

Sketching Data Sets for Large- Scale Learning

©SHUTTERSTOCK.COM/SYDA PRODUCTIONS

Keeping only what you need

Big data can be a blessing: with very large training data sets it becomes possible to perform complex learning tasks with unprecedented accuracy. Yet, this improved performance comes at the price of enormous computational challenges. Thus, one may wonder: Is it possible to leverage the information content of huge data sets while keeping computational resources under control? Can this also help solve some of the privacy issues raised by large-scale learning? This is the ambition of compressive learning, where the data set is massively compressed before learning. Here, a “sketch” is first constructed by computing carefully chosen nonlinear random features [e.g., random Fourier (RF) features] and averaging them over the whole data set. Parameters are then learned from the sketch, without access to the original data set. This article surveys the current state of the art in compressive learning, including the main concepts and algorithms,

their connections with established signal processing methods, existing theoretical guarantees on both information preservation and privacy preservation, and important open problems. For an extended version of this article that contains additional references and more in-depth discussions on a variety of topics, see [1].

Introduction to compressive learning

The overall principle of compressive learning is summarized in Figure 1. During the sketching phase, a potentially huge collection of d -dimensional data vectors $\{\mathbf{x}_i\}_{i=1}^n$ is summarized into a single m -dimensional vector $\tilde{\mathbf{z}}$ called the *sketch*. The sketch is constructed by transforming each data vector and then averaging the results:

$$\tilde{\mathbf{z}} := \frac{1}{n} \sum_{i=1}^n \Phi(\mathbf{x}_i). \quad (1)$$

Next, during the learning phase, an estimate $\hat{\boldsymbol{\theta}}$ of some essential statistical parameters $\boldsymbol{\theta}$ of the data set are extracted from

the sketch $\tilde{z}$. These parameters are application dependent. For example, as illustrated in Figure 2, they could represent principal data subspaces (in subspace fitting problems), prediction weights (in estimation and regression problems), distributional parameters (in density estimation problems), or centroids (in clustering problems); we give detailed examples in the following. Python notebooks that illustrate examples from the article are available at <https://gitlab.inria.fr/SketchedLearning/spm-notebook>.

The transformation $\Phi(\cdot)$, known as the *feature map*, is generally nonlinear and randomized. Although the use of $\Phi(\cdot)$ is related to “kernel methods” in machine learning (e.g., [2] and [3]), as will be discussed later, the act of sketching (1) also includes averaging over the n data samples. The advantages of compressive learning are the following:

- 1) By choosing a sketch of dimension $m \ll nd$, the data get massively compressed. This has obvious advantages for storage and transfer.
- 2) Sketching can speed up the learning phase, whose complexity becomes independent of the cardinality, n , of the original data set. This enables one to handle massive data sets while keeping computational resources under control.

- 3) Sketching can preserve privacy: the transformation $\Phi(\cdot)$ can be chosen so that individual user information is lost while aggregate user information is preserved.
- 4) The sketching mechanism in (1) is well matched to distributed implementations and streaming scenarios: the sketch of a concatenation of data sets is a simple mean of the sketches of those data sets.

Illustration using four worked examples

To illustrate the compressive learning framework and discuss the essential aspects of it, we now outline four canonical examples of machine learning tasks to which sketching can be readily applied. See Figure 2 and Table 1 for an illustration.

Principal component analysis

Principal component analysis (PCA) seeks to find the linear subspace of a fixed dimension $k < d$ that best fits the d -dimensional data $\{x_i\}_{i=1}^n$ in the least-squares (LS) sense. In this case, we assume the data to be centered so that the target parameters θ can be described by a k -dimensional orthonormal basis $\{u_\ell\}_{\ell=1}^k$ that maximizes $\sum_{\ell=1}^k \sum_{i=1}^n |u_\ell^T x_i|^2$.

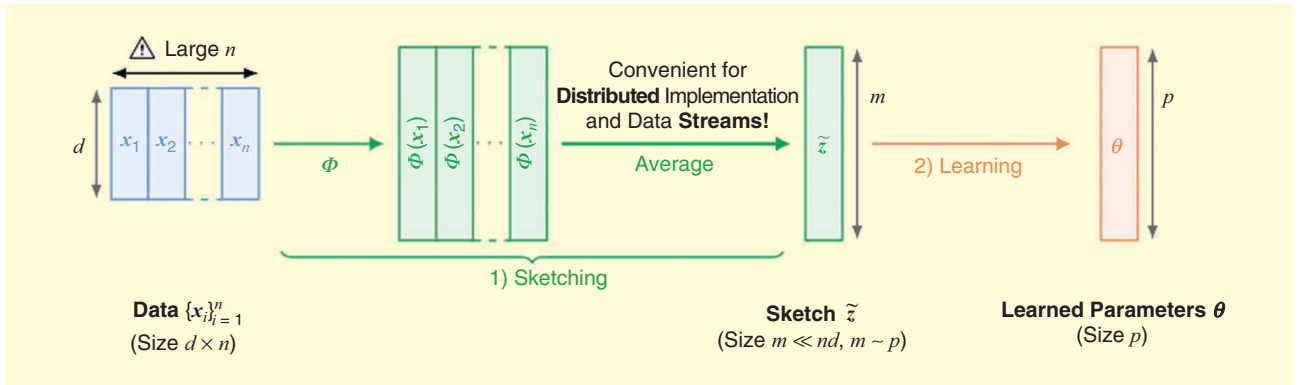


FIGURE 1. An overview of sketching and parameter learning.

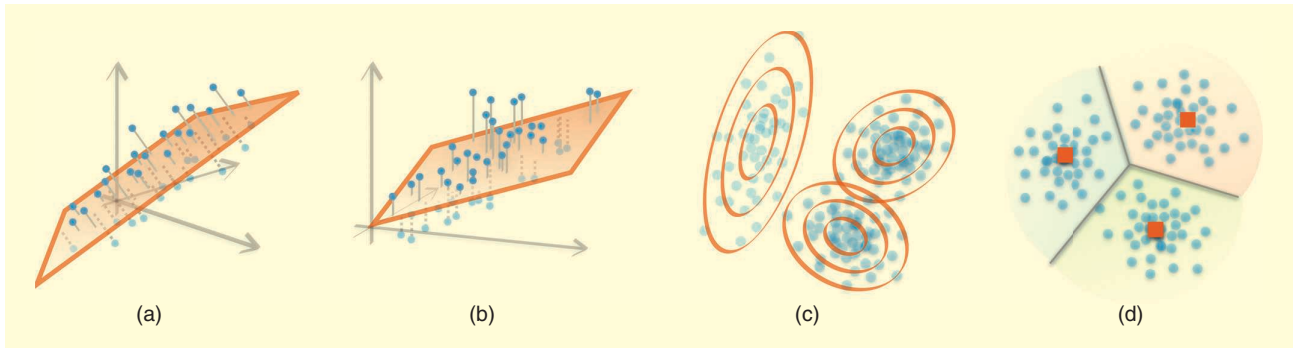


FIGURE 2. A schematic representation of the four running examples covered by this article. In the four figures, each x_i is associated with a blue colored disk; hence, the training collection corresponds to a point cloud, and the orange color geometrically represents the learned parameters. (a) Principal component analysis (PCA) learns the principal k -dimensional subspace of the data set for some $k \leq d$. (b) Linear regression fits observed data (the blue dot heights) as a linear model of the inputs (here, the 2D horizontal coordinates). For least squares, this amounts to minimizing the square of the differences between these data and the linear predictions. (c) In Gaussian mixture modeling, we learn the set of parameters (the mixture weight, mean, and covariance) characterizing each Gaussian term of the mixture; the probability-level sets are displayed here as orange ellipses. (d) Clustering methods (such as k -means) learn a set of centroids (the orange squares) defining a Voronoi partition grouping together similar data samples.

It is well known that one solution is given by the k principal eigenvectors of the empirical autocorrelation matrix $\hat{\mathbf{R}} = (1/n) \sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i^\top$. This $\hat{\mathbf{R}}$ can be interpreted as a sketch of the form (1) that uses the feature map $\Phi(\mathbf{x}) = \text{vec}(\mathbf{x} \mathbf{x}^\top) = (x_1 x_1, x_2 x_1, \dots, x_d x_d)^\top$ of dimension $m = d^2$. Here and in the sequel, the $\text{vec}(\cdot)$ operator vectorizes a matrix by stacking its columns. As soon as the cardinality n of the data set exceeds the dimension d , the dimension of this sketch is smaller than nd , the total size of the data set. Using techniques from matrix completion and compressive matrix sensing, the sketch dimension can be further reduced [4]. For example, one can sketch via (1) using $\Phi(\mathbf{x}) = ((\mathbf{w}_1^\top \mathbf{x})^2, \dots, (\mathbf{w}_m^\top \mathbf{x})^2)^\top$, where the d -dimensional vectors $\mathbf{w}_j$, $1 \leq j \leq m$ are the rows of a tall random matrix $\mathbf{W}$ of size $m \times d$. For m on the order of $kd < d^2$, low-rank matrix reconstruction techniques can estimate the k -dimensional principal subspace of $\hat{\mathbf{R}}$ with provable accuracy [5], [51].

LS linear regression

Suppose that the data vectors take the form $\mathbf{x}_i = [\mathbf{x}_{1i}^\top, \mathbf{x}_{2i}^\top]^\top$, and our goal is to linearly predict the d_1 -dimensional vector $\mathbf{x}_{1i}$ from the d_2 -dimensional vector $\mathbf{x}_{2i}$ (with $d = d_1 + d_2$). That is, we want to design a $d_1 \times d_2$ weight matrix $\boldsymbol{\theta}$ such that $\mathbf{x}_{1i} \approx \boldsymbol{\theta} \mathbf{x}_{2i}$ for all samples i . The LS approach to this supervised learning problem chooses $\hat{\boldsymbol{\theta}} = \arg\min_{\boldsymbol{\theta}} \sum_{i=1}^n \|\mathbf{x}_{1i} - \boldsymbol{\theta} \mathbf{x}_{2i}\|^2$.

Although it is possible to compute the LS solution using gradient descent, each iteration would involve the full data set

$\{\mathbf{x}_i\}_{i=1}^n$. A well-known alternative is to first build the empirical autocorrelation matrix $\hat{\mathbf{R}} := (1/n) \sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i^\top$ and then compute $\hat{\boldsymbol{\theta}}$ in closed form as

$$\hat{\boldsymbol{\theta}} = \hat{\mathbf{R}}_{12} \hat{\mathbf{R}}_{22}^{-1}, \text{ with } \begin{pmatrix} \hat{\mathbf{R}}_{11} & \hat{\mathbf{R}}_{12} \\ \hat{\mathbf{R}}_{21} & \hat{\mathbf{R}}_{22} \end{pmatrix} = \hat{\mathbf{R}}, \quad (2)$$

where each submatrix $\hat{\mathbf{R}}_{ij}$ has dimension $d_i \times d_j$. Here, we again use the sketch (1) with the feature map $\Phi(\mathbf{x}) = \text{vec}(\mathbf{x} \mathbf{x}^\top)$ and extract the target parameter $\hat{\boldsymbol{\theta}}$ from (2).

Gaussian mixture modeling

Here, the objective is to find the parameters $\boldsymbol{\theta} = \{\alpha_\ell, \boldsymbol{\mu}_\ell, \boldsymbol{\Sigma}_\ell\}_{\ell=1}^k \subset \boldsymbol{\Theta}$ that best fit a k -term Gaussian mixture model (GMM) $p(\mathbf{x} | \boldsymbol{\theta}) = \sum_{\ell=1}^k \alpha_\ell \mathcal{N}(\mathbf{x}; \boldsymbol{\mu}_\ell, \boldsymbol{\Sigma}_\ell)$ to the data $\{\mathbf{x}_i\}_{i=1}^n$. The parameter space $\boldsymbol{\Theta}$ demands that, for each ℓ : $\alpha_\ell > 0$, $\boldsymbol{\mu}_\ell$ is a d -dimensional centroid, $\boldsymbol{\Sigma}_\ell$ is a $d \times d$ positive definite matrix, and $\sum_{\ell} \alpha_\ell = 1$. Traditionally, expectation–maximization (EM) is used to approximate the maximum likelihood (ML) estimate $\hat{\boldsymbol{\theta}} = \arg\max_{\boldsymbol{\theta} \in \boldsymbol{\Theta}} \sum_{i=1}^n \log p(\mathbf{x}_i | \boldsymbol{\theta})$, but the EM algorithm processes all n data samples $\{\mathbf{x}_i\}_{i=1}^n$ at each iteration, which can be computationally burdensome when n is very large.

An alternative [6] is to compute a sketch $\tilde{\mathbf{z}}$ of the form (1) using RF features [7]:

$$\Phi(\mathbf{x}) = \exp(-j2\pi \mathbf{W} \mathbf{x}), \quad (3)$$

where $j = \sqrt{-1}$ and $\mathbf{W} \in \mathbb{R}^{m \times d}$ is a realization of a random matrix [e.g., i.i.d. Gaussian]. Here, $\exp(\cdot)$ is the componentwise exponential, not the matrix exponential. We will have a lot to say about the RF feature map later in this article (for example, the connection between the RF map and the Gaussian kernel is discussed in the sidebar “Approximating Kernel Methods With Random Feature Maps”). The parameters can then be extracted by optimizing a cost function, as explained later. This approach has been applied to audio source separation as well as speaker verification [6], where it was shown that 1,000 h of speech can be compressed down to a few kilobytes without loss of verification performance (see the sidebar “Compressive Clustering of Modified National Institute of Standards and Technology Digits”).

k-means clustering

The goal of clustering is to group together “similar” data samples from $\{\mathbf{x}_i\}_{i=1}^n$. In the k -means approach to clustering, one aims to find the set of k centroids $\{\mathbf{c}_\ell\}_{\ell=1}^k$ that minimizes the average squared distance from each sample to its nearest centroid; i.e., $\sum_{i=1}^n \min_{\ell} \|\mathbf{x}_i - \mathbf{c}_\ell\|^2$. The famous Lloyd algorithm [8] is typically used in an attempt to solve this problem. When n is very large, however, Lloyd’s algorithm becomes computationally demanding. Instead, one could sketch the data set using (1) and extract the centroids from the sketch [9], [10]. For this purpose, one could use RF features (3) and (as explained in the sequel) solve for the centroids $\{\mathbf{c}_\ell\}_{\ell=1}^k$ and the (nonnegative, sum-to-one) weights $\{\alpha_\ell\}_{\ell=1}^k$ that minimize

Table 1. The notations used in this article.

Notation	
$\mathcal{X} = \{\mathbf{x}_i\}_{i=1}^n$	Data set of n training samples $\mathbf{x}_i \in \mathbb{R}^d$
$\mathbf{z}$	(Empirical) sketch of the data set, a vector of dimension m
$\Phi(\cdot)$	Sketching feature map, a function mapping $\mathbb{R}^d$ to either $\mathbb{R}^m$ or $\mathbb{C}^m$
$\boldsymbol{\theta} \in \boldsymbol{\Theta}$	Target parameters to learn, of dimension p (e.g., principal component analysis matrix, centroids, and mixture model)
$\mathbf{W}$	Random matrix associated with certain feature maps
$\Phi(\cdot)$	Usually drawn i.i.d. Gaussian
$\varrho(\cdot)$	Componentwise nonlinearity associated with certain feature maps $\Phi(\cdot)$
w.h.p.	With high probability, i.e., with exponentially decaying failure probability
i.i.d.	Independent identically distributed
w.p.1	With probability 1
$\ \cdot\ , \langle \cdot, \cdot \rangle$	Euclidean norm of a vector, inner product between vectors
$\ \cdot\ _0$	ℓ^0 “norm” of a vector, i.e., its number of nonzero entries
$\ \cdot\ _F$	Frobenius norm of a matrix
$\ \cdot\ _{L^2}, \langle \cdot, \cdot \rangle_{L^2}$	L^2 norm of a function, L^2 inner product between functions.

$\|\tilde{z} - \sum_{\ell=1}^k \alpha_{\ell} \Phi(c_{\ell})\|$. In this setting, the weights allow us to model unbalanced clusters, yet only the centroids need to be recovered. On large data sets, this approach can be orders of magnitude better than Lloyd's algorithm in memory and runtime, provided that the sketch dimension m is large enough, i.e., that m is on the order of kd , where kd is the number of free parameters in $\{c_{\ell}\}_{\ell=1}^k$ [9], [10]. For example, this method allows us to cluster the Modified National Institute of Standards and Technology (MNIST) digit data set, of dimension $d = 784$ and cardinality $n = 70,000$, using a complex-valued sketch of dimension $m = 400$ (see the sidebar "Compressive Clustering of MNIST Digits").

Although this article focuses on these examples, the general compressive learning framework extends beyond them. It has been applied to, for example, independent component analysis [11] (by sketching the cumulant tensor) and subspace clustering [12] (using a polynomial feature map). Extending the framework to a new task raises essential questions, such as: How do we choose the feature map $\Phi(\cdot)$? How do we learn the essential parameters θ from the sketch $\tilde{z}$? Are there statistical learning guarantees? Can sketching and learning be made computationally efficient? Can we sketch while respecting privacy? This article sketches answers to these questions for the worked examples introduced in the preceding.

Historical background

The term *sketch* has different meanings depending on the field. Our use of the term comes from the literature on relational databases and more particularly from the subfield of approximate query processing (AQP) [13]. The goal of AQP is to build a short description of the content of a massive data set, called a *synopsis*, by analogy with the synopsis of a movie or a book,

such that certain queries can be efficiently performed to return answers with a controlled error and/or probability of failure. Well-known queries include the frequency of occurrence of a particular element in a stream of data (taken from a discrete collection) and the minimum of several of these frequencies, yielding the celebrated count-min sketch [13, Sec. 5.3.1] synopsis. In this context, the statistical parameters θ learned from a sketch are interpreted as the result of a particular query on the data set.

Despite the apparent similarities between AQP and compressive learning, both data types and learning tasks differ significantly between the two fields. Typically, AQP focuses on (multi)sets of elements taken from a discrete collection of objects and considers database queries and related operations, while compressive learning focuses on continuous-valued signals (e.g., images or audio signals) and considers machine learning tasks, such as density estimation or regression.

Sketches for streaming and distributed methods

In AQP, a popular class of synopses is that of linear sketches [13, Ch. 5]. A linear sketch, which maps a data set to a vector, must satisfy the single condition that the sketch of the concatenation of two data sets equals the sum of their sketches. In mathematical terms, if we denote by $\tilde{z}(\mathcal{X})$ the sketch of a data set $\mathcal{X} = \{\mathbf{x}_i\}_{i=1}^n$, then it is required that $\tilde{z}(\mathcal{X}) = \tilde{z}(\mathcal{X}_1) + \tilde{z}(\mathcal{X}_2)$ whenever $\mathcal{X}$ is the concatenation of $\mathcal{X}_1$ and $\mathcal{X}_2$. It thus follows that a linear sketch must be of the form $\tilde{z}(\mathcal{X}) = \sum_{i=1}^n \Phi(\mathbf{x}_i)$ for some possibly nonlinear feature map $\Phi(\cdot)$. That is the same definition that we adopt in (1) except that we normalize by the number of elements n . Linear sketches are popular in AQP mainly because they are well suited to streaming scenarios. That is,

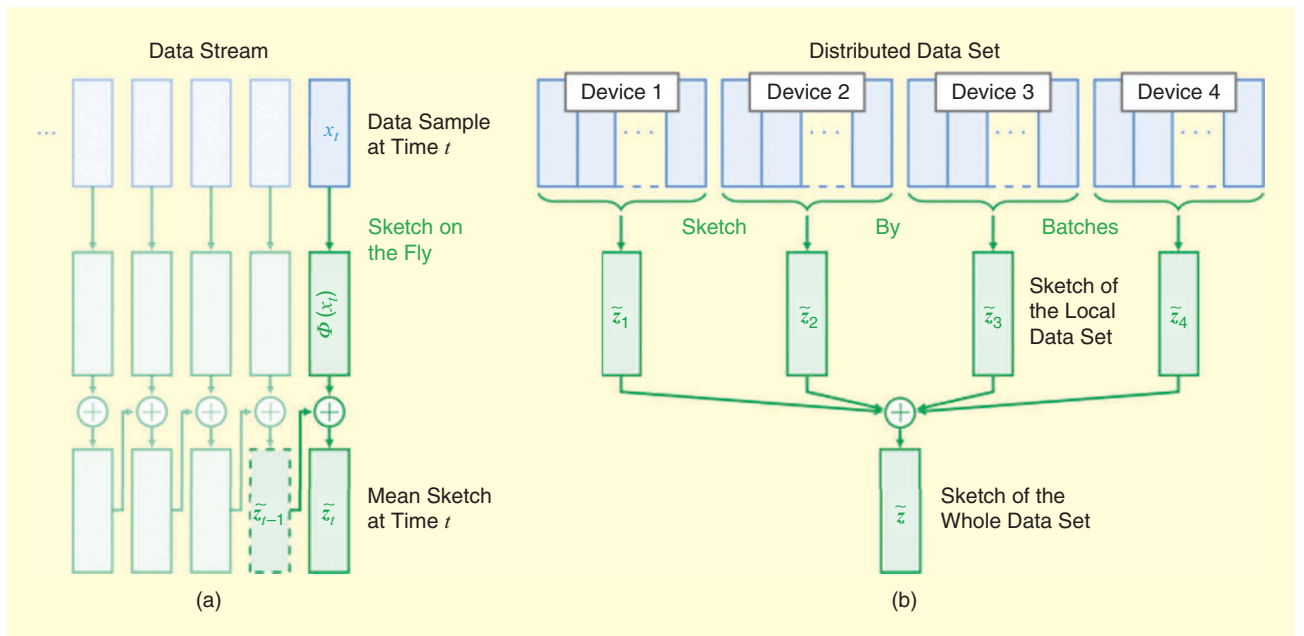


FIGURE 3. (a) A streaming scenario, where the data samples are sketched one by one and the mean sketch is updated at each time. (b) A distributed scenario, where each device computes a local sketch and a centralized entity further averages these local sketches.

inserting or deleting an element x_i from the data set corresponds to adding or subtracting $\Phi(x_i)$ from the sketch, respectively (see Figure 3). Note that some very simple sketches are not linear: for instance, it is easy to sketch the maximum running value in a stream of scalar data by computing the maximum of the current sketch and each new data point, but this sketching procedure is not linear [13, Ch. 5.2.1]. In particular, this “max” sketch facilitates

the insertion of a new element in the database but not the deletion of an existing one.

