

Geometry and Symmetry in Short-and-Sparse Deconvolution*

Han-Wen Kuo[†], Yuqian Zhang[‡], Yenson Lau[†], and John Wright^{†§}

Abstract. We study the *short-and-sparse (SaS) deconvolution* problem of recovering a short signal $\mathbf{a}_0$ and a sparse signal $\mathbf{x}_0$ from their convolution. We propose a method based on nonconvex optimization, which under certain conditions recovers the target short and sparse signals, up to a signed shift symmetry which is intrinsic to this model. This symmetry plays a central role in shaping the optimization landscape for deconvolution. We give a *regional analysis*, which characterizes this landscape geometrically, on a union of subspaces. Our geometric characterization holds when the length- p_0 short signal $\mathbf{a}_0$ has shift coherence μ , and $\mathbf{x}_0$ follows a random sparsity model with sparsity rate $\theta \in [\frac{c_1}{p_0}, \frac{c_2}{p_0\sqrt{\mu}+\sqrt{p_0}}] \cdot \frac{1}{\log^2 p_0}$. Based on this geometry, we give a provable method that successfully solves SaS deconvolution with high probability.

Key words. signal reconstruction, blind deconvolution, nonconvex geometry, nonconvex optimization

AMS subject classifications. 94A12, 90C26, 65Y20

DOI. 10.1137/19M1237569

1. Introduction. Datasets in a wide range of areas, including neuroscience [37], microscopy [15], and astronomy [49], can be modeled as superpositions of translations of a basic motif. Data of this nature can be modeled mathematically as a convolution $\mathbf{y} = \mathbf{a}_0 * \mathbf{x}_0$, between a *short* signal $\mathbf{a}_0$ (the motif) and a longer *sparse* signal $\mathbf{x}_0$, whose nonzero entries indicate where in the sample the motif is present. A very similar structure arises in image deblurring [14], where $\mathbf{y}$ is a blurry image, $\mathbf{a}_0$ the blur kernel, and $\mathbf{x}_0$ the (edge map) of the target sharp image.

Motivated by these and related problems in imaging and scientific data analysis, we study the *short-and-sparse (SaS) deconvolution* problem of recovering a short signal $\mathbf{a}_0 \in \mathbb{R}^{p_0}$ and a sparse signal $\mathbf{x}_0 \in \mathbb{R}^n$ ($n \gg p_0$) from their length- n cyclic convolution $\mathbf{y} = \mathbf{a}_0 * \mathbf{x}_0 \in \mathbb{R}^n$.¹ This SaS model exhibits a basic *scaled shift symmetry*: for any nonzero scalar α and cyclic shift $s_\ell[\cdot]$,

$$(1.1) \quad \left(\alpha s_\ell[\mathbf{a}_0] \right) * \left(\frac{1}{\alpha} s_{-\ell}[\mathbf{x}_0] \right) = \mathbf{y}.$$

Because of this symmetry, we only expect to recover $\mathbf{a}_0$ and $\mathbf{x}_0$ up to a signed shift (see

*Received by the editors January 8, 2019; accepted for publication (in revised form) December 2, 2019; published electronically February 27, 2020.

<https://doi.org/10.1137/19M1237569>

Funding: This work was funded by National Science Foundation grants NSF 1343282, NSF CCF 1527809, and NSF IIS 1546411.

[†]Department of Electronic Engineering and Data Science Institute, Columbia University, New York, NY 10027 USA (hk2673@columbia.edu, yl3027@columbia.edu, jw2966@columbia.edu).

[‡]Department of Computer Science, Cornell University, Ithaca, NY 14853 USA (yz2409@columbia.edu).

[§]Department of Applied Physics and Applied Mathematics, Columbia University, New York, NY 10027 USA

¹In this paper, the cyclic convolution $\mathbf{a}_0 * \mathbf{x}_0$ assumes $\mathbf{a}_0$ to be zero-padded $[\mathbf{a}_0, \mathbf{0}^{n-p_0}]$ to length n .

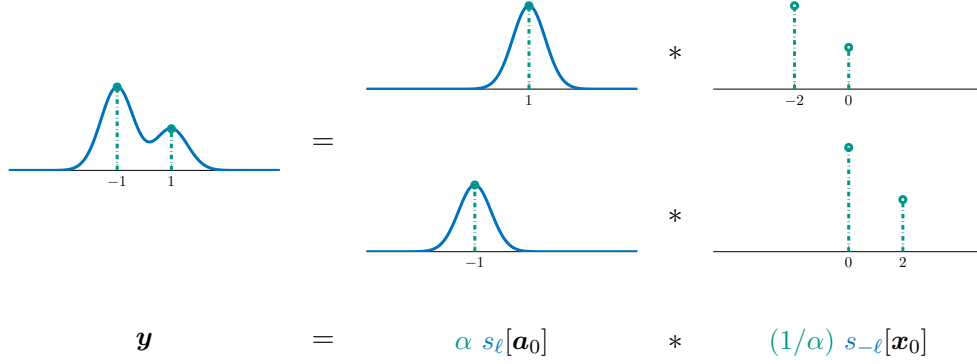


Figure 1. Shift symmetry in short-and-sparse deconvolution. An observation y (left) which is a convolution of a short signal a_0 and a sparse signal x_0 (top right) can be equivalently expressed as a convolution of $s_\ell[a_0]$ and $s_{-\ell}[x_0]$, where $s_\ell[\cdot]$ denotes shift ℓ samples. The ground truth signals a_0 and x_0 can only be identified up to a scaled shift.

Figure 1). Our problem of interest can be stated more formally as follows.

Problem 1.1 (short-and-sparse deconvolution). Given the cyclic convolution² $y = a_0 * x_0 \in \mathbb{R}^n$ of $a_0 \in \mathbb{R}^{p_0}$ short ($p_0 \ll n$) and $x_0 \in \mathbb{R}^n$ sparse, recover a_0 and x_0 , up to a scaled shift.

Despite a long history and many applications, until recently very little algorithmic theory was available for SaS deconvolution. Much of this difficulty can be attributed to the scale-shift symmetry: natural convex relaxations fail,³ and nonconvex formulations exhibit a complicated optimization landscape, with many equivalent global minimizers (scaled shifts of the ground truth), additional local minimizers (scaled shift truncations of the ground truth), and a variety of critical points [63, 64]. Currently available theory guarantees approximate recovery of a truncation⁴ of a shift $s_\ell[a_0]$, rather than guaranteeing recovery of a_0 as a whole, and requires certain (complicated) conditions on the convolution matrix associated with a_0 [63].

In this paper, we describe an algorithm which, under simpler conditions, *exactly* recovers a scaled shift of the pair (a_0, x_0) . Our algorithm is based on a formulation first introduced in [64], which casts the deconvolution problem as (nonconvex) optimization over the sphere. We characterize the geometry of this objective function and show that near a certain union of subspaces, every local minimizer is very close to a signed shift of a_0 . Based on this geometric analysis, we give provable methods for SaS deconvolution that exactly recover a scaled shift of (a_0, x_0) whenever a_0 is *shift-incoherent* and x_0 is a sufficiently sparse random vector. Our geometric analysis highlights the role of symmetry in shaping the objective landscape for SaS deconvolution.

The remainder of this paper is organized as follows. Section 2 introduces our optimization

²Our result can be applied to recovering direct convolutions. Let $y \in \mathbb{R}^{p_0+n-1}$ be the direct convolution between $a_0 \in \mathbb{R}^{p_0}$ and $x_0 \in \mathbb{R}^n$; then y can also be expressed as circular convolution between a_0 and $[x_0; \mathbf{0}^{p_0-1}]$.

³Such as matrix lifting relaxation [2, 39], in which a_0 or x_0 resides in random subspaces without shift symmetry.

⁴That is, the portion of the shifted signal $s_\ell[a_0]$ that falls in the window $\{0, \dots, p_0 - 1\}$.

approach and modeling assumptions. Section 3 introduces our main results—both geometric and algorithmic—and compares them to the literature. Sections 4 and 5 describes the main ideas of our analysis. Section 6 demonstrates the experimental performance of the analyzed algorithm. Finally, section 7 discusses two main limitations of our analysis and describes directions for future work.

2. Formulation and assumptions.

2.1. Nonconvex SaS over the sphere. Our starting point is the (natural) formulation

$$(2.1) \quad \min_{\mathbf{a}, \mathbf{x}} \underbrace{\frac{1}{2} \|\mathbf{a} * \mathbf{x} - \mathbf{y}\|_2^2}_{\text{Data Fidelity}} + \underbrace{\lambda \|\mathbf{x}\|_1}_{\text{Sparsity}} \quad \text{s.t.} \quad \|\mathbf{a}\|_2 = 1.$$

We term this optimization problem the *Bilinear Lasso*, for its resemblance to the Lasso estimator in statistics. Indeed, letting

$$(2.2) \quad \varphi_{\text{lasso}}(\mathbf{a}) \equiv \min_{\mathbf{x}} \left\{ \frac{1}{2} \|\mathbf{a} * \mathbf{x} - \mathbf{y}\|_2^2 + \lambda \|\mathbf{x}\|_1 \right\}$$

denote the optimal Lasso cost, we see that (2.1) simply optimizes φ_{lasso} with respect to $\mathbf{a}$:

$$(2.3) \quad \min_{\mathbf{a}} \varphi_{\text{lasso}}(\mathbf{a}) \quad \text{s.t.} \quad \|\mathbf{a}\|_2 = 1.$$

In (2.1)–(2.3), we constrain $\mathbf{a}$ to have unit ℓ^2 norm. This constraint breaks the scale ambiguity between $\mathbf{a}$ and $\mathbf{x}$. Moreover, the choice of constraint manifold has surprisingly strong implications for computation: if $\mathbf{a}$ is instead constrained to the simplex, the problem admits trivial global minimizers. In contrast, local minima of the sphere-constrained formulation often correspond to shifts (or shift truncations [64]) of the ground truth $\mathbf{a}_0$.

The problem (2.3) is defined in terms of the optimal Lasso cost. This function is challenging to analyze, especially far away from $\mathbf{a}_0$. The article [64] analyzes the local minima of a simplification of (2.3), obtained by approximating⁵ the data fidelity term as

$$(2.4) \quad \begin{aligned} \frac{1}{2} \|\mathbf{a} * \mathbf{x} - \mathbf{y}\|_2^2 &= \frac{1}{2} \|\mathbf{a} * \mathbf{x}\|_2^2 - \langle \mathbf{a} * \mathbf{x}, \mathbf{y} \rangle + \frac{1}{2} \|\mathbf{y}\|_2^2 \\ &\approx \frac{1}{2} \|\mathbf{x}\|_2^2 - \langle \mathbf{a} * \mathbf{x}, \mathbf{y} \rangle + \frac{1}{2} \|\mathbf{y}\|_2^2. \end{aligned}$$

This yields a simpler objective function,

$$(2.5) \quad \varphi_{\ell^1}(\mathbf{a}) = \min_{\mathbf{x}} \left\{ \frac{1}{2} \|\mathbf{x}\|_2^2 - \langle \mathbf{a} * \mathbf{x}, \mathbf{y} \rangle + \frac{1}{2} \|\mathbf{y}\|_2^2 + \lambda \|\mathbf{x}\|_1 \right\}.$$

We make one further simplification to this problem, replacing the nondifferentiable penalty $\|\cdot\|_1$ with a smooth approximation $\rho(\mathbf{x})$.⁶ Our analysis allows for a variety of smooth sparsity surrogates $\rho(\mathbf{x})$; for concreteness, we state our main results for the particular penalty⁷

$$(2.6) \quad \rho(\mathbf{x}) = \sum_i (\mathbf{x}_i^2 + \delta^2)^{1/2}.$$

⁵For a generic $\mathbf{a}$, we have $\langle s_i[\mathbf{a}], s_j[\mathbf{a}] \rangle \approx 0$ and hence $\|\mathbf{a} * \mathbf{x}\|_2^2 = \mathbf{x}^* \mathbf{C}_a^* \mathbf{C}_a \mathbf{x} \approx \mathbf{x}^* \mathbf{I} \mathbf{x} = \|\mathbf{x}\|_2^2$. The use of φ_ρ performs not as ideal compared to Bilinear Lasso when this approximation is inexact; see section 7.

⁶ φ_{ℓ^1} is not twice differentiable everywhere and hence can't be minimized with conventional second order methods.

⁷This particular surrogate is sometimes called the pseudo-Huber function.

For $\delta > 0$, this is a smooth function of $\mathbf{x}$; as $\delta \searrow 0$ it approaches $\|\mathbf{x}\|_1$. Replacing $\|\cdot\|_1$ with $\rho(\cdot)$, we obtain the objective function which will be our main object of study,

$$(2.7) \quad \varphi_\rho(\mathbf{a}) = \min_{\mathbf{x}} \left\{ \frac{1}{2} \|\mathbf{x}\|_2^2 - \langle \mathbf{a} * \mathbf{x}, \mathbf{y} \rangle + \frac{1}{2} \|\mathbf{y}\|_2^2 + \lambda \rho(\mathbf{x}) \right\}.$$

As in [64], we optimize $\varphi_\rho(\mathbf{a})$ over the sphere $\mathbb{S}^{p-1}$:

$$(2.8) \quad \boxed{\min_{\mathbf{a}} \varphi_\rho(\mathbf{a}) \quad \text{s.t.} \quad \mathbf{a} \in \mathbb{S}^{p-1}.}$$

Here, we set $p = 3p_0 - 2$. As we will see, optimizing over this slightly higher dimensional sphere enables us to recover a (full) shift of $\mathbf{a}_0$, rather than a *truncated* shift. Our approach will leverage the following fact: if we view $\mathbf{a} \in \mathbb{S}^{p-1}$ as indexed by coordinates $W = \{-p_0 + 1, \dots, 2p_0 - 1\}$, then for any shifts $\ell \in \{-p_0 + 1, \dots, p_0 - 1\}$, the support of ℓ -shifted short signal $s_\ell[\mathbf{a}_0]$ is entirely contained in interval W . We will give a provable method which recovers a scaled version of one of these canonical shifts.

