

Computational Hardness of Private Coreset

Badih Ghazi 

Google Research, Mountain View, CA, USA

Cristóbal Guzmán 

Pontificia Universidad Católica de Chile, Santiago, Chile

Pritish Kamath 

Google Research, Mountain View, CA, USA

Alexander Knop 

Google Research, Mountain View, CA, USA

Ravi Kumar 

Google Research, Mountain View, CA, USA

Pasin Manurangsi 

Google Research, Bangkok, Thailand

Abstract

We study the problem of differentially private (DP) computation of coreset for the k -means objective. For a given input set of points, a coreset is another set of points such that the k -means objective for any candidate solution is preserved up to a multiplicative $(1 \pm \alpha)$ factor (and some additive factor).

We prove the first computational lower bounds for this problem. Specifically, assuming the existence of one-way functions, we show that no polynomial-time $(\epsilon, 1/n^{\omega(1)})$ -DP algorithm can compute a coreset for k -means in the ℓ_∞ -metric for some constant $\alpha > 0$ (and some constant additive factor), even for $k = 3$. For k -means in the Euclidean metric, we show a similar result but only for $\alpha = \Theta(1/d^2)$, where d is the dimension.

2012 ACM Subject Classification Theory of computation $\rightarrow$ Computational complexity and cryptography; Security and privacy

Keywords and phrases Differentially Private Clustering, Coreset, Cryptographic Hardness

Digital Object Identifier 10.4230/LIPIcs.FORC.2026.1

1 Introduction

The widespread collection and use of personal data has necessitated rigorous frameworks for preserving individual privacy. Differential Privacy (DP), introduced by Dwork et al. [15, 14], has established itself as the gold standard framework, enabling data release while protecting the confidentiality of users. Roughly speaking, DP requires that the output distributions of the algorithm on two *neighboring* input datasets – those that differ on a single user’s input – to be similar. The similarity is measured by two parameters $\epsilon, \delta \geq 0$; the standard setting, which will also be our focus, is $\epsilon = O(1)$ and $\delta = 1/n^{\omega(1)}$, where n denotes the input dataset size (Definition 1). The DP framework has been successfully deployed in various large-scale applications, such as those by the U.S. Census Bureau [1] and Google’s RAPPOR [17].

Among the myriad of data analysis tasks, clustering remains a fundamental primitive in unsupervised learning. Consequently, a significant body of research has been dedicated to designing clustering algorithms that satisfy DP [6, 29, 4, 21, 39, 38, 33, 42, 26, 7, 24, 34]. While most of these works focus on developing efficient algorithms that output a final set of cluster centers, many of them employ a private *coreset* as a subroutine.



© Badih Ghazi, Cristóbal Guzmán, Pritish Kamath, Alexander Knop, Ravi Kumar, and Pasin Manurangsi;

licensed under Creative Commons License CC-BY 4.0

7th Symposium on Foundations of Responsible Computing (FORC 2026).

Editor: Huijia (Rachel) Lin; Article No. 1; pp. 1:1–1:14



Leibniz International Proceedings in Informatics

Schloss Dagstuhl – Leibniz-Zentrum für Informatik, Dagstuhl Publishing, Germany

Recall that a coreset is another (possibly weighted) set of points such that the clustering (e.g., k -means or k -median) objective for any candidate solution is preserved up to a multiplicative $(1 \pm \alpha)$ factor and some additive factor $\pm\beta$ (Definition 6).¹ In non-private settings, coreset computation has been extensively studied with the focus usually on finding a coreset of small size [31, 23, 30, 8, 35, 19, 22, 20, 41, 5, 32, 12, 11, 9]; from this perspective, a coreset can be viewed as a succinct summary of the original dataset. Coresets are highly useful in data analysis since they allow “big data” to be replaced by a “small summary”, thereby drastically reducing the computational cost of downstream algorithms. Furthermore, since the coreset approximates the cost function for any candidate solution, it enables hyperparameter tuning (e.g., k in k -means) and trying multiple heuristics (e.g., different initializations in k -means) efficiently, all without revisiting the massive original dataset.

In the context of privacy, a coreset offers a distinct advantage besides efficiency: it decouples the privacy cost from the analysis. Once a private coreset is computed, an analyst can run any clustering algorithm (e.g., k -means, k -median) on the coreset as many times as desired (for a fixed number of centers), without consuming additional privacy budget. Feldman et al. [18] first introduced the notion of private coresets and showed that they can be constructed information-theoretically with small errors. In particular, they give² an ε -DP algorithm for computing k -median and k -means coresets with $\beta = \tilde{O}_\alpha\left(\frac{kd}{\varepsilon n}\right)$, where d denotes the dimension. Alas, their algorithm is based on the exponential mechanism [37], which is not efficient; specifically, its running time is exponential in both k and d . They also give a more efficient algorithm that removes the exponential dependence on k in the running time, but both the running time and the additive error β now become exponential in d . Subsequent works have explored different DP coreset algorithms for k -means in various settings [26, 7, 34]; yet they all have this exponential dependence and efficient construction has remained a major open challenge.

In this work, we investigate if this computational efficiency barrier is fundamental. Namely,

*Does there exist a polynomial-time DP algorithm that returns
a coreset with any approximation factor arbitrarily close to 1?*

We remark that, for private coresets, the requirement of small size is not critical: given a privately generated coreset of arbitrary size, one can apply a non-private coreset construction scheme to compress it as a post-processing step. Due to this, we will not enforce the size restriction in the remainder of the paper. Note that this only strengthens our lower bounds.

1.1 Our Contributions

We provide a negative answer to the above question, establishing computational hardness results for private coreset construction. By reducing from the hardness of generating synthetic data [43], we show that efficient private coreset estimation cannot be achieved with standard cryptographic assumptions.

Our main results are the following:

¹ As formalized in Definition 6, the additive error β in our notation is normalized so that it always lies in $[0, 1]$, unlike some previous work (e.g., [18]) where β is unnormalized and can be as large as n .

² Strictly speaking, the bound stated in [18, Corollary 3.3] is $\beta = \tilde{O}_\alpha\left(\frac{k^2 d^2 \log n}{\varepsilon n}\right)$ and is only for the Euclidean metric. The bound we claim here is by replacing the use of [8] with [13], which works for any metric space with doubling dimension d (including the ℓ_∞ -metric).

- **Hardness in the ℓ_∞ -metric**³: We prove that, assuming the existence of one-way functions⁴, no polynomial-time (ε, δ) -DP algorithm can construct a coreset for k -means in the ℓ_∞ -metric that achieves an approximation factor $1 \pm \alpha$ for some constant $\alpha > 0$, even for $k = 3$ (Theorem 8).
- **Hardness in the Euclidean metric**: We extend our analysis to the Euclidean setting and show that for the ℓ_2 -metric, any efficient private coreset algorithm must incur an approximation error that scales inverse polynomially with the dimension. Specifically, we rule out polynomial-time algorithms achieving an approximation factor better than $1 \pm \Theta(1/d^2)$ (Theorem 12).

Our results show a dichotomy: while private coresets exist information-theoretically, they are computationally difficult to find under standard hardness assumptions. We further note that without the coreset size restriction, the non-private version of the problem is trivial, as one can simply output the input dataset. Thus, the computational challenges we present here are truly unique to the private setting.

1.2 Technical Overview

To show computational hardness, we use a reduction from the hardness of privately generating synthetic data for 3-literal disjunctions [43]. Roughly speaking, [43] shows that, for 3-literal disjunction queries, it is hard to sanitize a dataset even in the case where we “promise” that some subset of these queries are always evaluated to one (“satisfiable”) in the input and that we only wish to be accurate on these queries.

Our reduction is somewhat straightforward: We use the same dataset as the above problem (up to scaling)! The main property we need to show here is that, if we can find a DP coreset of such a dataset, then we can use the coreset to construct DP synthetic dataset for the 3-literal disjunction sanitization. Again, this latter step is somewhat straightforward: Given a coreset, we simply round each coordinate to -1 (false) or $+1$ (true) based on their signs. This completes our reduction.

