

CSC2420 Fall 2012: Algorithm Design, Analysis and Theory

Allan Borodin

November 15, 2012; Lecture 10

Randomized online bipartite matching and the adwords problem.

- We briefly return to online algorithms and algorithms in the random order model (ROM). Here we have already seen evidence of the power of randomization in the context of the MaxSat problem.
- Another nice sequence of results begins with a randomized online algorithm for bipartite matching due to Karp, Vazirani and Vazirani [1990]. We quickly overview some results in this area as it represents a topic of current interest. (The recent FOCS 2012 conference had a session of three papers related to this topic.)
- In the online bipartite matching problem, we have a bipartite graph G with nodes $U \cup V$. Nodes in U enter online revealing all their edges. A deterministic greedy matching produces a maximal matching and hence a $\frac{1}{2}$ approximation.
- It is easy to see that any deterministic online algorithm cannot be better than a $\frac{1}{2}$ approximation even when the degree of every $u \in U$ is at most (equal) 2

The randomized ranking algorithm

- The **Ranking** algorithm chooses a random permutation of the nodes in V and then when a node $u \in U$ appears, it matches u to the highest ranked unmatched $v \in V$ such that (u, v) is an edge (if such a v exists).
- Aside: making a random choice for each u is still only a $\frac{1}{2}$ approx.
- Equivalently, this algorithm can be viewed as a deterministic greedy algorithm (i.e. always matching when possible and breaking ties consistently) in the ROM model.
- That is, let $\{v_1, \dots, v_n\}$ be any fixed ordering of the vertices and let the nodes in U enter randomly, then match each u to the first unmatched $v \in V$ according to the fixed order.
- To argue this, consider fixed orderings of U and V ; the claim is that the matching will be the same whether U or V is entering online.
- Note: This is a comment about this particular algorithm and not some general equivalence between randomized online algorithms and deterministic algorithms in the ROM model.

The KVV result and recent progress

KVV Theorem

Ranking provides a $(1 - 1/e)$ approximation.

- Original analysis is not rigorous.
- There is an alternative proof (and extension) by Goel and Mehta [2008], and then another proof in Birnbaum and Mathieu [2008].
- Recall that this positive result can be stated either as the bound for a particular deterministic algorithm in the stochastic ROM model, or as the randomized Ranking algorithm in the (adversarial) online model.
- KVV show that the $(1 - 1/e)$ bound is essentially tight for any randomized online (i.e. adversarial input) algorithm. In the ROM model, Goel and Mehta state inapproximation bounds of $\frac{3}{4}$ (for deterministic) and $\frac{5}{6}$ (for randomized) algorithms.

Some comments on the Birnbaum and Mathieu proof

- The worst case example for any online algorithm is a (n, n) graph with a perfect matching.
- In particular, for $n = 2$, the precise expected competitive (i.e. approximation) ratio is $\frac{3}{4}$. The inapproximation can be seen by using the Yao principle for obtaining bounds on randomized algorithms.

The main lemma in the analysis

Let x_t be the probability (over the random permutations of the vertices in V) that the vertex of rank t is matched. Then $1 - x_t \leq \frac{1}{n} \sum_{s=1}^t x_s$

- Letting $S_t = \sum_{s=1}^t x_s$ the lemma can be restated as $S_t(1 + 1/n) \geq 1 + S_{t-1}$ for all t . Given that the graph has a perfect matching, the expected competitive ratio is S_n/n . It is shown that $\frac{1}{n} S_n \geq 1 - (1 - \frac{1}{n+1})^n \rightarrow 1 - 1/e$.

