

Sketching for Data Streams

Michael Kapralov

EPFL

August 30, 2023

Streaming model (Alon, Matias, Szegedy'96)

Observe a (very long) stream of data, e.g. IP packets, tweets, search queries....

Task: maintain (approximate) statistics of the stream

Streaming model

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

Typically, assume N is known

- ▶ Small (sublinear) storage: typically $N^\alpha, \alpha < 1$ or $\log^{O(1)} N$

Units of storage: bits, words or 'data items' (e.g., points, nodes/edges)

- ▶ Fast processing time per element
- ▶ Mostly randomized algorithms

Randomness often necessary

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!

1 2 3 4 5 6 7 8 9 10

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

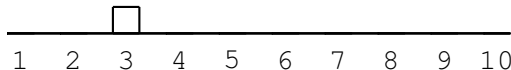
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

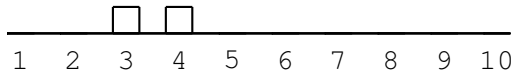
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

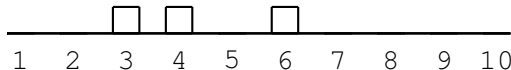
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

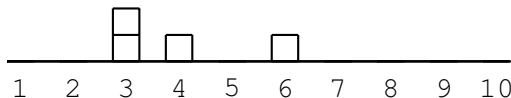
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

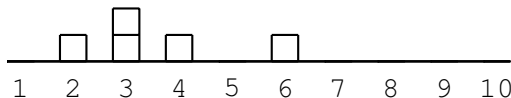
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3 2

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

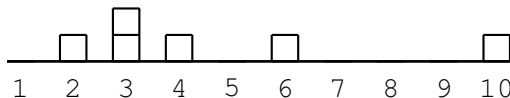
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3 2 10

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

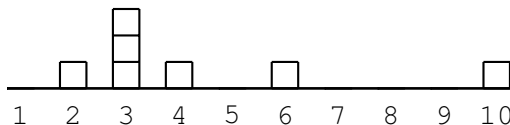
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3 2 10 3

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

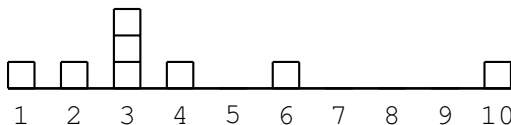
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3 2 10 3 1

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

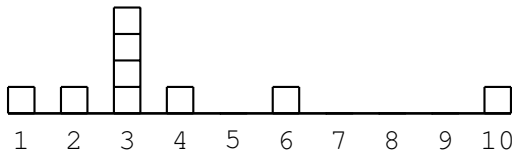
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3 2 10 3 1 3

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

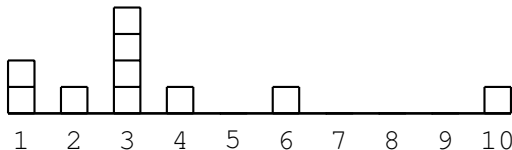
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3 2 10 3 1 3 1

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

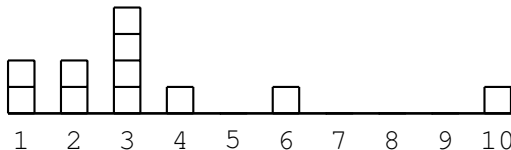
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3 2 10 3 1 3 1 2

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

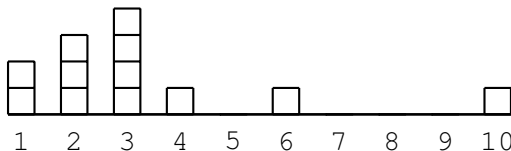
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3 2 10 3 1 3 1 2 2

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

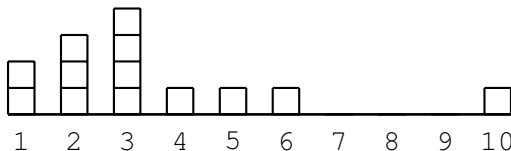
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3 2 10 3 1 3 1 2 2 5

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

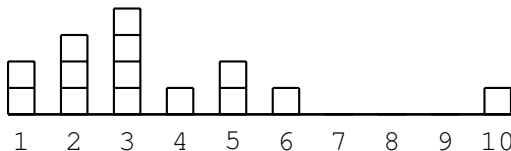
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3 2 10 3 1 3 1 2 2 5 5

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

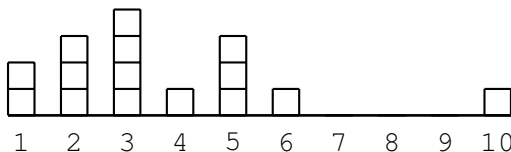
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3 2 10 3 1 3 1 2 2 5 5 5

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

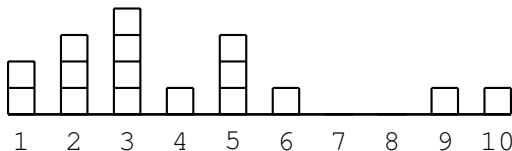
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3 2 10 3 1 3 1 2 2 5 5 5 9

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

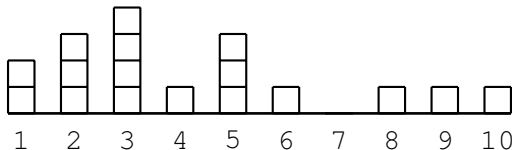
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3 2 10 3 1 3 1 2 2 5 5 5 9 8

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

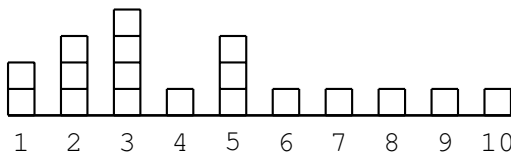
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3 2 10 3 1 3 1 2 2 5 5 5 9 8 7

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

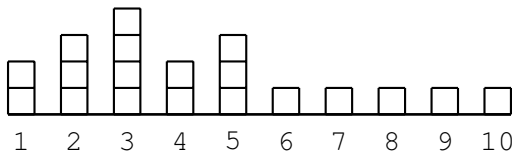
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3 2 10 3 1 3 1 2 2 5 5 5 9 8 7 4

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

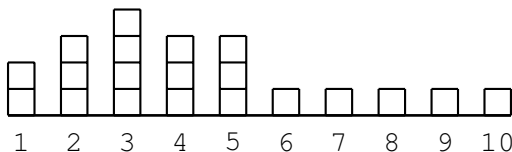
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3 2 10 3 1 3 1 2 2 5 5 5 9 8 7 4 4

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

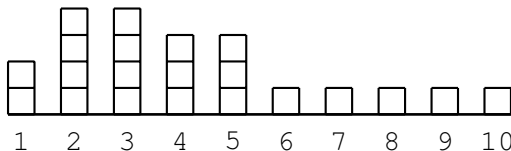
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3 2 10 3 1 3 1 2 2 5 5 5 9 8 7 4 4 2

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

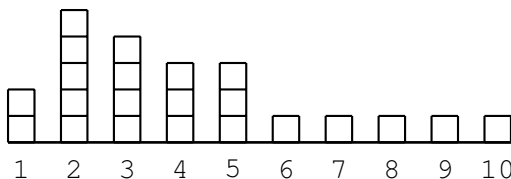
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3 2 10 3 1 3 1 2 2 5 5 5 9 8 7 4 4 2 2

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

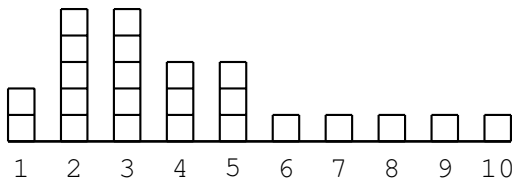
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3 2 10 3 1 3 1 2 2 5 5 5 9 8 7 4 4 2 2 3

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

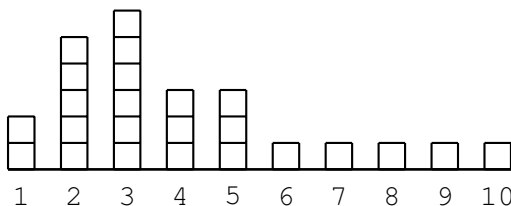
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3 2 10 3 1 3 1 2 2 5 5 5 9 8 7 4 4 2 2 3 3

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

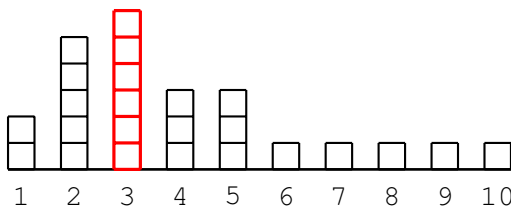
Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!



3 4 6 3 2 10 3 1 3 1 2 2 5 5 5 9 8 7 4 4 2 2 3 3

Estimating IP flows through a router



Estimate the dominant IP flows through a router

[illegible]

Estimating IP flows through a router

[illegible]

Estimating IP flows through a router



source	destination									
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	1	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0

Src	Dst
DATA	

Estimating IP flows through a router



source	destination									
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	1	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0

Src	Dst
DATA	

Estimating IP flows through a router



source	destination									
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	2	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	1	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0

Src	Dst
DATA	

Estimating IP flows through a router



source	destination									
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	2	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	1	0	0	0	0	0	0	0	0
	0	0	0	0	0	1	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0

Src	Dst
DATA	

Estimating IP flows through a router



source	destination									
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	2	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	1	0	0	0	0	0	0	0	0
	0	0	0	0	0	1	0	0	0	0
	0	0	1	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0

Src	Dst
DATA	

Estimating IP flows through a router



source	destination									
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	3	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	1	0	0	0	0	0	0	0	0
	0	0	0	0	0	1	0	0	0	0
	0	0	1	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0

Src	Dst
DATA	

Estimating IP flows through a router

[illegible]

Estimating IP flows through a router



		destination								
source	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	1	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	3	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	1	0	0	0	0	0	0	0	0
	0	0	0	0	0	1	0	0	0	0
	0	0	1	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0

Src	Dst
DATA	

Estimating IP flows through a router



	destination								
source	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	1	0	0
	0	0	0	0	0	0	0	0	0
	0	0	0	3	0	0	0	0	0
	0	0	0	0	0	0	0	0	0
	0	1	0	0	0	0	0	0	0
	0	0	0	0	0	1	0	0	0
	0	0	1	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0

source

Estimating IP flows through a router



source	destination								
	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	1	0	0
	0	0	0	0	0	0	0	0	0
	0	0	0	3	0	0	0	0	0
	0	0	0	0	0	0	0	0	0
	0	2	0	0	0	0	0	0	0
	0	0	0	0	0	1	0	0	0
	0	0	1	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0

Src	Dst
DATA	

Estimating IP flows through a router

[illegible]

Estimating IP flows through a router



source	destination									
	0	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	1	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	3	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	2	0	0	0	0	0	1	0	0
	0	0	0	0	0	1	0	0	0	0
	0	0	1	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0

Src	Dst
DATA	

Estimating IP flows through a router



source	destination								
	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	1	0	0
	0	0	0	0	0	0	0	0	0
	0	0	0	4	0	0	0	0	0
	0	0	0	0	0	0	0	0	0
	0	2	0	0	0	0	0	1	0
	0	0	0	0	0	1	0	0	0
	0	0	1	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0

Src	Dst
DATA	

Estimating IP flows through a router



source	destination								
	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	1	0	0
	0	0	0	0	0	0	0	0	0
	0	0	0	4	0	0	0	0	0
	0	0	0	0	0	0	0	0	0
	0	3	0	0	0	0	0	1	0
	0	0	0	0	0	1	0	0	0
	0	0	1	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0

Src	Dst
DATA	

Estimating IP flows through a router



	destination								
source	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	1	0	0
	0	0	0	0	0	0	0	0	0
	0	0	0	4	0	0	0	0	0
	0	0	0	0	0	0	0	0	0
	0	3	0	0	0	0	0	1	0
	0	0	0	0	0	1	0	0	0
	0	0	1	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0

source

Estimating IP flows through a router



source	destination								
	1	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	1	0	0
	0	0	0	0	0	0	0	0	0
	0	0	0	4	0	0	0	0	0
	0	0	0	0	0	0	0	0	0
	0	3	0	0	0	0	0	1	0
	0	0	0	0	0	1	0	0	0
	0	0	1	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0

Src	Dst
DATA	

Estimating IP flows through a router



source	destination								
	1	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	1	0	0
	0	0	0	0	0	0	0	0	0
	0	0	0	4	0	0	0	0	0
	0	0	0	0	0	0	0	0	0
	0	4	0	0	0	0	0	1	0
	0	0	0	0	0	1	0	0	0
	0	0	1	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0

Src	Dst
DATA	

Estimating IP flows through a router



source	destination									
	1	0	0	0	0	0	0	0	0	0
	0	0	0	0	0	0	1	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	0	0	5	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0
	0	4	0	0	0	0	0	1	0	0
	0	0	0	0	0	1	0	0	0	0
	0	0	1	0	0	0	0	0	0	0
	0	0	0	0	0	0	0	0	0	0

Src	Dst
DATA	

Estimating IP flows through a router



Estimate the **dominant IP flows** through a router



Estimating IP flows through a router



Estimate the **dominant IP flows** through a router



Trivial: store all distinct IP pairs

Space complexity: $\Theta(N)$

Estimating IP flows through a router



Estimate the **dominant IP flows** through a router



Trivial: store all distinct IP pairs

Space complexity: $\Theta(N)$

This lecture: solve in space $O(\log N)$

Exponential improvement!

