

# ROBUST COMPUTATION OF LINEAR MODELS, OR HOW TO FIND A NEEDLE IN A HAYSTACK

GILAD LERMAN\*, MICHAEL MCCOY†, JOEL A. TROPP‡, AND TENG ZHANG°

**ABSTRACT.** Consider a dataset of vector-valued observations that consists of a modest number of noisy inliers, which are explained well by a low-dimensional subspace, along with a large number of outliers, which have no linear structure. This work describes a convex optimization problem, called REAPER, that can reliably fit a low-dimensional model to this type of data. The paper provides an efficient algorithm for solving the REAPER problem, and it documents numerical experiments which confirm that REAPER can dependably find linear structure in synthetic and natural data. In addition, when the inliers are contained in a low-dimensional subspace, there is a rigorous theory that describes when REAPER can recover the subspace exactly.

## 1. INTRODUCTION

Low-dimensional linear models have applications in a huge array of data analysis problems. Let us highlight some examples from computer vision, machine learning, and bioinformatics.

**Illumination models:** Images of a face—or any Lambertian object—viewed under different illumination conditions are contained in a nine-dimensional subspace [EHY95, HYL<sup>+</sup>03, BJ03].

**Structure from motion:** Feature points on a moving rigid body lie on an affine space of dimension three, assuming the affine camera model [CK98]. More generally, estimating structure from motion involves estimating low-rank matrices [EvdH10, Sec. 5.2].

**Latent semantic indexing:** We can describe a large corpus of documents that concern a small number of topics using a low-dimensional linear model [DDL<sup>+</sup>88].

**Population stratification:** Low-dimensional models of single nucleotide polymorphism (SNP) data have been used to show that the genotype of an individual is correlated with her geographical ancestry [NJB<sup>+</sup>08]. More generally, linear models are used to assess differences in allele frequencies among populations [PPP<sup>+</sup>06].

In most of these applications, the datasets are noisy, and they contain a substantial number of outliers. Principal component analysis, the standard method for finding a low-dimensional linear model, is sensitive to these non-idealities. As a consequence, good robust modeling techniques would be welcome in a range of scientific and engineering disciplines.

The purpose of this work is to introduce a new method for fitting low-dimensional linear models. The approach is robust against noise in the inliers, and it can cope with a huge number of outliers. Our formulation involves a convex optimization problem, and we describe an efficient numerical algorithm that is guaranteed to produce an approximate solution after a modest number of spectral calculations. We include experiments with synthetic and natural data to verify that our technique reliably seeks out linear structure. Furthermore, under the ideal assumption that the inliers are contained in a subspace, we develop sufficient conditions for the method to identify the subspace without error.

---

*Date:* 17 February 2012.

\*Department of Mathematics, University of Minnesota, Twin Cities.

†Department of Computing and Mathematical Sciences, California Institute of Technology.

°Institute for Mathematics and its Applications, University of Minnesota, Twin Cities.

| Report Documentation Page                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                              |                                    |                                     |                                                           | Form Approved<br>OMB No. 0704-0188                  |                                 |
|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------|-------------------------------------|-----------------------------------------------------------|-----------------------------------------------------|---------------------------------|
| Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number. |                                    |                                     |                                                           |                                                     |                                 |
| 1. REPORT DATE<br><b>17 FEB 2012</b>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                   |                                    | 2. REPORT TYPE                      |                                                           | 3. DATES COVERED<br><b>00-00-2012 to 00-00-2012</b> |                                 |
| 4. TITLE AND SUBTITLE<br><b>Robust Computation of Linear Models, or How to Find a Needle in a Haystack</b>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                             |                                    |                                     |                                                           | 5a. CONTRACT NUMBER                                 |                                 |
|                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        |                                    |                                     |                                                           | 5b. GRANT NUMBER                                    |                                 |
|                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        |                                    |                                     |                                                           | 5c. PROGRAM ELEMENT NUMBER                          |                                 |
| 6. AUTHOR(S)                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                           |                                    |                                     |                                                           | 5d. PROJECT NUMBER                                  |                                 |
|                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        |                                    |                                     |                                                           | 5e. TASK NUMBER                                     |                                 |
|                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        |                                    |                                     |                                                           | 5f. WORK UNIT NUMBER                                |                                 |
| 7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES)<br><b>California Institute of Technology, Department of Computing and Mathematical Sciences, Pasadena, CA, 91125</b>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                |                                    |                                     |                                                           | 8. PERFORMING ORGANIZATION REPORT NUMBER            |                                 |
| 9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                |                                    |                                     |                                                           | 10. SPONSOR/MONITOR'S ACRONYM(S)                    |                                 |
|                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        |                                    |                                     |                                                           | 11. SPONSOR/MONITOR'S REPORT NUMBER(S)              |                                 |
| 12. DISTRIBUTION/AVAILABILITY STATEMENT<br><b>Approved for public release; distribution unlimited</b>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                  |                                    |                                     |                                                           |                                                     |                                 |
| 13. SUPPLEMENTARY NOTES                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                |                                    |                                     |                                                           |                                                     |                                 |
| 14. ABSTRACT<br><b>Consider a dataset of vector-valued observations that consists of a modest number of noisy inliers, which are explained well by a low-dimensional subspace, along with a large number of outliers, which have no linear structure. This work describes a convex optimization problem, called reaper, that can reliably fit a low-dimensional model to this type of data. The paper provides an efficient algorithm for solving the reaper problem, and it documents numerical experiments which confirm that reaper can dependably find linear structure in synthetic and natural data. In addition, when the inliers are contained in a low-dimensional subspace, there is a rigorous theory that describes when reaper can recover the subspace exactly.</b>                                                                      |                                    |                                     |                                                           |                                                     |                                 |
| 15. SUBJECT TERMS                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                      |                                    |                                     |                                                           |                                                     |                                 |
| 16. SECURITY CLASSIFICATION OF:                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        |                                    |                                     | 17. LIMITATION OF ABSTRACT<br><b>Same as Report (SAR)</b> | 18. NUMBER OF PAGES<br><b>45</b>                    | 19a. NAME OF RESPONSIBLE PERSON |
| a. REPORT<br><b>unclassified</b>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                       | b. ABSTRACT<br><b>unclassified</b> | c. THIS PAGE<br><b>unclassified</b> |                                                           |                                                     |                                 |

**1.1. Notation and Preliminaries.** We write  $\|\cdot\|$  for the  $\ell_2$  norm on vectors and the spectral norm on matrices;  $\|\cdot\|_F$  represents the Frobenius norm. Angle brackets  $\langle \cdot, \cdot \rangle$  denote the standard inner product on vectors and matrices, and  $\text{tr}$  refers to the trace. The orthogonal complement of a subspace  $L$  is expressed as  $L^\perp$ . The curly inequality  $\preceq$  denotes the semidefinite order: For symmetric matrices  $\mathbf{A}$  and  $\mathbf{B}$ , we write  $\mathbf{A} \preceq \mathbf{B}$  if and only if  $\mathbf{B} - \mathbf{A}$  is positive semidefinite.

An *orthoprojector* is a symmetric matrix  $\mathbf{\Pi}$  that satisfies  $\mathbf{\Pi}^2 = \mathbf{\Pi}$ . Each subspace  $L$  in  $\mathbb{R}^D$  is the range of a unique  $D \times D$  orthoprojector  $\mathbf{\Pi}_L$ . The trace of an orthoprojector equals the dimension of its range:  $\text{tr}(\mathbf{\Pi}_L) = \dim(L)$ . For each point  $\mathbf{x} \in \mathbb{R}^D$ , the image  $\mathbf{\Pi}_L \mathbf{x}$  is the best  $\ell_2$  approximation of  $\mathbf{x}$  in the subspace  $L$ .

Finally, we introduce the *spherization transform* for vectors:

$$\tilde{\mathbf{x}} := \begin{cases} \mathbf{x} / \|\mathbf{x}\|, & \mathbf{x} \neq \mathbf{0} \\ \mathbf{0}, & \text{otherwise.} \end{cases} \quad (1.1)$$

We extend the spherization transform to matrices by applying it separately to each column.

**1.2. Linear Modeling by Principal Component Analysis.** Let  $\mathcal{X}$  be a set of  $N$  distinct observations in  $\mathbb{R}^D$ . Suppose we wish to determine a  $d$ -dimensional subspace that best explains the data. For each point, we can measure the residual error in the approximation by computing the orthogonal distance from the point to the subspace. The classical method for fitting a subspace asks us to minimize the sum of the *squared* residuals:

$$\text{minimize} \quad \sum_{\mathbf{x} \in \mathcal{X}} \|\mathbf{x} - \mathbf{\Pi} \mathbf{x}\|^2 \quad \text{subject to} \quad \mathbf{\Pi} \text{ is an orthoprojector} \quad \text{and} \quad \text{tr} \mathbf{\Pi} = d. \quad (1.2)$$

This approach is equivalent with the method of *principal component analysis* (PCA) from the statistics literature [Jol02] and the *total least squares* (TLS) method from the linear algebra community [vHV87].

The mathematical program (1.2) is not convex because orthoprojectors do not form a convex set, so we have no right to expect that the problem is tractable. Nevertheless, we can compute an analytic solution using a singular value decomposition (SVD) of the data [EY39, vHV87]. Suppose that  $\mathbf{X}$  is a  $D \times N$  matrix whose columns are the data points, arranged in fixed order, and let  $\mathbf{X} = \mathbf{U} \mathbf{\Sigma} \mathbf{V}^t$  be an SVD of this matrix. Define  $\mathbf{U}_d$  by extracting the first  $d$  columns of  $\mathbf{U}$ ; the columns of  $\mathbf{U}_d$  are often called the *principal components* of the data. Then we can construct an optimal point  $\mathbf{\Pi}_\star$  for (1.2) using the formula  $\mathbf{\Pi}_\star = \mathbf{U}_d \mathbf{U}_d^t$ .

**1.3. Classical Methods for Achieving Robustness.** Imagine now that the dataset  $\mathcal{X}$  contains *inliers*, points we hope to explain with a linear model, as well as *outliers*, points that come from another process, such as a different population or noise. The data are not labeled, so it may be challenging to distinguish inliers from outliers. If we apply the PCA formulation (1.2) to fit a subspace to  $\mathcal{X}$ , the rogue points can interfere with the linear model for the inliers.

To guard the subspace estimation procedure against outliers, statisticians have proposed to replace the sum of squares in (1.2) with a figure of merit that is less sensitive to outliers. One possibility is to sum the *unsquared* residuals, which reduces the contribution from large residuals that may result from aberrant data points. This idea leads to the following optimization problem.

$$\text{minimize} \quad \sum_{\mathbf{x} \in \mathcal{X}} \|\mathbf{x} - \mathbf{\Pi} \mathbf{x}\| \quad \text{subject to} \quad \mathbf{\Pi} \text{ is an orthoprojector} \quad \text{and} \quad \text{tr} \mathbf{\Pi} = d. \quad (1.3)$$

In case  $d = D - 1$ , the problem (1.3) is sometimes called *orthogonal  $\ell_1$  regression* [SW87] or *least orthogonal absolute deviations* [Nyq88]. The extension to general  $d$  is apparently more recent [Wat01, DZHZ06]. See the books [HR09, RL87, MMY06] for an extensive discussion of other ways to combine residuals to obtain robust estimators.

Unfortunately, the mathematical program (1.3) is not convex, and, in contrast to (1.2), no *deus ex machina* emerges to make the problem tractable. Although there are many algorithms [Nyq88,

**Algorithm 1.1** *Prototype algorithm for robust computation of a linear model*

INPUT:

- A set  $\mathcal{X}$  of observations in  $\mathbb{R}^D$
- The target dimension  $d$  for the linear model, where  $0 < d < D$

OUTPUT:

- A  $\lceil d \rceil$ -dimensional subspace of  $\mathbb{R}^D$

PROCEDURE:

- 1 [Opt.] Solve (1.5) to obtain a center  $\mathbf{c}_\star$ , and center the data:  $\mathbf{x} \leftarrow \mathbf{x} - \mathbf{c}_\star$  for each  $\mathbf{x} \in \mathcal{X}$
- 2 [Opt.] Spherize the data:  $\mathbf{x} \leftarrow \mathbf{x} / \|\mathbf{x}\|$  for each nonzero  $\mathbf{x} \in \mathcal{X}$
- 3 Solve (1.4) with dataset  $\mathcal{X}$  and parameter  $d$  to obtain an optimal point  $\mathbf{P}_\star$
- 4 Output a dominant  $\lceil d \rceil$ -dimensional invariant subspace of  $\mathbf{P}_\star$

CM91, BH93, Wat02, DZHZ06, ZSL09] that attempt (1.3), none is guaranteed to return a global minimum. In fact, most of the classical proposals for robust linear modeling involve intractable optimization problems, which makes them poor options for computation in spite of their theoretical properties [MMY06].

**1.4. A Convex Program for Robust Linear Modeling.** The goal of this paper is to develop, analyze, and test a rigorous method for fitting robust linear models by means of convex optimization. We propose to *relax* the hard optimization problem (1.3) by replacing the nonconvex constraint set with a larger convex set. The advantage of this approach is that we can solve the resulting convex program completely using a variety of efficient algorithms.

The idea behind our relaxation is straightforward. Each eigenvalue of an orthoprojector  $\mathbf{\Pi}$  equals zero or one because  $\mathbf{\Pi}^2 = \mathbf{\Pi}$ . Although a 0–1 constraint on eigenvalues is hard to enforce, the symmetric matrices whose eigenvalues lie in the interval  $[0, 1]$  form a convex set. This observation leads us to frame the following convex optimization problem. Given a dataset  $\mathcal{X}$  in  $\mathbb{R}^D$  and a target dimension  $d \in (0, D)$  for the linear model, we solve

$$\text{minimize} \quad \sum_{\mathbf{x} \in \mathcal{X}} \|\mathbf{x} - \mathbf{P}\mathbf{x}\| \quad \text{subject to} \quad \mathbf{0} \preceq \mathbf{P} \preceq \mathbf{I} \quad \text{and} \quad \text{tr } \mathbf{P} = d. \quad (1.4)$$

We refer to (1.4) as REAPER because it harvests linear structure from data. The formulation (1.4) refines a proposal from the paper [ZL11], which describes a looser relaxation of (1.3).

In many cases, the optimal set of (1.4) consists of a single orthoprojector whose range provides an excellent fit for the data. The bulk of this paper provides theoretical and empirical support for this observation. In Section 4, we present a numerical algorithm for solving (1.4).

**1.4.1. Some Practical Matters.** Although the REAPER formulation is effective on its own, we can usually obtain better linear models if we preprocess the data before solving (1.4). Let us summarize the recommended procedure, which appears as Algorithm 1.1.

First, the REAPER problem assumes that the inliers are approximately centered. When they are not, it is important to identify a centering point  $\mathbf{c}_\star$  for the dataset and to work with the centered observations. We can compute a centering point  $\mathbf{c}_\star$  robustly by solving the Euclidean median problem [HR09, MMY06]:

$$\text{minimize} \quad \sum_{\mathbf{x} \in \mathcal{X}} \|\mathbf{x} - \mathbf{c}\|. \quad (1.5)$$

It is also possible to incorporate centering by modifying the optimization problem (1.4); see Section 6.1.5 for more information.

Second, the REAPER formulation can be sensitive to outliers with large magnitude. A simple but powerful method for addressing this challenge is to spherize the data points before solving the

optimization problem. For future reference, we write down the resulting convex program.

$$\text{minimize} \quad \sum_{\tilde{\mathbf{x}} \in \mathcal{X}} \|\tilde{\mathbf{x}} - \mathbf{P}\tilde{\mathbf{x}}\| \quad \text{subject to} \quad \mathbf{0} \preceq \mathbf{P} \preceq \mathbf{I} \quad \text{and} \quad \text{tr} \mathbf{P} = d. \quad (1.6)$$

The tilde denotes the spherization transform (1.1). We refer to (1.6) as the S-REAPER problem. In most (but not all) of our experimental work, we have found that S-REAPER outperforms REAPER. The idea of spherizing data before fitting a subspace was proposed in the paper [LMS<sup>+</sup>99], where it is called *spherical PCA*. Spherical PCA is regarded as one of the most effective current techniques for robust linear modeling [MMY06].

Finally, we regard the parameter  $d$  in REAPER and S-REAPER as a proxy for the dimension of the linear model. While the rank of an optimal solution  $\mathbf{P}_\star$  to (1.4) or (1.6) cannot be smaller than  $d$  because of the constraints  $\text{tr} \mathbf{P} = d$  and  $\mathbf{P} \preceq \mathbf{I}$ , the rank of  $\mathbf{P}_\star$  often exceeds  $d$ . Our numerical experience indicates that the dominant eigenvectors of  $\mathbf{P}_\star$  do not change much as we adjust  $d$ . Therefore, if a model with exact dimension  $d$  is required, we recommend using a dominant  $d$ -dimensional invariant subspace of  $\mathbf{P}_\star$ . See Section 6.1.6 for alternative approaches.

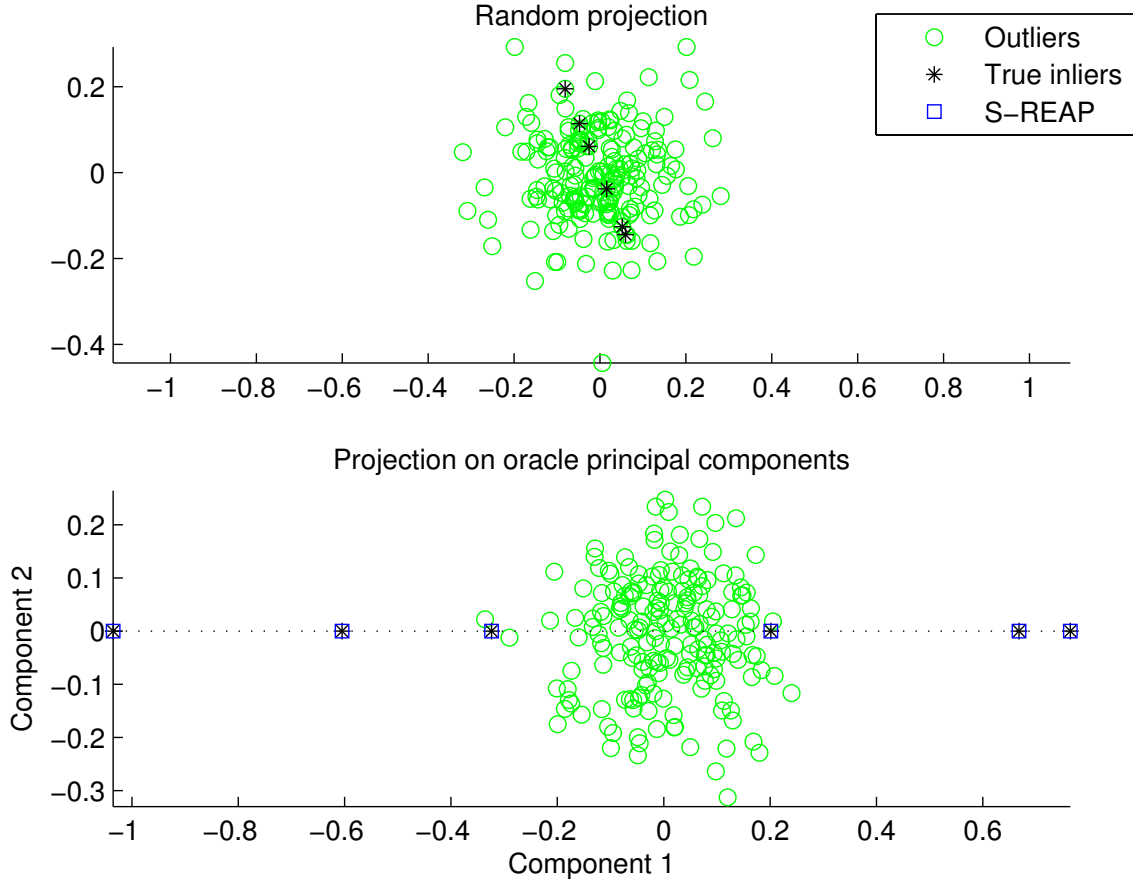


FIGURE 1.1. *A needle in a haystack*. The needle dataset is described in Section 1.5. [Top] The projection onto a random two-dimensional subspace. [Bottom] The projection onto the oracle two-dimensional subspace obtained by applying PCA to the inliers only. The blue squares mark the projection of the inliers onto the one-dimensional model computed with S-REAPER (1.6).

**1.5. Finding a Needle in a Haystack.** The REAPER and S-REAPER formulations work astonishingly well. They can reliably identify linear structure in situations where the task might seem impossible. As a first illustration, we examine the problem of “finding a needle in a haystack.” In Section 5, we present a more comprehensive suite of experiments.

Let the ambient dimension  $D = 100$ . Consider a dataset  $\mathcal{X}$  in  $\mathbb{R}^D$  that consists of 6 points arrayed arbitrarily on a line through the origin (the *needle*) and 200 points drawn from a centered normal distribution with covariance  $D^{-1}\mathbf{I}$  (the *haystack*). The 6 inliers have strong linear structure, but they comprise only 3% of the data. The 200 outliers, making up the bulk of the sample, typically have no linear structure at all. We solve the optimization problem S-REAPER with the parameter  $d = 1$  to obtain an optimal point  $\mathbf{P}_*$ , which happens to be a rank-one orthoprojector.

Figure 1.1 displays the outcome of this experiment. We see that the range of the S-REAPER model  $\mathbf{P}_*$  identifies the direction of the needle *with negligible error*! Of course, this feat is not magic. Even though the needle is impossible to see in most two-dimensional projections of the data, it becomes visible once we identify the direction of the needle. We find it remarkable that S-REAPER renders the search for this direction into a convex problem.

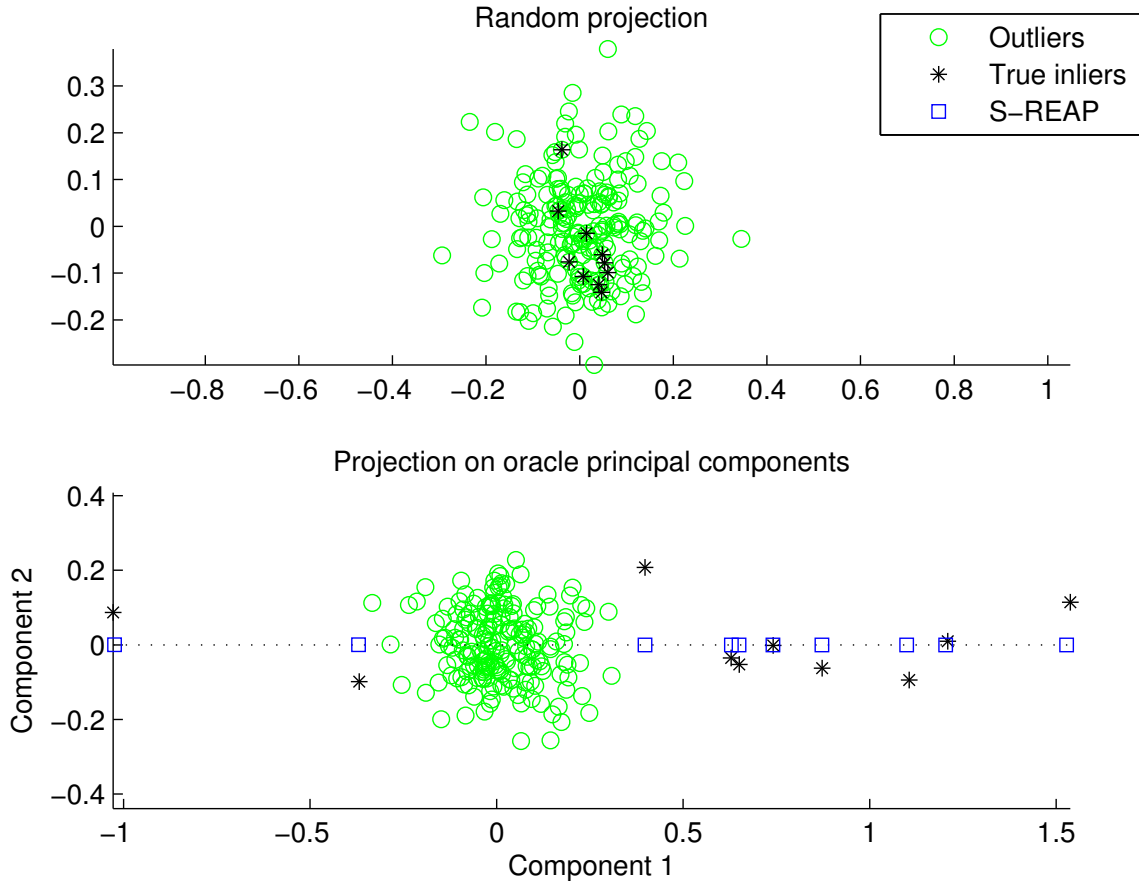


FIGURE 1.2. *A syringe in a haystack.* The syringe dataset is described in Section 1.5. [Top] The projection onto a random two-dimensional subspace. [Bottom] The projection onto the oracle two-dimensional model determined by applying PCA to the inliers only. The blue squares mark the projection of the inliers onto the one-dimensional model computed with S-REAPER (1.6).

Our method for fitting linear models is robust to noise in the inliers, although it may require more data to fit the model accurately. To illustrate, we consider a noisy version of the same problem. We retain the ambient dimension  $D = 100$ . Each of 10 inliers takes the form  $\mathbf{x}_i = g_i \mathbf{v} + \mathbf{z}_i$  where  $\mathbf{v}$  is a fixed vector,  $g_i$  are i.i.d. standard normal variables, and the noise vectors  $\mathbf{z}_i$  are i.i.d.  $\text{NORMAL}(\mathbf{0}, (16D)^{-1} \mathbf{I})$ . The inliers form a tube around the vector  $\mathbf{v}$ , so they look like a syringe. As before, the outliers consist of 200 points that are i.i.d.  $\text{NORMAL}(\mathbf{0}, D^{-1} \mathbf{I})$ . We apply S-REAPER with  $d = 1$  to the full dataset to obtain a solution  $\mathbf{P}_\star$ , and we use a dominant eigenvector  $\mathbf{u}_\star$  of  $\mathbf{P}_\star$  as our one-dimensional model. To find the optimal two-dimensional model for the inliers, we solve (1.2) with  $d = 2$  using the *inliers only* to obtain two oracle principal components. Figure 1.2 illustrates the outcome of this experiment. The angle between the first oracle principal component and the S-REAPER model  $\mathbf{u}_\star$  is approximately  $3.60^\circ$ .

