

Private Data Stream Analysis for Universal Symmetric Norm Estimation

Vladimir Braverman ✉

Rice University, Houston, TX, USA

Joel Manning ✉

Carnegie Mellon University, Pittsburgh, PA, USA

Zhiwei Steven Wu ✉

Carnegie Mellon University, Pittsburgh, PA, USA

Samson Zhou ✉

University of California Berkeley, CA, USA

Rice University, Houston, TX, USA

Abstract

We study how to release summary statistics on a data stream subject to the constraint of differential privacy. In particular, we focus on releasing the family of *symmetric norms*, which are invariant under sign-flips and coordinate-wise permutations on an input data stream and include L_p norms, k -support norms, top- k norms, and the box norm as special cases. Although it may be possible to design and analyze a separate mechanism for each symmetric norm, we propose a general parametrizable framework that differentially privately releases a number of sufficient statistics from which the approximation of all symmetric norms can be simultaneously computed. Our framework partitions the coordinates of the underlying frequency vector into different levels based on their magnitude and releases approximate frequencies for the “heavy” coordinates in important levels and releases approximate level sizes for the “light” coordinates in important levels. Surprisingly, our mechanism allows for the release of an *arbitrary* number of symmetric norm approximations without any overhead or additional loss in privacy. Moreover, our mechanism permits $(1 + \alpha)$ -approximation to each of the symmetric norms and can be implemented using sublinear space in the streaming model for many regimes of the accuracy and privacy parameters.

2012 ACM Subject Classification Security and privacy → Usability in security and privacy

Keywords and phrases Differential privacy, norm estimation

Digital Object Identifier 10.4230/LIPIcs.APPROX/RANDOM.2023.45

Category RANDOM

Related Version *Full Version*: <https://arxiv.org/pdf/2307.04249.pdf>

1 Introduction

The family of L_p norms represent important statistics on an underlying dataset, where the L_p norm¹ of an n -dimensional frequency vector x is defined as the number of nonzero coordinates of x for $p = 0$ and $L_p(x) = (x_1^p + \dots + x_n^p)^{1/p}$ for $p > 0$. Thus, the L_0 norm counts the number of distinct elements in the dataset and, e.g., is used to detect denial of service or port scan attacks in network monitoring [3, 32], to understand the magnitude of quantities such as search engine queries or internet graph connectivity in data mining [55], to manage workload in database design [33], and to select a minimum-cost query plan in

¹ L_p for $p \in (0, 1)$ does not satisfy the triangle inequality and therefore is not a norm, but is still well-defined/well-motivated and can be computed



query optimization [57]. The L_1 norm computes the total number of elements in the dataset and, e.g., is used for data mining [26] and hypothesis testing [39], while the L_2 norm, e.g., is used for training random forests in machine learning [20], computing the Gini index in statistics [50, 36], and network anomaly detection in traffic monitoring [44, 62], in particular in the context of heavy-hitters, e.g., [24, 16, 15, 17, 49, 14]. More generally, L_p norms for $p \in (0, 2)$ have been used for entropy estimation [37]. Consequently, L_p estimation has been extensively studied in the data stream model [4, 40, 38, 45, 41, 5, 18, 35, 65, 66]. The simplest streaming model is perhaps the insertion-only model, in which a sequence of m updates increments coordinates of an n -dimensional frequency vector x and the goal is to compute or approximate some statistic of x in space that is sublinear in both m and n . For a more formal introduction to the streaming model, see Section 2.1.

In many cases, the underlying dataset contains sensitive information that should not be leaked. Hence, an active line of work has focused on estimating L_p norms for various values of p , while preserving differential privacy [53, 12, 59, 21, 63].

► **Definition 1** (Differential privacy, [29]). *Given $\varepsilon > 0$ and $\delta \in (0, 1)$, a randomized algorithm $\mathcal{A} : \mathcal{U}^* \rightarrow \mathcal{Y}$ is (ε, δ) -differentially private if, for every neighboring streams $\mathfrak{S}$ and $\mathfrak{S}'$ and for all $E \subseteq \mathcal{Y}$,*

$$\Pr[\mathcal{A}(\mathfrak{S}) \in E] \leq e^\varepsilon \cdot \Pr[\mathcal{A}(\mathfrak{S}') \in E] + \delta.$$

For example, [12] showed that the Johnson-Lindenstrauss transformation preserves differential privacy (DP), thereby showing one of the main techniques in the streaming model for L_2 estimation already guarantees DP. Similarly, [59] showed that the Flajolet-Martin sketch, which is one of the main approaches for L_0 estimation in the streaming model, also preserves DP. However, algorithmic designs for L_p estimation in the streaming model differ greatly and require individual analysis to ensure DP, especially because it is known that for some problems, guaranteeing DP provably requires more space [28]. Unfortunately, the privacy and utility analysis can be quite difficult due to the complexity of the various techniques. This is especially pronounced in the work of [63], who studied the p -stable sketch [38], which estimates the L_p norm for $p \in (0, 2]$. [63] showed that for $p \in (0, 1]$, the p -stable sketch preserves DP, but was unable to show DP for $p \in (1, 2]$, even though the general algorithmic approach remains the same. Thus the natural question is whether differential privacy can be guaranteed for an approach that simultaneously estimates the L_p norm in the streaming model, for all p . More generally, the family of L_p norms are all symmetric norms, which are invariant under sign-flips and coordinate-wise permutations on an input data stream. Symmetric norms thus also include other important families of norms such as the k -support norms and the top- k norms.

1.1 Our Contributions

In this paper, we show that not only does there exist a differentially private algorithm for the estimation of symmetric norms in the streaming model, but also that there exists an algorithm that privately releases a set of statistics, from which estimates of all (properly parametrized) symmetric norms can be simultaneously computed. To illustrate the difference, suppose we wanted to release approximations of the L_p norm of the stream for k different values of p . To guarantee (ε, δ) -DP for the set of k statistics, we would need, by advanced composition, to demand $\left(\mathcal{O}\left(\frac{\varepsilon}{\sqrt{k}}\right), \mathcal{O}\left(\frac{\delta}{k}\right)\right)$ -DP from k instances of a single differentially private L_p -estimation algorithm, corresponding to the k different values of p . Due to accuracy-privacy tradeoffs, the quality of the estimation will degrade severely as k increases. For an

extreme example, consider when k is some large polynomial of n and m so that the added noise will also be polynomial in n and m , and then there is no utility at all – the private algorithm might as well just release 0 for all queries!

In contrast, our algorithm releases a single set C of private statistics. By post-processing, we can then estimate the L_p norms for k different values of p while only requiring (ε, δ) -DP from C . Hence, our algorithm can simultaneously handle any large number of estimations of symmetric norms without compromising the quality of approximation.

We first informally introduce the definition of the maximum modulus of concentration of a norm, which measures the worst-case ratio of the maximum value of a norm on the L_2 -unit sphere to the median value of a norm on the L_2 -unit sphere, where the median can be taken over any restriction of the coordinates. Intuitively, maximum modulus of concentration of a norm quantifies the complexity of computing a norm. For example, the L_1 norm is generally “easy” to compute and has maximum modulus of concentration $\mathcal{O}(\log n)$. See Definition 18 for a more formal definition. Then our main result can informally be stated as follows:

► **Theorem 2 (Informal).** *There exists a (ε, δ) -differentially private algorithm that outputs a set C , from which the $(1 + \alpha)$ -approximation to any norm, with maximum modulus of concentration at most M of a vector $x \in \mathbb{R}^n$ induced by a stream of length $\text{poly}(n)$ can be computed, with probability at least $1 - \delta$. The algorithm uses $M^2 \cdot \text{poly}(\frac{1}{\alpha}, \frac{1}{\varepsilon}, \log n, \log \frac{1}{\delta})$ bits of space.*

We remark that as is standard in differential privacy on data streams, both the privacy parameter ε and the accuracy parameter α cannot be too small or the additive noise will be too large and cannot be absorbed into the $(1 + \alpha)$ -multiplicative bounds. See Theorem 33 for the formal statement of Theorem 2 describing these bounds.

We also remark that in the statement of Theorem 2, the δ failure parameter of approximate DP is equal to the failure parameter δ of the utility guarantees of the algorithm. More generally, if the desired failure probability δ' of the utility guarantee is not equal to the privacy parameter δ , then the dependencies will change from $\log \frac{1}{\delta}$ to $\log \frac{1}{\delta\delta'}$.

We emphasize that prior to our work, there is no algorithm that can handle private symmetric norm estimation for arbitrary symmetric norms, much less simultaneously for all parametrized symmetric norms. Although there is specific analysis for various norm estimation algorithms, e.g., see the discussion on related work in Section 1.3, these algorithms require a specific predetermined norm for their input. Thus a separate private algorithm must be run for each estimation, which increases the overall space. Moreover, for a large number of queries, the privacy parameter will need to be much smaller due to the composition of privacy, and thus to ensure privacy, the utility of each algorithm is provably poor. Our algorithm sidesteps both the space and accuracy problems and is the first and only work to do so, as of yet.

Applications. We briefly describe a number of specific symmetric norms that are handled by Theorem 2 and commonly used across various applications in machine learning. We first note the following parameterization of the previously discussed L_p norms.

► **Lemma 3** ([52, 43]). *For L_p norms, we have that $\text{mmc}(L) = \mathcal{O}(\log n)$ for $p \in [1, 2]$ and $\text{mmc}(L) = \mathcal{O}(n^{1/2-1/p})$ for $p > 2$.*

Thus our algorithm immediately introduces a differentially private mechanism for the approximation of L_p norms that unlike previous work, e.g., [12, 58, 25, 59, 21, 63], does not need to provide separate analysis for specific values of p . Moreover for constant-factor approximation, the space complexity is tight with the *optimal* L_p -approximation algorithms that do not consider privacy, up to polylogarithmic factors [42, 46, 34, 65] in the universe size n .

► **Definition 4** (*Q-norm and Q'-norm*). We call a norm L a *Q-norm* if there exists a symmetric norm L' such that $L(x) = L'(x^2)^{1/2}$ for all $x \in \mathbb{R}^n$. Here, we use x^2 to denote the coordinate-wise square power of x . We also call a norm L' a *Q'-norm* if its dual norm is a *Q-norm*.

The family of *Q'-norms* includes the L_p norms for $1 \leq p \leq 2$, the k -support norm, and the box norm [10] and thus *Q'-norms* have been proposed to regularize sparse recovery problems in machine learning. For instance, [7] showed that *Q'* norms have tighter relaxations than elastic nets and can thus be more effective for sparse prediction. Similarly, [51] used *Q'* norms to optimize sparse prediction algorithms for multitask clustering.

► **Lemma 5** ([11]). $\text{mmc}(L) = \mathcal{O}(\log n)$ for every *Q'-norm* L .

Theorem 2 and Lemma 5 thus present a differentially private algorithm for *Q'-norm* approximation that uses polylogarithmic space.

► **Definition 6** (*Top- k norm*). The *top- k norm* for a vector $x \in \mathbb{R}^n$ is the sum of the largest k coordinates of $|x|$, where we use $|x|$ to denote the vector whose entries are the coordinate-wise absolute value of x .

The *top- k norm* is frequently used to understand the more general Ky Fan k -norm [67], which is used to regularize optimization problems in numerical linear algebra. Whereas the Ky Fan k norm is defined as the sum of the k largest singular values of a matrix, the *top- k norm* is equivalent to the Ky Fan k norm when the input vector x represents the vector of the singular values of the matrix.

► **Lemma 7** ([11]). $\text{mmc}(L) = \tilde{\mathcal{O}}\left(\sqrt{\frac{n}{k}}\right)$ for the *top- k norm* L .

In particular, the *top- k norm* for a vector of singular values when $k = n$ is equivalent to the Schatten-1 norm of a matrix, which is a common metric for matrix fitting problems such as low-rank approximation [47].

► **Definition 8** (*Shannon entropy*). For a frequency vector $v \in \mathbb{R}^n$, we define the *Shannon entropy* by $H(v) = -\sum_{i=1}^n v_i \log v_i$.

To achieve an additive approximation to the Shannon entropy, we instead compute a multiplicative approximation to the exponential form, as follows:

► **Observation 9**. A $(1 + \alpha)$ -multiplicative approximation of the function $h(v) := 2^{H(v)}$ corresponds to an α -additive approximation of the *Shannon Entropy* $H(v)$ (and vice versa).