Estimating search statistics

Given a set of items as a stream (e.g. queries on `google.com`
over a period of time)

Geneva to NYC, coffee in Geneva, Geneva to NYC

Estimating search statistics

Given a set of items as a stream (e.g. queries on `google.com` over a period of time)

Geneva to NYC, coffee in Geneva, Geneva to NYC

Find the most frequent items in the set

Geneva to NYC, coffee in Geneva

Estimating search statistics

Given a set of items as a stream (e.g. queries on `google.com` over a period of time)

Geneva to NYC, coffee in Geneva, Geneva to NYC

Find the most frequent items in the set

Geneva to NYC, coffee in Geneva

Estimating search statistics

Given a set of items as a stream (e.g. queries on `google.com` over a period of time)

Geneva to NYC, coffee in Geneva, Geneva to NYC

Find the most frequent items in the set

Geneva to NYC, coffee in Geneva

	Trivial	This lecture
Solution	<code>hash<string> h;</code>	COUNTSKETCH
Space	# of distinct items	$O(\log N)$

Streaming model

	Trivial	This lecture
Solution	<code>hash<string> h;</code>	COUNTSKETCH
Space	# of distinct items	$O(\log N)$

Streaming model

	Trivial	This lecture
Solution	<code>hash<string> h;</code>	COUNTSKETCH
Space	# of distinct items	$O(\log N)$

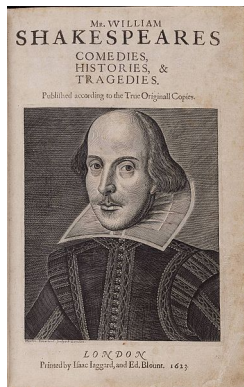
Are constants small?

Streaming model

	Trivial	This lecture
Solution	<code>hash<string> h;</code>	COUNTSKETCH
Space	# of distinct items	$O(\log N)$

Are constants small?

HyperLogLog: estimate
Shakespeare's vocabulary
using 128 bits of memory



Streaming model

Widely used in practice for **scalable data analytics**



most frequent searches on `google.com`
over a time period



most frequent tweets

Heavy hitters problem

- ▶ Single pass over the data: $i_1, i_2, \dots, i_N$

Assume N is known

- ▶ Output k most frequent items

(Heavy hitters)

- ▶ Small storage: will get $O(k \log N)$

Much better than storing all items!

Goal: design a small space data structure

FINDTOP(S, k): returns top k most frequent items seen so far

Goal: design a small space data structure

FINDTOP(S, k): returns top k most frequent items seen so far

Useful to first design

POINTQUERY(S, i): processes stream, then for any query item i
can return f_i =number of times item i appeared

Denote the number of times item i appears in the stream by f_i
(frequency of i)

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$

Denote the number of times item i appears in the stream by f_i
(frequency of i)

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$

POINTQUERY(S, i) in space $O(k \log N)$?

Denote the number of times item i appears in the stream by f_i
(frequency of i)

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$

POINTQUERY(S, i) in space $O(k \log N)$?

Impossible in general...

Imagine a stream where all elements occur with about the same frequency

FINDAPPROXTOP(S, k, ϵ): returns set of k items such that $f_i \geq (1 - \epsilon)f_k$ for all reported i

APPROXPOINTQUERY(S, i, ϵ): processes stream, then for any query item i can return **approximation** $\hat{f}_i \in [f_i - \epsilon f_k, f_i + \epsilon f_k]$

FINDAPPROXTOP(S, k, ϵ): returns set of k items such that $f_i \geq (1 - \epsilon)f_k$ for all reported i

APPROXPOINTQUERY(S, i, ϵ): processes stream, then for any query item i can return **approximation** $\hat{f}_i \in [f_i - \epsilon f_k, f_i + \epsilon f_k]$

In this lecture: find most frequent (**head**) items **if they contribute the bulk of the stream** under some measure

In what follows: APPROXPOINTQUERY in small space

Observe a stream of updates, maintain small space data structure

Task: after observing the stream, given $i \in \{1, 2, \dots, m\}$, compute estimate $\hat{f}_i$ of f_i

In what follows: APPROXPOINTQUERY in small space

Observe a stream of updates, maintain small space data structure

Task: after observing the stream, given $i \in \{1, 2, \dots, m\}$, compute estimate $\hat{f}_i$ of f_i

To be specified:

- ▶ space complexity?
- ▶ quality of approximation?
- ▶ success probability?

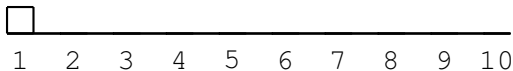
1. Finding top k elements via (APPROX)POINTQUERY
2. Basic version of APPROXPOINTQUERY
3. APPROXPOINTQUERY and the COUNTSKETCH algorithm

1. Finding top k elements via (APPROX)POINTQUERY
2. **Basic version of APPROXPOINTQUERY**
3. APPROXPOINTQUERY and the COUNTSKETCH algorithm

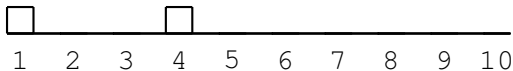
Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$

1 2 3 4 5 6 7 8 9 10

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$

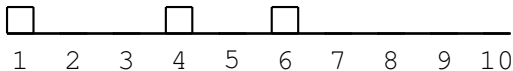


Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



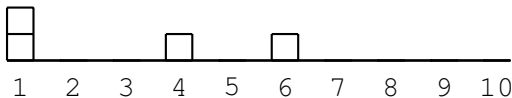
1 4

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



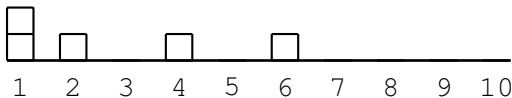
1 4 6

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



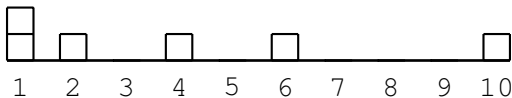
1 4 6 1

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



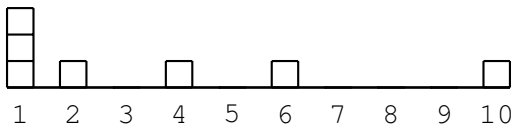
1 4 6 1 2

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



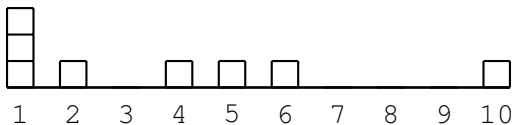
1 4 6 1 2 10

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



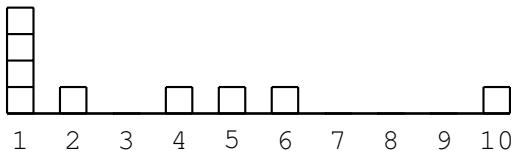
1 4 6 1 2 10 1

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



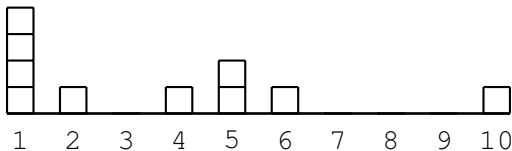
1 4 6 1 2 10 1 5

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



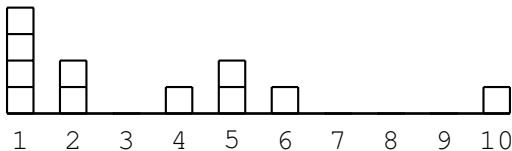
1 4 6 1 2 10 1 5 1

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



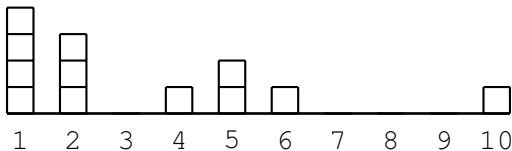
1 4 6 1 2 10 1 5 1 5

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



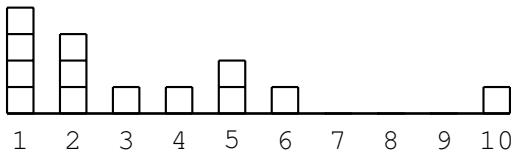
1 4 6 1 2 10 1 5 1 5 2

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



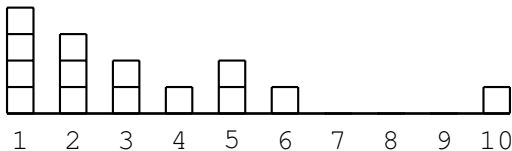
1 4 6 1 2 10 1 5 1 5 2 2

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



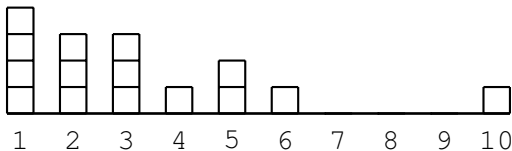
1 4 6 1 2 10 1 5 1 5 2 2 3

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



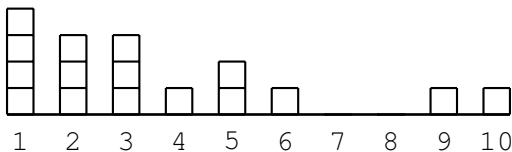
1 4 6 1 2 10 1 5 1 5 2 2 3 3

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



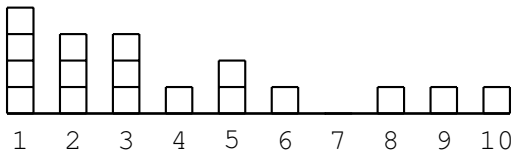
1 4 6 1 2 10 1 5 1 5 2 2 3 3 3

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



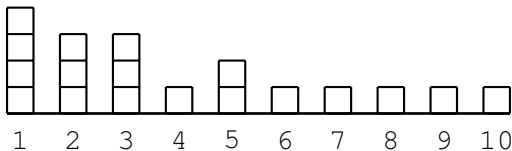
1 4 6 1 2 10 1 5 1 5 2 2 3 3 3 9

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



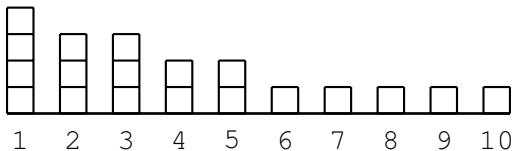
1 4 6 1 2 10 1 5 1 5 2 2 3 3 3 9 8

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



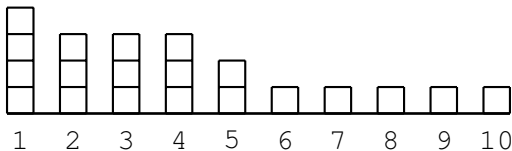
1 4 6 1 2 10 1 5 1 5 2 2 3 3 3 9 8 7

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



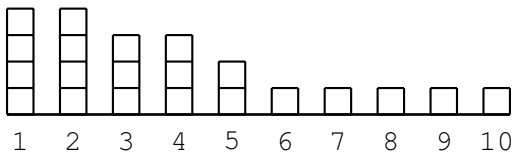
1 4 6 1 2 10 1 5 1 5 2 2 3 3 3 9 8 7 4

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



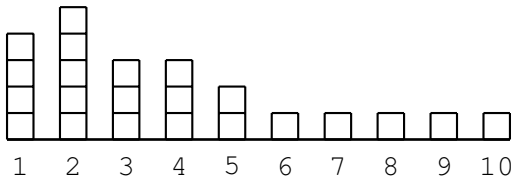
1 4 6 1 2 10 1 5 1 5 2 2 3 3 3 9 8 7 4 4

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



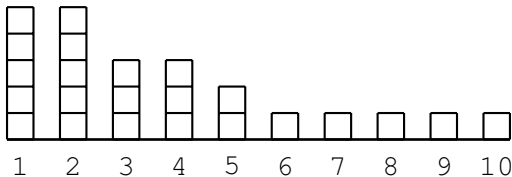
1 4 6 1 2 10 1 5 1 5 2 2 3 3 3 9 8 7 4 4 2

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



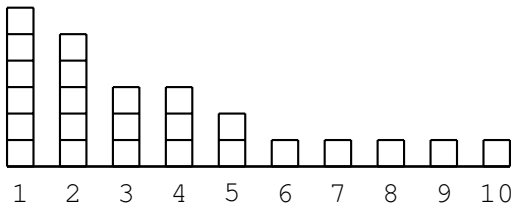
1 4 6 1 2 10 1 5 1 5 2 2 3 3 3 9 8 7 4 4 2 2

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



1 4 6 1 2 10 1 5 1 5 2 2 3 3 3 9 8 7 4 4 2 2 1

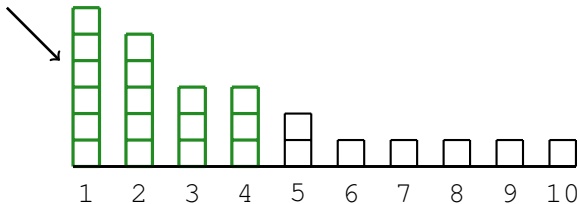
Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



1 4 6 1 2 10 1 5 1 1 2 2 3 3 3 9 8 7 4 4 2 2 1 5

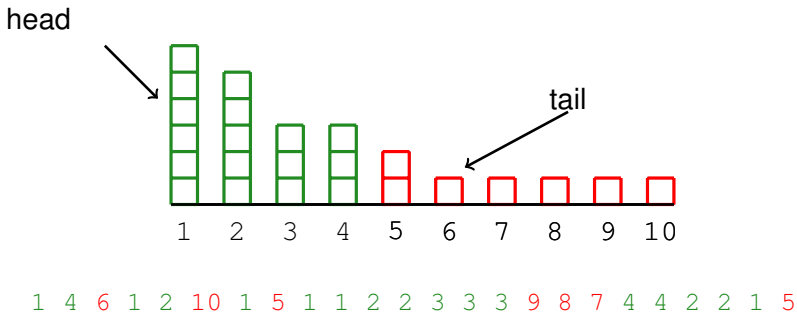
Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$

head



1 4 6 1 2 10 1 5 1 1 2 2 3 3 3 9 8 7 4 4 2 2 1 5

Assume elements are ordered by frequency: $f_1 \geq f_2 \geq \dots \geq f_m$



Basic estimate

Will design a basic estimate with $O(1)$ space complexity,
analyze precision

Basic estimate

Will design a basic estimate with $O(1)$ space complexity,
analyze precision

Choose a **hash function** $s : [m] \rightarrow \{-1, +1\}$ uniformly at random

INITIALIZE

$C \leftarrow 0$

UPDATE(C, i)

$C \leftarrow C + s(i)$

Basic estimate

Will design a basic estimate with $O(1)$ space complexity,
analyze precision

Choose a **hash function** $s : [m] \rightarrow \{-1, +1\}$ uniformly at random

INITIALIZE

$C \leftarrow 0$

UPDATE(C, i)

$C \leftarrow C + s(i)$

for every $p = 1, \dots, N$ (every element in the stream)

 UPDATE(C, i_p)

end for

Basic estimate

Will design a basic estimate with $O(1)$ space complexity,
analyze precision

Choose a **hash function** $s: [m] \rightarrow \{-1, +1\}$ uniformly at random

INITIALIZE

$C \leftarrow 0$

UPDATE(C, i)

$C \leftarrow C + s(i)$

for every $p = 1, \dots, N$ (every element in the stream)

 UPDATE(C, i_p)

end for

ESTIMATE(C, i)

return $C \cdot s(i)$

Show that $C \cdot s(i)$ is close to f_i 'with high probability'?

