

LECTURE 4

①

Data stream $A = (a_1, a_2, \dots, a_m)$

$$a_i \in N = \{1, 2, \dots, n\}$$

$$\text{Let } M_i = \{j \mid a_j = i\}$$

$$m_i = |M_i| = \# \text{ copies of value } i \text{ in stream}$$

$$\vec{m} = (m_1, m_2, \dots, m_n) \quad \text{FREQUENCY VECTOR}$$

The k th FREQUENCY MOMENT F_k of A is

$$F_k = \begin{cases} |\{i \mid m_i \neq 0\}| & \text{if } k=0 \\ \max_i m_i & \text{if } k=1 \\ \|\vec{m}\|_k^k = \sum_{i=1}^n m_i^k & \text{otherwise} \end{cases}$$

In particular $F_0 = \# \text{ distinct elements}$

$F_1 = \text{length of sequence } (=m)$

$F_2 = \text{REPEAT RATE or}$

GINI'S INDEX OF HOMOGENEITY

Motivation for Hon-Matias-Szegedy AMS99

FM important for database applications

Claim: Essentially all query optimization methods in relational DB systems require F_0

$F_k, k \geq 2$ indicate degree of SKEW

Important in many parallel DB applications

Can determine selection of algorithm.

(see intro of AMS for more sales talk...)

(2)

Thus, beneficial to maintain estimates for F_k

- 1) Needs to be done online as elements inserted into database, i.e., small time
- 2) Needs to stay in primary memory to be fast
large memory requirements $\Rightarrow$ store data on disk $\Rightarrow$ slow access and update times.
i.e., small memory

STREAMING ALGORITHMS ^(simplified) (~~simplified~~ version)

- One pass over data
- Sublinear memory (measured in input size)

This area of research essentially started in AMS 99

Influential paper — Gödel Prize 2005

Today whole area of research.

This and next lecture

- Algorithms [streaming] for frequency moments (upper bounds)
- Impossibility proofs (lower bounds; uses CC)

But first quick recap of probability theory

(SELECTIVE) RECAP OF PROBABILITY THEORY

③

Focus on discrete random variables X

Can take values $x_i, i = 1, 2, \dots, n$

$$\Pr[X = x_i] = p_i \geq 0 \quad \left(\sum_{i=1}^n p_i = 1 \right)$$

"Events" are subsets S, T of possible outcomes

INCLUSION-EXCLUSION PRINCIPLE

$$\Pr[S \cup T] = \Pr[S] + \Pr[T] - \Pr[S \cap T]$$

UNION BOUND

For any events $S_1, S_2, \dots, S_m$, the probability that one or more of these occur is upper-bounded by

$$\Pr\left[\bigcup_{i=1}^m S_i\right] \leq \sum_{i=1}^m \Pr[S_i]$$

CONDITIONAL PROBABILITY

Probability of event S given that T is known to have occurred

$$\Pr[S | T] = \frac{\Pr[S \cap T]}{\Pr[T]}$$

INDEPENDENCE

Two events S and T are independent if

$$\Pr[S \cap T] = \Pr[S] \cdot \Pr[T]$$

E.g. for discrete random variables X, Y they are independent if

$$\Pr[X = x_i, Y = y_j] = \Pr[X = x_i] \cdot \Pr[Y = y_j]$$

(4)

EXPECTATION $E[X] = \sum_{i=1}^n p_i \cdot x_i$

"expected value"
or "mean" $= \sum_{i=1}^n \Pr[X=x_i] \cdot x_i$

Linearity of expectation

$$E[\alpha X + \beta Y] = \alpha E[X] + \beta E[Y]$$

VARIANCE $\text{Var}(X) = E[(X - E[X])^2]$
 $= E[X^2] - (E[X])^2$

$$\text{Var}(c \cdot X) = c^2 \text{Var}(X)$$

For independent random variables $X_1, \dots, X_S$

$$\text{Var}\left(\sum_{i=1}^S X_i\right) = \sum_{i=1}^S \text{Var}(X_i)$$

Let all X_i be identically distributed and independent (i.e., repetitions of same experiment).