In the sections that follow, we describe linear sketches from other points of view. Before doing that, however, we want to clarify that the terms *sketching* and *linear sketching* appear in many other fields, although with meanings that differ considerably from ours. For instance, *sketching* is often used to indicate linear dimensionality reduction in n or d , as

Compressive Mixture Modeling for Speaker Verification

Given a fragment of speech and a candidate speaker, the goal of speaker verification is to assess whether the fragment was indeed spoken by that person. A classic approach to speaker verification is the Gaussian mixture model (GMM)–universal background model (UBM) [S1].

There, the idea is to train a model of a “universal” speaker from unlabeled training data and then compare it to a specialized model for the candidate speaker. Using the speech fragment, the likelihood ratio between the specialized and universal models is computed, and a positive

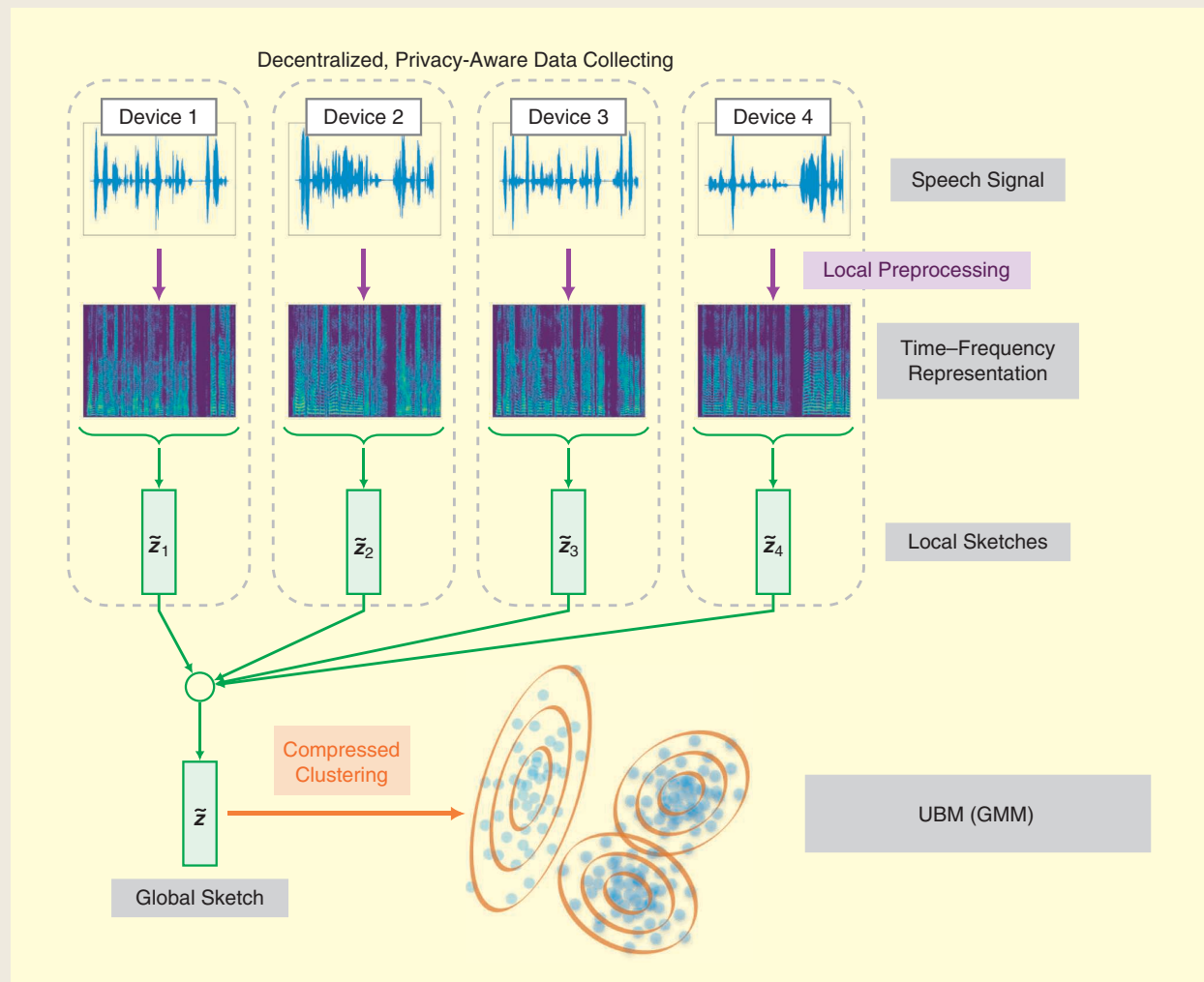


FIGURE S1. Speaker verification via compressive learning. Unlabeled training speech data are collected in a decentralized manner, preprocessed, and locally sketched. The local sketches are then merged into a global sketch, from which a UBM is learned.

attained by multiplying the data matrix $[\mathbf{x}_1, \dots, \mathbf{x}_n]^T \in \mathbb{R}^{n \times d}$ by a random matrix on the left or the right [14]–[17]. *Sketching* may also refer to the use of sampling-based approaches [18], such as core sets [19] or the Nyström method [20]. These latter methods differ from linear sketches in AQP, which is our notion of sketching, in that these latter methods generally do not respect the concatenation condition described earlier and are therefore less directly amenable to streaming

scenarios. Moreover, these methods typically do not involve a nonlinear feature map $\Phi(\cdot)$, which is a key component of our sketch.

Sketches as (randomized) generalized moments

In signal processing and machine learning, the data samples $\mathbf{x}_i$ generally live in the vector space $\mathbb{R}^d$ and are often modeled as i.i.d. random vectors having a probability distribution with

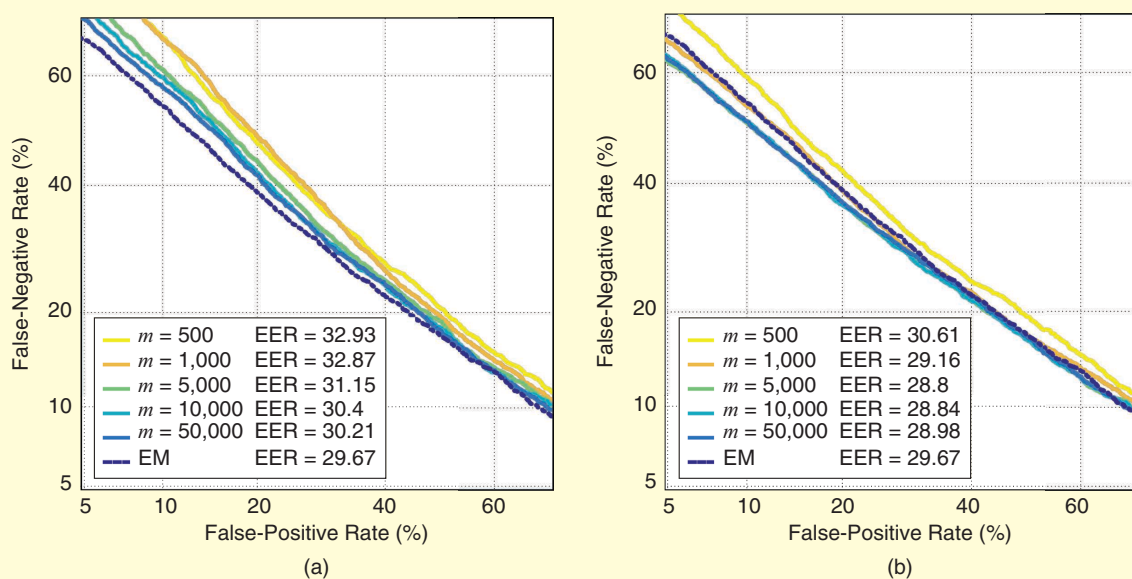


FIGURE S2. A detection error tradeoff curve (the false-positive/false-negative tradeoff obtained by varying the decision threshold) and an equal error rate (EER). On a single laptop, the EM was trained on 5 h of speech; the compressive GMM estimation used (a) 5 h and (b) 1,000 h for different sketch dimensions m .

decision is made if the ratio exceeds a certain threshold. In the GMM-UBM, the data are modeled using a GMM. That is, $p(\mathbf{x}|\boldsymbol{\theta})$ are multivariate Gaussian distributions applied to a suitable time–frequency transform of the raw audio signal [S1] (see Figure S1).

The training of the UBM is computationally demanding since it must be done on a large corpus of speech data. Moreover, the latter must be collected in a wide range of situations so that the final model may be as universal as possible. For this reason, the data collecting process is best performed in a decentralized manner. Finally, speech data collected in real-life situations is known to be sensitive information. For all of these reasons, compressive learning is well suited to speaker verification.

In [6], using a random Fourier feature map $\Phi(\cdot)$, the authors compressed 1,000 h of speech data (50 giga-

bytes) into a sketch of a few kilobytes on a single laptop. Subsequently, they performed GMM estimation, i.e., learning, using a greedy algorithm, thereby tackling the optimization problem (15) introduced in the main text. They compared this compressive learning approach to the expectation–maximization (EM) algorithm, which, on the same laptop, could only be trained using 5 h of speech given the available random-access memory. They observed that by enabling the use of more data within a fixed memory budget, sketching produced better results despite the tremendous compression factor (see Figure S2).

Reference

[S1] D. A. Reynolds, T. F. Quatieri, and R. B. Dunn, “Speaker verification using adapted gaussian mixture models,” *Dig. Sig. Proc.*, vol. 10, no. 1, pp. 1–3, 2000.

Compressive Clustering of MNIST Digits

Clustering is a classic machine learning task and a component of many machine learning pipelines. The most popular method is Lloyd's algorithm [8], which aims to solve the k -means problem. In many cases, the raw features are converted to a spectral embedding before clustering. As a proof of concept, the authors in [9] applied this technique to handwritten digits from the Modified National Institute of

Standards and Technology (MNIST) data set but used compressive k -means clustering in place of Lloyd's algorithm (see Figure S3). Using an $n = 10^7$ sample augmentation of the MNIST, they found that compressive k -means clustering gave approximately the same accuracy as Lloyd's algorithm but reduced the time and memory complexity by 1.5 and four orders of magnitude, respectively (see Figure S4).

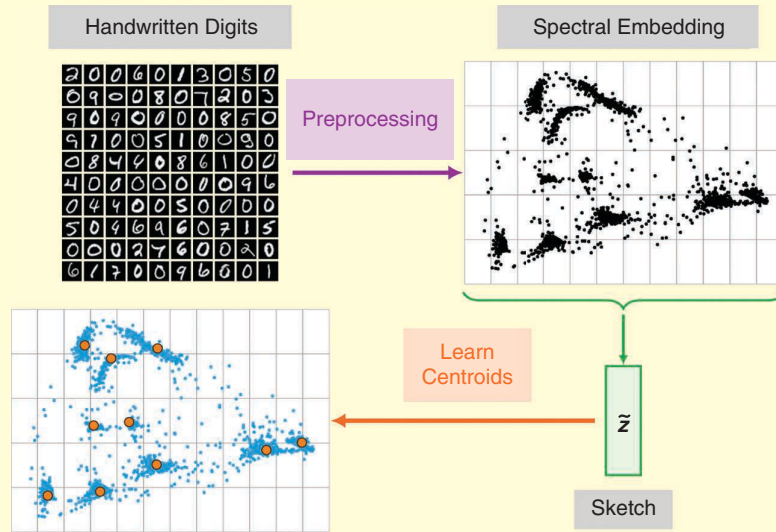


FIGURE S3. The clustering of handwritten digits via compressive k -means. First, scale-invariant feature transform descriptors are extracted from each image. Then, a similarity graph is constructed to obtain a so-called spectral embedding of the data set (using the first eigenvectors of the Laplacian of the graph). The spectral features are aggregated into a sketch, from which centroids are extracted.

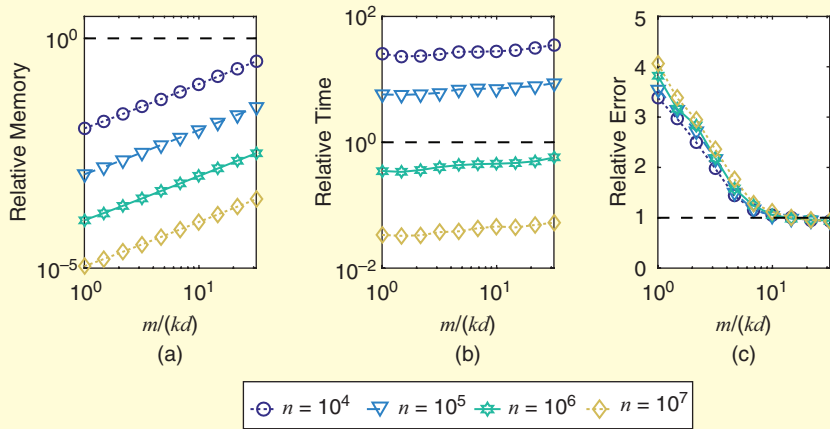


FIGURE S4. The (a) memory, (b) runtime, and (c) error of compressive k -means clustering relative to Lloyd's algorithm for various sample cardinalities n and sketch lengths m , using $k = 10$ centroids and $d = 10$ dimensional spectral features. Figure from [9].

Compressive Sensing and the Restricted Isometry Property

In compressive sensing, one observes a linear measurement $y = Ax \in \mathbb{R}^m$ of signal $x \in \mathbb{R}^d$ with $m \ll d$, i.e., with significantly reduced dimension. (For simplicity, we focus on the noiseless case for now.) To distinguish between different signals x in a given signal class, one desires that the distances between all signals in that class are preserved by the measurement operator A . For the class of k -sparse signals Σ_k , i.e., signals with at most k nonzero entries, this property is satisfied up to a tolerance $\delta \in [0, 1)$ when A obeys the restricted isometry property (RIP) [4], [22]

$$(1 - \delta)\|x - x'\|^2 \leq \|Ax - Ax'\|^2 \leq (1 + \delta)\|x - x'\|^2, \quad \forall x, x' \in \Sigma_k. \quad (S1)$$

One way to create a RIP-satisfying A is to draw it randomly. For example, if the coefficients of A are drawn as an i.i.d. zero-mean Gaussian, then A will satisfy the RIP with high probability when the sketch dimension m is at least on the order of $k \log(d/k)$ [4].

To recover k -sparse x from y , one might attempt to search for the sparsest signal among all of those that agree with the measurements, i.e., within the set $C_y := \{u : Au = y\}$. The complexity of this search, however, grows exponentially in k . Fortunately, when A satisfies the RIP, one can provably recover the true x using polynomial complexity methods [4]. One approach is to solve the convex problem of finding the signal with the smallest ℓ_1 norm within C_y . Another is to use a greedy algorithm, like orthogonal matching pursuit, which estimates x by progressively removing from y the k columns of A that best “align” with it, using a least-squares fit.

The RIP also provides guarantees on robust recovery. Suppose that we have noisy measurements $y = Ax + e$, with noise e of bounded norm $\|e\| \leq \varepsilon$. In addition, suppose that x is only approximately k -sparse, and use x_k to denote the best k -sparse approximation of x . Finally, consider recovering an estimate $\hat{x}$ of x by searching for

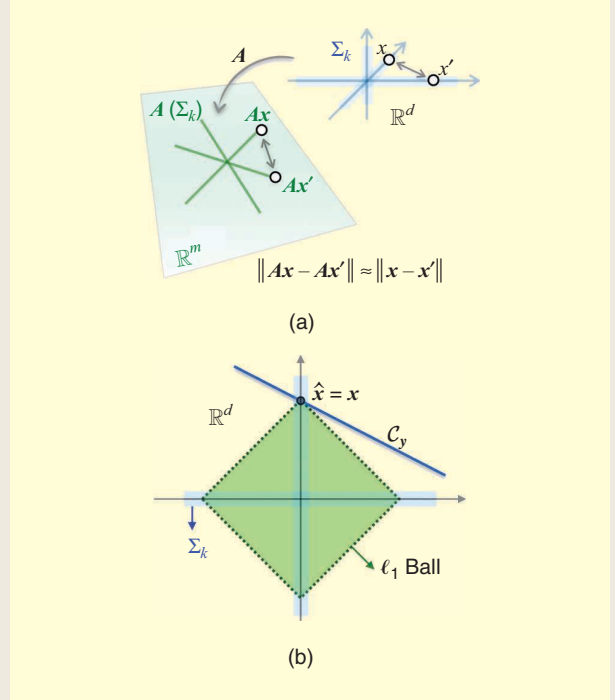


FIGURE S5. (a) A geometrical interpretation of the RIP. (b) A 2D illustration of the recoverability of x from $y = Ax$ by finding the vector $\hat{x}$ in C_y with the smallest ℓ_1 norm.

the signal with the smallest ℓ_1 norm that agrees with the measurements up to a tolerance of ε , i.e., within the set $C_{y,\varepsilon} := \{u : \|Au - y\| \leq \varepsilon\}$. Then, if the RIP (S1) holds, the estimation error $\hat{x} - x$ satisfies [4]

$$\|\hat{x} - x\| \leq C \frac{\|x - x_k\|_1}{\sqrt{k}} + D\varepsilon, \quad (S2)$$

where constants $C, D > 0$ depend only on the value of δ that appears in (S1). Thus, the estimation error increases linearly with the noise level ε and the deviation $\|x - x_k\|_1$ from perfect sparsity. Please see Figure S5.

density p_X . Consider what happens when the number of samples n goes to infinity. The strong law of large numbers says that

$$\lim_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n \Phi(x_i) \stackrel{\text{w.p.1}}{=} \mathbb{E}[\Phi(X)] = \int p_X(x) \Phi(x) dx, \quad (4)$$

where $\mathbb{E}[\cdot]$, denotes expectation with respect to the probability density p_X . If we consider the simple case of dimension $d = 1$ and the scalar transformation $\Phi(x) = x^k$ (so that $m = 1$), then $\mathbb{E}[\Phi(X)]$ is the (uncentered) k th moment of the random variable X , a quantity that has a long history in statistics. By analogy, with a generic vector-valued feature map $\Phi(\cdot)$ and in

dimension $d > 1$, quantities of the form $\mathbb{E}[\Phi(X)]$ are known as generalized moments of the random vector $X \in \mathbb{R}^d$.

Performing inference from generalized moments is often referred to as the generalized method of moments (GeMM) [21]. This method to “learn from a sketch” is very popular in, e.g., the field of econometrics [21, Ch. 1]. The GeMM can be seen as an alternative to ML estimation that avoids the need to work with the full likelihood function, which can have computational benefits. Indeed, for many classes of probability distributions, such as heavy-tailed α -stable distributions, the likelihood function is not given in closed form, but generalized moments are given in closed form for appropriately chosen

feature maps $\Phi(\cdot)$ (see the sidebar “Compressive Mixture Modeling for Speaker Verification”).