**1.6. The Haystack Model.** Inspired by the success of these experiments, let us attempt to place them on a firm theoretical foundation. To that end, we propose a very simple generative random model. We want to capture the idea that the inliers admit a low-dimensional linear model, while the outliers are totally unstructured.

TABLE 1.1. *The Haystack Model.* A generative random model for data with linear structure that is contaminated with outliers.

|                            |                                                                                                                    |
|----------------------------|--------------------------------------------------------------------------------------------------------------------|
| $D$                        | Dimension of the ambient space                                                                                     |
| $L$                        | A proper $d$ -dimensional subspace of $\mathbb{R}^D$ containing the inliers                                        |
| $N_{\text{in}}$            | Number of inliers                                                                                                  |
| $N_{\text{out}}$           | Number of outliers                                                                                                 |
| $\rho_{\text{in}}$         | Inlier sampling ratio $\rho_{\text{in}} := N_{\text{in}}/d$                                                        |
| $\rho_{\text{out}}$        | Outlier sampling ratio $\rho_{\text{out}} := N_{\text{out}}/D$                                                     |
| $\sigma_{\text{in}}^2$     | Variance of the inliers                                                                                            |
| $\sigma_{\text{out}}^2$    | Variance of the outliers                                                                                           |
| $\mathcal{X}_{\text{in}}$  | Set of $N_{\text{in}}$ inliers, drawn i.i.d. $\text{NORMAL}(\mathbf{0}, (\sigma_{\text{in}}^2/d) \mathbf{\Pi}_L)$  |
| $\mathcal{X}_{\text{out}}$ | Set of $N_{\text{out}}$ outliers, drawn i.i.d. $\text{NORMAL}(\mathbf{0}, (\sigma_{\text{out}}^2/D) \mathbf{I}_D)$ |
| $\mathcal{X}$              | The set $\mathcal{X}_{\text{in}} \cup \mathcal{X}_{\text{out}}$ containing all the data points                     |

There are a few useful intuitions associated with this model. As the inlier sampling ratio  $\rho_{\text{in}}$  increases, the inliers fill out the subspace  $L$  more completely so the linear structure becomes more evident. As the outlier sampling ratio  $\rho_{\text{out}}$  increases, the outliers become more distracting and they may even start to exhibit some linear structure due to chance. Next, observe that we have scaled the points so that their energy is independent of the dimensional parameters:

$$\mathbb{E} \|\mathbf{x}\|^2 = \sigma_{\text{in}}^2 \quad \text{for } \mathbf{x} \in \mathcal{X}_{\text{in}} \quad \text{and} \quad \mathbb{E} \|\mathbf{x}\|^2 = \sigma_{\text{out}}^2 \quad \text{for } \mathbf{x} \in \mathcal{X}_{\text{out}}.$$

As a result, when  $\sigma_{\text{in}}^2 = \sigma_{\text{out}}^2$ , we cannot screen outliers just by looking at their energy. The sampling ratios and the variances contain most of the information about the behavior of this model.

**1.7. Exact Recovery under the Haystack Model.** We have been able to establish conditions under which REAPER and S-REAPER provably find the low-dimensional subspace  $L$  in the Haystack Model with high probability. A sufficient condition for exact recovery is that  $\rho_{\text{in}}$ , the number of inliers *per subspace dimension*, should increase linearly with  $\rho_{\text{out}}$ , the number of outliers *per ambient*

*dimension.* As a consequence, we can find a low-dimensional linear structure in a high-dimensional space given a tiny number of examples, even when the number of outliers seems exorbitant.

**Theorem 1.1** (Exact Recovery for the Haystack Model). *Fix a number  $\beta > 0$ , and assume that  $1 \leq d \leq (D - 1)/2$ . Let  $L$  be an arbitrary  $d$ -dimensional subspace of  $\mathbb{R}^D$ , and draw the dataset  $\mathcal{X}$  at random according to the Haystack Model on page 6.*

(i) *Suppose that the sampling ratios and the variances satisfy the relation*

$$\rho_{\text{in}} \geq C_1 + C_2 \beta + C_3 \cdot \frac{\sigma_{\text{out}}}{\sigma_{\text{in}}} \cdot (\rho_{\text{out}} + 1 + 4\beta). \quad (1.7)$$

*Then  $\Pi_L$  is the unique solution to the REAPER problem (1.4) except with probability  $4e^{-\beta d}$ .*

(ii) *Suppose that the sampling ratios satisfy the relation*

$$\rho_{\text{in}} \geq \widetilde{C}_1 + \widetilde{C}_2 \beta + \widetilde{C}_3 \rho_{\text{out}}. \quad (1.8)$$

*Then  $\Pi_L$  is the unique solution to the S-REAPER problem (1.6) except with probability  $4e^{-\beta d}$ .*

*The numbers  $C_i$  and  $\widetilde{C}_i$  are positive universal constants.*

Theorem 1.1 is a consequence of a *deterministic* exact recovery condition that we discuss more fully in Section 3. We obtain Theorem 1.1 when we compute the values of some summary statistics, assuming that the data is drawn from the Haystack Model. See Appendix B for the details of this argument, including tighter conditions that hold for the full parameter range  $1 \leq d \leq D - 1$ .

For clarity, we have suppressed the values of the constants in (1.7) and (1.8). Our proof yields reasonable numerical estimates. For example,  $C_1 < 13$  and  $C_2 < 7$  and  $C_3 < 16$ . Experiments indicate that  $C_3$ , the constant of proportionality between  $\rho_{\text{in}}$  and  $\rho_{\text{out}}$ , is less than two.

**1.7.1. A Needle in a Haystack.** To complement the experiments in Section 1.5, we include a quantitative result on the behavior of S-REAPER for finding a needle in a haystack.

**Fact 1.2.** Assume the ambient dimension  $D \geq 50$ , and fix a one-dimensional subspace  $L$ . Suppose the number of inliers  $N_{\text{in}} \geq 13$  and the number of outliers  $N_{\text{out}} = 2D$ . Draw a dataset  $\mathcal{X}$  according to the Haystack Model on page 6. Then  $\Pi_L$  is the unique solution to the S-REAPER problem (1.6) at least 99% of the time.

Fact 1.2 follows from Theorem B.1. We highlight this *ad hoc* result to underscore the point that it takes only a few good observations for (1.6) to identify inliers that fall on a line. Indeed, Fact 1.2 is valid even in dimension  $D = 10^{37}$  with  $N_{\text{out}} = 2 \times 10^{37}$ .

**1.7.2. Is it Hard to Find a Needle in a Haystack?** Our low-dimensional intuition suggests that it might be difficult to identify the subspace in the Haystack Model, but there are several algorithms that can complete the task. For instance, the EM algorithm for normal mixtures has no trouble with this example [Rec12]. In contrast, for natural data, we have found that the S-REAPER method often outperforms other methods by a reasonable margin.

Let us emphasize that we do not intend for the Haystack Model to describe natural data. We have introduced this formulation as a way to probe the limitations of the REAPER and S-REAPER problems in experiments. The Haystack Model also gives us some theoretical insight on their exact recovery behavior. We believe that Theorem 1.1 is compelling in part because the optimization problems are derived without knowledge of the model. It would be valuable to analyze REAPER and S-REAPER with more sophisticated data models, but we leave this task for future work.



**1.8. Summary of Contributions.** This work partakes in a larger research vision: Given a difficult nonconvex optimization problem, it is often more effective to solve a convex variant than to accept a local minimizer of the original problem.

We believe that the main point of interest in this paper is our application of convex optimization to solve a problem involving subspaces. There are two key observations here. We parameterize subspaces by orthoprojectors, and then we replace the set of orthoprojectors with its convex hull  $\{\mathbf{P} : \mathbf{0} \preceq \mathbf{P} \preceq \mathbf{I}\}$ . Neither of these ideas is new, but this research provides clear evidence that they can be used to address important geometric questions in data analysis.

To support the claim that our approach works, we develop a deterministic analysis that describes when REAPER and S-REAPER can recover a subspace exactly. This analysis yields summary statistics that predict when the optimization problems succeed. For data that conform to the Haystack Model on page 6, we can evaluate these summary statistics to obtain the exact recovery guarantees of Theorem 1.1. See Section 3 for more information.

In Section 4, we present an efficient iterative algorithm for solving the semidefinite program (1.4). The algorithm typically requires a modest number of spectral computations, so it scales to relatively large problem instances.

In Section 5, we use our algorithm to perform some experiments that describe the behavior of REAPER and S-REAPER. First, we study synthetic data problems to understand the limitations of S-REAPER; these experiments indicate that Theorem 1.1 is qualitatively correct. We also describe some stylized applications involving natural data; these results show that S-REAPER dominates other robust PCA techniques in several situations. Furthermore, the linear models we obtain appear to generalize better than the robust models obtained with some other methods.

We describe a variety of future research directions in Section 6. In particular, we propose some alternative formulations of the REAPER problem. We also suggest some ways to extend our approach to manifold learning and hybrid linear modeling.

The technical details of our work appear in the appendices.

## 2. PREVIOUS WORK

Robust linear modeling has been an active subject of research for over three decades. Although many classical approaches have strong robustness properties *in theory*, the proposals usually involve intractable computational problems. More recently, researchers have developed several techniques, based on convex optimization, that are computationally efficient and admit some theoretical guarantees. In this section, we summarize classical and contemporary work on robust linear modeling, with a focus on the numerical aspects. We recommend the books [HR09, MMY06, RL87] for a comprehensive discussion of robust statistics.

**2.1. Classical Strategies for Robust Linear Modeling.** We begin with an overview of the major techniques that have been proposed in the statistics literature. The theoretical contributions in this area focus on breakdown points and influence functions of estimators. Researchers tend to devote less attention to the computational challenges inherent in these formulations.

**2.1.1. Robust Combination of Residuals.** Historically, one of the earliest approaches to linear regression is to minimize the sum of (nonorthogonal) residuals. This is the principle of *least absolute deviations* (LAD). Early proponents of this idea include Galileo, Boscovich, Laplace, and Edgeworth. See [Har74a, Har74b, Dod87] for some historical discussion. It appears that *orthogonal* regression with LAD was first considered in the late 1980s [OW85, SW87, Nyq88]; the extension from orthogonal regression to PCA seems to be even more recent [Wat01, DZHZ06]. LAD has also been considered as a method for hybrid linear modeling in [ZSL09, LZ11]. We are not aware of tractable algorithm for these formulations.

There are many other robust methods for combining residuals aside from LAD. An approach that has received wide attention is to minimize the median of the squared residuals [Rou84, RL87].

Other methods appear in the books [HR09, MMY06]. These formulations are generally not computationally tractable.

**2.1.2. Robust Estimation of Covariance Matrix.** Another standard technique for robust PCA is to form a robust estimate of the covariance matrix of the data [Mar76, HR09, MMY06, DGK81, Dav87, XY95, CH00]. The classical robust estimator of covariance is based on the maximum likelihood principle, but there are many other approaches, including S-estimators, the minimum covariance determinant (MCD), the minimum volume ellipsoid (MVE), and the Stahel–Donoho estimator. See [MMY06, Sec. 6] for a review. We are not aware of any scalable algorithm for implementing these methods that offers a guarantee of correctness.

**2.1.3. Projection Pursuit PCA.** Projection pursuit (often abbreviated PP-PCA) is a procedure that constructs principal components one at a time by finding a direction that maximizes a robust measure of scale, removing the component of the data in this direction, and repeating the process. The initial proposal appears in [HR09, 1st ed.], and it has been explored by many other authors [LC85, Amm93, CFO07, Kwa08, WTH09]. We are aware of only one formulation that provably (approximately) maximizes the robust measure of scale at each iteration [MT11], but there are no overall guarantees for this algorithm.

**2.1.4. Screening for Outliers.** Another common approach is to remove possible outliers and then estimate the underlying subspace by PCA [CW82, TB01, TB03]. The classical methods offer very limited guarantees. There are some recent algorithms that are provably correct [Bru09, XCM10] under model assumptions.

**2.1.5. RANSAC.** The randomized sample consensus (RANSAC) method is a randomized iterative procedure for fitting models to noisy data consisting of inliers and outliers [FB81]. Under some assumptions, RANSAC will eventually identify a linear model for the inliers, but there are no guarantees on the number of iterations required.

**2.1.6. Spherical PCA.** A useful method for fitting a robust linear model is to center the data robustly, project it onto a sphere, and then apply standard PCA. This approach is due to [LMS<sup>+</sup>99]. Maronna et al. [MMY06] recommend it as a preferred method for robust PCA. The technique is computationally practical, but it has limited theoretical guarantees.

**2.2. Approaches Based on Convex Optimization.** Recently, researchers have started to develop effective techniques for robust linear modeling that are based on convex optimization. These formulations invite a variety of tractable algorithms, and they have theoretical guarantees under appropriate model assumptions.

**2.2.1. Deconvolution Methods.** One class of techniques for robust linear modeling is based on splitting a data matrix into a low-rank model plus a corruption. The first approach of this form is due to Chandrasekaran et al. [CSPW11]. Given an observed matrix  $\mathbf{X}$ , they propose to solve the semidefinite problem

$$\text{minimize} \quad \|\mathbf{P}\|_{S_1} + \gamma \|\mathbf{C}\|_{\ell_1} \quad \text{subject to} \quad \mathbf{X} = \mathbf{P} + \mathbf{C}. \quad (2.1)$$

Minimizing the Schatten 1-norm  $\|\cdot\|_{S_1}$  promotes low rank, while minimizing the vector  $\ell_1$  norm promotes sparsity. The regularization parameter  $\gamma$  negotiates a tradeoff between the two goals. Candès et al. [CLMW11] study the performance of (2.1) for robust linear modeling in the setting where individual entries of the matrix  $\mathbf{X}$  are subject to error.

A related proposal is due to Xu et al. [XCS10a, XCS10b] and independently to McCoy and Tropp [MT11]. These authors recommend solving the decomposition problem

$$\text{minimize} \quad \|\mathbf{P}\|_{S_1} + \gamma \|\mathbf{C}\|_{1 \rightarrow 2}^* \quad \text{subject to} \quad \mathbf{X} = \mathbf{P} + \mathbf{C}. \quad (2.2)$$

The norm  $\|\cdot\|_{1 \rightarrow 2}^*$  returns the sum of Euclidean norms of the columns of its argument. This formulation is appropriate for inlier-outlier data models, where entire columns of the data matrix may be corrupted.

Both (2.1) and (2.2) offer some theoretical guarantees under appropriate model assumptions. The most common algorithmic framework for these problems depends on the alternating direction method of multipliers (ADMM), which is a type of augmented Lagrangian scheme. These algorithms converge slowly, so it takes excessive computation to obtain a high-accuracy solution.

**2.2.2. Precedents for the REAPER Problem.** The REAPER problem (1.4) is a semidefinite relaxation of the  $\ell_1$  orthogonal distance problem (1.3). Our work refines a weaker relaxation of (1.3) proposed by Zhang and Lerman:

$$\text{minimize } \sum_{\mathbf{x} \in \mathcal{X}} \|\mathbf{x} - \mathbf{P}\mathbf{x}\| \quad \text{subject to } \text{tr } \mathbf{P} = 1.$$

Some of the ideas in the theoretical analysis of REAPER are drawn from [LZ10, ZL11]. The IRLS algorithm for (1.4) and the convergence analysis that we present in Section 4 are also based on the earlier work.

The idea of relaxing a difficult nonconvex program to obtain a convex problem is well established in the literature on combinatorial optimization. Research on linear programming relaxations is summarized in [Vaz03]. Some significant works on semidefinite relaxation include [LS91, GW95].

### 3. DETERMINISTIC CONDITIONS FOR EXACT RECOVERY OF A SUBSPACE

In this section, we develop a deterministic model for a dataset that consists of inliers located in a fixed subspace and outliers that may appear anywhere in the ambient space. We introduce summary statistics that allow us to study when REAPER or S-REAPER can recover exactly the subspace containing the inliers. The result for the Haystack Model, Theorem 1.1, follows when we estimate the values for these summary statistics.

**3.1. A Deterministic Model for Studying Exact Recovery of a Subspace.** We begin with a deterministic model for a dataset that consists of low-dimensional inliers mixed with high-dimensional outliers.

TABLE 3.1. *The In & Out Model.* A deterministic model for data with linear structure that is contaminated with outliers.

|                            |                                                                                             |
|----------------------------|---------------------------------------------------------------------------------------------|
| $D$                        | Dimension of the ambient space                                                              |
| $L$                        | A proper $d$ -dimensional subspace of $\mathbb{R}^D$                                        |
| $N_{\text{in}}$            | Number of inliers                                                                           |
| $N_{\text{out}}$           | Number of outliers                                                                          |
| $\mathcal{X}_{\text{in}}$  | Set of $N_{\text{in}}$ inliers, at arbitrary locations in $L$                               |
| $\mathcal{X}_{\text{out}}$ | Set of $N_{\text{out}}$ outliers, at arbitrary locations in $\mathbb{R}^D \setminus L$      |
| $\mathcal{X}$              | Set $\mathcal{X}_{\text{in}} \cup \mathcal{X}_{\text{out}}$ containing all the observations |

The key point about this model is that the inliers are all contained within the subspace  $L$ , so it is reasonable to investigate when we can identify  $L$  exactly. For notational convenience, we assume that no observation is repeated, but this assumption plays no role in our analysis. The Haystack Model on page 6 almost surely produces a dataset that conforms to the In & Out Model.

**3.2. Summary Parameters for the In & Out Model.** In general, we cannot hope to identify the subspace  $L$  that appears in the In & Out Model without further assumptions on the data. This section describes some summary statistics that allow us to check when REAPER or S-REAPER succeeds in computing  $L$ . These quantities play a role analogous to the *coherence statistic* in the sparsity literature [DH01].

First, let us discuss what property of the inliers might allow us to determine the subspace  $L$  that contains them. Roughly, we need the inlying observations to permeate  $L$ . To quantify this idea, we define the *permeance statistic* (with respect to  $L$ ) by the formula

$$\mathcal{P}(L) := \inf_{\substack{\mathbf{u} \in L \\ \|\mathbf{u}\|=1}} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} |\langle \mathbf{u}, \mathbf{x} \rangle|. \quad (3.1)$$

Similarly, we define the *spherical permeance statistic* as

$$\widetilde{\mathcal{P}}(L) := \inf_{\substack{\mathbf{u} \in L \\ \|\mathbf{u}\|=1}} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} |\langle \mathbf{u}, \widetilde{\mathbf{x}} \rangle|. \quad (3.2)$$

When  $\mathcal{P}(L) = 0$  (equivalently,  $\widetilde{\mathcal{P}}(L) = 0$ ), there is a vector in the subspace  $L$  that is orthogonal to all the inliers. In this case, the data do not contain enough information for us to recover  $L$ . On the other hand, when one of the permeance statistics is large, the inliers corroborate each direction in the subspace. In this case, we expect our methods for linear modeling to be more effective.

To ensure recovery of  $L$ , we must also be certain that the outliers do not exhibit linear structure. Otherwise, there is no way to decide whether a subspace captures linear structure from the inliers or the outliers. Indeed, in the extreme case where the outliers are arranged on a line, any sensible estimation procedure would have to select a subspace that contains this line.

Let us introduce a quantity that measures the linear dependency of the outliers within a given subspace. The *linear structure statistic* is defined for each subspace  $M \subset \mathbb{R}^D$  by the formula

$$\mathcal{S}(M)^2 := \sup_{\substack{\mathbf{u} \in M \\ \|\mathbf{u}\|=1}} \sum_{\mathbf{x} \in \mathcal{X}_{\text{out}}} |\langle \mathbf{u}, \mathbf{x} \rangle|^2 = \|\Pi_M \mathbf{X}_{\text{out}}\|^2, \quad (3.3)$$

where  $\mathbf{X}_{\text{out}}$  is a  $D \times N_{\text{out}}$  matrix whose columns are the outliers, arranged in fixed order. Similarly, we can define the *spherical linear structure statistic*

$$\widetilde{\mathcal{S}}(M)^2 := \sup_{\substack{\mathbf{u} \in M \\ \|\mathbf{u}\|=1}} \sum_{\mathbf{x} \in \mathcal{X}_{\text{out}}} |\langle \mathbf{u}, \widetilde{\Pi_M \mathbf{x}} \rangle|^2 = \|\widetilde{\Pi_M \mathbf{X}_{\text{out}}}\|^2, \quad (3.4)$$

where  $\widetilde{\Pi_M \mathbf{X}_{\text{out}}}$  is a  $D \times N_{\text{out}}$  matrix whose columns are the vectors  $\widetilde{\Pi_M \mathbf{x}}$  for each  $\mathbf{x} \in \mathcal{X}_{\text{out}}$ . The linear structure statistics tend to be large when the outliers are nearly collinear (after projection onto  $M$ ), while the statistics tend to be small when the outliers are close to orthogonal.

**3.3. Conditions for Exact Recovery of a Subspace.** We are now prepared to state conditions when either REAPER or S-REAPER recovers the subspace  $L$  containing the inliers exactly.

**Theorem 3.1.** *Let  $L$  be a  $d$ -dimensional subspace of  $\mathbb{R}^D$ , and assume that  $\mathcal{X}$  is a dataset that follows the In & Out Model on page 10.*

(i) *Suppose that*

$$\mathcal{P}(L) > \sqrt{2d} \cdot \widetilde{\mathcal{S}}(L^\perp) \cdot \mathcal{S}(\mathbb{R}^D). \quad (3.5)$$

*Then  $\Pi_L$  is the unique minimizer of the REAPER problem (1.4).*

(ii) *Suppose that*

$$\widetilde{\mathcal{P}}(L) > \sqrt{2d} \cdot \widetilde{\mathcal{S}}(L^\perp) \cdot \widetilde{\mathcal{S}}(\mathbb{R}^D). \quad (3.6)$$

*Then  $\Pi_L$  is the unique minimizer of the S-REAPER problem (1.6).*

The permeance statistics  $\mathcal{P}$  and  $\tilde{\mathcal{P}}$  are defined in (3.1) and (3.2), while the linear structure statistics  $\mathcal{S}$  and  $\tilde{\mathcal{S}}$  are defined in (3.3) and (3.4).

Theorem 3.1 shows that the REAPER (resp. S-REAPER) problem can find the subspace  $L$  containing the inliers whenever the (spherical) permeance statistic is sufficiently large as compared with the (spherical) linear structure statistic. Be aware that the inequalities (3.5) and (3.6) are not comparable; neither one implies the other. The main difference is due to the fact that the condition (3.6) for S-REAPER is invariant to the scalings of the data points.

The proof of Theorem 3.1 appears in Appendix A.1. The main idea in the argument is to verify that, under the appropriate sufficient condition, the objective function of REAPER or S-REAPER increases if we perturb the decision variable  $\mathbf{P}$  away from  $\mathbf{\Pi}_L$ .

**3.4. Exact Recovery for the Haystack Model.** The conditions in Theorem 3.1 hold in highly nontrivial—even surprising—regimes. In this section, we make some heuristic calculations to indicate how the spherical summary statistics behave when  $\mathcal{X}$  is drawn from the Haystack Model on page 6. We prove rigorous results, including Theorem 1.1, in Appendix B.

First, let us consider the spherical permeance statistic  $\tilde{\mathcal{P}}(L)$ . Fix a unit vector  $\mathbf{u} \in L$ , and suppose that  $\mathbf{x}$  is a random unit vector in  $L$ . The random variable  $\langle \mathbf{u}, \mathbf{x} \rangle$  is approximately normal with mean zero and variance  $d^{-1}$ , so we have  $\mathbb{E} |\langle \mathbf{u}, \mathbf{x} \rangle| \approx \sqrt{2/(\pi d)}$ . It follows that

$$\tilde{\mathcal{P}}(L) \approx \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} |\langle \mathbf{u}, \mathbf{x} \rangle| \approx \sqrt{\frac{2}{\pi}} \cdot \frac{N_{\text{in}}}{\sqrt{d}}.$$

We have omitted the infimum over  $\mathbf{u}$  from (3.1) because it does not make a substantial difference. See Lemma B.3 for a rigorous computation.

Next, we examine the spherical linear structure statistic  $\tilde{\mathcal{S}}(M) = \|\widetilde{\mathbf{\Pi}_M \mathbf{X}_{\text{out}}}\|$ . This quantity is just the norm of a matrix with  $N_{\text{out}}$  columns that are i.i.d. uniform on the unit sphere in  $M$ . For many purposes, this type of matrix behaves like a  $\dim(M) \times N_{\text{out}}$  Gaussian matrix whose entries have variance  $\dim(M)^{-1}$ . Standard bounds for the norm of a Gaussian matrix [DS01, Thm. 2.13] give the estimate

$$\tilde{\mathcal{S}}(M) \approx \sqrt{\frac{N_{\text{out}}}{\dim(M)}} + 1.$$

Therefore, we can estimate the product

$$\tilde{\mathcal{S}}(L^\perp) \cdot \tilde{\mathcal{S}}(\mathbb{R}^D) \approx \left[ \sqrt{\frac{N_{\text{out}}}{D}} + 1 \right] \left[ \sqrt{\frac{N_{\text{out}}}{D-d}} + 1 \right] \leq \left[ \sqrt{\frac{N_{\text{out}}}{D-d}} + 1 \right]^2.$$

See Lemma B.5 for more careful calculations.

To conclude, when the data  $\mathcal{X}$  is drawn from the Haystack Model, the sufficient condition (3.6) for S-REAPER to succeed becomes

$$\frac{N_{\text{in}}}{d} \gtrsim \sqrt{\pi} \left[ \sqrt{\frac{N_{\text{out}}}{D-d}} + 1 \right]^2 \quad (3.7)$$

Assume that  $D-d$  is comparable with  $D$ . Then the heuristic (3.7) indicates that the inlier sampling ratio  $\rho_{\text{in}} = N_{\text{in}}/d$  should increase linearly with the outlier sampling ratio  $\rho_{\text{out}} = N_{\text{out}}/D$ . This scaling matches the rigorous bound in Theorem 1.1.

#### 4. AN ITERATIVE REWEIGHTED LEAST-SQUARES ALGORITHM FOR REAPER

In this section, we present a numerical algorithm for solving the REAPER problem (1.4). Of course, we can also solve S-REAPER by spherizing the data before applying this algorithm. Our approach is based on the iterative reweighted least squares (IRLS) framework [Bj96, Sec. 4.5.2].

At each step, we solve a least-squares problem with an adapted set of scaling parameters. The subproblem has a closed-form solution that we can obtain by computing an eigenvalue (or singular value) decomposition. Empirically, the IRLS approach exhibits linear convergence, so it can achieve high accuracy without an exorbitant number of iterations. It is possible to apply other types of optimization algorithms, such as interior-point methods, to solve REAPER, but we are not aware of an alternative approach that scales to large datasets.