The main challenge however is to show that this is a valid reduction. Namely, that the output gives accurate estimates of the satisfiable 3-literal disjunctions. It turns out that this can be viewed as a 3-means cost query. Namely, suppose that we have a disjunction consisting of $x_{i_1}, x_{i_2}, x_{i_3}$ with signs $s_{i_1}, s_{i_2}, s_{i_3}$ respectively. Then, we can create three centers, where the j th center only has the i_j coordinate set to $2s_{i_j}$ (and other coordinates being zero). Assuming that each output point $\tilde{x}$ belongs to $\{\pm 1\}^d$, it is simple to verify that the distance is small only when the clause is satisfied. Notice that the “gap” between satisfiable and unsatisfiable in the ℓ_∞ -metric is constant (i.e., 3) whereas that of the ℓ_2 -metric depends on d (i.e., $1 + \Theta(1/d)$). This is the reason that we get a constant α in Theorem 8 but requires α that decreases with d in Theorem 12. This completes the high-level overview of our reduction; in the actual proof, we of course cannot assume that each $\tilde{x}$ belongs to $\{\pm 1\}^d$ and a significant effort is required to deal with this issue.

³ While k -means was first studied in the Euclidean metric, it has by now been well studied on other metrics, including ℓ_∞ metric; see, e.g., [10] and the references therein.

⁴ Since we do not deal with this assumption directly, we refrain from defining one-way functions and discussing their importance. Nevertheless, we note that this is a central assumption in cryptography and detailed discussion can be found in any standard textbook on the topic, e.g., [27].

1.3 Related Work

Computational lower bounds intrinsic to DP are scarce. Dwork et al. [16] established such lower bounds for generating synthetic data when either the data universe or the number of queries is large (here large means exponential in some underlying size parameter). In particular, they provide a class of queries such that, assuming the existence of one-way functions, no polynomial-time $(\varepsilon, 1/n^{\omega(1)})$ -DP can generate a synthetic dataset that closely matches the query values of the input dataset. However, the class of queries produced in their arguments is specifically crafted based on a signature scheme, and does not correspond to a naturally studied query class. Subsequently, Ullman and Vadhan [43] extended the hardness to a large class of queries, including natural ones such as 2-way marginal queries. Their construction is based on the NP-hardness (via Karp reductions) of constraint satisfaction problems (CSPs), which focuses on discrete domains. From this perspective, our work expands the computational hardness landscape of private algorithms to the case of continuous domains, which we hope will stimulate further studies along this direction.

2 Preliminaries

Let $\mathbf{1}_S \in \{0, 1\}^n$ denote the indicator vector of a set $S \subseteq [n]$; for ease, we use $\mathbf{1}_i$ for $\mathbf{1}_{\{i\}}$. For $v \in \mathbb{R}$, let $\text{sgn}(v) \in \{-1, 1\}$ denote the sign of v (where $\text{sgn}(0) = 1$); moreover, for $x \in \mathbb{R}^d$, let $\text{sgn}(x) \in \{-1, 1\}^d$ denote the coordinate-wise sign of x , i.e., $\text{sgn}(x) = (\text{sgn}(x_1), \dots, \text{sgn}(x_d))$.

Let $\mathcal{X}$ be a domain. For a query $f : \mathcal{X} \rightarrow [0, 1]$ and a dataset $D = (x_1, \dots, x_n) \in \mathcal{X}^n$, let the linear query $f(D)$ be defined as $f(D) := \frac{1}{n} \sum_{i \in [n]} f(x_i) = \mathbb{E}_{x \sim D}[f(x)]$.

2.1 Differential Privacy

We quickly recall the definition of differential privacy (DP) here. Two datasets $D, D' \in \mathcal{X}^*$ are *neighbors* iff they are of the same size and they differ on a single value, i.e., $D = (x_1, \dots, x_n)$, $D' = (x'_1, \dots, x'_n) \in \mathcal{X}^n$ and there exists $i^* \in [n]$ such that $x_i = x'_i$ for all $i \neq i^*$.

► **Definition 1** (Differential Privacy [15, 14]). *Let $\varepsilon \geq 0, \delta \in [0, 1]$. An algorithm $\mathcal{A} : \mathcal{X}^* \rightarrow \mathcal{O}$ is (ε, δ) -differentially private (i.e., (ε, δ) -DP) iff, for any two neighboring datasets D, D' and for any subset $S \subseteq \mathcal{O}$ of outputs, we have*

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

2.2 Synthetic Data Generation

As alluded to above, both the starting point of our reduction and the coreset problem itself can be viewed as synthetic data generation problems. An algorithm for generating a synthetic data is also referred to as a “sanitizer”. More formally, a *sanitizer* $\mathcal{A}$ is a randomized algorithm that takes in the dataset $D \in \mathcal{X}^*$ and outputs another dataset $\tilde{D} \in \mathcal{X}^*$.

2.2.1 Accuracy and Efficiency

The accuracy of the output dataset is often measure against some family $\mathcal{F}$ of linear queries $f : \mathcal{X} \rightarrow [0, 1]$; the following is a standard definition used in literature (e.g. [16, 43]).

► **Definition 2** (Accuracy). A dataset $\tilde{D}$ is an $(\mathcal{F}, \alpha)$ -accurate estimate of D iff $|f(\tilde{D}) - f(D)| \leq \alpha$ for all $f \in \mathcal{F}$. A sanitizer $\mathcal{A}$ is $(\mathcal{F}, \alpha, n)$ -accurate if for any dataset⁵ $D \in \mathcal{X}^n$, $\Pr_{\tilde{D} \sim \mathcal{A}(D)}[\tilde{D} \text{ is an } (\mathcal{F}, \alpha)\text{-accurate estimate of } D] \geq \frac{2}{3}$.

Throughout our work, we will consider the setting where $\mathcal{X} = \mathcal{X}_d \subseteq \mathbb{R}^d$ and, thus, the family of queries $\mathcal{F}_d$ is also parameterized by d . We say that the sanitizer runs in polynomial time if it runs in $(nd)^{O(1)}$ time. Furthermore, the accuracy guarantee has to be against *all values of $d \in \mathbb{N}$* . In this setting, we define a *parameterized family* of queries to be a family $\mathcal{F} = \bigcup_{d \in \mathbb{N}} \mathcal{F}_d$ where $\mathcal{F}_d$ is the family of queries corresponding to $\mathcal{X}_d \subseteq \mathbb{R}^d$. We can define the accuracy for a parameterized family of queries as follows.

► **Definition 3** (Accuracy w.r.t. Parameterized Family of Queries). For a parameterized family of queries $\mathcal{F} = \bigcup_{d \in \mathbb{N}} \mathcal{F}_d$ where $\mathcal{F}_d$ is the family of queries corresponding to $\mathcal{X}_d \subseteq \mathbb{R}^d$, we say that a sanitizer $\mathcal{A}$ is $(\mathcal{F}, \alpha)$ -accurate if the following holds: There exists a constant $\nu > 0$ such that, for all $d \in \mathbb{N}$ and $n \geq \Theta(d^\nu)$, $\mathcal{A}$ is $(\mathcal{F}_d, \alpha, n)$ -accurate.

2.2.2 Promise Sanitizer

To prove our results, we will use the hardness result from [43, Theorem 4.4], which actually holds even against a weaker notion which we name *promise sanitizers*. In words, promise sanitizers only provide the accuracy guarantee on the queries $f \in \mathcal{F}$ whose values are one on the entire input dataset. For the queries that violate this condition, promise sanitizers do not provide any accuracy guarantee. This is formalized below.

