA, Society for Industrial and Applied Mathematics (2002) 657–666
6. Fredman, M.L., Komlós, J., Szemerédi, E.: Storing a sparse table with $O(1)$ worst case access time. J. ACM **31**(3) (1984) 538–544