Moreover, computing a $(1 + \alpha)$ -approximation to $2^{H(v)}$ can be achieved through computing a $(1 + \alpha)$ -approximation to various L_p norms for $p \in (0, 2)$.

► **Lemma 10** (Section 3.3 in [37]). Let $k = \log \frac{1}{\alpha} + \log \log m$ and $\alpha' = \frac{\alpha}{12(k+1)^3 \log m}$. There exists an explicit set $\{y_0, \dots, y_k\}$ with $y_i \in (0, 2)$ for all i and a post-processing function that takes $(1 + \alpha')$ -approximations to $F_{y_i}(x)$, i.e., the (y_i) -th frequency moment of x , and outputs a $(1 + \alpha)$ -approximation to $h(v) = 2^{H(x)}$. Furthermore, the set $\{y_0, \dots, y_k\}$ and post-processing function are both efficiently computable, i.e., polynomial runtime.

Since our mechanism releases a private set of statistics from which $(1 + \alpha)$ -approximations to L_p norms can be computed for any $p \in (0, 2)$, then our mechanism also privately achieves an additive α -approximation to Shannon entropy.

1.2 Algorithmic Intuition and Overview

Our starting point is the L_p estimation algorithm of [40], which was parametrized by [11] to handle symmetric norms. For a $(1 + \alpha)$ -approximation, the algorithm partitions the n coordinates of the frequency vector x into powers of ξ -based on their magnitudes, where $\xi > 1$ is a fixed function of α . Each partition forms a level set, so that the i -th level set consists of the coordinates of x with frequency $[\xi^i, \xi^{i+1})$, but [40, 11] showed that it suffices to accurately count the size of each *important* level set and zero out to the other level sets, where a level set is considered important if its size is large enough to contribute an $\frac{\alpha^2}{\log m}$ fraction of the symmetric norm. In other words, if $\tilde{x}$ is a vector whose coordinates match those of x in important levels sets and are 0 elsewhere, then $(1 - \alpha)L(x) \leq L(\tilde{x}) \leq (1 + \alpha)L(x)$. We formalize the definition of importance in Section 2.2.

Private symmetric norm estimation in the centralized setting. To preserve (ε, δ) -differential privacy, one initial approach would be to view the frequency vector as a histogram and add Laplacian noise with scale $\mathcal{O}(\frac{1}{\varepsilon})$ to the frequency of each element. However, the level sets consisting of elements with frequencies between $[\xi^i, \xi^{i+1})$ for small i , say $i = 0$, could be largely perturbed by such Laplacian noise. For example, it is possible that for some coordinate j in an important level set, we have $x_j = 1$, in which case adding Laplacian noise with scale $\mathcal{O}(\frac{1}{\varepsilon})$ to x_j will heavily distort the coordinate. This can happen to all coordinates in the important level set, which results in an inaccurate estimation of the norm.

Fortunately, if i is small, the corresponding level set must contain a large number of elements if it is important, so it seems possible to privately release the size Γ_i of the level set. Indeed, we can show that the L_1 sensitivity of the vector corresponding to level set sizes is small and so we can add Laplacian noise with scale $\mathcal{O}(\frac{1}{\varepsilon})$ to each level set size. Hence if the level set has size Γ_i roughly $\Omega(\frac{1}{\alpha\varepsilon})$, then the Laplacian noise will affect Γ_i by a $(1 + \alpha)$ -factor.

Unfortunately, there can be level sets that are both important and small in size. For example, if there is a single element with frequency m , then the size of the corresponding level set is just one. Then adding Laplacian noise with scale $\mathcal{O}(\frac{1}{\varepsilon})$ will severely affect the size of the level set and thus the estimation of the symmetric norm. On the other hand, for $m > \frac{1}{\alpha\varepsilon}$, the frequency of the coordinate is quite large so again it seems like we can just add Laplacian noise with scale $\mathcal{O}(\frac{1}{\varepsilon})$ and output the noisy frequency of the coordinate.

New approach: classifying and separately handling high, medium, and low frequency levels. The main takeaway from these challenges is that we should handle different level sets separately. For the level sets of small coordinates, the important level sets must have large size and thus we would like to release noisy sizes. For the important level sets of large coordinates, we would like to release noisy frequencies of the coordinates.

In that vein, we partition the levels into three groups after defining thresholds T_1 and T_2 , with $T_1 > T_2$. We define the “high frequency levels” as the levels whose coordinates exceed T_1 in frequency. The intuition is that because the high frequency levels have such large magnitude, their frequencies can be well-approximated by running an L_2 -heavy hitters algorithm on the stream S .

We define the “medium frequency levels” as the levels whose coordinates are between T_1 and T_2 in frequency. These coordinates are not large enough to be detected by running an L_2 -heavy hitters algorithm on the stream S . However, the sizes of these level sets must be large if the level set is important. Thus there exists a substream S_j for which a large number of these coordinates are subsampled and their frequencies can be well-approximated by running an L_2 -heavy hitters algorithm on the substream S_j .

Finally, we define the “low frequency levels” as the levels whose coordinates are less than T_2 in frequency. These coordinates are small enough that we cannot add Laplacian noise to their frequencies without affecting the level sets they are mapped to. Instead, we show that the L_1 sensitivity for the level set estimations is particularly small for the low frequency levels. Thus, for these frequency levels, we report the size of the frequency levels rather than the approximate frequencies of the heavy-hitters. We remark that if our goal was to just approximate the symmetric norms without preserving differential privacy, then it would suffice to just consider the high and medium frequency levels, since the low frequency levels are particularly problematic when Laplacian noise is added to the frequency vector. We also remark that we only use the thresholds T_1 and T_2 for the purposes of describing our algorithm – in the actual implementation of the algorithm, the thresholds T_1 and T_2 will be implicitly defined by each of the substreams.

Private symmetric norm estimation in the streaming model. Although the previously discussed intuition builds towards a working algorithm, the main caveat is that so far, we have mainly discussed the centralized model, where space is not restricted and so each coordinate and thus each level set size can be counted exactly. In the streaming model, we cannot explicitly track the frequency vector, or even the frequencies of a constant fraction of coordinates. Instead, to estimate the sizes of each level set, [40, 11] take the stream S and form $s = \mathcal{O}(\log n)$ substreams $S_1, \dots, S_s$, where the j -th substream is created by sampling the universe of size n at a rate of $\frac{1}{2^{j-1}}$. Then S_j will only consist of the stream updates to the particular coordinates of x that are sampled. Thus in expectation, the frequency vector induced by S_j will have sparsity $\frac{\|x\|_0}{2^{j-1}}$. Similarly, if a level set i has size Γ_i , then $\frac{\Gamma_i}{2^{j-1}}$ of its members will be sampled in S_j in expectation. It can then be shown through a variance argument that if level set i is important, then there exists an explicit substream j from which Γ_i can be well-approximated using the L_2 -heavy hitter algorithm COUNTSKETCH and as a result, the symmetric norm of x can be well-approximated. The main point of the subsampling approach is that if there exists a level set with large size consisting of small coordinates, then the coordinates will not be detected by the COUNTSKETCH on S , but because S_j has significantly smaller L_2 norm, then the coordinates will be detected by COUNTSKETCH on S_j .

However, adapting the subsampling and heavy-hitter approach introduces additional challenges for privacy. For instance, we can analyze the L_2 -heavy hitter algorithm COUNTSKETCH and show that although the L_1 sensitivity of the estimated frequency for a single coordinate is small, the L_1 sensitivity of the estimated frequency vector for all the coordinates may be large. Instead, we use the view that COUNTSKETCH is a composition function that first only estimates frequencies for the top $\text{poly}(\frac{1}{\alpha}, \frac{1}{\epsilon}, \log n)$ and then outputs only those estimates that are above a certain threshold. Similarly, the Laplacian noise added to privately use COUNTSKETCH can alter the sizes of a significant number of level sets for small coordinates. Thus for the small coordinates (corresponding to the substreams S_j with large j), we invoke COUNTSKETCH with much higher accuracy, so that with high probability, it will return *exactly* the frequencies for the small coordinates. For example, note that if the frequency x_k of a coordinate $k \in [n]$ is at most $\frac{1}{2\alpha^2\epsilon}$, then any $(1 + \alpha^2\epsilon)$ -approximation to x_k can be rounded to exactly recover x_k . This decreases the L_1 sensitivity of the vector of estimated level set sizes, therefore allowing us to add Laplacian noise without greatly affecting the quality of approximation.

1.3 Related Work

Non-private L_p norm estimation is one of the fundamental problems in the streaming model, beginning with [4]’s seminal work that tracks the inner product of the frequency vector with a random sign vector for L_2 estimation (as well as a telescoping argument for integer $p > 0$). [38, 45] later showed that this approach could be generalized for $p \in (0, 2]$ by tracking the inner product of the frequency vector with a vector with randomly generated p -stable variables, which only exist for $p \in (0, 2]$. For $p > 2$, [5] gave an L_p estimation algorithm using the max-stability property of exponential random variables. More generally, [40] introduced the framework of subsampling and using heavy-hitters for L_p estimation, which [11] parametrized to all symmetric norms. It should be emphasized that these techniques all handle the more general turnstile model, in which ± 1 updates are allowed to each coordinate, rather than single positive increments. Hence our techniques also extend to the turnstile model with a minor change on the conditions.

More recently, [13, 61] given a general framework for converting non-private approximation algorithms into private approximation algorithms, provided that the accuracy of these algorithms could be tuned with an input parameter $\varepsilon > 0$, i.e., the algorithms can achieve $(1 + \varepsilon)$ -approximation for a wide range of $\varepsilon > 0$. Their results presented a solution that addresses the difficulty of adapting privacy specifically to each non-private algorithm separately. However, their framework only applies to problems with scalar outputs and thus do not handle synthetic data release. Therefore, privately answering multiple norm queries while circumventing composition bounds is still a challenge that their results cannot handle.

Symmetric norms have also recently received attention in other big data models as well. [6] studied approximate near neighbors for general symmetric norms while [48] studied symmetric norm estimation for network monitoring. [60] considered Orlicz norm regression and other loss functions where the penalty is a symmetric norm. [19] gave an algorithm to approximate the symmetric norm in the sliding window model, where updates in the data stream implicitly expire after a fixed amount of time.

Specific cases of private L_p estimation in the streaming model have also been previously well-studied. [25, 59] studied private L_0 estimation using the Flajolet-Martin sketch, while [63] studied private L_p estimation for $p \in (0, 1]$ using the p -stable sketch and [12, 58, 25, 21] studied private L_2 estimation using the Johnson-Lindenstrauss projection. Specifically, [12] gave an (ε, δ) -DP algorithm for L_2 estimation that achieves a $(1 + \varepsilon)$ -approximation while using $\mathcal{O}(\frac{1}{\varepsilon^2} \log n \log \frac{1}{\delta})$ bits of space and [63] gave an (ε, δ) -DP algorithm for L_p estimation that achieves a $(1 + \alpha)$ -approximation while using $\mathcal{O}(\frac{1}{\alpha^2} \log n \log \frac{1}{\delta})$ bits of space for constant ε and $p \in (0, 1)$. For fractional $p > 1$, private distribution estimation algorithms [2, 68, 9, 64] can be used to approximate the L_p norm, but since the algorithms provide information over a much larger distribution, e.g., much larger histograms of frequencies, the privacy-accuracy trade-off is sub-optimal and the space complexity is exponentially worse.

The related problem of privately releasing heavy-hitters in big data models has also been well-studied. [23] studied the problem of continually releasing L_1 -heavy hitters in a stream, while [30] studied L_1 -heavy hitters and other problems in the pan-private streaming model. The heavy-hitter problem has also received significant attention in the local model, e.g., [9, 27, 1, 22, 8], where individual users should locally randomize their data before sending differentially private information to an untrusted server that aggregates the statistics across all users.

2 Preliminaries

In this section, we introduce definitions and simple or well-known results from differential privacy, sketching algorithms, and symmetric norms. For notation, we use $[n]$ for an integer $n > 0$ to denote the set $\{1, \dots, n\}$. We also use the notation $\text{poly}(n)$ to represent a constant degree polynomial in n and we say an event occurs *with high probability* if the event holds with probability $1 - \frac{1}{\text{poly}(n)}$. Similarly, we use $\text{polylog}(n)$ to denote $\text{poly}(\log n)$. Given a vector $x \in \mathbb{R}^n$, we define its second frequency moment $F_2(x) = x_1^2 + \dots + x_n^2$. Finally, for a parameter $c \geq 1$, we say that X provides a C -approximation to a quantity Y if $\frac{X}{C} \leq Y \leq C \cdot X$.