► **Definition 4** (Promise sanitizer). Let $\mathcal{F}$ be a class of queries and $\gamma \in [0, 1]$. $\mathcal{A}$ is an $(\mathcal{F}, \gamma, n)$ -promise sanitizer if $\Pr_{\tilde{D} \sim \mathcal{A}(D)}[f(\tilde{D}) \geq \gamma \text{ for all } f \in \tilde{\mathcal{F}}] \geq 2/3$, for any dataset $D \in \mathcal{X}^n$ and $\tilde{\mathcal{F}} \subseteq \mathcal{F}$ such that $f(D) = 1$ for all $f \in \tilde{\mathcal{F}}$.

Similar to Definition 3, for a parameterized family of queries $\mathcal{F} = \bigcup_{d \in \mathbb{N}} \mathcal{F}_d$, we say that $\mathcal{A}$ is an $(\mathcal{F}, \gamma)$ -promise sanitizer if there exists a constant $\nu > 0$ such that, for all $d \in \mathbb{N}$ and $n \geq \Theta(d^\nu)$, $\mathcal{A}$ is an $(\mathcal{F}_d, \gamma, n)$ -promise sanitizer.

Note that a $(\mathcal{F}, 1 - \gamma)$ -accurate sanitizer is a $(\mathcal{F}, \gamma)$ -promise sanitizer. In other words, a promise sanitizer is a weaker notion compared to an accurate sanitizer. While Ullman and Vadhan [43] only stated their results in term of an accurate sanitizer, we observe that their proof of [43, Theorem 4.4] already yields the hardness for a promise sanitizer as well⁶.

To state their result, let $\mathcal{F}_d^{3\text{Disj}}$ denote the class of all 3-literal disjunctions where the domain is $\{\pm 1\}^d$ (where we think of -1 as “false” and $+1$ as “true”). More formally, for every (literal indices) $\mathbf{i} = (i_1, i_2, i_3) \in [d]^3$ and (signs) $\mathbf{s} = (s_1, s_2, s_3) \in \{\pm 1\}^3$, we define $\psi_{\mathbf{i}, \mathbf{s}} : \{\pm 1\}^d \rightarrow \{\pm 1\}$ as follows:

$$\psi_{\mathbf{i}, \mathbf{s}}(x) = 1 - \frac{1}{4} \cdot (1 - s_1 x_{i_1})(1 - s_2 x_{i_2})(1 - s_3 x_{i_3})$$

In other words, this is the evaluation of the 3-literal disjunction $s_1 x_{i_1} \vee s_2 x_{i_2} \vee s_3 x_{i_3}$ on x . We then let $\mathcal{F}_d^{3\text{Disj}} := \{\psi_{\mathbf{i}, \mathbf{s}} \mid \mathbf{i} \in [d]^3, \mathbf{s} \in \{\pm 1\}^3\}$, and let $\mathcal{F}^{3\text{Disj}}$ be the parameterized family $\bigcup_{d \in \mathbb{N}} \mathcal{F}_d^{3\text{Disj}}$. The hardness result we need is as follows.

⁵ We note that the failure probability is usually parameterized in the accuracy notation but here we fix it to $2/3$ for simplicity, since we are focusing on proving lower bounds. We note that the probability can be easily reduced by employing DP hyperparameter tuning (e.g., [36, 40, 25]).

⁶ See Appendix A for more detail.

► **Theorem 5** ([43]). *Let $\varepsilon > 0$ be any constant. Assuming the existence of one-way functions, there is no polynomial-time $(\varepsilon, 1/n^{\omega(1)})$ -DP $(\mathcal{F}^{\text{3Disj}}, \gamma)$ -promise sanitizer for some constant $\gamma \in (0, 1)$.*

2.3 Coreset Estimation

Let $\mathbb{B}_{p,d}$ denote a unit ball in the ℓ_p -metric on $\mathbb{R}^d$; i.e., $\mathbb{B}_{p,d} := \{x \in \mathbb{R}^d : \|x\|_p \leq 1\}$. In the ℓ_p -coreset problem, we consider the domain $\mathcal{X} = \mathbb{B}_{p,d}$. Let $\mathcal{F}_{p,d}^{k\text{-means}}$ denote the family of coreset cost functions; i.e., $\mathcal{F}_{p,d}^{k\text{-means}} = \{\text{cost}_{\mathcal{C},p} : \mathbb{B}_{p,d} \rightarrow [0, 1] \mid \mathcal{C} \in \mathbb{B}_{p,d}^k\}$, where $\text{cost}_{\mathcal{C},p}(x) = \min_{c \in \mathcal{C}} \|x - c\|_p^2$. Finally, we let $\mathcal{F}_p^{k\text{-means}}$ be the parameterized family $\bigcup_{d \in \mathbb{N}} \mathcal{F}_{p,d}^{k\text{-means}}$.

► **Definition 6** (ℓ_p -coreset). *A dataset $\tilde{D} \in \mathbb{B}_{p,d}^*$ is a (k, p, α, β) -coreset of $D \in \mathbb{B}_{p,d}^*$ if⁷*

$$(1 - \alpha) \cdot \text{cost}_{\mathcal{C},p}(D) - \beta \leq \text{cost}_{\mathcal{C},p}(\tilde{D}) \leq (1 + \alpha) \cdot \text{cost}_{\mathcal{C},p}(D) + \beta,$$

for all $\mathcal{C} \in \mathbb{B}_{p,d}^k$. An algorithm $\mathcal{A}$ is a $(k, p, \alpha, \beta, d, n)$ -coreset estimator iff for any dataset $D \in \mathbb{B}_{p,d}^n$, it holds that

$$\Pr_{\tilde{D} \sim \mathcal{A}(D)} [\tilde{D} \text{ is a } (k, p, \alpha, \beta)\text{-coreset of } D] \geq 2/3.$$

Similar to Definition 3, we say that $\mathcal{A}$ is a (k, p, α, β) -coreset estimator if there exists a constant $\nu > 0$ such that, for all $d \in \mathbb{N}$ and $n \geq \Theta(d^\nu)$, $\mathcal{A}$ is a $(k, p, \alpha, \beta, d, n)$ -coreset estimator.

As explained in the introduction, we do not require that the private coreset be small. We also note that, unlike standard notations of coreset, we also require the coreset to be unweighted for notational convenience. Since there is no restriction on the coreset size and we allow an additive error β , this is w.l.o.g. because we can always rescale the weights by a large factor (e.g., $O(|\tilde{D}|/\beta)$) and round them to integers to obtain an unweighted dataset.

Also, to avoid dealing with two parameters α, β , it will be more convenient to work with the accuracy definition of the sanitizer (Definitions 2 and 3) instead. Indeed, it is simple to see that any algorithm that solves (k, p, α, β) -coreset is also an accurate sanitizer for the query family $\mathcal{F}_p^{k\text{-means}}$, as formalized below.

► **Observation 7.** *Any (k, p, α, β) -coreset estimator is an $(\mathcal{F}_p^{k\text{-means}}, 4\alpha + \beta)$ -accurate sanitizer.*

Proof. This follows immediately from the definition of coreset, since $\text{cost}_{\mathcal{C},p}(D) \in [0, 2]$. ◀

3 Hardness in the ℓ_∞ -metric

In this section, we prove the hardness of computing DP coreset in the ℓ_∞ -metric:

► **Theorem 8.** *Let $\varepsilon > 0$ be any constant. Assuming the existence of one-way functions, there is no polynomial-time $(\varepsilon, 1/n^{\omega(1)})$ -DP algorithm for the $(3, \infty, \alpha, \beta)$ -coreset estimation problem for some constant $\alpha, \beta > 0$.*

⁷ We use the same notation as linear queries: for any dataset D , let $\text{cost}_{\mathcal{C},p}(D) = \mathbb{E}_{x \sim D}[\text{cost}_{\mathcal{C},p}(x)]$.

As alluded to earlier, it is more convenient to prove the hardness in terms of sanitizer, as stated in Theorem 9 below. Due to Observation 7, Theorem 8 follows as an immediate corollary of Theorem 9.