2.2. Analysis setting and assumptions. For convenience, we assume that $\mathbf{a}_0$ has unit ℓ^2 norm, i.e., $\mathbf{a}_0 \in \mathbb{S}^{p_0-1}$.⁸ Our analysis makes two main assumptions, on the short motif $\mathbf{a}_0$ and the sparse map $\mathbf{x}_0$, respectively:

The first is that distinct shifts $\mathbf{a}_0$ have small inner product. We define the *shift coherence* of $\mu(\mathbf{a}_0)$ to be the largest inner product between distinct shifts:

$$(2.9) \quad \mu(\mathbf{a}_0) = \max_{\ell \neq 0} |\langle \mathbf{a}_0, s_\ell[\mathbf{a}_0] \rangle|.$$

The quantity $\mu(\mathbf{a}_0)$ is bounded between 0 and 1. Our theory allows any μ smaller than some numerical constant. Figure 2 shows three examples of families of $\mathbf{a}_0$ that satisfy this assumption:

- *Spiky.* When $\mathbf{a}_0$ is close to the Dirac delta δ_0 , the shift coherence $\mu(\mathbf{a}_0) \approx 0$.⁹ Here, the observed signal $\mathbf{y}$ consists of a superposition of sharp pulses. This is arguably the easiest instance of SaS deconvolution.
- *Generic.* If $\mathbf{a}_0$ is chosen uniformly at random from the sphere $\mathbb{S}^{p_0-1}$, its coherence is bounded as $\mu(\mathbf{a}_0) \lesssim \sqrt{1/p_0}$ with high probability.
- *Tapered generic low-pass.* Here, $\mathbf{a}_0$ is generated by taking a random conjugate symmetric superposition of the first L length- p_0 discrete Fourier transform (DFT) basis signals, windowing (e.g., with a Hamming window) and normalizing to unit ℓ^2 norm. When $L = p_0\sqrt{1-\beta}$, with high probability $\mu(\mathbf{a}_0) \lesssim \beta$. In this model, μ does not have to diminish as p_0 grows—it can be a fixed constant.¹⁰

⁸This is purely a technical convenience. Our theory guarantees recovery of a signed shift $(\pm s_\ell[\mathbf{a}_0], \pm s_{-\ell}[\mathbf{x}_0])$ of the truth. If $\mathbf{a}_0$ does not have unit norm, identical reasoning implies that our method recovers a scaled shift $(\alpha s_\ell[\mathbf{a}_0], \alpha^{-1} s_{-\ell}[\mathbf{x}_0])$ with $\alpha = \pm \frac{1}{\|\mathbf{a}_0\|_2}$.

⁹The use of “ $\approx$ ” here suppresses constant and logarithmic factors.

¹⁰The upper right panel of Figure 2 is generated using random DFT components with frequencies smaller than one-third Nyquist. Such a kernel is incoherent, with high probability. Many commonly occurring low-pass kernels have $\mu(\mathbf{a}_0)$ larger—very close to one. One of the most important limitations of our results is that they do not provide guarantees in this highly coherent situation. See [34].

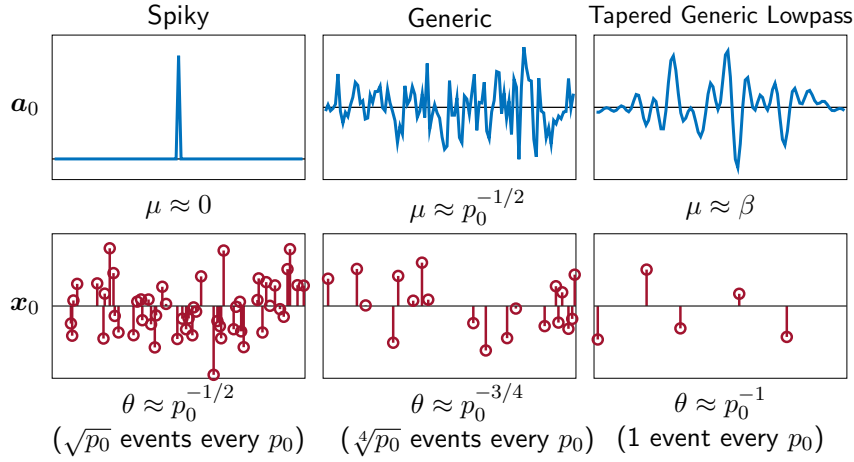


Figure 2. Sparsity-coherence tradeoff. *Top:* three families of motifs $\mathbf{a}_0$ with varying coherence μ . *Bottom:* maximum allowable sparsity θ and number of copies θp_0 within each length- p_0 window. Here, we suppress constants and logarithmic factors. When the target motif has smaller shift-coherence μ , our result allows larger θ , and vice versa. This sparsity-coherence tradeoff is made precise in our main result, [Theorem 3.1](#), which, loosely speaking, asserts that when $\theta \lesssim 1/(p_0\sqrt{\mu} + \sqrt{p_0})$, our method succeeds.

Intuitively speaking, problems with smaller μ are easier to solve, a claim which will be made precise in our technical results.

We assume that $\mathbf{x}_0$ is a sparse random vector. More precisely, we assume that $\mathbf{x}_0$ is Bernoulli–Gaussian, with rate θ :

$$(2.10) \quad \mathbf{x}_{0i} = \omega_i \mathbf{g}_i,$$

where $\omega_i \sim \text{Ber}(\theta)$, $\mathbf{g}_i \sim \mathcal{N}(0, 1)$, and all random variables are jointly independent. We write this as

$$(2.11) \quad \mathbf{x}_0 \sim_{\text{i.i.d.}} \text{BG}(\theta).$$

Here, θ is the probability that a given entry $\mathbf{x}_{0i}$ is nonzero. Problems with smaller θ are easier to solve. In the extreme case, when $\theta \ll 1/p_0$, the observation $\mathbf{y}$ contains many isolated copies of the motif $\mathbf{a}_0$, and $\mathbf{a}_0$ can be determined by direct inspection. Our analysis will focus on the nontrivial scenario, when $\theta \gtrsim 1/p_0$.

Our technical results will articulate *sparsity-coherence* tradeoffs, in which smaller coherence μ enables larger θ , and vice versa. More specifically, in our main theorem, the sparsity-coherence relationship is captured in the form

$$(2.12) \quad \theta \lesssim 1/(p_0\sqrt{\mu} + \sqrt{p_0}).$$

When the target $\mathbf{a}_0$ is very shift-incoherent ($\mu \approx 0$), our method succeeds when each length- p_0 window contains about $\sqrt{p_0}$ copies of $\mathbf{a}_0$. When μ is larger (as in the generic low-pass model), our method succeeds as long as relatively few copies of $\mathbf{a}_0$ overlap in the observed signal. In [Figure 2](#), we illustrate these tradeoffs for the three models described above.

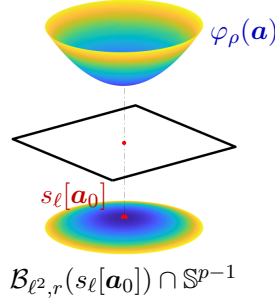


Figure 3. Geometry of φ_ρ near a shift of $\mathbf{a}_0$. *Bottom:* a portion of the sphere $\mathbb{S}^{p-1}$, colored according to φ_ρ . *Top:* φ_ρ visualized as height. φ_ρ is strongly convex in this region, and it has a minimizer very close to $s_\ell[\mathbf{a}_0]$.

3. Main results: Geometry and algorithms. In this section, we introduce our main results—on the geometry of φ_ρ (subsection 3.1) and its algorithmic implications (subsection 3.2). Finally, in subsection 3.3, we compare these results with the literature on deconvolution.

3.1. Geometry of the objective φ_ρ . The goal in SaS deconvolution is to recover $\mathbf{a}_0$ (and $\mathbf{x}_0$) up to a signed shift; i.e., we wish to recover some $\pm s_\ell[\mathbf{a}_0]$. The shifts $\pm s_\ell[\mathbf{a}_0]$ play a key role in shaping the landscape of φ_ρ . In particular, we will argue that over a certain subset of the sphere, *every local minimum of φ_ρ is close to some $\pm s_\ell[\mathbf{a}_0]$* .

To gain intuition into the properties of φ_ρ , we first visualize this function in the vicinity of a single shift $s_\ell[\mathbf{a}_0]$ of the ground truth $\mathbf{a}_0$. In Figure 3, we plot the function value of φ_ρ over

$$\mathcal{B}_{\ell^2, r}(s_\ell[\mathbf{a}_0]) \cap \mathbb{S}^{p-1},$$

where $\mathcal{B}_{\ell^2, r}(\mathbf{a})$ is a ball of radius r around $\mathbf{a}$. We make two observations:

- The objective function φ_ρ is strongly convex in this neighborhood of $s_\ell[\mathbf{a}_0]$.
- There is a local minimizer very close to $s_\ell[\mathbf{a}_0]$.

We next visualize the objective function φ_ρ near the linear span of *two* different shifts, $s_{\ell_1}[\mathbf{a}_0]$ and $s_{\ell_2}[\mathbf{a}_0]$. More precisely, we plot φ_ρ near the intersection (Figure 4, left) of the sphere $\mathbb{S}^{p-1}$ and the linear subspace

$$\mathcal{S}_{\{\ell_1, \ell_2\}} = \{ \alpha_1 s_{\ell_1}[\mathbf{a}_0] + \alpha_2 s_{\ell_2}[\mathbf{a}_0] \mid \alpha_1, \alpha_2 \in \mathbb{R} \}.$$

We make three observations:

- Again, there is a local minimizer near each shift $s_\ell[\mathbf{a}_0]$.
- These are the *only* local minimizers in the vicinity of $\mathcal{S}_{\{\ell_1, \ell_2\}}$. In particular, the objective function φ exhibits *negative curvature* along $\mathcal{S}_{\{\ell_1, \ell_2\}}$ at any superposition $\alpha_1 s_{\ell_1}[\mathbf{a}_0] + \alpha_2 s_{\ell_2}[\mathbf{a}_0]$ whose weights α_1 and α_2 are balanced, i.e., $|\alpha_1| \approx |\alpha_2|$.
- Furthermore, the function φ_ρ exhibits *positive curvature* in directions away from the subspace $\mathcal{S}_{\ell_1, \ell_2}$.

Finally, we visualize φ_ρ over the intersection (Figure 5, left) of the sphere $\mathbb{S}^{p-1}$ with the linear span of three shifts $s_{\ell_1}[\mathbf{a}_0]$, $s_{\ell_2}[\mathbf{a}_0]$, $s_{\ell_3}[\mathbf{a}_0]$ of the true kernel $\mathbf{a}_0$:

$$\mathcal{S}_{\{\ell_1, \ell_2, \ell_3\}} = \{ \alpha_1 s_{\ell_1}[\mathbf{a}_0] + \alpha_2 s_{\ell_2}[\mathbf{a}_0] + \alpha_3 s_{\ell_3}[\mathbf{a}_0] \mid \alpha_1, \alpha_2, \alpha_3 \in \mathbb{R} \}.$$

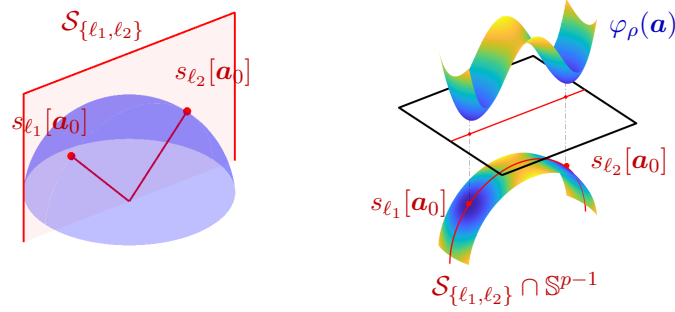


Figure 4. Geometry of φ_ρ near the span $\mathcal{S}_{\{\ell_1, \ell_2\}}$ of two shifts of $\mathbf{a}_0$. Left: each pair of shifts $s_{\ell_1}[\mathbf{a}_0]$, $s_{\ell_2}[\mathbf{a}_0]$ defines a linear subspace $\mathcal{S}_{\{\ell_1, \ell_2\}}$ of $\mathbb{R}^p$. Center/right: every local minimum of φ_ρ near $\mathcal{S}_{\{\ell_1, \ell_2\}}$ (red line) is close to either $s_{\ell_1}[\mathbf{a}_0]$ or $s_{\ell_2}[\mathbf{a}_0]$; there is a negative curvature in the middle of $s_{\ell_1}[\mathbf{a}_0]$ and $s_{\ell_2}[\mathbf{a}_0]$, and φ_ρ is convex in direction away from $\mathcal{S}_{\{\ell_1, \ell_2\}}$.

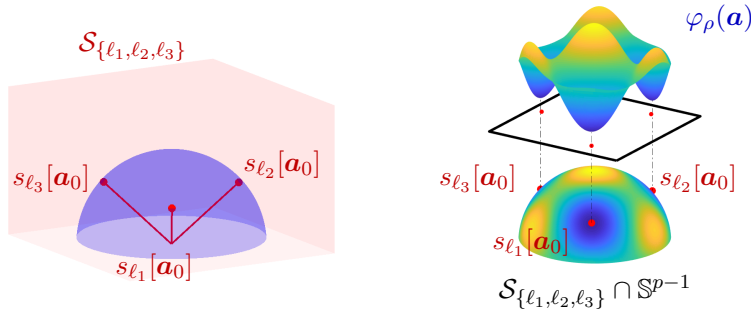


Figure 5. Geometry of φ_ρ over the span $\mathcal{S}_{\{\ell_1, \ell_2, \ell_3\}}$ of three shifts of $\mathbf{a}_0$. The subspace $\mathcal{S}_{\{\ell_1, \ell_2, \ell_3\}}$ is three-dimensional; its intersection with the sphere $\mathbb{S}^{p-1}$ is isomorphic to a two-dimensional sphere. On this set, φ_ρ has local minimizers near each of the $s_{\ell_i}[\mathbf{a}_0]$ and are the only minimizers near $\mathcal{S}_{\{\ell_1, \ell_2, \ell_3\}}$.