Getting past the $(1 - 1/e)$ bound

- The ROM model can be considered as an example of what is called stochastic optimization in the OR literature. There are other stochastic optimization models that are perhaps more natural, namely i.i.d sampling from known and unknown distributions.
- Feldman et al [2009] study the known distribution case and show a randomized algorithm that first computes an optimal offline solution (in terms of expectation) and uses that to guide an online allocation.
- They achieve a .67 approximation (improved to .699 Bahmani and Kapralov [2010] and also show that no online algorithm can achieve better than $26/27 \approx .99$ (improved to .902).
- Karande, Mehta, Tripathi [2011] show that an approximation in the ROM model implies the same approximation in the unknown distribution model. They show that the KVV Ranking algorithm achieves approximation .653 in the ROM model and is no better than .727.

The adwords problem: an extension of bipartite matching

- In the (single slot) adwords problem, the nodes in U are queries and the nodes in V are advertisers. For each query q and advertiser i , there is a bid $b_{q,i}$ representing the value of this query to the advertiser.
- Each advertiser also usually has a hard budget B_i which cannot be exceeded. The goal is to match the nodes in U to V so as to maximize the sum of the accepted bids without exceeding any budgets. Without budgets and when each advertiser will pay for at most one query, the problem then is edge weighted bipartite matching.
- In the online case, when a query arrives, all the relevant bids are revealed.

Some results for the adwords problem

- Here we are just considering the combinatorial problem and ignoring game theoretic aspects of the problem.
- The problem has been studied for the special (but well motivated case) that all bids are small relative to the budgets. As such this problem is incomparable to the matching problem where all bids are in $\{0,1\}$ and all budgets are 1.
- For this small bids case, Mehta et al [2005] provide a deterministic online algorithm achieving the $1 - 1/e$ bound and show that this is optimal for all randomized online algorithms (i.e. adversarial input).

Greedy for a class of adwords problems

- Goel and Mehta [2008] define a class of adwords problems which include the case of small budgets, bipartite matching and b -matching (i.e. when all budgets are equal to some b and all bids are equal to 1).
- For this class of problems, they show that a deterministic greedy algorithm achieves the familiar $1 - 1/e$ bound in the ROM model. Namely, the algorithm assigns each query (i.e. node in U) to the advertiser who values it most (truncating bids to keep them within budget and consistently breaking ties). Recall that Ranking can be viewed as greedy (with consistent tie breaking) in the ROM model.

Vertex weighted bipartite matching

- Aggarwal et al [2011] consider a vertex weighted version of the online bipartite matching problem. Namely, the vertices $v \in V$ all have a known weight w_v and the goal is now to maximize the weighted sum of matched vertices in V when again vertices in U arrive online.
- This problem can be shown to subsume the adwords problem when all bids $b_{q,i} = b_i$ from an advertiser are the same.
- It is easy to see that Ranking can be arbitrarily bad when there are arbitrary differences in the weight. Greedy (taking the maximum weight match) can be good in such cases. Can two such algorithms be somehow combined? Aggarwal et al are able to achieve the same $1-1/e$ bound for this class of vertex weighted bipartite matching.

The vertex weighted online algorithm

The perturbed greedy algorithm

For each $v \in V$, pick r_v randomly in $[0, 1]$

Let $f(x) = 1 - e^{1-x}$

When $u \in U$ arrives, match u to the unmatched v (if any) having the highest value of $w_v * f(x_v)$. Break ties consistently.

In the unweighted case when all w_v are identical this is the Ranking algorithm.

Some open problems in the ROM model

There are many open problems in the ROM model. In general any online problem can be studied with respect to this model. For example, relevant to what we have just discussed:

- The adwords problem without any restriction.
- Beating $1 - 1/e$ for the vertex weighted or b -matching problems.
- Perhaps, the first prominent use of this model is for the secretary problem; namely selecting the maximum element (or best k elements) in a randomly ordered sequence. Here again $1 - 1/e$ is the best approximation.
- This has been generalized to the matroid secretary problem by Babaioff, Immorlica, R. Kleinberg [2007]. For arbitrary matroids, the approximation ratio remains an open problem.

Sublinear time and sublinear space algorithms

We continue to consider contexts in which randomization is provably necessary. In particular, we will study sublinear time algorithms and then the (small space) streaming model.