Let $Y = \sum_{i=1}^S X_i / S$

Then $\text{Var}(Y) = \text{Var}(X_i) / S$

STANDARD DEVIATION $= \sqrt{\text{variance}}$

MARKOV'S INEQUALITY for $X \geq 0$

$$\Pr[X \geq a] \leq E[X] / a$$

Proof Define indicator random variable $I = \begin{cases} 1 & \text{if } X \geq a \\ 0 & \text{otherwise} \end{cases}$

Always $I \leq X/a$. Take expectation on both sides.

(5)

Useful reformulation of Markov's inequality

$$\Pr[X \geq c \cdot \mathbb{E}[X]] \leq 1/c$$

"The probability that X exceeds its expectation by factor c is at most $1/c$."

Not very strong bound, but holds for any X (≥ 0).

CHEBYSHEV'S INEQUALITY

"In any data sample (in particular random samples from probability distributions), nearly all values are close to the mean"

More precisely: "No more than $1/s^2$ ^{fraction} of distributions values can be more than s standard deviations away from mean."

Suppose X has mean μ and variance σ^2 ($\text{both} < \infty$). Then

$$\Pr[|X - \mu| \geq s\sigma] \leq 1/s^2$$

Proof Consider $Y = (X - \mathbb{E}[X])^2$ and apply Markov's inequality with $a = (s\sigma)^2$.

$$\Pr[(X - \mathbb{E}[X])^2 \geq s^2\sigma^2] \leq \frac{\mathbb{E}[(X - \mathbb{E}[X])^2]}{s^2\sigma^2} = 1/s^2.$$

"

$$\Pr[|X - \mathbb{E}[X]| \geq s\sigma]$$

(6)

CHERNOFF BOUND

In fact several flavours of bounds/inequalities.

"If r.v. is sampled many times, then the average value converges (quickly) to the expected value."

$x_1, \dots, x_n$ independent 0-1 r.v.

$$\Pr[x_i = 1] = p$$

$$\Pr[x_i = 0] = 1 - p$$

$\sum_{i=1}^n x_i / n$ should be close to $E[x_i] = p$

$$\Pr\left[\left|\frac{\sum_{i=1}^n x_i}{n} - p\right| \geq \delta\right]$$

$$\leq 2 \cdot \exp\left(-\frac{\delta^2 n}{2p(1-p)}\right)$$

$$= \exp(-\Omega(n))$$

exponential decrease in $n = \#$ repeated experiments.

* * * * *

Now back to frequency moments of data streams and how to estimate them...

7

Each moment F_k can be computed exactly and deterministically with $O(n \log m)$ bits.

Simply compute all numbers m_i

Use randomness to do better!

Example 1 (not very good) (and not very related)

Say we have a stream $A = (a_1, a_2, \dots, a_m)$
Want to compute average $\frac{1}{m} \sum_{i=1}^m a_i$

Pick uniformly and randomly $i \in [m] = \{1, \dots, m\}$
Answer a_i

Pros Correct in expectation
Uses very little memory

Cons: Can vary wildly (variance large)
How do we know we're even close
to being right

Want to define random variable X s.t.

- a) $E[X] = F_k$
- b) X can be computed in small space
- c) Variance of X small

This is our goal.

(8)

We will prove (or at least sketch proof) of following theorem

Thm 2 For every $k \geq 1$, every $\lambda > 0$, every $\epsilon > 0$, there is a randomized algorithm that given $A = (a_1, \dots, a_m)$, $a_i \in N = \{1, \dots, n\}$, makes one pass over A , computes Y s.t.

$$Pr [|Y - F_k| > \lambda F_k] \leq \epsilon$$

and uses

$$O \left(\frac{k \log(1/\epsilon)}{\lambda^2} n^{1-1/k} (\log n + \log m) \right)$$

bits of memory.

Asking for closer approximation or lower error probability will increase memory, but only by constant factors.