UPDATE(C, i)

$C \leftarrow C + s(i)$

ESTIMATE(C, i)

return $C \cdot s(i)$

Show that $C \cdot s(i)$ is close to f_i 'with high probability'?

UPDATE(C, i)
 $C \leftarrow C + s(i)$

ESTIMATE(C, i)
return $C \cdot s(i)$

Show that $C \cdot s(i)$ is close to f_i 'with high probability'?

Two steps:

- ▶ show that $\mathbf{E}_s[C \cdot s(i)] = f_i$
(so $C \cdot s(i)$ is an **unbiased estimate** of f_i)
- ▶ show that $\mathbf{Var}_s[C \cdot s(i)]$ is 'small'

Basic estimate:mean

UPDATE(C, i)

$C \leftarrow C + s(i)$

ESTIMATE(C, i)

return $C \cdot s(i)$

$$C \cdot s(i) = \sum_{p=1}^N s(i_p) s(i)$$

Basic estimate:mean

UPDATE(C, i)
 $C \leftarrow C + s(i)$

ESTIMATE(C, i)
return $C \cdot s(i)$

$$C \cdot s(i) = \sum_{p=1}^N s(i_p) s(i) = \sum_{j \in [m]} f_j \cdot s(j) s(i)$$

Basic estimate:mean

UPDATE(C, i)
 $C \leftarrow C + s(i)$

ESTIMATE(C, i)
return $C \cdot s(i)$

$$\begin{aligned} C \cdot s(i) &= \sum_{p=1}^N s(i_p) s(i) = \sum_{j \in [m]} f_j \cdot s(j) s(i) \\ &= f_i s(i)^2 + \sum_{j \in [m] \setminus i} f_j \cdot s(j) s(i) \end{aligned}$$

Basic estimate:mean

UPDATE(C, i)
 $C \leftarrow C + s(i)$

ESTIMATE(C, i)
return $C \cdot s(i)$

$$\begin{aligned} C \cdot s(i) &= \sum_{p=1}^N s(i_p) s(i) = \sum_{j \in [m]} f_j \cdot s(j) s(i) \\ &= f_i s(i)^2 + \sum_{j \in [m] \setminus i} f_j \cdot s(j) s(i) \\ &= f_i + \sum_{j \in [m] \setminus i} f_j \cdot s(j) s(i) \quad \leftarrow \text{random } \pm 1\text{'s} \end{aligned}$$

Basic estimate:mean

UPDATE(C, i)

$C \leftarrow C + s(i)$

ESTIMATE(C, i)

return $C \cdot s(i)$

$$C \cdot s(i) = f_i + \sum_{j \in [m] \setminus i} f_j \cdot s(j) s(i)$$

Basic estimate:mean

UPDATE(C, i)

$C \leftarrow C + s(i)$

ESTIMATE(C, i)

return $C \cdot s(i)$

$$\mathbf{E}[C \cdot s(i)] = f_i + \mathbf{E}\left[\sum_{j \in [m] \setminus i} f_j \cdot s(j)s(i)\right]$$

Basic estimate:mean

UPDATE(C, i)
 $C \leftarrow C + s(i)$

ESTIMATE(C, i)
return $C \cdot s(i)$

$$\begin{aligned}\mathbf{E}[C \cdot s(i)] &= f_i + \mathbf{E}\left[\sum_{j \in [m] \setminus i} f_j \cdot s(j) s(i)\right] \\ &= f_i + \sum_{j \in [m] \setminus i} f_j \cdot \mathbf{E}[s(j)] \mathbf{E}[s(i)] \quad (\text{by independence of } s(i))\end{aligned}$$

Basic estimate:mean

UPDATE(C, i)
 $C \leftarrow C + s(i)$

ESTIMATE(C, i)
return $C \cdot s(i)$

$$\begin{aligned}\mathbf{E}[C \cdot s(i)] &= f_i + \mathbf{E}\left[\sum_{j \in [m] \setminus i} f_j \cdot s(j) s(i)\right] \\ &= f_i + \sum_{j \in [m] \setminus i} f_j \cdot \mathbf{E}[s(j)] \mathbf{E}[s(i)] \quad (\text{by independence of } s(i)) \\ &= f_i\end{aligned}$$

Basic estimate:mean

UPDATE(C, i)
 $C \leftarrow C + s(i)$

ESTIMATE(C, i)
return $C \cdot s(i)$

$$\begin{aligned}\mathbf{E}[C \cdot s(i)] &= f_i + \mathbf{E}\left[\sum_{j \in [m] \setminus i} f_j \cdot s(j) s(i)\right] \\ &= f_i + \sum_{j \in [m] \setminus i} f_j \cdot \mathbf{E}[s(j)] \mathbf{E}[s(i)] \quad (\text{by independence of } s(i)) \\ &= f_i\end{aligned}$$

The mean is correct: our estimator is unbiased!

Basic estimate:mean

UPDATE(C, i)
 $C \leftarrow C + s(i)$

ESTIMATE(C, i)
return $C \cdot s(i)$

$$\begin{aligned}\mathbf{E}[C \cdot s(i)] &= f_i + \mathbf{E}\left[\sum_{j \in [m] \setminus i} f_j \cdot s(j) s(i)\right] \\ &= f_i + \sum_{j \in [m] \setminus i} f_j \cdot \mathbf{E}[s(j)] \mathbf{E}[s(i)] \quad (\text{by independence of } s(i)) \\ &= f_i\end{aligned}$$

The mean is correct: our estimator is unbiased!

Is the estimate $C \cdot s(i)$ close to f_i with high probability?

Basic estimate: variance

UPDATE(C, i)

$C \leftarrow C + s(i)$

We have

ESTIMATE(C, i)

return $C \cdot s(i)$

$$C \cdot s(i) = f_i + \sum_{j \in [m] \setminus i} f_j \cdot s(j) s(i)$$

and

$$\mathbf{E}[C \cdot s(i)] = f_i.$$

Basic estimate: variance

UPDATE(C, i)
 $C \leftarrow C + s(i)$

ESTIMATE(C, i)
return $C \cdot s(i)$

We have

$$C \cdot s(i) = f_i + \sum_{j \in [m] \setminus i} f_j \cdot s(j) s(i)$$

and

$$\mathbf{E}[C \cdot s(i)] = f_i.$$

We need to bound

$$\begin{aligned} \mathbf{Var}(C \cdot s(i)) &= \mathbf{E}[(C \cdot s(i) - \mathbf{E}[C \cdot s(i)])^2] \\ &= \mathbf{E}[(C \cdot s(i) - f_i)^2] \\ &= \mathbf{E} \left[\left(\sum_{j \in [m] \setminus i} f_j \cdot s(j) s(i) \right)^2 \right] \end{aligned}$$

Basic estimate: variance

UPDATE(C, i)
 $C \leftarrow C + s(i)$

ESTIMATE(C, i)
 return $C \cdot s(i)$

$$\begin{aligned}(C \cdot s(i) - f_i)^2 &= \left(\sum_{j \in [m] \setminus i} f_j \cdot s(j) s(i) \right)^2 \\&= \sum_{j \in [m] \setminus i} \sum_{j' \in [m] \setminus i} f_j f_{j'} \cdot s(j) s(j') \cdot s^2(i) \\&= \sum_{j \in [m] \setminus i} \sum_{j' \in [m] \setminus i} f_j f_{j'} \cdot s(j) s(j')\end{aligned}$$

Basic estimate: variance

UPDATE(C, i)
 $C \leftarrow C + s(i)$

ESTIMATE(C, i)
return $C \cdot s(i)$

$$\begin{aligned}\mathbf{E}[(C \cdot s(i) - f_i)^2] &= \mathbf{E}\left[\sum_{j \in [m] \setminus i} \sum_{j' \in [m] \setminus i} f_j f_{j'} \cdot s(j) s(j')\right] \\ &= \sum_{j \in [m] \setminus i} \sum_{j' \in [m] \setminus i} f_j f_{j'} \cdot \mathbf{E}[s(j) s(j')] \\ &= \sum_{j \in [m] \setminus i} f_j^2\end{aligned}$$

since

- ▶ $s(j)^2 = 1$ for all j
- ▶ $\mathbf{E}[s(j) s(j')] = \mathbf{E}[s(j)] \mathbf{E}[s(j')] = 0$ for $j \neq j'$.

Basic estimate: variance

UPDATE(C, i)

$C \leftarrow C + s(i)$

ESTIMATE(C, i)

return $C \cdot s(i)$

Basic estimate: variance

UPDATE(C, i)
 $C \leftarrow C + s(i)$

ESTIMATE(C, i)
return $C \cdot s(i)$

We have proved that

$$\mathbf{Var}(C \cdot s(i)) = \mathbf{E}[(C \cdot s(i) - f_i)^2] = \sum_{j \in [m] \setminus i} f_j^2$$

Basic estimate: variance

UPDATE(C, i)
 $C \leftarrow C + s(i)$

ESTIMATE(C, i)
return $C \cdot s(i)$

We have proved that

$$\mathbf{Var}(C \cdot s(i)) = \mathbf{E}[(C \cdot s(i) - f_i)^2] = \sum_{j \in [m] \setminus i} f_j^2$$

By Chebyshev's inequality

$$\mathbf{Pr} \left[|C \cdot s(i) - f_i| \geq 8 \cdot \sqrt{\sum_{j \in [m] \setminus i} f_j^2} \right] \leq 1/64$$

Basic estimate: variance

UPDATE(C, i)
 $C \leftarrow C + s(i)$

ESTIMATE(C, i)
return $C \cdot s(i)$

We have proved that

$$\mathbf{Var}(C \cdot s(i)) = \mathbf{E}[(C \cdot s(i) - f_i)^2] = \sum_{j \in [m] \setminus i} f_j^2$$

By Chebyshev's inequality

$$\mathbf{Pr} \left[|C \cdot s(i) - f_i| \geq 8 \cdot \sqrt{\sum_{j \in [m] \setminus i} f_j^2} \right] \leq 1/64$$

So $C \cdot s(i)$ is close (?) to f_i with high probability

Basic estimate: variance

UPDATE(C, i)
 $C \leftarrow C + s(i)$

ESTIMATE(C, i)
return $C \cdot s(i)$

We have proved that

$$\mathbf{Var}(C \cdot s(i)) = \mathbf{E}[(C \cdot s(i) - f_i)^2] = \sum_{j \in [m] \setminus i} f_j^2$$

By Chebyshev's inequality

$$\mathbf{Pr} \left[|C \cdot s(i) - f_i| > 8 \cdot \sqrt{\sum_{j \in [m] \setminus i} f_j^2} \right] \leq 1/64$$

So $C \cdot s(i)$ is close (?) to f_i with high probability

Basic estimate: summary

UPDATE(C, i)

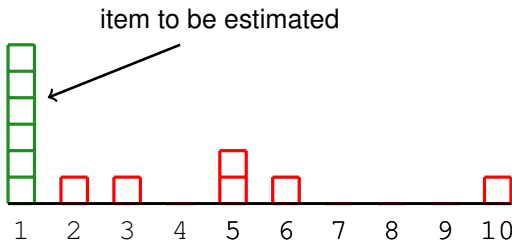
$C \leftarrow C + s(i)$

Estimate f_i up to

ESTIMATE(C, i)

return $C \cdot s(i)$

$$8 \cdot \sqrt{\sum_{j \in [m] \setminus i} f_j^2}$$



Pro: works well for most frequent item, if other items are small

Basic estimate: summary

UPDATE(C, i)

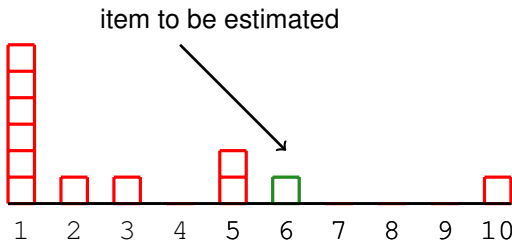
$C \leftarrow C + s(i)$

Estimate f_i up to

ESTIMATE(C, i)

return $C \cdot s(i)$

$$8 \cdot \sqrt{\sum_{j \in [m] \setminus i} f_j^2}$$



Pro: works well for most frequent item, if other items are small

Con: estimate for a small items contaminated by large items

1. Finding top k elements via (APPROX)POINTQUERY
2. Basic version of APPROXPOINTQUERY
3. APPROXPOINTQUERY and the COUNTSKETCH algorithm