**4.1. A Weighted Least-Squares Problem.** The REAPER problem (1.4) cannot be solved in closed form, but there is a closely related problem whose minimizer can be computed efficiently. We use this mathematical program and its solution to build an algorithm for solving REAPER.

For each  $\mathbf{x} \in \mathcal{X}$ , let  $\beta_{\mathbf{x}}$  be a nonnegative weight parameter. Consider the weighted least-squares problem with the same constraints as (1.4):

$$\text{minimize } \sum_{\mathbf{x} \in \mathcal{X}} \beta_{\mathbf{x}} \|\mathbf{x} - \mathbf{P}\mathbf{x}\|^2 \quad \text{subject to } \mathbf{0} \preceq \mathbf{P} \preceq \mathbf{I} \quad \text{and} \quad \text{tr } \mathbf{P} = d. \quad (4.1)$$

---

**Algorithm 4.1** *Solving the weighted least-squares problem (4.1)*

---

INPUT:

- A set  $\mathcal{X}$  of observations in  $\mathbb{R}^D$
- A nonnegative weight  $\beta_{\mathbf{x}}$  for each  $\mathbf{x} \in \mathcal{X}$
- The parameter  $d$  in (4.1), where  $0 < d < D$

OUTPUT:

- A  $D \times D$  matrix  $\mathbf{P}_\star$  that solves (4.1)

PROCEDURE:

- 1 Form the  $D \times D$  weighted covariance matrix

$$\mathbf{C} \leftarrow \sum_{\mathbf{x} \in \mathcal{X}} \beta_{\mathbf{x}} \mathbf{x} \mathbf{x}^t$$

- 2 Compute an eigenvalue decomposition  $\mathbf{C} = \mathbf{U} \cdot \text{diag}(\lambda_1, \dots, \lambda_D) \cdot \mathbf{U}^t$  with eigenvalues in nonincreasing order:  $\lambda_1 \geq \dots \geq \lambda_D \geq 0$

- 3 **if**  $\lambda_{\lfloor d \rfloor + 1} = 0$  **then**
  - a **for**  $i = 1, \dots, D$  **do**

$$\nu_i \leftarrow \begin{cases} 1, & i \leq \lfloor d \rfloor \\ d - \lfloor d \rfloor, & i = \lfloor d \rfloor + 1 \\ 0, & \text{otherwise} \end{cases}$$

- 4 **else**
  - a **for**  $i = \lfloor d \rfloor + 1, \dots, D$  **do**
    - i Set

$$\theta \leftarrow \frac{i - d}{\sum_{k=1}^i \lambda_k^{-1}}$$

- ii **if**  $\lambda_i > \theta \geq \lambda_{i+1}$  **then break for**

- b **for**  $i = 1, \dots, D$  **do**

$$\nu_i \leftarrow \begin{cases} 1 - \frac{\theta}{\lambda_i}, & \lambda_i > \theta \\ 0, & \lambda_i \leq \theta \end{cases}$$

- 5 **return**  $\mathbf{P}_\star := \mathbf{U} \cdot \text{diag}(\nu_1, \dots, \nu_D) \cdot \mathbf{U}^t$
-

Algorithm 4.1 describes a computational procedure for solving (4.1). The basic idea is to compute a weighted covariance matrix  $\mathbf{C}$  from the data. We use spectral shrinkage to scale the eigenvalues of  $\mathbf{C}$  to the range  $[0, 1]$ . The amount of shrinkage is determined by a water-filling calculation, which ensures that the output matrix has the correct trace. In Appendix C.1, we present a more mathematical formulation of Algorithm 4.1, and we use this statement to establish correctness.

**4.1.1. Discussion and Computational Costs.** The bulk of the computation in Algorithm 4.1 takes place during the spectral calculation in Steps 1 and 2. We need  $\mathcal{O}(ND^2)$  arithmetic operations to form the weighted covariance matrix, and the spectral calculation requires  $\mathcal{O}(D^3)$ . The remaining steps of the algorithm have order  $\mathcal{O}(D)$ .

We have presented Algorithm 4.1 in the most direct way possible. In practice, it is usually more efficient to form a  $D \times N$  matrix  $\mathbf{W}$  with columns  $\sqrt{\beta_{\mathbf{x}}} \mathbf{x}$  for  $\mathbf{x} \in \mathcal{X}$ , to compute a thin SVD  $\mathbf{W} = \mathbf{U}\mathbf{\Sigma}\mathbf{V}^t$ , and to set  $\mathbf{\Lambda} = \mathbf{\Sigma}^2$ . This approach is also more stable. In some situations, it is possible to accelerate the spectral calculations using randomized dimension reduction, as in [HMT11].

**4.2. Iterative Reweighted Least-Squares Algorithm.** We are prepared to develop an algorithm for solving the REAPER problem (1.4). Let us begin with the intuition. Suppose that  $\mathbf{P}_\star$  solves the weighted least-squares problem (4.1) where  $\beta_{\mathbf{x}} \approx \|\mathbf{x} - \mathbf{P}_\star \mathbf{x}\|^{-1}$  for each  $\mathbf{x} \in \mathcal{X}$ . Examining the objective function of (4.1), we see that

$$\sum_{\mathbf{x} \in \mathcal{X}} \beta_{\mathbf{x}} \|\mathbf{x} - \mathbf{P}_\star \mathbf{x}\|^2 \approx \sum_{\mathbf{x} \in \mathcal{X}} \|\mathbf{x} - \mathbf{P}_\star \mathbf{x}\|.$$

Therefore, it seems plausible that  $\mathbf{P}_\star$  is also close to the minimizer of the REAPER problem (1.4). This heuristic is classical, and it remains valid in our setting.

---

**Algorithm 4.2** *IRLS algorithm for solving the REAPER problem (1.4)*

---

INPUT:

- A set  $\mathcal{X}$  of observations in  $\mathbb{R}^D$
- The dimension parameter  $d$  in (1.4), where  $0 < d < D$
- A regularization parameter  $\delta > 0$
- A stopping tolerance  $\varepsilon > 0$

OUTPUT:

- A  $D \times D$  matrix  $\mathbf{P}_\star$  that satisfies  $\mathbf{0} \preceq \mathbf{P}_\star \preceq \mathbf{I}$  and  $\text{tr } \mathbf{P}_\star = d$

PROCEDURE:

- 1 Initialize the variables:
  - a Set the iteration counter  $k \leftarrow 0$
  - b Set the initial error  $\alpha^{(0)} \leftarrow +\infty$
  - c Set the weight  $\beta_{\mathbf{x}} \leftarrow 1$  for each  $\mathbf{x} \in \mathcal{X}$
- 2 **do**
  - a Increment  $k \leftarrow k + 1$
  - b Use Algorithm 4.1 to compute an optimal point  $\mathbf{P}^{(k)}$  of (4.1) with weights  $\beta_{\mathbf{x}}$
  - c Let  $\alpha^{(k)}$  be the optimal value of (4.1) at  $\mathbf{P}^{(k)}$
  - d Update the weights:

$$\beta_{\mathbf{x}} \leftarrow \frac{1}{\max\{\delta, \|\mathbf{x} - \mathbf{P}^{(k)} \mathbf{x}\|\}} \quad \text{for each } \mathbf{x} \in \mathcal{X}$$

**until** the objective fails to decrease:  $\alpha^{(k)} \geq \alpha^{(k-1)} - \varepsilon$

- 3 **Return**  $\mathbf{P}_\star = \mathbf{P}^{(k)}$
-

This observation motivates an iterative procedure for solving (1.4). Let  $\delta$  be a positive regularization parameter. Set the iteration counter  $k \leftarrow 0$ , and set the weight  $\beta_{\mathbf{x}} \leftarrow 1$  for each  $\mathbf{x} \in \mathcal{X}$ . We solve (4.1) with the weights  $\beta_{\mathbf{x}}$  to obtain a matrix  $\mathbf{P}^{(k)}$ , and then we update the weights according to the formula

$$\beta_{\mathbf{x}} \leftarrow \frac{1}{\max\{\delta, \|\mathbf{x} - \mathbf{P}^{(k)}\mathbf{x}\|\}} \quad \text{for each } \mathbf{x} \in \mathcal{X}.$$

In other words, we emphasize the observations that are explained well by the current model. The presence of the regularization  $\delta$  ensures that no single point can gain undue influence. We increment  $k$ , and we repeat the process until it has converged.

We summarize the computational procedure in Algorithm 4.2. The following result shows that Algorithm 4.2 is guaranteed to converge to a point whose value is close to the optimal value of the REAPER problem (1.4).

**Theorem 4.1** (Convergence of IRLS). *Assume that the set  $\mathcal{X}$  of observations does not lie in the union of two strict subspaces of  $\mathbb{R}^D$ . Then the iterates of Algorithm 4.2 with  $\varepsilon = 0$  converge to a point  $\mathbf{P}_{\delta}$  that satisfies the constraints of the REAPER problem (1.4). Moreover, the objective value at  $\mathbf{P}_{\delta}$  satisfies the bound*

$$\sum_{\mathbf{x} \in \mathcal{X}} \|\mathbf{x} - \mathbf{P}_{\delta}\mathbf{x}\| - \sum_{\mathbf{x} \in \mathcal{X}} \|\mathbf{x} - \mathbf{P}_{\star}\mathbf{x}\| \leq \frac{1}{2}\delta|\mathcal{X}|,$$

where  $\mathbf{P}_{\star}$  is an optimal point of REAPER.

The proof of Theorem 4.1 is similar to established convergence arguments [ZL11, Thms. 6 and 8], which follow the schema in [CM99, VE80]. See Appendix C for a summary of the proof. It can also be shown that the point  $\mathbf{P}_{\delta}$  is close to the optimal point  $\mathbf{P}_{\star}$  in Frobenius norm.

**4.2.1. Discussion and Computational Costs.** The bulk of the computational cost in Algorithm 4.2 is due the solution of the subproblem in Step 2b. Therefore, each iteration costs  $\mathcal{O}(ND^2)$ . The algorithm converges linearly in practice, so we need  $\mathcal{O}(\log(1/\eta))$  iterations to achieve an error of  $\eta$ .

Figure 4.1 shows that, empirically, Algorithm 4.2 exhibits linear convergence. In this experiment, we have drawn the data from the Haystack Model on page 6 with ambient dimension  $D = 100$  and  $N_{\text{out}} = 200$  outliers. The curves mark the performance for several choices of the model dimension  $d$  and the inlier sampling ratio  $\rho_{\text{in}} = N_{\text{in}}/d$ . For this plot, we run Algorithm 4.2 with the regularization parameter  $\delta = 10^{-10}$  and the error tolerance  $\varepsilon = 10^{-15}$  to obtain a sequence  $\{\mathbf{P}^{(k)}\}$  of iterates. We compare the computed iterates with an optimal point  $\mathbf{P}_{\star}$  of the REAPER problem (1.4) obtained by solving REAPER with the Matlab package CVX [GB08, GB10] at the highest-precision setting. The error is measured in Frobenius norm.

For synthetic data, the number of iterations required for Algorithm 4.2 seems to depend on the difficulty of the problem instance. Indeed, it may take as many as 200 iterations for the method to converge on challenging examples. In experiments with natural data, we usually obtain good performance after 20 iterations or so. In practice, Algorithm 4.2 is also much faster than algorithms [XCS10b, MT11] for solving the low-leverage decomposition problem (2.2).

The stopping criterion in Algorithm 4.2 is motivated by the fact that the objective value must decrease in each iteration. This result is a consequence of the proof of Theorem 4.1; see (C.7). Taking  $\varepsilon$  on the order of machine precision ensures that the algorithm terminates when the iterates are dominated by numerical errors. In practice, we achieve very precise results when  $\varepsilon = 10^{-15}$  or even  $\varepsilon = 0$ . In many applications, this degree of care is excessive, and we can obtain a reasonable solution for much larger values of  $\varepsilon$ .

Theorem 4.1 requires a technical condition, namely that the observations do not lie in the union of two strict subspaces. This condition cannot hold unless we have at least  $2D - 1$  observations. In practice, we find that this restriction is unnecessary for the algorithm to converge. It seems to be an artifact of the proof technique.



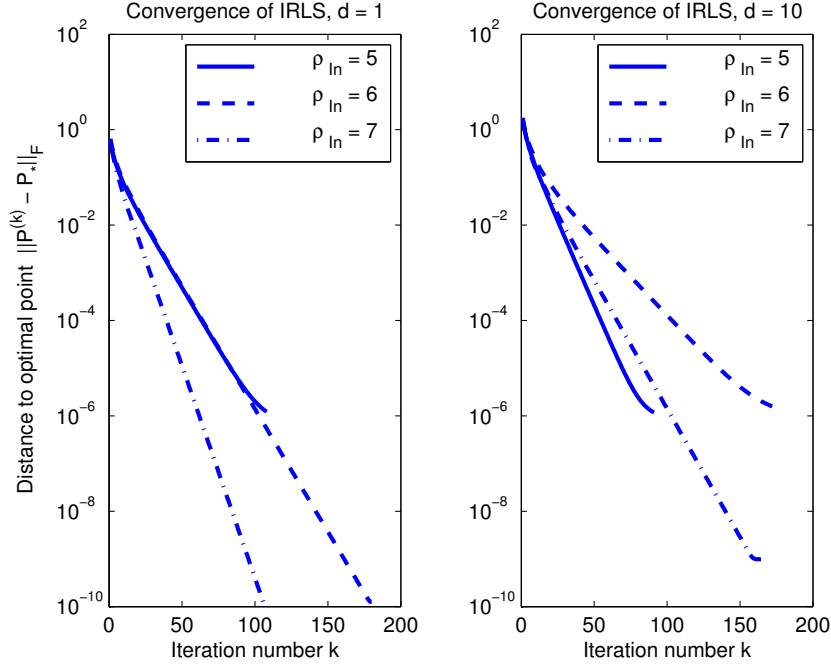


FIGURE 4.1. *Convergence of IRLS to an optimal point.* The data are drawn from the Haystack Model on page 6 with ambient dimension  $D = 100$  and  $N_{\text{out}} = 200$  outliers. Each curve is associated with a particular choice of model dimension  $d$  and inlier sampling ratio  $\rho_{\text{in}} = N_{\text{in}}/d$ . We use Algorithm 4.2 to compute a sequence  $\{P^{(k)}\}$  of iterates, which we compare to an optimal point  $P_*$  of the REAPER problem (1.4). See the text in Section 4.2 for more details of the experiment.

## 5. EXPERIMENTS

In this section, we present some numerical experiments involving REAPER and S-REAPER. To probe the limits of these optimization problems for fitting a linear model in the presence of outliers, we study how they behave on some simple random data models. To gauge their performance in more realistic settings, we consider some stylized problems involving natural data.

**5.1. Experimental Setup.** To solve the REAPER problem (1.4) and the S-REAPER problem (1.6), we use the IRLS method, Algorithm 4.2. We set the regularization parameter  $\delta = 10^{-10}$  and the stopping tolerance  $\varepsilon = 10^{-15}$ .

In all the cases considered here, the dimension parameter  $d$  is a positive integer. Unless we state otherwise, we postprocess the computed optimal point  $P_*$  of REAPER or S-REAPER to obtain a  $d$ -dimensional linear model. When  $P_*$  has rank greater than  $d$ , we construct a  $d$ -dimensional linear model by using the span of the  $d$  dominant eigenvectors of  $P_*$ . Our goal is to use a consistent methodology that involves no parameter tweaking. Other truncation strategies are also possible, as we discuss in Section 6.1.6.

**5.2. Comparisons.** By now, there are a huge number of proposals for robust linear modeling, so we have limited our attention to methods that are computationally tractable. That is, we consider only formulations that have a polynomial-time algorithm for constructing a global solution (up to some tolerance). We do not discuss techniques that involve Monte Carlo simulation, nonlinear programming, etc. because the success of these approaches depends largely on parameter settings and providence. As a consequence, it is hard to evaluate their behavior in a consistent way.

We consider two standard approaches, PCA [Jol02] and spherical PCA [LMS<sup>+</sup>99]. Spherical PCA rescales each observation so it lies on the Euclidean sphere, and then it applies standard PCA. In a recent book, Maronna et al. [MMY06] recommend spherical PCA as the most reliable classical algorithm, which provides a second justification for omitting many other established methods.

We also consider a more recent proposal [XCS10a, XCS10b, MT11], which is called *low-leverage decomposition* (LLD) or *outlier pursuit*. This method decomposes the  $D \times N$  matrix  $\mathbf{X}$  of observations by solving the optimization problem

$$\text{minimize} \quad \|\mathbf{P}\|_{S_1} + \gamma \|\mathbf{C}\|_{1 \rightarrow 2}^* \quad \text{subject to} \quad \mathbf{X} = \mathbf{P} + \mathbf{C} \quad (5.1)$$

where  $\|\cdot\|_{S_1}$  is the Schatten 1-norm and  $\|\cdot\|_{1 \rightarrow 2}^*$  returns the sum of Euclidean norms of the column. The idea is that the optimizer  $(\mathbf{P}_*, \mathbf{C}_*)$  will consist of a low-rank model  $\mathbf{P}_*$  for the data along with a column-sparse matrix  $\mathbf{C}_*$  that identifies the outliers. We always use the parameter choice  $\gamma = 0.8\sqrt{D/N}$ , which seems to be effective in practice.

We do not make comparisons with the rank-sparsity decomposition [CSPW11], which has also been considered for robust linear modeling in [CLMW11]. It is not effective for the problems that we consider here.

**5.3. Phase Transitions for Exact Recovery of Subspaces.** This section studies the limits of the S-REAPER problem by applying it to data drawn from the Haystack Model on page 6. The goal is to understand when S-REAPER can recover a low-dimensional subspace exactly, even when the dataset contains many unstructured outliers. Theorem 1.1 tells us that the inlier sampling ratio

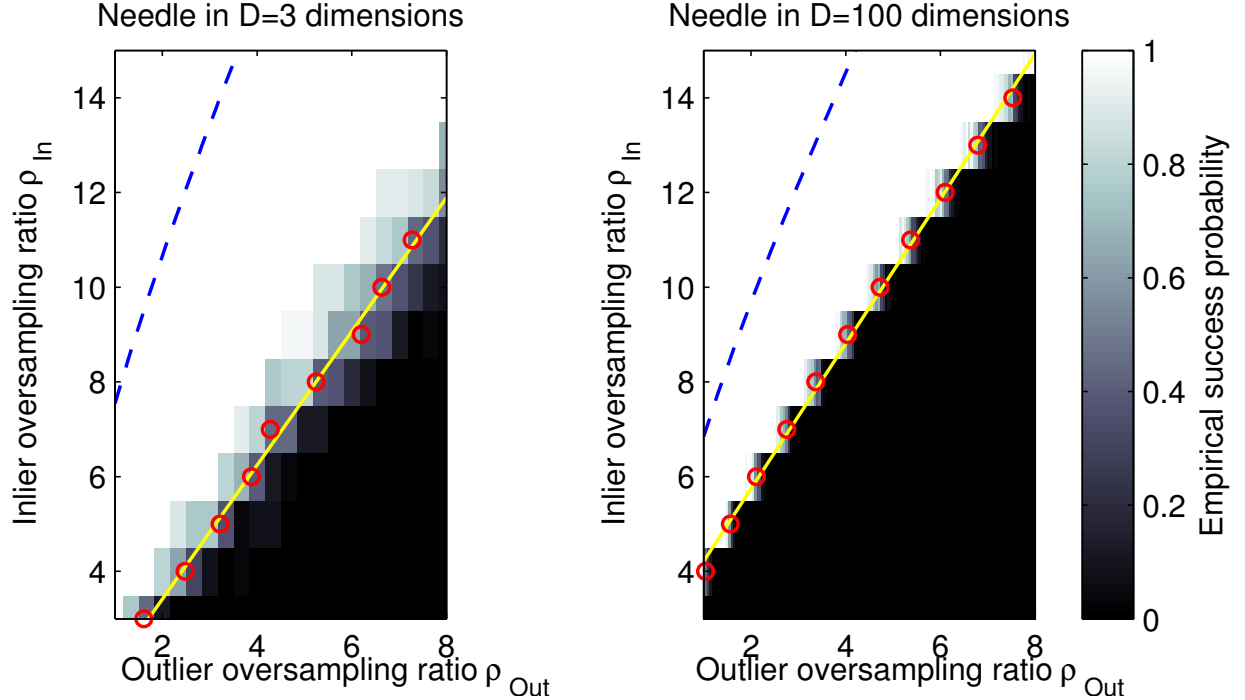


FIGURE 5.1. *Exact recovery of a needle in a haystack via S-REAPER.* The data is drawn from the Haystack Model on page 6 with subspace dimension  $d = 1$ . The ambient dimension  $D = 3$  [left] and  $D = 100$  [right]. The blue curves show the theoretical 50% success threshold given by (B.1c). The red circles mark the empirical 50% success threshold for each value of  $\rho_{in}$ , and the yellow lines indicate the least-squares fit. The trends are  $\rho_{in} = 1.41\rho_{out} + 0.59$  [left] and  $\rho_{in} = 1.53\rho_{out} + 2.68$  [right]. See Section 5.3 for details.

$\rho_{\text{in}} = N_{\text{in}}/d$  and the outlier sampling ratio  $\rho_{\text{out}} = N_{\text{out}}/D$  are relevant to exact recovery. Indeed, we expect to see a linear relationship between these two parameters.

In these experiments, we fix the ambient dimension  $D$  and the dimension  $d$  of the subspace  $L$ . The number  $N_{\text{in}}$  of inliers and the number  $N_{\text{out}}$  of outliers vary. Note that the specific choice of the subspace  $L$  is immaterial because the model is rotationally invariant. The variance parameters can be equated ( $\sigma_{\text{in}}^2 = \sigma_{\text{out}}^2 = 1$ ) because they do not play any role in the performance of S-REAPER. We want to determine when the solution  $\mathbf{P}_\star$  to the S-REAPER problem equals the projector  $\mathbf{\Pi}_L$  onto the model subspace, so we declare the experiment a success when the spectral-norm error  $\|\mathbf{P}_\star - \mathbf{\Pi}_L\| < 10^{-5}$ . For each pair  $(\rho_{\text{in}}, \rho_{\text{out}})$ , we repeat the experiment 25 times and calculate an empirical success probability.

Figure 5.1 charts the performance of S-REAPER for recovering a one-dimensional subspace in ambient dimension  $D = 3$  and  $D = 100$ . For each value of  $\rho_{\text{in}}$ , we identify the minimum empirical value of  $\rho_{\text{out}}$  where the success probability is 50%, and we use least-squares to fit a linear trend (with  $\rho_{\text{in}}$  the independent variable). This experiment shows that the linear relationship between  $\rho_{\text{in}}$  and  $\rho_{\text{out}}$  is valid for the needle problem. The theoretical bound in Theorem B.1 overestimates the empirical trend by approximately a factor of two.

We repeat the same experiment for larger values of the subspace dimension  $d$ . Figure 5.2 shows the results when  $(d, D)$  equals  $(2, 3)$  or  $(25, 100)$ . Once again, we see a linear relationship between  $\rho_{\text{in}}$  and  $\rho_{\text{out}}$ . Our theoretical bound in these cases is less satisfactory; Theorem B.1 overestimates the trend by approximately a factor of eight.

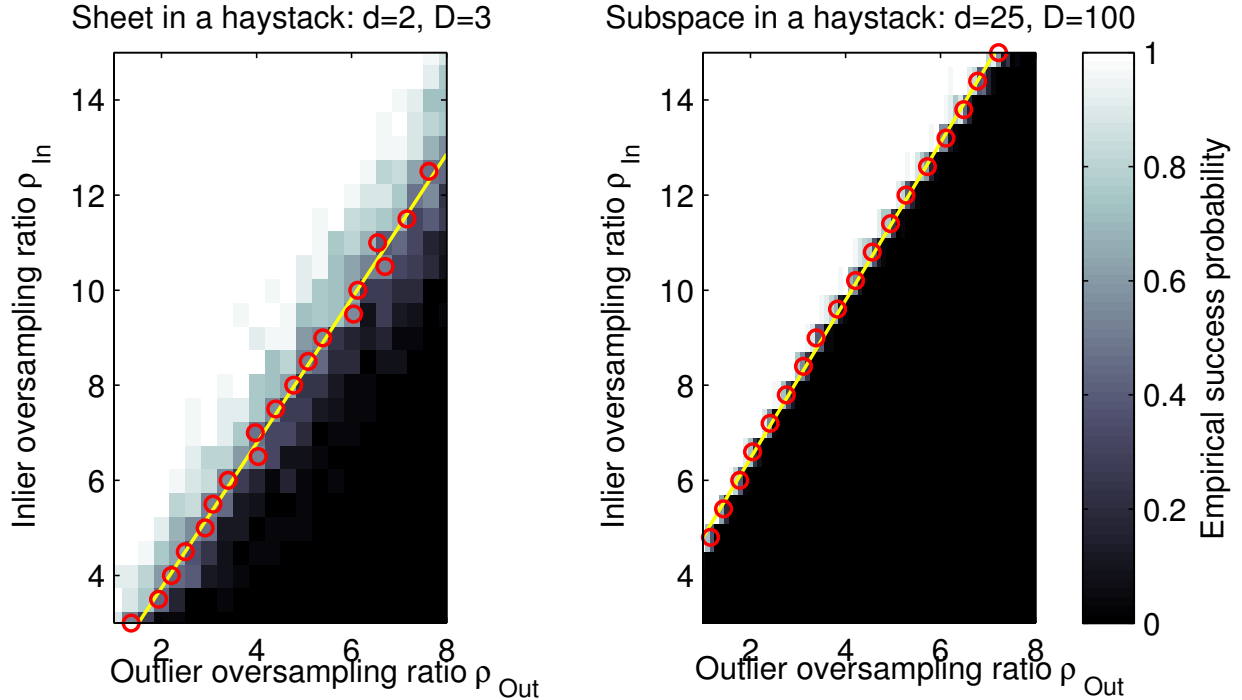


FIGURE 5.2. *Exact recovery of a subspace in a haystack via S-REAPER.* The data is drawn from the Haystack Model on page 6 with  $(d, D) = (2, 3)$  [left] and  $(d, D) = (25, 100)$  [right]. The red circles mark the empirical 50% success threshold for each value of  $\rho_{\text{in}}$ , and the yellow lines indicate the least-squares fit. The trends are  $\rho_{\text{in}} = 1.52\rho_{\text{out}} + 0.70$  [left] and  $\rho_{\text{in}} = 1.66\rho_{\text{out}} + 3.14$  [right]. See Section 5.3 for details.