► **Theorem 9.** *Let $\varepsilon > 0$ be any constant. Assuming the existence of one-way functions, there is no polynomial-time $(\varepsilon, 1/n^{\omega(1)})$ -DP $(\mathcal{F}_\infty^{k\text{-means}}, \xi)$ -accurate sanitizer for $k = 3$ and some constant $\xi \in [0, 1)$.*

The proof of this theorem follows closely the outline in Section 1.2. Namely, for an input dataset D for the $(\mathcal{F}^{3\text{Disj}}, \gamma)$ -promise sanitization problem (where each $x \in D$ belongs to $\{-1, 1\}^d$), we simply construct a dataset D' which is D scaled by a factor of $\kappa = \frac{1}{2}$. We then pass D' to the coreset estimator and round it to obtain the final output.

As explained in Section 1.2, for each $\psi \in \mathcal{F}^{3\text{Disj}}$, we can construct a 3-means query—where the centers are denoted by $\mathcal{C}^\psi$ in the proof below—to check if $\psi(D) = 1$. Since we do not have any guarantee that the coreset only contains points in $\{-\kappa, \kappa\}^d$, we additionally consider the origin as a center to prevent any “cheating”. Intuitively, if some coordinates have absolute values larger than κ , then its distance to the origin would be too large. Otherwise, if the absolute values are smaller than κ , then its distance to $\mathcal{C}^\psi$ is too large. This is argued formally below in Claim 11.

Proof of Theorem 9. We prove this by a reduction from the $(\mathcal{F}^{3\text{Disj}}, \gamma)$ -promise sanitization problem, which is known to be hard from Theorem 5. Assume for the sake of contradiction that there exists a polynomial-time $(\varepsilon, 1/n^{\omega(1)})$ -DP $(\mathcal{F}_\infty^{k\text{-means}}, \xi)$ -sanitizer $\mathcal{A}$, where $\xi = \frac{1-\gamma}{4}$.

We construct the $(\mathcal{F}^{3\text{Disj}}, \gamma)$ -promise sanitizer $\tilde{\mathcal{A}}$ as follows. Let $\kappa = \frac{1}{2}$.

- Let $D \in (\{\pm 1\}^d)^n$ be the input.
- Construct dataset $D' = \{x' = \kappa \cdot x \mid x \in D\}$.
- Run $\mathcal{A}(D')$ to get an output $\tilde{D}'$.
- Construct dataset $\tilde{D} = \{\tilde{x} = \text{sgn}(\tilde{x}') \mid \tilde{x}' \in \tilde{D}'\}$.
- Output $\tilde{D}$.

Since $\tilde{\mathcal{A}}$ is a post-processing of $\mathcal{A}$, the $(\varepsilon, 1/n^{\omega(1)})$ -DP guarantee continues to hold.

Next, we will show that $\tilde{\mathcal{A}}$ solves the $(\mathcal{F}^{3\text{Disj}}, \gamma)$ -promise sanitization problem. To do this, assume (the promise) that $\psi(D) = 1$ for all $\psi \in \tilde{\mathcal{F}} \subseteq \mathcal{F}^{3\text{Disj}}$. Since $\mathcal{A}$ is an $(\mathcal{F}_\infty^{k\text{-means}}, \xi)$ -accurate sanitizer, with probability at least $2/3$, the following holds for all $\mathcal{C} \in \mathbb{B}_p^k$:

$$|\text{cost}_{\mathcal{C}, \infty}(\mathcal{A}(D')) - \text{cost}_{\mathcal{C}, \infty}(D')| = |\text{cost}_{\mathcal{C}, \infty}(\tilde{D}') - \text{cost}_{\mathcal{C}, \infty}(D')| \leq \xi. \quad (1)$$

We will show that, when the above holds, we have $\psi(\tilde{D}) \geq \gamma$ for all $\psi \in \tilde{\mathcal{F}}$.

Let $\mathcal{C}^0 = \{0\}$. Furthermore, consider any 3-literal disjunction $\psi = \psi_{i,s}$ in $\tilde{\mathcal{F}}$. Recall that the indices of the three literals are i_1, i_2, i_3 and their signs are s_1, s_2, s_3 respectively. Consider centers $c^{\psi,1}, c^{\psi,2}, c^{\psi,3}$ defined by $c^{\psi,j} = 2\kappa \cdot s_j \cdot \mathbf{1}_{i_j}$. Finally, let $\mathcal{C}^\psi = (c^{\psi,1}, c^{\psi,2}, c^{\psi,3})$.

There are two claims crucial to the proof. The first is a claim that upper bounds $\text{cost}_{\mathcal{C}^0, \infty}(x') + \text{cost}_{\mathcal{C}^\psi, \infty}(x')$ for $x' \in D'$:

▷ **Claim 10.** For any $x' \in D'$ and $\psi \in \tilde{\mathcal{F}}$, we have $\text{cost}_{\mathcal{C}^0, \infty}(x') + \text{cost}_{\mathcal{C}^\psi, \infty}(x') \leq 2\kappa^2$.

Proof. Notice that $\text{cost}_{\mathcal{C}^0, \infty}(x') = \|x'\|_\infty^2 = \kappa^2$ for all $x' \in D'$.

Meanwhile, from the assumption that $\psi(D) = 1$, we also have that $\psi(x) = 1$ for all $x \in D$. That is, there exists a literal of ψ , say the j th literal for $j \in [3]$, that is set to true in x (i.e., $x_{i_j} = s_j$). Consider $x' - c^{\psi,j}$. We claim that all of its coordinates have absolute value κ , i.e., $x' - c^{\psi,j} \in \{\pm\kappa\}^d$. To see this, consider two cases based on the coordinate $i \in [d]$:

- Case I: $i \neq i_j$. In this case, we simply have $(x' - c^{\psi,j})_i = x'_i \in \{\pm\kappa\}$ as desired.
 - Case II: $i = i_j$. In this case, $(x' - c^{\psi,j})_{i_j} = \kappa \cdot x_{i_j} - 2\kappa \cdot s_j = -\kappa \cdot x_{i_j} \in \{\pm\kappa\}$.
- Hence, we have $x' - c^{\psi,j} \in \{\pm\kappa\}^d$, which implies that $\|x' - c^{\psi,j}\|_\infty = \kappa$. Thus, $\text{cost}_{\mathcal{C}^0,\infty}(x') + \text{cost}_{\mathcal{C}^\psi,\infty}(x') \leq \kappa^2 + \kappa^2 = 2\kappa^2$. $\triangleleft$

The second claim is a lower bound on $\text{cost}_{\mathcal{C}^0,\infty}(\tilde{x}') + \text{cost}_{\mathcal{C}^\psi,\infty}(\tilde{x}')$ for $\tilde{x}' \in \tilde{D}'$, as stated below. It should be noted that, when $\psi(\tilde{x}) = 1$, this lower bound is exactly the same as the bound in the above claim. However, if $\psi(\tilde{x}) = 0$, the lower bound here is larger. This will indeed allow us to bound the number of $\tilde{x}$'s that fall into the latter case.

▷ **Claim 11.** For any $\tilde{x}' \in \tilde{D}'$ and $\psi \in \tilde{\mathcal{F}}$, we have $\text{cost}_{\mathcal{C}^0,\infty}(\tilde{x}') + \text{cost}_{\mathcal{C}^\psi,\infty}(\tilde{x}') \geq 2\kappa^2(2 - \psi(\tilde{x}))$.