1. Finding top k elements via (APPROX)POINTQUERY
2. Basic version of APPROXPOINTQUERY
3. **APPROXPOINTQUERY and the COUNTSKETCH algorithm**

APPROXPOINTQUERY and COUNTSKETCH

COUNTSKETCH algorithm (Charikar, Chen, Farach-Colton'02)

Main ideas:

1. run basic estimate on subsampled/hashed stream
(reduces variance)

APPROXPOINTQUERY and COUNTSKETCH

COUNTSKETCH algorithm ([Charikar, Chen, Farach-Colton'02](#))

Main ideas:

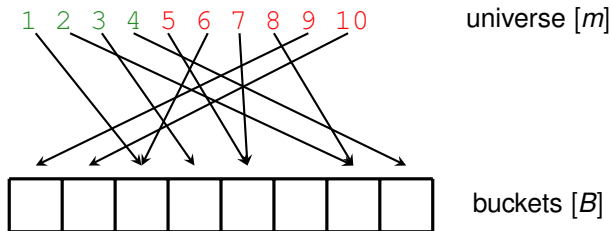
1. run basic estimate on **subsampled/hashed** stream
(reduces **variance**)
2. aggregate independent estimates to boost confidence
(take **medians**)

APPROXPOINTQUERY and COUNTSKETCH

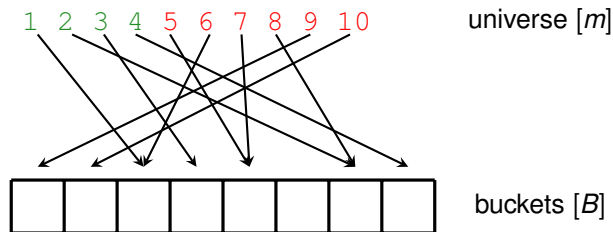
COUNTSKETCH algorithm (Charikar, Chen, Farach-Colton'02)

Main ideas:

1. run basic estimate on subsampled/hashed stream
(reduces variance)
2. aggregate independent estimates to boost confidence
(take medians)



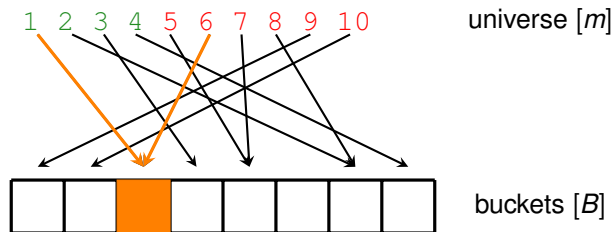
Hashing the items



Hashed into $B = 8$ buckets, get 8 subsampled streams

For item i its stream consists of $j \in [m]$ such that $h(j) = h(i)$

Hashing the items



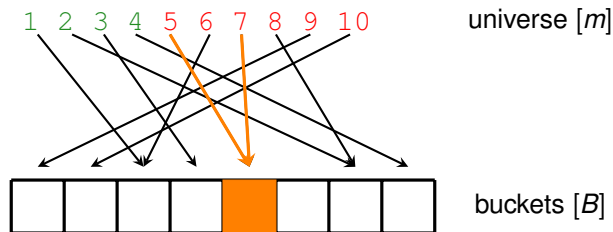
Hashed into $B = 8$ buckets, get 8 subsampled streams

For item i its stream consists of $j \in [m]$ such that $h(j) = h(i)$

For example,

- ▶ subsampled stream of item 1 is {1, 6}

Hashing the items



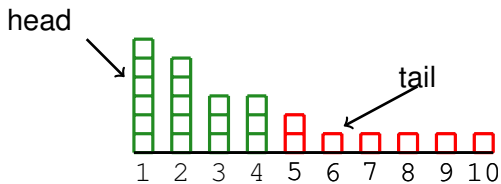
Hashed into $B = 8$ buckets, get 8 subsampled streams

For item i its stream consists of $j \in [m]$ such that $h(j) = h(i)$

For example,

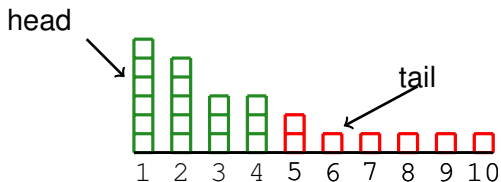
- ▶ subsampled stream of item 1 is {1, 6}
- ▶ subsampled stream of item 5 is {5, 7}

Note: hashing **the universe** $[m]$, not positions in the stream



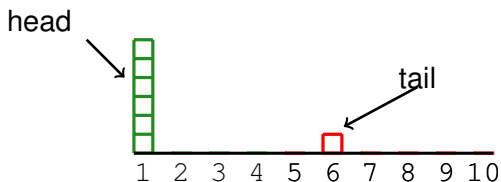
1 4 6 1 2 10 1 5 1 1 2 2 3 3 3 9 8 7 4 4 2 2 1 5

Note: hashing **the universe** $[m]$, not positions in the stream



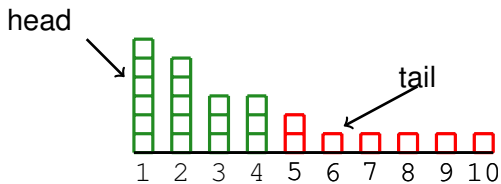
1 4 6 1 2 10 1 5 1 1 2 2 3 3 3 9 8 7 4 4 2 2 1 5

E.x. the subsampled stream of item 1 is {1, 6}



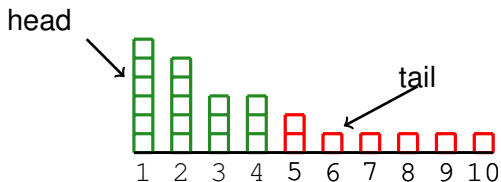
1 6 1 1 1 1 1

Note: hashing **the universe** $[m]$, not positions in the stream



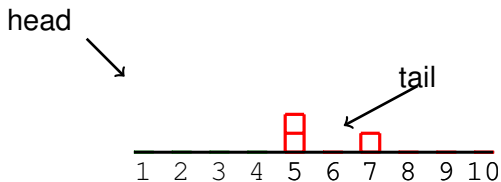
1 4 6 1 2 10 1 5 1 1 2 2 3 3 3 9 8 7 4 4 2 2 1 5

Note: hashing **the universe** $[m]$, not positions in the stream



1 4 6 1 2 10 1 5 1 1 2 2 3 3 3 9 8 7 4 4 2 2 1 5

E.x. the subsampled stream of item 5 is {5, 7}



5

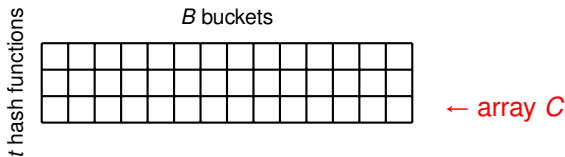
7

5

Final ApproxPointQuery

Choose

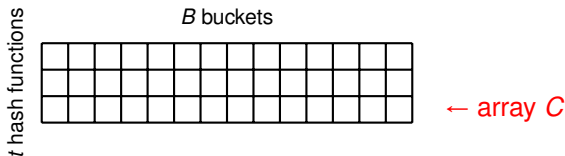
- ▶ t random hash functions $h_1, h_2, \dots, h_t$ from items $[m]$ to $B \approx k$ buckets $\{1, 2, \dots, B\}$
- ▶ t random hash functions $s_1, s_2, \dots, s_t$ from items $[m]$ to $\{-1, +1\}$



Final ApproxPointQuery

Choose

- ▶ t random hash functions $h_1, h_2, \dots, h_t$ from items $[m]$ to $B \approx k$ buckets $\{1, 2, \dots, B\}$
- ▶ t random hash functions $s_1, s_2, \dots, s_t$ from items $[m]$ to $\{-1, +1\}$



The algorithm runs t independent copies of basic estimate:

UPDATE(C, i)

for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

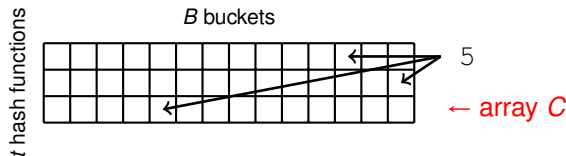
ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$

Final ApproxPointQuery

Choose

- ▶ t random hash functions $h_1, h_2, \dots, h_t$ from items $[m]$ to $B \approx k$ buckets $\{1, 2, \dots, B\}$
- ▶ t random hash functions $s_1, s_2, \dots, s_t$ from items $[m]$ to $\{-1, +1\}$



The algorithm runs t independent copies of basic estimate:

UPDATE(C, i)

for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** _{r} $\{C[r, h_r(i)] \cdot s_r(i)\}$

UPDATE(C, i)

for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$

```

UPDATE(C, i)
for  $r \in [1 : t]$ 
     $C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$ 
end for

ESTIMATE(C, i)
return median $_r \{C[r, h_r(i)] \cdot s_r(i)\}$ 

```

Lemma

If $B \geq 8 \max \left\{ k, \frac{32 \sum_{j \in \text{TAIL}} f_j^2}{(\epsilon f_k)^2} \right\}$ and $t = O(\log N)$, then for every $i \in [m]$

$$|\text{ESTIMATE}(C, i) - f_i| \leq \epsilon f_k$$

at every point in the stream whp.

UPDATE(C, i)	ESTIMATE(C, i)
for $r \in [1 : t]$	return median $_r \{C[r, h_r(i)] \cdot s_r(i)\}$
$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$	
end for	

Lemma

If $B \geq 8 \max \left\{ k, \frac{32 \sum_{j \in \text{TAIL}} f_j^2}{(\epsilon f_k)^2} \right\}$ and $t = O(\log N)$, then for every $i \in [m]$

$$|\text{ESTIMATE}(C, i) - f_i| \leq \epsilon f_k$$

at every point in the stream whp.

Space complexity is $O(B \log N)$

```

UPDATE(C, i)
for  $r \in [1 : t]$ 
     $C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$ 
end for

ESTIMATE(C, i)
return median $_r \{C[r, h_r(i)] \cdot s_r(i)\}$ 

```

Lemma

If $B \geq 8 \max \left\{ k, \frac{32 \sum_{j \in \text{TAIL}} f_j^2}{(\epsilon f_k)^2} \right\}$ and $t = O(\log N)$, then for every $i \in [m]$

$$|\text{ESTIMATE}(C, i) - f_i| \leq \epsilon f_k$$

at every point in the stream whp.

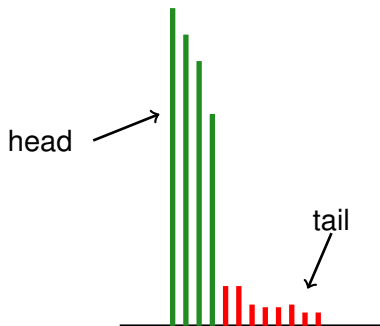
Space complexity is $O(B \log N)$

How large is B ?

Space complexity

$$\text{Set } B = 8 \max \left\{ k, \frac{32 \sum_{j \in \text{TAIL}} f_j^2}{(\epsilon f_k)^2} \right\}$$

Note that $B = O(k/\epsilon^2)$ if $\frac{1}{k} \sum_{j \in \text{TAIL}} f_j^2 = O(f_k^2)$



Note: if $B \geq k$, can detect elements with counts above

$$O\left(\sqrt{\frac{1}{B} \cdot \sum_{j \in \text{TAIL}} f_j^2}\right)$$

Space complexity

Set $k = 1$. Suppose that 1 appears $\sqrt{N}$ times in the stream, and other $N - \sqrt{N}$ elements are distinct

Space complexity

Set $k = 1$. Suppose that 1 appears $\sqrt{N}$ times in the stream, and other $N - \sqrt{N}$ elements are distinct

Then $f_1 = \sqrt{N}$, $f_i = 1$ for $i = 2, N - \sqrt{N}$.

$$\text{Set } B = 8 \max \left\{ 1, \frac{32 \sum_{j \in \text{TAIL}} f_j^2}{(\epsilon f_1)^2} \right\}$$

Space complexity

Set $k = 1$. Suppose that 1 appears $\sqrt{N}$ times in the stream, and other $N - \sqrt{N}$ elements are distinct

Then $f_1 = \sqrt{N}$, $f_i = 1$ for $i = 2, N - \sqrt{N}$.

$$\text{Set } B = 8 \max \left\{ 1, \frac{32 \sum_{j \in \text{TAIL}} f_j^2}{(\epsilon f_1)^2} \right\}$$

We have $\sum_{j \in \text{TAIL}} f_j^2 = N - \sqrt{N} \leq N$, and $f_1^2 = N$

Space complexity

Set $k = 1$. Suppose that 1 appears $\sqrt{N}$ times in the stream, and other $N - \sqrt{N}$ elements are distinct

Then $f_1 = \sqrt{N}$, $f_i = 1$ for $i = 2, N - \sqrt{N}$.

$$\text{Set } B = 8 \max \left\{ 1, \frac{32 \sum_{j \in \text{TAIL}} f_j^2}{(\epsilon f_1)^2} \right\}$$

We have $\sum_{j \in \text{TAIL}} f_j^2 = N - \sqrt{N} \leq N$, and $f_1^2 = N$

So $B = 8 \max \left\{ 1, \frac{32 \sum_{j \in \text{TAIL}} f_j^2}{(\epsilon f_1)^2} \right\} = O(1/\epsilon^2)$ suffices

Space complexity

Set $k = 1$. Suppose that 1 appears $\sqrt{N}$ times in the stream, and other $N - \sqrt{N}$ elements are distinct

Then $f_1 = \sqrt{N}$, $f_i = 1$ for $i = 2, N - \sqrt{N}$.

$$\text{Set } B = 8 \max \left\{ 1, \frac{32 \sum_{j \in \text{TAIL}} f_j^2}{(\epsilon f_1)^2} \right\}$$

We have $\sum_{j \in \text{TAIL}} f_j^2 = N - \sqrt{N} \leq N$, and $f_1^2 = N$

$$\text{So } B = 8 \max \left\{ 1, \frac{32 \sum_{j \in \text{TAIL}} f_j^2}{(\epsilon f_1)^2} \right\} = O(1/\epsilon^2) \text{ suffices}$$

Remarkable, as 1 appears only in $\sqrt{N}$ positions out of N : a vanishingly small fraction of positions!