Again, there is a local minimizer near each signed shift. At roughly balanced superpositions of shifts, the objective function exhibits negative curvature. As a result, again, the *only* local minimizers are close to signed shifts.

Our main geometric result will show that these properties are obtained from *every* subspace spanned by a few shifts of $\mathbf{a}_0$. Indeed, for each subset

$$(3.1) \quad \tau \subseteq \{-p_0 + 1, \dots, p_0 - 1\},$$

define a linear subspace

$$(3.2) \quad \mathcal{S}_\tau = \left\{ \sum_{\ell \in \tau} \alpha_\ell s_\ell[\mathbf{a}_0] \mid \alpha_{-p_0+1}, \dots, \alpha_{p_0-1} \in \mathbb{R} \right\}.$$

The subspace $\mathcal{S}_\tau$ is the linear span of the shifts $s_\ell[\mathbf{a}_0]$ indexed by ℓ in the set τ . Our geometric theory will show that with high probability the function φ_ρ has no spurious local minimizers

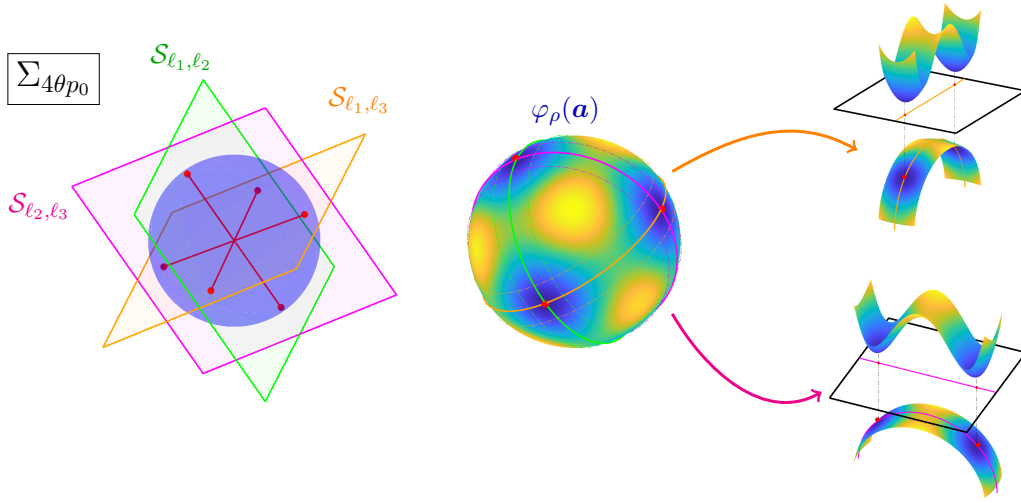


Figure 6. Geometry of φ_ρ over the union of subspaces $\Sigma_{2\theta p_0}$. Left: schematic representation of the union of subspaces $\Sigma_{4\theta p_0}$. For each set τ of at most $4\theta p_0$ shifts, we have a subspace $\mathcal{S}_\tau$. Right: φ_ρ has good geometry near this union of subspaces.

near any $\mathcal{S}_\tau$ for which τ is not too large—say, $|\tau| \leq 4\theta p_0$. Combining all of these subspaces into a single geometric object, define the union of subspaces

$$(3.3) \quad \Sigma_{4\theta p_0} = \bigcup_{|\tau| \leq 4\theta p_0} \mathcal{S}_\tau.$$

Figure 6 (left) gives a schematic representation of this set. We claim the following:

- In the neighborhood of $\Sigma_{4\theta p_0}$, all local minimizers are near signed shifts.
- The value of φ_ρ grows in any direction away from $\Sigma_{4\theta p_0}$.

Our main result formalizes the above observations under two key assumptions: first, that the sparsity rate θ is sufficiently small (relative to the shift coherence μ of p_0), and, second, that the signal length n is sufficiently large.

Theorem 3.1 (main geometric theorem). Let $\mathbf{y} = \mathbf{a}_0 * \mathbf{x}_0$ with $\mathbf{a}_0 \in \mathbb{S}^{p_0-1}$ μ -shift coherent¹¹ and $\mathbf{x}_0 \sim_{\text{i.i.d.}} \text{BG}(\theta) \in \mathbb{R}^n$ with sparsity rate

$$(3.4) \quad \theta \in \left[\frac{c_1}{p_0}, \frac{c_2}{p_0 \sqrt{\mu} + \sqrt{p_0}} \right] \cdot \frac{1}{\log^2 p_0}.$$

Choose $\rho(x) = \sqrt{x^2 + \delta^2}$ and set $\lambda = 0.1/\sqrt{p_0\theta}$ in φ_ρ . Then there exist $\delta > 0$ and numerical constant c such that if $n \geq \text{poly}(p_0)$, with high probability, every local minimizer $\bar{\mathbf{a}}$ of φ_ρ over $\Sigma_{4\theta p_0}$ satisfies $\|\bar{\mathbf{a}} - \sigma s_\ell[\mathbf{a}_0]\|_2 \leq c \max\{\mu, p_0^{-1}\}$ for some signed shift $\sigma s_\ell[\mathbf{a}_0]$ of the true kernel. Above, $c_1, c_2 > 0$ are positive numerical constants.

¹¹Typically it is possible to provide an overestimate $p'_0 \geq p_0$. Our theory and algorithm can be applied directly to the overestimate p'_0 , with the caveat that the sparsity rate θ now scales with p'_0 rather than p_0 .

Proof. This follows from [Theorem 4.1](#). ■

The upper bound on θ in (3.4) yields the tradeoff between coherence and sparsity described in [Figure 2](#). Simply put, when $\mathbf{a}_0$ is better conditioned (as a kernel), its coherence μ is smaller and $\mathbf{x}_0$ can be denser.

At a technical level, our proof of [Theorem 3.1](#) shows that (i) $\varphi_\rho(\mathbf{a})$ is strongly convex in the vicinity of each signed shift, and that at every other point $\mathbf{a}$ near $\Sigma_{4\theta p_0}$, there is either (ii) a nonzero gradient or (iii) a direction of strict negative curvature; furthermore (iv) the function φ_ρ grows away from $\Sigma_{4\theta p_0}$. Points (ii)–(iii) imply that near $\Sigma_{4\theta p_0}$ there are no “flat” saddles: every saddle point has a direction of strict negative curvature. We will leverage these properties to propose an efficient algorithm for finding a local minimizer near $\Sigma_{4\theta p_0}$. Moreover, this minimizer is close enough to a shift (here, $\|\bar{\mathbf{a}} - s_\ell[\mathbf{a}_0]\|_2 \lesssim \mu$) for us to exactly recover $s_\ell[\mathbf{a}_0]$: we will give a refinement algorithm that produces $(\pm s_\ell[\mathbf{a}_0], \pm s_{-\ell}[\mathbf{x}_0])$.

3.2. Provable algorithm for SaS deconvolution. The objective function φ_ρ has good geometric properties on (and near!) the union of subspaces $\Sigma_{4\theta p_0}$. In this section, we show an efficient method that exactly recovers $\mathbf{a}_0$ and $\mathbf{x}_0$ up to shift symmetry. Although our geometric analysis only controls φ_ρ near $\Sigma_{4\theta p_0}$, we will give a descent method which, with appropriate initialization $\mathbf{a}^{(0)}$, produces iterates $\mathbf{a}^{(1)}, \dots, \mathbf{a}^{(k)}, \dots$ that remain close to $\Sigma_{4\theta p_0}$ for all k . In short, it is easy to *start* near $\Sigma_{4\theta p_0}$ and easy to *stay* near $\Sigma_{4\theta p_0}$. After finding a local minimizer $\bar{\mathbf{a}}$, we refine it to produce a signed shift of $(\mathbf{a}_0, \mathbf{x}_0)$ using alternating minimization.

The next two paragraphs give the main ideas behind the principal steps of the algorithm. We then describe its components in more detail ([Algorithm 3.1](#)) and state our main algorithmic result ([Theorem 3.2](#)), which asserts that under appropriate conditions this method produces a signed shift of $(\mathbf{a}_0, \mathbf{x}_0)$.

Our algorithm starts with an initialization scheme which generates $\mathbf{a}^{(0)}$ near the union of subspaces $\Sigma_{4\theta p_0}$, which consists of linear combinations of just a few shifts of $\mathbf{a}_0$. How can we find a point near this union? Notice that *the data $\mathbf{y}$ also consists of a linear combination of just a few shifts of $\mathbf{a}_0$* . Indeed,

$$(3.5) \quad \mathbf{y} = \mathbf{a}_0 * \mathbf{x}_0 = \sum_{\ell \in \text{supp}(\mathbf{x}_0)} \mathbf{x}_{0\ell} s_\ell[\mathbf{a}_0].$$

A length- p_0 segment of data $\mathbf{y}_{0,\dots,p_0-1} = [\mathbf{y}_0, \dots, \mathbf{y}_{p_0-1}]^*$ captures portions of roughly $2\theta p_0 \ll 4\theta p_0$ shifts $s_\ell[\mathbf{a}_0]$.

Many of these copies of $\mathbf{a}_0$ are truncated by the restriction to $\{0, \dots, p_0 - 1\}$. A relatively simple remedy is as follows: First, we zero-pad $\mathbf{y}_{0,\dots,p_0-1}$ to length $p = 3p_0 - 2$, giving

$$(3.6) \quad [\mathbf{0}^{p_0-1}; \mathbf{y}_0; \dots; \mathbf{y}_{p_0-1}; \mathbf{0}^{p_0-1}].$$

Zero-padding provides enough space to accommodate any shift $s_\ell[\mathbf{a}_0]$ with $\ell \in \tau$. We then

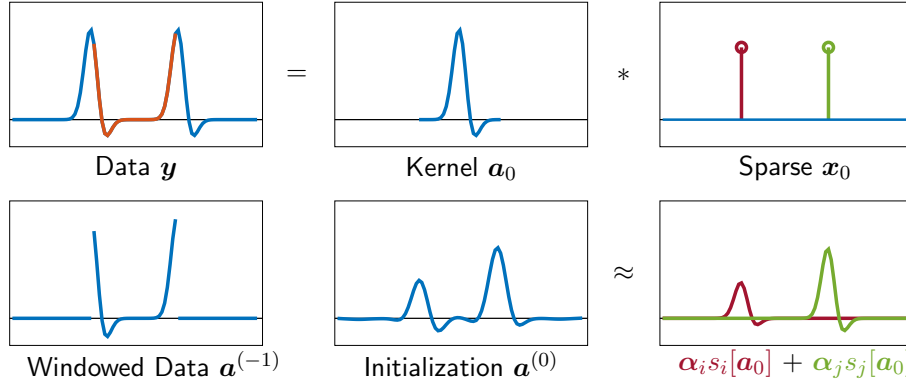


Figure 7. Data-driven initialization. Using a piece of the observed data $\mathbf{y}$ to generate an initial point $\mathbf{a}^{(0)}$ that is close to a superposition of shifts $s_\ell[\mathbf{a}_0]$ of the ground truth. Top: data $\mathbf{y} = \mathbf{a}_0 * \mathbf{x}_0$ is a superposition of shifts of the true kernel $\mathbf{a}_0$. Bottom: a length- p_0 window contains pieces of just a few shifts. Bottom middle: one step of the generalized power method approximately fills in the missing pieces, yielding a near superposition of shifts of $\mathbf{a}_0$ (right).

perform one step of the generalized power method,¹² writing

$$(3.7) \quad \mathbf{a}^{(0)} = -\mathbf{P}_{\mathbb{S}^{p-1}} \nabla \varphi_{\ell^1} \left(\mathbf{P}_{\mathbb{S}^{p-1}} [\mathbf{0}^{p_0-1}; \mathbf{y}_0; \dots; \mathbf{y}_{p_0-1}; \mathbf{0}^{p_0-1}] \right),$$

where $\mathbf{P}_{\mathbb{S}^{p-1}}$ projects onto the sphere. The reasoning behind this construction may seem obscure. We will explain it at a more technical level in section 5 after interpreting the gradient $\nabla \varphi_\rho$ in terms of its action on the shifts $s_\ell[\mathbf{a}_0]$ in section 4. For now, we note that this operation has the effect of (approximately) filling in the missing pieces of the truncated shifts $s_\ell[\mathbf{a}_0]$; see Figure 7 for an example. We will prove that with high probability $\mathbf{a}^{(0)}$ is indeed close to $\Sigma_{4\theta p_0}$.

The next key observation is that the function φ_ρ grows as we move away from the subspace $\mathcal{S}_\tau$; see Figure 8. Because of this, a small-stepping descent method will not move far away from $\Sigma_{4\theta p_0}$. For concreteness, we will analyze a variant of the curvilinear search method [23, 24], which moves in a linear combination of the negative gradient direction $-\mathbf{g}$ and a negative curvature direction $-\mathbf{v}$. At the k th iteration, the algorithm updates $\mathbf{a}^{(k+1)}$ as

$$(3.8) \quad \mathbf{a}^{(k+1)} \leftarrow \mathbf{P}_{\mathbb{S}^{p-1}} [\mathbf{a}^{(k)} - t\mathbf{g}^{(k)} - t^2\mathbf{v}^{(k)}]$$

with appropriately chosen step size t . The inclusion of a negative curvature direction allows the method to avoid stagnation near saddle points. Indeed, we will prove that, starting from initialization $\mathbf{a}^{(0)}$, this method produces a sequence $\mathbf{a}^{(1)}, \mathbf{a}^{(2)}, \dots$ which efficiently converges to a local minimizer $\bar{\mathbf{a}}$ that is near some signed shift $\pm s_\ell[\mathbf{a}_0]$ of the ground truth.