Proof. First, consider the case $\psi(\tilde{x}) = 1$. We have $\text{cost}_{\mathcal{C}^\psi,\infty}(\tilde{x}') = \|c^{\psi,j} - \tilde{x}'\|_\infty^2$ for some $j \in [3]$. We can lower bound this further by observing that $\|c^{\psi,j} - \tilde{x}'\|_\infty \geq |(c^{\psi,j})_{i_j} - \tilde{x}'_{i_j}| = |2\kappa \cdot s_j - \tilde{x}'_{i_j}|$. Meanwhile, $\text{cost}_{\mathcal{C}^0,\infty}(\tilde{x}') = \|\tilde{x}'\|_\infty^2 \geq (\tilde{x}'_{i_j})^2$. Thus, we have

$$\begin{aligned} \text{cost}_{\mathcal{C}^0,\infty}(\tilde{x}') + \text{cost}_{\mathcal{C}^\psi,\infty}(\tilde{x}') &\geq (2\kappa \cdot s_j - \tilde{x}'_{i_j})^2 + (\tilde{x}'_{i_j})^2 = 2\kappa^2 + 2(\kappa \cdot s_j - \tilde{x}'_{i_j})^2 \\ &\geq 2\kappa^2. \end{aligned}$$

Next, suppose that $\psi(\tilde{x}) = 0$. In this case, for all $j \in [3]$, we have that $(c^{\psi,j})_{i_j}$ and $\tilde{x}'_{i_j}$ have different signs. Thus, $\|c^{\psi,j} - \tilde{x}'\|_\infty \geq |(c^{\psi,j})_{i_j} - \tilde{x}'_{i_j}| \geq |(c^{\psi,j})_{i_j}| = 2\kappa$. As a result, we have $\text{cost}_{\mathcal{C}^\psi,\infty}(\tilde{x}') \geq (2\kappa)^2 = 4\kappa^2$ as desired. $\triangleleft$

Combining the above two claims and (1), we get

$$\begin{aligned} 2\kappa^2(2 - \psi(\tilde{D})) &= \mathbb{E}_{\tilde{x} \sim \tilde{D}} [2\kappa^2(2 - \psi(\tilde{x}))] \\ (\text{Claim 11}) &\leq \mathbb{E}_{\tilde{x}' \sim \tilde{D}'} [\text{cost}_{\mathcal{C}^0,\infty}(\tilde{x}') + \text{cost}_{\mathcal{C}^\psi,\infty}(\tilde{x}')] \\ &= \text{cost}_{\mathcal{C}^0,\infty}(\tilde{D}') + \text{cost}_{\mathcal{C}^\psi,\infty}(\tilde{D}') \\ \text{Using (1)} &\leq \text{cost}_{\mathcal{C}^0,\infty}(D') + \text{cost}_{\mathcal{C}^\psi,\infty}(D') + 2\xi \\ (\text{Claim 10}) &\leq 2\kappa^2 + 2\xi. \end{aligned}$$

Thus, from our choice of parameter $\xi = \frac{1-\gamma}{4}$, $\kappa = \frac{1}{2}$, we have that $\psi(\tilde{D}) \geq \gamma$. As a result, we can conclude that $\tilde{\mathcal{A}}$ solves the $(\mathcal{F}^{\text{3Disj}}, \gamma)$ -promise sanitization problem. Finally, by Theorem 5, this cannot hold assuming the existence of one-way functions. $\blacktriangleleft$

4 Hardness in the Euclidean metric

Next, we prove that hardness in the Euclidean metric, which is similar to the hardness for the ℓ_∞ -metric (Theorem 8) except with $\alpha, \beta = \Omega(\frac{1}{d^2})$ instead of constants.

► **Theorem 12.** *Let $\varepsilon > 0$ be any constant. Assuming the existence of one-way functions, there is no polynomial-time $(\varepsilon, 1/n^{\omega(1)})$ -DP algorithm for the $(3, 2, \alpha, \beta)$ -coreset estimation problem for some $\alpha = \Omega(\frac{1}{d^2})$ and $\beta = \Omega(\frac{1}{d^2})$.*

Again, we actually prove the sanitizer variant of the theorem, as stated below, from which Theorem 12 follows as a corollary (due to Observation 7).

► **Theorem 13.** *Let $\varepsilon > 0$ be any constant. Assuming the existence of one-way functions, there is no polynomial-time $(\varepsilon, 1/n^{\omega(1)})$ -DP $(\mathcal{F}_2^{k\text{-means}}, \xi)$ -sanitizer for $k = 3$ and some $\xi = \Omega(1/d^2)$.*

The proof follows a similar strategy as in Theorem 9 but it requires a more subtle sequence of inequalities to prove the accuracy bound. We note that we also need to set a smaller scaling factor κ : Setting $\kappa = \frac{1}{2}$ as before would result in $\|x'\|_2 = \kappa\sqrt{d} > 1$. This suggests setting $\kappa = \frac{1}{\sqrt{d}}$. However, we actually go with the setting $\kappa = \frac{1}{d}$ instead because, to facilitate one of our sum-of-squares rearrangement (in Claim 14), we need to select the centers (in $\mathcal{C}^i$) to be of norm $d\kappa$, which necessitates $\kappa \leq \frac{1}{d}$.

Proof of Theorem 13. Similar to the proof of Theorem 9, we reduce from $(\mathcal{F}^{\text{3Disj}}, \gamma)$ -promise sanitization problem. Assume for the sake of contradiction that there exists a polynomial-time $(\varepsilon, 1/n^{\omega(1)})$ -DP $(\mathcal{F}_2^{k\text{-means}}, \xi)$ -sanitizer $\mathcal{A}$, where $\xi = \frac{(1-\gamma)^2}{4d^2}$.

Let $\tilde{\mathcal{A}}$ be exactly the same as in the proof of Theorem 9 except that we set $x' = \kappa \cdot x$ where $\kappa = \frac{1}{d}$ (instead of $\kappa = 1/2$ as before), which ensures that $\|x'\|_2 \leq 1$. Again, $\tilde{\mathcal{A}}$ is $(\varepsilon, 1/n^{\omega(1)})$ -DP due to post-processing. We are left to show that $\tilde{\mathcal{A}}$ solves the $(\mathcal{F}^{\text{3Disj}}, \gamma)$ -promise sanitization problem. To do this, assume that $\psi(D) = 1$ for all $\psi \in \tilde{\mathcal{F}}$. Since $\mathcal{A}$ is an $(\mathcal{F}_2^{k\text{-means}}, \xi)$ -sanitizer, with probability at least $2/3$, the following holds for all $\mathcal{C} \in \mathbb{B}_2^k$:

$$|\text{cost}_{\mathcal{C},2}(\tilde{D}') - \text{cost}_{\mathcal{C},2}(D')| \leq \xi. \quad (2)$$

We will show that, when the above holds, we have $\psi(\tilde{D}) \geq \gamma$ for all $\psi \in \tilde{\mathcal{F}}$.

Consider any 3-literal disjunction ψ . We use the same notation (e.g., $\mathcal{C}^0, \mathcal{C}^\psi, c^{\psi,j}$) as in the proof of Theorem 9. Additionally, for each $i \in [d]$, let $\mathcal{C}^i = \{+d\kappa \cdot \mathbf{1}_i, -d\kappa \cdot \mathbf{1}_i\}$.

We begin by showing that most of the coordinates of $\tilde{x}' \in \tilde{D}'$ have to be close to $-\kappa$ or $+\kappa$. The exact bound is given below.

▷ **Claim 14.** $\mathbb{E}_{\tilde{x}' \in \tilde{D}'}[\|\tilde{x}' - \kappa \cdot \text{sgn}(\tilde{x}')\|_2^2] \leq \xi$.