The GeMM differs from compressive learning in several aspects. In the GeMM, the feature map $\Phi(\cdot)$ is typically constructed to make the parameter estimates computable in closed form. This narrows the range of learning tasks that the GeMM can handle. In compressive learning, $\Phi(\cdot)$ is designed with information preservation in mind. Consequently, the range of learning tasks is much broader in compressive learning, although the estimation procedure (i.e., learning from a sketch) may be algorithmic in nature. Another difference is that in compressive learning, $\Phi(\cdot)$ is typically randomized. This results in randomized generalized moments, which are rarely seen in the GeMM.

Compressive learning and compressive sensing

As we will show in this section, the sketching mechanism (1) can be interpreted as a dimensionality-reducing linear “projection” of the probability distribution underlying the data set $\{\mathbf{x}_i\}_{i=1}^n$. This differs from the traditional approach in signal processing, where dimensionality reduction is performed on the features $\mathbf{x}_i \in \mathbb{R}^d$ themselves rather than the distribution that generates them. Still, many of the intuitions, analyses, and tools designed for feature-based dimensionality reduction can be extended to distribution-based dimensionality reduction.

In the field of inverse problems and compressive sensing (CS) [4], [22], a signal of interest is modeled as a high-dimensional vector $\mathbf{x} \in \mathbb{R}^d$, and a physical measurement of that signal is approximated by a linear transformation plus additive noise: $\mathbf{y} = \mathbf{A}\mathbf{x} + \mathbf{e} \in \mathbb{R}^m$ (see the sidebar “Compressive Sensing and the Restricted Isometry Property”). Here, the linear measurement operator is represented by the $m \times d$ matrix $\mathbf{A}$. In many applications, there is a great motivation to make the measurement dimension much less than the signal dimension; i.e., $m \ll d$. In this case, the recovery of $\mathbf{x}$ from $\mathbf{y}$ is generally ill posed. Still, it is possible to accurately recover $\mathbf{x}$ from $\mathbf{y}$ when both $\mathbf{x}$ and $\mathbf{A}$ obey certain properties and the noise is small enough.

The celebrated Johnson–Lindenstrauss (JL) lemma and its variations specify, for instance, that with high probability one can nearly preserve the pairwise distances between N features in $\mathbb{R}^d$ by linearly projecting them in a $O(\log N)$ -dimensional domain. This is used, for instance, to accelerate nearest-neighbor searches in large databases. An important extension is to CS: if $\mathbf{x}$ is sparse, in that relatively few of its coefficients deviate significantly from zero, and if $\mathbf{A}$ satisfies the restricted isometry property (RIP)—an extension of the JL lemma to the continuous set of sparse signals—then one can pose a regularized inverse problem whose solution is close to the true $\mathbf{x}$ (see the sidebar “Compressive Sensing and the Restricted Isometry Property”).

Now that CS has been described, we can clearly connect it to compressive learning. Recall that the sketch (1) $\tilde{\mathbf{z}}$ converges to the generalized moment $\mathbf{z} := \mathbb{E}[\Phi(\mathbf{X})] = \int p_X(\mathbf{x})\Phi(\mathbf{x})d\mathbf{x}$ as the number of data samples n tends to infinity, as per (4).

A crucial observation is the following: due to the linearity of integration, the sketch $\mathbf{z}$ depends linearly on the probability density p_X . To see it another way, consider a mixture of two

densities, $p_X = \alpha p_{X_1} + (1 - \alpha)p_{X_2}$. The corresponding expectation satisfies

$$\mathbb{E}_{X \sim p_X}[\Phi(\mathbf{X})] = \alpha \mathbb{E}_{X_1 \sim p_{X_1}}[\Phi(\mathbf{X}_1)] + (1 - \alpha) \mathbb{E}_{X_2 \sim p_{X_2}}[\Phi(\mathbf{X}_2)], \quad (5)$$

which implies that the generalized moment $\mathbf{z}$ is linear in p_X . With this understanding, we can write

$$\mathbf{z} = \mathcal{A}(p_X) := \mathbb{E}_{X \sim p_X}[\Phi(\mathbf{X})], \quad (6)$$

where $\mathcal{A}$ is a linear operator mapping the probability distribution p_X to the m -dimensional sketch vector $\mathbf{z}$.

Although $\mathcal{A}$ is a linear function of p_X , we emphasize that the feature map $\Phi(\mathbf{x})$ is generally not a linear function of $\mathbf{x}$. This makes compressive learning concretely very different from the vast majority of existing “sketching” mechanisms, which use random linear projections of the data for dimensionality reduction.

With a small modification of the preceding arguments, we can handle the finite sample case. Consider the difference between the true generalized moment $\mathbf{z}$ and the empirical moment $\tilde{\mathbf{z}}$; i.e.,

$$\frac{1}{n} \sum_{i=1}^n \Phi(\mathbf{x}_i) - \mathbb{E}[\Phi(\mathbf{X})] =: \mathbf{e}. \quad (7)$$

As we discussed earlier, $\mathbf{e}$ converges to zero as n tends to infinity by the law of large numbers. Combining (1), (6), and (7), we obtain

$$\tilde{\mathbf{z}} = \mathcal{A}(p_X) + \mathbf{e}, \quad (8)$$

which shows that the sketch (1) can be interpreted as a “noisy” observation of the data distribution p_X through the linear measurement operator $\mathcal{A}$. Under mild conditions the central limit theorem can be used to show that $\|\mathbf{e}\|$ decays as $1/\sqrt{n}$ with high probability [4, Ch. 8]. From the assumed independence of the data samples, this happens, for instance, if $\mathbb{P}[\|\Phi(\mathbf{X})\| \geq t]$ decays exponentially fast when t increases. With the preceding interpretation of compressive learning, one recovers all of the traditional ingredients of CS:

- The measurement operator $\mathcal{A}$ is linear.
- The measurements $\tilde{\mathbf{z}}$ are drastically dimension reduced. In mathematical terms, probability distributions p_X belong to the infinite-dimensional vector space of so-called finite measures [23]. The operator $\mathcal{A}$ maps these infinite-dimensional objects to vectors of finite dimension m .
- The measurement operator $\mathcal{A}$ is typically designed using randomness. This is accomplished by choosing an appropriate randomized feature map $\Phi(\cdot)$, such as one based on RF features, as in (3).
- The measurements $\tilde{\mathbf{z}}$ are noisy, as per (6).

The analogy between compressive learning and CS is illustrated in Figure 4. The CS analog to the sketching phase (1) is the signal sensing phase, where the signal $\mathbf{x}$ is linearly mapped to the observation vector $\mathbf{y}$. The analog to sparse recovery, where an estimate of $\mathbf{x}$ is computed from the observation $\mathbf{y}$ by solving an optimization problem, is to learn from a sketch, where an estimate of the data distribution p_X (or of distributional

parameters θ of interest) is computed from the sketch $\tilde{z}$. Its expression as an optimization problem will be further discussed in the “Learning From a Sketch” section.

RF sampling and superresolution recovery

In this section, we consider the specific case of compressive clustering, which allows us to forge a concrete connection between CS and compressive learning using RF sampling and superresolution recovery.

We begin by considering the goal of recovering k centroids $\{c_\ell\}_{\ell=1}^k$ in $\mathbb{R}^d$ with Euclidean norm $\leq r$ that are separated from each other in Euclidean distance by $\geq \epsilon$. A naive approach would be to discretize the d -dimensional cube of side length $2r$ with a grid spacing of ϵ , leading to $N = (2r/\epsilon)^d$ bins. A valid sketch of the data $\mathcal{X}$ is obtained by simply computing the histogram $\hat{p} \in \mathbb{R}_+^N$ over these bins (i.e., by using the “binning” feature map). However, the dimension of this sketch, N , grows exponentially in the feature dimension, d . To construct a smaller sketch, one might reason that if the data clusters tightly around k points, then $\hat{p}$ is close to a k -sparse vector. In this case, ideas from CS can be directly exploited. In particular, one could use a sketched histogram [13] of the form $\tilde{z} = A\hat{p}$, where matrix $A \in \mathbb{R}^{m \times N}$ is randomly drawn with i.i.d. Gaussian components. Here, *learning from a sketch* means recovering the centroids. For this, one would first search for the best nonnegative, k -sparse, sum-to-one vector using

$$\tilde{p} = \underset{p \in \Sigma_k^+}{\operatorname{argmin}} \|\tilde{z} - Ap\|, \quad (9)$$

(or a convex or greedy relaxation of this problem) where Σ_k^+ here denotes the set of k -sparse, nonnegative, sum-to-one vectors. Then, one would identify the k grid locations in $\mathbb{R}^d$ corresponding to the nonzero indices of $\tilde{p}$. CS theory [4] says that the support of $\tilde{p}$ will be accurate for sketch dimensions m at least on the order of $k \log N$, i.e., at least on the order of $kd \log(r/\epsilon)$. Although this latter approach substantially reduces the dimension of the sketch, practical challenges remain when the feature dimension d is large. For example, the number of columns, N , in the compression matrix A grows exponentially in d , making storage and multiplication by A impractical.

An alternative approach could be to construct A using m rows of the (d -dimensional, in this case) discrete Fourier transform (DFT) matrix, i.e., by sampling the DFT at m (d -dimensional) frequencies. In this case, multiplication by A could be implemented by the fast Fourier transform (FFT) algorithm, and thus the matrix A would not need to be explicitly stored. Fourier domain sampling is a familiar operation in the context of signal processing, as it forms the cornerstone for radar, medical imaging, and radio interferometry; see, e.g., [24] and [25]. When the m frequencies are drawn uniformly at random, CS theory [4], [22] has established that accurate recovery of an N -length k -sparse signal can be accomplished (with high probability) when m is on the order of $k \log^3(k) \log N$ [4, Corollary 12.38]. Since $N = (2r/\epsilon)^d$ here, this would mandate sketch

dimensions m at least on the order of $kd \log^3(k) \log(r/\epsilon)$. Although the FFT avoids the need to store A as an explicit matrix and allows efficient computation of the sketch, the cost of solving the optimization problem (9) using existing convex relaxations or greedy approaches is impractical due to the need to manipulate N -dimensional vectors, where N grows exponentially in d .

Until now, we considered discretizing the d -dimensional feature space on an ϵ -spaced hypergrid within a $2r$ -sidelength hypercube but found that this requires manipulating vectors (e.g., a histogram $\hat{p}$) whose dimension grows exponentially in d . We can avoid this discretization (i.e., take $\epsilon \rightarrow 0$ and $r \rightarrow \infty$) by replacing the DFT with the continuous Fourier transform (CFT), in which case we are Fourier transforming the empirical distribution $\hat{p}_\mathcal{X}(\mathbf{x}) := (1/n) \sum_{i=1}^n \delta(\mathbf{x} - \mathbf{x}_i)$ rather than its N -bin histogram $\hat{p}$. The sketch $\tilde{z}$ then has components

$$\tilde{z}_j = \int_{\mathbb{R}^d} \hat{p}_\mathcal{X}(\mathbf{x}) \exp(-j2\pi \mathbf{w}_j^\top \mathbf{x}) d\mathbf{x} \quad (10)$$

$$= \frac{1}{n} \sum_{i=1}^n \exp(-j2\pi \mathbf{w}_j^\top \mathbf{x}_i) \quad j = 1, \dots, m, \quad (11)$$

where $\{\mathbf{w}_j\}_{j=1}^m$ are d -dimensional frequencies that are drawn at random. Typically, $\mathbf{w}_j$ are drawn i.i.d. Gaussian, but other distributions can be used, as discussed in the sections “The Challenge of Designing a Feature Map Given a Learning Task” and “Sketching With Structured Random Matrices.” Note that (11) corresponds precisely to the sketch (1) with the RF feature map $\Phi(\cdot)$ from (3). Taking a statistical perspective, the components $\tilde{z}_j$ in (10) can also be recognized as samples of the characteristic function of $\hat{p}_\mathcal{X}$. Recall that for a density p , the characteristic function Ψ_p is defined as

$$\Psi_p(\mathbf{w}) := \int_{\mathbb{R}^d} p(\mathbf{x}) \exp(j2\pi \mathbf{w}^\top \mathbf{x}) d\mathbf{x} = \mathbb{E}_{\mathbf{x} \sim p} [\exp(j2\pi \mathbf{w}^\top \mathbf{x})]. \quad (12)$$

Intuitively, when a probability distribution p has a “simple” structure, one can recover it (with high probability) from enough randomly chosen samples of its CFT. Centroid recovery from the sketch (10) is premised on the empirical distribution $\hat{p}_\mathcal{X}$ being well approximated by a mixture of k Diracs;

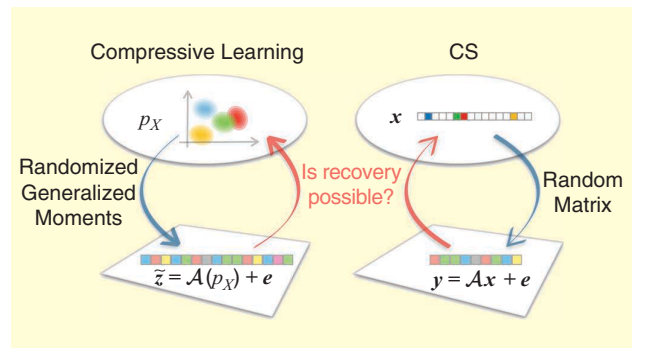


FIGURE 4. The analogy between compressive learning, which uses the dimensionality-reducing linear measurement $\mathcal{A}(p_\mathcal{X}) = \mathbb{E}[\Phi(\mathbf{X})]$ of a distribution $p_\mathcal{X}$, and CS, which uses the dimensionality-reducing linear measurement $A\mathbf{x}$ of a signal $\mathbf{x}$.

i.e., $\hat{p}_{\mathcal{X}}(\mathbf{x}) \approx \sum_{\ell=1}^k \alpha_{\ell} \delta(\mathbf{x} - \mathbf{c}_{\ell})$. In this case, centroid recovery parallels the “superresolution” recovery problem (see the sidebar “Superresolution Recovery”) through an optimization problem that is the continuous analog to (9) (or a convex or greedy relaxation for the continuous case) and will be further elaborated in the next section.

In both problems, recovery guarantees are possible when the frequencies $\mathbf{w}_j$ are randomly drawn. For example, when the centroids $\{\mathbf{c}_{\ell}\}_{\ell=1}^k$ are ϵ -separated and r -bounded, centroid recovery guarantees have been established provided the sketch dimension m is on the order of $k^2 d \log(r/\epsilon)$, omitting for simplicity some log factors involving k and d [5], [51]. Similar guarantees hold when $\hat{p}_{\mathcal{X}}$ is approximately a sum of spatially localized components (e.g., in compressive GMM) [5], [51].

Learning from a sketch

Until now, we have primarily focused on the first stage of the compressive learning pipeline (see Figure 1), where the data set $\mathcal{X}$ is sketched down to $\tilde{\mathbf{z}}$, a compressed and noisy representation of the underlying data-generating distribution $p_{\mathcal{X}}$. We now discuss the second stage of the pipeline, where the distributional parameters of interest, $\boldsymbol{\theta}$, are recovered from the sketch $\tilde{\mathbf{z}}$. The close analogy between sketching and CS allows us to cast this “parameter learning” stage as an optimization problem; i.e.,

$$\tilde{\boldsymbol{\theta}} = \underset{\boldsymbol{\theta}}{\operatorname{argmin}} C(\boldsymbol{\theta} | \tilde{\mathbf{z}}), \quad (13)$$

where the cost function $C(\cdot | \tilde{\mathbf{z}})$ is adapted to the considered learning task. As in CS, many candidate distributions $p_{\mathcal{X}}$ (and hence many candidate parameters $\boldsymbol{\theta}$) can yield the same sketch $\tilde{\mathbf{z}}$. Thus, to make the inverse problem well posed, one needs to employ concrete modeling assumptions and regular-

ization, both of which can take several forms. As in CS, we will assume that the sketched quantity $p_{\mathcal{X}}$ is of low intrinsic complexity, i.e., close to some family of “simple” probability distributions.

As a first example, we consider the problem of learning a mixture model from a sketch. Similar to a sparse vector $\mathbf{x}$, which is a linear combination of a few elements of the standard basis, a mixture model $p_{\mathcal{X}}$ is a linear combination of a few “simple” densities $\{p_{\boldsymbol{\theta}_{\ell}}\}_{\ell=1}^k$. Concretely, $p_{\mathcal{X}} = \sum_{\ell=1}^k \alpha_{\ell} p_{\boldsymbol{\theta}_{\ell}}$, where the mixture weights $\{\alpha_{\ell}\}_{\ell=1}^k$ are nonnegative and sum to one. For example, with a GMM, we have that $p_{\boldsymbol{\theta}_{\ell}} = \mathcal{N}(\boldsymbol{\mu}_{\ell}, \boldsymbol{\Sigma}_{\ell})$, where $\boldsymbol{\theta}_{\ell} = (\boldsymbol{\mu}_{\ell}, \boldsymbol{\Sigma}_{\ell})$ contains the mean $\boldsymbol{\mu}_{\ell}$ and covariance $\boldsymbol{\Sigma}_{\ell}$. If $p_{\mathcal{X}}$ is well approximated by the mixture model $\sum_{\ell=1}^k \alpha_{\ell} p_{\boldsymbol{\theta}_{\ell}}$, then, according to (8), the sketch $\tilde{\mathbf{z}}$ is well approximated by the linear combination $\sum_{\ell=1}^k \alpha_{\ell} \mathcal{A}(p_{\boldsymbol{\theta}_{\ell}})$. Hence, one could try to extract the mixture parameters, $\boldsymbol{\theta} = \{\alpha_{\ell}, \boldsymbol{\theta}_{\ell}\}_{\ell=1}^k$, from the sketch $\tilde{\mathbf{z}}$ by solving the (nonconvex) optimization problem (13) with

$$C(\boldsymbol{\theta} | \tilde{\mathbf{z}}) := \left\| \tilde{\mathbf{z}} - \sum_{\ell=1}^k \alpha_{\ell} \mathcal{A}(p_{\boldsymbol{\theta}_{\ell}}) \right\|^2. \quad (14)$$

The cost $C(\boldsymbol{\theta} | \tilde{\mathbf{z}})$ can be interpreted as the negative log likelihood (up to a shift and scale) of $\boldsymbol{\theta}$ given the sketch $\tilde{\mathbf{z}}$, under the classic modeling assumption of i.i.d. Gaussian measurement noise $\mathbf{e}$ in (8).

When the RF feature map $\Phi(\cdot)$ is used to compute the sketch $\tilde{\mathbf{z}}$ and the component densities $p_{\boldsymbol{\theta}_{\ell}}$ are Gaussian or α stable, there exist analytic expressions for $\mathcal{A}(p_{\boldsymbol{\theta}_{\ell}})$ and for the gradient of $\mathcal{A}(p_{\boldsymbol{\theta}_{\ell}})$ with respect to the mixture parameters in $\boldsymbol{\theta}_{\ell}$ [6]. These expressions are convenient when numerically optimizing (14). For instance, greedy approaches, similar to the orthogonal matching pursuit (OMP) algorithm for CS (recall the sidebar

Superresolution Recovery

Superresolution is a general class of techniques to enhance the resolution of a sensing system, e.g., to observe subwavelength features in astronomy or medical imaging [23]. The problem addressed by superresolution is to recover a continuous-time (when dimension $d = 1$) or continuous-space (when dimension $d \geq 2$) sparse signal $s(\mathbf{t})$ from a few, possibly noisy, Fourier measurements $\{y_j\}_{j=1}^m$. This amounts to recovering a weighted sum $s(\mathbf{t}) = \sum_{\ell=1}^k \alpha_{\ell} \delta(\mathbf{t} - \mathbf{t}_{\ell})$ of k Diracs with amplitudes α_{ℓ} and locations $\mathbf{t}_{\ell} \in \mathbb{R}^d$ from

$$y_j = \int_{\mathbb{R}^d} s(\mathbf{t}) \exp(-j2\pi \mathbf{w}_j^T \mathbf{t}) d\mathbf{t} + e_j, \quad j = 1, \dots, m, \quad (S3)$$

with measurement noise e_j and frequency vectors $\{\mathbf{w}_j\}_{j=1}^m$ in $\mathbb{R}^d$. Recovery of $s(\cdot)$ can be posed as an infinite-dimensional convex problem on measures [23]. However, most reconstruction algorithms involve nonconvex steps [26]. When the frequencies $\mathbf{w}_j$ are drawn randomly, the

signal can be accurately recovered with high probability when m is of the order of at least kd^3 , up to log factors. Proving this usually requires additional assumptions, such as a minimal separation between the locations $\mathbf{t}_{\ell}$ [S2] or positivity of the amplitudes α_{ℓ} [26].