UPDATE(C, i)

for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$

Lemma

If $B \geq 8 \max \left\{ k, \frac{32 \sum_{j \in \text{TAIL}} f_j^2}{(\epsilon f_k)^2} \right\}$ and $t \geq A \log N$ for an absolute constant $A > 0$, then for every $i \in [m]$

$$|\text{ESTIMATE}(C, i) - f_i| \leq \epsilon f_k$$

with high probability.

(f_i is the frequency of i)

UPDATE(C, i)	ESTIMATE(C, i)
for $r \in [1 : t]$	return median _{r} $\{C[r, h_r(i)] \cdot s_r(i)\}$
$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$	
end for	

Variance of estimate for i from r -th row:

$$\sum_{j \neq i: \mathbf{h}_r(\mathbf{j}) = \mathbf{h}_r(\mathbf{i})} f_j^2$$

```

UPDATE(C, i)
for  $r \in [1 : t]$ 
     $C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$ 
end for

ESTIMATE(C, i)
return  $\text{median}_r \{C[r, h_r(i)] \cdot s_r(i)\}$ 

```

Variance of estimate for i from r -th row:

$$\sum_{j \neq i: \mathbf{h}_r(\mathbf{j}) = \mathbf{h}_r(\mathbf{i})} f_j^2$$

Show that

$$\sum_{j \neq i: \mathbf{h}_r(\mathbf{j}) = \mathbf{h}_r(\mathbf{i})} f_j^2 = O(1/B) \sum_{j \in \text{TAIL}, j \neq i} f_j^2$$

with high constant probability.

Consider contribution of **head** and **tail** items separately:

$$\sum_{j \neq i: h_r(j) = h_r(i)} f_j^2 = \sum_{\substack{j \in \text{HEAD}, j \neq i \\ \mathbf{h}_r(\mathbf{j}) = \mathbf{h}_r(\mathbf{i})}} f_j^2 + \sum_{\substack{j \in \text{TAIL}, j \neq i \\ \mathbf{h}_r(\mathbf{j}) = \mathbf{h}_r(\mathbf{i})}} f_j^2$$

Consider contribution of **head** and **tail** items separately:

$$\sum_{j \neq i: h_r(j) = h_r(i)} f_j^2 = \sum_{\substack{j \in \text{HEAD}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2 + \sum_{\substack{j \in \text{TAIL}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2$$

For each $r \in [1 : t]$ and each item $i \in [m]$ define three events:

- $\text{No-COLLISIONS}_r(i)$ – i does not collide with **any** of the **head** items under hashing r

Consider contribution of **head** and **tail** items separately:

$$\sum_{j \neq i: h_r(j) = h_r(i)} f_j^2 = \sum_{\substack{j \in \text{HEAD}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2 + \sum_{\substack{j \in \text{TAIL}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2$$

For each $r \in [1 : t]$ and each item $i \in [m]$ define three events:

- ▶ $\text{NO-COLLISIONS}_r(i)$ – i does not collide with **any** of the **head** items under hashing r
- ▶ $\text{SMALL-VARIANCE}_r(i)$ – i does not collide with **too many** of **tail** items under hashing r

Consider contribution of **head** and **tail** items separately:

$$\sum_{j \neq i: h_r(j) = h_r(i)} f_j^2 = \sum_{\substack{j \in \text{HEAD}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2 + \sum_{\substack{j \in \text{TAIL}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2$$

For each $r \in [1 : t]$ and each item $i \in [m]$ define three events:

- ▶ $\text{NO-COLLISIONS}_r(i)$ – i does not collide with **any** of the **head** items under hashing r
- ▶ $\text{SMALL-VARIANCE}_r(i)$ – i does not collide with **too many** of **tail** items under hashing r
- ▶ $\text{SMALL-DEVIATION}_r(i)$ – success event from basic analysis

Consider contribution of **head** and **tail** items separately:

$$\sum_{j \neq i: h_r(j) = h_r(i)} f_j^2 = \sum_{\substack{j \in \text{HEAD}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2 + \sum_{\substack{j \in \text{TAIL}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2$$

For each $r \in [1 : t]$ and each item $i \in [m]$ define three events:

- ▶ NO-COLLISIONS $_r(i)$ – i does not collide with **any** of the **head** items under hashing r
- ▶ SMALL-VARIANCE $_r(i)$ – i does not collide with **too many** of **tail** items under hashing r
- ▶ SMALL-DEVIATION $_r(i)$ – success event from basic analysis

Show that **all three events** hold simultaneously with probability strictly bigger than $1/2$ – so median gives good estimate

(No) collisions with head items

$\text{No-COLLISIONS}_r(i) :=$ event that

$$\{j \in \text{HEAD} \setminus i : h_r(j) = h_r(i)\} = \emptyset,$$

i.e. that i collides with none of top k elements under h_r .

(No) collisions with head items

$\text{No-COLLISIONS}_r(i) :=$ event that

$$\{j \in \text{HEAD} \setminus i : h_r(j) = h_r(i)\} = \emptyset,$$

i.e. that i collides with none of top k elements under h_r .

For every $j \neq i$ and every $r \in [1 : t]$

$$\Pr[h_r(i) = h_r(j)] \leq 1/B$$

(No) collisions with head items

$\text{No-COLLISIONS}_r(i)$:= event that

$$\{j \in \text{HEAD} \setminus i : h_r(j) = h_r(i)\} = \emptyset,$$

i.e. that i collides with none of top k elements under h_r .

For every $j \neq i$ and every $r \in [1 : t]$

$$\Pr[h_r(i) = h_r(j)] \leq 1/B$$

Suppose that $B \geq 8k$. Then by the union bound

$$\begin{aligned}\Pr[\text{No-COLLISIONS}_r(i)] &\geq 1 - k/B \\ &\geq 1 - 1/8\end{aligned}$$

Consider contribution of **head** and **tail** items separately:

$$\sum_{j \neq i: h_r(j) = h_r(i)} f_j^2 = \sum_{\substack{j \in \text{HEAD}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2 + \sum_{\substack{j \in \text{TAIL}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2$$

For each $r \in [1 : t]$ and each item $i \in [m]$ define three events:

- ▶ NO-COLLISIONS $_r(i)$ – i does not collide with **any** of the **head** items under hashing r
- ▶ SMALL-VARIANCE $_r(i)$ – i does not collide with **too many** of **tail** items under hashing r
- ▶ SMALL-DEVIATION $_r(i)$ – success event from basic analysis

Show that **all three events** hold simultaneously with probability strictly bigger than $1/2$ – so median gives good estimate

Consider contribution of **head** and **tail** items separately:

$$\sum_{j \neq i: h_r(j) = h_r(i)} f_j^2 = \sum_{\substack{j \in \text{HEAD}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2 + \sum_{\substack{j \in \text{TAIL}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2$$

For each $r \in [1 : t]$ and each item $i \in [m]$ define three events:

- ▶ **NO-COLLISIONS_r(i)** – i does not collide with **any** of the **head** items under hashing r
- ▶ **SMALL-VARIANCE_r(i)** – i does not collide with **too many** of **tail** items under hashing r
- ▶ **SMALL-DEVIATION_r(i)** – success event from basic analysis

Show that **all three events** hold simultaneously with probability strictly bigger than $1/2$ – so median gives good estimate

Small variance from **tail** elements

SMALL-VARIANCE_r(i):=event that

$$\sum_{\substack{j \in \text{TAIL}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2 \leq \frac{8}{B} \sum_{j \in \text{TAIL}} f_j^2$$

Small variance from **tail** elements

SMALL-VARIANCE_r(i) := event that

$$\sum_{\substack{j \in \text{TAIL}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2 \leq \frac{8}{B} \sum_{j \in \text{TAIL}} f_j^2$$

For every $i, j \in [m], i \neq j$ and $r \in [1 : t]$

$$\Pr_{h_r}[h_r(i) = h_r(j)] = 1/B \quad (B \text{ is the number of buckets})$$

Small variance from **tail** elements

SMALL-VARIANCE_r(i) := event that

$$\sum_{\substack{j \in \text{TAIL}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2 \leq \frac{8}{B} \sum_{j \in \text{TAIL}} f_j^2$$

For every $i, j \in [m], i \neq j$ and $r \in [1 : t]$

$$\Pr_{h_r}[h_r(i) = h_r(j)] = 1/B \quad (B \text{ is the number of buckets})$$

So by linearity of expectation

$$\mathbf{E} \left[\sum_{\substack{j \in \text{TAIL}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2 \right] = \sum_{j \in \text{TAIL}, j \neq i} f_j^2 \cdot \Pr_{h_r}[h_r(i) = h_r(j)]$$

Small variance from **tail** elements

$\text{SMALL-VARIANCE}_r(i) := \text{event that}$

$$\sum_{\substack{j \in \text{TAIL}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2 \leq \frac{8}{B} \sum_{j \in \text{TAIL}} f_j^2$$

For every $i, j \in [m], i \neq j$ and $r \in [1 : t]$

$$\Pr_{h_r}[h_r(i) = h_r(j)] = 1/B \quad (B \text{ is the number of buckets})$$

So by linearity of expectation

$$\begin{aligned} \mathbf{E} \left[\sum_{\substack{j \in \text{TAIL}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2 \right] &= \sum_{j \in \text{TAIL}, j \neq i} f_j^2 \cdot \Pr_{h_r}[h_r(i) = h_r(j)] \\ &\leq \frac{1}{B} \sum_{j \in \text{TAIL}} f_j^2 \end{aligned}$$

We proved that

$$\mathbf{E} \left[\sum_{\substack{j \in \text{TAIL}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2 \right] \leq \frac{1}{B} \sum_{j \in \text{TAIL}} f_j^2$$

By Markov's inequality one has, for every i and every r ,

$$\Pr[\text{SMALL-VARIANCE}_r(i)] \geq 1 - 1/8$$

No-COLLISIONS_r(i) and SMALL-VARIANCE_r(i): recap

Consider contribution of **head** and **tail** items separately:

$$\sum_{j \neq i: h_r(j) = h_r(i)} f_j^2 = \sum_{\substack{j \in \text{HEAD}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2 + \sum_{\substack{j \in \text{TAIL}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2$$

Conditioned on No-COLLISIONS_r(i) and SMALL-VARIANCE_r(i)

No-COLLISIONS_r(i) and SMALL-VARIANCE_r(i): recap

Consider contribution of **head** and **tail** items separately:

$$\sum_{j \neq i: h_r(j) = h_r(i)} f_j^2 = \sum_{\substack{j \in \text{HEAD}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2 + \sum_{\substack{j \in \text{TAIL}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2$$

Conditioned on No-COLLISIONS_r(i) and SMALL-VARIANCE_r(i)

- first term is zero

No-COLLISIONS_r(i) and SMALL-VARIANCE_r(i): recap

Consider contribution of **head** and **tail** items separately:

$$\sum_{j \neq i: h_r(j) = h_r(i)} f_j^2 = \sum_{\substack{j \in \text{HEAD}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2 + \sum_{\substack{j \in \text{TAIL}, j \neq i \\ h_r(j) = h_r(i)}} f_j^2$$

Conditioned on No-COLLISIONS_r(i) and SMALL-VARIANCE_r(i)

- ▶ first term is zero
- ▶ second term is at most

$$\frac{8}{B} \sum_{j \in \text{TAIL}} f_j^2$$

Small deviation event

$\text{SMALL-DEVIATION}_r(i)$ = event that

$$(C[r, h_r(i)] \cdot s_r(i) - f_i)^2 \leq 8 \mathbf{Var}(C[r, h_r(i)] \cdot s_r(i)).$$

Small deviation event

$\text{SMALL-DEVIATION}_r(i)$ = event that

$$(C[r, h_r(i)] \cdot s_r(i) - f_i)^2 \leq 8 \mathbf{Var}(C[r, h_r(i)] \cdot s_r(i)).$$

By Markov's inequality one has, for every i and every r ,

$$\mathbf{Pr}[\text{SMALL-DEVIATION}_r(i)] \geq 1 - 1/8$$

$$\Pr[\text{SMALL-VARIANCE}_r(i)] \geq 1 - 1/8$$

$$\Pr[\text{NO-COLLISIONS}_r(i)] \geq 1 - 1/8$$

$$\Pr[\text{SMALL-DEVIATION}_r(i)] \geq 1 - 1/8$$

So by the union bound

$$\Pr[\text{SMALL-VARIANCE}_r(i) \text{ and NO-COLLISIONS}_r(i) \\ \text{and SMALL-DEVIATION}_r(i)] \geq 5/8.$$

For every $p \in [1 : N]$ let $f_i(p) :=$ frequency of i up to position p

Lemma

If $B \geq 8 \max \left\{ k, \frac{32 \sum_{j \in \text{TAIL}} f_j^2}{(\epsilon f_k)^2} \right\}$ and $t \geq A \log N$ for an absolute constant $A > 0$, then with probability $\geq 1 - 1/N^3$ for every $i \in [m]$

$$|\text{ESTIMATE}(C, i) - f_i(p)| \leq \epsilon f_k$$

at the end of the stream.