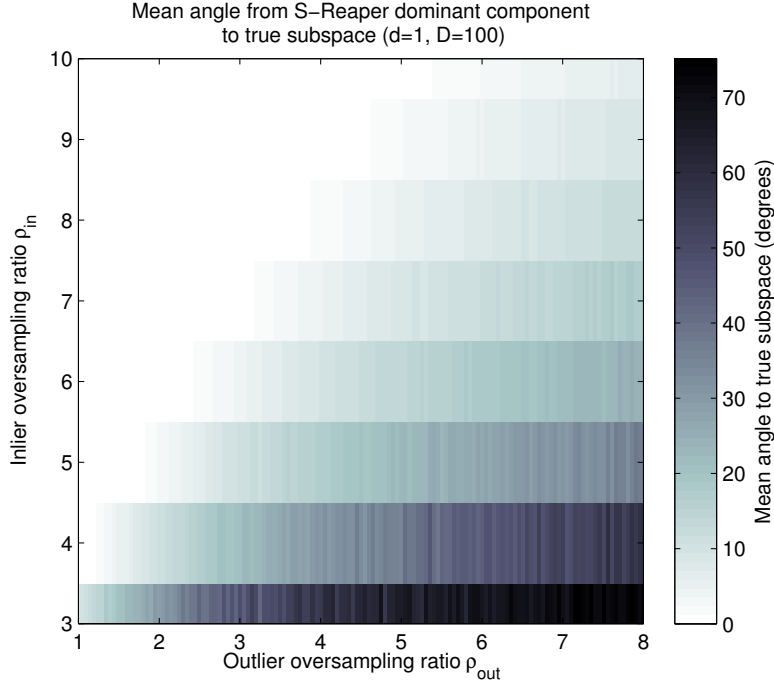


FIGURE 5.3. *Mean angle between the needle and the computed model.* The data is drawn from the Haystack Model on page 6 with  $(d, D) = (1, 100)$ . The image shows the empirical mean, over 25 trials, of the angle between the true subspace and a dominant eigenvector of a solution to the S-REAPER problem (1.6). See Section 5.3.1 for details.

5.3.1. *Graceful Degradation after Recovery Threshold.* The phase transition plots in Figure 5.1 and 5.2 are misleading because S-REAPER continues to work well, even when it does not recover the true subspace exactly. A different view of the data from the phase transition experiments in Section 5.3 captures how our procedure breaks down. Figure 5.3 shows the empirical mean angle between the true subspace  $L$  and the dominant eigenvector of the solution  $\mathbf{P}_\star$  to the S-REAPER problem. The angle increases gradually as the outlier sampling ratio increases. We have also found that the solution of S-REAPER degrades gracefully when the dimension  $d$  of the subspace is larger.

5.4. **Faces in a Crowd.** The next experiment is designed to test how well several robust methods are able to fit a linear model to face images that are dispersed in a collection of random images. This setup allows us to study how well the robust model generalizes to faces we have not seen.

We draw 64 images of a single face under different illuminations from the Extended Yale Face Database [LHK05]. We use the first 32 faces for the sample, and we reserve the other 32 for the out-of-sample test. Next, we add 400 additional random images from the BACKGROUND/Google folder of the Caltech101 database [Cal06, FFFP04]. Each image is converted to grayscale and downsampled to  $20 \times 20$  pixels. We center the images by subtracting the Euclidean median (1.5). Then we apply PCA, spherical PCA, LLD, REAPER, and S-REAPER to fit a nine-dimensional subspace to the data. See [BJ03] for justification of the choice  $d = 9$ . This experiment is similar to work reported in [LLY<sup>+</sup>10, Sec. VI].

Figure 5.4 displays several images from the sample projected onto the computed nine-dimensional subspace (with the centering added back after projection). For every method, the projection of an in-sample face image onto the subspace is recognizable as a face. Meanwhile, the out-of-sample

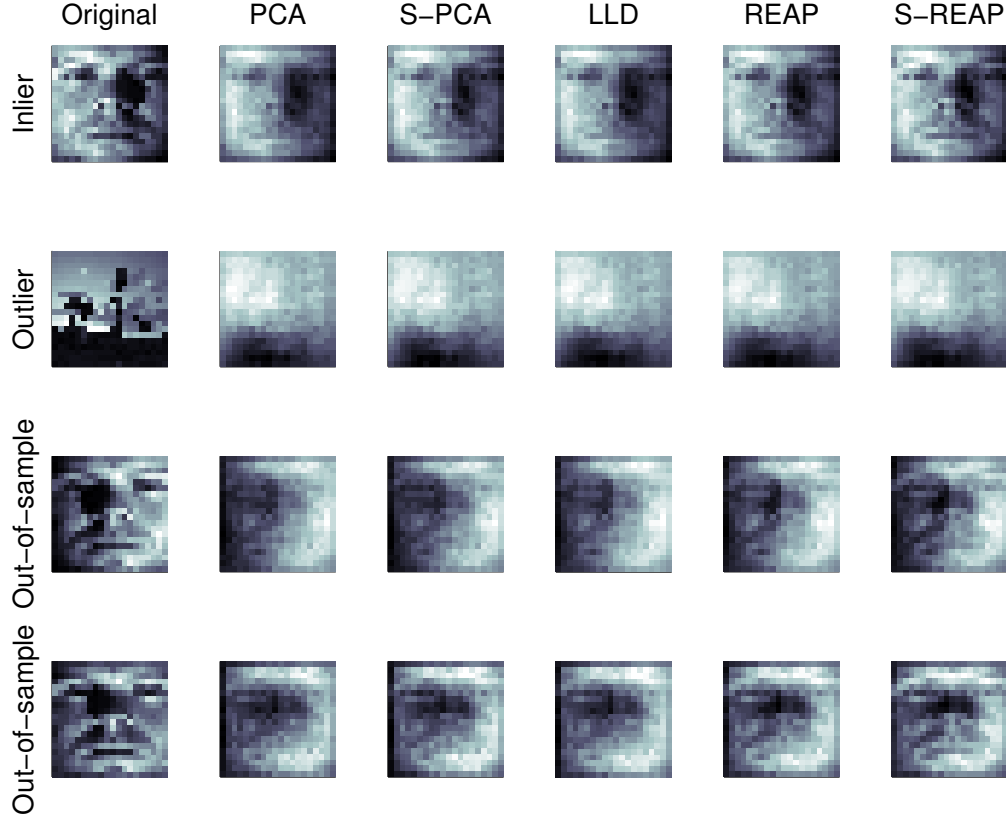


FIGURE 5.4. *Face images projected onto nine-dimensional linear model.* The dataset consists of 32 images of a single face under different illuminations and 400 random images from the Caltech101 database. The original images [left column] are projected onto the nine-dimensional subspaces computed using five different modeling techniques. The first two rows indicate how well the models explain the in-sample faces versus the random images. The last two rows show projections of two out-of-sample faces, which were not used to compute the linear models. See Section 5.4 for details.

faces are described poorly by the PCA subspace. All of the robust subspaces capture the facial features better, with S-REAPER producing the clearest images.

Figure 5.5 shows the ordered distances of the 32 out-of-sample faces to the robust linear model as a function of the ordered distances to the model computed with PCA. A point below the 1:1 line means that the  $i$ th closest point is closer to the robust model than the  $i$ th closest point is to the PCA model. Under this metric, S-REAPER is the dominant method, which explains the qualitative behavior seen in Figure 5.4. This plot clearly demonstrates that S-REAPER computes a subspace that generalizes better than the subspaces obtained with the other robust methods. Indeed, LLD performs only slightly better than PCA.

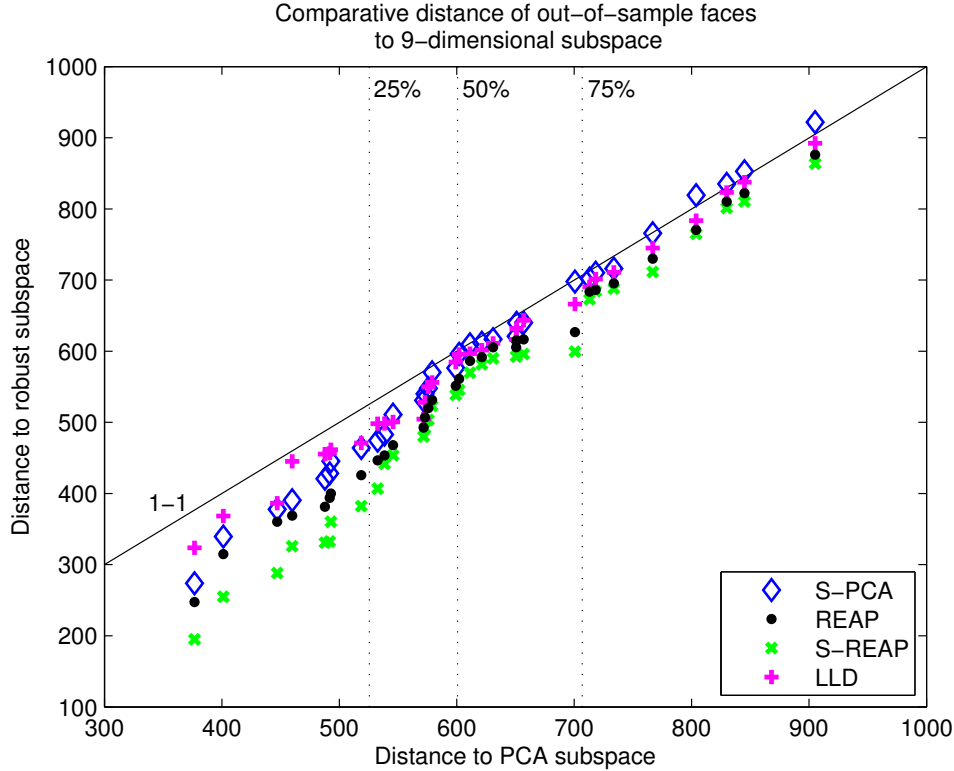


FIGURE 5.5. *Approximation of out-of-sample face images by several linear models.* The ordered distances of the out-of-sample face images to each robust model as a function of the ordered distances to the PCA model. The model was computed from 32 in-sample images; this graph shows how the model generalizes to the 32 out-of-sample images. Lower is better. See Section 5.4 for details.

**5.5. Low-Dimensional Modeling of Textures.** Our next experiment involves an image of a structured texture (a pegboard [SIP12, File 1.5.02]) that has been occluded with an unstructured texture (sand [SIP12, File 1.5.04]). Shifts of a structured texture can be explained well with a low-dimensional model, while unstructured textures look like (anisotropic) noise. This type of model might be interesting for inpainting applications.

See Figure 5.6 for the images we use. The occlusion covers over 80% of the original image, and it is scaled so that the original image is faint in comparison with the occlusion. We block the occluded image into  $24 \times 24$  patches, and we vectorize the resulting 441 patches to obtain observations in a 576-dimensional space. We center the data about the Euclidean median (1.5) of the patches. We compute a three-dimensional linear model for the patches in the occluded image using PCA, spherical PCA, LLD, REAPER, and S-REAPER.

Figure 5.7 shows the reconstruction of some patches from the original, unoccluded image. It appears that S-REAPER and REAPER reproduce the structured texture better than PCA or spherical PCA. Figure 5.8 shows a plot of the ordered distances of patches in the *original* image to the robust linear model versus the ordered distances to the model computed by PCA. Note that the models are all computed from the occluded image.

The plot confirms the qualitative behavior seen in Figure 5.7: the regression planes determined by REAPER and S-REAPER capture the energy of the patches in the original image better than PCA and spherical PCA. Here, the advantage of spherization coupled with REAPER is clear: the performance of S-REAPER is significantly better than either spherical PCA or REAPER.

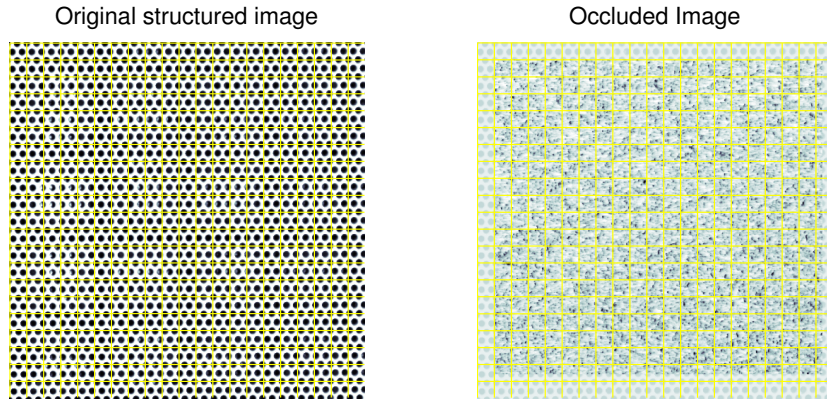


FIGURE 5.6. *Original image and occluded image.* The original [left] and occluded [right] images used in this experiment. The yellow gridlines demarcate the boundaries of the  $24 \times 24$  patches.

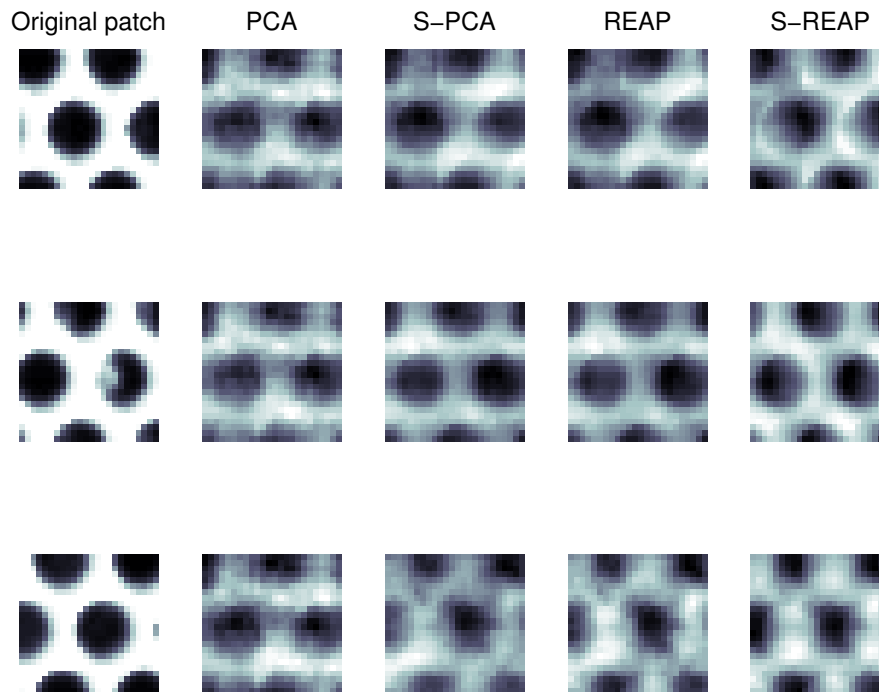


FIGURE 5.7. *Reconstruction of some patches from the original texture.* The left column shows three patches from the original texture that are completely hidden in the occluded image. The next four columns show the projections of these patches onto the three-dimensional subspace computed from the occluded image. We omit the reconstruction obtained with LLD because it is substantially the same as the result with PCA. See Section 5.5 for details.



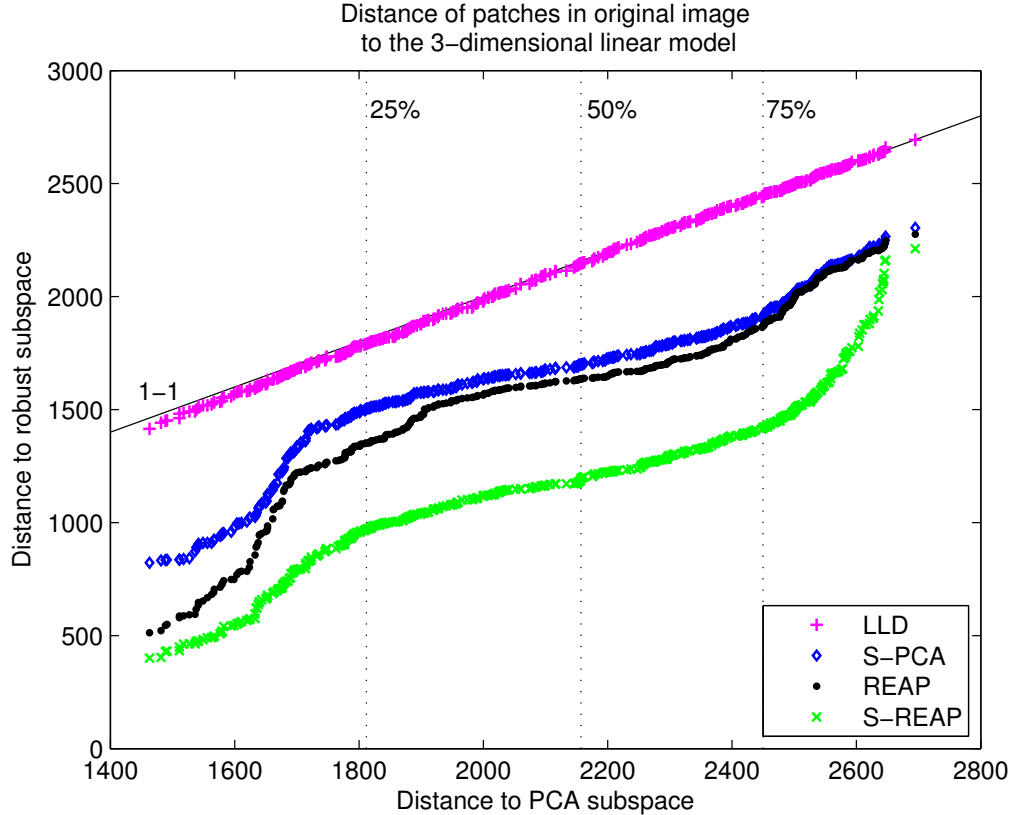


FIGURE 5.8. *Approximation of original image patches by several linear models.* We use four robust methods to determine a three-dimensional linear model for the patches in the occluded image, Figure 5.6. The plot shows the ordered distances of the patches in the original image to each robust models as a function of the ordered distances to the PCA model. Lower is better. See Section 5.5 for details.

## 6. EXTENSIONS AND MODIFICATIONS

There are many opportunities for research on robust linear models. For example, it would also be valuable to extend the analysis in this work to more sophisticated data models and to consider more realistic experimental setups. There are also many variants of the REAPER problem that may be effective, and we believe that our formulations could be extended to more challenging problems. In this section, we survey some of the possibilities for future work.

**6.1. Variations on the REAPER Problem.** The REAPER formulation (1.4) is based on a relaxation of the  $\ell_1$  orthogonal distance problem (1.3). First, we examine what happens if we try to strengthen or weaken the constraints in REAPER. Next, we discuss some alternatives for the  $\ell_1$  objective function that we have applied to the residuals in both (1.3) and (1.4). Third, we show that it is possible to incorporate a centering step directly into the optimization. Last, we describe several methods for constructing a linear model from the solution to the REAPER problem.

**6.1.1. A Stronger Relaxation?** The constraint set in (1.3) consists of all trace- $d$  orthoprojectors, where  $d$  is a positive integer. In the REAPER problem, we relax this constraint to include all trace- $d$  matrices whose eigenvalues fall in  $[0, 1]$ . One may wonder whether a tighter relaxation is possible. If we restrict our attention to convex sets, the answer is negative.



**Fact 6.1.** For each integer  $d \in [0, D]$ , the set  $\{\mathbf{P} \in \mathbb{R}^{D \times D} : \mathbf{0} \preceq \mathbf{P} \preceq \mathbf{I} \text{ and } \text{tr } \mathbf{P} = d\}$  is the convex hull of the  $D \times D$  orthoprojectors with trace  $d$ .

To prove this fact, it suffices to apply a diagonalization argument and to check that the set  $\{\boldsymbol{\lambda} \in \mathbb{R}^D : \sum_{i=1}^D \lambda_i = d \text{ and } 0 \leq \lambda_i \leq 1\}$  is the convex hull of the set of vectors that have  $d$  ones and  $D - d$  zeros. See [OW92] for a discussion of this result.

**6.1.2. A Weaker Relaxation?** Another question is whether a weaker relaxation of REAPER might be effective. As it turns out, only one of the semidefinite constraints is really necessary. Consider the modified optimization problem

$$\text{minimize } \sum_{\mathbf{x} \in \mathcal{X}} \|\mathbf{x} - \mathbf{P}\mathbf{x}\| \quad \text{subject to } \mathbf{0} \preceq \mathbf{P} \quad \text{and} \quad \text{tr } \mathbf{P} = d. \quad (6.1)$$

This problem may be substantially easier to solve in practice with general-purpose software such as CVX [GB08, GB10]. The modification (6.1) is equivalent to (1.4) in the following sense.

**Theorem 6.2** (Weak REAPER). *The problem (6.1) has an optimal point  $\mathbf{P}_\star$  that satisfies  $\mathbf{P}_\star \preceq \mathbf{I}$ .*

We provide a constructive proof of Theorem 6.2 in Appendix D that converts any optimal point of (6.1) into an optimal point that satisfies the extra semidefinite constraint.

**6.1.3. Deadzone Formulation.** In the next two sections, we consider some ways to adapt REAPER that preserve convexity and yet may improve robustness.

First, we consider a modified objective function that ignores small residuals entirely so that the large residuals are penalized even more severely.

$$\text{minimize } \sum_{\mathbf{x} \in \mathcal{X}} \max\{0, \|\mathbf{x} - \mathbf{P}\mathbf{x}\| - \delta\} \quad \text{subject to } \mathbf{0} \preceq \mathbf{P} \preceq \mathbf{I} \quad \text{and} \quad \text{tr } \mathbf{P} = d. \quad (6.2)$$

In other words, no point contributes to the objective unless the model explains it poorly. The positive parameter  $\delta$  controls the width of the *deadzone*. This approach is similar in spirit to a soft-margin classifier [HTF09]. One conceptual disadvantage of the deadzone formulation (6.2) is that it probably does not admit an exact recovery guarantee.

**6.1.4. Outlier Identification and Model Fitting.** In practice, a two-step approach can lead to substantially better linear models for data. First, we use a *robust* method to find a subspace that fits the data. Next, we identify and discard points that are poorly explained by the model. Finally, we fit a subspace to the reduced data, either using a robust method or a classical method. See [RL87] for an exposition of this idea.

We can amplify this approach into an iterative reweighted  $\ell_1$  method [CWB08, KXAH10]. Set the iteration counter  $k \leftarrow 0$ , and set  $\beta_{\mathbf{x}} \leftarrow 1$  for each  $\mathbf{x} \in \mathcal{X}$ . At each step, we solve a weighted version of the REAPER problem:

$$\text{minimize } \sum_{\mathbf{x} \in \mathcal{X}} \beta_{\mathbf{x}} \|\mathbf{x} - \mathbf{P}\mathbf{x}\| \quad \text{subject to } \mathbf{0} \preceq \mathbf{P} \preceq \mathbf{I} \quad \text{and} \quad \text{tr } \mathbf{P} = d. \quad (6.3)$$

Then we update the weights according to some rule. For example,

$$\beta_{\mathbf{x}} \leftarrow \frac{1}{\max\{\delta, \|\mathbf{x} - \mathbf{P}^{(k)}\mathbf{x}\|\}} \quad \text{for each } \mathbf{x} \in \mathcal{X}.$$

By canceling off the magnitude of the residual in (6.3), we hope to minimize the *number* of points that are poorly explained by the model. We repeat this procedure until we do not see any additional improvement. In a small experiment, we found that a single reweighting step already provided most of the benefit.

6.1.5. *Incorporating Centering into REAPER.* In Section 1.4.1, we discussed the importance of centering data before fitting a subspace. In this paragraph, we consider a direct method for finding an *affine space* that explains the data instead of using a separate centering step. There is a natural convex formulation for fitting an affine model to data.

$$\text{minimize} \quad \sum_{\mathbf{x} \in \mathcal{X}} \|\mathbf{x} - \mathbf{P}\mathbf{x} - \mathbf{c}\| \quad \text{subject to} \quad \mathbf{0} \preceq \mathbf{P} \preceq \mathbf{I} \quad \text{and} \quad \text{tr } \mathbf{P} = d. \quad (6.4)$$

The decision variables here are the center  $\mathbf{c}$  and the symmetric matrix  $\mathbf{P}$ . A variant of Algorithm 4.2 can be used to solve this problem efficiently. We have obtained deterministic recovery guarantees for (6.4) when the inliers are contained inside an affine space. For brevity, we omit the details.

The main shortcoming of (6.4) is that we do not see a way to *spherize* the centered data within the framework of convex optimization. We have found that S-REAPER is often more powerful than REAPER, so we believe it is better to center and spherize the data *before* fitting a model.

6.1.6. *Rounding the Solution.* For real data, it is common that the solution  $\mathbf{P}_\star$  to the REAPER problem has rank that exceeds  $d$ . In our numerical work, we select a  $d$ -dimensional model by forming the span of the  $d$  dominant eigenvectors of  $\mathbf{P}_\star$ . This approach is effective in practice, but it hardly seems canonical. It is likely that there are other good ways to achieve the same end.

Let us mention three alternative ideas. First, we might take the *entire* range of  $\mathbf{P}_\star$  as our model. It may have more dimensions than we specified, but it will surely provide a better fit than the lower-dimensional subspace.

Second, we can adjust the parameter in the REAPER problem until the solution has rank  $d$ . A quick way to accomplish this is to apply bisection on the interval  $[0, d]$ . This method also works well in practice.

Third, we can randomly round the matrix  $\mathbf{P}_\star$  to a subspace with expected dimension  $d$ . To do so, let  $(\lambda_i, \mathbf{u}_i)$  be the eigenpairs of  $\mathbf{P}_\star$ . Draw Bernoulli variables  $\delta_i$  with success probability  $\lambda_i$ , and form  $\text{span}\{\mathbf{u}_i : \delta_i = 1\}$ . This approach is inspired by work on randomized rounding of linear programming solutions to combinatorial problems.

6.2. **Extensions of the Analysis.** The theoretical analysis in this paper is designed to provide an explanation for the numerical experiments that we have conducted. Nevertheless, there are many ways in which the simple results here could be improved.

6.2.1. *Refining the Data Models.* The In & Out Model on page 10 is designed to capture the essential attributes of a dataset where the inliers are restricted to a subspace while the outliers can appear anywhere in the ambient space. The permeance statistic and the linear structure statistic capture the basic properties of the data that allow us to establish exact recovery. It would be interesting to identify geometric properties of the data that allow us to control these statistics.