The link between superresolution and compressive clustering follows from rewriting (10) as

$$\tilde{\mathbf{z}}_j = \int_{\mathbb{R}^d} p_{\mathcal{X}}(\mathbf{x}) \exp(-j2\pi \mathbf{w}_j^T \mathbf{x}) d\mathbf{x} + e_j, \quad j = 1, \dots, m, \quad (S4)$$

where e_j captures the “noise” due to finite-sample effects [recall (7)]. Comparing (S4) to (S3), we see that they are mathematically equivalent when $p_{\mathcal{X}}(\mathbf{x}) = \sum_{\ell=1}^k \alpha_{\ell} \delta(\mathbf{x} - \mathbf{c}_{\ell})$ except for the fact that, in the case of compressive clustering, α_{ℓ} are nonnegative and sum to one.

Reference

[S2] C. Poon, N. Keriven, and G. Peyré, “The geometry of off-the-grid compressed sensing,” 2020. arXiv: 1802.08464.

“Compressive Sensing and the Restricted Isometry Property”), can be used [6] to estimate the parameters $\{\alpha_\ell, \theta_\ell\}_{\ell=1}^k$. These approaches sequentially estimate and subtract, from $\tilde{\mathbf{z}}$, each of the k components $\alpha_\ell \mathcal{A}(p_{\theta_\ell})$ that best align with it, where “best” is measured via the correlation between $\tilde{\mathbf{z}}$ and $\mathcal{A}(p_{\theta_\ell})$. As another example, an iterative approach [10] was proposed that exploits the log likelihood interpretation of $C(\theta|\tilde{\mathbf{z}})$ in (14) and the i.i.d. random nature of the linear transform $\mathbf{W}$ in the RF map (3). In [6], applications of compressive mixture modeling are demonstrated on speaker verification (see the sidebar “Compressive Mixture Modeling for Speaker Verification”) and source separation.

As a second example, we consider compressive k -means clustering. Here, the goal is to recover, from the sketch $\tilde{\mathbf{z}}$ the centroids $\theta = \{\mathbf{c}_\ell\}_{\ell=1}^k$ that minimize the average squared Euclidean distance from each sample to its nearest centroid; i.e., $(1/n) \sum_{i=1}^n \min_\ell \|\mathbf{x}_i - \mathbf{c}_\ell\|^2$. To tackle this k -means problem, we view it as an approximation of a particular GMM fitting problem. In particular, suppose that the probability distribution p_X is a GMM with weights α_ℓ , mean vectors $\mathbf{c}_\ell$, and covariance matrices Σ_ℓ . Then, in the special case that $\alpha_\ell = 1/k$ and $\Sigma_\ell = \sigma^2 \mathbf{I}$ for all components $1 \leq \ell \leq k$, we can write the likelihood as $\prod_{i=1}^n p(\mathbf{x}_i|\theta) \propto \prod_{i=1}^n \sum_{\ell=1}^k \exp(-(1/2\sigma^2) \|\mathbf{x}_i - \mathbf{c}_\ell\|^2)$, and so the negative log likelihood becomes $-\sum_{i=1}^n \log \sum_{\ell=1}^k \exp(-(1/2\sigma^2) \|\mathbf{x}_i - \mathbf{c}_\ell\|^2)$ up to an additive constant. We can then use the log-sum-exp approximation $\log \sum_{\ell=1}^k \exp(f_\ell(\mathbf{x})) \approx \max_{\ell} f_\ell(\mathbf{x})$ to approximate this latter expression as $(1/2\sigma^2) \sum_{i=1}^n \min_{\ell} \|\mathbf{x}_i - \mathbf{c}_\ell\|^2$, which agrees with the k -means cost up to a scaling. If we furthermore consider the case of a vanishing variance $\sigma^2 \rightarrow 0$, then the component density p_{θ_ℓ} reduces to a point mass; i.e., $p_{\theta_\ell}(\mathbf{x}) \rightarrow \delta(\mathbf{x} - \mathbf{c}_\ell)$. In this limiting case, the linear measurement of the point mass p_{θ_ℓ} is $\mathcal{A}(p_{\theta_\ell}) = \mathbb{E}_{X \sim p_{\theta_\ell}}[\Phi(X)] = \int \Phi(\mathbf{x}) p_{\theta_\ell}(\mathbf{x}) d\mathbf{x} = \Phi(\mathbf{c}_\ell)$, according to (6).

Thus, with these justifications, the cost function (14) for the compressive GMM would change to

$$C(\theta|\tilde{\mathbf{z}}) := \left\| \tilde{\mathbf{z}} - \frac{1}{k} \sum_{\ell=1}^k \Phi(\mathbf{c}_\ell) \right\|^2 \quad (15)$$

for compressive k -means. If we do not want to assume that $\alpha_\ell = 1/k$ for each ℓ , we could instead estimate $\{\alpha_\ell\}_{\ell=1}^k$ from the sketch, leading to the cost function suggested in [9]:

$$C(\theta|\tilde{\mathbf{z}}) := \min_{\alpha} \left\| \tilde{\mathbf{z}} - \sum_{\ell=1}^k \alpha_\ell \Phi(\mathbf{c}_\ell) \right\|^2. \quad (16)$$

Similar to the compressive GMM problem described earlier, minimization of the cost function (16) for compressive k -means can be tackled by greedy approaches, as described in [9]. Despite the fact that (16) does not directly minimize the k -means cost, it has been shown empirically [9] that the centroids estimated by such greedy algorithms nearly minimize this cost; see, e.g., the sidebar “Compressive Clustering of MNIST Digits” for an example on MNIST data. This claim is also supported by theoretical results guaranteeing that the minimizer of (16) is endowed with statistical learning guar-

antees with respect to the original k -means cost [5], [51]. Such guarantees will be discussed shortly.

Depending on the choice of parameters θ and the feature map $\Phi(\cdot)$, the form of the optimization problem posed to learn from a sketch can differ considerably from that for GMMs in (14) and that for the k -means in (15) and (16). Consider, for example, learning from a sketch for PCA. As described earlier, the parameter θ of interest is the k -dimensional subspace that best fits the d -dimensional data $\{\mathbf{x}_i\}_{i=1}^n$ in an LS sense. It is well known that this subspace is spanned by k principal eigenvectors of the empirical autocorrelation matrix $\hat{\mathbf{R}} = (1/n) \sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i^\top$, or, equivalently, the column space of the (positive semidefinite symmetric) matrix $\hat{\mathbf{R}}_k$ that is closest to $\hat{\mathbf{R}}$ in the Frobenius norm:

$$\hat{\mathbf{R}}_k := \underset{\mathbf{R}: \text{rank}(\mathbf{R}) \leq k}{\text{argmin}} \|\hat{\mathbf{R}} - \mathbf{R}\|_F = \underset{\mathbf{R}: \text{rank}(\mathbf{R}) \leq k}{\text{argmin}} \|\text{vec}(\hat{\mathbf{R}}) - \text{vec}(\mathbf{R})\|^2. \quad (17)$$

When sketching using quadratic features $\Phi(\mathbf{x}) = ((\mathbf{w}_1^\top \mathbf{x})^2, \dots, (\mathbf{w}_m^\top \mathbf{x})^2)$ with random $\mathbf{w}_j \in \mathbb{R}^d$, the j th component of the sketch becomes $\tilde{z}_j = (1/n) \sum_{i=1}^n \mathbf{w}_j^\top \mathbf{x}_i \mathbf{x}_i^\top \mathbf{w}_j = \mathbf{w}_j^\top \hat{\mathbf{R}} \mathbf{w}_j$. Importantly, this $\tilde{z}_j$ is a linear function of $\hat{\mathbf{R}}$, and so there exists an $m \times d^2$ matrix $\mathbf{A}$ such that $\tilde{\mathbf{z}} = \mathbf{A} \text{vec}(\hat{\mathbf{R}})$. Thus, by analogy with (17), one could first fit a positive semidefinite symmetric low-rank matrix to the sketch $\tilde{\mathbf{z}}$ via

$$\tilde{\mathbf{R}} = \underset{\mathbf{R}: \text{rank}(\mathbf{R}) \leq k, \mathbf{R}^\top = \mathbf{R}, \mathbf{R} \geq 0}{\text{argmin}} \|\tilde{\mathbf{z}} - \mathbf{A} \text{vec}(\mathbf{R})\|^2 \quad (18)$$

and then set the parameter estimate $\tilde{\theta}$ equal to the column space of $\tilde{\mathbf{R}}$. While the optimum of (17) is automatically symmetric and positive semidefinite, this property needs to be enforced explicitly in (18).

The low-rank matrix recovery problem (18) has been thoroughly investigated by the signal processing and machine learning communities (e.g., [27]). It arises, for example, in applications such as collaborative filtering for recommender systems and signal reconstruction from phaseless measurements. Early approaches to solving the nonconvex problem (18) involved convex relaxation via nuclear norm regularization [4, Ch. 4]. More recent approaches exploit the nonconvex geometry of (18) [28].

Compressive learning with theoretical guarantees

In the previous section, learning from a sketch was posed as the optimization problem (13), repeated here for convenience:

$$\tilde{\theta} = \underset{\theta}{\text{argmin}} C(\theta|\tilde{\mathbf{z}}). \quad (19)$$

The quantity $C(\cdot|\tilde{\mathbf{z}})$ is a real-valued cost function whose minimizer $\tilde{\theta}$ is the “best” (in some sense) estimate of θ from the sketch $\tilde{\mathbf{z}}$. This approach contrasts with traditional statistical learning [29], [30], which computes the parameter estimate

$$\hat{\theta} := \underset{\theta}{\text{argmin}} R(\theta|\mathcal{X}) \quad \text{with} \quad R(\theta|\mathcal{X}) := \frac{1}{n} \sum_{i=1}^n L(\theta|\mathbf{x}_i), \quad (20)$$

where $R(\cdot|\mathcal{X})$ is the “empirical risk” and $L(\cdot|x_i)$ is the “loss function” for the i th sample. Table 2 presents typical loss functions for our four running examples. In most cases, performing the minimization in (20) involves querying the full training data set $\mathcal{X}$ many times, e.g., for stochastic gradient descent. In compressive learning, the cost $C(\cdot|\tilde{z})$, which depends only on the low-dimensional sketch $\tilde{z}$, is used as a surrogate for the empirical risk. Because the dimension of the sketch is so much smaller than the cardinality of the full data set $\mathcal{X}$, the minimization in (13) can be made much more efficient than that in (20).

Of course, the estimate $\tilde{\theta}$, which minimizes the cost $C(\cdot|\tilde{z})$, is not, in general, the same as $\hat{\theta}$, which minimizes the empirical risk $R(\cdot|\mathcal{X})$. Both, however, are approximations of the ideal estimate

$$\theta^* := \underset{\theta}{\operatorname{argmin}} R^*(\theta) \quad \text{with} \quad R^*(\theta) := \mathbb{E}[L(\theta|X)], \quad (21)$$

where $R^*(\theta)$ is known as the “true risk.” Indeed, the feature map $\Phi(\cdot)$ and cost $C(\cdot|\tilde{z})$ are designed precisely to ensure that the surrogate minimization (19) approximates the true risk minimization (21).

Excess risk guarantees

To establish a guarantee on the goodness of an estimate θ , one must prove that the “excess risk” $R^*(\theta) - R^*(\theta^*) \geq 0$ is small. Indeed, the true risk can be interpreted as the expected loss on test samples that have the same generating distribution p_X as the training samples $\mathcal{X}$ but that are drawn independently and not accessible at training time. In statistical learning tasks, the true risk $R^*(\cdot)$ is the primary metric by which one judges the quality of an arbitrary estimate θ . Note that controlling the excess risk is different from proving that θ is close to the ideal estimate θ^* in, e.g., the Euclidean distance $\|\theta - \theta^*\|$. Consider, for example, the problem of fitting a 1D linear subspace to a data set (i.e., PCA with $k = 1$). When the two largest eigenvalues of the autocorrelation matrix R are equal, any nontrivial linear combination of the corresponding eigenvectors generates a 1D subspace θ with minimum true risk. The problem of estimating a subspace “close to the optimal one” is thus ill defined, yet finding a subspace with close-to-optimal performance is well defined and achievable, both by classic PCA and compressive PCA.

Table 2. The loss functions for our four running examples.

Running Example	Parameters θ	Loss Function $L(\theta x_i)$
PCA	Orthonormal basis $\theta = \{u_\ell\}_{\ell=1}^k$	$\sum_{\ell=1}^k x_i^\top u_\ell ^2$
LS linear regression	A weight matrix θ	$\ x_{1i} - \theta x_{2i}\ ^2, x_i := (x_{1i}^\top, x_{2i}^\top)^\top$
Gaussian mixture modeling	GMM parameters $\theta = \{\alpha_\ell, \mu_\ell, \Sigma_\ell\}_{\ell=1}^k$	$-\log \sum_{\ell=1}^k \alpha_\ell \mathcal{N}(x_i; \mu_\ell, \Sigma_\ell)$
k-means clustering	A set of centroids $\theta = \{c_\ell\}_{\ell=1}^k$	$\min_\ell \ x_i - c_\ell\ ^2$

Despite the fact that the estimates $\tilde{\theta}$, $\hat{\theta}$, and θ^* minimize different objective functions, they often yield similar true risks; i.e., the excess risks of $\tilde{\theta}$ and $\hat{\theta}$ are provably small. For $\hat{\theta}$, the proofs use classic results from statistical learning [29], [30], under assumptions that we will briefly discuss in the sequel. For $\tilde{\theta}$, with an appropriately designed feature map $\Phi(\cdot)$ and cost function $C(\cdot|\tilde{z})$, the proofs informally follow from the fact that $C(\cdot|\tilde{z})$ and $R^*(\cdot)$ have a similar shape. This fact is illustrated in the sidebar “Traditional Statistical Learning Versus Compressive Learning” for a toy example of compressive clustering. There it can be seen that, with a properly designed feature map $\Phi(\cdot)$, a well-chosen cost function $C(\cdot|\tilde{z})$, and a sufficiently large sketch dimension m , the minimizer $\tilde{\theta}$ of the cost yields a nearly minimal true risk $R^*(\tilde{\theta})$ and lives in a large basin of attraction in $C(\cdot|\tilde{z})$. If the sketch dimension m is chosen too small, however, then the excess risk $R^*(\tilde{\theta}) - R^*(\theta^*)$ increases.

A theory of compressive learning [5], [51] has been developed to better understand how (random) feature maps $\Phi(\cdot)$ and cost functions $C(\cdot|\tilde{z})$ can be designed to ensure that the sketch $\tilde{z}$ captures sufficient information to control the excess risk on $\tilde{\theta}$. In the following four sections, we aim to summarize the key aspects of this theory using broadly accessible language. For those interested, the full technical details can be found in [5], [51].

In a nutshell, the theory relies on interpreting $\tilde{\theta}$ as the minimizer of the risk attributed to a surrogate probability distribution $\tilde{p}$ with the following properties:

- $\tilde{p}$ is “close” to the distribution p_X , as measured by a task-driven distance $d(\tilde{p}, p_X)$.
- $\tilde{p}$ is a “simple” distribution (e.g., for the compressive GMM, $\tilde{p}$ is a Gaussian mixture).
- $\tilde{p}$ minimizes $\|\tilde{z} - \mathcal{A}(p)\|$ among all simple distributions p . For example, in the compressive GMM, this holds when $\tilde{p} = p_{\tilde{\theta}}$, i.e., the Gaussian mixture distribution parameterized by $\tilde{\theta}$. Given such a $\tilde{p}$, the theory can be explained in the following step-by-step manner:
- Just as in traditional statistical learning, excess risk bounds can be established if the task-driven distance $d(\tilde{p}, p_X)$ is controlled (see the section “Task-Driven Distances and Excess Risk Bounds”).
- Given a family of “simple” probability distributions (see the section “Exploiting Simplified Models to Learn From a Sketch”), this distance can indeed be controlled provided a certain “lower RIP” (LRIP) holds (see the section “The LRIP and Excess Risk Control”). This is analogous to traditional CS, where the RIP allows one to control the performance of sparse signal recovery.
- To establish the LRIP for random feature maps $\Phi(\cdot)$, one can build on connections between kernel methods and the JL lemma (see the section “Establishing the LRIP via Kernels and the JL Lemma”).
- Finally, to design task-specific feature maps $\Phi(\cdot)$, one can appeal to the problem of kernel design (see the section “The Challenge of Designing a Feature Map Given a Learning Task”).

Task-driven distances and excess risk bounds

In traditional statistical learning, it is common to define a distance between the true distribution p_X and an arbitrary surrogate p'_X as

$$d(p'_X, p_X) = \sup_{\theta} |R^*(\theta | p'_X) - R^*(\theta | p_X)|, \quad (22)$$

where $R^*(\theta | p) := \mathbb{E}_{X \sim p}[L(\theta | X)]$ denotes the risk under an arbitrary distribution p . This distance is task specific since the loss function $L(\cdot | \cdot)$ depends on the estimation task. The utility of $d(p'_X, p_X)$ for assessing the goodness of parameter estimation follows from the inequalities

$$\begin{aligned} R^*(\theta^* | p_X) &\leq R^*(\theta'' | p_X) \leq R^*(\theta'' | p'_X) + d(p'_X, p_X) \\ &\leq R^*(\theta^* | p'_X) + d(p'_X, p_X) \\ &\leq R^*(\theta^* | p_X) + 2d(p'_X, p_X), \end{aligned} \quad (23)$$

where $\theta^* := \operatorname{argmin}_{\theta} R^*(\theta | p_X)$ and $\theta'' := \operatorname{argmin}_{\theta} R^*(\theta | p'_X)$. In particular, these inequalities yield the following bounds on the excess risk of θ'' :

$$0 \leq R^*(\theta'' | p_X) - R^*(\theta^* | p_X) \leq 2d(p'_X, p_X). \quad (24)$$

For example, if p'_X equals the empirical distribution $\hat{p}_X(\mathbf{x}) := (1/n) \sum_{i=1}^n \delta(\mathbf{x} - \mathbf{x}_i)$, then (24) bounds the “generalization error” of empirical risk minimization, i.e., the error incurred when training with $\mathcal{X} = \{\mathbf{x}_i\}_{i=1}^n$ but testing with independent samples drawn from p_X . This is the reasoning behind many classic statistical learning guarantees, where so-called uniform convergence results establish that, under appropriate conditions, $d(\hat{p}_X, p_X) = O(1/\sqrt{n})$ with high probability on the draw of i.i.d. training samples $\mathcal{X} = \{\mathbf{x}_i\}_{i=1}^n$. The term *uniform convergence* stems from the analogy between the distance (22) and the ℓ_{∞} norm.

Traditional Statistical Learning Versus Compressive Learning

In statistical learning, the ideal parameter θ^* minimizes the true risk $R^*(\theta) = \mathbb{E}_{X \sim p_X}[L(\theta | X)]$. The performance of any estimate θ is measured according to its excess risk, i.e., $R^*(\theta) - R^*(\theta^*)$. In compressive learning, one obtains $\tilde{\theta}$ by minimizing a cost function $C(\theta | \tilde{z})$, where the sketch $\tilde{z}$ is a compressed version of a finite-size data set $\mathcal{X}$ with samples drawn from distribution p_X .

To gain intuition into how these quantities manifest in compressive learning, Figure S6 shows a simple example: com-

pressive k -means clustering of 1D data $\mathcal{X} = \{\mathbf{x}_i\}_{i=1}^n$ with two centroids, θ_1 and θ_2 . The true risk $R^*(\cdot)$ and the cost $C(\cdot | \tilde{z})$ are plotted versus $\theta = (\theta_1, \theta_2)$ and versus a 1D slice in the θ plane. As the sketch dimension m increases, it can be seen that the excess risk $R^*(\tilde{\theta}) - R^*(\theta^*)$ decreases. Moreover, although the cost function $C(\cdot | \tilde{z})$ is nonconvex, it is approximately quadratic in a large basin of attraction around its global minimizer, suggesting that gradient descent algorithms will behave well when properly initialized.