First try [assume for simplicity m known; can be fixed later]
let $p \in \{1, \dots, m\}$ uniformly random.

Suppose $a_p = \ell$

$$\text{let } r = | \{ q \mid q \geq p, a_q = \ell \} |$$

= number of occurrences of $a_p = \ell$ starting from position p . Define

$$X = m (r^k - (r-1)^k)$$

To compute X , need $\log n$ bits for a
 $+$ $\log m$ bits for r
 What is $E[X]$?

(Manually) sort A by occurrences/copies of distinct elements. Note that all are equally likely to be chosen, so divide by $1/m$ to get expectation.

$$\begin{array}{ccccccc}
 \frac{m(m_1^k - (m_1 - 1)^k)}{\text{1st copy of 1}} & + & \frac{m((m_1 - 1)^k - (m_1 - 2)^k)}{\text{2nd copy of 1}} & + & \dots & + & \frac{m(2^k - 1^k)}{\text{next to last copy of 1}} & + & \frac{m 1^k}{\text{last copy of 1}} \\
 \\
 \frac{m(m_2^k - (m_2 - 1)^k)}{\text{1st copy of 2}} & + & \frac{m((m_2 - 1)^k - (m_2 - 2)^k)}{\text{2nd copy of 2}} & + & \dots & + & \frac{m(2^k - 1^k)}{\text{next to last}} & + & \frac{m 1^k}{\text{last}} \\
 \\
 \vdots & & & & & & & & \\
 \\
 \frac{m(m_n^k - (m_n - 1)^k)}{\text{1st copy of } n} & + & \frac{m((m_n - 1)^k - (m_n - 2)^k)}{\text{2nd copy of } n} & + & \dots & + & \frac{m(2^k - 1^k)}{\text{next to last}} & + & \frac{m 1^k}{\text{last}}
 \end{array}$$

Sum all and divide by m — telescoping sum

$$\frac{m}{m} \left(m_1^k + m_2^k + \dots + m_n^k \right) = F_k$$

So expectation correct!

What about variance?

Use inequality for $a > b > 0$

(10)

$$a^k - b^k = (a-b)(a^{k-1} + a^{k-2}b + \dots + ab^{k-2} + b^{k-1}) \\ \leq (a-b) k a^{k-1}$$

To estimate variance, bound $E[x^2] =$

$$\frac{m^2}{m} \left[(m_1^k - (m_1 - 1)^k)^2 + \dots + (2^{k-1})^2 + 1^{2k} \right. \\ \left. + (m_2^k - (m_2 - 1)^k)^2 + \dots + (2^{k-1})^2 + 1^{2k} + \dots \right. \\ \left. + (m_n^k - (m_n - 1)^k)^2 + \dots + (2^{k-1})^2 + 1^{2k} \right] \\ \leq \dots \quad [\text{see AMS '99 page 139}] \\ \leq m \left[k m_1^{2k-1} + k m_2^{2k-1} + \dots + k m_n^{2k-1} \right] \\ = k m F_{2k-1} = k F, F_{2k-1} = *$$

Would like sth in terms of F_k , though...

Fact: For any $m_i > 0$

$$\left(\sum_{i=1}^n m_i \right) \left(\sum_{i=1}^n m_i^{2k-1} \right) \leq n^{1-1/k} \left(\sum_{i=1}^n m_i^k \right)^2$$

[Again see AMS '99 page 139]

$$\text{Var}(x) \leq E[x^2] \leq k \cdot n^{1-1/k} F_k^2$$

(11)

This variance is too high... How to decrease it? Take many independent samples and average!

Namely $s_1 = 8 k n^{1-1/k} / \lambda^2$ samples

$$\text{Let } Y = \frac{\sum_{i=1}^{s_1} X_i}{s_1}$$

$$\text{Var}(Y) \leq \frac{k \cdot n^{1-1/k} F_k^2}{8 k n^{1-1/k} / \lambda^2} = \frac{\lambda^2}{8} \cdot F_k^2$$

How close should we expect that Y gets to F_k ?