The Haystack Model on page 6 is a very special case of the In & Out Model, where the inliers are isotropic normal variables supported on a subspace and the outliers are isotropic normal variables on the ambient space. It would be more realistic to allow anisotropic data distributions and to consider extensions beyond the Gaussian case.

Another worthwhile direction is to study the performance of REAPER and S-REAPER when the inliers are contaminated with noise. In this case, we cannot hope for exact recovery guarantees, but experiments indicate that the optimization methods are still effective at finding linear models that are comparable with the performance of oracle PCA on the inliers.

6.2.2. *Connections with Other Algorithms.* Recht [Rec12] has observed that the IRLS algorithm is very similar to an EM algorithm for the normal mixture model. There are also evident parallels with the multiplicative weight framework [Kal07] that is studied by the online algorithms community. It would be interesting to understand whether there are deeper connections among these methods.

**6.3. Harder Problems.** A very large number of applications require that we find a linear model for a noisy dataset. It is also common that we need to fit a more complicated model that is built out of simpler linear structures. The methods in this work are also relevant in these situations.

**6.3.1. Manifold Learning.** A common component in methods for manifold learning is a local PCA step (see, e.g., [LV07]). The idea is to look at small neighborhoods in the data and to fit a linear model to each neighborhood. The algorithms then stitch these pieces together to form a manifold. In this application, it is common for the data to contain outliers because the manifold usually has some curvature and the neighborhood may also contain points from different parts of the manifold. Robust methods for linear modeling could be very helpful in this setting.

**6.3.2. Hybrid Linear Modeling.** There are many problems where the data lies on a union of subspaces, rather than a single subspace. The techniques in this paper are very relevant for this problem because, from the point of view of each subspace, the observations from the other subspaces are outliers. Suppose that  $\mathcal{X}$  is a set of observations, and we would like to fit  $m$  subspaces with dimensions  $d_1, \dots, d_m$ . The following optimization offers one way to approach this problem.

$$\begin{aligned} \text{minimize} \quad & \sum_{\mathbf{x} \in \mathcal{X}} \sum_{i=1}^m \|\mathbf{x} - \mathbf{P}_i \mathbf{x}\| \quad \text{subject to} \quad \mathbf{0} \preceq \mathbf{P}_i \preceq \mathbf{I}, \quad \text{tr } \mathbf{P}_i = d_i, \quad \text{and} \\ & \left\| [\mathbf{P}_1 \quad \mathbf{P}_2 \quad \dots \quad \mathbf{P}_m] \right\| \leq 1 + \alpha. \end{aligned} \quad (6.5)$$

We need the third constraint to prevent the decision variables  $\mathbf{P}_i$  from coalescing. The parameter  $\alpha \geq 0$  reflects how much overlap among the subspaces we are willing to tolerate.

It would be very interesting to study the empirical performance of this approach and to develop theory that explains when it can reliably fit hybrid linear models to data. The optimization problem (6.5) is inspired by the formulations in [ZSL09, LZ11].

**6.3.3. Generalized Linear Models.** In this paper, we have focused on a robust method for fitting a standard linear model. Generalized linear models (GLMs) also play an important role in statistical practice. Unfortunately, the specific methods in this paper do not apply directly to GLMs. To make the extension, we need to replace the squared Euclidean loss with a Bregman divergence that is derived from the exponential family of the GLM [BMDG05]. In this setting, Bregman projectors play the role of orthogonal projectors. It would be interesting to develop a method for parameterizing the Bregman projector onto a subspace as a function of the subspace. Then, we might study how to relax this formulation to obtain a convex optimization problem.

**6.4. Conclusions.** We have found that the REAPER and S-REAPER formulations provide a rigorous and effective method for fitting low-dimensional linear models to data. In future work, we hope to improve our theoretical understanding of these methods, and we plan to extend them to related problems. We hope that practitioners will try out these techniques to enhance data analysis in science and engineering applications.

## APPENDIX A. EXACT RECOVERY CONDITIONS FOR THE IN &amp; OUT MODEL

In this appendix, we establish Theorem 3.1, which provides recovery conditions for the deterministic In & Out Model. The proof is an easy consequence of the following sharper result, which we prove in Section A.1 below.

**Theorem A.1.** *Let  $L$  be a  $d$ -dimensional subspace of  $\mathbb{R}^D$ , and let  $\mathcal{X}$  be a dataset that conforms to the In & Out Model on page 10. Define a matrix  $\mathbf{X}_{\text{out}}$  whose columns are the points in  $\mathcal{X}_{\text{out}}$ , arranged in fixed order. Suppose that*

$$\inf_{\substack{\mathbf{u} \in L \\ \|\mathbf{u}\|=1}} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} |\langle \mathbf{u}, \mathbf{x} \rangle| > \sqrt{2d} \cdot \max \left\{ \left\| \widetilde{(\Pi_{L^\perp} \mathbf{X}_{\text{out}})} (\Pi_{L^\perp} \mathbf{X}_{\text{out}})^\mathbf{t} \right\|, \left\| \widetilde{(\Pi_{L^\perp} \mathbf{X}_{\text{out}})} (\Pi_L \mathbf{X}_{\text{out}})^\mathbf{t} \right\| \right\}. \quad (\text{A.1})$$

Then  $\Pi_L$  is the unique minimizer of the REAPER problem (1.4).

*Proof of Theorem 3.1 from Theorem A.1.* First, we check that the sufficient condition (3.5) for REAPER follows from (A.1). The permeance statistic (3.1) equals the left-hand side of (A.1):

$$\mathcal{P}(L) := \inf_{\substack{\mathbf{u} \in L \\ \|\mathbf{u}\|=1}} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} |\langle \mathbf{u}, \mathbf{x} \rangle|.$$

We can bound the right-hand side of (A.1) using the linear structure statistics (3.3) and (3.4).

$$\left\| \widetilde{(\Pi_{L^\perp} \mathbf{X}_{\text{out}})} (\Pi_{L^\perp} \mathbf{X}_{\text{out}})^\mathbf{t} \right\| \leq \left\| \widetilde{\Pi_{L^\perp} \mathbf{X}_{\text{out}}} \right\| \|\mathbf{X}_{\text{out}}\| = \tilde{\mathcal{S}}(L^\perp) \cdot \mathcal{S}(\mathbb{R}^D).$$

The inequality follows because the spectral norm is submultiplicative and  $\|\Pi_{L^\perp}\| = 1$ . Similarly,

$$\left\| \widetilde{(\Pi_{L^\perp} \mathbf{X}_{\text{out}})} (\Pi_L \mathbf{X}_{\text{out}})^\mathbf{t} \right\| \leq \tilde{\mathcal{S}}(L^\perp) \cdot \mathcal{S}(\mathbb{R}^D).$$

Introduce the latter three displays into (A.1) to conclude that (3.5) is a sufficient condition for  $\Pi_L$  to be the unique minimizer of REAPER.

To obtain the sufficient condition for S-REAPER, we need to apply the sufficient condition for REAPER to the spherized data. In this case, the scale-invariant permeance statistic (3.2) matches the left-hand side of (A.1):

$$\tilde{\mathcal{P}}(L) = \inf_{\substack{\mathbf{u} \in L \\ \|\mathbf{u}\|=1}} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} |\langle \mathbf{u}, \tilde{\mathbf{x}} \rangle|.$$

On the right-hand side of (A.1), we replace each instance of  $\mathbf{X}_{\text{out}}$  with the spherized matrix  $\widetilde{\mathbf{X}_{\text{out}}}$ . To complete this step, observe that two separate normalizations are redundant:

$$\widetilde{\Pi_{L^\perp} \widetilde{\mathbf{X}_{\text{out}}}} = \widetilde{\Pi_{L^\perp} \mathbf{X}_{\text{out}}}.$$

Therefore,

$$\left\| \widetilde{(\Pi_{L^\perp} \widetilde{\mathbf{X}_{\text{out}}})} (\Pi_{L^\perp} \widetilde{\mathbf{X}_{\text{out}}})^\mathbf{t} \right\| \leq \left\| \widetilde{\Pi_{L^\perp} \mathbf{X}_{\text{out}}} \right\| \cdot \left\| \widetilde{\mathbf{X}_{\text{out}}} \right\| = \tilde{\mathcal{S}}(L^\perp) \cdot \tilde{\mathcal{S}}(\mathbb{R}^D).$$

Likewise,

$$\left\| \widetilde{(\Pi_{L^\perp} \widetilde{\mathbf{X}_{\text{out}}})} (\Pi_L \widetilde{\mathbf{X}_{\text{out}}})^\mathbf{t} \right\| \leq \tilde{\mathcal{S}}(L^\perp) \cdot \tilde{\mathcal{S}}(\mathbb{R}^D).$$

Introduce the latter three displays into (A.1) to conclude that (3.6) is a sufficient condition for  $\Pi_L$  to be the unique minimizer of S-REAPER.  $\square$

**A.1. Proof of Theorem A.1.** Throughout this section, we use the notation from the In & Out Model on page 10. In particular, recall that  $L$  is a  $d$ -dimensional subspace of  $\mathbb{R}^D$ , and let  $\mathcal{X} = \mathcal{X}_{\text{in}} \cup \mathcal{X}_{\text{out}}$  be the set formed from a set  $\mathcal{X}_{\text{in}}$  of inliers located in  $L$  and a set  $\mathcal{X}_{\text{out}}$  of outliers located in  $\mathbb{R}^D \setminus L$ .

The proof of Theorem A.1 is based on a perturbative argument. Observe that  $\mathbf{P} = \mathbf{\Pi}_L + \mathbf{\Delta}$  is feasible for the REAPER problem (1.4) if and only if the perturbation  $\mathbf{\Delta}$  satisfies

$$\mathbf{0} \preceq \mathbf{\Pi}_L + \mathbf{\Delta} \preceq \mathbf{I} \quad \text{and} \quad \text{tr } \mathbf{\Delta} = 0. \quad (\text{A.2})$$

We demonstrate that, under the condition (A.1), the objective of REAPER evaluated at  $\mathbf{\Pi}_L + \mathbf{\Delta}$  is strictly larger than the objective evaluated at  $\mathbf{\Pi}_L$ .

To establish Theorem A.1, we proceed through a series of (very) technical results. The first lemma identifies a variational inequality that is sufficient for the overall result to hold. The proof appears below in Section A.2.

**Lemma A.2** (Sufficient Conditions for Optimality). *Decompose the matrix  $\mathbf{\Delta}$  into blocks:*

$$\mathbf{\Delta} = \underbrace{\mathbf{\Pi}_L \mathbf{\Delta} \mathbf{\Pi}_L}_{=:\mathbf{\Delta}_1} + \underbrace{\mathbf{\Pi}_{L^\perp} \mathbf{\Delta} \mathbf{\Pi}_L}_{=:\mathbf{\Delta}_2} + \underbrace{\mathbf{\Pi}_L \mathbf{\Delta} \mathbf{\Pi}_{L^\perp}}_{=\mathbf{\Delta}_2^\text{t}} + \underbrace{\mathbf{\Pi}_{L^\perp} \mathbf{\Delta} \mathbf{\Pi}_{L^\perp}}_{=:\mathbf{\Delta}_3} \quad (\text{A.3})$$

Suppose that, for all  $\mathbf{\Delta} \neq \mathbf{0}$  that satisfy (A.2), we have the conditions

$$\sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} \|\mathbf{\Delta}_1 \mathbf{x}\| > \sqrt{2} \sum_{\mathbf{x} \in \mathcal{X}_{\text{out}}} \left\langle \mathbf{\Delta}_3, \left( \widetilde{\mathbf{\Pi}_{L^\perp} \mathbf{x}} \right) (\mathbf{\Pi}_{L^\perp} \mathbf{x})^\text{t} \right\rangle, \quad \text{and} \quad (\text{A.4})$$

$$\sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} \|\mathbf{\Delta}_2 \mathbf{x}\| \geq \sqrt{2} \sum_{\mathbf{x} \in \mathcal{X}_{\text{out}}} \left\langle \mathbf{\Delta}_2, \left( \widetilde{\mathbf{\Pi}_{L^\perp} \mathbf{x}} \right) (\mathbf{\Pi}_L \mathbf{x})^\text{t} \right\rangle. \quad (\text{A.5})$$

Then  $\mathbf{\Pi}_L$  is the unique minimizer of the REAPER problem (1.4).

We must check that the two conditions in Lemma A.2 hold under the hypothesis (A.1) of Theorem A.1. The proof of this result appears below in Section A.5.

**Lemma A.3** (Checking the Conditions). *The inequality (A.4) holds when*

$$\frac{1}{\sqrt{d}} \cdot \inf_{\substack{\mathbf{u} \in L \\ \|\mathbf{u}\|=1}} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} |\langle \mathbf{u}, \mathbf{x} \rangle| > \sqrt{2} \left\| \left( \widetilde{\mathbf{\Pi}_{L^\perp} \mathbf{X}_{\text{out}}} \right) (\mathbf{\Pi}_{L^\perp} \mathbf{X}_{\text{out}})^\text{t} \right\|. \quad (\text{A.6})$$

The inequality (A.5) holds when

$$\frac{1}{\sqrt{d}} \cdot \inf_{\substack{\mathbf{u} \in L \\ \|\mathbf{u}\|=1}} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} |\langle \mathbf{u}, \mathbf{x} \rangle| \geq \sqrt{2} \left\| \left( \widetilde{\mathbf{\Pi}_{L^\perp} \mathbf{X}_{\text{out}}} \right) (\mathbf{\Pi}_L \mathbf{X}_{\text{out}})^\text{t} \right\|. \quad (\text{A.7})$$

We may now complete the proof of Theorem A.1.

*Proof of Theorem A.1 from Lemmas A.2 and A.3.* The hypothesis (A.1) implies the two conditions in Lemma A.3. Therefore, the two conditions in Lemma A.2 are in force, which means that  $\mathbf{\Pi}_L$  is the unique minimizer of REAPER.  $\square$

**A.2. Proof of Lemma A.2.** The REAPER problem (1.4) is convex, so it suffices to check that the objective increases for every *sufficiently small* nonzero perturbation  $\mathbf{\Delta}$  that satisfies the feasibility condition (A.2). The difference between the perturbed and the unperturbed objective values can be expressed as

$$\eta(\mathbf{\Delta}) := \sum_{\mathbf{x} \in \mathcal{X}} \left[ \|\mathbf{\Pi}_{L^\perp} \mathbf{x} - \mathbf{\Delta} \mathbf{x}\| - \|\mathbf{\Pi}_{L^\perp} \mathbf{x}\| \right] \quad (\text{A.8})$$

because  $\mathbf{\Pi}_{L^\perp} = \mathbf{I} - \mathbf{\Pi}_L$ . Our goal is to show that  $\eta(\mathbf{\Delta}) > 0$  whenever  $\mathbf{\Delta}$  is small and feasible.

We split the expression for  $\eta(\Delta)$  into a sum over inliers and a sum over outliers. Each inlier  $\mathbf{x} \in \mathcal{X}_{\text{in}}$  falls in the subspace  $L$ , so that  $\Pi_{L^\perp} \mathbf{x} = \mathbf{0}$ . Therefore,

$$\|(\Pi_{L^\perp} - \Delta)\mathbf{x}\| - \|\Pi_{L^\perp} \mathbf{x}\| = \|\Delta \mathbf{x}\| \quad \text{for } \mathbf{x} \in \mathcal{X}_{\text{in}}. \quad (\text{A.9})$$

To evaluate the summands in (A.8) where  $\mathbf{x} \in \mathcal{X}_{\text{out}}$ , recall that each outlier has a nontrivial projection on the subspace  $L^\perp$ . When  $\Delta \rightarrow \mathbf{0}$ , we compute that

$$\begin{aligned} \|(\Pi_{L^\perp} - \Delta)\mathbf{x}\| &= [\|\Pi_{L^\perp} \mathbf{x}\|^2 - 2\langle \Delta \mathbf{x}, \Pi_{L^\perp} \mathbf{x} \rangle + \|\Delta \mathbf{x}\|^2]^{1/2} \\ &= \|\Pi_{L^\perp} \mathbf{x}\| \left[ 1 - \frac{2}{\|\Pi_{L^\perp} \mathbf{x}\|} \langle \Delta \mathbf{x}, \widetilde{\Pi_{L^\perp} \mathbf{x}} \rangle + \mathcal{O}(\|\Delta\|^2) \right]^{1/2} \\ &= \|\Pi_{L^\perp} \mathbf{x}\| - \langle \Delta \mathbf{x}, \widetilde{\Pi_{L^\perp} \mathbf{x}} \rangle + \mathcal{O}(\|\Delta\|^2) \quad \text{for } \mathbf{x} \in \mathcal{X}_{\text{out}}. \end{aligned}$$

The last identity follows from the Taylor development of the square root. In summary, as  $\Delta \rightarrow \mathbf{0}$ ,

$$\|(\Pi_{L^\perp} - \Delta)\mathbf{x}\| - \|\Pi_{L^\perp} \mathbf{x}\| = -\langle \Delta \mathbf{x}, \widetilde{\Pi_{L^\perp} \mathbf{x}} \rangle + \mathcal{O}(\|\Delta\|^2) \quad \text{for } \mathbf{x} \in \mathcal{X}_{\text{out}}. \quad (\text{A.10})$$

Introduce the inequality (A.9) for the inlier terms and the inequality (A.10) for the outlier terms into the definition (A.8) to obtain

$$\eta(\Delta) = \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} \|\Delta \mathbf{x}\| - \sum_{\mathbf{x} \in \mathcal{X}_{\text{out}}} \langle \Delta \mathbf{x}, \widetilde{\Pi_{L^\perp} \mathbf{x}} \rangle + \mathcal{O}(\|\Delta\|^2) \quad \text{when } \Delta \rightarrow \mathbf{0}.$$

Therefore,  $\eta(\Delta) > 0$  for all nonzero feasible perturbations provided that

$$\sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} \|\Delta \mathbf{x}\| > \sum_{\mathbf{x} \in \mathcal{X}_{\text{out}}} \langle \Delta \mathbf{x}, \widetilde{\Pi_{L^\perp} \mathbf{x}} \rangle \quad \text{for each } \Delta \neq \mathbf{0} \text{ that satisfies (A.2)}. \quad (\text{A.11})$$

In fact, the same argument demonstrates that (A.11) becomes a *necessary* condition for  $\Pi_L$  to minimize (1.4) as soon as we replace the strong inequality by a weak inequality.

To continue, we use the block decomposition (A.3) to break the sufficient condition (A.11) into two pieces that we can check separately. First, note that  $\Pi_{L^\perp} \Delta = \Delta_2 \Pi_L + \Delta_3 \Pi_{L^\perp}$ , so the right-hand side of (A.11) splits up as

$$\begin{aligned} \sum_{\mathbf{x} \in \mathcal{X}_{\text{out}}} \langle \Delta \mathbf{x}, \widetilde{\Pi_{L^\perp} \mathbf{x}} \rangle &= \sum_{\mathbf{x} \in \mathcal{X}_{\text{out}}} \langle \Delta_2 \Pi_L \mathbf{x}, \widetilde{\Pi_{L^\perp} \mathbf{x}} \rangle + \sum_{\mathbf{x} \in \mathcal{X}_{\text{out}}} \langle \Delta_3 \Pi_{L^\perp} \mathbf{x}, \widetilde{\Pi_{L^\perp} \mathbf{x}} \rangle \\ &= \sum_{\mathbf{x} \in \mathcal{X}_{\text{out}}} \langle \Delta_2, (\widetilde{\Pi_{L^\perp} \mathbf{x}})(\Pi_L \mathbf{x})^\top \rangle + \sum_{\mathbf{x} \in \mathcal{X}_{\text{out}}} \langle \Delta_3, (\widetilde{\Pi_{L^\perp} \mathbf{x}})(\Pi_{L^\perp} \mathbf{x})^\top \rangle. \end{aligned} \quad (\text{A.12})$$

The second identity holds because  $\langle \mathbf{A} \mathbf{b}, \mathbf{c} \rangle = \langle \mathbf{A}, \mathbf{c} \mathbf{b}^\top \rangle$ . Next, recall that each inlier falls in  $L$ , and use orthogonality to check that

$$\|\Delta \mathbf{x}\|^2 = \|(\Pi_L + \Pi_{L^\perp}) \Delta \Pi_L \mathbf{x}\|^2 = \|\Delta_1 \mathbf{x}\|^2 + \|\Delta_2 \mathbf{x}\|^2 \quad \text{for } \mathbf{x} \in \mathcal{X}_{\text{in}}.$$

Therefore, we can bound the left-hand side of (A.11) below.

$$\sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} \|\Delta \mathbf{x}\| = \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} [\|\Delta_1 \mathbf{x}\|^2 + \|\Delta_2 \mathbf{x}\|^2]^{1/2} \geq \frac{1}{\sqrt{2}} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} [\|\Delta_1 \mathbf{x}\| + \|\Delta_2 \mathbf{x}\|] \quad (\text{A.13})$$

Together, the assumptions (A.4) and (A.5) imply that the right-hand side of (A.13) exceeds the right-hand side of (A.12). Therefore, the perturbation bound (A.11) holds.

**A.3. Properties of a Feasible Perturbation.** To simplify the conditions in Lemma A.2, we need to collect some information about the properties of a feasible perturbation.

**Sublemma A.4.** *Suppose that  $\Delta$  satisfies the condition (A.2) for a feasible perturbation. Then the following properties hold.*

- (i)  $\text{tr } \Delta_1 + \text{tr } \Delta_3 = 0$ ,
- (ii)  $\Delta_1 \preceq \mathbf{0}$ , and
- (iii)  $\Delta_3 \succeq \mathbf{0}$ .

*In addition, assume that  $\Delta_2 \neq \mathbf{0}$ . Then*

- (iv)  $\Delta_1 \neq \mathbf{0}$ ,
- (v)  $\Delta_3 \neq \mathbf{0}$ , and
- (vi)  $\text{tr } \Delta_3 = -\text{tr } \Delta_1 > 0$ .

*In particular,  $\Delta_1 = \mathbf{0}$  or  $\Delta_3 = \mathbf{0}$  implies that  $\Delta_2 = \mathbf{0}$ .*

*Proof.* Assume that  $\Delta$  satisfies the condition (A.2) for a feasible perturbation. To prove Claim (i), apply the cyclicity of the trace to see that

$$\text{tr } \Delta_2 = \text{tr}(\Pi_{L^\perp} \Delta \Pi_L) = \text{tr}((\Pi_L \Pi_{L^\perp}) \Delta) = \text{tr}(\mathbf{0}) = 0.$$

Take the trace of the block decomposition (A.3) to verify that  $\text{tr } \Delta_1 + \text{tr } \Delta_3 = \text{tr } \Delta = 0$ .

To check (ii), conjugate the relation  $\mathbf{0} \preceq \Pi_L + \Delta$  by the orthoprojector  $\Pi_{L^\perp}$  to see that

$$\mathbf{0} \preceq \Pi_{L^\perp}(\Pi_L + \Delta)\Pi_{L^\perp} = \Delta_3.$$

Similarly, Claim (iii) follows when we conjugate  $\Pi_L + \Delta \preceq \mathbf{I}$  by the orthoprojector  $\Pi_L$ .

We verify (iv) and (v) by establishing the contrapositive. If  $\Delta_1 = \mathbf{0}$  or  $\Delta_3 = \mathbf{0}$ , then  $\Delta_2 = \mathbf{0}$ . First, we argue that  $\Delta_1$  and  $\Delta_3$  are zero together. Indeed, assume that  $\Delta_1 = \mathbf{0}$ . Claim (i) implies that  $\text{tr}(\Delta_3) = 0$ . But then (iii) forces  $\Delta_3 = \mathbf{0}$ . Likewise,  $\Delta_3 = \mathbf{0}$  implies that  $\Delta_1 = \mathbf{0}$ .

Now, suppose  $\Delta_1 = \Delta_3 = \mathbf{0}$ . According to the Schur complement lemma [And05, Thm. 5.3],

$$\mathbf{0} \preceq \Pi_L + \Delta \iff \Delta_1 \succeq \mathbf{0} \quad \text{and} \quad \Delta_3 \succeq \Delta_2(\Pi_L + \Delta_1)^\dagger \Delta_2^\dagger,$$

where the dagger  $^\dagger$  denotes the Moore–Penrose pseudoinverse. Since  $\Delta$  is feasible, the left-hand side of the implication is true, and so

$$\Delta_3 \succeq \Delta_2 \Pi_L^\dagger \Delta_2^\dagger = \Delta_2 \Pi_L \Delta_2^\dagger = \Delta_2 \Delta_2^\dagger.$$

Therefore,  $\Delta_2 = \mathbf{0}$  because  $\Delta_3 = \mathbf{0}$ .

Finally, Claim (vi) holds because (iii) and (v) together imply that  $\text{tr}(\Delta_3) > 0$ . Invoke (i) to complete the proof.  $\square$

**A.4. Lower Bounds for the Sum over Inliers.** The next result allows us to obtain lower bounds for the left-hand sides of the conditions in Lemma A.2. The statement involves the Schatten 1-norm  $\|\cdot\|_{S_1}$ , which returns the sum of the singular values of a matrix [Bha97, p. 92].

**Sublemma A.5.** *The following relations hold for every matrix  $\Delta$ .*

$$\inf_{\Pi_{L^\perp} \Delta \Pi_L \|_{S_1} = 1} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} \|\Pi_{L^\perp} \Delta \Pi_L \mathbf{x}\| \geq \inf_{\substack{\text{tr}(\Pi_L \Delta \Pi_L) = 1 \\ \Pi_L \Delta \Pi_L \succeq \mathbf{0}}} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} \|\Pi_L \Delta \Pi_L \mathbf{x}\| \quad (\text{A.14})$$

$$\geq \frac{1}{\sqrt{d}} \cdot \inf_{\substack{\mathbf{u} \in L \\ \|\mathbf{u}\|=1}} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} |\langle \mathbf{u}, \mathbf{x} \rangle|. \quad (\text{A.15})$$

*Proof.* To verify the relation (A.14), we show that each feasible point  $\Delta$  for the left-hand infimum can be converted into a feasible point  $\tilde{\Delta}$  for the right-hand infimum without changing the objective value. Let  $\Delta$  be an arbitrary matrix, and fix a singular value decomposition  $\Pi_{L^\perp} \Delta \Pi_L = U \Sigma V^\dagger$ .