¹²The power method for minimizing a quadratic form $\xi(\mathbf{a}) = \frac{1}{2}\mathbf{a}^* \mathbf{M} \mathbf{a}$ over the sphere consists of the iteration $\mathbf{a} \mapsto -\mathbf{P}_{\mathbb{S}^{p-1}} \mathbf{M} \mathbf{a}$. Notice that in this mapping, $-\mathbf{M} \mathbf{a} = -\nabla \xi(\mathbf{a})$. The generalized power method, for minimizing a function φ over the sphere, consists of repeatedly projecting $-\nabla \varphi$ onto the sphere, giving the iteration $\mathbf{a} \mapsto -\mathbf{P}_{\mathbb{S}^{p-1}} \nabla \varphi(\mathbf{a})$. Equation (3.7) can be interpreted as one step of the generalized power method for the objective function φ_ρ .

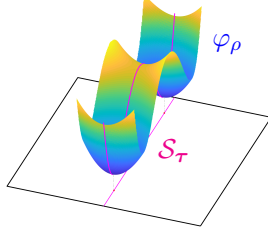


Figure 8. Growth of φ_ρ away from S_τ . Because φ_ρ grows away from S_τ , small-stepping descent methods stay near S_τ .

The second step of our algorithm *rounds* the local minimizer $\bar{\mathbf{a}} \approx \sigma s_\ell[\mathbf{a}_0]$ to produce an exact solution $\hat{\mathbf{a}} = \sigma s_\ell[\mathbf{a}_0]$. As a by-product, it also exactly recovers the corresponding signed shift of the true sparse signal, $\hat{\mathbf{x}} = \sigma s_{-\ell}[\mathbf{x}_0]$.

Our rounding algorithm is an alternating minimization scheme, which alternates between minimizing the Lasso cost over $\mathbf{a}$ with $\mathbf{x}$ fixed, and minimizing the Lasso cost over $\mathbf{x}$ with $\mathbf{a}$ fixed. We make two modifications to this basic idea, both of which are important for obtaining exact recovery. First, unlike the standard Lasso cost, which penalizes all of the entries of $\mathbf{x}$, we maintain a running estimate $I^{(k)}$ of the support of $\mathbf{x}_0$ and only penalize those entries that are not in $I^{(k)}$:

$$(3.9) \quad \frac{1}{2} \|\mathbf{a} * \mathbf{x} - \mathbf{y}\|_2^2 + \lambda \sum_{i \notin I^{(k)}} |\mathbf{x}_i|.$$

This can be viewed as an extreme form of *reweighting* [11]. Second, our algorithm gradually decreases penalty variable λ to 0, so that eventually

$$(3.10) \quad \hat{\mathbf{a}} * \hat{\mathbf{x}} \approx \mathbf{y}.$$

This can be viewed as a *homotopy* or *continuation* method [46, 19]. For concreteness, at the k th iteration the algorithm reads

$$(3.11) \quad \text{Update } \mathbf{x}: \quad \mathbf{x}^{(k+1)} \leftarrow \underset{\mathbf{x}}{\operatorname{argmin}} \frac{1}{2} \|\mathbf{a}^{(k)} * \mathbf{x} - \mathbf{y}\|_2^2 + \lambda^{(k)} \sum_{i \notin I^{(k)}} |\mathbf{x}_i|;$$

$$(3.12) \quad \text{Update } \mathbf{a}: \quad \mathbf{a}^{(k+1)} \leftarrow \underset{\mathbf{a}}{P_{\mathbb{S}^{p-1}}} \left[\underset{\mathbf{a}}{\operatorname{argmin}} \frac{1}{2} \|\mathbf{a} * \mathbf{x}^{(k+1)} - \mathbf{y}\|_2^2 \right];$$

$$(3.13) \quad \text{Update } \lambda \text{ and } I: \quad \lambda^{(k+1)} \leftarrow \frac{1}{2} \lambda^{(k)}, \quad I^{(k+1)} \leftarrow \operatorname{supp}(\mathbf{x}^{(k+1)}).$$

We prove that the iterates produced by this sequence of operations converge to the ground truth at a linear rate, as long as the initializer $\bar{\mathbf{a}}$ is sufficiently nearby.

Our overall algorithm is summarized as [Algorithm 3.1](#). [Figure 9](#) illustrates the main steps of this algorithm. Our main algorithmic result states that under essentially the same hypotheses as above, [Algorithm 3.1](#) produces a signed shift of the ground truth $(\mathbf{a}_0, \mathbf{x}_0)$.

Algorithm 3.1. Short-and-sparse deconvolution.

Input: Observation $\mathbf{y}$, motif length p_0 , sparsity θ , shift-coherence μ , and curvature threshold $-\eta_v$.

Minimization:

Set $\mathbf{a}^{(0)} \leftarrow -\mathbf{P}_{\mathbb{S}^{p-1}} \nabla \varphi_\rho(\mathbf{P}_{\mathbb{S}^{p-1}} [\mathbf{0}^{p_0-1}; \mathbf{y}_0; \dots; \mathbf{y}_{p_0-1}; \mathbf{0}^{p_0-1}])$.

Set $\lambda = 0.1/\sqrt{p_0\theta}$ ¹³ and $\delta > 0$ in φ_ρ . For $k = 1, 2, \dots, K_1$, let

$$(3.14) \quad \mathbf{a}^{(k+1)} \leftarrow \mathbf{P}_{\mathbb{S}^{p-1}} [\mathbf{a}^{(k)} - t\mathbf{g}^{(k)} - t^2\mathbf{v}^{(k)}],$$

where $\mathbf{g}^{(k)}$ is the Riemannian gradient; $\mathbf{v}^{(k)}$ is the eigenvector of smallest Riemannian Hessian eigenvalue if less than $-\eta_v$ with $\langle \mathbf{v}^{(k)}, \mathbf{g}^{(k)} \rangle \geq 0$, otherwise let $\mathbf{v}^{(k)} = \mathbf{0}$; and $t \in (0, 0.1/n\theta]$ satisfies

$$(3.15) \quad \varphi_\rho(\mathbf{a}^{(k+1)}) < \varphi_\rho(\mathbf{a}^{(k)}) - \frac{1}{2}t\|\mathbf{g}^{(k)}\|_2^2 - \frac{1}{4}t^4\eta_v\|\mathbf{v}^{(k)}\|_2^2$$

to obtain a near local minimizer $\bar{\mathbf{a}} \leftarrow \mathbf{a}^{(K_1)}$.

Refinement:

Set $\mathbf{a}^{(0)} \leftarrow \bar{\mathbf{a}}$, $\lambda^{(0)} \leftarrow 10(p\theta + \log n)(\mu + 1/p)$, and $I^{(0)} \leftarrow \mathcal{S}_{\lambda^{(0)}}[\text{supp}(\check{\mathbf{y}} * \bar{\mathbf{a}})]$. For $k = 1, 2, \dots, K_2$, let

$$(3.16) \quad \mathbf{x}^{(k+1)} \leftarrow \arg\min_{\mathbf{x}} \frac{1}{2}\|\mathbf{a}^{(k)} * \mathbf{x} - \mathbf{y}\|_2^2 + \lambda^{(k)} \sum_{i \notin I^{(k)}} |\mathbf{x}_i|,$$

$$(3.17) \quad \mathbf{a}^{(k+1)} \leftarrow \mathbf{P}_{\mathbb{S}^{p-1}} \left[\arg\min_{\mathbf{a}} \frac{1}{2}\|\mathbf{a} * \mathbf{x}^{(k+1)} - \mathbf{y}\|_2^2 \right],$$

$$(3.18) \quad \lambda^{(k+1)} \leftarrow \lambda^{(k)}/2, \quad I^{(k+1)} \leftarrow \text{supp}(\mathbf{x}^{(k+1)}),$$

to obtain $(\hat{\mathbf{a}}, \hat{\mathbf{x}}) \leftarrow (\mathbf{a}^{(K_2)}, \mathbf{x}^{(K_2)})$.

Output: Return $(\hat{\mathbf{a}}, \hat{\mathbf{x}})$.

Theorem 3.2 (main algorithmic theorem). Suppose $\mathbf{y} = \mathbf{a}_0 * \mathbf{x}_0$, where $\mathbf{a}_0 \in \mathbb{S}^{p_0-1}$ is μ -truncated shift coherent¹⁴ such that $\max_{i \neq j} |\langle \mathbf{u}_{p_0}^* s_i[\mathbf{a}_0], \mathbf{u}_{p_0}^* s_j[\mathbf{a}_0] \rangle| \leq \mu$ and $\mathbf{x}_0 \sim_{\text{i.i.d.}} \text{BG}(\theta) \in \mathbb{R}^n$ with θ, μ satisfying

$$(3.19) \quad \theta \in \left[\frac{c_1}{p_0}, \frac{c_2}{(p_0\sqrt{\mu} + \sqrt{p_0}) \log^2 p_0} \right], \quad \mu \leq \frac{c_3}{\log^2 n}$$

for some constant $c_1, c_2, c_3 > 0$. If the signal lengths n, p_0 satisfy $n > \text{poly}(p_0)$ and $p_0 > \text{polylog}(n)$, then there exist $\delta, \eta_v > 0$ such that with high probability Algorithm 3.1 produces $(\hat{\mathbf{a}}, \hat{\mathbf{x}})$ that are equal to the ground truth up to signed shift symmetry:

$$(3.20) \quad \|(\hat{\mathbf{a}}, \hat{\mathbf{x}}) - \sigma(s_\ell[\mathbf{a}_0], s_{-\ell}[\mathbf{x}_0])\|_2 \leq \varepsilon$$

for $\sigma \in \{\pm 1\}$ and $\ell \in \{-p_0 + 1, \dots, p_0 - 1\}$ if $K_1 > \text{poly}(n, p_0)$ and $K_2 > \text{polylog}(n, p_0, \varepsilon^{-1})$.

Proof. See Theorems 5.1 and 5.2. ■

¹³In practice, we suggest setting $\lambda = c_\lambda/\sqrt{p_0\theta}$ with $c_\lambda \in [0.5, 0.8]$.

¹⁴The truncated shift coherence is a stronger condition than natural shift coherence. The statement appears mainly due to the limitation of the proof strategy for the algorithm.

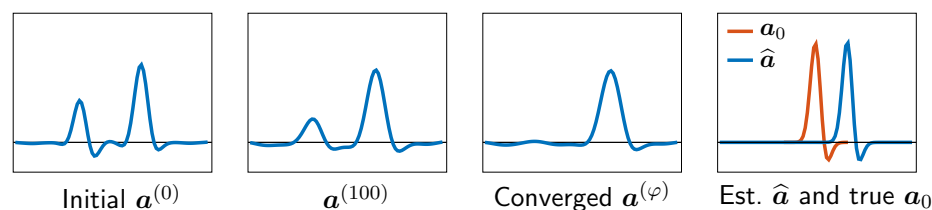


Figure 9. Local minimization and refinement. *Left: data-driven initialization $\mathbf{a}^{(0)}$ consisting of a near superposition of two shifts. Middle: minimizing φ_ρ produces a near shift of $\mathbf{a}_0$. Right: rounded solution $\hat{\mathbf{a}}$ using the Lasso. $\hat{\mathbf{a}}$ is very close to a shift of $\mathbf{a}_0$.*

When solving SaS deconvolution via minimizing Bilinear Lasso objective (2.2) in practice, the algorithm is analogous to the provable method introduced in Algorithm 3.1, where the curvilinear descent and the refinement step can be realized as alternating gradient descent of both variables $\mathbf{a}, \mathbf{x}$ in (2.2). Unlike Algorithm 3.1, this alternating gradient method has yet to come with theoretical guarantees, but has shown to be an effective and efficient method for SaS deconvolution problems both in simulation and in reality [34].

3.3. Relationship to the literature. Blind deconvolution is a classical problem in signal processing [54, 12] and has been studied under a variety of hypotheses. In this section, we first discuss the relationship between our results and the existing literature on the SaS version of this problem, and then briefly discuss other deconvolution variants in the theoretical literature.

The SaS model arises in a number of applications. One class of applications involves finding basic motifs (repeated patterns) in datasets. This *motif discovery* problem arises in extracellular spike sorting [37, 20] and calcium imaging [48], where the observed signal exhibits repetitive *short* neuron excitation patterns occurring *sparingly* across time and/or space. Similarly, electron microscopy images [15] arising in study of nanomaterials often exhibit repeated motifs.

Another significant application of SaS deconvolution is *image deblurring*. Typically, the blur kernel is small relative to the image size (*short*) [3, 62, 13, 35, 36]. In natural image deblurring, the target image is often assumed to have relatively few sharp edges [21, 27, 36], and hence have *sparse* derivatives. In scientific image deblurring, e.g., in astronomy [33, 25, 9] and geophysics [28], the target image is often sparse, either in the spatial or wavelet domains, again leading to variants of the SaS model. The literature on blind image deconvolution is large; see, e.g., [31, 10] for surveys.

Variants of the SaS deconvolution problem arise in many other areas of engineering as well. Examples include *blind equalization* in communications [50, 51, 26], *dereverberation* in sound engineering [44, 45], and image *superresolution* [4, 53, 61].

These applications have motivated a great deal of algorithmic work on variants of the SaS problem [32, 8, 6, 31, 43, 10, 56]. In contrast, relatively little theory is available to explain when and why algorithms succeed. Our algorithm minimizes φ_ρ as an approximation to the Lasso cost over the sphere. Our formulation and results have strong precedent in the literature. Lasso-like objective functions have been widely used in image deblurring [62, 14, 21, 35, 52, 60, 18, 30, 36, 59, 47, 64]. A number of insights have been obtained into the

geometry of sparse deconvolution—in particular, into the effect of various constraints on $\mathbf{a}$ on the presence or absence of spurious local minimizers. In image deblurring, a simplex constraint ($\mathbf{a} \geq \mathbf{0}$ and $\|\mathbf{a}\|_1 = 1$) arises naturally from the physical structure of the problem [62, 14]. Perhaps surprisingly, simplex-constrained deconvolution admits trivial global minimizers, at which the recovered kernel $\mathbf{a}$ is a spike, rather than the target blur kernel [7, 36].