Privately releasing multiple statistics that are individually differentially private can also be done, but comes at a slight cost.

► **Theorem 11** (Composition and post-processing of differential privacy, [31]). *Let $\mathcal{A}_i : \mathcal{U}_i \rightarrow X_i$ be an $(\varepsilon_i, \delta_i)$ -differential private algorithm for $i \in [k]$. Then $\mathcal{A}_{[k]}(x) = (\mathcal{A}_1(x), \dots, \mathcal{A}_k(x))$ is $(\sum_{i=1}^k \varepsilon_i, \sum_{i=1}^k \delta_i)$ -differentially private. Furthermore, if $g_i : X_i \rightarrow X'_i$ is an arbitrary random mapping, then $g_i(\mathcal{M}_i(x))$ is $(\varepsilon_i, \delta_i)$ -differentially private.*

Although there exists more sophisticated approaches for composition, such as advanced composition, we do not need them for our purposes.

2.1 Streaming and Sketching Algorithms

In the streaming model, a frequency vector $x \in \mathbb{R}^n$ is induced by a sequence of updates. In the insertion-only streaming model, x is defined through a stream of m updates $u_1, \dots, u_m$, where $u_t \in [n]$ for each $t \in [m]$ so that $x_i = |\{t \in [m] \mid u_t = i\}|$ for all $i \in [n]$. In other words, x_i is the number of times that $i \in [n]$ appears in the stream. We remark that our techniques generalize to some degree to turnstile streams, where each update is an ordered pair $u_t = (\Delta_t, c_t)$, so that the t -th update changes the c_t -th coordinate by Δ_t , i.e., $c_t \in [n]$ is a coordinate and $\Delta_t \in [-M, M]$ for some parameter $M > 0$. In this turnstile model, the vector x is defined so that $x_i = \sum_{t: c_t = i} \Delta_t$ for all $i \in [n]$. Although our techniques can apply to the general turnstile model with a minor change on the conditions and assumptions, we shall work with the insertion-only streaming model throughout the remainder of the paper.

Given a frequency vector $x \in \mathbb{R}^n$ on a data stream, the AMS algorithm for L_2 -estimation first generates a sign vector $\sigma \in \{-1, +1\}^n$ and sets $S_1 = (\langle \sigma, x \rangle)^2$. We remark that to maintain σ in small space, it suffices for the coordinates of the sign vector σ to be 4-wise independent and therefore it suffices to randomly generate and store a 4-wise independent hash function. The AMS algorithm then repeats this process $b = \frac{6}{\alpha^2}$ independent times to obtain dot products $S_1, \dots, S_b$, sets Z^2 to be the arithmetic mean of $S_1, \dots, S_b$, and reports Z . We define the L_2 norm of a vector $x \in \mathbb{R}^n$ by $L_2(x) = \sqrt{x_1^2 + \dots + x_n^2}$.

► **Definition 12** (ν -approximate η L_2 -heavy hitters problem). *Given an accuracy parameter $\nu \in (0, 1)$, a threshold parameter η , and a frequency vector $x \in \mathbb{R}^n$, compute a set $H \subseteq [n]$ and a set of approximations $\widehat{x}_k$ for all $k \in H$ such that:*

- (1) *If $x_k \geq \eta L_2(x)$ for any $k \in [n]$, then $k \in H$, so that H contains all η L_2 -heavy hitters of x .*
- (2) *There exists a universal constant $C \in (0, 1)$ so that if $x_k \leq \frac{C\eta}{2} L_2(x)$ for any $k \in [n]$, then $k \notin H$, so that H does not contain any index that is not an $\frac{C\eta}{2}$ L_2 -heavy hitter of x .*
- (3) *If $k \in H$ for any $k \in [n]$, then compute $(1 \pm \nu)$ -approximation to the frequency x_k , i.e., a value $\widehat{x}_k$ such that $(1 - \nu)x_k \leq \widehat{x}_k \leq (1 + \nu)x_k$.*

The well-known COUNTSKETCH algorithm can be parametrized to provide an estimated frequency to each item and then releases the approximate frequencies of each item that surpasses a threshold proportional to the output of AMS:

► **Theorem 13** (CountSketch for ν -approximate η L_2 -heavy hitters, [24]). *There exists a one-pass streaming algorithm COUNTSKETCH that takes an accuracy parameter $\nu \in (0, 1)$ and a threshold parameter η^2 and outputs a list H that contains all indices $k \in [n]$ of an underlying frequency vector x with $x_k \geq \eta L_2(x)$ and no index $k \in [n]$ with $x_k \leq \eta(1 - \nu) L_2(x)$. For each $k \in H$, COUNTSKETCH also reports a estimated frequency $\widehat{x}_k$ such that $(1 - \nu)x_k \leq \widehat{x}_k \leq (1 + \nu)x_k$. The algorithm uses $\mathcal{O}\left(\frac{1}{\eta^2 \nu^2} \log^2 n\right)$ bits of space and succeeds with probability $1 - \frac{1}{\text{poly}(m)}$.*

■ **Algorithm 1** Heavy-hitter algorithm COUNTSKETCH.

Input: Stream $\mathfrak{S}$ inducing frequency vector $x \in \mathbb{R}^n$, accuracy parameter $\nu \in (0, 1)$, and threshold parameter $\eta \in (0, 1)$

Output: L_2 Heavy-hitter algorithm

```

1:  $r \leftarrow \mathcal{O}(\log n)$ ,  $b \leftarrow \mathcal{O}\left(\frac{1}{\eta^2 \nu^2}\right)$ 
2: Pick hash functions  $h^{(1)}, \dots, h^{(r)} : [n] \rightarrow [b]$  and  $s^{(1)}, \dots, s^{(r)} : [n] \rightarrow \{-1, +1\}$ 
3:  $S_{i,j} \leftarrow 0$  for  $(i, j) \in [r] \times [b]$ 
4: for each update  $u_i \in [n]$ ,  $i \in [m]$  do
5:   for each  $j \in [r]$  do
6:      $b_{i,j} \leftarrow h^{(j)}(u_i)$  and  $s_{i,j} \leftarrow s^{(j)}(u_i)$ 
7:      $S_{j,b_{i,j}} \leftarrow S_{j,b_{i,j}} + s_{i,j}$ 
8: for each  $i \in [n]$  do
9:    $b_{i,j} \leftarrow h^{(j)}(u_i)$  for each  $j \in [r]$ 
10: return  $\text{median}_{j \in [r]} |S_{j,b_{i,j}}|$  as the estimated frequency for  $x_i$ 

```

We recall the following sensitivity analysis of CountSketch.

► **Lemma 14** (Sensitivity of CountSketch). *Let $x, x' \in \mathbb{R}^n$ with $\max(\|x - x'\|_0, \|x - x'\|_1) \leq 2$.*

There exists a private variant PRIVCOUNTSKETCH of COUNTSKETCH that adds noise to each coordinate and then uses a standard private threshold routine to ensure differential privacy, giving the following guarantees:

► **Lemma 15.** *There exists a one-pass streaming algorithm PRIVCOUNTSKETCH that takes an accuracy parameter $\nu \in (0, 1)$ and a threshold parameter η^2 and outputs a list H that contains all indices $k \in [n]$ of an underlying frequency vector x with $x_k \geq \eta L_2(x)$ and no index $k \in [n]$ with $x_k \leq \eta(1 - \nu) L_2(x)$. For each $k \in H$, PRIVCOUNTSKETCH also reports a estimated frequency $\widehat{x}_k$ such that $(1 - \nu)x_k - \mathcal{O}\left(\frac{\log m}{\eta \nu}\right) \leq \widehat{x}_k \leq (1 + \nu)x_k + \mathcal{O}\left(\frac{\log m}{\eta \nu}\right)$. The algorithm uses $\mathcal{O}\left(\frac{1}{\eta^2 \nu^2} \log^2 n\right)$ bits of space and succeeds with probability $1 - \frac{1}{\text{poly}(m)}$.*

2.2 Symmetric Norms

In this section, we provide necessary preliminaries for symmetric norm estimation.

► **Definition 16** (Symmetric norm). *A function $L : \mathbb{R}^n \rightarrow \mathbb{R}$ is a symmetric norm if L is a norm and for all $x \in \mathbb{R}^n$ and any vector $y \in \mathbb{R}^n$ that is a permutation of the coordinates of x , we have $L(x) = L(y)$. Moreover, we have $L(x) = L(|x|)$, where $|x|$ is the coordinate-wise absolute value of x .*

► **Definition 17** (Modulus of concentration). Let $x \in \mathbb{R}^n$ be a random variable drawn from the uniform distribution on the L_2 -unit sphere S^{n-1} and let b_L denote the maximum value of $L(x)$ over S^{n-1} . The median of a symmetric norm L is the unique value M_L such that $\Pr[L(x) \geq M_L] \geq \frac{1}{2}$ and $\Pr[L(x) \leq M_L] \geq \frac{1}{2}$. Then the ratio $\text{mc}(L) := \frac{b_L}{M_L}$ is the modulus of concentration of the norm L .

Although the modulus of concentration quantifies the “average” behavior of the norm L on $\mathbb{R}^n$, norms with challenging behavior can still be embedded in lower-dimensional subspaces. For instance, the L_1 norm satisfies $\text{mc}(L) = \mathcal{O}(1)$, but when $x \in \mathbb{R}^n$ has fewer than $\sqrt{n}$ nonzero coordinates, the norm $\max(L_\infty(x), L_1(x)/\sqrt{n})$ on the unit ball becomes identically $L_\infty(x)$ [11], which requires $\Omega(\sqrt{n})$ space [4] to estimate. Hence, we further quantify the behavior of a norm L by examining its behavior on all lower dimensions.

► **Definition 18** (Maximum modulus of concentration). For a norm $L : \mathbb{R}^n \rightarrow \mathbb{R}$ and every $k \leq n$, define the norm $L^{(k)} : \mathbb{R}^k \rightarrow \mathbb{R}$ by $L^{(k)}((x_1, \dots, x_k)) := L((x_1, \dots, x_k, 0, \dots, 0))$. Then the maximum modulus of concentration of the norm L is $\text{mmc}(L) := \max_{k \leq n} \text{mc}(L^{(k)}) = \max_{k \leq n} \frac{b_{L^{(k)}}}{M_{L^{(k)}}}$.

► **Definition 19** (Important Levels). For $x \in \mathbb{R}^n$ and $\xi > 1$, we define the level i as the set $B_i = \{k \in [n] : \xi^{i-1} \leq |x_k| \leq \xi^i\}$. We define $b_i := |B_i|$ as the size of level i . For $\beta \in (0, 1]$, we say level i is β -important if

$$b_i > \beta \sum_{j>i} b_j, \quad b_i \xi^{2i} \geq \beta \sum_{j \leq i} b_j \xi^{2j}.$$

Informally, level i is β -important if (1) its size is at least a β -fraction of the total sizes of the higher levels and (2) its contribution is roughly a β -fraction of the total contribution of all the lower levels. We would like to show that to approximate a symmetric norm $L(x)$, it suffices to identify the β -important levels and their sizes for a fixed base $\xi > 1$.

► **Definition 20** (Level Vectors and Buckets). For $x \in \mathbb{R}^n$ and $\xi > 1$, the level vector for x is

$$V(x) := (\underbrace{\xi^1, \dots, \xi^1}_{b_1 \text{ times}}, \underbrace{\xi^2, \dots, \xi^2}_{b_2 \text{ times}}, \dots, \underbrace{\xi^k, \dots, \xi^k}_{b_k \text{ times}}, 0, \dots, 0) \in \mathbb{R}^n,$$

where each b_i is the size of level i . The i -th bucket of $V(x)$ is

$$V_i(x) := (\underbrace{0, \dots, 0}_{b_1 + \dots + b_{i-1} \text{ times}}, \underbrace{\xi^i, \dots, \xi^i}_{b_i \text{ times}}, \underbrace{0, \dots, 0}_{b_{i+1} + \dots + b_k \text{ times}}, 0, \dots, 0) \in \mathbb{R}^n.$$

We similarly define the approximate level vectors $\widehat{V}(x)$ and $\widehat{V}_i(x)$ using approximations $\widehat{b}_1, \dots, \widehat{b}_k$ for $b_1, \dots, b_k$. We write $V(x) \setminus V_i(x)$ to denote the vector that replaces the i -th bucket in $V(x)$ with all zeros and we write $V(x) \setminus V_i(x) \cup \widehat{V}_i(x)$ to denote the vector that replaces the i -th bucket in $V(x)$ with $\widehat{b}_i$ instances of ξ^i .