Define the positive semidefinite matrix  $\mathring{\Delta} = \mathbf{V}\Sigma\mathbf{V}^t$ , and observe that  $\mathring{\Delta} = \mathbf{\Pi}_L\mathring{\Delta}\mathbf{\Pi}_L$  because  $\Sigma\mathbf{V}^t\mathbf{\Pi}_L = \Sigma\mathbf{V}^t$ . First, the unitary invariance of the trace shows that

$$\|\mathbf{\Pi}_{L^\perp}\mathbf{\Delta}\mathbf{\Pi}_L\|_{S_1} = \text{tr}\Sigma = \text{tr}\mathring{\Delta} = \text{tr}(\mathbf{\Pi}_L\mathring{\Delta}\mathbf{\Pi}_L).$$

Next, let  $\mathbf{x}$  be an arbitrary vector, and calculate that

$$\|\mathbf{\Pi}_{L^\perp}\mathbf{\Delta}\mathbf{\Pi}_L\mathbf{x}\| = \|\Sigma\mathbf{V}^t\mathbf{x}\| = \|\mathring{\Delta}\mathbf{x}\| = \|\mathbf{\Pi}_L\mathring{\Delta}\mathbf{\Pi}_L\mathbf{x}\|.$$

Therefore, inequality (A.14) is valid.

The second relation (A.15) is more difficult. The basic idea is to replace the trace constraint in the left-hand infimum with a Frobenius-norm constraint. Subject to this new constraint, the infimum has a simpler form. First, factor out the Frobenius norm  $\|\mathbf{\Pi}_L\mathbf{\Delta}\mathbf{\Pi}_L\|_F$  from the left-hand side of (A.15):

$$\begin{aligned} \inf_{\substack{\text{tr}(\mathbf{\Pi}_L\mathbf{\Delta}\mathbf{\Pi}_L)=1 \\ \mathbf{\Pi}_L\mathbf{\Delta}\mathbf{\Pi}_L \succcurlyeq \mathbf{0}}} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} \|\mathbf{\Pi}_L\mathbf{\Delta}\mathbf{\Pi}_L\mathbf{x}\| \\ \geq \inf_{\substack{\text{tr}(\mathbf{\Pi}_L\mathbf{\Delta}\mathbf{\Pi}_L)=1 \\ \mathbf{\Pi}_L\mathbf{\Delta}\mathbf{\Pi}_L \succcurlyeq \mathbf{0}}} \|\mathbf{\Pi}_L\mathbf{\Delta}\mathbf{\Pi}_L\|_F \times \inf_{\|\mathbf{\Pi}_L\mathbf{\Delta}\mathbf{\Pi}_L\|_F=1} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} \|\mathbf{\Pi}_L\mathbf{\Delta}\mathbf{\Pi}_L\mathbf{x}\|. \end{aligned} \quad (\text{A.16})$$

To compute the first infimum on the right-hand side of (A.16), note that the rank of  $\mathbf{\Pi}_L\mathbf{\Delta}\mathbf{\Pi}_L$  cannot exceed  $d$  because  $L$  is  $d$ -dimensional subspace. Therefore, by diagonalization,

$$\inf_{\substack{\text{tr}(\mathbf{\Pi}_L\mathbf{\Delta}\mathbf{\Pi}_L)=1 \\ \mathbf{\Pi}_L\mathbf{\Delta}\mathbf{\Pi}_L \succcurlyeq \mathbf{0}}} \|\mathbf{\Pi}_L\mathbf{\Delta}\mathbf{\Pi}_L\|_F = \inf \left\{ \|\boldsymbol{\lambda}\| : \|\boldsymbol{\lambda}\|_{\ell_1} = 1, \boldsymbol{\lambda} \in \mathbb{R}_+^d \right\} = \frac{1}{\sqrt{d}}. \quad (\text{A.17})$$

To compute the second infimum, we use a more careful diagonalization argument. We can parameterize the feasible set using the spectral decomposition  $\mathbf{\Pi}_L\mathbf{\Delta}\mathbf{\Pi}_L = \sum_{i=1}^d \lambda_i \mathbf{u}_i \mathbf{u}_i^t$ , where the eigenvalues  $\lambda_i$  are nonnegative and the eigenvectors  $\mathbf{u}_i$  form an orthobasis for  $L$ . Thus,

$$\inf_{\|\mathbf{\Pi}_L\mathbf{\Delta}\mathbf{\Pi}_L\|_F=1} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} \|\mathbf{\Pi}_L\mathbf{\Delta}\mathbf{\Pi}_L\mathbf{x}\| = \inf_{\text{orthobases } \{\mathbf{u}_j\} \text{ for } L} \inf_{\|\boldsymbol{\lambda}\|=1} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} \left[ \sum_{i=1}^d \lambda_i^2 |\langle \mathbf{u}_i, \mathbf{x} \rangle|^2 \right]^{1/2}.$$

This expression can be simplified dramatically. Make the identification  $p_i = \lambda_i^2$  for each  $i = 1, \dots, d$ . Observe that the objective on the right-hand side is a *concave* function of the vector  $\mathbf{p}$ . Furthermore,  $\mathbf{p}$  is constrained to the convex set of nonnegative vectors that sum to one (i.e., probability densities). The minimum of a concave function on a compact set is achieved at an extreme point, so the infimum occurs when  $\mathbf{p}$  is a standard basis vector. Therefore,

$$\inf_{\|\boldsymbol{\lambda}\|=1} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} \left[ \sum_{i=1}^d \lambda_i^2 |\langle \mathbf{u}_i, \mathbf{x} \rangle|^2 \right]^{1/2} = \min_{i=1, \dots, d} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} |\langle \mathbf{u}_i, \mathbf{x} \rangle|.$$

Combining the last two relations, we obtain

$$\inf_{\|\mathbf{\Pi}_L\mathbf{\Delta}\mathbf{\Pi}_L\|_F=1} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} \|\mathbf{\Pi}_L\mathbf{\Delta}\mathbf{\Pi}_L\mathbf{x}\| = \inf_{\substack{\mathbf{u} \in L \\ \|\mathbf{u}\|=1}} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} |\langle \mathbf{u}, \mathbf{x} \rangle|. \quad (\text{A.18})$$

To complete the proof of (A.15), we introduce (A.17) and (A.18) into (A.16).  $\square$

**A.5. Proof of Lemma A.3.** We are now prepared to check that the conditions in Lemma A.3 imply the conditions in Lemma A.2. Throughout the argument, we assume that  $\mathbf{\Delta}$  is a nonzero perturbation that satisfies the feasibility condition (A.2), so we have access to the results in Sublemma A.4.

First, observe that the inequality (A.4) is homogeneous in  $\mathbf{\Delta}$ . In view of Sublemma A.4, we can scale  $\mathbf{\Delta}$  so that  $\text{tr}(\mathbf{\Delta}_3) = -\text{tr}(\mathbf{\Delta}_1) = 1$ . Furthermore, we may assume that  $\mathbf{\Delta}_1$  is negative semidefinite and  $\mathbf{\Delta}_3$  is positive semidefinite. Therefore, it suffices to check the constrained infimum



of the left-hand side of (A.4) exceeds the constrained supremum of the right-hand side. It follows immediately from Sublemma A.5 that

$$\inf_{\substack{\text{tr } \Delta_1 = -1 \\ \Delta_1 \preceq \mathbf{0}}} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} \|\Delta_1 \mathbf{x}\| \geq \frac{1}{\sqrt{d}} \cdot \inf_{\substack{\mathbf{u} \in L \\ \|\mathbf{u}\|=1}} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} |\langle \mathbf{u}, \mathbf{x} \rangle|.$$

Meanwhile, the duality between the Schatten 1-norm and the spectral norm gives

$$\begin{aligned} \sqrt{2} \sup_{\substack{\text{tr } \Delta_3 = 1 \\ \Delta_3 \succeq \mathbf{0}}} \sum_{\mathbf{x} \in \mathcal{X}_{\text{out}}} \langle \Delta_3, (\widetilde{\Pi_{L^\perp} \mathbf{x}})(\Pi_{L^\perp} \mathbf{x})^\mathbf{t} \rangle &= \sqrt{2} \left\| \sum_{\mathbf{x} \in \mathcal{X}_{\text{out}}} (\widetilde{\Pi_{L^\perp} \mathbf{x}})(\Pi_{L^\perp} \mathbf{x})^\mathbf{t} \right\| \\ &= \sqrt{2} \left\| (\widetilde{\Pi_{L^\perp} \mathbf{X}_{\text{out}}})(\Pi_{L^\perp} \mathbf{X}_{\text{out}})^\mathbf{t} \right\|. \end{aligned}$$

The first equality holds because the matrix inside the spectral norm happens to be positive semi-definite. For the second identity, we simply identify the sum as the product of the two matrices. Combine the last two displays to see that the condition (A.4) follows from (A.6).

The second part of the proof is similar. Observe that the inequality (A.5) is homogeneous in  $\Delta_2$ , so we may assume that  $\|\Delta_2\|_{S_1} = 1$ . Sublemma A.5 delivers the bound

$$\inf_{\|\Delta_2\|_{S_1}=1} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} \|\Delta_1 \mathbf{x}\| \geq \frac{1}{\sqrt{d}} \cdot \inf_{\substack{\mathbf{u} \in L \\ \|\mathbf{u}\|=1}} \sum_{\mathbf{x} \in \mathcal{X}_{\text{in}}} |\langle \mathbf{u}, \mathbf{x} \rangle|.$$

The duality argument we used for the first condition now delivers

$$\sqrt{2} \sup_{\|\Delta_2\|_{S_1}=1} \sum_{\mathbf{x} \in \mathcal{X}_{\text{out}}} \langle \Delta_2, (\widetilde{\Pi_{L^\perp} \mathbf{x}})(\Pi_{L^\perp} \mathbf{x})^\mathbf{t} \rangle = \sqrt{2} \left\| (\widetilde{\Pi_{L^\perp} \mathbf{X}_{\text{out}}})(\Pi_{L^\perp} \mathbf{X}_{\text{out}})^\mathbf{t} \right\|.$$

Combine these two displays to see that (A.5) is a consequence of (A.7).

## APPENDIX B. EXACT RECOVERY CONDITIONS FOR THE HAYSTACK MODEL

In this appendix, we establish exact recovery conditions for the Haystack Model. To accomplish this goal, we study the probabilistic behavior of the (spherical) permeance statistic and the (spherical) linear structure statistic. Our main result for the Haystack Model, Theorem B.1, follows when we introduce the probability bounds into the deterministic recovery result, Theorem 3.1. Theorem 1.1, the simplified result for the Haystack Model, is a consequence of the more detailed theory we develop here.

**Theorem B.1.** *Fix a parameter  $c > 0$ . Let  $L$  be an arbitrary  $d$ -dimensional subspace of  $\mathbb{R}^D$ , and draw the dataset  $\mathcal{X}$  at random according to the Haystack Model on page 6.*

(i) *Let  $1 \leq d \leq D - 1$ . Suppose that*

$$\sqrt{\frac{2}{\pi}} \rho_{\text{in}} - (2 + c) \sqrt{\rho_{\text{in}}} \geq \frac{\sigma_{\text{out}}}{\sigma_{\text{in}}} \cdot \sqrt{\frac{2D}{D - d - 0.5}} \cdot \left( \sqrt{\rho_{\text{out}}} + 1 + c \sqrt{\frac{d}{D}} \right)^2. \quad (\text{B.1a})$$

*Then  $\Pi_L$  is the unique minimum of REAPER (1.4), except with probability  $3.5 e^{-c^2 d/2}$ .*

(ii) *Let  $2 \leq d \leq D - 1$ . Suppose that*

$$\sqrt{\frac{2}{\pi}} \rho_{\text{in}} - (2 + c\sqrt{2}) \sqrt{\rho_{\text{in}}} \geq \sqrt{\frac{D}{D - 0.5}} \cdot \sqrt{\frac{2D}{D - d - 0.5}} \cdot \left( \sqrt{\rho_{\text{out}}} + 1 + c \sqrt{\frac{d}{D}} \right)^2. \quad (\text{B.1b})$$

*Then  $\Pi_L$  is the unique minimum of S-REAPER (1.6), except with probability  $4 e^{-c^2 d/2}$ .*

(iii) Let  $d = 1$  and  $D \geq 2$ . Suppose that

$$\rho_{\text{in}} \geq \sqrt{\frac{D}{D-0.5}} \cdot \sqrt{\frac{2D}{D-1.5}} \cdot \left( \sqrt{\rho_{\text{out}}} + 1 + c\sqrt{\frac{1}{D}} \right)^2. \quad (\text{B.1c})$$

Then  $\Pi_L$  is the unique minimum of S-REAPER (1.6), except with probability  $3e^{-c^2/2}$ .

We prove these rather daunting expressions below in Sections B.1 and B.2. First, let us examine the content of Theorem B.1 in a specific regime. Suppose that

$$D \rightarrow \infty, \quad d \rightarrow \infty, \quad \text{and} \quad d/D \rightarrow 0.$$

Then the condition (B.1b) simplifies to

$$\rho_{\text{in}} \geq \sqrt{\pi} \cdot (\sqrt{\rho_{\text{out}}} + 1)^2.$$

This inequality matches our heuristic (3.7) exactly! Now, we demonstrate that Theorem B.1 contains the simplified result for the Haystack model, Theorem 1.1.

*Proof of Theorem 1.1 from Theorem B.1.* To begin, we collect some basic inequalities. For  $\alpha > 0$ , the function  $f(x) = x - \alpha\sqrt{x}$  is convex when  $x \geq 0$ , so that

$$x - \alpha\sqrt{x} = f(x) \geq f(\alpha^2) + f'(\alpha^2)(x - \alpha^2) = \frac{1}{2}(x - \alpha^2).$$

For  $1 \leq d \leq (D-1)/2$ , we have the numerical bounds

$$\frac{D}{D-d-0.5} \leq 2, \quad \frac{D}{D-0.5} \leq 1.2, \quad \text{and} \quad \frac{d}{D} \leq 2.$$

Finally, recall that  $(a+b+c)^2 \leq 3(a^2+b^2+c^2)$  as a consequence of Hölder's inequality.

To prove the claim (1.7) of Theorem 1.1, we apply our numerical inequalities and the fact  $(2+c)^2 \leq 2(4+c^2)$  to weaken (B.1a) to the condition

$$\rho_{\text{in}} - \pi(4+c^2) \geq 12\sqrt{\frac{\pi}{2}} \cdot \frac{\sigma_{\text{out}}}{\sigma_{\text{in}}} \cdot (\rho_{\text{out}} + 1 + 2c^2)$$

Apply Theorem B.1(i) with  $c = \sqrt{2\beta}$  to obtain (1.7) with constants  $C_1 = 4\pi < 13$  and  $C_2 = 2\pi < 7$  and  $C_3 = 12\sqrt{\pi/2} < 16$ .

To prove the claim (1.8) of Theorem 1.1, we use the same reasoning to weaken (B.1b) to the condition

$$\rho_{\text{in}} - \pi(4+2c^2) \geq 12\sqrt{\frac{3\pi}{5}}(\rho_{\text{out}} + 1 + 2c^2)$$

Apply Theorem B.1(ii) with  $c = \sqrt{2\beta}$  to obtain (1.8) in the range  $2 \leq d \leq (D-1)/2$ . The case  $d = 1$  follows from a similar argument using Theorem B.1(iii).  $\square$

**B.1. Probability Inequalities for the Summary Statistics.** The proof of Theorem B.1 requires probability bounds on the permeance statistic  $\mathcal{P}$ , the linear structure statistic  $\mathcal{S}$ , and their spherical counterparts  $\tilde{\mathcal{P}}$  and  $\tilde{\mathcal{S}}$ . These bounds follow from tail inequalities for Gaussian and spherically distributed random vectors that we develop in the next two subsections.

**B.1.1. The Permeance Statistic.** Our first result is used to provide a high-probability lower bound for the permeance statistic  $\mathcal{P}(L)$  under the Haystack Model.

**Lemma B.2.** Suppose  $\mathbf{g}_1, \dots, \mathbf{g}_n$  are i.i.d.  $\text{NORMAL}(\mathbf{0}, \mathbf{I})$  vectors in  $\mathbb{R}^d$ . For each  $t \geq 0$ ,

$$\inf_{\|\mathbf{u}\|=1} \sum_{i=1}^n |\langle \mathbf{u}, \mathbf{g}_i \rangle| > \sqrt{\frac{2}{\pi}} \cdot n - 2\sqrt{nd} - t\sqrt{n}, \quad (\text{B.2})$$

except with probability  $e^{-t^2/2}$ .

*Proof.* Add and subtract the mean from each summand on the left-hand side of (B.2) to obtain

$$\inf_{\|\mathbf{u}\|=1} \sum_{i=1}^n |\langle \mathbf{u}, \mathbf{g}_i \rangle| \geq \inf_{\|\mathbf{u}\|=1} \sum_{i=1}^n [|\langle \mathbf{u}, \mathbf{g}_i \rangle| - \mathbb{E} |\langle \mathbf{u}, \mathbf{g}_i \rangle|] + \inf_{\|\mathbf{u}\|=1} \sum_{i=1}^n \mathbb{E} |\langle \mathbf{u}, \mathbf{g}_i \rangle| \quad (\text{B.3})$$

The second sum on the right-hand side has a closed form expression because each term is the expectation of a half-Gaussian random variable:  $\mathbb{E} |\langle \mathbf{u}, \mathbf{g}_i \rangle| = \sqrt{2/\pi}$  for every unit vector  $\mathbf{u}$ . Therefore,

$$\inf_{\|\mathbf{u}\|=1} \sum_{i=1}^n \mathbb{E} |\langle \mathbf{u}, \mathbf{g}_i \rangle| = \sqrt{\frac{2}{\pi}} \cdot n. \quad (\text{B.4})$$

To control the first sum on the right-hand side of (B.3), we use a standard argument. To bound the mean, we symmetrize the sum and invoke a comparison theorem. To control the probability of a large deviation, we apply a measure concentration argument.

To proceed with the calculation of the mean, we use the Rademacher symmetrization lemma [LT91, Lem. 6.3] to obtain

$$\mathbb{E} \sup_{\|\mathbf{u}\|=1} \sum_{i=1}^n [(\mathbb{E} |\langle \mathbf{u}, \mathbf{g}_i \rangle|) - |\langle \mathbf{u}, \mathbf{g}_i \rangle|] \leq 2 \mathbb{E} \sup_{\|\mathbf{u}\|=1} \sum_{i=1}^n \varepsilon_i |\langle \mathbf{u}, \mathbf{g}_i \rangle|.$$

The random variables  $\varepsilon_1, \dots, \varepsilon_n$  are i.i.d. Rademacher random variables that are independent from the Gaussian sequence. Next, invoke the Rademacher comparison theorem [LT91, Eqn. (4.20)] with the function  $\varphi(\cdot) = |\cdot|$  to obtain the further bound

$$\mathbb{E} \sup_{\|\mathbf{u}\|=1} \sum_{i=1}^n [(\mathbb{E} |\langle \mathbf{u}, \mathbf{g}_i \rangle|) - |\langle \mathbf{u}, \mathbf{g}_i \rangle|] \leq 2 \mathbb{E} \sup_{\|\mathbf{u}\|=1} \sum_{i=1}^n \varepsilon_i \langle \mathbf{u}, \mathbf{g}_i \rangle = 2 \mathbb{E} \left\| \sum_{i=1}^n \varepsilon_i \mathbf{g}_i \right\|.$$

The identity follows when we draw the sum into the inner product and maximize over all unit vectors. From here, the rest of the argument is very easy. Use Jensen's inequality to bound the expectation by the root-mean-square, which has a closed form:

$$\mathbb{E} \sup_{\|\mathbf{u}\|=1} \sum_{i=1}^n [(\mathbb{E} |\langle \mathbf{u}, \mathbf{g}_i \rangle|) - |\langle \mathbf{u}, \mathbf{g}_i \rangle|] \leq 2 \left[ \mathbb{E} \left\| \sum_{i=1}^n \varepsilon_i \mathbf{g}_i \right\|^2 \right]^{1/2} = 2\sqrt{nd}. \quad (\text{B.5})$$

Note that the mean fluctuation (B.5) is dominated by the centering term (B.4) when  $n \gg d$ .

To control the probability that the fluctuation term is large, we use a standard concentration inequality [Bog98, Theorem 1.7.6] for a Lipschitz function of independent Gaussian variables. Define a real-valued function on  $d \times n$  matrices:  $f(\mathbf{Z}) = \sup_{\|\mathbf{u}\|=1} \sum_{i=1}^n (\sqrt{2/\pi} - |\langle \mathbf{u}, \mathbf{z}_i \rangle|)$ , where  $\mathbf{z}_i$  denotes the  $i$ th column of  $\mathbf{Z}$ . Compute that

$$|f(\mathbf{Z}) - f(\mathbf{Z}')| \leq \sup_{\|\mathbf{u}\|=1} \sum_{i=1}^n |\langle \mathbf{u}, \mathbf{z}_i - \mathbf{z}'_i \rangle| \leq \sum_{i=1}^n \|\mathbf{z}_i - \mathbf{z}'_i\| \leq \sqrt{n} \|\mathbf{Z} - \mathbf{Z}'\|_F.$$

Therefore,  $f$  has Lipschitz constant  $\sqrt{n}$  with respect to the Frobenius norm. In view of the estimate (B.5) for the mean, the Gaussian concentration bound implies that

$$\mathbb{P} \left\{ \sup_{\|\mathbf{u}\|=1} \sum_{i=1}^n [(\mathbb{E} |\langle \mathbf{u}, \mathbf{g}_i \rangle|) - |\langle \mathbf{u}, \mathbf{g}_i \rangle|] \geq 2\sqrt{nd} + t\sqrt{n} \right\} \leq e^{-t^2/2}. \quad (\text{B.6})$$

Introduce the bound (B.6) and the identity (B.4) into (B.3) to complete the proof.  $\square$

To control the spherical permeance statistic, we need a variant of Lemma B.2. The following observations play a critical role in this argument and elsewhere. Suppose that  $\tilde{\mathbf{g}}$  is uniformly distributed on the unit sphere in  $\mathbb{R}^d$ , and let  $r$  be a  $\chi$ -distributed random variable with  $d$  degrees

of freedom that is independent from  $\tilde{\mathbf{g}}$ . Then  $r\tilde{\mathbf{g}} \sim \mathbf{g}$ , where  $\mathbf{g}$  is a standard normal vector. To apply this fact, we need some bounds for the mean of a  $\chi$ -distributed variable:

$$\sqrt{d-0.5} \leq \mathbb{E} r \leq \sqrt{d} \quad \text{when } r \sim \chi_d. \quad (\text{B.7})$$

The upper bound follows from Jensen's inequality, while the lower bound is a refinement of Gautschi's inequality due to Kershaw [Ker83].

**Lemma B.3.** *Suppose  $\tilde{\mathbf{g}}_1, \dots, \tilde{\mathbf{g}}_n$  are i.i.d. random vectors uniformly distributed on the unit sphere  $\mathbb{S}^{d-1}$  in  $\mathbb{R}^d$ . When  $d = 1$ ,*

$$\inf_{\|\mathbf{u}\|=1} \sum_{i=1}^n |\langle \mathbf{u}, \tilde{\mathbf{g}}_i \rangle| = n. \quad (\text{B.8})$$

Assume that  $d \geq 2$ . For all  $t \geq 0$ ,

$$\inf_{\|\mathbf{u}\|=1} \sum_{i=1}^n |\langle \mathbf{u}, \tilde{\mathbf{g}}_i \rangle| > \sqrt{\frac{2}{\pi}} \cdot \frac{n}{\sqrt{d}} - 2\sqrt{n} - t\sqrt{\frac{n}{d-1}} \quad (\text{B.9})$$

except with probability  $e^{-t^2/2}$ .

*Proof.* When  $d = 1$ , the argument is particularly simple. A random variable uniformly distributed on the 0-dimensional sphere has a Rademacher distribution. The result (B.8) is immediate.

Assume that  $d \geq 2$ . The proof is analogous to the argument for the Gaussian case, Lemma B.2. Add and subtract the mean from the left-hand side of (B.9) to reach

$$\inf_{\|\mathbf{u}\|=1} \sum_{i=1}^n |\langle \mathbf{u}, \tilde{\mathbf{g}}_i \rangle| \geq \inf_{\|\mathbf{u}\|=1} \sum_{i=1}^n [|\langle \mathbf{u}, \tilde{\mathbf{g}}_i \rangle| - \mathbb{E} |\langle \mathbf{u}, \tilde{\mathbf{g}}_i \rangle|] + \inf_{\|\mathbf{u}\|=1} \sum_{i=1}^n \mathbb{E} |\langle \mathbf{u}, \tilde{\mathbf{g}}_i \rangle|. \quad (\text{B.10})$$

To bound the second term on right-hand side, we introduce independent  $\chi$ -variables  $r_i$  to facilitate comparison with a Gaussian. Indeed,

$$\mathbb{E} |\langle \mathbf{u}, \tilde{\mathbf{g}}_i \rangle| = \frac{1}{\mathbb{E} r_i} \cdot \mathbb{E}_{r_i} \mathbb{E}_{\tilde{\mathbf{g}}_i} |\langle \mathbf{u}, r_i \tilde{\mathbf{g}}_i \rangle| \geq \frac{1}{\mathbb{E} r_i} \cdot \mathbb{E} |\langle \mathbf{u}, \mathbf{g}_i \rangle| \geq \sqrt{\frac{2}{\pi d}}.$$

The last inequality follows from (B.7). Therefore,

$$\inf_{\|\mathbf{u}\|=1} \sum_{i=1}^n \mathbb{E} |\langle \mathbf{u}, \tilde{\mathbf{g}}_i \rangle| \geq \sqrt{\frac{2}{\pi}} \cdot \frac{n}{\sqrt{d}}. \quad (\text{B.11})$$

Next, we bound the mean of the first term on the right-hand of (B.10). By repeating the symmetrization and comparison argument from Lemma B.2, we reach

$$\mathbb{E} \sup_{\|\mathbf{u}\|=1} \sum_{i=1}^n [(\mathbb{E} |\langle \mathbf{u}, \tilde{\mathbf{g}}_i \rangle|) - |\langle \mathbf{u}, \tilde{\mathbf{g}}_i \rangle|] \leq 2 \left[ \mathbb{E} \left\| \sum_{i=1}^n \varepsilon_i \tilde{\mathbf{g}}_i \right\|^2 \right]^{1/2} = 2\sqrt{n}.$$

The identity depends on the fact that each vector  $\tilde{\mathbf{g}}_i$  has unit norm.