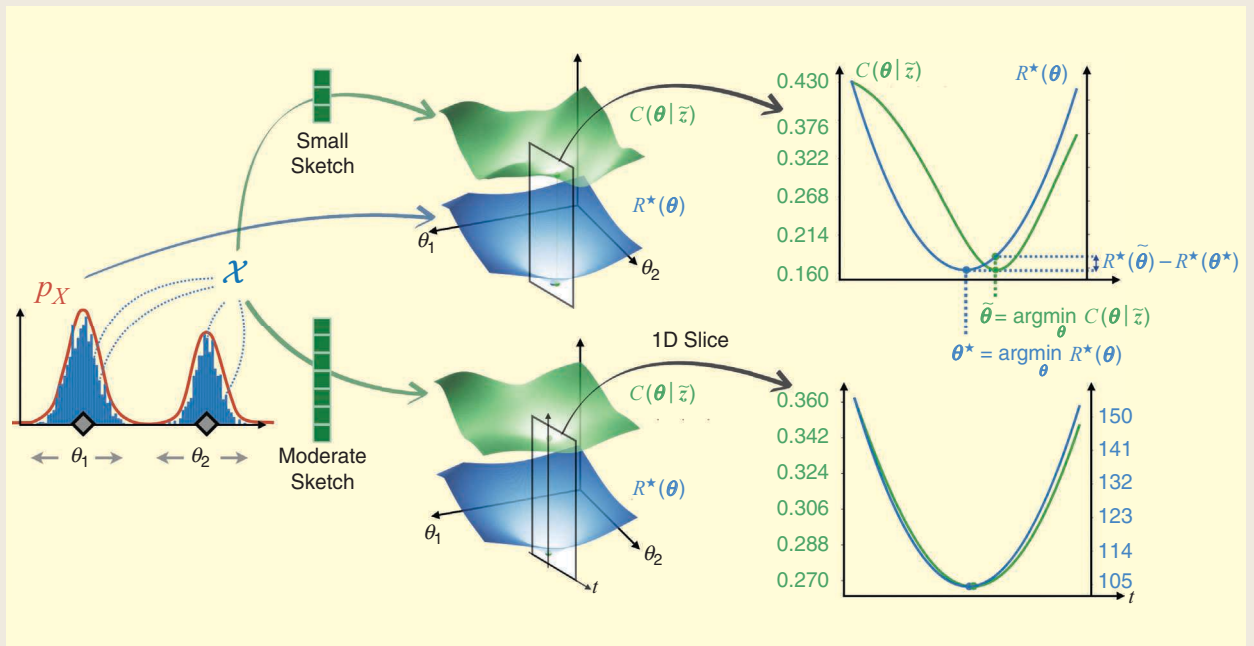


FIGURE S6. The k -means clustering of 1D data $\mathcal{X}$, showing cost function $C(\theta | \tilde{z})$, risk $R^*(\theta)$, ideal estimate θ^* , compressive learning estimate $\tilde{\theta}$, and excess risk $R^*(\tilde{\theta}) - R^*(\theta^*)$, for a sketch of a small-dimension m (top) and for a sketch of a moderate-size m (bottom).

Similarly, excess risk bounds for compressive learning are achieved by interpreting the minimizer $\tilde{\theta}$ of (19) as the minimizer of the risk $R^*(\theta | \tilde{p})$ for some “simple” probability distribution $\tilde{p}$ that minimizes $\|\tilde{z} - \mathcal{A}(\tilde{p})\|$ and by controlling the distance $d(\tilde{p}, p_X)$. To control this distance, key geometric intuitions exploit the connections between compressive learning and CS, as illustrated in Figure 4.

Exploiting simplified models to learn from a sketch

Recall that in CS, the goal is to recover the best k -sparse approximation of an unknown vector $\mathbf{x}_0 \in \mathbb{R}^d$ given the noisy linear measurement $\mathbf{y} = \mathbf{A}\mathbf{x}_0 + \mathbf{e} \in \mathbb{R}^m$, where $m \ll d$. Ideally, recovery aims to find the k -sparse vector whose (noiseless) measurement is closest to the observed $\mathbf{y}$; i.e.,

$$\hat{\mathbf{x}} = \operatorname{argmin}_{\mathbf{x} \in \Sigma_k} \|\mathbf{y} - \mathbf{A}\mathbf{x}\|^2. \quad (25)$$

In practice, convex or greedy approximations of this combinatorial approach are often used. Recovery guarantees can be established using the RIP in (S1), which says that the Euclidean distance $\|\mathbf{A}\mathbf{x} - \mathbf{A}\mathbf{x}'\|$ between the (noiseless) measurements of two k -sparse vectors is almost the same as the Euclidean distance $\|\mathbf{x} - \mathbf{x}'\|$ between the vectors themselves (see the sidebar “Compressive Sensing and the Restricted Isometry Property”). These guarantees require that the sparsity k is sufficiently small for a given m and d .

To connect compressive learning to CS, we first reframe the goal of learning from a sketch as that of recovering a parametric model distribution $p_\theta \in \Sigma_\theta$ from the sketch $\tilde{z}$ rather than recovering the model parameters $\theta \in \Theta$ themselves. Here, $\Sigma_\theta := \{p_\theta : \theta \in \Theta\}$ denotes some set of admissible distributions p_θ . For example, in the compressive GMM, the goal becomes that of recovering a k -term GMM $p_\theta = \sum_{\ell=1}^k \alpha_\ell \mathcal{N}(\boldsymbol{\mu}_\ell, \boldsymbol{\Sigma}_\ell)$ rather than the GMM parameters $\theta = \{\alpha_\ell, \boldsymbol{\mu}_\ell, \boldsymbol{\Sigma}_\ell\}_{\ell=1}^k$. Likewise, in compressive PCA, the goal becomes that of recovering a distribution p_θ whose autocorrelation matrix has a rank of at most k , rather than a k -dimensional orthonormal basis θ for the column space of that autocorrelation matrix. Then, mirroring the CS recovery approach (25), compressive learning recovery aims to find a parametric distribution $p_\theta \in \Sigma_\theta$ whose (noiseless) sketch is closest to the observed sketch $\tilde{z}$; i.e.,

$$\tilde{p} = p_{\tilde{\theta}} := \operatorname{argmin}_{p_\theta \in \Sigma_\theta} \|\tilde{z} - \mathcal{A}(p_\theta)\|^2. \quad (26)$$

The construction of Σ_θ implies that the parameters $\tilde{\theta}$ defining $p_{\tilde{\theta}}$ minimize a certain cost function; i.e.,

$$\tilde{\theta} = \operatorname{argmin}_{\theta \in \Theta} C(\theta | \tilde{z}) \quad \text{with} \quad C(\theta | \tilde{z}) = \|\tilde{z} - \mathcal{A}(p_\theta)\|^2. \quad (27)$$

Note the similarity between (27) and (13) and (14). Moreover, $\tilde{\theta}$ also minimizes the risk $R^*(\theta | \tilde{p})$ [5], [51].

From the preceding description, we see that a central theme of both CS and compressive learning is that measurement compression makes it impossible to recover the full object of interest (i.e., $\mathbf{x}_0$ in CS or p_X in compressive learning) with-

out additional side information. For this reason, both seek to recover the parameters of a simplified model of the object of interest. In CS, this is accomplished by seeking to recover the best k -sparse approximation to $\mathbf{x}_0$ rather than $\mathbf{x}_0$ itself. In compressive learning, this is accomplished by seeking to recover the best model distribution $p_\theta \in \Sigma_\theta$ rather than the true distribution p_X . In both cases, if the linear operator [i.e., $\mathbf{A}$ in CS or $\mathcal{A}(\cdot)$ in compressive learning] is well designed, then the model parameters (in Σ_k for CS or in Σ_θ for compressive learning) can be accurately recovered.

The LRIP and excess risk control

Recovery guarantees can be established using a tool analogous to the RIP (S1) in CS.

Take, for example, the case of compressive PCA. As discussed around (18), one could use a sketch of the form $\tilde{z} = \operatorname{Avec}(\hat{\mathbf{R}}) \in \mathbb{R}^m$, where $\hat{\mathbf{R}} = (1/n) \sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i^\top$ is the $d \times d$ empirical autocorrelation matrix and $m \ll d^2$, and then search for $\hat{\mathbf{R}}$, the symmetric matrix with a rank of at most k that minimizes the cost $C(\mathbf{R} | \tilde{z}) = \|\tilde{z} - \operatorname{Avec}(\mathbf{R})\|^2$. Recovery guarantees can be established using a RIP of the form

$$(1 - \delta) \|\mathbf{R} - \mathbf{R}'\|_F^2 \leq \|\operatorname{Avec}(\mathbf{R}) - \operatorname{Avec}(\mathbf{R}')\|^2 \leq (1 + \delta) \|\mathbf{R} - \mathbf{R}'\|_F^2, \quad \forall \mathbf{R}, \mathbf{R}' \in \Sigma_k, \quad (28)$$

where $\delta \in (0, 1)$ is a tolerance and Σ_k is now the set of $d \times d$ symmetric rank- k matrices. It is known [4, Ch. 9] that the i.i.d. Gaussian $\mathbf{A}$ satisfies the RIP (28) with high probability when m is on the order of kd times a constant that depends on the tolerance δ . Furthermore, the cost-minimizing $\hat{\mathbf{R}}$ is near optimal in the sense of minimizing the excess risk, even when a convex relaxation of the cost is used [5], [51].

These compressive PCA guarantees hold even under model mismatch: although Σ_k constrains $\hat{\mathbf{R}}$ to have a rank of at most k , the empirical autocorrelation matrix $\hat{\mathbf{R}}$ used to construct $\tilde{z}$ tends to have a full-rank d in practice. To explore this idea in more detail, suppose for the moment that the data $\mathcal{X}$ lie in a rank- k subspace. In this case, $\hat{\mathbf{R}}$ would have a rank of at most k , and there exists a symmetric $\mathbf{R}$ with a rank of at most k that drives the cost $C(\mathbf{R} | \tilde{z}) = \|\operatorname{Avec}(\hat{\mathbf{R}}) - \operatorname{Avec}(\mathbf{R})\|^2$ to zero. Furthermore, when $\mathbf{A}$ satisfies the RIP (28), the left inequality in (28) implies that this cost-minimizing $\hat{\mathbf{R}}$ must equal $\hat{\mathbf{R}}$. When the data $\mathcal{X}$ do not lie in a rank- k subspace, the minimal cost will be nonzero. In this case, an upper bound on the error $\|\hat{\mathbf{R}} - \tilde{\mathbf{R}}\|_F$ of the cost-minimizing estimate $\tilde{\mathbf{R}}$ as well as an upper bound on the excess risk of the principal subspace $\tilde{\theta}$ of $\hat{\mathbf{R}}$ can both be obtained as a consequence of the RIP (28).

In place of the RIP, compressive learning uses the so-called LRIP,

$$d(p_\theta, p_{\theta'}) \leq C_0 \|\mathcal{A}(p_\theta) - \mathcal{A}(p_{\theta'})\|, \quad (29)$$

assumed to be valid for each pair $p_\theta, p_{\theta'} \in \Sigma_\theta$, where C_0 is a positive constant. The LRIP says that the Euclidean distance $\|\mathcal{A}(p_\theta) - \mathcal{A}(p_{\theta'})\|$ between the (noiseless) sketches of two

distributions is—up to a scaling—controlling the distance $d(p_\theta, p_{\theta'})$ between the distributions themselves (22). Note that the lower bound $d(\cdot, \cdot)$ in (29) is task specific and not Euclidean as in the CS case (S1).

The first main theoretical guarantee about compressive learning [5], [51] is that when the LRIP (29) holds, the estimate $\tilde{\theta}$ obtained by minimizing (27) automatically has controlled excess risk. In particular, using $\tilde{\theta} = \operatorname{argmin}_\theta R^*(\theta | p_\theta)$ [5], [51], the excess risk bound (24) can be combined with the LRIP (29) and the definition of $\tilde{\theta}$ (27) as follows:

$$\begin{aligned} R^*(\tilde{\theta} | p_X) - R^*(\theta^* | p_X) &\leq 2d(p_{\tilde{\theta}}, p_X) \\ &\leq 2[d(p_{\tilde{\theta}}, p_{\theta'}) + d(p_{\theta'}, p_X)] \\ &\leq 2[C_0 \|\mathcal{A}(p_{\tilde{\theta}}) - \mathcal{A}(p_{\theta'})\| + d(p_{\theta'}, p_X)] \\ &\leq 2[C_0 \|\mathcal{A}(p_{\tilde{\theta}}) - \tilde{z}\| + C_0 \|\tilde{z} - \mathcal{A}(p_{\theta'})\| + d(p_{\theta'}, p_X)] \\ &\leq 2[2C_0 \|\mathcal{A}(p_{\theta'}) - \tilde{z}\| + d(p_{\theta'}, p_X)] \\ &\leq 2[2C_0 \|\mathcal{A}(p_X) - \tilde{z}\| + 2C_0 \|\mathcal{A}(p_{\theta'}) - \mathcal{A}(p_X)\| \\ &\quad + d(p_{\theta'}, p_X)], \end{aligned} \quad (30)$$

for any distribution $p_{\theta'}$ in Σ_θ [because $d(\cdot, \cdot)$ is a valid distance metric; i.e., $d(p_1, p_2) \leq d(p_1, p_3) + d(p_3, p_2) \forall p_3$]. One option is to choose $p_{\theta'} = p_\theta$, in which case the term $2C_0 \|\mathcal{A}(p_{\theta'}) - \mathcal{A}(p_X)\| + d(p_{\theta'}, p_X)$ in the upper bound reflects the excess risk due to modeling and $2C_0 \|\mathcal{A}(p_X) - \tilde{z}\|$

reflects the excess risk due to sketching from a finite data set $\mathcal{X}$. For a tighter bound, we could choose $p_{\theta'}$ as the distribution in Σ_θ that minimizes the right side of (30). These approaches and more refined variants have been applied to analyze various learning tasks in, e.g., [5], [11], [12], and [51].

It should be emphasized that recovery guarantees based on the RIP (S1) or LRIP (29) hold even in the presence of measurement noise e and/or modeling errors. In CS, measurement noise arises due to, e.g., thermal noise or interference, while modeling errors arise when x_0 is not truly k -sparse. Many sparse recovery techniques are provably robust to such noise and modeling errors; see (S2) in the sidebar “Compressive Sensing and the Restricted Isometry Property.” In compressive learning, measurement noise arises due to the finite cardinality of the data set $\mathcal{X}$ [recall (7)], while modeling errors arise when $p_X \notin \Sigma_\theta$. For example, GMM recovery guarantees can be established even when p_X is not truly a GMM.

Establishing the LRIP via kernels and the JL lemma

In light of the fact that the LRIP (29) yields statistical learning guarantees (30) for compressive learning, a key question becomes: How can we choose the sketch dimension m so that the LRIP (29) holds? Similar to how the RIP is proven in CS (see, e.g., [4, Lemma 9.33]), refinements of a mathematical tool called the *JL lemma* [31] can be used to obtain a value of m

Kernel Methods and Kernel Embeddings of Probability Distributions

Sketching shares connections with kernel methods [3], a family of machine learning techniques that produces decisions or insights using a kernel function, $\kappa(x, x') \in \mathbb{R}$, which measures the “similarity” between x and x' . A kernel is said to be “positive definite” if the $n \times n$ matrix $\mathbf{K}$, constructed with entries $\kappa(x_i, x_j)$ for $1 \leq i, j \leq n$, is positive semidefinite for every possible $\{x_i\}_{i=1}^n$. The celebrated “kernel trick” states that any positive definite kernel implicitly amounts to an inner product in some higher-dimensional (and potentially infinite-dimensional) feature space $\mathcal{H}$ and vice versa. That is, $\kappa(x, x') = \langle \Phi(x), \Phi(x') \rangle$ for some (not necessarily explicitly known) mapping $\Phi(\cdot)$ from the signal space to $\mathcal{H}$. Any machine learning method that relies only on the evaluation of inner products—such as ridge regression, support vector machine classification, PCA [29], and dictionary learning [S3]—can be “kernelized” by using a kernel in place of the inner product. Kernelizing a method is thus tantamount to applying that method in a transformed, higher-dimensional feature space. In this way, more complex estimation and/or decision functions can be implemented.

Given a positive definite kernel $\kappa(x, x')$ operating on signals x and x' in some set, it is possible to “lift” $\kappa(\cdot, \cdot)$ to a positive definite kernel operating on probability distributions over this set by defining the so-called mean kernel

$$k(p, q) := \mathbb{E}_{x \sim p, x' \sim q} [\kappa(x, x')]. \quad (S5)$$

This defines an embedding of probability distributions into a kernel space, which is analogous to the finite-dimensional embedding $\mathcal{A}(p)$ of probability distribution p from (6). The maximum mean discrepancy (MMD) [S4], [S5]

$$\text{MMD}(p, q) := \sqrt{k(p, p) + k(q, q) - 2k(p, q)} \quad (S6)$$

is the Euclidean metric naturally induced by the mean kernel. It is analogous to the Euclidean distance between sketches, $\|\mathcal{A}(p) - \mathcal{A}(q)\|$. The MMD, originally introduced in the context of two-sample hypothesis testing [S4], is now well known in machine learning. When the MMD behaves as a true metric, i.e., when $\text{MMD}(p, q) = 0 \Leftrightarrow p = q$, the mean kernel $k(\cdot, \cdot)$ is said to be “characteristic.” In $\mathbb{R}^d$, many classic kernels $\kappa(\cdot, \cdot)$, such as the Gaussian and Laplace kernels, yield characteristic mean kernels [S5].

References

- [S3] R. Rubinstein, A. M. Bruckstein, and M. Elad, “Dictionaries for sparse representation modeling,” *Proc. IEEE*, vol. 98, no. 6, pp. 1045–1057, 2010. doi: 10.1109/JPROC.2010.2040551.
- [S4] A. Gretton, K. M. Borgwardt, M. J. Rasch, B. Schölkopf, and A. J. Smola, “A Kernel method for the two-sample problem,” in *Proc. Adv. Neural Inform. Process. Syst. (NIPS)*, 2007.
- [S5] B. K. Sriperumbudur, A. Gretton, K. Fukumizu, B. Schölkopf, and G. R. Lanckriet, “Hilbert space embeddings and metrics on probability measures,” *J. Mach. Learning Res.*, vol. 11, no. 10, pp. 1517–1561, Mar. 2010.

sufficient for the LRIP to hold with high probability on the draw of the random feature map $\Phi(\cdot)$. We now summarize this approach.

To begin, we establish connections to kernel methods (see the sidebar “Kernel Methods and Kernel Embeddings of Probability Distributions” for a brief review of kernels). Some key observations are that any feature map explicitly defines a positive definite kernel and that the expectation of this kernel defines another useful positive definite kernel. To see why, consider, for example, the RF feature map (3). First, for a fixed set of frequency vectors $\{\mathbf{w}_j\}_{j=1}^m$, we have, for all $\mathbf{x}, \mathbf{x}'$,

$$\left\langle \frac{1}{\sqrt{m}} \Phi(\mathbf{x}), \frac{1}{\sqrt{m}} \Phi(\mathbf{x}') \right\rangle = \frac{1}{m} \sum_{j=1}^m \exp(-j2\pi \mathbf{w}_j^\top (\mathbf{x} - \mathbf{x}')), \quad (31)$$

where $\langle \cdot, \cdot \rangle$ is the standard Euclidean inner product in $\mathbb{R}^m$ or $\mathbb{C}^m$. The left-hand side defines a positive definite kernel according to the terminology reviewed in the sidebar “Kernel Methods and Kernel Embeddings of Probability Distributions.” This kernel is also “shift-invariant” in that it depends only on the difference $\mathbf{x} - \mathbf{x}'$. Next, imagine that the d -dimensional frequency vectors $\mathbf{w}_j$, which constitute the m rows of the random matrix $\mathbf{W}$, are drawn i.i.d. from some probability distribution p_W . In the limit of large sketch dimension m , the law of large numbers combined with (31) says

$$\left\langle \frac{1}{\sqrt{m}} \Phi(\mathbf{x}), \frac{1}{\sqrt{m}} \Phi(\mathbf{x}') \right\rangle \xrightarrow{\text{w.p.1}} \mathbb{E}_W \exp(-j2\pi \mathbf{W}^\top (\mathbf{x} - \mathbf{x}')). \quad (32)$$

The right-hand side of (32) is the CFT of the probability density p_W evaluated at $\mathbf{x} - \mathbf{x}'$, i.e., the characteristic function $\Psi_{p_W}(\mathbf{x}' - \mathbf{x})$ using the notation from (12). Because $\Psi_{p_W}(\cdot)$ is the CFT of a nonnegative function, it is a so-called positive definite function, which means that for any $\{\mathbf{x}_i\}_{i=1}^n$, the $n \times n$ matrix Ψ defined with elements $\Psi_{ij} = \Psi_{p_W}(\mathbf{x}_i - \mathbf{x}_j)$ for $1 \leq i, j \leq n$ will be positive semidefinite. Thus, if we construct a kernel as $\kappa(\mathbf{x}, \mathbf{x}') := \Psi_{p_W}(\mathbf{x}' - \mathbf{x})$, then it will be a positive definite kernel. When $p_W = \mathcal{N}(\mathbf{0}, \sigma_w^2 \mathbf{I}_d)$, this approach yields the familiar Gaussian kernel (a particular type of “radial basis function”); i.e., $\kappa_\sigma(\mathbf{x}, \mathbf{x}') := \exp(-\|\mathbf{x} - \mathbf{x}'\|^2 / 2\sigma^2)$, here of width $\sigma = 1/\sigma_w$.