Rather than directly handle the important levels, we define the β -contributing levels and instead work toward estimating the contribution of the β -contributing levels.

► **Definition 21** (Contributing Levels). Given $x \in \mathbb{R}^n$, a level i defined by base $\xi > 1$ is β -contributing if $L(V_i(x)) \geq \beta L(V(x))$.

[11] showed that even if all levels that are not β -contributing are removed, the contribution of the remaining levels forms a good approximation to $L(x)$.

► **Lemma 22** ([11]). *Given $x \in \mathbb{R}^n$ and levels defined by a base $\xi > 1$, let $V'(x)$ be the vector obtained by removing all levels that are not β -contributing from $V(x)$. Then $(1 - \mathcal{O}(\log_\xi n) \cdot \beta)L(V(x)) \leq L(V'(x)) \leq L(V(x))$.*

Hence for appropriate $\xi > 1$ and $\beta \in (0, 1]$, it suffices to identify the β -contributing levels, zero out the remaining levels, and determine the contribution of the resulting vector to approximate the symmetric norm $L(x)$.

► **Lemma 23** ([11]). *Given an accuracy parameter $\alpha \in (0, 1]$, let base $\xi = (1 + \mathcal{O}(\alpha))$, importance parameter $\beta = \mathcal{O}\left(\frac{\alpha^5}{\text{mmc}(\ell)^2 \cdot \log^5 m}\right)$, and $\alpha' = \mathcal{O}\left(\frac{\alpha^2}{\log n}\right)$. Let $\hat{b}_i \leq b_i$ for all i and $\hat{b}_i \geq (1 - \alpha')b_i$ for all β -important levels. Let $\hat{V}$ be the level vector constructed using the estimates $\hat{b}_1, \hat{b}_2, \dots$ and let V' be the level vector constructed by removing all the buckets that are not β -contributing in $\hat{V}$. Then $(1 - \alpha)L(V(x)) \leq L(V'(x)) \leq L(V(x))$.*

To identify the β -contributing levels, [11] first notes that the size of the level must be at least a significant fraction of the total size of the higher levels.

► **Lemma 24** ([11]). *Given $x \in \mathbb{R}^n$, let the level sets be defined by a base $\xi > 1$. If level i is β -contributing, then there exists some fixed constant $\lambda > 0$ such that $b_i \geq \frac{\lambda \beta^2}{\text{mmc}(\ell)^2 \log^2 n} \cdot \sum_{j>i} b_j$.*

Moreover, [11] observes that the squared mass of a β -contributing level must be at least a significant fraction of the total squared mass of the lower levels.

► **Lemma 25** ([11]). *Given $x \in \mathbb{R}^n$, let the level sets be defined by a base $\xi > 1$. If level i is β -contributing, then there exists some fixed constant $\lambda > 0$ such that $b_i \xi^{2i} \geq \frac{\lambda \beta^2}{\text{mmc}(\ell)^2 (\log_\xi n) \log^2 n} \cdot \sum_{j \leq i} b_j \xi^{2j}$.*

Observe that together, Lemma 24 and Lemma 25 imply that a β -contributing level i must also be an important level as defined in Definition 19. Crucially, since Lemma 25 states that the squared mass (or the F_2 frequency moment) of the β -contributing levels must be a significant fraction of the total squared mass of the lower levels, then it suggests we might be able to identify the β -contributing levels through an L_2 -heavy hitters algorithm after removing the higher levels. Indeed, [11] show that the problem of identifying the size (and thus the contribution) of the β -contributing levels can be reduced to the task of finding ν -approximate η -heavy hitters for specific parameters of ν and η .

► **Lemma 26** ([11]). *Let $s = \mathcal{O}(\log n)$. If a level i is β -important, then either $\xi^{2i} \geq \frac{\alpha^2 \beta \varepsilon^2}{\log^2 m} F_2(x)$ or there exists $j \in [s]$ such that $b_i \geq \frac{2^j \log^2 m}{\alpha^2 \varepsilon^2}$ and $\xi^{2i} \in \left[\frac{\alpha^2 \beta \varepsilon^2}{\log^2 m} \cdot \frac{F_2(x)}{2^j}, \frac{\alpha^2 \beta \varepsilon^2}{\log^2 m} \cdot \frac{F_2(x)}{2^{j-1}} \right]$.*

Lemma 26 implies that if level i is β -important, then either (1) it will be identified by using PRIVCOUNTSKETCH, i.e., Lemma 15, with threshold $\frac{\alpha^2 \beta}{\log^2 m}$ on the stream or (2) its contribution can be well-approximated by using PRIVCOUNTSKETCH with threshold $\frac{\alpha^2 \beta \varepsilon^2}{\log^2 m}$ on a substream formed by sampling coordinates of the universe with probability $\frac{1}{2^j}$. We thus split our algorithm and analysis to handle these cases. In particular, we call a frequency level i “high” if $\xi^{2i} \geq \frac{\alpha^2 \beta \varepsilon^2}{\log^2 m} F_2(x)$. We call a frequency level i “medium” if $\xi^{2i} \geq \frac{\alpha^2 \beta' \varepsilon^2}{2^j} F_2(x) > T$ and $b_i \geq \mathcal{O}\left(\frac{2^j \log^2 m}{\alpha^2 \varepsilon^2}\right)$ for a certain $\beta' > 0$ and a threshold T . We call a frequency level i “low” if $\xi^{2i} \geq \frac{\alpha^2 \beta' \varepsilon^2}{2^j} F_2(x)$ and $b_i \geq \mathcal{O}\left(\frac{2^j \log^2 m}{\alpha^2 \varepsilon^2}\right)$, but $T \geq \frac{\alpha^2 \beta' \varepsilon^2}{2^j} F_2(x)$.

3 Private Symmetric Norm Estimation Algorithm

In this section, we give our algorithm that releases a set of private statistics from which an arbitrary number of symmetric norms can be well-approximated. In particular, recall that Lemma 23 suggests that it suffices to approximate the sizes of the important levels and identify the non-important levels, so that the contributions of the non-important levels can be set to zero. We partition the levels into three groups after defining explicit thresholds T_1 and T_2 , with $T_1 > T_2$. Recall that we define the “high frequency levels” as the levels whose coordinates exceed T_1 in frequency, the “medium frequency levels” as the levels whose coordinates are between T_1 and T_2 in frequency, and the “low frequency levels” as the levels whose coordinates are less than T_2 in frequency.

The intuition is that because the high frequency levels have such large magnitude, their frequencies can be well-approximated by running an L_2 -heavy hitters algorithm on the stream S . On the other hand, the medium frequency level coordinates are not large enough to be detected by running an L_2 -heavy hitters algorithm on the stream S , but the sizes of these level sets must be large if the level set is important and therefore, there exists a substream S_j for which a large number of these coordinates are subsampled and their frequencies can be well-approximated by running an L_2 -heavy hitters algorithm on the substream S_j . Here we form substreams $S_0, S_1, \dots$ so that S_j first samples elements of the universe $[n]$ at a rate $\frac{1}{2^j}$ and then only contains the stream updates that are relevant to the sampled elements. Finally, the low frequency level coordinates are small enough that we cannot add Laplacian noise to their frequencies without affecting the level sets they are mapped to. We instead show that L_1 sensitivity for the level set estimations is particularly small for the low frequency levels and thus, we report the size of the level sets of the low frequency levels rather than the approximate frequencies of the heavy-hitters.

We emphasize that we only use the thresholds T_1 and T_2 for the purposes of describing our algorithm – in the actual implementation of the algorithm, the thresholds T_1 and T_2 will be implicitly defined by each of the substreams. For example, the items with threshold larger than T_1 will automatically be revealed through the stream S , while the items with thresholds between T_1 and T_2 will be revealed through the substreams S_j with $2^j > \frac{\log n}{\beta' \alpha \varepsilon}$ for explicit parameters α , β' , and ε . More specifically, note that Algorithm 2 sets $\beta' = \mathcal{O}\left(\frac{\alpha^2 \beta \varepsilon^2}{\log^2 m}\right)$ or more specifically $\beta' = \frac{\alpha^2 \beta \varepsilon^2}{2 \log^2 m}$. Then $\beta' \cdot F_2(x)$ corresponds to the threshold T_1 , which is utilized in the proofs of Section 3.1. Similarly, Algorithm 3 leverages the quantity $\frac{\log n}{\beta' \alpha \varepsilon}$ to define the threshold T_2 , which is then utilized in the proofs of Section 3.2.

3.1 Recovery of High Frequency Levels

In this section, we describe our algorithm for recovering the high frequency levels, whose coordinates have sufficiently large magnitude and thus their frequencies can be well-approximated by running an L_2 -heavy hitters algorithm on the stream S . Moreover, with high probability, adding Laplacian noise will not affect the level sets because the frequencies are so large. Thus it simply suffices to return the noisy estimated frequencies of each of the elements in the high frequency levels. This algorithm is the simplest of our cases and we give the algorithm in full in Algorithm 2.

We first show that coordinates in high frequency levels are identified and their frequencies are accurately estimated. Similarly, we show that if a coordinate does not have high frequency, it will not be output by Algorithm 2.

■ **Algorithm 2** Algorithm to privately estimate the high levels.

Input: Privacy parameter $\varepsilon > 0$, accuracy parameter $\alpha \in (0, 1)$

Output: Private estimation of the frequencies of the coordinates of the high frequency levels

- 1: $\beta \leftarrow \mathcal{O}\left(\frac{\alpha^5}{\text{mmc}(L)^2 \log^5 m}\right)$, $\beta' \leftarrow \mathcal{O}\left(\frac{\alpha^2 \beta \varepsilon^2}{\log^2 m}\right)$
- 2: Run PRIVCOUNTSKETCH on the stream S with threshold $\alpha^2 \beta'$ and failure probability $\frac{1}{\text{poly}(m)}$
- 3: **for** each heavy-hitter $k \in [n]$ reported by PRIVCOUNTSKETCH **do**
- 4: Let $\widetilde{x}_k$ be the frequency estimated by PRIVCOUNTSKETCH
- 5: $\widehat{x}_k \leftarrow \widetilde{x}_k + \text{Lap}\left(\frac{8}{\beta' \varepsilon}\right)$
- 6: **return** $\widehat{x}_k$

► **Lemma 27.** Suppose $m = \frac{\Omega(\log^5 m)}{\alpha^5 \beta^2 \varepsilon^5}$. Then with high probability Algorithm 2 outputs $\widehat{x}_k$ such that if $x_k^2 \geq \frac{\alpha^2 \beta \varepsilon^2}{\log^2 m} F_2(x)$, then $(1 - \alpha^2)x_k \leq \widehat{x}_k \leq x_k$ and if $x_k^2 < \frac{\alpha^2 \beta \varepsilon^2}{2 \log^2 m} F_2(x)$, then $\widehat{x}_k < \frac{3\alpha^2 \beta \varepsilon^2}{4 \log^2 m} F_2(x)$.

We then show that Algorithm 2 preserves differential privacy and analyze its space complexity.

► **Lemma 28.** Algorithm 2 is $(\frac{\varepsilon}{4}, \frac{\delta}{4})$ -differentially private for $\delta = \frac{1}{\text{poly}(m)}$ and uses space $\text{mmc}(L)^2 \cdot \text{poly}\left(\frac{1}{\alpha}, \frac{1}{\varepsilon}, \log m\right)$.

3.2 Recovery of Medium Frequency Levels

In this section, we describe our algorithm for recovering the medium frequency levels, whose coordinates do not have sufficiently large magnitude to be detected by running an L_2 -heavy hitters algorithm on the stream S , but have sufficiently large size, so that there exists some $j \in [s]$ across the s subsampling levels such that the coordinates can be detected by running an L_2 -heavy hitters algorithm on the stream S_j . On the other hand, their magnitudes are sufficiently large so that with high probability, adding Laplacian noise will not affect the level sets. We give the algorithm in full in Algorithm 3.