Remarks, related results, open problems

UPDATE(C, i)

for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$

1 2 3 4 5 6 7 8 9 10

UPDATE(C, i)

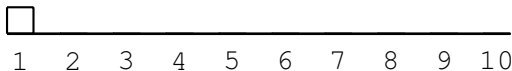
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



UPDATE(C, i)

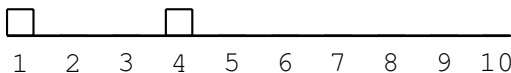
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



UPDATE(C, i)

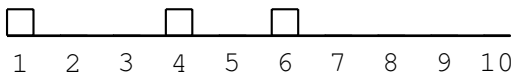
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6

UPDATE(C, i)

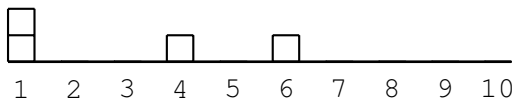
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1

UPDATE(C, i)

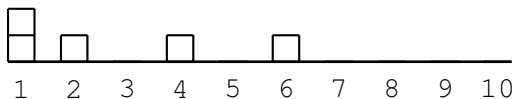
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2

UPDATE(C, i)

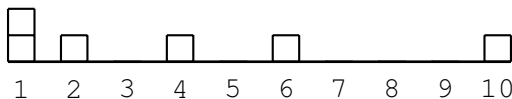
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10

UPDATE(C, i)

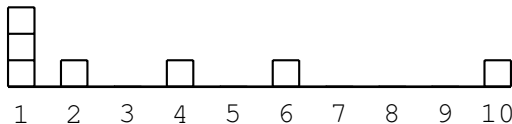
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10 1

UPDATE(C, i)

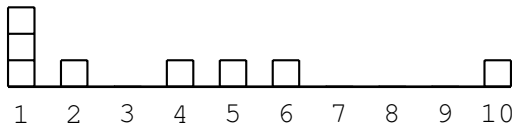
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10 1 5

UPDATE(C, i)

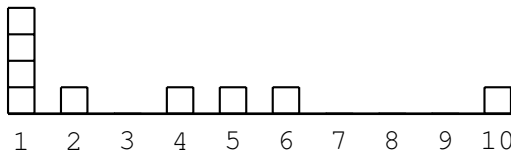
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return $\text{median}_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10 1 5 1

UPDATE(C, i)

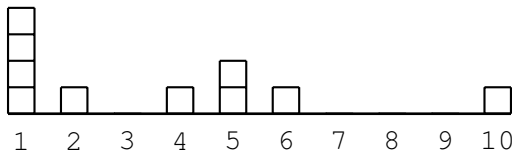
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return $\text{median}_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10 1 5 1 5

UPDATE(C, i)

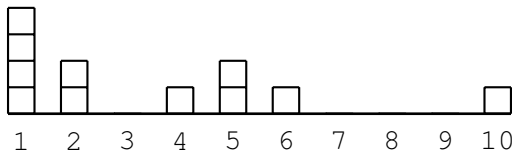
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return $\text{median}_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10 1 5 1 5 2

UPDATE(C, i)

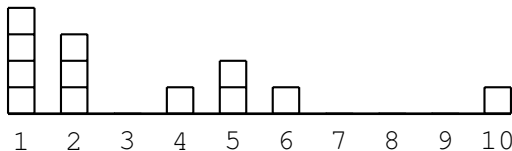
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10 1 5 1 5 2 2

UPDATE(C, i)

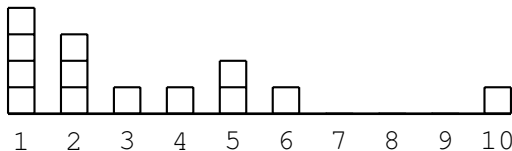
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10 1 5 1 5 2 2 3

UPDATE(C, i)

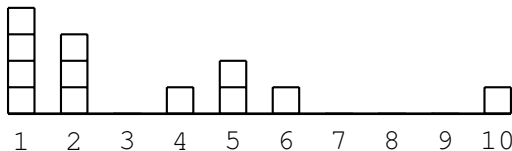
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10 1 5 1 5 2 2 3 -3

UPDATE(C, i)

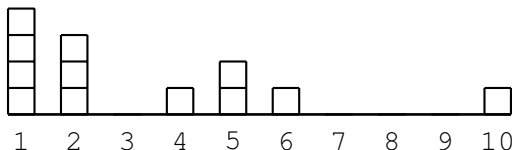
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median**_r $\{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10 1 5 1 5 2 2 3

-3

UPDATE(C, i)

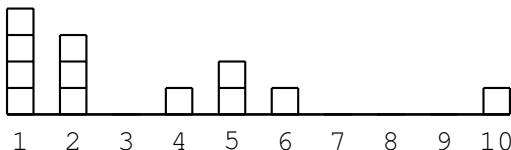
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10 1 5 1 5 2 2 3 -3

UPDATE(C, i)

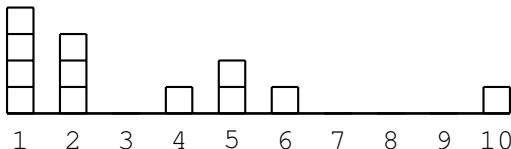
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] - s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10 1 5 1 5 2 2 3 -3

UPDATE(C, i)

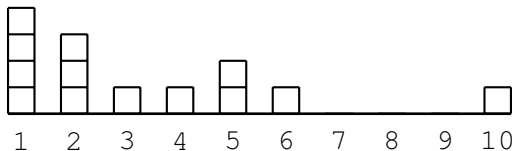
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10 1 5 1 5 2 2 3 -3 3

UPDATE(C, i)

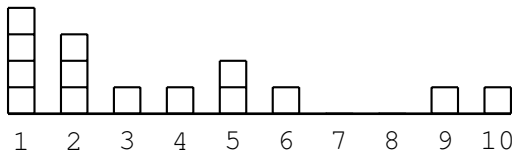
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10 1 5 1 5 2 2 3 -3 3 9

UPDATE(C, i)

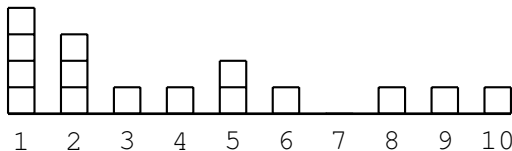
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10 1 5 1 5 2 2 3 -3 3 9 8

UPDATE(C, i)

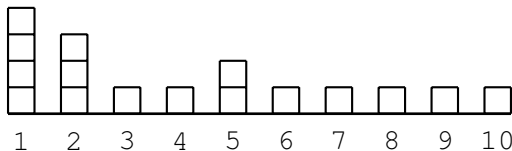
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10 1 5 1 5 2 2 3 -3 3 9 8 7

UPDATE(C, i)

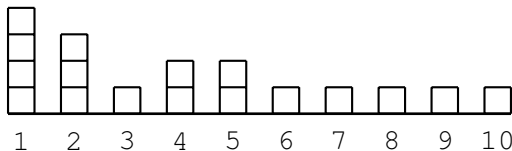
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10 1 5 1 5 2 2 3 -3 3 9 8 7 4

UPDATE(C, i)

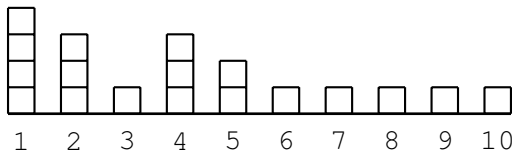
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10 1 5 1 5 2 2 3 -3 3 9 8 7 4 4

UPDATE(C, i)

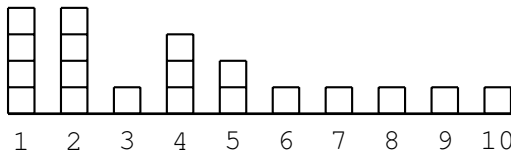
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10 1 5 1 5 2 2 3 -3 3 9 8 7 4 4 2

UPDATE(C, i)

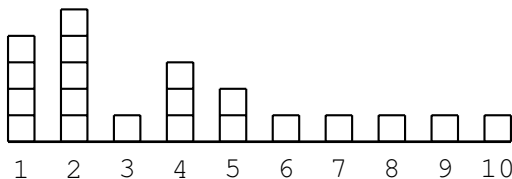
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10 1 5 1 5 2 2 3 -3 3 9 8 7 4 4 2 2

UPDATE(C, i)

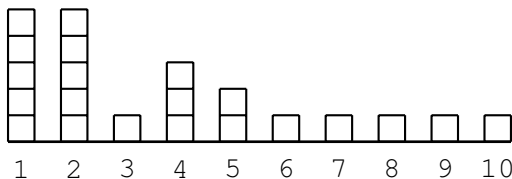
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10 1 5 1 5 2 2 3 -3 3 9 8 7 4 4 2 2 1

UPDATE(C, i)

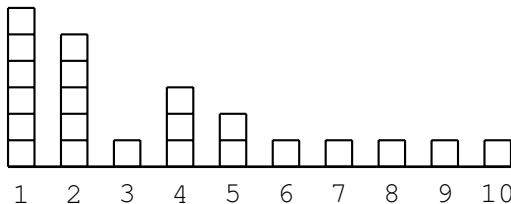
for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$



1 4 6 1 2 10 1 5 1 1 2 2 3 -3 3 9 8 7 4 4 2 2 1 5

UPDATE(C, i)

for $r \in [1 : t]$

$C[r, h_r(i)] \leftarrow C[r, h_r(i)] + s_r(i)$

end for

ESTIMATE(C, i)

return **median** $_r \{C[r, h_r(i)] \cdot s_r(i)\}$

Sketching: take (randomized) linear measurements of the input

$$\boxed{S} \cdot \boxed{x} = \boxed{b}$$

sketching matrix

space=number of rows

Easy to maintain sketch in dynamic streams
(insertions and deletions)

Sparse recovery

$$\boxed{S} \cdot \boxed{x} = \boxed{b}$$

sketching matrix
 $O(k \log n)$ rows

Let S be a COUNTSKETCH matrix with $O(k \log n)$ rows

Lemma

For every $x \in \mathbb{R}^n$ if $\hat{x} = \text{EST}(Sx)$, then whp

$$\|x - \hat{x}\|_{\infty} \leq \frac{1}{\sqrt{k}} \|x_{\text{TAIL}}\|_2.$$

(x_{TAIL} – x with largest k elements zeroed out)

Sparse recovery

Let S be a COUNTSKETCH matrix with $O(k \log n)$ rows

Lemma

For every $x \in \mathbb{R}^n$ if $\hat{x} = \text{EST}(Sx)$, then whp

$$\|x - \hat{x}\|_{\infty} \leq \frac{1}{\sqrt{k}} \|x_{\text{TAIL}}\|_2.$$

(x_{TAIL} – x with largest k elements zeroed out)

Sparse recovery

Let $\mathbf{S}$ be a COUNTSKETCH matrix with $O(k \log n)$ rows

Lemma

For every $x \in \mathbb{R}^n$ if $\hat{x} = \text{EST}(\mathbf{S}x)$, then whp

$$\|x - \hat{x}\|_{\infty} \leq \frac{1}{\sqrt{k}} \|x_{\text{TAIL}}\|_2.$$

(x_{TAIL} – x with largest k elements zeroed out)

Observation 1: # of measurements is optimal for ℓ_{∞}/ℓ_2 guarantee above

(capacity of the Gaussian channel, i.e. Shannon-Hartley theorem, – see Do Ba, Indyk, Price, Woodruff'10)

Sparse recovery

Let S be a COUNTSKETCH matrix with $O(k \log n)$ rows

Lemma

For every $x \in \mathbb{R}^n$ if $\hat{x} = \text{EST}(Sx)$, then whp

$$\|x - \hat{x}\|_{\infty} \leq \frac{1}{\sqrt{k}} \|x_{\text{TAIL}}\|_2.$$

(x_{TAIL} – x with largest k elements zeroed out)

Observation 1: # of measurements is optimal for ℓ_{∞}/ℓ_2 guarantee above

(capacity of the Gaussian channel, i.e. Shannon-Hartley theorem, – see Do Ba, Indyk, Price, Woodruff'10)

Observation 2: ℓ_2/ℓ_2 sparse recovery guarantee follows:

$$\|x - \hat{x}\|_2 = O(1) \cdot \|x_{\text{TAIL}}\|_2.$$

Sparse recovery

Let S be a COUNTSKETCH matrix with $O(k \log n)$ rows

Lemma

For every $x \in \mathbb{R}^n$ if $\hat{x} = \text{EST}(Sx)$, then whp

$$\|x - \hat{x}\|_{\infty} \leq \frac{1}{\sqrt{k}} \|x_{\text{TAIL}}\|_2.$$

$(x_{\text{TAIL}} - x \text{ with largest } k \text{ elements zeroed out})$

Observation 1: # of measurements is optimal for ℓ_{∞}/ℓ_2 guarantee above

(capacity of the Gaussian channel, i.e. Shannon-Hartley theorem, – see Do Ba, Indyk, Price, Woodruff'10)

Sparse recovery

Let S be a COUNTSKETCH matrix with $O(k \log n)$ rows

Lemma

For every $x \in \mathbb{R}^n$ if $\hat{x} = \text{EST}(Sx)$, then whp

$$\|x - \hat{x}\|_{\infty} \leq \frac{1}{\sqrt{k}} \|x_{\text{TAIL}}\|_2.$$

$(x_{\text{TAIL}} - x \text{ with largest } k \text{ elements zeroed out})$

Observation 1: # of measurements is optimal for ℓ_{∞}/ℓ_2 guarantee above

(capacity of the Gaussian channel, i.e. Shannon-Hartley theorem, – see Do Ba, Indyk, Price, Woodruff'10)

Observation 2: ℓ_2/ℓ_2 sparse recovery guarantee follows:

$$\|x - \hat{x}\|_2 = O(1) \cdot \min_{k\text{-sparse } x'} \|x - x'\|_2.$$

Sparse recovery

Let $\mathbf{S}$ be a COUNTSKETCH matrix with $O(\frac{1}{\epsilon^2} k \log n)$ rows

Lemma

For every $x \in \mathbb{R}^n$ if $\hat{x} = \text{EST}(\mathbf{S}x)$, then whp

$$\|x - \hat{x}\|_{\infty} \leq \frac{1}{\sqrt{k}} \|x_{\text{TAIL}}\|_2.$$

(x_{TAIL} – x with largest k elements zeroed out)

Observation 1: # of measurements is optimal for ℓ_{∞}/ℓ_2 guarantee above

(capacity of the Gaussian channel, i.e. Shannon-Hartley theorem, – see Do Ba, Indyk, Price, Woodruff'10)

Observation 2: ℓ_2/ℓ_2 sparse recovery guarantee follows:

$$\|x - \hat{x}\|_2 = (1 + \epsilon) \cdot \min_{k\text{-sparse } x'} \|x - x'\|_2.$$

Sparse recovery

Let $\mathbf{S}$ be a COUNTSKETCH matrix with $O(k \log n)$ rows

Lemma

For every $x \in \mathbb{R}^n$ if $\hat{x} = \text{EST}(\mathbf{S}x)$, then whp

$$\|x - \hat{x}\|_{\infty} \leq \frac{1}{\sqrt{k}} \|x_{\text{TAIL}}\|_2.$$

$(x_{\text{TAIL}} - x \text{ with largest } k \text{ elements zeroed out})$

Observation 3: if $\|x\|_0 \leq k$, then $x_{\text{TAIL}} = 0$ and $\text{EST}(\mathbf{S}x) = x$ whp

Exact sparse recovery: a k -sparse vector can be recovered from $O(k)$ linear measurements

CountSketch for matrices?