Use Chebyshev's inequality

$$\begin{aligned} \Pr[|Y - F_k| > \lambda F_k] &\leq \Pr[|Y - F_k| > \sqrt{8 \cdot \text{Var}(Y)}] \\ &\leq 1/8 \quad (+) \end{aligned}$$

So with constant probability we are close.
Not good enough! Want ϵ -small error probability!

Do many independent samples of Y ,
and upper-bound probability that
half of these samples fails to get
within λF_k of correct answer.

Do $s_2 = 4 \log(1/\epsilon)$ samples $Y_1, \dots, Y_{s_2}$ of Y and pick median. This median is within λF_k of correct answer unless $(+)$ fails to hold for $1/2$ of the experiments, but should only fail for $\approx 1/8$ fraction.

$$\begin{aligned}
 & \Pr [\text{half of } s_2 \text{ } Y\text{'s wrong}] \\
 &= \Pr [3/8 \text{ more than expected } Y\text{'s wrong}] \\
 &\leq \Pr \left[\left| \frac{\sum_{i=1}^{s_2} Y_i}{s_2} - \frac{1}{8} \right| \geq \frac{3}{8} \right] \\
 &\leq 2 \exp \left(- \frac{\frac{3}{8} \cdot \frac{3}{8} s_2}{2 \cdot \frac{1}{8} \cdot \frac{7}{8}} \right) \\
 &\leq 2 \exp (-2 \log(1/\epsilon)) \\
 &= 2 \exp(2 \log \epsilon) = 2 \epsilon^2 < \epsilon.
 \end{aligned}$$

So the median of $4 \log(1/\epsilon)$ Y 's will get us within λF_k of right answer with prob $> 1 - \epsilon$.

Need a total of to compute

$$4 \log(1/\epsilon) \cdot 8 k n^{1-1/\epsilon} / \lambda^2$$

X samples, each requiring space $\approx \log n + \log m$.
 $\Omega(n^2)$

Except we don't know $m = \text{length of } A$!
What to do?

Suppose we have an algorithm for upto m elements. Now $(m+1)$ st element comes along

We have computed all X_{ij} so far for
 $i = 1, \dots, s_1, j = 1, \dots, s_2$
Every X_{ij} has its a_p and a_r ← (uniform over $a_1, \dots, a_m$)

With probability $1/(m+1)$ (independently for every X_{ij}), replace a_p by a_{m+1} and update $r = 1$

With prob $1 - 1/(m+1)$, keep a_p and increment r if $a_{m+1} = a_p$

Then after this step all $a_1, a_2, \dots, a_{m+1}$ are equally likely to have been chosen, and r is correct.

This completes the proof (modulo computations that we skipped).

In particular, can compute F_2
(or at least approximate closely with arbitrarily
high probability) in

$$O(\sqrt{n} (\log n + \log m)) \text{ memory}$$

(fixing closeness λ and probability $1-\epsilon$).

But AMS '99 shows possible
to do much better

Thm 3 For every $\lambda > 0$, $\epsilon > 0$ $\exists$ randomized
algorithm that given $A = (a_1, \dots, a_m)$
 $a_i \in N = \{1, 2, \dots, n\}$, makes one pass over
 A , computes Y s.t.

$$\Pr[|Y - F_2| > \lambda F_2] \leq \epsilon$$

and uses $O\left(\frac{\log(1/\epsilon)}{\lambda^2} (\log n + \log m)\right)$

bits of memory.

Proof uses same general idea, but needs
some neat ideas about somewhat-but-not-quite
independent ~~distributions~~ random variables and
some nontrivial abstract algebra.

OK, nice algorithms. But let's not ignore the disadvantages:

- ① Uses lots of memory for large k
- ② Doesn't give correct answers, only approximation
- ③ And even this only with high probability over random choices. But computers are deterministic, so shouldn't there be a deterministic algorithm.

Next lecture Show that limitations ① - ③ are inherent.

It is impossible to get around them!