Proof. For any $y \in \mathbb{R}^d$, observe that $\text{cost}_{\mathcal{C}^i,2}(y) = \|y - d\kappa \cdot \text{sgn}(y_i) \cdot \mathbf{1}_i\|^2 = (y_i - d\kappa \cdot \text{sgn}(y_i))^2 + \sum_{j \neq i} y_j^2$. Thus, aggregating this across all $i \in [d]$, we thus have

$$\begin{aligned} \sum_{i \in [d]} \text{cost}_{\mathcal{C}^i,2}(y) &= \sum_{i \in [d]} ((d-1)y_i^2 + (y_i - d\kappa \cdot \text{sgn}(y_i))^2) \\ &= \sum_{i \in [d]} ((d^2 - d)\kappa^2 + d(y_i - \kappa \cdot \text{sgn}(y_i))^2) \\ &= (d^3 - d^2)\kappa^2 + d \cdot \|y - \kappa \cdot \text{sgn}(y)\|_2^2. \end{aligned}$$

From (2) and since $x' = \kappa \cdot \text{sgn}(x)$ for all $x \in D$, we thus have

$$d \cdot \xi \geq \sum_{i \in [d]} (\text{cost}_{\mathcal{C}^i,2}(\tilde{D}') - \text{cost}_{\mathcal{C}^i,2}(D')) = \mathbb{E}_{\tilde{x}' \in \tilde{D}'}[d \cdot \|\tilde{x}' - \kappa \cdot \text{sgn}(\tilde{x}')\|_2^2].$$

Dividing both sides by d yields the claimed inequality. $\triangleleft$

The next two claims are in the same spirit as Claim 10 and Claim 11 in the proof of Theorem 9.

▷ **Claim 15.** For any $x' \in D'$ and $\psi \in \tilde{\mathcal{F}}$, we have $\text{cost}_{\mathcal{C}^\psi,2}(x') - \text{cost}_{\mathcal{C}^0,2}(x') \leq 0$.

Proof. From our assumption, there exists a literal of ψ , say the j th literal for $j \in [3]$, that is set to true in x . From the proof of Claim 10, we have $x' - c^{\psi,j} \in \{\pm\kappa\}^d$. Recall that x' also belongs to $\{\pm\kappa\}^d$. Thus, $\text{cost}_{\mathcal{C}^0,2}(x') = \|x'\|_2^2 = \kappa^2 \cdot d = \|x' - c^{\psi,j}\|_2^2 \geq \text{cost}_{\mathcal{C}^\psi,2}(x')$. $\triangleleft$

▷ **Claim 16.** For any $\tilde{x}' \in \tilde{D}'$ and $\psi \in \tilde{\mathcal{F}}$, we have

$$\text{cost}_{\mathcal{C}^\psi,2}(\tilde{x}') - \text{cost}_{\mathcal{C}^0,2}(\tilde{x}') \geq 4\kappa^2(1 - \psi(\tilde{x})) - 4\kappa \cdot \|\tilde{x}' - \kappa \cdot \text{sgn}(\tilde{x}')\|_2.$$

Proof. We have $\text{cost}_{\mathcal{C}^\psi,2}(\tilde{x}') = \|c^{\psi,j} - \tilde{x}'\|_2^2$ for some $j \in [3]$. Thus,

$$\begin{aligned} \text{cost}_{\mathcal{C}^\psi,2}(\tilde{x}') - \text{cost}_{\mathcal{C}^0,2}(\tilde{x}') &= \|c^{\psi,j} - \tilde{x}'\|_2^2 - \|\tilde{x}'\|_2^2 \\ &= \|c^{\psi,j}\|_2^2 - 2\langle c^{\psi,j}, \tilde{x}' \rangle \\ &= 4\kappa^2 - 4\kappa \cdot s_j \cdot \tilde{x}'_{i_j}. \end{aligned}$$

First, consider the case $\psi(\tilde{x}) = 1$. In this case, we can bound the above further by

$$\begin{aligned} \text{cost}_{\mathcal{C}^\psi,2}(\tilde{x}') - \text{cost}_{\mathcal{C}^0,2}(\tilde{x}') &\geq 4\kappa^2 - 4\kappa \cdot |\tilde{x}'_{i_j}| \\ &= -4\kappa \left(|\tilde{x}'_{i_j}| - \kappa \right) \\ &\geq -4\kappa \cdot |\tilde{x}'_{i_j} - \kappa \cdot \text{sgn}(\tilde{x}'_{i_j})| \\ &\geq -4\kappa \cdot \|\tilde{x}' - \kappa \cdot \text{sgn}(\tilde{x}')\|_2, \end{aligned}$$

which implies the desired bound.

Finally, for the case $\psi(\tilde{x}) = 0$, we simply have $s_j \cdot \tilde{x}'_{i_j} \leq 0$ and thus $\text{cost}_{\mathcal{C}^\psi,2}(\tilde{x}') - \text{cost}_{\mathcal{C}^0,2}(\tilde{x}') \geq 4\kappa^2$. $\triangleleft$

Now, from the above three claims, we have

$$\begin{aligned} 0 &\geq \mathbb{E}_{x' \sim D'} [\text{cost}_{\mathcal{C}^\psi,2}(x') - \text{cost}_{\mathcal{C}^0,2}(x')] \\ &= \text{cost}_{\mathcal{C}^\psi,2}(D') - \text{cost}_{\mathcal{C}^0,2}(D') \\ \text{Using (2)} &\geq -2\xi + \text{cost}_{\mathcal{C}^\psi,2}(\tilde{D}') - \text{cost}_{\mathcal{C}^0,2}(\tilde{D}') \\ &= -2\xi + \mathbb{E}_{\tilde{x}' \sim \tilde{D}'} [\text{cost}_{\mathcal{C}^\psi,2}(\tilde{x}') - \text{cost}_{\mathcal{C}^0,2}(\tilde{x}')] \\ \text{(Claim 16)} &\geq -2\xi + \mathbb{E}_{\tilde{x}' \sim \tilde{D}'} [4\kappa^2(1 - \psi(\tilde{x})) - 4\kappa \cdot \|\tilde{x}' - \kappa \cdot \text{sgn}(\tilde{x}')\|_2] \\ &= 4\kappa^2(1 - \psi(\tilde{D})) - 2\xi - 4\kappa \cdot \mathbb{E}_{\tilde{x}' \sim \tilde{D}'} [\|\tilde{x}' - \kappa \cdot \text{sgn}(\tilde{x}')\|_2] \\ \text{(Cauchy-Schwarz Inequality)} &\geq 4\kappa^2(1 - \psi(\tilde{D})) - 2\xi - 4\kappa \cdot \sqrt{\mathbb{E}_{\tilde{x}' \sim \tilde{D}'} [\|\tilde{x}' - \kappa \cdot \text{sgn}(\tilde{x}')\|_2^2]} \\ \text{(Claim 14)} &\geq 4\kappa^2(1 - \psi(\tilde{D})) - 2\xi - 4\kappa \cdot \sqrt{\xi} \\ &\geq 4\kappa^2 \left(1 - \frac{\xi}{2\kappa^2} - \frac{\sqrt{\xi}}{\kappa} - \psi(\tilde{D}) \right), \end{aligned}$$

where the first inequality is from Claim 15.

The above inequality implies that $\psi(\tilde{D}) \geq 1 - \frac{\xi}{2\kappa^2} - \frac{\sqrt{\xi}}{\kappa}$, which is at least γ due to our choice of $\xi = \frac{(1-\gamma)^2}{4d^2}$, $\kappa = \frac{1}{d}$. As a result, $\tilde{\mathcal{A}}$ solves the $(\mathcal{F}^{\text{3Disj}}, \gamma)$ -promise sanitization problem; by Theorem 5, this cannot hold assuming the existence of one-way functions. $\blacktriangleleft$

5 Conclusion and Open Questions

In this paper, we showed computational lower bounds for efficiently constructing DP coresets, assuming the one-way functions exist. Our result should be contrasted against the work of Feldman et al. [18], who showed that small private coresets exist information-theoretically.

Our work leaves open several questions for future research. While our hardness result for the ℓ_2 -metric rules out approximation factors $1 \pm \Theta(1/d^2)$, it is an interesting question if a dimension-independent constant-factor hardness (similar to our ℓ_∞ -metric result) can be obtained for the Euclidean metric. Additionally, our hardness results hold for $k = 3$ due to our “gadget” to reduce from 3-literal disjunctions, leaving the case $k = 2$ as an intriguing open question. We remark that, while Ullman and Vadhan [43] also proved the hardness

of generating synthetic data for 2-way marginals, their result does not achieve “perfect completeness” (in the promise sense that $f(D) = 1$ for all $f \in \mathcal{F}$ in Definition 4). This property is crucial in our reduction, preventing us from obtaining hardness for the $k = 2$ case.