To control the probability of a large deviation, we need an appropriate concentration inequality for the function  $f$  defined on page 34. We apply a result for a Lipschitz function on a product of spheres [Led01, Thm. 2.4], with the parameter  $c = d - 1$ , along with the definition [Led01, Prop. 1.11] of the concentration function to reach

$$\mathbb{P} \left\{ \sup_{\|\mathbf{u}\|=1} \sum_{i=1}^n [(\mathbb{E} |\langle \mathbf{u}, \tilde{\mathbf{g}}_i \rangle|) - |\langle \mathbf{u}, \tilde{\mathbf{g}}_i \rangle|] \geq 2\sqrt{n} + t\sqrt{\frac{n}{d-1}} \right\} \leq e^{-t^2/2}. \quad (\text{B.12})$$

Introduce (B.12) and (B.11) into (B.10) to complete the proof of (B.9).  $\square$

**B.1.2. Probability Bounds for the Linear Structure Statistics.** To control the linear structure statistic  $\mathcal{S}$ , we need a tail bound for the maximum singular value of a Gaussian matrix. The following inequality is a well-known consequence of Slepian's lemma. See [DS01, Thm. 2.13] and the errata [DS03] for details.

**Proposition B.4.** *Let  $\mathbf{G}$  be an  $m \times n$  matrix whose entries are i.i.d. standard normal random variables. For each  $t \geq 0$ ,*

$$\mathbb{P} \left\{ \|\mathbf{G}\| \geq \sqrt{m} + \sqrt{n} + t \right\} < 1 - \Phi(t) < e^{-t^2/2},$$

where  $\Phi(t)$  is the Gaussian cumulative density function

$$\Phi(t) := \frac{1}{\sqrt{2\pi}} \int_{-\infty}^t e^{-\tau^2/2} d\tau.$$

To bound the spherical linear structure statistic  $\tilde{\mathcal{S}}$ , we need a related result for random matrices with independent columns that are uniformly distributed on the sphere. The argument bootstraps from Proposition B.4.

**Lemma B.5.** *Let  $\mathbf{S}$  be an  $m \times n$  matrix whose columns are i.i.d. random vectors, uniformly distributed on the sphere  $\mathbb{S}^{m-1}$  in  $\mathbb{R}^m$ . For each  $t \geq 0$ ,*

$$\mathbb{P} \left\{ \|\mathbf{S}\| \geq \frac{\sqrt{n} + \sqrt{m} + t}{\sqrt{m - 0.5}} \right\} \leq 1.5 e^{-t^2/2}. \quad (\text{B.13})$$

*Proof.* Fix  $\theta > 0$ . The Laplace transform method shows that

$$P := \mathbb{P} \left\{ \|\mathbf{S}\| \geq \frac{\sqrt{n} + \sqrt{m} + t}{\sqrt{m - 0.5}} \right\} \leq e^{-\theta(\sqrt{n} + \sqrt{m} + t)} \cdot \mathbb{E} e^{\theta \sqrt{m - 0.5} \|\mathbf{S}\|}.$$

We compare  $\|\mathbf{S}\|$  with the norm of a Gaussian matrix by introducing a diagonal matrix of  $\chi$ -distributed variables. The rest of the argument is purely technical.

Let  $\mathbf{r} = (r_1, \dots, r_n)$  be a vector of i.i.d.  $\chi$ -distributed random variables with  $m$  degrees of freedom. Recall that  $r_i \tilde{\mathbf{g}}_i \sim \mathbf{g}_i$ , where  $\tilde{\mathbf{g}}_i$  is uniform on the sphere and  $\mathbf{g}_i$  is standard normal. Using the bound (B.7), which states that  $\sqrt{m - 0.5} \leq \mathbb{E} r_i$ , and Jensen's inequality, we obtain

$$\mathbb{E} e^{\theta \sqrt{m - 0.5} \|\mathbf{S}\|} \leq \mathbb{E} e^{\theta \|\mathbb{E} \mathbf{r} \text{diag}(\mathbf{r}) \mathbf{S}\|} \leq \mathbb{E} e^{\theta \|\mathbf{G}\|},$$

where  $\mathbf{G}$  is an  $m \times n$  matrix with i.i.d. standard normal entries.

Define a random variable  $Z := \|\mathbf{G}\| - \sqrt{n} - \sqrt{m}$ , and let  $Z_+ := \max\{Z, 0\}$  denote its positive part. Then

$$e^{\theta t} \cdot P \leq \mathbb{E} e^{\theta Z} \leq \mathbb{E} e^{\theta Z_+} = 1 + \int_0^\infty e^{\theta \tau} \cdot \mathbb{P}\{Z_+ > \tau\} d\tau.$$

Apply the cdf bound in Proposition B.4, and identify the complementary error function  $\text{erfc}$ .

$$e^{\theta t} \cdot P \leq 1 + \frac{\theta}{2} \int_0^\infty e^{\theta \tau} \cdot \text{erfc} \left( \frac{\tau}{\sqrt{2}} \right) d\tau,$$

A computer algebra system will report that this frightening integral has a closed form:

$$\theta \int_0^\infty e^{\theta \tau} \cdot \text{erfc} \left( \frac{\tau}{\sqrt{2}} \right) d\tau = e^{\theta^2/2} (\text{erf}(\theta) + 1) - 1 \leq 2 e^{\theta^2/2} - 1.$$

We have used the simple bound  $\text{erf}(\theta) \leq 1$  for  $\theta \geq 0$ . In summary,

$$P \leq e^{-\theta t} \cdot \left[ \frac{1}{2} + e^{\theta^2/2} \right]$$

Select  $\theta = t$  to obtain the advertised bound (B.13). □

**B.2. Proof of Theorem B.1.** Let  $\mathbf{X}_{\text{out}}$  be a  $D \times N_{\text{out}}$  matrix whose columns are the outliers  $\mathbf{x} \in \mathcal{X}_{\text{out}}$ , arranged in fixed order. Recall that the inlier sampling ratio  $\rho_{\text{in}} := N_{\text{in}}/d$  and the outlier sampling ratio  $\rho_{\text{out}} := N_{\text{out}}/D$ .

First, we establish Claim (i). The  $N_{\text{in}}$  inliers are drawn from a centered Gaussian distribution on the  $d$ -dimensional space  $L$  with covariance  $(\sigma_{\text{in}}^2/d) \mathbf{I}_L$ . Rotational invariance and Lemma B.2, with  $t = c\sqrt{d}$ , together imply that the permeance statistic (3.1) satisfies

$$\mathcal{P}(L) > \frac{\sigma_{\text{in}}}{\sqrt{d}} \left[ \sqrt{\frac{2}{\pi}} \cdot N_{\text{in}} - (2+c)\sqrt{N_{\text{in}}d} \right] = \sigma_{\text{in}}\sqrt{d} \left[ \sqrt{\frac{2}{\pi}}\rho_{\text{in}} - (2+c)\sqrt{\rho_{\text{in}}} \right],$$

except with probability  $e^{-c^2d/2}$ . The  $N_{\text{out}}$  outliers are independent, centered Gaussian vectors in  $\mathbb{R}^D$  with covariance  $(\sigma_{\text{out}}^2/D) \mathbf{I}$ . Proposition B.4, with  $t = c\sqrt{d}$ , delivers a bound for the linear structure statistic (3.3):

$$\mathcal{S}(\mathbb{R}^D) = \|\mathbf{X}_{\text{out}}\| \leq \frac{\sigma_{\text{out}}}{\sqrt{D}} [\sqrt{N_{\text{out}}} + \sqrt{D} + c\sqrt{d}] = \sigma_{\text{out}} \left[ \sqrt{\rho_{\text{out}}} + 1 + c\sqrt{\frac{d}{D}} \right],$$

except with probability  $e^{-c^2d/2}$ . Rotational invariance implies that the columns of  $\widetilde{\mathbf{\Pi}_{L^\perp} \mathbf{X}_{\text{out}}}$  are independent vectors that are uniformly distributed on the unit sphere of a  $(D-d)$ -dimensional space. Lemma B.5 yields

$$\tilde{\mathcal{S}}(L^\perp) = \|\widetilde{\mathbf{\Pi}_{L^\perp} \mathbf{X}_{\text{out}}}\| \leq \frac{\sqrt{N_{\text{out}}} + \sqrt{D-d} + c\sqrt{d}}{\sqrt{D-d-0.5}} < \sqrt{\frac{D}{D-d-0.5}} \left[ \sqrt{\rho_{\text{out}}} + 1 + c\sqrt{\frac{d}{D}} \right],$$

except with probability  $1.5e^{-c^2d/2}$ . Under the assumption (B.1a), we conclude that

$$\mathcal{P}(L) > \sqrt{2d} \cdot \tilde{\mathcal{S}}(L^\perp) \cdot \tilde{\mathcal{S}}(\mathbb{R}^D)$$

except with probability  $3.5e^{-c^2d/2}$ . Apply Theorem 3.1 to complete the argument.

To establish Claim (ii), we also need to bound the spherical permeance statistic (3.2) and the spherical linear structure statistic (3.4). Assume that  $d \geq 2$ . The  $N_{\text{in}}$  spherized inliers are uniformly distributed over the unit-sphere in the  $d$ -dimensional space  $L$ . Lemma B.3, with  $t = c\sqrt{d}$ , implies that the spherical permeance statistic satisfies

$$\tilde{\mathcal{P}}(L) > \sqrt{\frac{2}{\pi}} \cdot \frac{N_{\text{in}}}{\sqrt{d}} - 2\sqrt{N_{\text{in}}} - c\sqrt{\frac{N_{\text{in}}d}{d-1}} \geq \sqrt{d} \left[ \sqrt{\frac{2}{\pi}}\rho_{\text{in}} - (2+c\sqrt{2})\sqrt{\rho_{\text{in}}} \right],$$

except with probability  $e^{-c^2d/2}$ . The second inequality follows from the numerical bound  $d/(d-1) \leq 2$  for  $d \geq 2$ . For the outliers, note that the  $N_{\text{out}}$  columns of  $\widetilde{\mathbf{X}_{\text{out}}}$  are independent random vectors, uniformly distributed on the unit sphere in  $\mathbb{R}^D$ , so Lemma B.5 shows that

$$\tilde{\mathcal{S}}(\mathbb{R}^D) = \|\widetilde{\mathbf{X}_{\text{out}}}\| \leq \frac{\sqrt{N_{\text{out}}} + \sqrt{D} + c\sqrt{d}}{\sqrt{D-0.5}} = \sqrt{\frac{D}{D-0.5}} \left[ \sqrt{\rho_{\text{out}}} + 1 + c\sqrt{\frac{d}{D}} \right],$$

except with probability  $1.5e^{-c^2d/2}$ . Under the assumption (B.1b), we conclude that

$$\tilde{\mathcal{P}}(L) > \sqrt{2d} \cdot \tilde{\mathcal{S}}(L^\perp) \cdot \tilde{\mathcal{S}}(\mathbb{R}^D),$$

except with probability  $4e^{-c^2d/2}$ . Theorem 3.1 establishes the result.

Finally, Claim (iii) follows from a refined bound on  $\tilde{\mathcal{P}}(L)$  that holds when  $d = 1$ . In this case, Lemma B.3 yields  $\tilde{\mathcal{P}}(L) = N_{\text{in}}$ . The result follows in the same manner as part (ii).

## APPENDIX C. ANALYSIS OF THE IRLS ALGORITHM

This appendix contains the details of our analysis of the IRLS method, Algorithm 4.2. First, we verify that Algorithm 4.1 reliably solves the weighted least-squares subproblem (4.1). Then, we argue that IRLS converges to a point near the true optimum of the REAPER problem (1.4).

**C.1. Solving the Weighted Least-Squares Problem.** In this section, we verify that Algorithm 4.1 correctly solves the weighted least-squares problem (4.1). The following lemma provides a more mathematical statement of the algorithm, along with the proof of correctness.

**Lemma C.1** (Solving the Weighted Least-Squares Problem). *Assume that  $0 < d < D$ , and suppose that  $\mathcal{X}$  is a set of observations in  $\mathbb{R}^D$ . For each  $\mathbf{x} \in \mathcal{X}$ , let  $\beta_{\mathbf{x}}$  be a nonnegative weight. Form the weighted sample covariance matrix  $\mathbf{C}$ , and compute its eigenvalue decomposition:*

$$\mathbf{C} := \sum_{\mathbf{x} \in \mathcal{X}} \beta_{\mathbf{x}} \mathbf{x} \mathbf{x}^t = \mathbf{U} \mathbf{\Lambda} \mathbf{U}^t \quad \text{where } \lambda_1 \geq \dots \geq \lambda_D \geq 0.$$

When  $\text{rank}(\mathbf{C}) \leq d$ , construct a vector  $\boldsymbol{\nu} \in \mathbb{R}^D$  via the formula

$$\boldsymbol{\nu} := (\underbrace{1, \dots, 1}_{[d] \text{ times}}, d - [d], 0, \dots, 0)^t. \quad (\text{C.1})$$

When  $\text{rank}(\mathbf{C}) > d$ , define the positive quantity  $\theta$  implicitly by solving the equation

$$\sum_{i=1}^D \frac{(\lambda_i - \theta)_+}{\lambda_i} = d. \quad (\text{C.2})$$

Construct a vector  $\boldsymbol{\nu} \in \mathbb{R}^D$  whose components are

$$\nu_i := \frac{(\lambda_i - \theta)_+}{\lambda_i} \quad \text{for } i = 1, \dots, D. \quad (\text{C.3})$$

In either case, an optimal solution to (4.1) is given by

$$\mathbf{P}_{\star} := \mathbf{U} \cdot \text{diag}(\boldsymbol{\nu}) \cdot \mathbf{U}^t. \quad (\text{C.4})$$

In this statement,  $(a)_+ := \max\{0, a\}$ , we enforce the convention  $0/0 := 0$ , and  $\text{diag}$  forms a diagonal matrix from a vector.

*Proof of Lemma C.1.* First, observe that the construction (C.4) yields a matrix  $\mathbf{P}_{\star}$  that satisfies the constraints of (4.1) in both cases.

When  $\text{rank}(\mathbf{C}) \leq d$ , we can verify that our construction of the vector  $\boldsymbol{\nu}$  yields a optimizer of (4.1) by showing that the objective value is zero, which is minimal. Evaluate the objective function (4.1) at the point  $\mathbf{P}_{\star}$  to see that

$$\sum_{\mathbf{x} \in \mathcal{X}} \|(\mathbf{I} - \mathbf{P}_{\star}) \mathbf{x}\|^2 = \text{tr}[(\mathbf{I} - \mathbf{P}_{\star}) \mathbf{C} (\mathbf{I} - \mathbf{P}_{\star})] = \sum_{i=1}^D (1 - \nu_i)^2 \lambda_i \quad (\text{C.5})$$

by definition of  $\mathbf{C}$  and the fact that  $\mathbf{C}$  and  $\mathbf{P}_{\star}$  are simultaneously diagonalizable. The nonzero eigenvalues of  $\mathbf{C}$  appear among  $\lambda_1, \dots, \lambda_{[d]}$ . At the same time,  $1 - \nu_i = 0$  for each  $i = 1, \dots, [d]$ . Therefore, the value of (C.5) equals zero at  $\mathbf{P}_{\star}$ .

Next, assume that  $\text{rank}(\mathbf{C}) > d$ . The objective function in (4.1) is convex, so we can verify that  $\mathbf{P}_{\star}$  solves the optimization problem if the directional derivative of the objective at  $\mathbf{P}_{\star}$  is nonnegative in every feasible direction. A matrix  $\mathbf{\Delta}$  is a feasible perturbation if and only if

$$\mathbf{0} \preceq \mathbf{P}_{\star} + \mathbf{\Delta} \preceq \mathbf{I} \quad \text{and} \quad \text{tr } \mathbf{\Delta} = 0.$$

Let  $\mathbf{\Delta}$  be an arbitrary matrix that satisfies these constraints. By expanding the objective of (4.1) about  $\mathbf{P}_{\star}$ , easily compute the derivative in the direction  $\mathbf{\Delta}$ . Therefore, the condition

$$-\langle \mathbf{\Delta}, (\mathbf{I} - \mathbf{P}_{\star}) \mathbf{C} \rangle \geq 0 \quad (\text{C.6})$$

ensures that the derivative increases in the direction  $\Delta$ .

First, note that the quantity  $\theta$  can be defined. Indeed, the left-hand side of (C.2) equals  $\text{rank}(\mathbf{C})$  when  $\theta = 0$ , and it equals zero when  $\theta \geq \lambda_1$ . By continuity, there exists a value of  $\theta$  that solves the equation. Let  $i_*$  be the largest index where  $\lambda_{i_*} > \theta$ , so that  $\nu_i = 0$  for each  $i > i_*$ . Next, define  $M$  to be the subspace spanned by the eigenvectors  $\mathbf{u}_{i_*+1}, \dots, \mathbf{u}_D$ . Since  $\nu_i$  is the eigenvalue of  $\mathbf{P}_*$  with eigenvector  $\mathbf{u}_i$ , we must have  $\Pi_M \mathbf{P}_* \Pi_M = \mathbf{0}$ . It follows that  $\Pi_M \Delta \Pi_M \succcurlyeq \mathbf{0}$  because  $\Pi_M (\mathbf{P}_* + \Delta) \Pi_M \succcurlyeq \mathbf{0}$ .

To complete the argument, observe that

$$(1 - \nu_i) \lambda_i = \lambda_i - (\lambda_i - \theta)_+ = \min\{\lambda_i, \theta\}.$$

Therefore,  $(\mathbf{I} - \mathbf{P}_*)\mathbf{C} = \mathbf{U} \cdot \text{diag}(\min\{\lambda_i, \theta\}) \cdot \mathbf{U}^t$ . Using the fact that  $\text{tr} \Delta = 0$ , we obtain

$$\begin{aligned} \langle \Delta, (\mathbf{I} - \mathbf{P}_*)\mathbf{C} \rangle &= \langle \Delta, \mathbf{U} \cdot \text{diag}(\min\{\lambda_i, \theta\}) \cdot \mathbf{U}^t \rangle \\ &= \langle \Delta, \underbrace{\mathbf{U} \cdot \text{diag}(0, \dots, 0, \lambda_{i_*+1} - \theta, \dots, \lambda_D - \theta) \cdot \mathbf{U}^t}_{=: \mathbf{Z}} \rangle \end{aligned}$$

Since  $\lambda_i \leq \theta$  for each  $i > i_*$ , each eigenvalue of  $\mathbf{Z}$  is nonpositive. Furthermore,  $\Pi_M \mathbf{Z} \Pi_M = \mathbf{Z}$ . We see that

$$\langle \Delta, (\mathbf{I} - \mathbf{P}_*)\mathbf{C} \rangle = \langle \Delta, \Pi_M \mathbf{Z} \Pi_M \rangle = \langle \Pi_M \Delta \Pi_M, \mathbf{Z} \rangle \leq 0,$$

because the compression of  $\Delta$  on  $M$  is positive semidefinite and  $\mathbf{Z}$  is negative semidefinite. In other words, (C.6) is satisfied for every feasible perturbation  $\Delta$  about  $\mathbf{P}_*$ .  $\square$

**C.2. Convergence of IRLS.** In this section, we argue that the IRLS method of Algorithm 4.2 converges to a point whose value is nearly optimal for the REAPER problem (1.4). The proof consists of two phases. First, we explain how to modify the argument from [CM99] to show that the iterates  $\mathbf{P}^{(k)}$  converge to a matrix  $\mathbf{P}_\delta$ , which is characterized as the solution to a regularized counterpart of REAPER. The fact that the limit point  $\mathbf{P}_\delta$  achieves a near-optimal value for REAPER follows from the characterization.

*Proof sketch for Theorem 4.1.* We find it more convenient to work with the variables  $\mathbf{Q} := \mathbf{I} - \mathbf{P}$  and  $\mathbf{Q}^{(k)} := \mathbf{I} - \mathbf{P}^{(k)}$ . First, let us define a regularized objective. For a parameter  $\delta > 0$ , consider the Huber-like function

$$H_\delta(x, y) = \begin{cases} \frac{1}{2} \left( \frac{x^2}{\delta} + \delta \right), & 0 \leq y \leq \delta \\ \frac{1}{2} \left( \frac{x^2}{y} + y \right), & y \geq \delta. \end{cases}$$

We introduce the convex function

$$\begin{aligned} F(\mathbf{Q}) &:= \sum_{\mathbf{x} \in \mathcal{X}} H_\delta(\|\mathbf{Q}\mathbf{x}\|, \|\mathbf{Q}\mathbf{x}\|) \\ &= \sum_{\{\mathbf{x}: \|\mathbf{Q}\mathbf{x}\| \geq \delta\}} \|\mathbf{Q}\mathbf{x}\| + \frac{1}{2} \sum_{\{\mathbf{x}: \|\mathbf{Q}\mathbf{x}\| < \delta\}} \left( \frac{\|\mathbf{Q}\mathbf{x}\|^2}{\delta} + \delta \right). \end{aligned}$$

The second identity above highlights the interpretation of  $F$  as a regularized objective function for (1.4) under the assignment  $\mathbf{Q} = \mathbf{I} - \mathbf{P}$ . Note that  $F$  is continuously differentiable at each matrix  $\mathbf{Q}$ , and the gradient

$$\nabla F(\mathbf{Q}) = \sum_{\mathbf{x} \in \mathcal{X}} \frac{\mathbf{Q}\mathbf{x}\mathbf{x}^t}{\max\{\|\mathbf{Q}\mathbf{x}\|, \delta\}}.$$

The technical assumption that the observations do not lie in the union of two strict subspaces of  $\mathbb{R}^D$  implies that  $F$  is *strictly* convex; compare with the proof [ZL11, Thm. 1]. We define  $\mathbf{Q}_\delta$  to be the solution of a constrained optimization problem:

$$\mathbf{Q}_\delta := \arg \min_{\substack{\mathbf{0} \preceq \mathbf{Q} \preceq \mathbf{I} \\ \text{tr } \mathbf{Q} = D-d}} F(\mathbf{Q}).$$



The strict convexity of  $F$  implies that  $\mathbf{Q}_\delta$  is well defined.

The key idea in the proof is to show that the iterates  $\mathbf{Q}^{(k)}$  of Algorithm 4.2 converge to the optimizer  $\mathbf{Q}_\delta$  of the regularized objective function  $F$ . We demonstrate that Algorithm 4.2 is a generalized Weiszfeld method in the sense of [CM99, Sec. 4]. After defining some additional auxiliary functions and facts about these functions, we explain how the argument of [CM99, Lem. 5.1] can be adapted to prove that the iterates of  $\mathbf{Q}^{(k)} = \mathbf{I} - \mathbf{P}^{(k)} \rightarrow \mathbf{Q}_\delta$ . The only innovation required is an inequality from convex analysis that lets us handle the constraints  $\mathbf{0} \preceq \mathbf{Q} \preceq \mathbf{I}$  and  $\text{tr } \mathbf{Q} = D - d$ .

Now for the definitions. We introduce the potential function

$$G(\mathbf{Q}, \mathbf{Q}^{(k)}) := \sum_{\mathbf{x} \in \mathcal{X}} H_\delta(\|\mathbf{Q}\mathbf{x}\|, \|\mathbf{Q}^{(k)}\mathbf{x}\|).$$

Then  $G(\cdot, \mathbf{Q}^{(k)})$  is a smooth quadratic function. By collecting terms, we may relate  $G$  and  $F$  through the expansion

$$G(\mathbf{Q}, \mathbf{Q}^{(k)}) = F(\mathbf{Q}^{(k)}) + \langle \mathbf{Q} - \mathbf{Q}^{(k)}, \nabla F(\mathbf{Q}^{(k)}) \rangle + \frac{1}{2} \langle \mathbf{Q} - \mathbf{Q}^{(k)}, C(\mathbf{Q}^{(k)})(\mathbf{Q} - \mathbf{Q}^{(k)}) \rangle,$$

where  $C$  is the continuous function

$$C(\mathbf{Q}^{(k)}) := \sum_{\mathbf{x} \in \mathcal{X}} \frac{\mathbf{x}\mathbf{x}^\text{t}}{\max\{\|\mathbf{Q}\mathbf{x}\|, \delta\}}.$$

Next, we verify some facts related to Hypothesis 4.2 and 4.3 of [CM99, Sec. 4]. Note that  $F(\mathbf{Q}) = G(\mathbf{Q}, \mathbf{Q})$ . Furthermore,  $F(\mathbf{Q}) \leq G(\mathbf{Q}, \mathbf{Q}^{(k)})$  because  $H_\delta(x, x) \leq H_\delta(x, y)$ , which is a direct consequence of the AM–GM inequality.

We now relate the iterates of Algorithm 4.2 to the definitions above. Given that  $\mathbf{Q}^{(k)} = \mathbf{I} - \mathbf{P}^{(k)}$ , Step 2b of Algorithm 4.2 is equivalent to the iteration

$$\mathbf{Q}^{(k+1)} = \arg \min_{\substack{\mathbf{0} \preceq \mathbf{Q} \preceq \mathbf{I} \\ \text{tr } \mathbf{Q} = D - d}} G(\mathbf{Q}, \mathbf{Q}^{(k)}).$$

From this characterization, we have the monotonicity property

$$F(\mathbf{Q}^{(k+1)}) \leq G(\mathbf{Q}^{(k+1)}, \mathbf{Q}^{(k)}) \leq G(\mathbf{Q}^{(k)}, \mathbf{Q}^{(k)}) = F(\mathbf{Q}^{(k)}). \quad (\text{C.7})$$

This fact motivates the stopping criterion for Algorithm 4.2 because it implies the objective values are decreasing:  $\alpha^{(k+1)} = F(\mathbf{Q}^{(k+1)}) \leq F(\mathbf{Q}^{(k)}) = \alpha^{(k)}$ .

We also require some information regarding the bilinear form induced by  $C$ . Introduce the quantity  $m := \max\{\delta, \max_{\mathbf{x} \in \mathcal{X}}\{\|\mathbf{x}\|\}\}$ . Then, by symmetry of the matrix  $\mathbf{Q}$ , and the fact that the inner product between positive semidefinite matrices is nonnegative we have

$$\langle \mathbf{Q}, C(\mathbf{Q}^{(k)})\mathbf{Q} \rangle \geq \frac{1}{m} \text{tr} \left( \mathbf{Q}^2 \sum_{\mathbf{x} \in \mathcal{X}} \mathbf{x}\mathbf{x}^\text{t} \right) \geq \|\mathbf{Q}\|_F^2 \underbrace{\left( \frac{\lambda_{\min}(\sum_{\mathbf{x} \in \mathcal{X}} \mathbf{x}\mathbf{x}^\text{t})}{m} \right)}_{=: \mu}.$$

The technical assumption that the observations do not lie in two strict subspaces of  $\mathbb{R}^D$  implies in particular that the observations span  $\mathbb{R}^D$ . We deduce that  $\mu > 0$ .