We first upper bound the second frequency moment (and hence the L_2 norm) of each substream. This is necessary because we want to detect the coordinates of the medium frequency levels as L_2 -heavy hitters for each substream, but if the substream has overwhelmingly large L_2 norm, then we will not be able to find coordinates of the medium frequency levels. However, it may not be true that $F_2(S_j)$ is significantly smaller than $F_2(S)$ with high probability. For example, if there were a single large element, then the probability it is sampled at level s is $\frac{1}{2^s}$, which is roughly $\frac{1}{n} > \frac{1}{\text{poly}(m)}$. Instead, we note that PRIVCOUNTSKETCH benefits from the stronger *tail guarantee*, which states that not only does PRIVCOUNTSKETCH with threshold $\eta < 1$ detect the elements k such that $(x_k)^2 \geq \eta F_2(S)$, but it also detects the elements k such that $(x_k)^2 \geq \eta F_2(S_{\text{tail}(1/\eta)})$, where $S_{\text{tail}(1/\eta)}$ is the frequency vector x induced by S , with the largest $\frac{1}{\eta}$ entries instead set to zero [15, 17].

► **Lemma 29.** Consider a β -important level i with $\xi^{2i} \in \left[\frac{\beta \alpha^2 \varepsilon^2}{\log^2 m} \cdot \frac{F_2(x)}{2^j}, \frac{\beta \alpha^2 \varepsilon^2}{\log^2 m} \cdot \frac{F_2(x)}{2^{j-1}}\right]$ for some integer $j > 0$ and $\xi^i > \frac{\log n}{\beta' \alpha \varepsilon}$. If $F_2((S_j)_{1/(\alpha^2 \beta' \varepsilon^2)}) \leq \frac{200 \log m}{2^j} F_2(x)$ for all $j \in [s]$, then with high probability, Algorithm 3 outputs $\widehat{b}_i$ such that $(1 - \mathcal{O}(\alpha))b_i \leq \widehat{b}_i \leq b_i$, where b_i is the size of level i .

■ **Algorithm 3** Algorithm to privately estimate the medium levels.

Input: Privacy parameter $\varepsilon > 0$, accuracy parameter $\alpha \in (0, 1)$
Output: Private estimations of the sizes of the medium frequency levels

- 1: $\beta \leftarrow \mathcal{O}\left(\frac{\alpha^5}{\text{mmc}(L)^2 \log^5 m}\right)$, $\beta' \leftarrow \mathcal{O}\left(\frac{\alpha^3 \beta \varepsilon^2}{\log^2 m}\right)$, $\xi \leftarrow (1 + \mathcal{O}(\varepsilon))$
- 2: $\gamma \leftarrow (1/2, 1)$ uniformly at random, $\ell \leftarrow \lceil \log_\xi(2m) \rceil$, $s \leftarrow \mathcal{O}(\log n)$
- 3: **for** $j \in [s]$ with $2^j > \frac{\log n}{\beta' \alpha \varepsilon}$ **do**
- 4: Form stream S_j by sampling elements of $[n]$ with probability $\frac{1}{2^j}$
- 5: Run PRIVCOUNTSKETCH $_j$ on stream S_j with threshold $\alpha^2 \beta' \varepsilon^2$ and failure probability $\frac{1}{\text{poly}(m)}$
- 6: **for** each heavy-hitter $k \in [n]$ reported by PRIVCOUNTSKETCH $_j$ **do**
- 7: Let $\widehat{x}_k$ be the frequency estimated by PRIVCOUNTSKETCH $_j$
- 8: **if** $\widehat{x}_k > \frac{\log n}{\beta' \alpha \varepsilon}$ **then**
- 9: $\widetilde{x}_k \leftarrow \widehat{x}_k + \text{Lap}\left(\frac{8}{\beta' \varepsilon}\right)$
- 10: **for** $i \in [\ell]$ with $\frac{m^2}{2^{j+1}} > \gamma \xi^{2i} \geq 2^j > \mathcal{O}\left(\frac{\log n}{\beta' \alpha^2 \varepsilon}\right)$ **do**
- 11: Let $\widetilde{b}_i$ be the number of indices $k \in [n]$ such that $\gamma \xi^{2i} \leq \widetilde{x}_k < \gamma \xi^{2i+2}$
- 12: $\widehat{b}_i \leftarrow \frac{2^j}{(1 + \mathcal{O}(\alpha))} \widetilde{b}_i$
- 13: **return** $\widehat{b}_i$

We now show that Algorithm 3 preserves differential privacy and analyze its space complexity.

► **Lemma 30.** *Algorithm 3 is $(\frac{\varepsilon}{4}, \frac{\delta}{4})$ -differentially private for $\delta = \frac{1}{\text{poly}(m)}$ and uses space $\text{mmc}(L)^2 \cdot \text{poly}\left(\frac{1}{\alpha}, \frac{1}{\varepsilon}, \log m\right)$.*

3.3 Recovery of Low Frequency Levels

In this section, we describe our algorithm for recovering the low frequency levels, whose coordinates have magnitude small enough that we cannot add Laplacian noise to their frequencies without affecting the corresponding level set sizes. We instead report the sizes of the level sets for the low frequency levels rather than the approximate frequencies of the heavy-hitters. Thus we must add Laplacian noise to the sizes of the level sets; we show that the L_1 sensitivity for the level set estimations is particularly small for the low frequency levels and thus the Laplacian noise does not greatly affect the estimates of the level set sizes. We note that this approach does not work for the high frequency levels because the high frequency levels may have small level set sizes, so that adding Laplacian noise to the sizes can significantly affect the resulting estimates of the level set sizes. Similarly, it is more challenging to argue the low L_1 sensitivity for the level set estimations for the medium frequency levels. Hence, both the algorithm and analysis are especially well-catered to the low frequency levels. We give the algorithm in full in Algorithm 4.

We first show that the estimates of the level set sizes for the low frequency levels are accurate.

► **Lemma 31.** *Consider a β -important level i with $\xi^{2i} \in \left[\frac{\beta \alpha^2 \varepsilon^2}{\log^2 m} \cdot \frac{F_2(x)}{2^j}, \frac{\beta \alpha^2 \varepsilon^2}{\log^2 m} \cdot \frac{F_2(x)}{2^{j-1}}\right]$ for some integer $j > 0$ and $\xi^i \leq \frac{\log n}{\beta' \alpha \varepsilon}$. If $F_2((S_j)_{1/(\alpha^2 \beta' \varepsilon^2)}) \leq \frac{200 \log m}{2^j} F_2(x)$ for all $j \in [s]$, then with high probability, Algorithm 4 outputs $\widehat{b}_i$ such that*

$$(1 - \mathcal{O}(\alpha))b_i \leq \widehat{b}_i \leq b_i,$$

where b_i is the size of level set i .

■ **Algorithm 4** Algorithm to privately estimate the low levels.

Input: Privacy parameter $\varepsilon > 0$, accuracy parameter $\alpha \in (0, 1)$
Output: Private estimations of the sizes of the low frequency levels

- 1: $\beta \leftarrow \mathcal{O}\left(\frac{\alpha^5}{\text{mmc}(L)^2 \log^5 m}\right)$, $\beta' \leftarrow \mathcal{O}\left(\frac{\alpha^2 \beta \varepsilon}{\log n}\right)$, $\xi \leftarrow (1 + \mathcal{O}(\varepsilon))$
- 2: $\gamma \leftarrow (1/2, 1)$ uniformly at random, $\ell \leftarrow \lceil \log_\xi(2m) \rceil$, $s \leftarrow \mathcal{O}(\log n)$
- 3: **for** $j \in [s]$ with $2^j \leq \frac{\log n}{\beta' \alpha \varepsilon}$ **do**
- 4: Form stream S_j by sampling elements of $[n]$ with probability $\frac{1}{2^j}$
- 5: Run PRIVCOUNTSKETCH $_j$ on stream S_j with threshold $\beta'' := \mathcal{O}\left(\frac{\beta' \alpha^2 \varepsilon^3}{\log^2 n}\right)$
- 6: **for** each heavy-hitter $k \in [n]$ reported by PRIVCOUNTSKETCH $_j$ **do**
- 7: Let $\widehat{x}_k$ be the frequency estimated by PRIVCOUNTSKETCH $_j$
- 8: **for** $i \in [\ell]$ with $\mathcal{O}\left(\frac{\log n}{\beta' \alpha^2 \varepsilon}\right) \geq 2^{j+1} > \gamma \xi^{2i} \geq 2^j$ **do**
- 9: Let $\widetilde{b}_i$ be the number of indices $k \in [n]$ such that $\gamma \xi^{2i} \leq \widehat{x}_k < \gamma \xi^{2i+2}$
- 10: $\widehat{b}_i \leftarrow \frac{2^j}{(1+\mathcal{O}(\alpha))} \left(\widetilde{b}_i + \text{Lap}\left(\frac{8}{\varepsilon}\right) \right)$
- 11: **return** $\widehat{b}_i$

We then show that Algorithm 4 is differentially private and analyze its space complexity.

► **Lemma 32.** *Algorithm 4 is $(\frac{\varepsilon}{4}, \frac{\delta}{4})$ -differentially private for $\delta = \frac{1}{\text{poly}(m)}$ and uses space $\text{mmc}(L)^2 \cdot \text{poly}\left(\frac{1}{\alpha}, \frac{1}{\varepsilon}, \log m\right)$.*

3.4 Putting Things Together

We would like to combine the subroutines from the previous sections to output a private dataset for symmetric norm estimation. Thus it remains to describe how to privately partition the coordinates into the high, medium, and low frequency levels. To that end, we remark that by Lemma 14, the sensitivity of PRIVCOUNTSKETCH in Algorithm 1 is at most 2. Moreover, although PRIVCOUNTSKETCH actually provides an estimated frequency for each coordinate, for our purposes, we only need estimated frequencies for the L_2 -heavy hitters and there are at most $K := \mathcal{O}\left(\frac{1}{\eta^2}\right)$ possible L_2 -heavy hitters with whichever threshold η that we choose, e.g., $\eta = \alpha^2 \beta'$ in Algorithm 2. Thus it suffices to observe that we can privately partition the coordinates into the high, medium, and low frequency levels by first privately outputting the top K estimated frequencies and then partitioning the coordinates according to their noisy estimated frequencies, which can be viewed as post-processing. In particular, [56] observes that it suffices to add Laplacian noise with scale $\frac{8}{\eta \varepsilon}$ to each of the frequencies and then outputting the top K noisy estimated frequencies to achieve $\frac{\varepsilon}{4}$ -differential privacy.

We now finally put together the results from the previous sections to show the following result. We remark that we set $\varepsilon, \alpha = \tilde{\Omega}\left(\left(\frac{M^2}{m}\right)^{\frac{1}{30}}\right)$ so that along with the assumption that $m \geq n$, the conditions of the previous statements are satisfied, e.g., Lemma 34, we obtain the following formalization of Theorem 2.