- An algorithm is sublinear time if its running time is $o(n)$, where n is the length of the input. As such an algorithm must provide an answer without reading the entire input.
- Thus to achieve non-trivial tasks, we almost always have to use randomness in sublinear time algorithms to **sample** parts of the inputs.
- The subject of sublinear time algorithms is a big topic and we will only present a very small selection of hopefully representative results.
- The general flavour of results will be a tradeoff between the accuracy of the solution and the time bound.
- This topic will take us beyond search and optimization problems.

A deterministic exception: estimating the diameter in a finite metric space

- We first consider an exception of a “sublinear time” algorithm that does not use randomization. (Comment: “sublinear in a weak sense”.)
- Suppose we are given a finite metric space M (with say n points x_i) where the input is given as n^2 distance values $d(x_i, x_j)$. The problem is to compute the diameter D of the metric space, that is, the maximum distance between any two points.
- For this maximum diameter problem, there is a simple $O(n)$ time (and hence sublinear) algorithm; namely, choose an arbitrary point $x \in M$ and compute $D = \max_j d(x, x_j)$. By the triangle inequality, D is a 2-approximation of the diameter.
- I say sublinear time in a weak sense because in an explicitly presented space (such as d dimensional Euclidean space), the points could be explicitly given as inputs and then the input size is n and not n^2 .

Sampling the inputs: some examples

- The goal in this area is to minimize execution time while still being able to produce a reasonable answer with sufficiently high probability.
- We will consider the following examples:
 - ① Finding an element in an (anchored) sorted list [Chazelle,Liu,Magen]
 - ② Estimating the average degree in a graph [Feige 2006]
 - ③ Estimating the size of some maximal (and maximum) matching [Nguyen and Onak 2008] in bounded degree graphs.
 - ④ Examples of property testing, a major topic within the area of sublinear time algorithms. See Dana Ron's DBLP for many results and surveys.

Finding an element in an (anchored) sorted list

- Suppose we have an array $A[i]$ for $1 \leq i \leq n$ where each $A[i]$ is a pair (x_i, p_i) with $x_1 = \min\{x_i\}$ and p_i being a pointer to the next smallest value in the linked list.
- That is, $x_{p_i} = \min\{x_j \mid x_j > x_i\}$. (For simplicity we are assuming all x_j are distinct.)
- We would like to determine if a given value x occurs in the linked list and if so, output the index j such that $x = x_j$.

A $\sqrt{n}$ algorithm for searching in anchored sorted linked list

Let $R = \{j_i\}$ be $\sqrt{n}$ randomly chosen indices plus the index 1.

Access these $\{A[j_i]\}$ to determine k such that x_k is the largest of the accessed array elements less than or equal to x .

Search forward $2\sqrt{n}$ steps in the linked list to see if and where x exists

Claim:

This is a one-sided error algorithm that (when $x \in \{A[i]\}$) will fail to return j such that $x = A[j]$ with probability at most $1/2$.

Estimating average degree in a graph

- Given a graph $G = (V, E)$ with $|V| = n$, we want to estimate the average degree d of the vertices.
- We want to construct an algorithm that approximates the average degree within a factor less than $(2 + \epsilon)$ with probability at least $3/4$ in time $O(\frac{\sqrt{n}}{\text{poly}(\epsilon)})$. We will assume that we can access the degree d_i of any vertex v_i in one step.
- Like a number of results in this area, the algorithm is simple but the analysis requires some care.

The Feige algorithm

Sample $8/\epsilon$ random subsets S_i of V each of size (say) $\frac{\sqrt{n}}{\epsilon^3}$

Compute the average degree a_i of nodes in each S_i .

The output is the minimum of these $\{a_i\}$.