More generally, by considering any parametric feature map of the form $\Phi(\mathbf{x} | \mathbf{W})$, where the parameter $\mathbf{W}$ is drawn at random according to some probability distribution, one can define [in many papers, the $1/\sqrt{m}$ scaling in (31)–(33) is subsumed in the feature map $\Phi(\cdot)$] the “expected kernel”

$$\kappa(\mathbf{x}, \mathbf{x}') := \mathbb{E}_W \left\langle \frac{1}{\sqrt{m}} \Phi(\mathbf{x} | \mathbf{W}), \frac{1}{\sqrt{m}} \Phi(\mathbf{x}' | \mathbf{W}) \right\rangle, \quad (33)$$

not to be confused with the “mean kernel” defined in (S5) (see the sidebar “Kernel Methods and Kernel Embeddings of Probability Distributions”). This setting includes RF features (3) with i.i.d. frequencies $\mathbf{w}_j$ (as in the preceding) or with a

frequency matrix $\mathbf{W}$ that includes structured blocks of rows, as we will soon discuss.

Now that the kernel connections have been established, we return to our original objective of understanding when the LRIP (29) holds. Importantly, the law of large numbers shows that, in the limit of large sketch dimension m , the right side of (29) is related to the maximum mean discrepancy (MMD) from (S6), which is a kernel-based distance between distributions (see the sidebar “Kernel Methods and Kernel Embeddings of Probability Distributions”). Indeed, for arbitrary probability distributions p and q , we have

$$\lim_{m \rightarrow \infty} \frac{1}{\sqrt{m}} \|\mathcal{A}(p) - \mathcal{A}(q)\| \stackrel{\text{w.p.1}}{=} \text{MMD}(p, q), \quad (34)$$

where the maximum mean discrepancy is implicitly the one obtained from the kernel used to build the sketching operator $\mathcal{A}$. Thus, when the LRIP (29) holds, it must also be true that, for a sufficiently large sketch dimension m ,

$$d(p_\theta, p_{\theta'}) \leq C'_0 \text{MMD}(p_\theta, p_{\theta'}), \quad \text{for all } p_\theta, p_{\theta'} \in \Sigma_\theta, \quad (35)$$

where C'_0 is a positive constant. Property (35) connects two different metrics on probability distributions: the left-hand side of (35) is defined by the learning task [recall (22)], while the right-hand side of (35) is defined by the expected kernel κ , from (33), associated with the randomized feature map. Note that (35) is a deterministic property; it does not depend on the draw of the randomized feature map, unlike the LRIP (29). In the literature, (35) is called the *kernel LRIP* [5], [51] because it is a kernel-based analog to the LRIP. Being deterministic, the expected kernel is often easier to manipulate than the random feature map, and thus it eases the proof of the kernel LRIP. This is important because, if the kernel LRIP (35) holds, then, using arguments based on the JL lemma, one can also establish [5], [51] the LRIP (29), as we show in the following.

The JL lemma [31] is a precursor of the RIP that is specialized to finite sets (S1). It states that given N arbitrary d -dimensional vectors $\mathbf{x}_i$, there exists an $m \times d$ matrix $\mathbf{A}$, with m on the order of $\log N$, such that $\|\mathbf{A}\mathbf{x}_i - \mathbf{A}\mathbf{x}_j\| \approx \|\mathbf{x}_i - \mathbf{x}_j\|$ for all i, j .

To illustrate how the JL lemma is useful in the context of compressive learning, let us momentarily restrict our attention to a learning problem where the collection Σ_θ of parametric distributions is of finite cardinality, noting that we can extend this approach to continuous families of parametric distributions through discretization arguments involving the notion of covering numbers [4, Appendix C], as described in the following. As a concrete example, let us consider a discretized variant of the compressive GMM using the RF feature map (3). As in [5], [6], and [51], we assume that the d -dimensional frequency vectors $\mathbf{w}_j$ are drawn i.i.d. from the normal distribution $\mathcal{N}(\mathbf{0}, \sigma_w^2 \mathbf{I}_d)$, and we consider learning a mixture of Gaussian components $p_{\theta_\ell} = \mathcal{N}(\boldsymbol{\mu}_\ell, \mathbf{I})$ for $\ell = 1, \dots, k$, with equal weights $\alpha_\ell = 1/k$ for each ℓ , where the means $\boldsymbol{\mu}_\ell$ are assumed to be bounded, separated, and discretized on a regular grid. In this setting, there exists a finite number, N , of possible parametric mixture distributions $p_\theta \in \Sigma_\theta$. As long as the MMD

is a true metric (as defined in the sidebar “Kernel Methods and Kernel Embeddings of Probability Distributions”), the ratio $d(p_\theta, p_{\theta'})/\text{MMD}(p_\theta, p_{\theta'})$ is finite for $p_\theta \neq p_{\theta'}$; hence, whenever Σ_Θ is a finite collection, as in this concrete example, the kernel LRIP (35) holds (for a sufficiently large C'_0). Moreover, specializing to an arbitrary pair of mixtures $p_\theta, p_{\theta'} \in \Sigma_\Theta$ and a given sketch dimension m , refinements of (34) (using measure concentration) enable one to show that the lower bound

$$\frac{1}{\sqrt{2}}\text{MMD}(p_\theta, p_{\theta'}) \leq \frac{1}{\sqrt{m}}\|\mathcal{A}(p_\theta) - \mathcal{A}(p_{\theta'})\| \quad (36)$$

holds with a probability of at least $1 - \exp(-c_0 m)$, where c_0 is a positive concentration constant. Since there are only N^2 pairs of parametric mixtures $p_\theta, p_{\theta'} \in \Sigma_\Theta$, the lower bound (36) is valid uniformly for all of them, except with a probability of at most $N^2 \exp(-c_0 m)$. This failure probability can be made smaller than any $\epsilon > 0$ by choosing an m larger than $2c_0^{-1} \log(N/\sqrt{\epsilon})$. Combining (35) and (36) using a union bound, it can be shown [5], [51] that the LRIP (29) holds with high probability when C_0 is on the order of $C'_0 \sqrt{m}$ and the sketch dimension m grows logarithmically with N , the number of parametric distributions in the finite set Σ_Θ .

For infinite collections Σ_Θ , proving the kernel LRIP (35), and eventually the LRIP (29), is more technical and can require some additional assumptions. As an example, for compressive clustering with RF features, it can be proven that there is a constant C'_0 such that (35) holds, provided that the centroids are sufficiently separated and bounded [5], [51]. The JL lemma can be extended by refining techniques used to establish the RIP (S1). The main idea is to first prove the LRIP on some finite collection $\Sigma' \subset \Sigma_\Theta$ of N probability distributions and then to extrapolate it (with slightly worse constants) to Σ_Θ . Technically, this involves the notion of covering numbers, and the cardinality $N = |\Sigma'|$ is typically exponential in the number of parameters needed to describe Σ_Θ . For example, for the GMM, one needs kd parameters to describe the means $\mu_\ell \in \mathbb{R}^d$, $1 \leq \ell \leq k$ of the mixture $p_\theta = \sum_{\ell=1}^k (1/k) \mathcal{N}(\mu_\ell, \mathbf{I})$, and $\log N$ essentially depends linearly on kd . The dimension m of the sketch for which the LRIP holds with high probability is thus on the order of kd/c_0 , up to some additional factors due to the proof technique [5], [51].

Empirical studies of compressive clustering [9], [10] and the compressive GMM [6] suggest that a sketch dimension m on the order of kd (the number of parameters in these settings) is sufficient to yield accurate learning performance. The best-known bounds on provably good sketch dimensions [5], [51] remain pessimistic compared to these empirically validated sketch dimensions. This is most likely related to suboptimal bounds for the concentration constant c_0 and/or shortcomings in the techniques used to extend the LRIP from a finite collection Σ' to an infinite collection Σ_Θ .

The challenge of designing a feature map given a learning task

In the existing literature, the LRIP has been established [5], [51] for randomized feature maps $\Phi(\cdot)$ (e.g., RF features and

random quadratic features) that mimic related constructions from CS, developed either for sparse vector recovery or low-rank matrix recovery.

When sketching with RF features (e.g., for compressive clustering and the compressive GMM), the main design choice for $\Phi(\cdot)$ is the distribution from which to draw the random frequencies $\mathbf{w}_j$ [i.e., the rows of $\mathbf{W}$ in (3)]. In light of the connections to shift-invariant kernels [recall (31)], this design task is a particular instance of the difficult problem of kernel design [3, Sec. 4.4.5].

For example, when the rows of $\mathbf{W}$ are drawn i.i.d. $\mathcal{N}(\mathbf{0}, \sigma_w^2 \mathbf{I})$, the choice of the variance σ_w^2 determines the choice of the width $\sigma = 1/\sigma_w$ of the corresponding Gaussian kernel $\kappa_\sigma(\cdot, \cdot)$. Indeed, from a signal processing standpoint, the corresponding mean kernel $k_\sigma(\cdot, \cdot)$ [recall (S5)] acts to low-pass-filter the underlying data distributions. To see why, observe that $\kappa_\sigma(\mathbf{x}, \mathbf{x}') = g_\sigma(\mathbf{x}, \mathbf{x}')$ with $g_\sigma(\mathbf{x}) := e^{-\|\mathbf{x}\|^2/2\sigma^2}$, and so

$$\begin{aligned} k_\sigma(p, q) &= \mathbb{E}_{X \sim p, X' \sim q} [\kappa_\sigma(X, X')] \\ &= \iint e^{-\|X - X'\|^2/2\sigma^2} dp(X) dq(X') \end{aligned} \quad (37)$$

$$= \langle g_\sigma * p, q \rangle_{L^2} = \langle g_\sigma * p, g_\sigma * q \rangle_{L^2}, \quad (38)$$

where $*$ denotes convolution and $\bar{\sigma} = \sigma/\sqrt{2}$. Hence, the associated MMD (S6) satisfies $\text{MMD}(p, q) = \|g_\sigma * p - g_\sigma * q\|_{L^2}$. Recall that learning from a sketch is often performed by minimizing a “sketch matching” cost $\|\tilde{\mathbf{z}} - \mathcal{A}(p_\theta)\|^2 = \|\mathcal{A}(\hat{p}_\mathcal{X}) - \mathcal{A}(p_\theta)\|^2$, as in (26), where $\hat{p}_\mathcal{X}$ denotes the empirical distribution of the data $\mathcal{X}$. In the limit of large sketch dimension m , this cost compares the smoothed versions of the probability distributions $\hat{p}_\mathcal{X}$ and p_θ since $(1/m) \|\tilde{\mathbf{z}} - \mathcal{A}(p_\theta)\|^2 \approx \|g_\sigma * \hat{p}_\mathcal{X} - g_\sigma * p_\theta\|_{L_2}^2$. Similarly, when using a greedy algorithm to learn a mixture model (or cluster centroids) from a sketch, the normalized inner product $(1/m) \langle \tilde{\mathbf{z}}, \mathcal{A}(p_\theta) \rangle$ approximates the correlation $\langle g_\sigma * \hat{p}_\mathcal{X}, g_\sigma * p_\theta \rangle_{L_2}$ between the low-passed versions of the empirical data distribution $\hat{p}_\mathcal{X}$ and the candidate mixture component p_θ , respectively.

This latter idea is illustrated for compressive clustering in Figure 5. There, since the mixture component associated to a candidate centroid $\mathbf{c}$ is the Dirac $p_c(\mathbf{x}) = \delta(\mathbf{x} - \mathbf{c})$ [recall the discussion before (15)], we have that $\langle \tilde{\mathbf{z}}_\sigma, \mathcal{A}(p_c) \rangle = \langle \tilde{\mathbf{z}}_\sigma, \Phi_\sigma(\mathbf{c}) \rangle$, where the dependence on the kernel width $\sigma = 1/\sigma_w$ has been made explicit. Meanwhile,

$$\langle g_\sigma * \hat{p}_\mathcal{X}, g_\sigma * p_c \rangle_{L_2} = (g_\sigma * \hat{p}_\mathcal{X})(\mathbf{c}) = \frac{1}{n} \sum_{i=1}^n g_\sigma(\mathbf{c} - \mathbf{x}_i). \quad (39)$$

In (39), we recognize a Parzen window density estimator [32], whose computation for a given $\mathbf{c}$ requires access to the n training samples $\mathbf{x}_i$. In contrast, its surrogate $\langle \tilde{\mathbf{z}}_\sigma, \Phi_\sigma(\mathbf{c}) \rangle$ only requires access to the m -dimensional sketch $\tilde{\mathbf{z}}$. In a large-scale setting, this can save huge amounts of memory and computation. As can be seen when comparing the top and bottom rows in Figure 5, $\langle \tilde{\mathbf{z}}_\sigma, \Phi_\sigma(\mathbf{c}) \rangle$ well approximates $(g_\sigma * \hat{p}_\mathcal{X})(\mathbf{c})$ for a sufficiently large kernel width σ . By comparing the different columns of Figure 5, it can also be seen that the kernel width σ should be chosen compatible with the cluster width and separation. This choice involves

a tradeoff between smoothing the unwanted gaps between data samples and oversmoothing the (desired) gaps between clusters. Existing theory [5], [51] identifies sufficient conditions on the choice of σ related to the number k of candidate clusters and their minimum separation.

Although i.i.d. Gaussian frequencies w_j were used with the RF map in the preceding, one may also consider the use of i.i.d. non-Gaussian frequencies. Such designs, as proposed in [7], can yield improved empirical behavior. As we will see in the next section, it is also possible to deviate from the RF map, with the goal of improving computational efficiency. Such constructions also yield non-Gaussian expected kernels (33). Although they work well in practice, there is currently no proof that these latter kernels satisfy the kernel LRIP.

While existing theory focuses on proving the LRIP for a given random feature map and learning task, an important open question is: How should one design the feature map to best match a given learning task? In particular, can we design a random feature map that satisfies the LRIP (29) for a given learning task defined by a loss function $L(\theta | x_i)$ and embodied by a task-driven distance (22)? A promising yet still challenging avenue would be to first identify a positive definite kernel $\kappa_0(x, x')$ for which the corresponding MMD satisfies the kernel LRIP (35) and then use Bochner's theorem or Mercer's theorem (see the sidebar "Approximating Kernel Methods With Random

Feature Maps") to design a random feature map $\Phi(\cdot)$ whose expected kernel (33) is precisely κ_0 .

Sketching with reduced computational resources

The computational cost of sketching via (1) is heavily dependent on the feature map $\Phi(\cdot)$. The computational cost of parameter estimation via (13) is also heavily dependent on $\Phi(\cdot)$, since it often involves iterative application of $\Phi(\cdot)$.

Often, the feature map is constructed as a randomized linear operation followed by a componentwise nonlinear operation; i.e.,

$$\Phi(x) = \varrho(Wx), \quad (40)$$

where W is a (randomly drawn) matrix of size $m \times d$ and $\varrho(\cdot)$ applies a scalar nonlinear function identically to each element of the vector Wx . For example, the feature maps described earlier for compressive PCA, the GMM, and clustering all have this form. Reducing the computational cost of each stage has been the goal of several studies. For example, using a fast transform for W drastically reduces the memory and computational complexity demands relative to an explicit matrix. Also, quantized versions of the nonlinearity $\varrho(\cdot)$ are much more easily implemented in hardware than, say, the complex exponential nonlinearity $\varrho_{\text{RF}}(\cdot) := \exp(-j2\pi \cdot)$ used in the RF map (3). We discuss such constructions of W and $\varrho(\cdot)$ in the following.

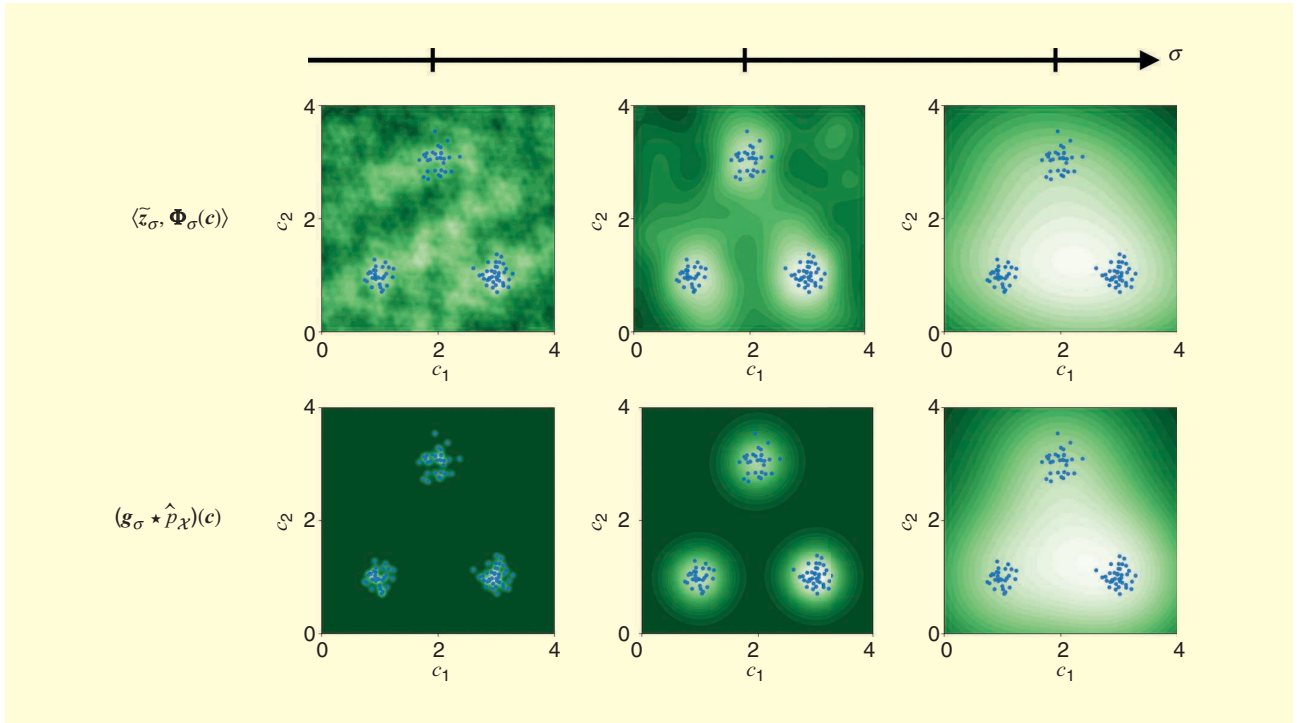


FIGURE 5. The criterion $\langle \tilde{z}_\sigma, \Phi_\sigma(c) \rangle$ used by greedy parameter estimation algorithms in compressive clustering with RF features (top row) versus its expected value $(g_\sigma \star \hat{p}_\lambda)(c)$ (bottom row) as a function of the centroid hypothesis location $c = [c_1, c_2]^T$. The data set (in blue) consists of $n = 100$ points drawn according to a mixture of $k = 3$ isotropic Gaussians. The frequencies w_j used to define the feature map $\Phi_\sigma(\cdot) = \exp(-j2\pi W \cdot)$ are drawn according to a standard Gaussian $\mathcal{N}(\mathbf{0}, \sigma_w^2 I)$ with $\sigma_w = 1/\sigma$. The sketch $\tilde{z}_\sigma$ is computed with the feature map $\Phi_\sigma(\cdot)$. With $\sigma^2 = 1/500$ (left), there is insufficient smoothing, and the criterion displays many spurious local maxima; with $\sigma^2 = 1/10$ (middle), there is appropriate smoothing, and local maxima are in good correspondence with the true cluster centers; with $\sigma^2 = 1$ (right), there is oversmoothing, and the criterion displays only a single maximum.

Sketching with structured random matrices

In most of our previous examples, we constructed $\mathbf{W}$ by drawing its m rows i.i.d. from the normal distribution $\mathcal{N}(\mathbf{0}, \sigma_w^2 \mathbf{I}_d)$ with some variance $\sigma_w^2 > 0$. Recall that, with the RF map $\varrho_{\text{RF}}(\cdot)$, the rows of $\mathbf{W}$ correspond to the (d -dimensional) frequencies used when sampling the CFT of the (empirical) data distribution. In any case, when $\mathbf{W}$ is an explicit matrix, the computational complexity of computing $\Phi(\mathbf{x})$ of the form (40) is dominated by the matrix–vector product $\mathbf{W}\mathbf{x}$. Thus, it is on the order of md , which is also the order of the memory needed to store $\mathbf{W}$.