► **Theorem 33.** *Given a parameter $M > 1$, let $\varepsilon, \alpha = \tilde{\Omega}\left(\left(\frac{M^2}{m}\right)^{\frac{1}{30}}\right)$. There exists a (ε, δ) -differentially private algorithm that outputs a set C , from which the $(1 + \alpha)$ -approximation to any norm, with maximum modulus of concentration at most M of a vector $x \in \mathbb{R}^n$ induced by a stream of length $\text{poly}(n)$ can be computed, with probability at least $1 - \delta$. The algorithm uses $M^2 \cdot \text{poly}\left(\frac{1}{\alpha}, \frac{1}{\varepsilon}, \log n, \log \frac{1}{\delta}\right)$ bits of space.*

References

- 1 Jayadev Acharya and Ziteng Sun. Communication complexity in locally private distribution estimation and heavy hitters. In *Proceedings of the 36th International Conference on Machine Learning, ICML*, pages 51–60, 2019.
- 2 Gergely Ács, Claude Castelluccia, and Rui Chen. Differentially private histogram publishing through lossy compression. In *12th IEEE International Conference on Data Mining, ICDM*, pages 1–10, 2012.
- 3 Aditya Akella, Ashwin Bharambe, Mike Reiter, and Srinivasan Seshan. Detecting ddos attacks on isp networks. In *Proceedings of the Twenty-Second ACM SIGMOD/PODS Workshop on Management and Processing of Data Streams*, pages 1–3, 2003.
- 4 Noga Alon, Yossi Matias, and Mario Szegedy. The space complexity of approximating the frequency moments. *J. Comput. Syst. Sci.*, 58(1):137–147, 1999.
- 5 Alexandr Andoni. High frequency moments via max-stability. In *2017 IEEE International Conference on Acoustics, Speech and Signal Processing, ICASSP*, pages 6364–6368, 2017.
- 6 Alexandr Andoni, Huy L. Nguyen, Aleksandar Nikolov, Ilya P. Razenshteyn, and Erik Waingarten. Approximate near neighbors for general symmetric norms. In *Proceedings of the 49th Annual ACM SIGACT Symposium on Theory of Computing, STOC*, pages 902–913, 2017.
- 7 Andreas Argyriou, Rina Foygel, and Nathan Srebro. Sparse prediction with the k -support norm. In *Advances in Neural Information Processing Systems 25: Annual Conference on Neural Information Processing Systems*, pages 1466–1474, 2012.
- 8 Raef Bassily, Kobbi Nissim, Uri Stemmer, and Abhradeep Thakurta. Practical locally private heavy hitters. *J. Mach. Learn. Res.*, 21:16:1–16:42, 2020.
- 9 Raef Bassily and Adam D. Smith. Local, private, efficient protocols for succinct histograms. In *Proceedings of the Forty-Seventh Annual ACM on Symposium on Theory of Computing, STOC*, pages 127–135, 2015.
- 10 Rajendra Bhatia. *Matrix analysis*, volume 169. Springer Science & Business Media, 2013.
- 11 Jaroslaw Blasiok, Vladimir Braverman, Stephen R. Chestnut, Robert Krauthgamer, and Lin F. Yang. Streaming symmetric norms via measure concentration. In *Proceedings of the 49th Annual ACM SIGACT Symposium on Theory of Computing, STOC*, pages 716–729, 2017.
- 12 Jeremiah Blocki, Avrim Blum, Anupam Datta, and Or Sheffet. The johnson-lindenstrauss transform itself preserves differential privacy. In *53rd Annual IEEE Symposium on Foundations of Computer Science, FOCS*, pages 410–419, 2012.
- 13 Jeremiah Blocki, Elena Grigorescu, Tamalika Mukherjee, and Samson Zhou. How to make your approximation algorithm private: A black-box differentially-private transformation for tunable approximation algorithms of functions with low sensitivity. *CoRR*, abs/2210.03831, 2022.
- 14 Jeremiah Blocki, Seunghoon Lee, Tamalika Mukherjee, and Samson Zhou. Differentially private L_2 -heavy hitters in the sliding window model. In *The Eleventh International Conference on Learning Representations, ICLR*, 2023.
- 15 Vladimir Braverman, Stephen R. Chestnut, Nikita Ivkin, Jelani Nelson, Zhengyu Wang, and David P. Woodruff. Bptree: An l_2 heavy hitters algorithm using constant memory. In *Proceedings of the 36th ACM SIGMOD-SIGACT-SIGAI Symposium on Principles of Database Systems, PODS*, pages 361–376, 2017.
- 16 Vladimir Braverman, Stephen R. Chestnut, Nikita Ivkin, and David P. Woodruff. Beating counts sketch for heavy hitters in insertion streams. In *Proceedings of the 48th Annual ACM SIGACT Symposium on Theory of Computing, STOC*, pages 740–753, 2016.
- 17 Vladimir Braverman, Elena Grigorescu, Harry Lang, David P. Woodruff, and Samson Zhou. Nearly optimal distinct elements and heavy hitters on sliding windows. In *Approximation, Randomization, and Combinatorial Optimization. Algorithms and Techniques, APPROX/RANDOM*, pages 7:1–7:22, 2018.

- 18 Vladimir Braverman, Emanuele Viola, David P. Woodruff, and Lin F. Yang. Revisiting frequency moment estimation in random order streams. In *45th International Colloquium on Automata, Languages, and Programming, ICALP*, pages 25:1–25:14, 2018.
- 19 Vladimir Braverman, Viska Wei, and Samson Zhou. Symmetric norm estimation and regression on sliding windows. In *Computing and Combinatorics – 27th International Conference, COCOON, Proceedings*, pages 528–539, 2021.
- 20 Leo Breiman. Random forests. *Machine learning*, 45(1):5–32, 2001.
- 21 Zhiqi Bu, Sivakanth Gopi, Janardhan Kulkarni, Yin Tat Lee, Judy Hanwen Shen, and Uthaipon Tantipongpipat. Fast and memory efficient differentially private-sgd via JL projections. *CoRR*, abs/2102.03013, 2021.
- 22 Mark Bun, Jelani Nelson, and Uri Stemmer. Heavy hitters and the structure of local privacy. *ACM Trans. Algorithms*, 15(4):51:1–51:40, 2019.
- 23 T.-H. Hubert Chan, Mingfei Li, Elaine Shi, and Wenchang Xu. Differentially private continual monitoring of heavy hitters from distributed streams. In *Privacy Enhancing Technologies – 12th International Symposium, PETS Proceedings*, pages 140–159, 2012.
- 24 Moses Charikar, Kevin C. Chen, and Martin Farach-Colton. Finding frequent items in data streams. *Theor. Comput. Sci.*, 312(1):3–15, 2004.
- 25 Seung Geol Choi, Dana Dachman-Soled, Mukul Kulkarni, and Arkady Yerukhimovich. Differentially-private multi-party sketching for large-scale statistics. *Proc. Priv. Enhancing Technol.*, 2020(3):153–174, 2020.
- 26 Graham Cormode, S. Muthukrishnan, and Irina Rozenbaum. Summarizing and mining inverse distributions on data streams via dynamic inverse sampling. In *Proceedings of the 31st International Conference on Very Large Data Bases*, pages 25–36, 2005.
- 27 Bolin Ding, Janardhan Kulkarni, and Sergey Yekhanin. Collecting telemetry data privately. In *Advances in Neural Information Processing Systems 30: Annual Conference on Neural Information Processing Systems*, pages 3571–3580, 2017.
- 28 Itai Dinur, Uri Stemmer, David P. Woodruff, and Samson Zhou. On differential privacy and adaptive data analysis with bounded space. In *Advances in Cryptology – EUROCRYPT – 42nd Annual International Conference on the Theory and Applications of Cryptographic Techniques, Proceedings, Part III*, pages 35–65, 2023.
- 29 Cynthia Dwork, Frank McSherry, Kobbi Nissim, and Adam D. Smith. Calibrating noise to sensitivity in private data analysis. In *Theory of Cryptography, Third Theory of Cryptography Conference, TCC, Proceedings*, pages 265–284, 2006.
- 30 Cynthia Dwork, Moni Naor, Toniann Pitassi, Guy N. Rothblum, and Sergey Yekhanin. Pan-private streaming algorithms. In *Innovations in Computer Science – ICS. Proceedings*, pages 66–80, 2010.
- 31 Cynthia Dwork and Aaron Roth. The algorithmic foundations of differential privacy. *Found. Trends Theor. Comput. Sci.*, 9(3-4):211–407, 2014.
- 32 Cristian Estan, George Varghese, and Mike Fisk. Bitmap algorithms for counting active flows on high speed links. In *Proceedings of the 3rd ACM SIGCOMM conference on Internet measurement*, pages 153–166, 2003.
- 33 Sheldon J. Finkelstein, Mario Schkolnick, and Paolo Tiberio. Physical database design for relational databases. *ACM Trans. Database Syst.*, 13(1):91–128, 1988.
- 34 Sumit Ganguly. Taylor polynomial estimator for estimating frequency moments. In *Automata, Languages, and Programming – 42nd International Colloquium, ICALP Proceedings, Part I*, pages 542–553, 2015.
- 35 Sumit Ganguly and David P. Woodruff. High probability frequency moment sketches. In *45th International Colloquium on Automata, Languages, and Programming, ICALP*, pages 58:1–58:15, 2018.
- 36 Corrado Gini. Variabilità e mutabilità. *Reprinted in Memorie di metodologica statistica*, 1912.

- 37 Nicholas J. A. Harvey, Jelani Nelson, and Krzysztof Onak. Sketching and streaming entropy via approximation theory. In *49th Annual IEEE Symposium on Foundations of Computer Science, FOCS*, pages 489–498, 2008.
- 38 Piotr Indyk. Stable distributions, pseudorandom generators, embeddings, and data stream computation. *J. ACM*, 53(3):307–323, 2006.
- 39 Piotr Indyk and Andrew McGregor. Declaring independence via the sketching of sketches. In *Proceedings of the Nineteenth Annual ACM-SIAM Symposium on Discrete Algorithms, SODA*, pages 737–745, 2008.
- 40 Piotr Indyk and David P. Woodruff. Optimal approximations of the frequency moments of data streams. In *Proceedings of the 37th Annual ACM Symposium on Theory of Computing*, pages 202–208, 2005.
- 41 Daniel M. Kane, Jelani Nelson, Ely Porat, and David P. Woodruff. Fast moment estimation in data streams in optimal space. In *Proceedings of the 43rd ACM Symposium on Theory of Computing, STOC*, pages 745–754, 2011.
- 42 Daniel M. Kane, Jelani Nelson, and David P. Woodruff. On the exact space complexity of sketching and streaming small norms. In *Proceedings of the Twenty-First Annual ACM-SIAM Symposium on Discrete Algorithms, SODA*, pages 1161–1178, 2010.
- 43 Bo’az Klartag and Roman Vershynin. Small ball probability and dvoretzky’s theorem. *Israel Journal of Mathematics*, 157(1):193–207, 2007.
- 44 Balachander Krishnamurthy, Subhabrata Sen, Yin Zhang, and Yan Chen. Sketch-based change detection: Methods, evaluation, and applications. In *Proceedings of the 3rd ACM SIGCOMM conference on Internet measurement*, pages 234–247, 2003.
- 45 Ping Li. Estimators and tail bounds for dimension reduction in l_α ($0 < \alpha \leq 2$) using stable random projections. In *Proceedings of the Nineteenth Annual ACM-SIAM Symposium on Discrete Algorithms, SODA*, pages 10–19, 2008.
- 46 Yi Li and David P. Woodruff. A tight lower bound for high frequency moment estimation with small error. In *Approximation, Randomization, and Combinatorial Optimization. Algorithms and Techniques – 16th International Workshop, APPROX 2013, and 17th International Workshop, RANDOM. Proceedings*, pages 623–638, 2013.
- 47 Yi Li and David P. Woodruff. Input-sparsity low rank approximation in Schatten norm. In *Proceedings of the 37th International Conference on Machine Learning, ICML*, pages 6001–6009, 2020.
- 48 Zaoxing Liu, Antonis Manousis, Gregory Vorsanger, Vyas Sekar, and Vladimir Braverman. One sketch to rule them all: Rethinking network flow monitoring with univmon. In *Proceedings of the ACM SIGCOMM 2016 Conference*, pages 101–114, 2016.
- 49 Zaoxing Liu, Samson Zhou, Ori Rottenstreich, Vladimir Braverman, and Jennifer Rexford. Memory-efficient performance monitoring on programmable switches with lean algorithms. In *1st Symposium on Algorithmic Principles of Computer Systems, APOCS*, pages 31–44, 2020.
- 50 Max O Lorenz. Methods of measuring the concentration of wealth. *Publications of the American statistical association*, 9(70):209–219, 1905.
- 51 Andrew M. McDonald, Massimiliano Pontil, and Dimitris Stamos. Spectral k-support norm regularization. In *Advances in Neural Information Processing Systems 27: Annual Conference on Neural Information Processing Systems*, pages 3644–3652, 2014.
- 52 Vitali D Milman and Gideon Schechtman. *Asymptotic theory of finite dimensional normed spaces: Isoperimetric inequalities in riemannian manifolds*, volume 1200. Springer, 2009.
- 53 Darakhshan J. Mir, S. Muthukrishnan, Aleksandar Nikolov, and Rebecca N. Wright. Pan-private algorithms via statistics on sketches. In *Proceedings of the 30th ACM SIGMOD-SIGACT-SIGART Symposium on Principles of Database Systems, PODS*, pages 37–48, 2011.
- 54 Noam Nisan. Pseudorandom generators for space-bounded computation. *Comb.*, 12(4):449–461, 1992.
- 55 Christopher R Palmer, Georgos Siganos, Michalis Faloutsos, Christos Faloutsos, and Phillip B Gibbons. The connectivity and fault-tolerance of the internet topology, 2001.