Input: r parties hold vectors $x_1, \dots, x_r \in \mathbb{R}^n$

each party sends $O(B \log n)$ bits to coordinator

(assume shared randomness)

CountSketch for matrices?

Input: r parties hold vectors $x_1, \dots, x_r \in \mathbb{R}^n$

each party sends $O(B \log n)$ bits to coordinator

(assume shared randomness)

Output: find largest entries in $A = \sum_{i=1}^r x_i x_i^T$

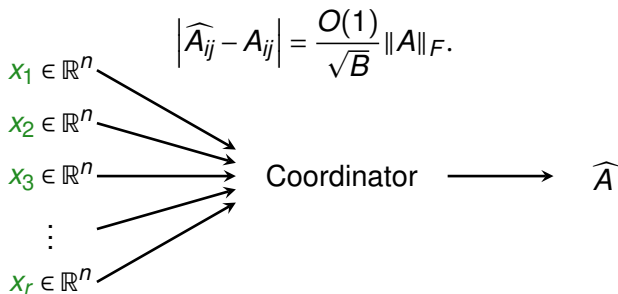
more precisely, output approximation $\widehat{A}$

$$|\widehat{A}_{ij} - A_{ij}| = \frac{O(1)}{\sqrt{B}} \|A\|_F.$$

CountSketch for matrices?

Input: r parties hold vectors $x_1, \dots, x_r \in \mathbb{R}^n$
each party sends $O(B \log n)$ bits to coordinator
(assume shared randomness)

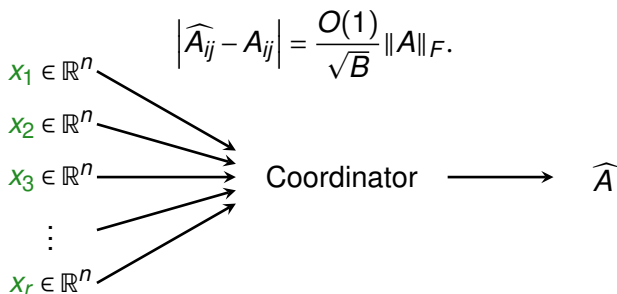
Output: find largest entries in $A = \sum_{i=1}^r x_i x_i^T$
more precisely, output approximation $\hat{A}$



CountSketch for matrices?

Input: r parties hold vectors $x_1, \dots, x_r \in \mathbb{R}^n$
each party sends $O(B \log n)$ bits to coordinator
(assume shared randomness)

Output: find largest entries in $A = \sum_{i=1}^r x_i x_i^T$
more precisely, output approximation $\hat{A}$



Every party i sends $\text{COUNTSKETCH}(x_i x_i^T)$ into $O(B)$ buckets
(slow)

Define

$$A = \sum_{i=1}^r x_i x_i^T \in \mathbb{R}^{n \times n}.$$

Define

$$A = \sum_{i=1}^r x_i x_i^T \in \mathbb{R}^{n \times n}.$$

Hash function

$$h: [n] \times [n] \rightarrow [B]$$

and random signs

$$s: [n] \times [n] \rightarrow \{-1, +1\}.$$

Define

$$A = \sum_{i=1}^r x_i x_i^T \in \mathbb{R}^{n \times n}.$$

Hash function

$$h: [n] \times [n] \rightarrow [B]$$

and random signs

$$s: [n] \times [n] \rightarrow \{-1, +1\}.$$

$$(Sx)_b = \sum_{i,j \in [n]: h(i,j)=b} s(i,j) \cdot x_i x_j.$$

Define

$$A = \sum_{i=1}^r x_i x_i^T \in \mathbb{R}^{n \times n}.$$

Hash function

$$h: [n] \times [n] \rightarrow [B]$$

and random signs

$$s: [n] \times [n] \rightarrow \{-1, +1\}.$$

$$(Sx)_b = \sum_{i,j \in [n]: h(i,j)=b} s(i,j) \cdot x_i x_j.$$

COUNTSKETCH($x_i x_i^T$) takes n^2 time to compute...

Define

$$A = \sum_{i=1}^r x_i x_i^T \in \mathbb{R}^{n \times n}.$$

Hash function

$$h: [n] \times [n] \rightarrow [B]$$

and random signs

$$s: [n] \times [n] \rightarrow \{-1, +1\}.$$

$$(Sx)_b = \sum_{i,j \in [n]: h(i,j)=b} s(i,j) \cdot x_i x_j.$$

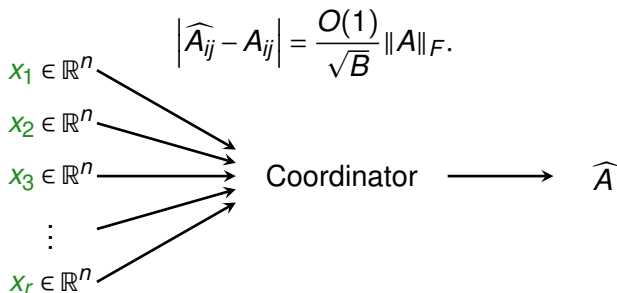
COUNTSKETCH($x_i x_i^T$) takes n^2 time to compute...

Make hash functions 'separable'?

CountSketch for matrices?

Input: r parties hold vectors $x_1, \dots, x_r \in \mathbb{R}^n$
each party sends $O(B \log n)$ bits to coordinator
(assume shared randomness)

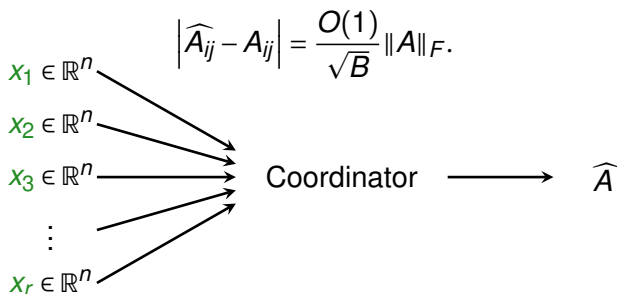
Output: find largest entries in $A = \sum_{i=1}^r x_i x_i^T$
more precisely, output approximation $\hat{A}$



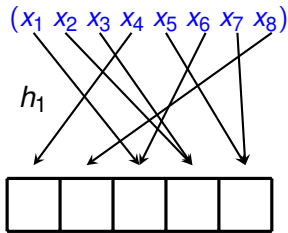
CountSketch for matrices?

Input: r parties hold vectors $x_1, \dots, x_r \in \mathbb{R}^n$
each party sends $O(B \log n)$ bits to coordinator
(assume shared randomness)

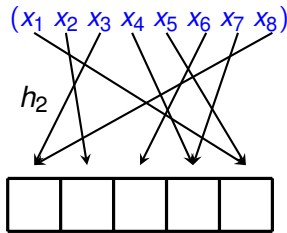
Output: find largest entries in $A = \sum_{i=1}^r x_i x_i^T$
more precisely, output approximation $\hat{A}$



Every party i sends $\text{SOMESKETCH}(x_i)$ into B buckets?
(fast)



$S_1 x$



$S_2 x$

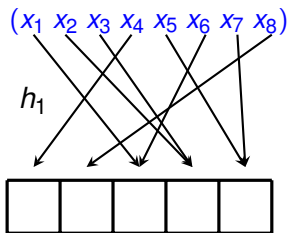
Take two **independent** instances of COUNTSKETCH: hash functions

$$h_1, h_2 : [n] \rightarrow [B],$$

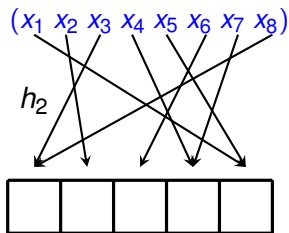
random signs

$$s_1, s_2 : [n] \rightarrow \{-1, +1\}$$

Tensor COUNTSKETCH₁ and COUNTSKETCH₂!



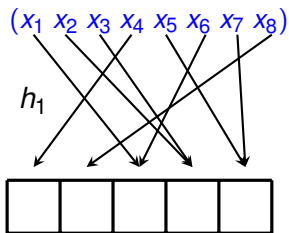
$S_1 x$



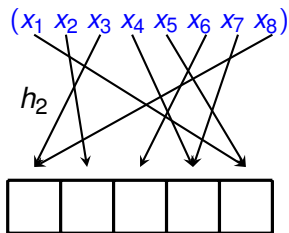
$S_2 x$

Define tensoring of COUNTSKETCH₁ and COUNTSKETCH₂:

$$h(i, j) = h_1(i) + h_2(j) \pmod{B}.$$



$S_1 x$

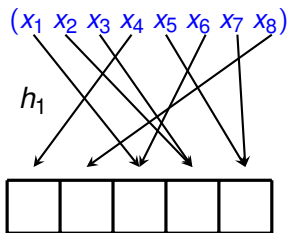


$S_2 x$

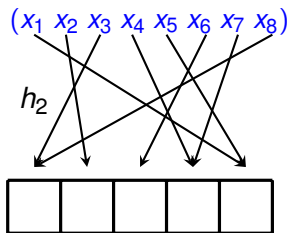
Define tensoring of COUNTSKETCH₁ and COUNTSKETCH₂:

$$h(i, j) = h_1(i) + h_2(j) \pmod{B}.$$

and $s(i, j) = s_1(i) \cdot s_2(j)$.



$S_1 x$



$S_2 x$

Define tensoring of COUNTSKETCH₁ and COUNTSKETCH₂:

$$h(i, j) = h_1(i) + h_2(j) \pmod{B}.$$

and $s(i, j) = s_1(i) \cdot s_2(j)$.

$$\begin{aligned} (Sx)_b &= \sum_{i, j \in [n]: h(i, j) = b} s(i, j) \cdot x_i x_j \\ &= \sum_{i, j \in [n]: h_1(i) + h_2(j) = b} s_1(i) \cdot s_2(j) \cdot x_i x_j \end{aligned}$$

$$(\textcolor{red}{S}x)_b = \sum_{i,j \in [n]: h_1(i) + h_2(j) = b} \textcolor{green}{s}_1(i) \cdot \textcolor{green}{s}_2(j) \cdot x_i x_j.$$

Can find $\textcolor{red}{S}x$ from $\text{COUNTSKETCH}_1(x)$ and $\text{COUNTSKETCH}_2(x)$
fast! ($\textcolor{red}{e}x\textcolor{red}{e}r\textcolor{red}{c}i\textcolor{red}{s}e$)

Stronger analysis of COUNTSKETCH

The bound

$$\|x - \hat{x}\|_{\infty} \leq \frac{1}{\sqrt{k}} \|x_{TAIL}\|_2$$

is optimal for sketches with $O(k \log n)$ rows, **for worst case x**

Stronger analysis of COUNTSKETCH

The bound

$$\|x - \hat{x}\|_{\infty} \leq \frac{1}{\sqrt{k}} \|x_{TAIL}\|_2$$

is optimal for sketches with $O(k \log n)$ rows, **for worst case x**

If x is drawn from a distribution (e.g., power law, Zipfian), one can do better by about $\log n$ factor: [Minton-Price'14](#)

Stronger analysis of COUNTSKETCH

The bound

$$\|x - \hat{x}\|_{\infty} \leq \frac{1}{\sqrt{k}} \|x_{\text{TAIL}}\|_2$$

is optimal for sketches with $O(k \log n)$ rows, **for worst case x**

If x is drawn from a distribution (e.g., power law, Zipfian), one can do better by about $\log n$ factor: [Minton-Price'14](#)

[Minton-Price'14](#) assumes uniformly random hashing. A very recent improvement:

Pseudorandom Hashing for Space-bounded Computation
with Applications in Streaming

Praneeth Kacham* Rasmus Pagh[†] Mikkel Thorup[‡] David P. Woodruff[§]

Abstract

We revisit Nisan's classical pseudorandom generator (PRG) for space-bounded computation (STOC 1990) and its applications in streaming algorithms. We describe a new generator, Hash-PRG, that can be thought of as a symmetric version of Nisan's generator over larger alphabets.