As an alternative to these approaches, it has been suggested to construct $\mathbf{W}$ as a structured random matrix with a fast implementation that mimics an i.i.d. Gaussian matrix. Multiple ways of accomplishing this goal have been proposed in the literature. We focus on the approach suggested in [33], which was successfully applied to compressive learning in [34]. There, the idea is to construct $\mathbf{W}$ as a vertical concatenation of $b = \lceil m/d \rceil$ blocks $\{\mathbf{B}_j\}_{j=1}^b$, each of size $d \times d$. These blocks have the form

$\mathbf{B}_j = \mathbf{D}_j^{(0)} \mathbf{H} \mathbf{D}_j^{(1)} \mathbf{H} \mathbf{D}_j^{(2)} \mathbf{H} \mathbf{D}_j^{(3)}$, where $\mathbf{H}$ is the Walsh–Hadamard matrix and $\mathbf{D}_j^{(k)}$ are random diagonal matrices. In particular, the diagonal elements of $\mathbf{D}_j^{(1)}$, $\mathbf{D}_j^{(2)}$, and $\mathbf{D}_j^{(3)}$ are drawn i.i.d. from the uniform distribution over $\{-1, 1\}$ and the diagonal elements of $\mathbf{D}_j^{(0)}$ are drawn i.i.d. from the $\mathcal{X}$ distribution with d degrees of freedom, which is the distribution of the norm of a d -variate Gaussian vector. This construction is depicted in Figure 6. The fast Walsh–Hadamard transform offers an order $d \log d$ -complexity implementation of the matrix–vector multiplication $\mathbf{H}\mathbf{x}$ and prevents the need to explicitly store $\mathbf{H}$. With this structured and fast incarnation of $\mathbf{W}$, the sketching complexity shrinks from order md to order $m \log d$. Moreover, since only the diagonal matrices need to be stored, the storage cost shrinks from order md to order m .

Sketching with quantized contributions

With the RF map (3), which is commonly used in compressive clustering and the GMM, the nonlinear operation $\varrho(\cdot)$ in (40) becomes $\varrho_{\text{RF}}(\cdot) := \exp(-j2\pi \cdot)$. Since implementing

Approximating Kernel Methods With Random Feature Maps

While each random feature map $\Phi(\cdot)$ implicitly defines a positive definite kernel by taking the expectation (33), the converse is also true. For shift-invariant kernels, i.e., kernels for which $\kappa(x, x') = \kappa(x - x', 0)$ only depends on the difference $x - x'$ (such as the Gaussian kernel), this is a consequence of Bochner’s theorem [7], which states that if κ is a positive definite kernel such that $\kappa(x, x) = 1$ for all x , then its CFT yields a probability distribution $p_w(w) = \int \kappa(x, 0) \exp(-j2\pi w^\top x) dx$. Conversely, the kernel can be obtained by the inverse CFT, which can also be phrased as an expectation:

$$\begin{aligned} \kappa(x, x') &= \int \exp(j2\pi w^\top (x - x')) p_w(w) dw \\ &= \mathbb{E}_{w \sim p_w} \exp(j2\pi W^\top (x - x')). \end{aligned}$$

Hence, drawing i.i.d. frequency vectors w_j according to p_w yields a random Fourier feature map (3) whose expected kernel is precisely κ . For instance, a Laplace kernel can be approximated if the rows of $\mathbf{W}$ are drawn i.i.d. from the Cauchy distribution [S6].

More generally, under mild assumptions on a positive definite kernel κ , one can invoke Mercer’s theorem [S7] to similarly show the existence of a random feature map whose expected kernel, in the sense of (33), matches κ .

References

- [S6] P. T. Boufounos, S. Rane, and H. Mansour, “Representation and coding of signal geometry,” *Inform. Inference*, vol. 6, p. 4, Dec. 2017.
- [S7] F. Bach, “On the equivalence between Kernel quadrature rules and random feature expansions,” *J. Mach. Learning Res.*, vol. 18, no. 1, pp. 714–751, 2017.

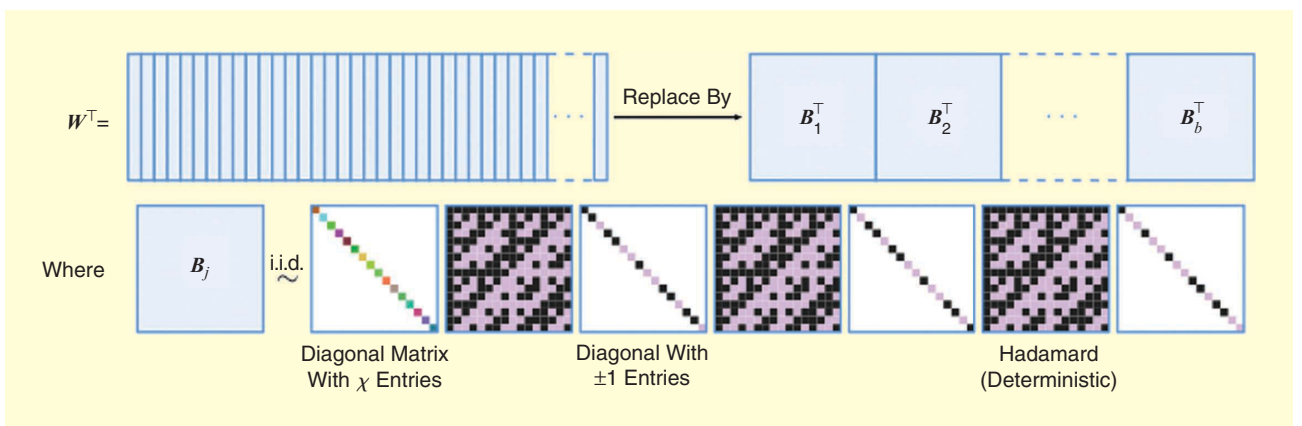


FIGURE 6. The structured random matrix design from [40]. Each block is a composition of several Hadamard and diagonal matrices. For convenience, we draw $\mathbf{W}^T$ instead of $\mathbf{W}$.

$\varrho_{\text{RF}}(\cdot)$ with high accuracy is somewhat costly and not amenable to easy hardware acceleration, one might consider quantizing it. For example, dropping the imaginary part of $\varrho_{\text{RF}}(\cdot)$ and quantizing the real part to 1 bit of precision yields the 2π -periodic function $\varrho_q(\cdot) := \text{sign}(\cos(2\pi \cdot))$ which is simply a “square wave.”

To alleviate the effects of quantization, one can apply dithering [8] to the input. In this case, the feature map becomes $\Phi_q(\mathbf{x}) = \varrho_q(\mathbf{W}\mathbf{x} + \boldsymbol{\xi}) \in \{-1, +1\}^m$, with i.i.d. dither components ξ_j drawn uniformly over $[0, 1)$. The effect of dithering is to make the quantized Φ_q behave similarly to the nonquantized Φ_{RF} , on average. For instance, it was shown in [35] that for each $\mathbf{W}$, $\mathbf{x}$, $\mathbf{x}'$, and $\boldsymbol{\xi}$,

$$\begin{aligned} \langle \Phi_q(\mathbf{x}), \Phi_{\text{RF}}(\mathbf{x}') \rangle &\approx \mathbb{E}_{\boldsymbol{\xi}} \langle \Phi_q(\mathbf{x}), \Phi_{\text{RF}}(\mathbf{x}') \rangle \\ &= c \langle \Phi_{\text{RF}}(\mathbf{x}), \Phi_{\text{RF}}(\mathbf{x}') \rangle, \end{aligned} \quad (41)$$

where c is a constant. The approximation in the preceding is accurate for a typical draw of the m -dimensional dither vector $\boldsymbol{\xi}$ when the sketch dimension m is large enough [36]. Note that, when this dithered quantizer is used to compute a sketch $\tilde{\mathbf{z}}_q := (1/n) \sum_{i=1}^n \Phi_q(\mathbf{x}_i) = (1/n) \sum_{i=1}^n \varrho_q(\mathbf{W}\mathbf{x}_i + \boldsymbol{\xi})$, it is important to use the same dither realization $\boldsymbol{\xi}$ for all samples i .

The question then arises: When estimating parameters $\boldsymbol{\theta}$ via (27), how should we account for quantization in the sketch? Simply replacing the $\mathcal{A}(p_{\boldsymbol{\theta}_i})$ term with $\mathcal{A}_q(p_{\boldsymbol{\theta}_i}) := \mathbb{E}_{\mathbf{x} \sim p_{\boldsymbol{\theta}_i}}[\Phi_q(\mathbf{x})]$ may sound appealing. For example, with the compressive GMM (14) or compressive clustering (16), this would mean minimizing $\|\tilde{\mathbf{z}}_q - \sum_{\ell=1}^k \alpha_{\ell} \mathcal{A}_q(p_{\boldsymbol{\theta}_{\ell}})\|^2$. However, the optimization problem (27) would become more challenging, as suggested by comparing the left two panels to the right two panels in Figure 7, which plots the centroid selection criterion $\langle \tilde{\mathbf{z}}, \Phi(\mathbf{c}) \rangle = \langle \tilde{\mathbf{z}}, \mathcal{A}(p_{\mathbf{c}}) \rangle$ used by greedy algorithms. Also, this approach would not inherit the theoretical guarantees that were carefully established using the LRIP (29), which does not easily translate to the quantized case.

Instead, we suggest to use $\mathcal{A}_{\text{RF}}(p_{\boldsymbol{\theta}_i}) := \mathbb{E}_{\mathbf{x} \sim p_{\boldsymbol{\theta}_i}}[\Phi_{\text{RF}}(\mathbf{x})]$ for the $\mathcal{A}(p_{\boldsymbol{\theta}_i})$ term in (14) and to rescale $\tilde{\mathbf{z}}'_q = \tilde{\mathbf{z}}_q/c$ with the con-

stant c from (41). Indeed, thanks to (41), the resulting cost function $C(\boldsymbol{\theta} | \tilde{\mathbf{z}}'_q) = \|\tilde{\mathbf{z}}'_q - \sum_{\ell=1}^k \alpha_{\ell} \mathcal{A}_{\text{RF}}(p_{\boldsymbol{\theta}_{\ell}})\|^2$ is, in expectation over $\boldsymbol{\xi}$, exactly the same (up a constant additive bias term) as

the nonquantized cost function $C(\boldsymbol{\theta} | \tilde{\mathbf{z}}_{\text{RF}})$ [35]. The similarity between $C(\boldsymbol{\theta} | \tilde{\mathbf{z}}'_q)$ and $C(\boldsymbol{\theta} | \tilde{\mathbf{z}}_{\text{RF}})$, even for a single realization of the m -dimensional dither vector $\boldsymbol{\xi}$, is suggested by comparing the right two panels in Figure 7.

Empirical results [35] suggest that, when the sketch dimension is inflated by about 25%, this quantized compressive learning procedure yields the same performance as

the nonquantized procedure. Moreover, accurate probabilistic bounds for approximation (41), established in [36], allow one to extend the theoretical compressive learning guarantees in (30) to this new cost function.

Privacy preservation

In addition to its efficient use of computational resources, sketching is a promising tool for privacy-preserving machine learning. In numerous applications, such as when working with medical records, online surveys, or measurements coming from personal devices, data samples contain sensitive personal information, and data providers ask that individuals' contributions to the data set remain private, i.e., not publicly discoverable. Learning from such data collections while protecting the privacy of individual contributors has become a crucial challenge [37], [38], [39].

A common way to preserve privacy is to have a trusted data holder (or “curator”) corrupt the response to each query of the data set [37] in a controlled manner. A query may ask for something as simple as counting the number of times a given event occurred, or it may ask for more sophisticated information that requires the data holder to run an inference algorithm. As the corruption becomes more significant, the privacy guarantee gets stronger, but the quality of the response to the query (called the *utility*) degrades. This can be conceptualized by a privacy–utility tradeoff [37], [39].

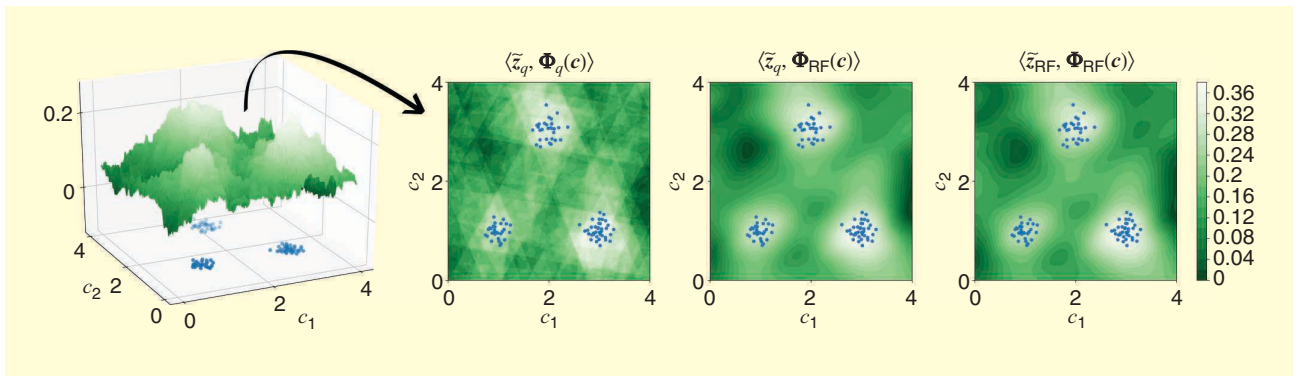


FIGURE 7. The criterion $\langle \tilde{\mathbf{z}}, \Phi(\mathbf{c}) \rangle$ versus the centroid hypothesis $\mathbf{c} = [c_1, c_2]^T$ used by greedy parameter estimation algorithms in compressive clustering with $k = 3$ and a 2D data set (in blue), with and without quantization of the sketch $\tilde{\mathbf{z}}$ and/or feature map Φ . The surface plot on the left highlights the irregularity of $\mathbf{c} \mapsto \langle \tilde{\mathbf{z}}_q, \Phi_q(\mathbf{c}) \rangle$.

When sketching a data set via (1), there is a very simple way to preserve privacy in any subsequent learning task: simply add i.i.d. noise (with appropriate distribution) to the sketch. The privacy level can be adjusted by changing the variance of that noise, as described in the following. This one-time approach to privacy preservation (more generally known as *privacy-preserving data publishing* [37]) has several benefits over the query-based approach to privacy preservation discussed in the previous paragraph (see also the sidebar “Potential of Sketching for Privacy-Aware Learning and Alternative Approaches”).

Sketching with differential privacy guarantees

Differential privacy [38] is a standard framework for privacy preservation that has a precise mathematical definition and is well known in machine learning and signal processing (e.g., [39]). When a given (randomized) learning pipeline is differentially private, its output depends negligibly on the presence or absence of any individual sample in the data set. Differential privacy is robust to many forms of attack, such as when the adversary can access side information that nullifies privacy guarantees based on anonymization or mutual information measures (e.g., when the adversary can control some of the data vectors $\mathbf{x}_i$ or can access additional databases that are correlated with the primary database).

For compressive learning methods, enforcing differential privacy guarantees is as simple as adding well-calibrated noise $\mathbf{v}$ to the usual sketch $\tilde{\mathbf{z}}$, i.e., constructing

$$\tilde{\mathbf{s}}(\mathcal{X}) := \tilde{\mathbf{z}}(\mathcal{X}) + \mathbf{v}, \quad (42)$$

where we find it helpful to explicitly denote the dependence of the data set $\mathcal{X}$. We will assume that the realization $\Phi(\cdot)$ of the random feature map is fixed and publicly known, in contrast to other approaches, like [40] and [41], that use linear mixing matrices as encryption keys to ensure privacy preservation. As a result, when we treat $\tilde{\mathbf{s}}(\mathcal{X})$ as random, this is due to the randomness in $\mathbf{v}$, not the randomness in $\Phi(\cdot)$ or $\mathcal{X}$.

Formally, the sketching mechanism $\tilde{\mathbf{s}}(\cdot)$ is said to be ϵ -differentially private [38] if, for any data set $\mathcal{X} = \{\mathbf{x}_i\}_{i=1}^n$ and “neighboring” data set $\mathcal{X}' = \mathcal{X} \setminus \{\mathbf{x}_j\} \cup \{\mathbf{x}'_j\}$ that replaces the individual sample $\mathbf{x}_j$ by another sample $\mathbf{x}'_j$, and for any possible sketch outcome $\mathbf{s}$, we have that

$$\exp(-\epsilon) \leq \frac{p_{\tilde{\mathbf{s}}(\mathcal{X})}(\mathbf{s})}{p_{\tilde{\mathbf{s}}(\mathcal{X}')}(\mathbf{s})} \leq \exp(\epsilon). \quad (43)$$

Here, $\epsilon > 0$ plays the role of a privacy level: a smaller ϵ implies a stronger privacy guarantee. In words, (43) says that, when ϵ is small, the densities of $\tilde{\mathbf{s}}(\mathcal{X})$ and $\tilde{\mathbf{s}}(\mathcal{X}')$ are almost indistinguishable, as depicted in Figure 8.

Potential of Sketching for Privacy-Aware Learning and Alternative Approaches

Given a data set, privacy preservation can be achieved by asking a trusted data holder to corrupt all queries of the data set [37], [38] in a controlled manner. There are, however, challenges to this so-called interactive approach. For example, because the privacy-preserving effects of this corruption can often be diminished through the mining of multiple query responses (especially if the queries are adaptive), the per-query corruption levels must be designed with the type and total number of queries in mind. These corruption levels are often designed using a so-called privacy budget, which is expended over multiple queries to meet an overall privacy level. Once the entire privacy budget has been used up, the data can no longer be accessed by a given data user. Also, the data holder must ensure that responses to different data users cannot be combined in a way that circumvents the intended privacy preservation.

In contrast, the noninteractive approach [37], [38] is to publish an intermediate privacy-preserving synopsis of the data set, to which the public is allowed unlimited access. For example, with a low-dimensional data set, one could publish a privacy-preserving histogram of the data [S8],

from which aggregate statistics could be subsequently extracted. The noninteractive approach is attractive for several reasons. For example, there is no need to formulate or allocate a privacy budget; it is sufficient to set an overall privacy level. Also, there is no need to worry about data users sharing/combining data.

By adding noise to a sketch of the form (1), one can easily generate a privacy-preserving synopsis of a data set. By construction, such sketches capture the global statistics of the data set $\mathcal{X} = \{\mathbf{x}_i\}_{i=1}^n$ while being relatively insensitive to each individual data sample $\mathbf{x}_i$, especially when the sample cardinality n is large. Also, when the original data are distributed across multiple devices, a privacy-preserving global sketch can be constructed by first locally sketching at each device and then averaging those local sketches at a fusion center, as illustrated in the sidebar “Compressive Mixture Modeling for Speaker Verification.” In this scenario, the local sketches will themselves be privacy preserving, which alleviates concerns about privacy leaks during data fusion.

Reference

[S8] W. Gardaj, W. Yang, and N. Li, “Differentially private grids for geospatial data,” in *Proc. IEEE 29th Int. Conf. Data Eng. (ICDE)*, 2013, pp. 757–768.

The condition (43) can be interpreted as bounding a “likelihood ratio” [42], a familiar quantity in signal processing. Consider two hypotheses: one that the data set equals $\mathcal{X}$ (i.e., it includes x_j) and the other that the data set equals $\mathcal{X}'$ (i.e., it includes x'_j instead). Then, $p_{\tilde{s}(\mathcal{X})}(s)$ would be the likelihood of observing private sketch s under the first hypothesis, while $p_{\tilde{s}(\mathcal{X}')} (s)$ would be the same for the second hypothesis. Say an adversary wanted to detect whether or not the data set contains x_j . By appropriately thresholding the likelihood ratio $p_{\tilde{s}(\mathcal{X})}(s)/p_{\tilde{s}(\mathcal{X}')} (s)$, one can obtain hypothesis tests that are optimal from various perspectives (e.g., Bayes, minimax, and Neyman–Pearson) [42, Ch. 2]. Thus, when (43) holds with a small ϵ , it is fundamentally difficult for an adversary to determine whether x_j or x'_j was present in the sketch. Even if the adversary had nontrivial prior knowledge of the true hypothesis (as in so-called linkage attacks, which make use of a second public data set to which the target user contributed), (43) implies that—for any method—the probability of recovering the true hypothesis from the sketch is only slightly higher than that which is achievable without observing the sketch.

To ensure that the noisy sketch (42) is differentially private, it is sufficient to draw the noise $\mathbf{v}$ as i.i.d. Laplacian with appropriate variance. The variance needed to achieve a given privacy level ϵ can be determined by analyzing the so-called sensitivity of the noiseless sketch, i.e., the biggest possible change that can result from removing one sample. When using the RF feature map (3), which generates a complex-valued $\tilde{z}(\mathcal{X})$, it has been established [43], [44] that it is sufficient for the real and imaginary components of $\mathbf{v}$ to be i.i.d. Laplacian with a standard deviation of $\sigma_v \propto \frac{m}{n\epsilon}$.

A weaker form of privacy, known as *approximate differential privacy* or (ϵ, δ) -differential privacy [38], can be attained by adding Gaussian noise $\mathbf{v}$ with a smaller variance. For example, with RF features, it is sufficient for the real and imaginary components of $\mathbf{v}$ to be i.i.d. Gaussian with a standard deviation of $\sigma_v \propto \frac{\sqrt{m}}{n\epsilon}$.