- 56 Gang Qiao, Weijie J. Su, and Li Zhang. Oneshot differentially private top-k selection. In *Proceedings of the 38th International Conference on Machine Learning, ICML*, pages 8672–8681, 2021.
- 57 Patricia G. Selinger, Morton M. Astrahan, Donald D. Chamberlin, Raymond A. Lorie, and Thomas G. Price. Access path selection in a relational database management system. In *Proceedings of the 1979 ACM SIGMOD International Conference on Management of Data*, pages 23–34, 1979.
- 58 Or Sheffet. Differentially private ordinary least squares. *J. Priv. Confidentiality*, 9(1), 2019.
- 59 Adam D. Smith, Shuang Song, and Abhradeep Thakurta. The flajolet-martin sketch itself preserves differential privacy: Private counting with minimal space. In *Advances in Neural Information Processing Systems 33: Annual Conference on Neural Information Processing Systems, NeurIPS*, 2020.
- 60 Zhao Song, Ruosong Wang, Lin F. Yang, Hongyang Zhang, and Peilin Zhong. Efficient symmetric norm regression via linear sketching. In *Advances in Neural Information Processing Systems 32: Annual Conference on Neural Information Processing Systems, NeurIPS*, pages 828–838, 2019.
- 61 Jakub Tetek. Additive noise mechanisms for making randomized approximation algorithms differentially private. *CoRR*, abs/2211.03695, 2022.
- 62 Mikkel Thorup and Yin Zhang. Tabulation based 4-universal hashing with applications to second moment estimation. In *Proceedings of the Fifteenth Annual ACM-SIAM Symposium on Discrete Algorithms, SODA*, pages 615–624, 2004.
- 63 Lun Wang, Iosif Pinelis, and Dawn Song. Differentially private fractional frequency moments estimation with polylogarithmic space. In *The Tenth International Conference on Learning Representations, ICLR*, 2022.
- 64 Tianhao Wang, Milan Lopuhaä-Zwakenberg, Zitao Li, Boris Skoric, and Ninghui Li. Locally differentially private frequency estimation with consistency. In *27th Annual Network and Distributed System Security Symposium, NDSS*, 2020.
- 65 David P. Woodruff and Samson Zhou. Separations for estimating large frequency moments on data streams. In *48th International Colloquium on Automata, Languages, and Programming, ICALP*, volume 198, pages 112:1–112:21, 2021.
- 66 David P. Woodruff and Samson Zhou. Tight bounds for adversarially robust streams and sliding windows via difference estimators. In *62nd IEEE Annual Symposium on Foundations of Computer Science, FOCS*, pages 1183–1196, 2021.
- 67 Bin Wu, Chao Ding, Defeng Sun, and Kim-Chuan Toh. On the moreau-yosida regularization of the vector k-norm related functions. *SIAM J. Optim.*, 24(2):766–794, 2014.
- 68 Jia Xu, Zhenjie Zhang, Xiaokui Xiao, Yin Yang, Ge Yu, and Marianne Winslett. Differentially private histogram publication. *VLDB J.*, 22(6):797–822, 2013.

A Missing Proofs

We first show that coordinates in high frequency levels are identified and their frequencies are accurately estimated.

► **Lemma 34.** *Suppose $x_k^2 \geq \frac{\alpha^2 \beta \varepsilon^2}{\log^2 m} F_2(x)$ and $m = \frac{\Omega(\log^5 m)}{\alpha^5 \beta^2 \varepsilon^5}$. Then with high probability, Algorithm 2 outputs $\widehat{x}_k$ such that*

$$(1 - \alpha^2)x_k \leq \widehat{x}_k \leq x_k.$$

Proof. Consider Algorithm 2. Since $x_k^2 \geq \frac{\alpha^2 \beta \varepsilon^2}{2 \log^2 m} F_2(x)$ and we call PRIVCOUNTSKETCH with threshold $\alpha^2 \beta'$ with $\beta' := \mathcal{O}\left(\frac{\alpha^2 \beta \varepsilon^2}{\log^2 m}\right)$, then with high probability, the output $\widetilde{x}_k$ satisfies

$$(1 - \mathcal{O}(\alpha^2))x_k \leq \widetilde{x}_k \leq x_k.$$

We then add Laplacian noise $\text{Lap}\left(\frac{8}{\beta'\varepsilon}\right)$ to $\widetilde{x}_k$ to form $\widehat{x}_k$. Since $x_k^2 \geq \frac{\alpha^2\beta\varepsilon^2}{2\log^2 m} F_2(x) = \beta' F_2(x)$ and $F_2(x) \geq m$, then with high probability, the Laplacian noise is at most an α^2 fraction of $\widehat{x}_k$ for $\frac{\mathcal{O}(\log m)}{\beta'\varepsilon} \leq \alpha^2 m$ or equivalently, $m \geq \frac{\Omega(\log m)}{\alpha(\beta')^2\varepsilon} \geq \frac{\Omega(\log^5 m)}{\alpha^5\beta^2\varepsilon^5}$. Hence with high probability,

$$(1 - \alpha^2)x_k \leq \widehat{x}_k \leq x_k. \quad \blacktriangleleft$$

Similarly, we show that if a coordinate does not have high frequency, it will not be output by Algorithm 2.

► **Lemma 35.** *Suppose $x_k^2 < \frac{\alpha^2\beta\varepsilon^2}{2\log^2 m} F_2(x)$ and $m = \frac{\Omega(\log^5 m)}{\alpha^5\beta^2\varepsilon^5}$. Then with high probability, Algorithm 2 outputs $\widehat{x}_k$ such that*

$$\widehat{x}_k < \frac{3\alpha^2\beta\varepsilon^2}{4\log^2 m} F_2(x).$$

Proof. Since $x_k^2 < \frac{\alpha^2\beta\varepsilon^2}{2\log^2 m} F_2(x)$ and we call PRIVCOUNTSKETCH with threshold $\alpha^2\beta'$ with $\beta' := \mathcal{O}\left(\frac{\alpha^2\beta\varepsilon^2}{\log^2 m}\right)$, then the output $\widetilde{x}_k$ satisfies

$$|(\widetilde{x}_k)^2 - (x_k)^2| \leq 2\alpha^2\beta' F_2(x).$$

We then add Laplacian noise $\text{Lap}\left(\frac{8}{\beta'\varepsilon}\right)$ to $\widetilde{x}_k$ to form $\widehat{x}_k$. Since $F_2(x) \geq m$, then with high probability, the Laplacian noise is at most an $\alpha^2\beta'$ fraction of $F_2(x)$ for $\frac{\mathcal{O}(\log m)}{\beta'\varepsilon} \leq \alpha^2 m$ or equivalently, $m \geq \frac{\Omega(\log m)}{\alpha(\beta')^2\varepsilon} \geq \frac{\Omega(\log^5 m)}{\alpha^5\beta^2\varepsilon^5}$. Hence with high probability,

$$|(\widehat{x}_k)^2 - (x_k)^2| \leq \frac{\alpha^2\beta\varepsilon^2}{4\log^2 m} F_2(x).$$

Since $x_k^2 < \frac{\alpha^2\beta\varepsilon^2}{2\log^2 m} F_2(x)$, then it follows that

$$\widehat{x}_k < \frac{3\alpha^2\beta\varepsilon^2}{4\log^2 m} F_2(x). \quad \blacktriangleleft$$

We now show that Algorithm 2 preserves differential privacy.

► **Lemma 36.** *Algorithm 2 is $(\frac{\varepsilon}{4}, \frac{\delta}{4})$ -differentially private for $\delta = \frac{1}{\text{poly}(m)}$. Algorithm 2 uses space $\text{mmc}(L)^2 \cdot \text{poly}\left(\frac{1}{\alpha}, \frac{1}{\varepsilon}, \log m\right)$.*

Proof. By Lemma 14, the sensitivity of PRIVCOUNTSKETCH is at most 2 and the failure probability is $\frac{1}{\text{poly}(m)}$. Thus by adding Laplacian noise $\text{Lap}\left(\frac{8}{\beta'\varepsilon}\right)$ to $\widetilde{x}_k$, each estimated frequency is $(\frac{\beta'\varepsilon}{4}, \frac{\delta}{4\beta})$ -differentially private for $\delta = \frac{1}{\text{poly}(m)}$. Since PRIVCOUNTSKETCH with threshold β' can release at most $\frac{1}{\beta}$ estimated frequencies and post-processing does not cause loss in privacy, then by Theorem 11, Algorithm 2 is $(\frac{\varepsilon}{4}, \frac{\delta}{4})$. $\blacktriangleleft$

Finally, we analyze the space complexity of Algorithm 2.

► **Lemma 37.** *Algorithm 2 uses space $\text{mmc}(L)^2 \cdot \text{poly}\left(\frac{1}{\alpha}, \frac{1}{\varepsilon}, \log m\right)$.*

Proof. The space complexity follows from running a single instance of PRIVCOUNTSKETCH with threshold $\alpha^2\beta'$ and failure probability $\frac{1}{\text{poly}(m)}$, where $\beta' = \mathcal{O}\left(\frac{\alpha^2\beta\varepsilon^2}{\log^2 m}\right)$ and $\beta = \mathcal{O}\left(\frac{\alpha^5}{\text{mmc}(L)^2 \log^5 m}\right)$. $\blacktriangleleft$

► **Lemma 38.** *With high probability, we have that $F_2((S_j)_{1/(\alpha^2 \beta' \varepsilon^2)}) \leq \frac{200 \log m}{2^j} F_2(x)$ for all $j \in [s]$.*

Proof. For each $j \in [s]$, we have that $\mathbb{E}[F_2(S_j)] = \frac{F_2(x)}{2^j}$. By Chernoff bounds with $\mathcal{O}(\log n)$ -wise limited independence, we have that

$$\Pr \left[F_2((S_j)_{1/(\alpha^2 \beta' \varepsilon^2)}) > \frac{200 \log m}{2^j} F_2(x) \right] \leq \frac{1}{\text{poly}(m)}.$$

Since $s \leq 2 \log m$, then by a union bound over all $j \in [s]$, we have that $F_2(S_j) \leq (200 \log m) F_2(x)$ for all $j \in [s]$. ◀

We now show that conditioned on the event that the L_2 norm of the subsampled streams are not too large, then we can well-approximate the frequency of any coordinate of the medium frequency levels, provided that they are sampled in the substream.

► **Lemma 39.** *Suppose i is a β -important level and $k \in [n]$ is in level i , so that $x_k \in [\xi^i, \xi^{i+1})$. If $F_2((S_j)_{1/(\alpha^2 \beta' \varepsilon^2)}) \leq \frac{200 \log m}{2^j} F_2(x)$ for all $j \in [s]$ and k is sampled in stream S_j with $2^j > \frac{\log n}{\beta' \alpha \varepsilon}$, then with high probability, Algorithm 3 outputs $\widehat{x}_k$ such that*

$$(1 - \alpha^2)x_k \leq \widehat{x}_k \leq x_k.$$

Proof. Consider Algorithm 3. By Lemma 26, $x_k^2 \in \left[\frac{\alpha^2 \beta \varepsilon^2}{\log^2 m} \cdot \frac{F_2(x)}{2^j}, \frac{\alpha^2 \beta \varepsilon^2}{\log^2 m} \cdot \frac{F_2(x)}{2^{j-1}} \right]$. Conditioned on the event that $F_2((S_j)_{1/(\alpha^2 \beta' \varepsilon^2)}) \leq \frac{200 \log m}{2^j} F_2(x)$ for all $j \in [s]$, then $x_k^2 \geq \frac{\alpha^2 \beta \varepsilon^2}{200 \log m} F_2(S_j)$. We call PRIVCOUNTSKETCH with threshold $\alpha^2 \beta' \varepsilon^2 = \mathcal{O}\left(\frac{\alpha^4 \beta \varepsilon^3}{\log^2 m}\right)$. Thus with high probability, the output $\widetilde{x}_k$ satisfies

$$(1 - \mathcal{O}(\alpha^2))x_k \leq \widetilde{x}_k \leq x_k.$$

We then add Laplacian noise $\text{Lap}\left(\frac{8}{\beta' \varepsilon}\right)$ to $\widetilde{x}_k$ to form $\widehat{x}_k$. Since $x_k^2 \geq \mathcal{O}\left(\frac{\log n}{\beta' \alpha^2 \varepsilon}\right)$, then with high probability, the Laplacian noise is at most an α^2 fraction of $\widehat{x}_k$. Hence with high probability,

$$(1 - \alpha^2)x_k \leq \widehat{x}_k \leq x_k. \quad \blacktriangleleft$$

Unfortunately, Lemma 39 only provides guarantees for the coordinates of the medium frequency levels that are sampled. Thus, we still need to use Lemma 39 to show that a good estimator to the sizes of the medium frequency levels can be obtained from the estimates of the coordinates of the medium frequency levels that are sampled. In particular, we show that rescaling the empirical sizes of the medium frequency levels forms a good estimator to the actual sizes of the medium frequency levels.