Furthermore, our proofs rely on the fact that we are dealing with the k -means objective (by rearranging terms into sum-of-squares); it would be interesting to see if our results can be extended to the case of, e.g., k -median as well. Finally, understanding the computational landscape of private coresets in other metric spaces or for other objective functions is an interesting research direction.

References

- 1 John M. Abowd. The U.S. Census Bureau adopts differential privacy. In *KDD*, pages 6353–6356, 2018.
- 2 Sanjeev Arora, Carsten Lund, Rajeev Motwani, Madhu Sudan, and Mario Szegedy. Proof verification and the hardness of approximation problems. *J. ACM*, 45(3):501–555, 1998. doi:10.1145/278298.278306.
- 3 Sanjeev Arora and Shmuel Safra. Probabilistic checking of proofs: A new characterization of NP. *J. ACM*, 45(1):70–122, 1998. doi:10.1145/273865.273901.
- 4 Maria-Florina Balcan, Travis Dick, Yingyu Liang, Wenlong Mou, and Hongyang Zhang. Differentially private clustering in high-dimensional Euclidean spaces. In *ICML*, pages 322–331, 2017. URL: <http://proceedings.mlr.press/v70/balcan17a.html>.
- 5 Luca Becchetti, Marc Bury, Vincent Cohen-Addad, Fabrizio Grandoni, and Chris Schwiegelshohn. Oblivious dimension reduction for k -means: beyond subspaces and the Johnson–Lindenstrauss lemma. In *STOC*, pages 1039–1050, 2019. doi:10.1145/3313276.3316318.
- 6 Avrim Blum, Cynthia Dwork, Frank McSherry, and Kobbi Nissim. Practical privacy: the SuLQ framework. In *PODS*, pages 128–138, 2005. doi:10.1145/1065167.1065184.
- 7 Alisa Chang, Badih Ghazi, Ravi Kumar, and Pasin Manurangsi. Locally private k -means in one round. In *ICML*, pages 1441–1451, 2021. URL: <http://proceedings.mlr.press/v139/chang21a.html>.
- 8 Ke Chen. On coresets for k -median and k -means clustering in metric and Euclidean spaces and their applications. *SICOMP*, 39(3):923–947, 2009. doi:10.1137/070699007.
- 9 Vincent Cohen-Addad, Andrew Draganov, Matteo Russo, David Saulpic, and Chris Schwiegelshohn. A tight VC-dimension analysis of clustering coresets with applications. In *SODA*, pages 4783–4808, 2025. doi:10.1137/1.9781611978322.162.
- 10 Vincent Cohen-Addad and Karthik C. S. Inapproximability of clustering in lp metrics. In *FOCS*, pages 519–539, 2019. doi:10.1109/FOCS.2019.00040.
- 11 Vincent Cohen-Addad, Kasper Green Larsen, David Saulpic, and Chris Schwiegelshohn. Towards optimal lower bounds for k -median and k -means coresets. In *STOC*, pages 1038–1051, 2022. doi:10.1145/3519935.3519946.
- 12 Vincent Cohen-Addad, Kasper Green Larsen, David Saulpic, Chris Schwiegelshohn, and Omar Ali Sheikh-Omar. Improved coresets for Euclidean k -means. In *NeurIPS*, 2022.
- 13 Vincent Cohen-Addad, David Saulpic, and Chris Schwiegelshohn. A new coreset framework for clustering. In *STOC*, pages 169–182, 2021. doi:10.1145/3406325.3451022.
- 14 Cynthia Dwork, Krishnaram Kenthapadi, Frank McSherry, Ilya Mironov, and Moni Naor. Our data, ourselves: Privacy via distributed noise generation. In *EUROCRYPT*, pages 486–503, 2006. doi:10.1007/11761679_29.
- 15 Cynthia Dwork, Frank McSherry, Kobbi Nissim, and Adam Smith. Calibrating noise to sensitivity in private data analysis. In *TCC*, pages 265–284, 2006. doi:10.1007/11681878_14.
- 16 Cynthia Dwork, Moni Naor, Omer Reingold, Guy N Rothblum, and Salil Vadhan. On the complexity of differentially private data release: efficient algorithms and hardness results. In *STOC*, pages 381–390, 2009. doi:10.1145/1536414.1536467.

- 17 Úlfar Erlingsson, Vasyl Pihur, and Aleksandra Korolova. RAPPOR: Randomized aggregatable privacy-preserving ordinal response. In *CCS*, pages 1054–1067, 2014. doi:10.1145/2660267.2660348.
- 18 Dan Feldman, Amos Fiat, Haim Kaplan, and Kobbi Nissim. Private coresets. In *STOC*, pages 361–370, 2009. doi:10.1145/1536414.1536465.
- 19 Dan Feldman and Michael Langberg. A unified framework for approximating and clustering data. In *STOC*, pages 569–578, 2011. doi:10.1145/1993636.1993712.
- 20 Dan Feldman, Melanie Schmidt, and Christian Sohler. Turning big data into tiny data: Constant-size coresets for k -means, PCA, and projective clustering. *SICOMP*, 49(3):601–657, 2020. doi:10.1137/18M1209854.
- 21 Dan Feldman, Chongyuan Xiang, Ruihao Zhu, and Daniela Rus. Coresets for differentially private k -means clustering and applications to privacy in mobile sensor networks. In *IPSN*, pages 3–16, 2017. doi:10.1145/3055031.3055090.
- 22 Hendrik Fichtenberger, Marc Gillé, Melanie Schmidt, Chris Schwiegelshohn, and Christian Sohler. BICO: BIRCH meets coresets for k -means clustering. In *ESA*, pages 481–492, 2013. doi:10.1007/978-3-642-40450-4_41.
- 23 Gereon Frahling and Christian Sohler. Coresets in dynamic geometric data streams. In *STOC*, pages 209–217, 2005. doi:10.1145/1060590.1060622.
- 24 Badih Ghazi, Junfeng He, Kai Kohlhoff, Ravi Kumar, Pasin Manurangsi, Vidhya Navalpakkam, and Nachiappan Valliappan. Differentially private heatmaps. In *AAAI*, pages 7696–7704, 2023. doi:10.1609/AAAI.V37I6.25933.
- 25 Badih Ghazi, Pritish Kamath, Alexander Knop, Ravi Kumar, Pasin Manurangsi, and Chiyuan Zhang. Private hyperparameter tuning with ex-post guarantee. In *NeurIPS*, 2025.
- 26 Badih Ghazi, Ravi Kumar, and Pasin Manurangsi. Differentially private clustering: Tight approximation ratios. In *NeurIPS*, 2020.
- 27 Oded Goldreich. *The Foundations of Cryptography - Volume 1: Basic Techniques*. Cambridge University Press, 2001. doi:10.1017/CB09780511546891.
- 28 Oded Goldreich. *The Foundations of Cryptography - Volume 2: Basic Applications*. Cambridge University Press, 2004. doi:10.1017/CB09780511721656.
- 29 Anupam Gupta, Katrina Ligett, Frank McSherry, Aaron Roth, and Kunal Talwar. Differentially private combinatorial optimization. In *SODA*, pages 1106–1125, 2010. doi:10.1137/1.9781611973075.90.
- 30 Sarel Har-Peled and Akash Kushal. Smaller coresets for k -median and k -means clustering. *Discret. Comput. Geom.*, 37(1):3–19, 2007. doi:10.1007/S00454-006-1271-X.
- 31 Sarel Har-Peled and Soham Mazumdar. On coresets for k -means and k -median clustering. In *STOC*, pages 291–300, 2004. doi:10.1145/1007352.1007400.
- 32 Lingxiao Huang and Nisheeth K. Vishnoi. Coresets for clustering in Euclidean spaces: importance sampling is nearly optimal. In *STOC*, pages 1416–1429, 2020. doi:10.1145/3357713.3384296.
- 33 Zhiyi Huang and Jinyan Liu. Optimal differentially private algorithms for k -means clustering. In *PODS*, pages 395–408, 2018. doi:10.1145/3196959.3196977.
- 34 Max Dupré la Tour, Monika Henzinger, and David Saulpic. Making old things new: A unified algorithm for differentially private clustering. In *ICML*, 2024.
- 35 Michael Langberg and Leonard J. Schulman. Universal epsilon-approximators for integrals. In *SODA*, pages 598–607, 2010. doi:10.1137/1.9781611973075.50.
- 36 Jingcheng Liu and Kunal Talwar. Private selection from private candidates. In *STOC*, pages 298–309, 2019. doi:10.1145/3313276.3316377.
- 37 Frank McSherry and Kunal Talwar. Mechanism design via differential privacy. In *FOCS*, pages 94–103, 2007. doi:10.1109/FOCS.2007.66.
- 38 Kobbi Nissim and Uri Stemmer. Clustering algorithms for the centralized and local models. In *ALT*, pages 619–653, 2018. URL: <http://proceedings.mlr.press/v83/nissim18a.html>.