Now we discuss the challenge imposed by the constraint set. When the minimizer  $\mathbf{Q}^{(k+1)}$  lies on the boundary of the constraint set, the equality [CM99, Eqn. (4.3)] may not hold. However, if we denote the gradient of  $G$  with respect to its first argument by  $G_{\mathbf{Q}}$ , the inequality

$$\begin{aligned} 0 &\leq \langle \mathbf{Q} - \mathbf{Q}^{(k+1)}, G_{\mathbf{Q}}(\mathbf{Q}^{(k+1)}, \mathbf{Q}^{(k)}) \rangle \\ &= \langle \mathbf{Q} - \mathbf{Q}^{(k+1)}, \nabla F(\mathbf{Q}^{(k)}) + C(\mathbf{Q}^{(k)})(\mathbf{Q}^{(k+1)} - \mathbf{Q}^{(k)}) \rangle \end{aligned} \quad (\text{C.8})$$

holds for every  $\mathbf{Q}$  in the feasible set. This is simply the first-order necessary and sufficient condition for the constrained minimum of a smooth convex function over a convex set.

With the facts above, a proof that the iterates  $\mathbf{Q}^{(k)}$  converge to  $\mathbf{Q}_\delta$  follows the argument of [CM99, Lem. 5.1] nearly line-by-line. However, due to inequality (C.8), the final conclusion is that, at the limit point  $\overline{\mathbf{Q}}$ , the inequality  $\langle \mathbf{Q} - \overline{\mathbf{Q}}, \nabla F(\overline{\mathbf{Q}}) \rangle \geq 0$  holds for all feasible  $\mathbf{Q}$ . This inequality characterizes the global minimum of a convex function over a convex set, so the limit point must indeed be a global minimizer. That is,  $\overline{\mathbf{Q}} = \mathbf{Q}_\delta$ . In particular, this argument shows that the iterates  $\mathbf{P}^{(k)}$  converge to  $\mathbf{P}_\delta := \mathbf{I} - \mathbf{Q}_\delta$  as  $k \rightarrow \infty$ .

The only remaining claim is that  $\mathbf{P}_\delta = \mathbf{I} - \mathbf{Q}_\delta$  nearly minimizes (1.4). We abbreviate the objective of (1.4) under the identification  $\mathbf{Q} = \mathbf{I} - \mathbf{P}$  by

$$F_0(\mathbf{Q}) := \sum_{x \in \mathcal{X}} \|\mathbf{Q}x\|.$$

Define  $\mathbf{Q}_\star := \arg \min F_0(\mathbf{Q})$  with respect to the feasible set  $\mathbf{0} \preceq \mathbf{Q} \preceq \mathbf{I}$  and  $\text{tr}(\mathbf{Q}) = D - d$ . From the easy inequalities  $x \leq H_\delta(x, x) \leq x + \frac{1}{2}\delta$  for  $x \geq 0$ , we see that

$$0 \leq F(\mathbf{Q}) - F_0(\mathbf{Q}) \leq \frac{1}{2}\delta |\mathcal{X}|.$$

Evaluate the latter inequality at  $\mathbf{Q}_\delta$ , and subtract the result from the inequality evaluated at  $\mathbf{Q}_\star$  to reach

$$(F(\mathbf{Q}_\star) - F(\mathbf{Q}_\delta)) + (F_0(\mathbf{Q}_\delta) - F_0(\mathbf{Q}_\star)) \leq \frac{1}{2}\delta |\mathcal{X}|.$$

Since  $\mathbf{Q}_\delta$  and  $\mathbf{Q}_\star$  are optimal for their respective problems, both terms in parenthesis above are positive, and we deduce that  $F_0(\mathbf{Q}_\delta) - F_0(\mathbf{Q}_\star) \leq \frac{1}{2}\delta |\mathcal{X}|$ . Since  $F_0$  is the objective function of (1.4) under the map  $\mathbf{P} = \mathbf{I} - \mathbf{Q}$ , the proof is complete.  $\square$

#### APPENDIX D. PROOF OF THEOREM 6.2

We need to verify that the weak REAPER problem (6.1) has a solution that satisfies the constraints of the REAPER problem. The argument is constructive. Given a feasible point  $\mathbf{P}$  for (6.1), we convert it into a matrix  $\mathbf{P}'$  that has a smaller objective value and satisfies the constraint  $\mathbf{P}' \preceq \mathbf{I}$ . Since the weak REAPER problem is a relaxation of the REAPER problem, we can convert a solution of the weak REAPER problem into a solution to the REAPER problem.

Let  $\mathbf{P}$  be feasible for (6.1), and suppose that  $\mathbf{P} = \mathbf{U}\mathbf{\Lambda}\mathbf{U}^\mathbf{t}$  is an eigenvalue decomposition with eigenvalues arranged in weakly decreasing order:  $\lambda_1 \geq \dots \geq \lambda_D \geq 0$ . We construct a new diagonal matrix  $\mathbf{\Lambda}'$  whose diagonal elements satisfy

- (i)  $0 \leq \lambda'_i \leq 1$  for each  $i = 1, \dots, D$ .
- (ii)  $\sum_{i=1}^D \lambda'_i = \text{tr } \mathbf{P}$ .
- (iii)  $(1 - \lambda'_i)^2 \leq (1 - \lambda_i)^2$

Then we set  $\mathbf{P}' = \mathbf{U}\mathbf{\Lambda}'\mathbf{U}^\mathbf{t}$ . Let us demonstrate that this matrix has the required properties.

Conditions (i) and (ii) imply that  $\mathbf{P}'$  is feasible for REAPER. We only need to check that the objective value of (weak) REAPER is the same at  $\mathbf{P}$  and  $\mathbf{P}'$ . By unitary invariance, we compute that, for each vector  $\mathbf{x}$ ,

$$\|(\mathbf{I} - \mathbf{P}')\mathbf{x}\|^2 = \|(\mathbf{I} - \mathbf{\Lambda}')(\mathbf{U}^\mathbf{t}\mathbf{x})\|^2 = \sum_{i=1}^D (1 - \lambda'_i)^2 (\mathbf{U}^\mathbf{t}\mathbf{x})_i^2 \leq \sum_{i=1}^D (1 - \lambda_i)^2 (\mathbf{U}^\mathbf{t}\mathbf{x})_i^2 = \|(\mathbf{I} - \mathbf{P})\mathbf{x}\|^2.$$

Summing over  $\mathbf{x} \in \mathcal{X}$ , we see that the objective value of (weak) REAPER at  $\mathbf{P}'$  does not exceed the value at  $\mathbf{P}$ .

We construct the desired matrix  $\mathbf{\Lambda}'$  by a water-filling argument. Let  $i_\star$  be the smallest integer such that  $\sum_{j=1}^{i_\star} \lambda_j \leq i_\star$ . This number is well defined because  $\sum_{j=1}^D \lambda_j = d$  and  $0 \leq d \leq D$ .

Furthermore, since the eigenvalues are listed in decreasing order,

$$\lambda_i \leq \frac{1}{i_\star} \sum_{j=1}^{i_\star} \lambda_j \leq 1 \quad \text{for } i > i_\star.$$

Set

$$\lambda'_i = \begin{cases} 1, & i < i_\star \\ 1 - i_\star + \sum_{j=1}^{i_\star} \lambda_j, & i = i_\star \\ \lambda_i, & i > i_\star. \end{cases} \quad (\text{D.1})$$

We just need to verify that  $\lambda'_i$  satisfy Conditions (i)–(iii).

Condition (i) is evident for  $i \neq i_\star$ . When  $i = i_\star$ ,

$$\lambda'_{i_\star} = 1 - i_\star + \sum_{j=1}^{i_\star} \lambda_j \leq 1 \quad (\text{D.2})$$

because the sum is at most  $i_\star$ . Furthermore, since  $i_\star$  is the smallest integer with the specified property, we have

$$\lambda'_{i_\star} = \lambda_{i_\star} + \sum_{j=1}^{i_\star-1} \lambda_j - (i_\star - 1) \geq \lambda_{i_\star} \geq 0. \quad (\text{D.3})$$

The last inequality holds because  $\lambda_i \geq 0$  for each  $i$ .

For Condition (ii), we split the sum of  $\lambda'_i$  into three pieces according to the cases in (D.1), and we compute

$$\sum_{i=1}^D \lambda'_i = [i_\star - 1] + \left[ 1 - i_\star + \sum_{j=1}^{i_\star} \lambda_j \right] + \left[ \sum_{j=i_\star+1}^D \lambda_j \right] = \sum_{j=1}^D \lambda_j = \text{tr } \mathbf{P}.$$

Condition (iii) clearly holds for each  $i \neq i_\star$ . For  $i = i_\star$ , observe that  $0 \leq \lambda_{i_\star} \leq \lambda'_{i_\star} \leq 1$  because of (D.2) and (D.3). The claim follows.

#### ACKNOWLEDGMENTS

Lerman and Zhang were supported in part by the IMA and by NSF grants DMS-09-15064 and DMS-09-56072. McCoy and Tropp were supported by ONR awards N00014-08-1-0883 and N00014-11-1002, AFOSR award FA9550-09-1-0643, DARPA award N66001-08-1-2065, and a Sloan Research Fellowship. The authors thank Ben Recht, Eran Halperin, John Wright, and Yi Ma for helpful conversations.

#### REFERENCES

- [Amm93] L. P. Ammann. Robust singular value decompositions: A new approach to projection pursuit. *J. Amer. Statist. Assoc.*, 88(422):505–514, 1993.
- [And05] T. Ando. Schur complements and matrix inequalities: Operator-theoretic approach. In F. Zhang, editor, *The Schur Complement and its Applications*, volume 4 of *Numerical Methods and Algorithms*, chapter 5. Springer, New York, NY, 2005.
- [BH93] A. Bargiela and J. K. Hartley. Orthogonal linear regression algorithm based on augmented matrix formulation. *Comput. Oper. Res.*, 20:829–836, Oct. 1993.
- [Bha97] R. Bhatia. *Matrix Analysis*. Number 169 in Graduate Texts in Mathematics. Springer, New York, 1997.
- [BJ03] R. Basri and D. Jacobs. Lambertian reflectance and linear subspaces. *IEEE Trans. Pattern Anal. Mach. Intell.*, 25(2):218–233, Feb. 2003.
- [Bj96] Å. Björck. *Numerical Methods for Least Squares Problems*. Society for Industrial and Applied Mathematics, Philadelphia, PA, 1996.
- [BMDG05] A. Banerjee, S. Merugu, I. S. Dhillon, and J. Ghosh. Clustering with Bregman divergences. *J. Mach. Learn. Res.*, 6:1705–1749, 2005.
- [Bog98] V. I. Bogachev. *Gaussian Measures*, volume 62 of *Mathematical Surveys and Monographs*. American Mathematical Society, Providence, RI, 1998.

- [Bru09] S. C. Brubaker. Robust PCA and clustering in noisy mixtures. In *Proc. 20th Ann. ACM-SIAM Symp. Discrete Algorithms*, SODA '09, pages 1078–1087, Philadelphia, PA, USA, 2009. Society for Industrial and Applied Mathematics.
- [Cal06] Caltech 101. Online, April 2006. [http://www.vision.caltech.edu/Image\\_Datasets/Caltech101/](http://www.vision.caltech.edu/Image_Datasets/Caltech101/).
- [CFO07] C. Croux, P. Filzmoser, and M. Oliveira. Algorithms for projection pursuit robust principal component analysis. *Chemometrics Intell. Lab. Sys.*, 87(2):218–225, 2007.
- [CH00] C. Croux and G. Haesbroeck. Principal component analysis based on robust estimators of the covariance or correlation matrix: Influence functions and efficiencies. *Biometrika*, 87:603–618, 2000.
- [CK98] J. Costeira and T. Kanade. A multibody factorization method for independently moving objects. *Int. J. Comput. Vision*, 29(3):159–179, 1998.
- [CLMW11] E. J. Candès, X. Li, Y. Ma, and J. Wright. Robust principal component analysis? *J. Assoc. Comput. Mach.*, 58(3), 2011.
- [CM91] T. M. Cavalier and B. J. Melloy. An iterative linear programming solution to the Euclidean regression model. *Comput. Oper. Res.*, 18:655–661, 1991.
- [CM99] T. F. Chan and P. Mulet. On the convergence of the lagged diffusivity fixed point method in total variation image restoration. *SIAM J. Numer. Anal.*, 36:354–367, Feb. 1999.
- [CSPW11] V. Chandrasekaran, S. Sanghavi, P. A. Parrilo, and A. S. Willsky. Rank-sparsity incoherence for matrix decomposition. *SIAM J. Optim.*, 21(2):572–596, 2011.
- [CW82] R. D. Cook and S. Weisberg. *Residuals and influence in regression*. Chapman and Hall, New York, 1982.
- [CWB08] E. J. Candès, M. B. Wakin, and S. P. Boyd. Enhancing sparsity by reweighted  $l_1$  minimization. *J. Fourier Anal. Appl.*, 14(5):877–905, Dec. 2008.
- [Dav87] P. L. Davies. Asymptotic behaviour of S-estimates of multivariate location parameters and dispersion matrices. *Ann. Statist.*, 15(3):1269–1292, 1987.
- [DDL<sup>+</sup>88] S. Deerwester, S. Dumais, T. Landauer, G. Furna, and L. Beck. Improving Information Retrieval with Latent Semantic Indexing. In Christine L. Borgman and Edward Y. H. Pai, editors, *Information & Technology Planning for the Second 50 Years Proceedings of the 51st Annual Meeting of the American Society for Information Science*, volume 25, Atlanta, Georgia, October 1988. Learned Information, Inc.
- [DGK81] S. J. Devlin, R. Gnandesikan, and J. R. Kettenring. Robust estimation of dispersion matrices and principal components. *J. Amer. Statist. Assoc.*, 76(374):354–362, 1981.
- [DH01] D. L. Donoho and X. Huo. Uncertainty principles and ideal atomic decomposition. *IEEE Trans. Inform. Theory*, 47:2845–2862, Nov. 2001.
- [Dod87] Y. Dodge. An introduction to  $l_1$ -norm based statistical data analysis. *Comput. Statist. Data Anal.*, 5(4):239–253, 1987.
- [DS01] K. R. Davidson and S. J. Szarek. Local operator theory, random matrices and Banach spaces. In *Handbook of the geometry of Banach spaces, Vol. I*, pages 317–366. North-Holland, Amsterdam, 2001.
- [DS03] K. R. Davidson and S. J. Szarek. Addenda and corrigenda to: “Local operator theory, random matrices and Banach spaces” [in *Handbook of the geometry of Banach spaces, Vol. I*, 317–366, North-Holland, Amsterdam, 2001; MR1863696 (2004f:47002a)]. In *Handbook of the geometry of Banach spaces, Vol. 2*, pages 1819–1820. North-Holland, Amsterdam, 2003.
- [DZH06] C. Ding, D. Zhou, X. He, and H. Zha. R1-PCA: Rotational invariant  $L_1$ -norm principal component analysis for robust subspace factorization. In *ICML '06: Proc. 23rd Int. Conf. Machine Learning*, pages 281–288, Pittsburgh, PA, 2006. Association for Computing Machinery.
- [EHY95] R. Epstein, P.W. Hallinan, and A. L. Yuille.  $5 \pm 2$  eigenimages suffice: An empirical investigation of low-dimensional lighting models. In *Physics-Based Modeling in Computer Vision, 1995., Proceedings of the Workshop on*, page 108, jun. 1995.
- [EvdH10] A. Eriksson and A. van den Hengel. Efficient computation of robust low-rank matrix approximations in the presence of missing data using the  $l_1$  norm. In *Proc. 2010 IEEE Conf. Computer Vision and Pattern Recognition*, pages 771–778, june 2010.
- [EY39] C. Eckart and G. Young. A principal axis transformation for non-hermitian matrices. *Bull. Amer. Math. Soc.*, 45(2):118–121, 1939.
- [FB81] M. Fischler and R. Bolles. Random sample consensus: A paradigm for model fitting with applications to image analysis and automated cartography. *Comm. Assoc. Comput. Mach.*, 24(6):381–395, June 1981.
- [FFFP04] L. Fei-Fei, R. Fergus, and P. Perona. Learning generative visual models from a few training examples: an incremental bayesian approach tested on 101 object categories. In *CVPR 2004, Workshop on Generative-Model Based Vision*. IEEE, 2004.
- [GB08] M. Grant and S. Boyd. Graph implementations for nonsmooth convex programs. In V. Blondel, S. Boyd, and H. Kimura, editors, *Recent Advances in Learning and Control*, Lecture Notes in Control and Information Sciences, pages 95–110. Springer, London, 2008. [http://stanford.edu/~boyd/graph\\_dcp.html](http://stanford.edu/~boyd/graph_dcp.html).

- [GB10] M. Grant and S. Boyd. CVX: Matlab software for disciplined convex programming, version 1.21. <http://cvxr.com/cvx>, Oct. 2010.
- [GW95] M. X. Goemans and D. P. Williamson. Improved approximation for maximum cut and satisfiability problems using semidefinite programming. *J. Assoc. Comput. Mach.*, 42:1115–1145, 1995.
- [Har74a] H. L. Harter. The method of least squares and some alternatives: Part I. *Int. Statist. Rev.*, 42(2):147–174, Aug. 1974.
- [Har74b] H. L. Harter. The method of least squares and some alternatives: Part II. *Int. Statist. Rev.*, 42(3):235–264+282, Dec. 1974.
- [HMT11] N. Halko, P.-G. Martinsson, and J. A. Tropp. Finding structure with randomness: Stochastic algorithms for constructing approximate matrix decompositions. *SIAM Rev.*, 53(2):217–288, June 2011.
- [HR09] P. J. Huber and E. M. Ronchetti. *Robust Statistics*. Wiley Series in Probability and Statistics. Wiley, Hoboken, NJ, 2nd edition, 2009.
- [HTF09] T. Hastie, R. Tibshirani, and J. Friedman. *The Elements of Statistical Learning: Data Mining, Inference, and Prediction*. Springer, Berlin, 2nd edition, 2009.
- [HYL<sup>+</sup>03] J. Ho, M. Yang, J. Lim, K. Lee, and D. Kriegman. Clustering appearances of objects under varying illumination conditions. In *Proc. 2003 IEEE Int. Conf. Computer Vision and Pattern Recognition*, volume 1, pages 11–18, 2003.
- [Jol02] I. T. Jolliffe. *Principal Component Analysis*. Springer, Berlin, 2nd edition, 2002.
- [Kal07] S. Kale. *Efficient algorithms using the multiplicative weights update method*. PhD Dissertation, Computer Science, Princeton Univ., Princeton, NJ, 2007.
- [Ker83] D. Kershaw. Some extensions of W. Gautschi’s inequalities for the gamma function. *Math. Comput.*, 41(164):pp. 607–611, 1983.
- [Kwa08] N. Kwak. Principal component analysis based on  $L_1$ -norm maximization. *IEEE Trans. Pattern Anal. Mach. Intell.*, 30(9):1672–1680, 2008.
- [KXAH10] A. Khajehnejad, W. Xu, S. Avestimehr, and B. Hassibi. Improving the thresholds of sparse recovery: An analysis of a two-step reweighted Basis Pursuit algorithm. Submitted, 2010.
- [LC85] G. Li and Z. Chen. Projection-pursuit approach to robust dispersion matrices and principal components: Primary theory and Monte Carlo. *J. Amer. Statist. Assoc.*, 80(391):759–766, 1985.
- [Led01] M. Ledoux. *The Concentration of Measure Phenomenon*. Number 89 in Mathematical Surveys and Monographs. American Mathematical Society, Providence, Rhode Island, 2001.
- [LHK05] K. C. Lee, J. Ho, and D. Kriegman. Acquiring linear subspaces for face recognition under variable lighting. *IEEE Trans. Pattern Anal. Mach. Intelligence*, 27(5):684–698, 2005.
- [LLY<sup>+</sup>10] G. Liu, Z. Lin, S. Yan, J. Sun, Y. Yu, and Y. Ma. Robust recovery of subspace structures by low-rank representation. Available at [arxiv:1010.2955](https://arxiv.org/abs/1010.2955), 2010.
- [LMS<sup>+</sup>99] N. Locantore, J. S. Marron, D. G. Simpson, N. Tripoli, J. T. Zhang, and K. L. Cohen. Robust principal component analysis for functional data. *Test*, 8(1):1–73, 1999. With discussion and a rejoinder by the authors.
- [LS91] L. Lovasz and A. Schrijver. Cones of matrices and set-functions and 0–1 optimization. *SIAM J. Optim.*, 1(2):166–190, 1991.
- [LT91] M. Ledoux and M. Talagrand. *Probability in Banach spaces*, volume 23 of *Ergebnisse der Mathematik und ihrer Grenzgebiete (3) [Results in Mathematics and Related Areas (3)]*. Springer, Berlin, 1991. Isoperimetry and processes.
- [LV07] J. A. Lee and M. Verleysen. *Nonlinear Dimensionality Reduction*. Information Science and Statistics. Springer, Berlin, 2007.
- [LZ10] G. Lerman and T. Zhang.  $\ell_p$ -Recovery of the most significant subspace among multiple subspaces with outliers. Available at [arxiv:1012.4116](https://arxiv.org/abs/1012.4116), Dec. 2010.
- [LZ11] G. Lerman and T. Zhang. Robust recovery of multiple subspaces by geometric  $\ell_p$  minimization. *Ann. Statist.*, 39(5):2686–2715, 2011.
- [Mar76] R. A. Maronna. Robust M-estimators of multivariate location and scatter. *Ann. Statist.*, 4(1):51–67, 1976.
- [MMY06] R. A. Maronna, D. R. Martin, and V. J. Yohai. *Robust Statistics*. Wiley Series in Probability and Statistics. Wiley, Chichester, 2006. Theory and methods.
- [MT11] M. McCoy and J. A. Tropp. Two proposals for robust PCA using semidefinite programming. *Electron. J. Statist.*, 5:1123–1160, 2011.
- [NJB<sup>+</sup>08] J. Novembre, T. Johnson, K. Bryc, Z. Kutalik, A. R. Boyko, A. Auton, A. Indap, K. S. King, S. Bergmann, M.R. Nelson, M. Stephens, and C. D. Bustamante. Genes mirror geography within Europe. *Nature*, 456(7218):98–101, November 2008.
- [Nyq88] H. Nyquist. Least orthogonal absolute deviations. *Comput. Statist. Data Anal.*, 6(4):361–367, 1988.

- [OW85] M. R. Osborne and G. A. Watson. An analysis of the total approximation problem in separable norms, and an algorithm for the total  $l_1$  problem. *SIAM J. Sci. Statist. Comput.*, 6(2):410–424, 1985.
- [OW92] M. L. Overton and R. S. Womersley. On the sum of the largest eigenvalues of a symmetric matrix. *SIAM J. Matrix Anal. Appl.*, 13(1):41–45, Jan. 1992.
- [PPP<sup>+</sup>06] A. L. Price, N. J. Patterson, R. M. Plenge, M. E. Weinblatt, N. A. Shadick, and D. Reich. Principal components analysis corrects for stratification in genome-wide association studies. *Nature Genetics*, 38(8):904–909, 2006.
- [Rec12] B. Recht. Expectation–Maximization for fitting a mixture of subspaces. Personal communication, 2012.
- [RL87] P. J. Rousseeuw and A. M. Leroy. *Robust Regression and Outlier Detection*. Wiley Series in Probability and Mathematical Statistics: Applied Probability and Statistics. Wiley, New York, 1987.
- [Rou84] P. J. Rousseeuw. Least median of squares regression. *J. Amer. Statist. Assoc.*, 79(388):871–880, Dec. 1984.
- [SIP12] SIPI image database, volume 1: Textures. Online, January 2012. <http://sipi.usc.edu/database/database.php?volume=textures>.
- [SW87] H. Späth and G. A. Watson. On orthogonal linear approximation. *Numer. Math.*, 51:531–543, October 1987.
- [TB01] F. De La Torre and M. J. Black. Robust principal component analysis for computer vision. In *Proc. 8th IEEE Conf. Computer Vision*, volume 1, pages 362–369 vol.1, 2001.
- [TB03] F. De La Torre and M. J. Black. A framework for robust subspace learning. *Int. J. Comput. Vision*, 54:117–142, 2003. 10.1023/A:1023709501986.
- [Vaz03] V. V. Vazirani. *Approximation Algorithms*. Springer, Berlin, 2003.
- [VE80] H. Voss and U. Eckhardt. Linear convergence of generalized Weiszfeld’s method. *Computing*, 25:243–251, 1980. 10.1007/BF02242002.
- [vHV87] S. van Huffel and J. Vandewalle. *Total Least Squares: Computational Aspects and Analysis*. Society for Industrial and Applied Mathematics, Philadelphia, PA, 1987.
- [Wat01] G. A. Watson. Some problems in orthogonal distance and non-orthogonal distance regression. In *Proc. 2001 Symp. Algorithms for Approximation IV*. Defense Technical Information Center, 2001.
- [Wat02] G. A. Watson. On the Gauss–Newton method for  $l_1$  orthogonal distance regression. *IMA J. Numer. Anal.*, 22(3):345–357, 2002.
- [WTH09] D. Witten, R. Tibshirani, and T. Hastie. A penalized matrix decomposition, with applications to sparse principal components and canonical correlation analysis. *Biostat.*, 10(3):515–534, 2009.
- [XCM10] H. Xu, C. Caramanis, and S. Mannor. Principal Component Analysis with Contaminated Data: The High Dimensional Case. In *Proc. 2010 Conf. Learning Theory*. OmniPress, Haifa, 2010.
- [XCS10a] H. Xu, C. Caramanis, and S. Sanghavi. Robust PCA via outlier pursuit. In J. Lafferty, C. K. I. Williams, J. Shawe-Taylor, R.S. Zemel, and A. Culotta, editors, *Neural Information Processing Systems 23*, pages 2496–2504. MIT Press, Vancouver, 2010.
- [XCS10b] H. Xu, C. Caramanis, and S. Sanghavi. Robust PCA via outlier pursuit. *IEEE Trans. Inform. Theory*, pages 1–24, 2010.
- [XY95] L. Xu and A. L. Yuille. Robust principal component analysis by self-organizing rules based on statistical physics approach. *IEEE Trans. Neural Networks*, 6(1):131–143, Jan. 1995.
- [ZL11] T. Zhang and G. Lerman. A novel M-estimator for robust PCA. Available at [arXiv:1112.4863](https://arxiv.org/abs/1112.4863), 2011.
- [ZSL09] T. Zhang, A. Szlam, and G. Lerman. Median  $K$ -flats for hybrid linear modeling with many outliers. In *Proc. 12th IEEE Int. Conf. Computer Vision*, pages 234–241, Kyoto, 2009.