The analysis of the approximation

Since we are sampling subsets to estimate the average degree, we might have estimates that are too low or too high. But we will show that with high probability these estimates will not be too bad. More precisely, we need:

- ❶ Lemma 1: $\text{Prob}[a_i < \frac{1}{2}(1 - \epsilon)\bar{d}] \leq \frac{\epsilon}{64}$
- ❷ Lemma 2: $\text{Prob}[a_i > (1 + \epsilon)\bar{d}] \leq 1 - \frac{\epsilon}{2}$

The probability bound in Lemma 2 is amplified as usual by repeated trials. For Lemma 1, we fall outside the desired bound if any of the repeated trials gives a very small estimate of the average degree but by the union bound this is no worse than the sum of the probabilities for each trial.

Understanding the input query model

- As we initially noted, sublinear time algorithms almost invariably sample (i.e. query) the input in some way. The nature of these queries will clearly influence what kinds of results can be obtained.
- Feige's algorithm for estimating the average degree uses only "degree queries"; that is, "what is the degree of a vertex v ".
- Feige shows that in this degree query model, that any algorithm that achieves a $(2 - \epsilon)$ approximation (for any $\epsilon > 0$) requires time $\Omega(n)$.
- In contrast, Goldreich and Ron [2008] consider the same average degree problem in the "neighbour query" model; that is, upon a query (v, j) , the query oracle returns the j^{th} neighbour of v or a special symbol indicating that v has degree less than j . A degree query can be simulated by $\log n$ neighbour queries.
- Goldreich and Ron show that in the neighbour query model, that the average degree $\bar{d}$ can be $(1 + \epsilon)$ approximated (with one sided error probability $2/3$) in time $O(\sqrt{(n/\bar{d})} \text{poly}(\log n, \frac{1}{\epsilon}))$
- They also show that this $\sqrt{n}$ time bound is essentially optimal.

Approximating the size of a maximum matching in a bounded degree graph

- We recall that the size of any *maximal* matching is within a factor of 2 of the size of a maximum matching.
- Our goal is to compute with high probability a *maximal* matching in time depending only on the maximum degree D .

Nguyen and Onak Algorithm

Choose a random permutation p of the edges $\{e_j\}$

% Note: this will be done “on the fly” as needed

The permutation determines a maximal matching M as given by the greedy algorithm that adds an edge whenever possible.

Choose $r = O(d/\epsilon^2)$ nodes $\{v_i\}$ at random

Using an “oracle” let X_i be the indicator random variable for whether or not vertex v_i is in the maximal matching.

Output $\tilde{m} = \sum_{i=1\dots r} X_i$

Performance and time for maximal matching

Claims

- ❶ $m \leq \tilde{m} \leq m + \epsilon n$ where $m = |M|$.
 - ❷ The algorithm runs in time $2^{O(D)}/\epsilon^2$
- This immediately gives an approximation of the *maximum* matching m^* such that $m^* \leq \tilde{m} \leq 2m^* + \epsilon n$
 - A more involved algorithm by Nguyen and Onak yields the following result:

Nguyen and Onak maximum matching result

Let $\delta, \epsilon > 0$ and let $k = \lceil 1/\delta \rceil$. There is a randomized one sided algorithm (with probability $2/3$) running in time $\frac{2^{O(D^k)}}{\epsilon^{2k+1}}$ that outputs a maximum matching estimate $\tilde{m}$ such that $m^* \leq \tilde{m} \leq (1 + \delta)m^* + \epsilon n$.

Property Testing

- Perhaps the most prevalent and useful aspect of sublinear time algorithms is for the concept of property testing. This is its own area of research with many results.
- Here is the concept: Given an object G (e.g. a function, a graph), test whether or not G has some property P (e.g. G is bipartite).
- The tester determines with sufficiently high probability (say $2/3$) if G has the property or is “ ϵ -far” from having the property. The tester can answer either way if G does not have the property but is “ ϵ -close” to having the property.
- We will usually have a 1-sided error in that we will always answer YES if G has the property.
- We will see what it means to be “ ϵ -far” (or close) from a property by some examples.