The work [59] imposes the ℓ^2 regularization on $\mathbf{a}$ and observes that this alternative constraint gives a more reliable algorithm. In [64], the geometry of the simplified objective φ_{ℓ^1} over the sphere is studied, and it is proved that in the dilute limit in which $\mathbf{x}_0$ has one nonzero entry, all strict local minima of φ_{ℓ^1} are close to signed shifts truncations of $\mathbf{a}_0$. By adopting a different objective function (based on ℓ^4 maximization) over the sphere, [63] proves that on a certain region of the sphere every local minimum is near a *truncated* signed shift of $\mathbf{a}_0$, i.e., the restriction of $s_\ell[\mathbf{a}_0]$ to the window $\{0, \dots, p_0 - 1\}$. The analysis of [63] allows the sparse sequence $\mathbf{x}_0$ to be denser ($\theta \sim p_0^{-2/3}$ for a generic kernel $\mathbf{a}_0$, as opposed to $\theta \lesssim p_0^{-3/4}$ in our result). Both [64] and [63] guarantee *approximate* recovery of a portion of $s_\ell[\mathbf{a}_0]$, under complicated conditions on the kernel $\mathbf{a}_0$. Our core optimization problem is very similar to that of [64]. However, we obtain *exact* recovery of both $\mathbf{a}_0$ and relatively dense $\mathbf{x}_0$ under the much simpler assumption of shift incoherence.

Other aspects of the SaS problem have been studied theoretically. One basic question is under what circumstances the problem is identifiable up to the scaled shift ambiguity. The paper [17] shows that the problem is ill-posed for worst case $(\mathbf{a}_0, \mathbf{x}_0)$, particularly for certain support patterns in which $\mathbf{x}_0$ does not have any isolated nonzero entries. This demonstrates that *some* modeling assumptions on the support of the sparse term are needed. At the same time, this worst-case structure is unlikely to occur, either under the Bernoulli model or in practical deconvolution problems.

Motivated by a variety of applications, much research has focused on low-dimensional deconvolution models in the theoretical literature. In communication applications, the signals $\mathbf{a}_0$ and $\mathbf{x}_0$ either live in known low-dimensional subspaces or are sparse in some known dictionary [2, 16, 29, 39, 40, 41, 42]. These theoretical works assume that the subspace/dictionary are chosen at random. This low-dimensional deconvolution model does not exhibit the signed shift ambiguity; nonconvex formulations for this model exhibit a different structure from that studied here. In fact, the variant in which both signals belong to known subspaces can be solved by convex relaxation [2]. The SaS model does not appear to be amenable to convexification and exhibits a more complicated nonconvex geometry due to the shift ambiguity. The main motivation for tackling this model lies in the aforementioned applications in imaging and data analysis.

In [38, 57] the related *multi-instance* sparse blind deconvolution problem (MISBD) is studied, where there are K observations $\mathbf{y}_i = \mathbf{a}_0 * \mathbf{x}_i$ consisting of multiple convolutions $i = 1, \dots, K$ of a kernel $\mathbf{a}_0$ and different sparse vectors $\mathbf{x}_i$. Both works develop provable algorithms. There are several key differences with our work. First, both the proposed algorithms and their analysis require the kernel to be invertible. Second, despite the apparent similarity between the SaS model and MISBD, these problems are not equivalent. It might seem possible to reduce SaS to MISBD by dividing the single observation $\mathbf{y}$ into K pieces; this apparent reduction fails due to boundary effects.

3.4. Notation. All vectors/matrices are written in bold font, $\mathbf{a}/\mathbf{A}$; indexed values are written as $\mathbf{a}_i, \mathbf{A}_{ij}$. Zero or one vectors are defined as $\mathbf{0}$ or $\mathbf{1}$, and i th canonical basis vector defined as $\mathbf{e}_i$. The indices for vectors/matrices all start from 0 and are taken modulo- n , and thus a vector of length n should have its indices labeled as $\{0, 1, \dots, n-1\}$. We write $[n] = \{0, \dots, n-1\}$. We often use the capital italic symbols I, J for subsets of $[n]$. We abuse notation slightly and write $[-p] = \{n-p+1, \dots, n-1, 0\}$ and $[\pm p] = \{n-p+1, \dots, n-1, 0, 1, \dots, p-1\}$. Index sets can be labels for vectors; $\mathbf{a}_I \in \mathbb{R}^{|I|}$ denotes the restriction of the vector $\mathbf{a}$ to coordinates I . Also, we use a check symbol to denote the reversal operator on index set $\check{I} = -I$ and vectors $\check{\mathbf{a}}_i = \mathbf{a}_{-i}$.

We let $\mathbf{P}_C$ denote the projection operator associated with a compact set C . The zero-filling operator $\iota_I : \mathbb{R}^{|I|} \rightarrow \mathbb{R}^n$ injects the input vector to higher dimensional Euclidean space, via $(\iota_I \mathbf{x})_i = \mathbf{x}_{I^{-1}(i)}$ for $i \in I$, and 0 otherwise. Its adjoint operator ι_I^* can be understood as a subset selection operator which picks up entries of coordinates I . A common zero-filling operator throughout this paper, ι , is an abbreviation of $\iota_{[p]}$, which is often addressed as the zero-padding operator and its adjoint ι^* as truncation operator.

The convolution operators are all circular with modulo- n : $(\mathbf{a} * \mathbf{x})_i = \sum_{j \in [n]} \mathbf{a}_j \mathbf{x}_{i-j}$; also, the convolution operator works on the index set: $I * J = \text{supp}(\mathbf{1}_I * \mathbf{1}_J)$. Similarly, the shift operator $s_\ell[\cdot] : \mathbb{R}^p \rightarrow \mathbb{R}^n$ is circular with modulo- n without specification: $(s_\ell[\mathbf{a}])_j = (\iota_{[p]} \mathbf{a})_{j-\ell}$. Notice that here $\mathbf{a}$ can be shorter, $p \leq n$. Let $\mathbf{C}_\mathbf{a} \in \mathbb{R}^{n \times n}$ denote a circulant matrix (with modulo- n) for vector $\mathbf{a}$, whose j th column is the cyclic shift of $\mathbf{a}$ by j : $\mathbf{C}_\mathbf{a} \mathbf{e}_j = s_j[\mathbf{a}]$. It satisfies for any $\mathbf{b} \in \mathbb{R}^n$,

$$(3.21) \quad \mathbf{C}_\mathbf{a} \mathbf{b} = \mathbf{a} * \mathbf{b}.$$

The correlation between $\mathbf{a}$ and $\mathbf{b}$ can be also written in similar form of convolution operators which reverse one vector before convolution. Define two correlation matrices $\mathbf{C}_\mathbf{a}^*$ and $\check{\mathbf{C}}_\mathbf{a}$ as $\mathbf{C}_\mathbf{a}^* \mathbf{e}_j = s_j[\check{\mathbf{a}}]$ and $\check{\mathbf{C}}_\mathbf{a} \mathbf{e}_j = s_{-j}[\mathbf{a}]$. The two operators will satisfy

$$(3.22) \quad \mathbf{C}_\mathbf{a}^* \mathbf{b} = \check{\mathbf{a}} * \mathbf{b}, \quad \check{\mathbf{C}}_\mathbf{a} \mathbf{b} = \mathbf{a} * \check{\mathbf{b}}.$$

4. Geometry of φ_ρ in shift space. Underlying our main geometric and algorithmic results is a relationship between the geometry of the function φ_ρ and the symmetries of the deconvolution problem. In this section, we describe this relationship at a more technical level by interpreting the gradient and Hessian of the function φ_ρ in terms of the shifts $s_\ell[\mathbf{a}_0]$ and stating a key lemma which asserts that a certain neighborhood of the union of subspaces $\Sigma_{4\theta p_0}$ can be decomposed into regions of negative curvature, strong gradient, and strong convexity near the target solutions $\pm s_\ell[\mathbf{a}_0]$.

4.1. Shifts and correlations. The set $\Sigma_{4\theta p_0}$ is a union of subspaces. Any point $\mathbf{a}$ in one of these subspaces $\mathcal{S}_\tau$ is a superposition of shifts of $\mathbf{a}_0$:

$$(4.1) \quad \mathbf{a} = \sum_{\ell \in \tau} \alpha_\ell s_\ell[\mathbf{a}_0].$$

This representation can be extended to a general point $\mathbf{a} \in \mathbb{S}^{p-1}$ by writing

$$(4.2) \quad \mathbf{a} = \sum_{\ell \in \tau} \alpha_\ell s_\ell[\mathbf{a}_0] + \sum_{\ell \notin \tau} \alpha_\ell s_\ell[\mathbf{a}_0].$$

The vector α can be viewed as the coefficients of a decomposition of $\mathbf{a}$ into different shifts of $\mathbf{a}_0$. This representation is not unique. For $\mathbf{a}$ close to $\mathcal{S}_\tau$, we can choose a particular α for which α_{τ^c} is small, a notion that we will formalize below.

For convenience, we introduce a closely related vector $\beta \in \mathbb{R}^n$, whose entries are the inner products between $\mathbf{a}$ and the shifts of $\mathbf{a}_0$: $\beta_\ell = \langle \mathbf{a}, s_\ell[\mathbf{a}_0] \rangle$. Since the columns of $\mathbf{C}_{\mathbf{a}_0}$ are the shifts of $\mathbf{a}_0$, we can write

$$\begin{aligned} (4.3) \quad \beta &= \mathbf{C}_{\mathbf{a}_0}^* \iota \mathbf{a} \\ (4.4) \quad &= \mathbf{C}_{\mathbf{a}_0}^* \iota^* \mathbf{C}_{\mathbf{a}_0} \alpha =: \mathbf{M} \alpha. \end{aligned}$$

The matrix $\mathbf{M}$ is the Gram matrix of the truncated shifts: $M_{ij} = \langle \iota^* s_i[\mathbf{a}_0], \iota^* s_j[\mathbf{a}_0] \rangle$. When μ is small, the off-diagonal elements of $\mathbf{M}$ are small. In particular, on $\mathcal{S}_\tau$ we may take $\alpha_{\tau^c} = \mathbf{0}$, and $\beta \approx \alpha$ in the sense that $\beta_\tau \approx \alpha_\tau$ and the entries of β_{τ^c} are small. For detailed elaboration, see [section SM2](#) in the supplementary material.

4.2. Shifts and the calculus of φ_{ℓ^1} . Our main geometric claims pertain to the function φ_ρ , which is based on a smooth sparsity surrogate $\rho(\cdot) \approx \|\cdot\|_1$. In this section, we sketch the main ideas of the proof as if $\rho(\cdot) = \|\cdot\|_1$ by relating the geometry of the function φ_{ℓ^1} to the vectors α, β introduced above. Working with φ_{ℓ^1} simplifies the exposition; it is also faithful to the structure of our proof, which relates the derivatives of the smooth function φ_ρ to similar quantities associated with the nonsmooth function φ_{ℓ^1} .

The function φ_{ℓ^1} has a relatively simple closed form:

$$(4.5) \quad \varphi_{\ell^1}(\mathbf{a}) = -\frac{1}{2} \|\mathcal{S}_\lambda[\check{\mathbf{y}} * \mathbf{a}]\|_2^2.$$

Here, $\mathcal{S}_\lambda$ is the *soft thresholding operator*, which is defined for scalars t as

$$(4.6) \quad \mathcal{S}_\lambda[t] = \text{sign}(t) \max\{|t| - \lambda, 0\}$$

and is extended to vectors by applying it elementwise. The operator $\mathcal{S}_\lambda[\mathbf{x}]$ shrinks the elements of $\mathbf{x}$ toward zero. Small elements become identically zero, resulting in a sparse vector.