- 39 Kobbi Nissim, Uri Stemmer, and Salil P. Vadhan. Locating a small cluster privately. In *PODS*, pages 413–427, 2016. doi:10.1145/2902251.2902296.
- 40 Nicolas Papernot and Thomas Steinke. Hyperparameter tuning with Rényi differential privacy. In *ICLR*, 2022.
- 41 Christian Sohler and David P. Woodruff. Strong coresets for k -median and subspace approximation: Goodbye dimension. In *FOCS*, pages 802–813, 2018. doi:10.1109/FOCS.2018.00081.
- 42 Uri Stemmer. Locally private k -means clustering. In *SODA*, pages 548–559, 2020. doi:10.1137/1.9781611975994.33.
- 43 Jonathan R. Ullman and Salil P. Vadhan. PCPs and the hardness of generating synthetic data. *J. Cryptol.*, 33(4):2078–2112, 2020. doi:10.1007/S00145-020-09363-Y.

A

 Hardness from [43] in terms of Promise Sanitizer

In this section, we briefly explain how to interpret the hardness result of Ullman and Vadhan [43] in terms of a promise sanitizer (Theorem 5).

Informally, the proof is by contradiction, assuming that an efficient DP promise sanitizer exists. The idea is to construct a dataset by taking a set of valid digital signatures and encoding them using a PCP of the signature’s verification circuit. Now, if the sanitizer can produce a synthetic dataset that satisfies a good fraction of the PCP predicates, the PCP’s soundness guarantee means we can decode a valid signature from the synthetic dataset. But, this yields a contradiction: either the decoded signature is new, which would violate the security of the super-secure signature scheme, or it matches one of the original signatures, which would violate the DP guarantees.

To avoid repeating their proof verbatim, we will only sketch the high-level arguments here.

A.1 Background

We review a couple of key tools required for the construction of [43].

Super-Secure Digital Signature Scheme

A *Digital Signature Scheme (DSS)* consists of three polynomial-time algorithms:

- **Gen** takes in the security parameter 1^κ and output a secret key $\mathbf{sk}$ and a verification key $\mathbf{vk}$,
- **Sign** takes in⁸ a secret key $\mathbf{sk}$ and outputs a (random) signature σ ,
- **Ver** takes in the verification key $\mathbf{vk}$ and a signature σ . For any valid signature σ produced by **Sign**, $\mathbf{Ver}(\mathbf{vk}, \sigma)$ always accepts (i.e., returns 1).

The security guarantee is that, an adversary that is given the verification key $\mathbf{vk}$ together with a set Σ of any polynomially large number of signatures, cannot produce a new signature $\sigma^* \notin \Sigma$ that is accepted by **Ver** with at least $1/n^{\Omega(1)}$ probability.

It is known that super-secure digital signature schemes exist, assuming one-way functions [28].

⁸ We use a “no-message” version of the signature scheme, which can be initiated from the standard notion by setting the message m to be fixed (e.g., empty string) for every signature.

Probabilistic Checkable Proofs

Another ingredient required in the construction is the Probabilistic Checkable Proofs (PCPs) under Levin reductions. We state this only for Max-3SAT.

PCPs with $\mathcal{F}^{3\text{Disj}}$ predicates consist of three (deterministic) algorithms:

- **R** takes in the circuit C and outputs a set $\tilde{\mathcal{F}} \subseteq \mathcal{F}_d^{3\text{Disj}}$ of predicates
- **Enc** takes in any circuit C and an input y for the circuit C , and produces $\pi \in \{0, 1\}^d$,
- **Dec** takes in $\pi \in \{0, 1\}^d$ and C and produces an input y' for the circuit C .

The guarantees are as follows:

- (Completeness) If $C(y) = 1$, then $\psi(\text{Enc}(y, C)) = 1$ for all $\psi \in \tilde{\mathcal{F}}$,
- (Soundness) If $\mathbb{E}_{\psi \sim \tilde{\mathcal{F}}}[\psi(\pi)] \geq \gamma$, then $C(\text{Dec}(\pi, C)) = 1$.

In the papers that originally proved the PCP theorem [2, 3], it was also shown that PCPs of the above form exist for some constant $\gamma \in (0, 1)$.

A.2 The Construction and Hardness Argument

Formally, the construction of [43] (which builds on [16]) works as follows:

- Let sk, vk be the key pair generated by **Gen**.
- Let $\sigma_1, \dots, \sigma_n$ be outputs from n runs of **Sign**(sk).
- Let C_{vk} denote the circuit for **Ver**($\text{vk}, \cdot$).
- Let $\tilde{\mathcal{F}} = \mathbf{R}(C_{\text{vk}})$.
- For $i \in [n]$, let $\pi_i = \text{Enc}(\sigma_i, C_{\text{vk}})$.
- Output the dataset $D = (\pi_1, \dots, \pi_n)$ together with $\tilde{\mathcal{F}}$.

Now, suppose for the sake of contradiction that there is an $(\varepsilon, 1/n^{\omega(1)})$ -DP $(\mathcal{F}^{3\text{Disj}}, \gamma)$ -promise sanitizer that runs in polynomial time. Then, we can run the sanitizer on D to produce a synthetic dataset $\tilde{D} = (\tilde{x}_1, \dots, \tilde{x}_{\tilde{n}})$. Now, from the completeness of the PCP, we have $\mathbb{E}_{\psi \sim \tilde{\mathcal{F}}}[\psi(D)] = 1$. Thus, by the guarantee of the promise sanitizer, with probability $2/3$, we have $\mathbb{E}_{\psi \sim \tilde{\mathcal{F}}}[\psi(\tilde{D})] \geq \gamma$. When this happens, there must exist $\tilde{x}_i \in \tilde{D}$ such that $\mathbb{E}_{\psi \sim \tilde{\mathcal{F}}}[\psi(\tilde{x}_i)] \geq \gamma$. Finally, we can run **Dec**($\tilde{x}_i, C_{\text{vk}}$) to get a new signature σ^* .

By the guarantee of the PCP, we have $C_{\text{vk}}(\sigma^*) = 1$, i.e., **Ver**(vk, σ^*) = 1. Now, one of the following two cases can happen: (i) if $\sigma^* \notin \{\sigma_1, \dots, \sigma_n\}$, then we violate the security of the signature scheme, or (ii) if $\sigma^* \in \{\sigma_1, \dots, \sigma_n\}$, then $\pi^* = \text{Enc}(\sigma^*, C_{\text{vk}})$ belongs to the input dataset and we violate $(\varepsilon, 1/n^{\omega(1)})$ -DP.