Tester for linearity of a function

- Let $f : Z_n \rightarrow Z_n$; f is linear if $\forall x, y \ f(x + y) = f(x) + f(y)$.
- Note: this will really be a test for group homomorphism
- f is said to be ϵ -close to linear if its values can be changed in at most a fraction ϵ of the function domain arguments (i.e. at most ϵn elements of Z_n) so as to make it a linear function. Otherwise f is said to be ϵ -far from linear.

The tester

Repeat $4/\epsilon$ times

Choose $x, y \in Z_n$ at random

If $f(x) + f(y) \neq f(x + y)$

then Output f is not linear

End Repeat If all these $4/\epsilon$ tests succeed then Output linear

- Clearly if f is linear, the tester says linear.
- If f is ϵ -far from being linear then the probability of detecting this is at least $2/3$.

Testing a list for monotonicity

- Given a list $A[i] = x_i, i = 1 \dots n$ of distinct elements, determine if A is a monotone list (i.e. $i < j \Rightarrow A[i] < A[j]$) or is ϵ -far from being monotone in the sense that more than $\epsilon * n$ list values need to be changed in order for A to be monotone.
- The algorithm randomly chooses $2/\epsilon$ random indices i and performs binary search on x_i to determine if x_i is in the list. The algorithm reports that the list is monotone if and only if all binary searches succeed.
- Clearly the time bound is $O(\log n/\epsilon)$ and clearly if A is monotone then the tester reports monotone.
- If A is ϵ -far from monotone, then the probability that a random binary search will succeed is at most $(1 - \epsilon)$ and hence the probability of the algorithm failing to detect non-monotonicity is at most $(1 - \epsilon)^{\frac{2}{\epsilon}} \leq \frac{1}{e^2}$

Graph Property testing

- Graph property testing is an area by itself. There are several models for testing graph properties.
- Let $G = (V, E)$ with $n = |V|$ and $m = |E|$.
- Dense model: Graphs represented by adjacency matrix. Say that graph is ϵ -far from having a property P if more than ϵn^2 matrix entries have to be changed so that graph has property P .
- Sparse model, bounded degree model: Graphs represented by vertex adjacency lists. Graph is ϵ -far from property P if at least ϵm edges have to be changed.
- In general there are substantially different results for these two graph models.

The property of being bipartite

- In the dense model, there is a constant time one-sided error tester. The tester is (once again) conceptually what one might expect but the analysis is not at all immediate.

Goldreich, Goldwasser, Ron bipartite tester

Pick a random subset S of vertices of size $r = \Theta\left(\frac{\log(\frac{1}{\epsilon})}{\epsilon^2}\right)$

Output bipartite iff the induced subgraph is bipartite

- Clearly if G is bipartite then the algorithm will always say that it is bipartite.
- The claim is that if G is ϵ -far from being bipartite then the algorithm will say that it is not bipartite with probability at least $2/3$.
- The algorithm runs in time quadratic in the size of the induced subgraph (i.e. the time needed to create the induced subgraph).

Testing bipartiteness in the bounded degree model

- Even for degree 3 graphs, $\Omega(\sqrt{n})$ queries are required to test for being bipartite or ϵ -far from being bipartite. Goldreich and Ron [1997]
- There is a nearly matching algorithm that uses $O(\sqrt{n} \text{poly}(\log n/\epsilon))$ queries. The algorithm is based on random walks in a graph and utilizes the fact that a graph is bipartite iff it has no odd length cycles.

Goldreich and Ron [1999] bounded degree algorithm

Repeat $O(1/\epsilon)$ times

Randomly select a vertex $s \in V$

If algorithm $OddCycle(s)$ returns cycle found then REJECT

End Repeat

If case the algorithm did not already reject, then ACCEPT

- $OddCycle$ performs $\text{poly}(\log n/\epsilon)$ random walks from s each of length $\text{poly}(\log n/\epsilon)$. If some vertex v is reached by both an even length and an odd length prefix of a walk then report cycle found; else report odd cycle not found