Gradient: Sparsifying the correlations β . Our goal is to understand the local minimizers of the function φ_{ℓ^1} over the sphere. The function φ_{ℓ^1} is differentiable. Clearly, any point $\mathbf{a}$ at which its gradient (over the sphere) is nonzero cannot be a local minimizer. We first give an expression for the gradient of φ_{ℓ^1} over Euclidean space $\mathbb{R}^p$, and then extend it to the sphere $\mathbb{S}^{p-1}$. Using $\mathbf{y} = \mathbf{a}_0 * \mathbf{x}_0$ and calculus gives

$$\begin{aligned} (4.7) \quad \nabla \varphi_{\ell^1}(\mathbf{a}) &= -\iota^* \mathbf{C}_{\mathbf{a}_0} \check{\mathbf{C}}_{\mathbf{x}_0} \mathcal{S}_\lambda[\check{\mathbf{C}}_{\mathbf{x}_0} \mathbf{C}_{\mathbf{a}_0}^* \iota \mathbf{a}] \\ &= -\iota^* \mathbf{C}_{\mathbf{a}_0} \check{\mathbf{C}}_{\mathbf{x}_0} \mathcal{S}_\lambda[\check{\mathbf{C}}_{\mathbf{x}_0} \beta] \\ &= -\iota^* \mathbf{C}_{\mathbf{a}_0} \chi[\beta], \end{aligned}$$

where we have simplified the notation by introducing an operator $\chi: \mathbb{R}^n \rightarrow \mathbb{R}^n$ as $\chi[\beta] = \check{\mathbf{C}}_{\mathbf{x}_0} \mathcal{S}_\lambda[\check{\mathbf{C}}_{\mathbf{x}_0} \beta]$. This representation exhibits the (negative) gradient as a superposition of shifts of $\mathbf{a}_0$ with coefficients given by the entries of $\chi[\beta]$:

$$(4.8) \quad -\nabla \varphi_{\ell^1}(\mathbf{a}) = \sum_{\ell} \chi[\beta]_{\ell} s_{\ell}[\mathbf{a}_0].$$

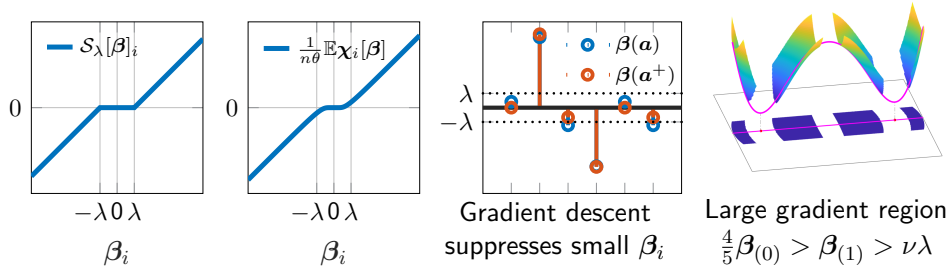


Figure 10. Gradient sparsifies correlations. *Left: the soft thresholding operator $S_\lambda[\beta]$ shrinks the entries of β toward zero, making it sparser. Middle left: the negative gradient $-\nabla\varphi_{\ell^1}$ is a superposition of shifts $s_\ell[\mathbf{a}_0]$, with coefficients $\chi_\ell[\beta] \approx S_\lambda[\beta]_\ell$. Because of this, gradient descent sparsifies β . Middle right: $\beta(\mathbf{a})$ before, and $\beta(\mathbf{a}^+)$ after, one projected gradient step $\mathbf{a}^+ = \mathbf{P}_{\mathbb{S}^{p-1}}[\mathbf{a} - t \cdot \text{grad}[\varphi_{\ell^1}](\mathbf{a})]$. Notice that the small entries of β are shrunk towards zero. Right: the gradient $\text{grad}[\varphi_{\ell^1}](\mathbf{a})$ is large whenever it is easy to sparsify β , particularly when the largest entry $\beta_{(0)} \gg \beta_{(1)} \gg 0$.*

The operator χ appears complicated. However, its effect is relatively simple: when $\mathbf{x}_0$ is a long random vector, $\chi[\beta]$ acts like a soft thresholding operator on the vector β . That is,

$$(4.9) \quad \frac{1}{n\theta} \cdot \chi[\beta]_\ell \approx \begin{cases} \beta_\ell - \lambda, & \beta_\ell > \lambda, \\ \beta_\ell + \lambda, & \beta_\ell < -\lambda, \\ 0 & \text{otherwise.} \end{cases}$$

We show this rigorously below in the proof of our main theorems. Here, we support this claim pictorially by plotting the ℓ th entry $\chi[\beta]_\ell$ as β_ℓ varies; see Figure 10 (middle left) and compare to Figure 10 (left). Because $\chi[\beta]$ suppresses small entries of β , the strongest contributions to $-\nabla\varphi_{\ell^1}$ in (4.8) will come from shifts $s_\ell[\mathbf{a}_0]$ with large β_ℓ . In particular, the Euclidean gradient is large whenever there is a single preferred shift $s_\ell[\mathbf{a}_0]$, i.e., the largest entry of β is significantly larger than the second largest entry.

The (Euclidean) gradient $\nabla\varphi_{\ell^1}$ measures the slope of φ_{ℓ^1} over $\mathbb{R}^n$. We are interested in the slope of φ_{ℓ^1} over the sphere $\mathbb{S}^{p-1}$, which is measured by the Riemannian gradient

$$(4.10) \quad \begin{aligned} \text{grad}[\varphi_{\ell^1}](\mathbf{a}) &= \mathbf{P}_{\mathbf{a}^\perp} \nabla\varphi_{\ell^1}(\mathbf{a}) \\ &= -\mathbf{P}_{\mathbf{a}^\perp} \sum_{\ell} \chi_\ell[\beta] s_\ell[\mathbf{a}_0]. \end{aligned}$$

The Riemannian gradient simply projects the Euclidean gradient onto the tangent space $\mathbf{a}^\perp$ to $\mathbb{S}^{p-1}$ at $\mathbf{a}$. The Riemannian gradient is large whenever

- (i) *the negative gradient points to one particular shift:* there is a single preferred shift $s_\ell[\mathbf{a}_0]$ so that the Euclidean gradient is large; and
- (ii) *$\mathbf{a}$ is not too close to any shift:* it is possible to move in the tangent space in the direction of this shift.¹⁵ Since the tangent space consists of those vectors orthogonal to $\mathbf{a}$, this is possible whenever $s_\ell[\mathbf{a}_0]$ is not too aligned with $\mathbf{a}$, i.e., $\mathbf{a}$ is not too close to $s_\ell[\mathbf{a}_0]$.

¹⁵...so the projection of the Euclidean gradient onto the tangent space does not vanish.

Our technical lemma quantifies this situation in terms of the ordered entries of β . Write $|\beta_{(0)}| \geq |\beta_{(1)}| \geq \dots$, with corresponding shifts $s_{(0)}[\mathbf{a}_0], s_{(1)}[\mathbf{a}_0], \dots$. There is a strong gradient whenever $|\beta_{(0)}|$ is significantly larger than $|\beta_{(1)}|$ and $|\beta_{(1)}|$ is not too small compared to λ : in particular, when $\frac{4}{5}|\beta_{(0)}| > |\beta_{(1)}| > \frac{\lambda}{4 \log^2 \theta - 1}$. In this situation, gradient descent drives $\mathbf{a}$ toward $s_{(0)}[\mathbf{a}_0]$, reducing $|\beta_{(1)}|, \dots$, and making the vector β sparser. We establish the technical claim that the (Euclidean) gradient of φ_{ℓ^1} sparsifies vectors in shift space in [section SM3](#).

Hessian: Negative curvature breaks symmetry. When there is no single preferred shift, i.e., when $|\beta_{(1)}|$ is close to $|\beta_{(0)}|$, the gradient can be small. Similarly, when $\mathbf{a}$ is very close to $\pm s_{(0)}[\mathbf{a}_0]$, the gradient can be small. In either of these situations, we need to study the curvature of the function φ to determine whether there are local minimizers.

Strictly speaking, the function φ_{ℓ^1} is not twice differentiable, due to the nonsmoothness of the soft thresholding operator $\mathcal{S}_\lambda[t]$ at $t = \pm\lambda$. Indeed, φ_{ℓ^1} is nonsmooth at any point $\mathbf{a}$ for which some entry of $\tilde{\mathbf{y}} * \mathbf{a}$ has magnitude λ . At other points $\mathbf{a}$, φ_{ℓ^1} is twice differentiable, and its Hessian is given by

$$(4.11) \quad \tilde{\nabla}^2 \varphi_{\ell^1}(\mathbf{a}) = -\iota^* C_{\mathbf{a}_0} \tilde{C}_{x_0} P_I \tilde{C}_{x_0}^* C_{\mathbf{a}_0}^* \iota,$$

with $I = \text{supp}(\mathcal{S}_\lambda[\tilde{C}_{\mathbf{y}} \iota \mathbf{a}])$. We (formally) extend this expression to *every* $\mathbf{a} \in \mathbb{R}^n$, terming $\tilde{\nabla}^2 \varphi_{\ell^1}$ the *pseudo-Hessian* of φ_{ℓ^1} . For appropriately chosen smooth sparsity surrogate ρ , we will see that the (true) Hessian of the smooth function $\nabla^2 \varphi_\rho$ is close to $\tilde{\nabla}^2 \varphi_{\ell^1}$, and so $\tilde{\nabla}^2 \varphi_{\ell^1}$ yields useful information about the curvature of φ_ρ .

As with the gradient, the Hessian is complicated, but becomes simpler when the sample size is large. The approximation

$$(4.12) \quad \tilde{\nabla}^2 \varphi_{\ell^1}(\mathbf{a}) \approx - \sum_{\ell} s_{\ell}[\mathbf{a}_0] s_{\ell}[\mathbf{a}_0]^* \left(\frac{\partial}{\partial \beta_{\ell}} \chi_{\ell}[\beta] \right)$$

can be obtained from [\(4.8\)](#) by noting that $\frac{\partial}{\partial \mathbf{a}} \chi_{\ell}[\beta] = \sum_j s_j[\mathbf{a}_0] \frac{\partial}{\partial \beta_j} \chi_{\ell}[\beta]$, that $\frac{\partial}{\partial \beta_j} \chi_{\ell}[\beta] \approx 0$ for $j \neq \ell$, and that

$$(4.13) \quad \frac{1}{n\theta} \cdot \frac{\partial \chi_{\ell}[\beta]}{\partial \beta_{\ell}} \approx \begin{cases} 0, & |\beta_{\ell}| \ll \lambda, \\ 1, & |\beta_{\ell}| \gg \lambda. \end{cases}$$

Again, we corroborate this approximation pictorially; see [Figure 11](#).

From this approximation, we can see that the quadratic form $\mathbf{v}^* \tilde{\nabla}^2 \varphi_{\ell^1} \mathbf{v}$ takes on a large negative value whenever $\mathbf{v}$ is a shift $s_{\ell}[\mathbf{a}_0]$ corresponding to some $|\beta_{\ell}| \geq \lambda$, or whenever $\mathbf{v}$ is a linear combination of such shifts. *In particular, if for some j , $|\beta_{(0)}|, |\beta_{(1)}|, \dots, |\beta_{(j)}| \gg \lambda$, then φ_{ℓ^1} will exhibit negative curvature in any direction $\mathbf{v} \in \text{span}(s_{(0)}[\mathbf{a}_0], s_{(1)}[\mathbf{a}_0], \dots, s_{(j)}[\mathbf{a}_0])$.*

The (Euclidean) Hessian measures the curvature of the function φ_{ℓ^1} over $\mathbb{R}^n$. The Riemannian Hessian

$$(4.14) \quad \widetilde{\text{Hess}}[\varphi_{\ell^1}](\mathbf{a}) = P_{\mathbf{a}^\perp} \left(\begin{array}{c} \tilde{\nabla}^2 \varphi_{\ell^1}(\mathbf{a}) \\ \text{Curvature of } \varphi_{\ell^1} \end{array} + \begin{array}{c} \langle -\nabla \varphi_{\ell^1}(\mathbf{a}), \mathbf{a} \rangle \cdot I \\ \text{Curvature of the sphere} \end{array} \right) P_{\mathbf{a}^\perp}$$

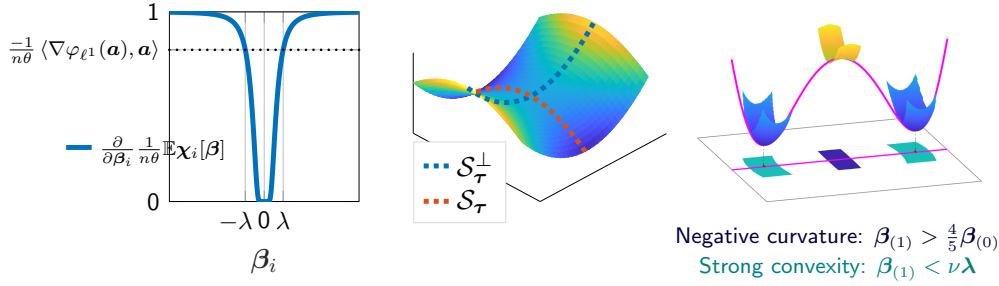


Figure 11. Hessian breaks symmetry. *Left:* contribution of $-s_i[\mathbf{a}_0]s_i[\mathbf{a}_0]^*$ to the Euclidean Hessian. If $|\beta_i| \gg \lambda$, the Euclidean Hessian exhibits a strong negative component in the $s_i[\mathbf{a}_0]$ direction. The Riemannian Hessian exhibits negative curvature in directions spanned by $s_i[\mathbf{a}_0]$ with corresponding $|\beta_i| \gg \lambda$ and positive curvature in directions spanned by $s_i[\mathbf{a}_0]$ with $|\beta_i| \ll \lambda$. *Middle:* this creates negative curvature along the subspace $\mathcal{S}_\tau$ and positive curvature orthogonal to this subspace. *Right:* our analysis shows that there is always a direction of negative curvature when $\beta_{(1)} > \frac{4}{5}\beta_{(0)}$; conversely, when $\beta_{(1)} \ll \lambda$ there is positive curvature in every feasible direction and the function is strongly convex.

measures the curvature of φ_{ℓ^1} over the sphere. The projection $\mathbf{P}_{\mathbf{a}^\perp}$ restricts its action to directions $\mathbf{v} \perp \mathbf{a}$ that are tangent to the sphere. The additional term $\langle -\nabla \varphi_{\ell^1}(\mathbf{a}), \mathbf{a} \rangle$ accounts for the curvature of the sphere. This term is always positive. The net effect is that directions of strong negative curvature of φ_{ℓ^1} over $\mathbb{R}^n$ become directions of moderate negative curvature over the sphere. Directions of nearly zero curvature over $\mathbb{R}^n$ become directions of positive curvature over the sphere. This has three implications for the geometry of φ_{ℓ^1} over the sphere:

- (i) *Negative curvature in symmetry breaking directions:* If $|\beta_{(0)}|, |\beta_{(1)}|, \dots, |\beta_{(j)}| \gg \lambda$, then φ_{ℓ^1} will exhibit negative curvature in any tangent direction $\mathbf{v} \perp \mathbf{a}$ which is in the linear span

$$\text{span}(s_{(0)}[\mathbf{a}_0], s_{(1)}[\mathbf{a}_0], \dots, s_{(j)}[\mathbf{a}_0])$$

of the corresponding shifts of $\mathbf{a}_0$.

- (ii) *Positive curvature in directions away from $\mathcal{S}_\tau$:* The Euclidean Hessian quadratic form $\mathbf{v}^* \tilde{\nabla}^2 \varphi_{\ell^1} \mathbf{v}$ takes on relatively small values in directions orthogonal to the subspace $\mathcal{S}_\tau$. The Riemannian Hessian is positive in these directions, creating positive curvature orthogonal to the subspace $\mathcal{S}_\tau$.
- (iii) *Strong convexity around minimizers:* Around a minimizer $s_\ell[\mathbf{a}_0]$, only a single entry β_ℓ is large. Any tangent direction $\mathbf{v} \perp \mathbf{a}$ is nearly orthogonal to the subspace $\text{span}(s_\ell[\mathbf{a}_0])$, and hence is a direction of positive (Riemannian) curvature. The objective function φ_ρ is strongly convex around the target solutions $\pm s_\ell[\mathbf{a}_0]$.