Proof of Lemma 29. Suppose i is a β -important level. Then by Lemma 26 and a shifting of the index j , $b_i \geq \mathcal{O}\left(\frac{2^j \log^2 m}{\alpha^2 \varepsilon^2}\right)$. Thus in S_j , the expected number of items E_j from level i is at least $\frac{\log^2 m}{\alpha^2 \varepsilon^2}$ and the variance V_j is at most E_j . Hence by Chernoff bounds with $\mathcal{O}(\log n)$ -wise limited independence, we have that the number of items N_j from level i satisfies

$$(1 - \mathcal{O}(\alpha))b_i \leq 2^j \cdot N_j \leq (1 + \mathcal{O}(\alpha))b_i,$$

with high probability. [11] show that due to the uniformly random chosen $\gamma \in (1/2, 1)$, we further have

$$(1 - \mathcal{O}(\alpha))N_j \leq (1 + \mathcal{O}(\alpha))\widehat{b}_i \leq (1 + \mathcal{O}(\alpha))N_j,$$

with high probability. Since $s \leq 2 \log m$, then by a union bound over all $j \in [s]$, we have that with high probability, Algorithm 3 outputs $\widehat{b}_i$ such that

$$(1 - \mathcal{O}(\alpha))b_i \leq \widehat{b}_i \leq b_i. \quad \blacktriangleleft$$

We now show that Algorithm 3 preserves differential privacy.

► **Lemma 40.** *Algorithm 3 is $(\frac{\varepsilon}{4}, \frac{\delta}{4})$ -differentially private for $\delta = \frac{1}{\text{poly}(m)}$.*

Proof. By Lemma 14, the sensitivity of PRIVCOUNTSKETCH is at most 2 and the failure probability is $\frac{1}{\text{poly}(m)}$. Thus by adding Laplacian noise $\text{Lap}\left(\frac{8}{\beta'\varepsilon}\right)$ to $\widehat{x}_k$, each estimated frequency is $(\frac{\beta'\varepsilon}{4}, \frac{\delta}{4\beta})$ -differentially private for $\delta = \frac{1}{\text{poly}(m)}$. Since PRIVCOUNTSKETCH with threshold β' can release at most $\frac{1}{\beta}$ estimated frequencies, then by Theorem 11, Algorithm 3 is $(\frac{\varepsilon}{4}, \frac{\delta}{4})$. $\blacktriangleleft$

It remains to analyze the space complexity of Algorithm 3.

► **Lemma 41.** *Algorithm 3 uses space $\text{mmc}(L)^2 \cdot \text{poly}\left(\frac{1}{\alpha}, \frac{1}{\varepsilon}, \log m\right)$.*

Proof. The space complexity follows from running s instances of PRIVCOUNTSKETCH with threshold $\alpha^2\beta'$ and failure probability $\frac{1}{\text{poly}(m)}$, where $\beta' = \mathcal{O}\left(\frac{\alpha^2\beta\varepsilon^2}{\log^2 m}\right)$ and $\beta = \mathcal{O}\left(\frac{\alpha^5}{\text{mmc}(L)^2 \log^5 m}\right)$. Since $s = \mathcal{O}(\log n)$ and we assume $n \leq m$ so that $\mathcal{O}(\log n) = \mathcal{O}(\log m)$, then the space complexity follows. $\blacktriangleleft$

Proof of Lemma 31. Suppose i is a β -important level. Hence by a shifting of the index j in Lemma 26, we have that $b_i \geq \mathcal{O}\left(\frac{2^j \log^2 m}{\alpha^2 \varepsilon^2}\right)$. Therefore, the expected number of items E_j from level i sampled in the substream S_j is at least $\frac{\log^2 m}{\alpha^2 \varepsilon^2}$ and the variance V_j is at most E_j . Thus by Chernoff bounds with $\mathcal{O}(\log n)$ -wise limited independence, the number of items N_j from level i satisfies

$$(1 - \mathcal{O}(\alpha))b_i \leq 2^j \cdot N_j \leq (1 + \mathcal{O}(\alpha))b_i,$$

with high probability. [11] show that due to the uniformly random chosen $\gamma \in (1/2, 1)$, we further have

$$(1 - \mathcal{O}(\alpha))N_j \leq (1 + \mathcal{O}(\alpha))\widehat{b}_i \leq (1 + \mathcal{O}(\alpha))N_j,$$

with high probability. Since $s \leq 2 \log m$ and $\text{Lap}\left(\frac{8}{\varepsilon}\right)$ is at most an ε -fraction of $b_i \geq \mathcal{O}\left(\frac{2^j \log^2 m}{\alpha^2 \varepsilon^2}\right)$ with high probability, then by a union bound over all $j \in [s]$, we have that with high probability, Algorithm 3 outputs $\widehat{b}_i$ such that

$$(1 - \mathcal{O}(\alpha))b_i \leq \widehat{b}_i \leq b_i. \quad \blacktriangleleft$$

We then show that Algorithm 4 is differentially private.

► **Lemma 42.** *Algorithm 4 is $(\frac{\varepsilon}{4}, \frac{\delta}{4})$ -differentially private for $\delta = \frac{1}{\text{poly}(m)}$.*

Proof. Note that since each instance of PRIVCOUNTSKETCH _{j} uses threshold $\beta'' := \mathcal{O}\left(\frac{\beta'\alpha^2\varepsilon^3}{\log^2 n}\right)$ on a stream S_j with $F_2(S_j) \leq \frac{200 \log m}{2^j} F_2(x)$, then for any $k \in [n]$ with $x_k \leq \mathcal{O}\left(\frac{\log n}{\beta'\alpha^2\varepsilon}\right)$, we have that PRIVCOUNTSKETCH _{j} outputs x_k exactly. Hence, at most two estimates of the sizes of the level sets $\widehat{b}_i$ can change, and then can change by at most one. Thus the sensitivity is at most 2, so it suffices to add Laplacian noise $\text{Lap}\left(\frac{8}{\varepsilon}\right)$ to each estimate $\widehat{b}_i$ to obtain $(\frac{\varepsilon}{4}, \frac{\delta}{4})$ -differentially private for $\delta = \frac{1}{\text{poly}(m)}$. $\blacktriangleleft$

Finally, we argue the space complexity of Algorithm 4.

► **Lemma 43.** *Algorithm 4 uses space $\text{mmc}(L)^2 \cdot \text{poly}\left(\frac{1}{\alpha}, \frac{1}{\varepsilon}, \log m\right)$.*

Proof. Similar to Algorithm 3, the space complexity follows as a result of running s instances of PRIVCOUNTSKETCH with threshold $\alpha^2\beta'$ and failure probability $\frac{1}{\text{poly}(m)}$, where $\beta' = \mathcal{O}\left(\frac{\alpha^2\beta\varepsilon^2}{\log^2 m}\right)$ and $\beta = \mathcal{O}\left(\frac{\alpha^5}{\text{mmc}(L)^2 \log^5 m}\right)$. Since $s = \mathcal{O}(\log n)$ and we assume $n \leq m$ so that $\mathcal{O}(\log n) = \mathcal{O}(\log m)$, then the space complexity follows. ◀

► **Theorem 44.** *Given a parameter $M > 1$, let $\varepsilon, \alpha = \tilde{\Omega}\left(\left(\frac{M^2}{m}\right)^{\frac{1}{30}}\right)$. There exists a (ε, δ) -differentially private algorithm that outputs a set C , for $\delta = \frac{1}{\text{poly}(m)}$. From C , the $(1 + \alpha)$ -approximation to any norm with maximum modulus of concentration at most M can be computed, with probability at least $1 - \delta$. The algorithm uses $M^2 \cdot \text{poly}\left(\frac{1}{\alpha}, \frac{1}{\varepsilon}, \log m\right)$ bits of space.*

Proof. Note that from Lemma 34 and Lemma 35, the frequencies of the coordinates in the high frequency levels are well-approximated with high probability. Similarly, from Lemma 29 and Lemma 31, the sizes of the level sets of the medium and low frequency levels are well-approximated with high probability. Moreover, all the level sets are partitioned into the high, medium, or low frequency levels. We would like to say that by Lemma 23, these statistics are sufficient to recover a $(1 + \alpha)$ -approximation to any norm with *maximum modulus of concentration* at most M and so we achieve a $(1 + \alpha)$ -approximation to any norm with maximum modulus of concentration at most M that with high probability. Indeed, in an idealized process where $\xi^i \leq \widehat{x}_k \leq \xi^{i+1}$ if and only if k is sampled by the substream j assigned to level i and $\xi^i \leq x_k < \xi^{i+1}$, Lemma 23 would show that we achieve a $(1 + \alpha)$ -approximation to any norm with maximum modulus of concentration at most M that with high probability. However, this may not always be the case because the frequency x_k may lie near the boundary of the interval $[\xi^i, \xi^{i+1})$ and the estimate $\widehat{x}_k$ may lie outside of the interval, in which case $\widehat{x}_k$ is used toward the estimation of some other level set. Thus, our algorithm randomizes the boundaries of the level sets by instead defining the level sets as $[\gamma\xi^i, \gamma\xi^{i+1})$ for some $\gamma \in (1/2, 1)$ chosen uniformly at random. Since we call PRIVCOUNTSKETCH with threshold at most $\alpha^2\beta'$, then the probability that item $k \in [n]$ is misclassified over the choice of γ is at most $\mathcal{O}(\alpha^2\beta')$. Furthermore, if k in level set i is misclassified, it can only be classified into level set $i - 1$ or $i + 1$, causing at most an incorrect multiplicative factor of two. Then in expectation across all $k \in [n]$, the error due to the misclassification is at most an $\mathcal{O}(\alpha^2\beta')$ fraction of the symmetric norm. Hence by Markov's inequality, the error due to the misclassification is at most an additive $\frac{\alpha}{2}$ fraction of the symmetric norm with probability at least 0.99. To obtain high probability of success, it then suffices to take the median across $\mathcal{O}(\log m)$ independent instances, finally showing correctness of our algorithm.

The private partitioning of the coordinates into the high, medium, and low frequency levels is $\frac{\varepsilon}{4}$ -differentially private. Each of the three sets of statistics released by the high, medium, and low frequency levels are $(\frac{\varepsilon}{4}, \frac{\delta}{4})$ -differentially private, by Lemma 36, Lemma 40, and Lemma 42. Then (ε, δ) -differential privacy follows from the composition of differential privacy, i.e., Theorem 11.

Finally, the space complexity follows from Lemma 37, Lemma 41, and Lemma 43. ◀

We remark that our algorithm is presented as having unlimited access to random bits but is analyzed using $\mathcal{O}(\log m)$ -wise independence, so it can be properly derandomized to provide the space guarantees without needing to store a large number of random bits. Alternatively, our algorithm can also be derandomized using Nisan's pseudorandom generator, which induces an extra multiplicative factor of $\mathcal{O}(\log m)$ in the space overhead [54].

Finally, we remark that the failure probability can be raised from $\delta = \frac{1}{\text{poly}(m)}$ to arbitrarily $\delta > 0$ using additional space overhead $\text{polylog } \frac{1}{\delta}$, since the space dependency in each subroutine on the failure probability δ is $\text{polylog } \frac{1}{\delta}$.