Non-asymptotic measurement complexity?

Good constants are achieved by ℓ_1 -minimization and related (non-sublinear) methods. Get best of both worlds?

Non-asymptotic measurement complexity?

Good constants are achieved by ℓ_1 -minimization and related (non-sublinear) methods. Get best of both worlds?

Ex: in LDPC codes, the Tanner graphs needs to be irregular to achieve capacity – similar effects here?

Non-asymptotic measurement complexity?

Good constants are achieved by ℓ_1 -minimization and related (non-sublinear) methods. Get best of both worlds?

Ex: in LDPC codes, the Tanner graphs needs to be irregular to achieve capacity – similar effects here?

Information Theoretic Limits of Cardinality Estimation: Fisher Meets Shannon*

Seth Pettie
pettie@umich.edu
University of Michigan
Computer Science and Engineering
Ann Arbor, MI, USA

Dingyu Wang
wangdy@umich.edu
University of Michigan
Computer Science and Engineering
Ann Arbor, MI, USA

ABSTRACT

Estimating the *cardinality* (number of distinct elements) of a large multiset is a classic problem in streaming and sketching, dating back to Flajolet and Martin's classic Probabilistic Counting (PCSA) algorithm from 1983.

In this paper we study the intrinsic tradeoff between the *space complexity* of the sketch and its *estimation error* in the RANDOM ORACLE model. We define a new measure of efficiency for cardinality estimators called the *Fisher-Shannon* (Fish) number $\mathcal{H}/I$. It captures the tension between the limiting Shannon entropy ($\mathcal{H}$) of the sketch and its normalized Fisher information (I), which characterizes the variance of a statistically efficient, asymptotically unbiased estimator.

KEYWORDS

cardinality estimation, streaming algorithm

ACM Reference Format:

Seth Pettie and Dingyu Wang. 2021. Information Theoretic Limits of Cardinality Estimation: Fisher Meets Shannon. In *Proceedings of the 53rd Annual ACM SIGACT Symposium on Theory of Computing (STOC '21)*, June 21–25, 2021, Virtual, Italy. ACM, New York, NY, USA, 14 pages. <https://doi.org/10.1145/3406325.3451032>

1 INTRODUCTION

Cardinality Estimation (aka *Distinct Elements* or F_0 -estimation) is a fundamental problem in streaming/sketching, with widespread industrial deployments in databases, networking, and sensing.

Non-asymptotic measurement complexity?

Good constants are achieved by ℓ_1 -minimization and related (non-sublinear) methods. Get best of both worlds?

Ex: in LDPC codes, the Tanner graphs needs to be irregular to achieve capacity – similar effects here?

Information Theoretic Limits of Cardinality Estimation: Fisher Meets Shannon*

Seth Pettie
pettie@umich.edu
University of Michigan
Computer Science and Engineering
Ann Arbor, MI, USA

Dingyu Wang
wangdy@umich.edu
University of Michigan
Computer Science and Engineering
Ann Arbor, MI, USA

ABSTRACT

Estimating the cardinality of a multiset is a classical problem. We go back to Flajolet and Martin's algorithm from 1989.

In this paper we study the complexity of the ORACLE model. We show that the cardinality estimators in this model captures the tensor norm of the sketch and characterizes the error of unbiased estimators.

Abstract

KEYWORDS

Peeling Close to the Orientability Threshold –
Spatial Coupling in Hashing-Based Data Structures

Stefan Walzer*

1 Introduction

Non-asymptotic measurement complexity?

Good constants are achieved by ℓ_1 -minimization and related (non-sublinear) methods. Get best of both worlds?

Ex: in LDPC codes, the Tanner graphs needs to be irregular to achieve capacity – similar effects here?

Information Theoretic Limits of Cardinality Estimation: Fisher Meets Shannon*

Seth Pettie
pettie@umich.edu
University of Michigan
Computer Science and Engineering
Ann Arbor, MI, USA

Dingyu Wang
wangdy@umich.edu
University of Michigan
Computer Science and Engineering
Ann Arbor, MI, USA

ABSTRACT

Estimating the cardinality of a multiset is a classical problem. We go back to Flajolet and Martin's algorithm from 1989.

In this paper we study the complexity of the ORACLE model. We analyze cardinality estimators that capture the variance of the sketch and characterize the unbiased estimator.

KEYWORDS

Peeling Close to the Orientability Threshold –
Spatial Coupling in Hashing-Based Data Structures

Stefan Walzer*

Simple Set Sketching

Abstract

Learning-augmented sketching: learn the hash function h in COUNTSKETCH (and more!) from data

Adversarially robust sketching: what if x is chosen by an adversary with (partial) knowledge of the data structure?

Learning-augmented sketching: learn the hash function h in COUNTSKETCH (and more!) from data

Adversarially robust sketching: what if x is chosen by an adversary with (partial) knowledge of the data structure?

ON THE ROBUSTNESS OF COUNTSKETCH TO ADAPTIVE INPUTS

Edith Cohen* Xin Lyu† Jelani Nelson‡ Tamás Sarlós§ Moshe Shechner¶ Uri Stemmer||

March 1, 2022

ABSTRACT

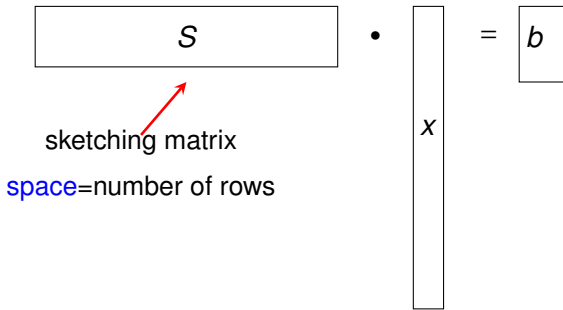
CountSketch is a popular dimensionality reduction technique that maps vectors to a lower dimension using randomized linear measurements. The sketch supports recovering ℓ_2 -heavy hitters of a vector (entries with $v[i]^2 \geq \frac{1}{k} \|v\|_2^2$). We study the robustness of the sketch in *adaptive* settings where input vectors may depend on the output from prior inputs. Adaptive settings arise in processes with feedback or with adversarial attacks. We show that the classic estimator is not robust, and can be attacked with a number of queries of the order of the sketch size. We propose a robust estimator (for a slightly modified sketch) that allows for quadratic number of queries in the sketch size, which is an improvement factor of $\sqrt{k}$ (for k heavy hitters) over prior work.

Take (randomized) linear measurements of the input

$$\boxed{S} \cdot \boxed{x} = \boxed{b}$$

sketching matrix

space=number of rows



Distribution of the sketching matrix?

Distribution of the sketching matrix?

+1	0	-1	0	0	0	0	0	0	0	+1	0
0	0	0	0	+1	0	0	0	0	-1	0	0
0	-1	0	0	0	0	0	+1	0	0	0	0
0	0	0	0	0	-1	0	0	+1	0	0	-1
0	0	0	+1	0	0	+1	0	0	0	0	0

Bernoulli(± 1) or Gaussian linear measurements on random subsets of the universe

The nonzeros are specified by the hash function $h: [n] \rightarrow [B]$

Can compute Sx in time $\text{nnz}(x)$!

Distribution of the sketching matrix?

+1	0	-1	0	0	0	0	0	0	0	+1	0
0	0	0	0	+1	0	0	0	0	-1	0	0
0	-1	0	0	0	0	0	+1	0	0	0	0
0	0	0	0	0	-1	0	0	+1	0	0	-1
0	0	0	+1	0	0	+1	0	0	0	0	0

Bernoulli(± 1) or Gaussian linear measurements on random subsets of the universe

The nonzeros are specified by the hash function $h: [n] \rightarrow [B]$

Can compute Sx in time $\text{nnz}(x)!$

Distribution of the sketching matrix?

+1	0	-1	0	0	0	0	0	0	0	+1	0
0	0	0	0	+1	0	0	0	0	-1	0	0
0	-1	0	0	0	0	0	+1	0	0	0	0
0	0	0	0	0	-1	0	0	+1	0	0	-1
0	0	0	+1	0	0	+1	0	0	0	0	0

Bernoulli(± 1) or Gaussian linear measurements on random subsets of the universe

The nonzeros are specified by the hash function $h: [n] \rightarrow [B]$

Can compute Sx in time $\text{nnz}(x)!$

Distribution of the sketching matrix?

+1	0	-1	0	0	0	0	0	0	0	+1	0
0	0	0	0	+1	0	0	0	0	-1	0	0
0	-1	0	0	0	0	0	+1	0	0	0	0
0	0	0	0	0	-1	0	0	+1	0	0	-1
0	0	0	+1	0	0	+1	0	0	0	0	0

Bernoulli(± 1) or Gaussian linear measurements on random subsets of the universe

The nonzeros are specified by the hash function $h: [n] \rightarrow [B]$

Can compute Sx in time $\text{nnz}(x)$!

Distribution of the sketching matrix?

+1	0	-1	0	0	0	0	0	0	0	+1	0
0	0	0	0	+1	0	0	0	0	-1	0	0
0	-1	0	0	0	0	0	+1	0	0	0	0
0	0	0	0	0	-1	0	0	+1	0	0	-1
0	0	0	+1	0	0	+1	0	0	0	0	0

Bernoulli(± 1) or Gaussian linear measurements on random subsets of the universe

The nonzeros are specified by the hash function $h: [n] \rightarrow [B]$

Can compute Sx in time $\text{nnz}(x)$!

Distribution of the sketching matrix?

+1	0	-1	0	0	0	0	0	0	0	+1	0
0	0	0	0	+1	0	0	0	0	-1	0	0
0	-1	0	0	0	0	0	+1	0	0	0	0
0	0	0	0	0	-1	0	0	+1	0	0	-1
0	0	0	+1	0	0	+1	0	0	0	0	0

Bernoulli(± 1) or Gaussian linear measurements on random subsets of the universe

The nonzeros are specified by the hash function $h: [n] \rightarrow [B]$

Can compute Sx in time $\text{nnz}(x)$!

Random restrictions (hashing)

What can we learn from Sx , where S is just random restrictions?

+1	0	0	0	0	0	0	0	0	0	+1	0
0	+1	+1	0	+1	0	0	0	0	+1	0	0
+1	0	0	+1	+1	+1	0	+1	0	0	0	0
+1	+1	0	0	+1	+1	0	0	+1	0	0	+1
0	+1	0	+1	+1	+1	+1	0	+1	0	+1	0

Can learn $\|x\|_0$, i.e. number of nonzeros in x

Johnson-Lindenstrauss transform

Sketching matrix S = a row of i.i.d. Gaussians of unit variance

Johnson-Lindenstrauss transform

Sketching matrix S = a row of i.i.d. Gaussians of unit variance

Measures ℓ_2^2 norm of x in expectation:

$$\mathbb{E} \left[\|Sx\|_2^2 \right] = \|x\|_2^2$$

Johnson-Lindenstrauss transform

Sketching matrix $S = \frac{1}{\sqrt{m}}$ rows of i.i.d. Gaussians of unit variance

Johnson-Lindenstrauss transform

Sketching matrix $S = \frac{1}{\sqrt{m}}$ rows of i.i.d. Gaussians of unit variance

Measures ℓ_2^2 norm of x with high probability:

$$\mathbb{P} \left[\left| \|Sx\|_2^2 - \|x\|_2^2 \right| \geq \epsilon^2 m \right] = 1 - \exp(-\Omega(\epsilon^2 m))$$

Johnson-Lindenstrauss transform

Sketching matrix $S = \frac{1}{\sqrt{m}}$ rows of i.i.d. Gaussians of unit variance

Measures ℓ_2^2 norm of x with high probability:

$$\mathbb{P} \left[\left| \|Sx\|_2^2 - \|x\|_2^2 \right| \geq \epsilon^2 \|x\|_2^2 \right] \leq \exp(-\Omega(\epsilon^2 m))$$

Downside: Sx takes $m \cdot n$ time to compute

(Faster) Johnson-Lindenstrauss transform

Subsampled randomized Hadamard transform

(Faster) Johnson-Lindenstrauss transform

Subsampled randomized Hadamard transform

Sketching matrix $S = P \cdot H \cdot D$

(Faster) Johnson-Lindenstrauss transform

Subsampled randomized Hadamard transform

$$\text{Sketching matrix } \mathbf{S} = \mathbf{P} \cdot \mathbf{H} \cdot \mathbf{D}$$

$\mathbf{D}$ =diagonal random sign matrix, $\mathbf{H}$ =Hadamard transform,
 $\mathbf{P}$ =random sampling matrix

(Faster) Johnson-Lindenstrauss transform

Subsampled randomized Hadamard transform

$$\text{Sketching matrix } \mathbf{S} = \mathbf{P} \cdot \mathbf{H} \cdot \mathbf{D}$$

$\mathbf{D}$ =diagonal random sign matrix, $\mathbf{H}$ =Hadamard transform,
 $\mathbf{P}$ =random sampling matrix

$\mathbf{S}x$ can be computed in $O(m + n \log n)$ time

Frequency moments

The p -th frequency moment

$$F_p = \sum_{i \in [n]} f_i^p$$

Frequency moments

The p -th frequency moment

$$F_p = \sum_{i \in [n]} f_i^p$$

Theorem

*Can approximate F_p for all $p \in [0, 2]$ in polylogarithmic space,
but need $\Omega(n^{1-2/p})$ space for all $p > 2$*

Frequency moments

The p -th frequency moment

$$F_p = \sum_{i \in [n]} f_i^p$$

Theorem

*Can approximate F_p for all $p \in [0, 2]$ in polylogarithmic space,
but need $\Omega(n^{1-2/p})$ space for all $p > 2$*

Bar-Yossef et al. Information complexity approach to data
stream lower bounds