Figure 11 visualizes these regions of negative and positive curvatures, and the technical claim of positivity/negativity of curvature in shift space is presented in detail in [section SM4](#).

4.3. Any local minimizer is a near shift. We close this section by stating a key theorem, which makes the above discussion precise. We will show that a certain neighborhood of any subspace $\mathcal{S}_\tau$ can be covered by regions of *negative curvature*, of *large gradient*, and of *strong convexity* containing target solutions $\pm s_\ell[\mathbf{a}_0]$. Furthermore, at the boundary of this neighborhood, the negative gradient points back—*retracts*—toward the subspace $\mathcal{S}_\tau$, due to

the (directional) convexity of φ_ρ away from the subspace.

To formally state the result, we need a way of measuring how close $\mathbf{a}$ is to the subspace $\mathcal{S}_\tau$. For technical reasons, it turns out to be convenient to do this in terms of the coefficients $\boldsymbol{\alpha}$ in the representation

$$(4.15) \quad \mathbf{a} = \sum_{\ell \in \tau} \alpha_\ell s_\ell[\mathbf{a}_0] + \sum_{\ell' \in \tau^c} \alpha_{\ell'} s_{\ell'}[\mathbf{a}_0].$$

If $\mathbf{a} \in \mathcal{S}_\tau$, we can take $\boldsymbol{\alpha}$ with $\alpha_{\tau^c} = \mathbf{0}$. We can view the energy $\|\alpha_{\tau^c}\|_2$ as a measure of the distance from $\mathbf{a}$ to $\mathcal{S}_\tau$. A technical wrinkle arises, because the representation (4.15) is not unique. We resolve this issue by choosing the $\boldsymbol{\alpha}$ that minimizes $\|\alpha_{\tau^c}\|_2$, writing

$$(4.16) \quad d_\alpha(\mathbf{a}, \mathcal{S}_\tau) = \inf \left\{ \|\alpha_{\tau^c}\|_2 : \sum_{\ell} \alpha_\ell s_\ell[\mathbf{a}_0] = \mathbf{a} \right\}.$$

The distance $d_\alpha(\mathbf{a}, \mathcal{S}_\tau)$ is zero for $\mathbf{a} \in \mathcal{S}_\tau$. Our analysis controls the geometric properties of φ_ρ over the set of $\mathbf{a}$ for which $d_\alpha(\mathbf{a}, \mathcal{S}_\tau)$ is not too large. Similar to (3.3), we define an object which contains all points that are close to some $\mathcal{S}_\tau$ in the above sense:

$$(4.17) \quad \Sigma_{4\theta p_0}^\gamma := \bigcup_{|\tau| \leq 4\theta p_0} \{\mathbf{a} : d_\alpha(\mathbf{a}, \mathcal{S}_\tau) \leq \gamma\}.$$

The aforementioned geometric properties hold over this set.

Theorem 4.1 (geometry of φ_ρ over union of subspaces). *Suppose that $\mathbf{y} = \mathbf{a}_0 * \mathbf{x}_0$, where $\mathbf{a}_0 \in \mathbb{S}^{p_0-1}$ is μ -shift coherent and $\mathbf{x}_0 \sim_{\text{i.i.d.}} \text{BG}(\theta) \in \mathbb{R}^n$ satisfying*

$$(4.18) \quad \theta \in \left[\frac{c'}{p_0}, \frac{c}{p_0 \sqrt{\mu} + \sqrt{p_0}} \right] \cdot \frac{1}{\log^2 p_0}$$

for some constants $c', c > 0$. Set $\lambda = 0.1/\sqrt{p_0 \theta}$ in φ_ρ , where $\rho(x) = \sqrt{x^2 + \delta^2}$. There exist numerical constants $C, c'', c''', c_1, c_4 > 0$ such that if $\delta \leq \frac{c'' \lambda \theta^8}{p^2 \log^2 n}$ and $n > Cp_0^5 \theta^{-2} \log p_0$, then with probability at least $1 - c'''/n$, for every $\mathbf{a} \in \Sigma_{4\theta p_0}^\gamma$, we have the following:

(Negative curvature.) If $|\beta_{(1)}| \geq \nu_1 |\beta_{(0)}|$, then

$$(4.19) \quad \lambda_{\min}(\text{Hess}[\varphi_\rho](\mathbf{a})) \leq -c_1 n \theta \lambda.$$

(Large gradient.) If $\nu_1 |\beta_{(0)}| \geq |\beta_{(1)}| \geq \nu_2(\theta) \lambda$, then

$$(4.20) \quad \|\text{grad}[\varphi_\rho](\mathbf{a})\|_2 \geq c_2 n \theta \frac{\lambda^2}{\log^2 \theta^{-1}}.$$

(Convex near shifts.) If $\nu_2(\theta) \lambda \geq |\beta_{(1)}|$, then

$$(4.21) \quad \text{Hess}[\varphi_\rho](\mathbf{a}) \succ c_3 n \theta \mathbf{P}_{\mathbf{a}^\perp}.$$

(Retraction to subspace.) If $\frac{\gamma}{2} \leq d_\alpha(\mathbf{a}, \mathcal{S}_\tau) \leq \gamma$, then for every $\boldsymbol{\alpha}$ satisfying $\mathbf{a} = \boldsymbol{\iota}^* \mathbf{C}_{\mathbf{a}_0} \boldsymbol{\alpha}$, there exists $\boldsymbol{\zeta}$ satisfying $\text{grad}[\varphi_\rho](\mathbf{a}) = \boldsymbol{\iota}^* \mathbf{C}_{\mathbf{a}_0} \boldsymbol{\zeta}$, such that

$$(4.22) \quad \langle \boldsymbol{\zeta}_{\tau^c}, \boldsymbol{\alpha}_{\tau^c} \rangle \geq c_4 \|\boldsymbol{\zeta}_{\tau^c}\|_2 \|\boldsymbol{\alpha}_{\tau^c}\|_2.$$

(Local minimizers.) If $\mathbf{a}$ is a local minimizer,

$$(4.23) \quad \min_{\substack{\ell \in [\pm p] \\ \sigma \in \{\pm 1\}}} \|\mathbf{a} - \sigma s_\ell[\mathbf{a}_0]\|_2 \leq \frac{1}{2} \max\{\mu, p_0^{-1}\},$$

where $\nu_1 = \frac{4}{5}$, $\nu_2(\theta) = \frac{1}{4 \log^2 \theta^{-1}}$, and $\gamma = \frac{c \cdot \text{poly}(\sqrt{1/\theta}, \sqrt{1/\mu})}{\log^2 \theta^{-1}} \cdot \frac{1}{\sqrt{p_0}}$.

Proof. See subsection SM6.5. ■

The retraction property elaborated upon in (4.22) implies that the negative gradient at $\mathbf{a}$ points in a direction that decreases $d_\alpha(\mathbf{a}, \mathcal{S}_\tau)$. This is a consequence of positive curvature away from $\mathcal{S}_\tau$. It essentially implies that the gradient is monotone in α_{τ^c} space: choose any $\underline{\mathbf{a}} \in \mathcal{S}_\tau \cap \mathbb{S}^{p-1}$, write $\underline{\alpha}$ to be its coefficient, and let $\underline{\zeta}$ be the coefficient of $\text{grad}[\varphi_\rho](\underline{\mathbf{a}})$. Then $\underline{\alpha}_{\tau^c} = \mathbf{0}$, $\underline{\zeta}_{\tau^c} \approx \mathbf{0}$, and

$$\langle \zeta_{\tau^c} - \underline{\zeta}_{\tau^c}, \alpha_{\tau^c} - \underline{\alpha}_{\tau^c} \rangle \approx \langle \zeta_{\tau^c} - \mathbf{0}, \alpha_{\tau^c} - \mathbf{0} \rangle = \langle \zeta_{\tau^c}, \alpha_{\tau^c} \rangle > 0.$$

Our main geometric claim in Theorem 3.1 is a direct consequence of Theorem 4.1. Moreover, it suggests that as long as we can minimize φ_ρ within the region $\Sigma_{4\theta p_0}^\gamma$, we will solve the SaS deconvolution problem.

5. Provable algorithm. In light of Theorem 4.1, in this section we introduce a two-part algorithm, Algorithm 3.1, which first applies the curvilinear descent method to find a local minimum of φ_ρ within $\Sigma_{4\theta p_0}^\gamma$, followed by a refinement algorithm that uses alternating minimization to exactly recover the ground truth. This algorithm exactly solves the SaS deconvolution problem.

5.1. Minimization. There are three major issues in finding a local minimizer within $\Sigma_{4\theta p_0}^\gamma$:

- (i) *Initialization.* The initializer $\mathbf{a}^{(0)}$ to reside within $\Sigma_{4\theta p_0}^\gamma$.
- (ii) *Negative curvature.* The method to avoid stagnating near saddle points of φ_ρ .
- (iii) *No exit.* The descent method to remain inside $\Sigma_{4\theta p_0}^\gamma$.

In the following paragraphs, we describe how our proposed algorithm achieves the above desiderata.

Initialization within $\Sigma_{4\theta p_0}^\gamma$. Our data-driven initialization scheme produces $\mathbf{a}^{(0)}$, where

$$\begin{aligned} \mathbf{a}^{(0)} &= -\mathbf{P}_{\mathbb{S}^{p-1}} \nabla \varphi_\rho (\mathbf{P}_{\mathbb{S}^{p-1}} [\mathbf{0}^{p_0-1}; \mathbf{y}_0; \dots; \mathbf{y}_{p_0-1}; \mathbf{0}^{p_0-1}]) \\ &= -\mathbf{P}_{\mathbb{S}^{p-1}} \nabla \varphi_\rho \mathbf{P}_{\mathbb{S}^{p-1}} [\mathbf{P}_{[p_0]}(\mathbf{a}_0 * \mathbf{x}_0)] \\ &\approx -\mathbf{P}_{\mathbb{S}^{p-1}} \nabla \varphi_\rho [\mathbf{P}_{[p_0]}(\mathbf{a}_0 * \tilde{\mathbf{x}}_0)] \end{aligned}$$

is the normalized gradient vector from a chunk of data $\mathbf{a}^{(-1)} := \mathbf{P}_{[p_0]}(\mathbf{a}_0 * \tilde{\mathbf{x}}_0)$ with $\tilde{\mathbf{x}}_0$ a normalized Bernoulli-Gaussian random vector of length $2p_0 - 1$. Since $\nabla \varphi_\rho \approx \nabla \varphi_{\ell^1}$, expanding the gradient $\nabla \varphi_{\ell^1}$ and rewriting the gradient $\nabla_{\ell^1}(\mathbf{a}^{(-1)})$ in shift space gives us

$$\begin{aligned} -\nabla \varphi_{\rho^1}(\mathbf{a}^{(-1)}) &\approx \iota^* \mathbf{C}_{\mathbf{a}_0} \tilde{\mathbf{C}}_{\mathbf{x}_0} \mathcal{S}_\lambda \left[\tilde{\mathbf{C}}_{\mathbf{x}_0} \mathbf{C}_{\mathbf{a}_0}^* \mathbf{P}_{[p_0]}(\mathbf{a}_0 * \tilde{\mathbf{x}}_0) \right] \\ &= \iota^* \mathbf{C}_{\mathbf{a}_0} \mathcal{X} \left[\mathbf{C}_{\mathbf{a}_0}^* \mathbf{P}_{[p_0]} \mathbf{C}_{\mathbf{a}_0} \tilde{\mathbf{x}}_0 \right] \\ &\approx \iota^* \mathbf{C}_{\mathbf{a}_0} \mathcal{X} [\tilde{\mathbf{x}}_0] \\ &\approx n\theta \cdot \iota^* \mathbf{C}_{\mathbf{a}_0} \mathcal{S}_\lambda [\tilde{\mathbf{x}}_0], \end{aligned}$$

where the approximation in the third equation is accurate if the truncated shifts are incoherent:

$$(5.1) \quad \max_{i \neq j} |\langle \boldsymbol{\iota}_{p_0}^* s_i[\mathbf{a}_0], \boldsymbol{\iota}_{p_0}^* s_j[\mathbf{a}_0] \rangle| \leq \mu \ll 1.$$

With this simple approximation, it becomes clear that the coefficients (in shift space) of initializer $\mathbf{a}^{(0)}$,

$$(5.2) \quad \mathbf{a}^{(0)} \approx \mathbf{P}_{\mathbb{S}^{p-1}} \boldsymbol{\iota}^* \mathbf{C}_{\mathbf{a}_0} \mathcal{S}_\lambda [\tilde{\mathbf{x}}_0],$$

approximate $\mathcal{S}_\lambda [\tilde{\mathbf{x}}_0]$, which resides near the subspace $\mathcal{S}_\tau$, in which τ contains the nonzero entries of $\tilde{\mathbf{x}}_0$ on $\{-p_0 + 1, \dots, p_0 - 1\}$. With high probability, the number of nonzero entries is $|\tau| \lesssim 4\theta p_0$, and we therefore conclude that our initializer $\mathbf{a}^{(0)}$ satisfies

$$(5.3) \quad \mathbf{a}^{(0)} \in \Sigma_{4\theta p_0}^\gamma.$$

Furthermore, since $\tilde{\mathbf{x}}_0$ is normalized, the largest magnitude for entries of $|\tilde{\mathbf{x}}_0|$ is likely to be around $1/\sqrt{2p_0\theta}$. To ensure that $\mathcal{S}_\lambda [\tilde{\mathbf{x}}_0]$ does not annihilate all nonzero entries of $\tilde{\mathbf{x}}_0$ (otherwise our initializer $\mathbf{a}^{(0)}$ will become $\mathbf{0}$), the ideal λ should be slightly less than the largest magnitude of $|\tilde{\mathbf{x}}_0|$. We suggest setting λ in φ_ρ as

$$(5.4) \quad \lambda = \frac{c}{\sqrt{p_0\theta}}$$

for some $c \in (0, 1)$.