The privacy–utility tradeoff facilitates the comparison of different privacy-preserving learning strategies. For example, given two strategies, one could match the privacy levels ϵ and compare utilities, or one could match utilities and compare privacy levels. For compressive learning, the

utility of interest is the risk [recall (21)].

In this section, we focused on differential privacy. Other definitions of privacy exist in the literature, such as information-theoretic ones [49]. Likewise, cryptographic methods, such as fully homomorphic encryption [50], can be used to transmit and manipulate data in a secure and private manner, although this notion of “privacy” is quite different from the former ones. Additional work is needed to understand whether compressive learning is “private” according to definitions other than differential privacy.

Perspectives and open challenges

By averaging well-chosen randomized feature transformations over large training collections, sketching significantly compresses data in a way that facilitates provably accurate yet scalable learning from huge and/or streaming data sets, while simultaneously preserving privacy.

In this article, we described several approaches to accelerate the sketching process, including feature quantization and the use

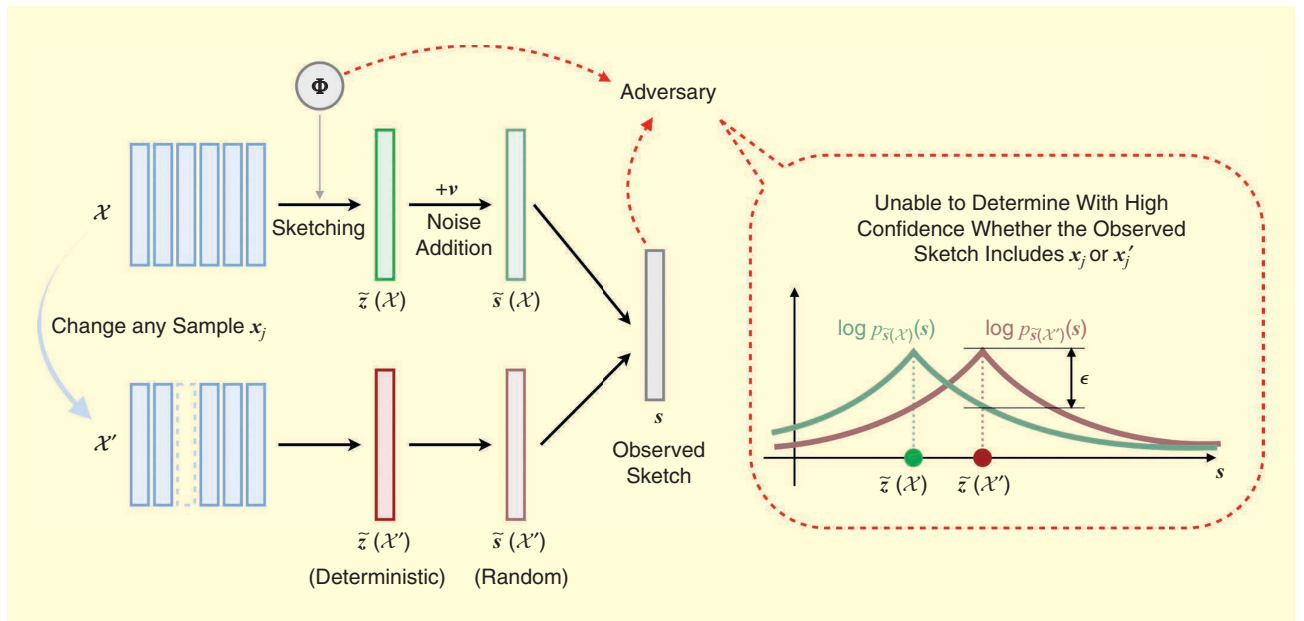


FIGURE 8. When sketching with differential privacy, the output log density of the sketch s remains close when changing one sample in the data set [since (49) is equivalent to $|\log(p_{\tilde{s}(\mathcal{X})}(s)) - \log(p_{\tilde{s}(\mathcal{X}')} (s))| \leq \epsilon$ for all possible s]. An adversary with knowledge of $\Phi(\cdot)$ and s , as symbolized by the red arrows, could then hardly decide whether a given sample x_j was used to compute the sketch or not.

of randomized fast transforms. Another approach is to randomly mask each feature vector $\Phi(\mathbf{x}_i)$ prior to averaging, i.e., setting a random subset of its components to zero. It has been established [44] that such random masking does not increase or decrease the differential privacy level ϵ . But it does reduce the need to compute all entries of each feature vector, and thus it reduces sketching complexity. Another promising approach consists of mixed analog–digital sketches, where, e.g., optical processing units are used to significantly improve the energy efficiency of the linear stage [45].

When discussing methods to learn from a sketch, we focused on optimization-based approaches. Although heuristics based on the OMP [6] and approximate message passing [10] have been proposed that yield promising empirical results, performance guarantees for these approaches have yet to be established. Alternatives, such as total variation minimization over the space of signed measures [23], [26], [46], principled greedy methods [47], and gradient flows on systems of particles [48], could be leveraged to make progress on this front, and black-box optimization could be used to learn from sketches computed by optical processing units.

As for applications of compressive learning, most of the current literature, and hence most of our article, has focused on unsupervised learning tasks. Further work is needed to develop compressive learning methods for supervised tasks like regression and classification [30]. For example, one approach to compressive classification was proposed in [49]: for each class $\ell = 1, \dots, k$, one computes a sketch $\tilde{\mathbf{z}}_\ell$ using only the training examples with label ℓ (i.e., we sketch the k conditional distributions of the data). From those sketches, one could estimate the conditional densities of each class (using a mixture model, for example), from which an ML or maximum a posteriori classifier could be derived. Another approach is to perform classification directly in the compressed domain: to an unseen example $\mathbf{x}'$, we would assign the class ℓ that maximizes the correlation $\langle \tilde{\mathbf{z}}_\ell, \Phi(\mathbf{x}') \rangle$. This strategy can be interpreted [49] as compressively evaluating a Parzen window classifier [32]. Further work is also needed on unsupervised matrix factorization tasks like dictionary learning, low-rank matrix completion, and nonnegative matrix factorization [50].

Acknowledgments

The authors are grateful to Luc Giffon and the anonymous reviewers for their comments on the manuscript. The authors are funded, in part, by the Fonds de la Recherche Scientifique (FNRS); FNRS Projet de Recherche grant T.0136.20 (Learn-2Sense); Agence Nationale de la Recherche (ANR), under grant ANR-19-CHIA-0009 (AllegroAssai); and National Science Foundation, under grant CCF-1955587.

Authors

Rémi Gribonval (remi.gribonval@inria.fr) received his Habilitation to Conduct Research in applied mathematics from University of Rennes 1, France. He is a senior researcher

with Inria at École Normale Supérieure de Lyon, Lyon, 69007, France. His research, at the interface of mathematical signal processing, machine learning, and approximation theory, focuses on the versatile notion of sparsity and its applications,

with an emphasis on dimension reduction, inverse problems, and high-dimensional statistics. He founded the series of international workshops Signal Processing With Adaptive/Sparse Representations and received the 2011 Blaise Pascal Award in Applied Mathematics and Scientific Engineering from the Société de Mathématiques Appliquées et Industrielles through the French National Academy of Sciences. He is a Fellow of IEEE and the

European Association of Signal Processing.

Antoine Chatalic (antoine.chatalic@dibris.unige.it) received his Ph.D. degree in signal processing in 2020 from the University of Rennes 1, France. He is a postdoctoral researcher at the Machine Learning Genoa Center, Genoa, 16146, Italy, under the supervision of Lorenzo Rosasco. His research interests include sketching and dimensionality-reduction techniques, kernel methods, compressive sensing, online algorithms, and privacy-aware learning.

Nicolas Keriven (nicolas.keriven@cnsr.fr) received his Ph.D. degree from the Université de Rennes 1 in 2017, under the supervision of Rémi Gribonval. He is a researcher at the Centre National de la Recherche Scientifique, Paris, 75000, France, and the Grenoble Images Parole Signal Automatique lab, Grenoble, 38402, France. He received the 2017 Best Student Paper award at the Signal Processing With Adaptive/Sparse Representations workshop, for “Random Moments for Sketched Mixture Learning.”

Vincent Schellekens (vincent.schellekens@uclouvain.be) received his Ph.D. degree in signal processing in 2021 from the UCLouvain, Louvain-la-Neuve, B1348, Belgium. He is a postdoctoral researcher funded by the Fonds de la Recherche Scientifique, under the supervision of Laurent Jacques, at the Information and Communication Technologies, Electronics, and Applied Mathematics, UCLouvain, Louvain-la-Neuve, B1348, Belgium. His research focuses on compressive sensing, general random embeddings, and privacy-preservation methods for large-scale learning.

Laurent Jacques (laurent.jacques@uclouvain.be) received his Ph.D. degree in mathematics and physics from UCLouvain, Louvain-la-Neuve, Belgium, in 2004. He is a professor and senior research associate of the Fonds de la Recherche Scientifique with the Image and Signal Processing Group, Institute of Information and Communication Technologies, Electronics, and Applied Mathematics, UCLouvain, Louvain-la-Neuve, B1348, Belgium. His research focuses on sparse signal representations, quantized compressive sensing theory, and computational imaging. He has coauthored more than 40 papers in international journals, 80 conference proceedings and presentations, and four book chapters.

Philip Schniter (schniter.1@osu.edu) received his Ph.D. degree in electrical engineering from Cornell University,

Ithaca, New York, USA, in 2000. He is a professor in the Department of Electrical and Computer Engineering, The Ohio State University, Columbus, Ohio, 43210, USA. He received the IEEE Signal Processing Society (SPS) Best Paper Award in 2016 and has served on several IEEE SPS technical committees, including those on signal processing for communications and networking, sensor arrays and multi-channel communication, and computational imaging. He is a Fellow of IEEE.

References

- [1] R. Gribonval, A. Chatalic, N. Keriven, V. Schellekens, L. Jacques, and P. Schniter, "Sketching datasets for large-scale learning (long version)," 2021, arXiv:2008.01839.
- [2] F. Pérez-Cruz and O. Bousquet, "Kernel methods and their potential use in signal processing," *IEEE Signal Process. Mag.*, vol. 21, no. 3, p. 3, 2004.
- [3] B. Schölkopf and A. Smola, "Learning with kernels," in *Adaptive Computation and Machine Learning*. Cambridge: MIT Press, 2002. [Online]. Available: <https://mitpress.mit.edu/books/learning-kernels>
- [4] S. Foucart and H. Rauhut, *A Mathematical Introduction to Compressive Sensing*. BerlSpringer-Verlag, May 2012.
- [5] R. Gribonval, G. Blanchard, N. Keriven, and Y. Traonmilin, "Compressive statistical learning with random feature moments," *Math. Stat. Learn.*, to be published.
- [6] N. Keriven, A. Bourrier, R. Gribonval, and P. Pérez, "Sketching for large-scale learning of mixture models," *Inform. Inference*, vol. 7, no. 3, p. 3, 2017.
- [7] A. Rahimi and B. Recht, "Random features for large scale kernel machines," in *Proc. Neural Inform. Process. Syst. Conf.*, 2007, 1177–1184.
- [8] R. M. Gray and D. L. Neuhoff, "Quantization," *IEEE Trans. Inform. Theory*, vol. 44, no. 6, pp. 2325–2383, 1998. doi: 10.1109/18.720541.
- [9] N. Keriven, N. Tremblay, Y. Traonmilin, and R. Gribonval, "Compressive K-means," in *Proc. IEEE Int. Conf. Acoust. Speech & Signal Process*, 2017, 6369–6373.
- [10] E. Byrne, A. Chatalic, R. Gribonval, and P. Schniter, "Sketched clustering via hybrid approximate message passing," *IEEE Trans. Signal Process*, vol. 67, no. 17, p. 17, 2019. doi: 10.1109/TSP.2019.2924585.
- [11] M. P. Sheehan, M. S. Kotzagiannidis, and M. E. Davies, "Compressive independent component analysis," in *Proc. 27th IEEE European Signal Process. Conf. (EUSIPCO)*, 2019, pp. 1–5. doi: 10.23919/EUSIPCO.2019.8903095.
- [12] M. P. Sheehan, A. Gonon, and M. E. Davies, "Compressive learning for semi-parametric models," 2019, arXiv:1910.10024.
- [13] G. Cormode, M. Garofalakis, P. J. Haas, and C. Jermaine, "Synopses for massive data," *Found. Trends Databases*, vol. 4, nos. 1–3, p. 1, 2011. doi: 10.1561/1900000004.
- [14] D. Achlioptas, "Database-friendly random projections: Johnson–Lindenstrauss with binary coins," *J. Comput. Syst. Sci.*, vol. 66, no. 4, p. 4, 2003. doi: 10.1016/S0022-0000(03)00025-4.
- [15] M. W. Mahoney, "Randomized algorithms for matrices and data," *Found. Trends Mach. Learning*, vol. 3, no. 2, pp. 123–224, 2010.
- [16] C. Boutsidis, A. Zouzias, and P. Drineas, "Random projections for k-means clustering," in *Proc. Adv. Neural Inform. Process. Syst.*, 2010, pp. 1–9.
- [17] D. P. Woodruff, "Sketching as a tool for numerical linear algebra," *Found. Trends Theoretical Comput. Sci.*, vol. 10, nos. 1–2, pp. 1–157, 2014. doi: 10.1561/0400000060.
- [18] D. Achlioptas, F. McSherry, and B. Schölkopf, "Sampling techniques for kernel methods," in *Proc. Advances Neural Inform. Process. Syst.*, 2002, pp. 335–342.
- [19] D. Feldman, M. Monemizadeh, and C. Sohler, "Coresets and sketches for high dimensional subspace approximation problems," in *Proc. Annu. ACM-SIAM SY-P. Discrete Algorithms*, 2010, vol. 1, pp. 630–649. doi: 10.1137/1.9781611973075.53.
- [20] C. Williams and M. W. Seeger, "Using the Nystrom method to speed up kernel machines," in *Proc. Advances Neural Inform. Processing Syst.*, vol. 13, 2001, pp. 628–688.
- [21] A. R. Hall, *Generalized Method of Moments*. London: Oxford Univ. Press, 2005.
- [22] E. J. Candès and M. B. Wakin, "An introduction to compressive sampling," *IEEE Signal Process. Mag.*, vol. 25, no. 2, p. 2, 2008. doi: 10.1109/MSP.2007.914731.
- [23] E. J. Candès and C. Fernandez-Granda, "Towards a mathematical theory of super-resolution," *Commun. Pure Appl. Math.*, vol. 67, p. 6, 2014.
- [24] M. Lustig, D. L. Donoho, J. M. Santos, and J. M. Pauly, "Compressed sensing MRI," *IEEE Signal Process. Mag.*, vol. 25, no. 2, p. 2, 2008. doi: 10.1109/MSP.2007.914728.
- [25] J. A. Fessler, "Optimization methods for magnetic resonance image reconstruction: key models and optimization algorithms," *IEEE Signal Process. Mag.*, vol. 37, no. 1, p. 1, 2020. doi: 10.1109/MSP.2019.2943645.
- [26] Q. Denoyelle, V. Duval, G. Peyr  e, and E. Soubies, "The sliding Frank–Wolfe algorithm and its application to super-resolution microscopy," *Inverse Problems*, vol. 36, no. 1, p. 1, Jan. 2020. doi: 10.1088/1361-6420/ab2a29.
- [27] Y. Chen and Y. Chi, "Harnessing structures in big data via guaranteed low-rank matrix estimation: Recent theory and fast algorithms via convex and nonconvex optimization," *IEEE Signal Process. Mag.*, vol. 35, no. 4, p. 4, 2018. doi: 10.1109/MSP.2018.2821706.
- [28] R. Ge, J. D. Lee, and T. Ma, "Matrix completion has no spurious local minimum," in *Proc. Adv. Neural Inf. Process. Syst.*, New York: Curran Associates, 2016.
- [29] V. Vapnik, *Nature Statistical Learning Theory*, vol. 6, Berlin, Germany: Springer, 2013.
- [30] S. Shalev-Shwartz and S. Ben-David, *Understanding Machine Learning*. Cambridge, MA: Cambridge Univ. Press, 2009.
- [31] W. B. Johnson and J. Lindenstrauss, "Extensions of Lipschitz mappings into a Hilbert space," *Contemporary Math.*, vol. 26, pp. 189–206, 1984.
- [32] Z.-W. Pan, D.-H. Xiang, Q.-W. Xiao, and D.-X. Zhou, "Parzen windows for multi-class classification," *J. Complexity*, vol. 24, nos. 5–6, pp. 5–6, 2008. doi: 10.1016/j.jco.2008.07.001.
- [33] F. X. Yu, A. T. Suresh, K. M. Choromanski, D. N. Holtmann-Rice, and S. Kumar, "Orthogonal random features," in *Proc. Neural Inform. Process. Syst. Conf.*, 2016, pp. 1975–1983.
- [34] A. Chatalic, R. Gribonval, and N. Keriven, "Large-scale high-dimensional clustering with fast sketching," in *Proc. IEEE Int. Conf. Acoust. Speech & Signal Process*, 2018, pp. 4714–4718.
- [35] V. Schellekens and L. Jacques, "Quantized compressive k-means," *IEEE Signal Process. Lett.*, vol. 25, no. 8, pp. 1211–1215, Aug. 2018. doi: 10.1109/LSP.2018.2847908.
- [36] V. Schellekens and L. Jacques, "Breaking the waves: Asymmetric random periodic features for low-bitrate kernel machines," *Inf. Inference*, to be published. doi: <https://doi.org/10.1093/imaia/iaab008>
- [37] B. C. M. Fung, K. Wang, R. Chen, and P. S. Yu, "Privacy-preserving data publishing: A survey of recent developments," *ACM Comput. Survey*, vol. 42, no. 4, p. 4, June 2010. doi: 10.1145/1749603.1749605.
- [38] C. Dwork and A. Roth, "The algorithmic foundations of differential privacy," *Theoretical Comput. Sci.*, vol. 9, no. 3–4, pp. 3–4, 2014.
- [39] A. D. Sarwate and K. Chaudhuri, "Signal processing and machine learning with differential privacy: Algorithms and challenges for continuous data," *IEEE Signal Process. Mag.*, vol. 30, no. 5, p. 5, 2013. doi: 10.1109/MSP.2013.2259911.
- [40] M. Testa, D. Valsesia, T. Bianchi, and E. Magli, *Compressed Sensing for Privacy-Preserving Data Processing*. Berlin: Springer-Verlag, 2019.
- [41] S. Rane and P. T. Boufounos, "Privacy-preserving nearest neighbor methods: Comparing signals without revealing them," *IEEE Signal Process. Mag.*, vol. 30, no. 2, p. 2, 2013. doi: 10.1109/MSP.2012.2230221.
- [42] H. V. Poor, *An Introduction to Signal Detection and Estimation*. Springer Science & Business Media, 2013.
- [43] V. Schellekens, A. Chatalic, F. Houssiau, Y.-A. de Montjoye, L. Jacques, and R. Gribonval, "Differentially private compressive k-means," in *Proc. IEEE Int. Conf. Acoust. Speech & Signal Process*, 2019, pp. 7933–7937.
- [44] A. Chatalic, V. Schellekens, F. Houssiau, Y.-A. de Montjoye, L. Jacques, and R. Gribonval, "Compressive learning with privacy guarantees," *Information and Inference: A Journal of the IMA*, May 2021. doi: 10.1093/imaia/iaab005. [Online]. Available: <https://doi.org/10.1093/imaia/iaab005>
- [45] A. Saade et al., "Random projections through multiple optical scattering," in *Proc. IEEE Int. Conf. Acoust. Speech & Signal Process*, 2016, 6215–6219.
- [46] N. Boyd, G. Schiebinger, and B. Recht, "The alternating descent conditional gradient method for sparse inverse problems," *SIAM J. Optim.*, vol. 27, no. 2, p. 2, 2017. doi: 10.1137/15M1035793.
- [47] Y. Traonmilin and J.-F. Aujol, "The basins of attraction of the global minimizers of the non-convex sparse spikes estimation problem," 2018, arXiv:1811.12000.
- [48] L. Chizat and F. Bach, "On the global convergence of gradient descent for over-parameterized models using optimal transport," in *Proc. Adv. Neural Inform. Process. Syst.*, 2018.
- [49] V. Schellekens and L. Jacques, "Compressive classification (machine learning without learning)," 2018, arXiv:1812.01410.
- [50] M. Udell, C. Horn, R. Zadeh, and S. Boyd, "Generalized low rank models," *Found. Trends Mach. Learn.*, vol. 9, no. 1, p. 1, 2016. doi: 10.1561/22000000055.
- [51] R. Gribonval, G. Blanchard, N. Keriven and Y. Traonmilin, "Statistical learning guarantees for compressive clustering and compressive mixture modeling," *Math. Stat. Learn.*, to be published.