Many methods have been proposed to optimize functions whose saddle points exhibit strict negative curvature, including the noisy gradient method [22], trust region methods [1, 55], and curvilinear search [58]. Any of the above methods can be adapted to minimize φ_ρ . In this paper, we use the *curvilinear method with restricted stepsize* to demonstrate how to analyze an optimization problem using the geometric properties of φ_ρ over $\Sigma_{4\theta p_0}^\gamma$ —in particular, negative curvature in symmetry breaking directions and positive curvature away from $\mathcal{S}_\tau$.

Curvilinear search uses an update strategy that combines the gradient $\mathbf{g}$ and a direction of negative curvature $\mathbf{v}$, which here we choose as an eigenvector of the Hessian $\mathbf{H}$ with smallest eigenvalue, scaled such that $\mathbf{v}^* \mathbf{g} \geq 0$. In particular, we set

$$(5.5) \quad \mathbf{a}^+ \leftarrow \mathbf{P}_{\mathbb{S}^{p-1}} [\mathbf{a} - t\mathbf{g} - t^2\mathbf{v}].$$

For small t ,

$$(5.6) \quad \varphi(\mathbf{a}^+) \approx \varphi(\mathbf{a}) + \langle \mathbf{g}, \boldsymbol{\xi} \rangle + \frac{1}{2} \boldsymbol{\xi}^* \mathbf{H} \boldsymbol{\xi}.$$

Since $\boldsymbol{\xi}$ converges to $\mathbf{0}$ only if $\mathbf{a}$ converges to the local minimizer (otherwise either gradient $\mathbf{g}$ is nonzero or there is a negative curvature direction $\mathbf{v}$), this iteration produces a local minimizer for φ_ρ , whose saddle points near any $\mathcal{S}_\tau$ have negative curvature, we just need to ensure all iterates stay near some such subspace. We prove this by showing the following:

- When $d_\alpha(\mathbf{a}, \mathcal{S}_\tau) \leq \gamma$, curvilinear steps move a small distance away from the subspace:

$$(5.7) \quad |d_\alpha(\mathbf{a}^+, \mathcal{S}_\tau) - d_\alpha(\mathbf{a}, \mathcal{S}_\tau)| \leq \frac{\gamma}{2}.$$

- When $d_\alpha(\mathbf{a}, \mathcal{S}_\tau) \in [\frac{\gamma}{2}, \gamma]$, curvilinear steps retract toward subspace:

$$(5.8) \quad d_\alpha(\mathbf{a}^+, \mathcal{S}_\tau) \leq d_\alpha(\mathbf{a}, \mathcal{S}_\tau).$$

Together, we can prove that the iterates $\mathbf{a}^{(k)}$ converge to a minimizer, and

$$(5.9) \quad \forall k = 1, 2, \dots, \quad \mathbf{a}^{(k)} \in \Sigma_{4\theta p_0}^\gamma.$$

We conclude this section with the following theorem.

Theorem 5.1 (convergence of retractive curvilinear search). *Suppose signals $\mathbf{a}_0, \mathbf{x}_0$ satisfy the conditions of Theorem 4.1, $\theta > 10^3 c/p_0$ ($c > 1$), and $\mathbf{a}_0$ is μ -truncated shift coherent $\max_{i \neq j} |\langle \boldsymbol{\iota}_{p_0}^* s_i[\mathbf{a}_0], \boldsymbol{\iota}_{p_0}^* s_j[\mathbf{a}_0] \rangle| \leq \mu$. Write $\mathbf{g} = \text{grad}[\varphi_\rho](\mathbf{a})$ and $\mathbf{H} = \text{Hess}[\varphi_\rho](\mathbf{a})$. When the smallest eigenvalue of $\mathbf{H}$ is strictly smaller than $-\eta_v$, let $\mathbf{v}$ be the unit eigenvector of smallest eigenvalue, scaled so $\mathbf{v}^* \mathbf{g} \geq 0$; otherwise let $\mathbf{v} = \mathbf{0}$. Define a sequence $\{\mathbf{a}^{(k)}\}_{k \in \mathbb{N}}$ where $\mathbf{a}^{(0)}$ equals (3.7) and for $k = 1, 2, \dots, K_1$*

$$(5.10) \quad \mathbf{a}^{(k+1)} \leftarrow \mathbf{P}_{\mathbb{S}^{p-1}} \left[\mathbf{a}^{(k)} - t\mathbf{g}^{(k)} - t^2\mathbf{v}^{(k)} \right],$$

with largest $t \in (0, \frac{0.1}{n\theta}]$ satisfying Armijo steplength

$$(5.11) \quad \varphi_\rho(\mathbf{a}^{(k+1)}) < \varphi_\rho(\mathbf{a}^{(k)}) - \frac{1}{2} \left(t \|\mathbf{g}^{(k)}\|_2^2 + \frac{1}{2} t^4 \eta_v \|\mathbf{v}^{(k)}\|_2^2 \right).$$

Then with probability at least $1 - 1/c$, there exists some signed shift $\bar{\mathbf{a}} = \pm s_i[\mathbf{a}_0]$, where $i \in [\pm p_0]$, such that $\|\mathbf{a}^{(k)} - \bar{\mathbf{a}}\|_2 \leq \mu + 1/p$ for all $k \geq K_1 = \text{poly}(n, p)$. Here, $\eta_v = c'n\theta\lambda$ for some $c' < c_1$ in Theorem 4.1.

Proof. See subsection SM7.2. ■

5.2. Local refinement. In this section, we describe and analyze an algorithm which refines an estimate $\bar{\mathbf{a}} \approx \mathbf{a}_0$ of the kernel to exactly recover $(\mathbf{a}_0, \mathbf{x}_0)$. Set

$$(5.12) \quad \mathbf{a}^{(0)} \leftarrow \bar{\mathbf{a}}, \quad \lambda^{(0)} \leftarrow C(p\theta + \log n)(\mu + 1/p), \quad I^{(0)} \leftarrow \text{supp}(\mathcal{S}_\lambda[\mathbf{C}_{\bar{\mathbf{a}}}^* \mathbf{y}]).$$

We alternatively minimize the Lasso objective with respect to $\mathbf{a}$ and $\mathbf{x}$:

$$(5.13) \quad \mathbf{x}^{(k+1)} \leftarrow \underset{\mathbf{x}}{\text{argmin}} \frac{1}{2} \|\mathbf{a}^{(k)} * \mathbf{x} - \mathbf{y}\|_2^2 + \lambda^{(k)} \sum_{i \notin I^{(k)}} |\mathbf{x}_i|,$$

$$(5.14) \quad \mathbf{a}^{(k+1)} \leftarrow \mathbf{P}_{\mathbb{S}^{p-1}} \left[\underset{\mathbf{a}}{\text{argmin}} \frac{1}{2} \|\mathbf{a} * \mathbf{x}^{(k+1)} - \mathbf{y}\|_2^2 \right],$$

$$(5.15) \quad \lambda^{(k+1)} \leftarrow \frac{1}{2} \lambda^{(k)}, \quad I^{(k+1)} \leftarrow \text{supp}(\mathbf{x}^{(k+1)}).$$

One departure from standard alternating minimization procedures is our use of a continuation method, which (i) decreases λ , and (ii) maintains a running estimate $I^{(k)}$ of the support set. Our analysis will show that $\mathbf{a}^{(k)}$ converges to one of the signed shifts of $\mathbf{a}_0$ at a linear rate, in the sense that

$$(5.16) \quad \min_{\sigma \in \pm 1, \ell \in [\pm p_0]} \|\mathbf{a}^{(k)} - \sigma \cdot s_\ell[\mathbf{a}_0]\|_2 \leq C' 2^{-k}.$$

It should be clear that exact recovery is unlikely if $\mathbf{x}_0$ contains many consecutive nonzero entries: in fact in this situation, even *nonblind* deconvolution fails. Therefore to obtain exact recovery it is necessary to put an upper bound on signal dimension n . Here, we introduce the notation κ_I as an upper bound for the number of nonzero entries of $\mathbf{x}_0$ in a length- p window:

$$(5.17) \quad \kappa_I := 6 \max \{ \theta p, \log n \},$$

where the indexing and addition should be interpreted modulo n . We will denote the support sets of true sparse vector $\mathbf{x}_0$ and recovered $\mathbf{x}^{(k)}$ in the intermediate k th steps as

$$(5.18) \quad I = \text{supp}(\mathbf{x}_0), \quad I^{(k)} = \text{supp}(\mathbf{x}^{(k)}).$$

Then in the Bernoulli–Gaussian model, with high probability,

$$(5.19) \quad \max_{\ell} |I \cap ([p] + \ell)| \leq \kappa_I.$$

The $\log n$ term reflects the fact that as n becomes enormous (exponential in p), eventually it becomes likely that some length- p window of $\mathbf{x}_0$ is densely occupied. In our main theorem statement, we preclude this possibility by putting an upper bound on signal length n with respect to window length p and shift coherence μ . We will assume

$$(5.20) \quad (\mu + 1/p) \cdot \kappa_I^2 < c$$

for some numerical constant $c \in (0, 1)$.

Recall that (4.23) in Theorem 3.1 provides that

$$(5.21) \quad \|\bar{\mathbf{a}} - \mathbf{a}_0\|_2 \leq (\mu + 1/p),$$

which is sufficiently close to $\mathbf{a}_0$ as long as (5.19) holds true. Here, we will elaborate upon this by showing that a single iteration of alternating minimization algorithm (5.13)–(5.15) is a contraction mapping for $\mathbf{a}$ toward $\mathbf{a}_0$.

To this end, at k th iteration, write $T = I^{(k)}$, $J = I^{(k+1)}$, and $\boldsymbol{\sigma}^{(k)} = \text{sign}(\mathbf{x}^{(k)})$; then first observe that the solution to the reweighted Lasso problem (5.13) can be written as

$$(5.22) \quad \mathbf{x}^{(k+1)} = \boldsymbol{\iota}_J \left(\boldsymbol{\iota}_J^* \mathbf{C}_{\mathbf{a}^{(k)}}^* \mathbf{C}_{\mathbf{a}^{(k)}} \boldsymbol{\iota}_J \right)^{-1} \boldsymbol{\iota}_J^* \left(\mathbf{C}_{\mathbf{a}^{(k)}}^* \mathbf{C}_{\mathbf{a}_0} \mathbf{x}_0 - \lambda^{(k)} \mathbf{P}_{J \setminus T} \boldsymbol{\sigma}^{(k+1)} \right),$$

and the solution to least squares problem (5.14) will be

$$(5.23) \quad \mathbf{a}^{(k+1)} = \left(\boldsymbol{\iota}^* \mathbf{C}_{\mathbf{x}^{(k+1)}}^* \mathbf{C}_{\mathbf{x}^{(k+1)}} \boldsymbol{\iota} \right)^{-1} \left(\boldsymbol{\iota}^* \mathbf{C}_{\mathbf{x}^{(k+1)}}^* \mathbf{C}_{\mathbf{x}_0} \boldsymbol{\iota} \mathbf{a}_0 \right).$$

Here, we are going to illustrate the relationship between $\mathbf{a}^{(k+1)} - \mathbf{a}_0$ and $\mathbf{a}^{(k)} - \mathbf{a}_0$ using simple approximations. First, let us assume that $\mathbf{a}^{(k)} \approx \mathbf{a}_0$, $\mathbf{C}_{\mathbf{a}_0}^* \mathbf{C}_{\mathbf{a}_0} \approx \mathbf{I}$, and $I \approx J \approx T$. Then (5.22) gives

$$(5.24) \quad \mathbf{x}^{(k+1)} \approx \mathbf{x}_0,$$

$$(5.25) \quad \begin{aligned} (\mathbf{x}^{(k+1)} - \mathbf{x}_0) &\approx \mathbf{P}_I \left(\mathbf{C}_{\mathbf{a}_0}^* \mathbf{C}_{\mathbf{a}_0} \mathbf{x}_0 - \mathbf{C}_{\mathbf{a}_0}^* \mathbf{C}_{\mathbf{a}^{(k)}} \mathbf{x}_0 \right) \\ &\approx \mathbf{P}_I \left[\mathbf{C}_{\mathbf{a}_0}^* \mathbf{C}_{\mathbf{x}_0} \boldsymbol{\iota} (\mathbf{a}_0 - \mathbf{a}^{(k)}) \right], \end{aligned}$$

which implies, while assuming $\mathbf{C}_{\mathbf{x}_0}^* \mathbf{C}_{\mathbf{x}_0} \approx n\theta \mathbf{I}$, that from (5.23),

$$\begin{aligned}
 (\mathbf{a}^{(k+1)} - \mathbf{a}_0) &\approx (n\theta)^{-1} \boldsymbol{\iota}^* \mathbf{C}_{\mathbf{x}^{(k+1)}}^* \mathbf{C}_{\mathbf{x}_0} \boldsymbol{\iota} \mathbf{a}_0 - \boldsymbol{\iota}^* \mathbf{C}_{\mathbf{x}^{(k+1)}}^* \mathbf{C}_{\mathbf{x}^{(k+1)}} \boldsymbol{\iota} \mathbf{a}_0 \\
 &\approx (n\theta)^{-1} \boldsymbol{\iota}^* \mathbf{C}_{\mathbf{x}_0}^* \mathbf{C}_{\mathbf{a}_0} (\mathbf{x}_0 - \mathbf{x}^{(k+1)}) \\
 &\approx (n\theta)^{-1} \boldsymbol{\iota}^* \mathbf{C}_{\mathbf{x}_0}^* \mathbf{C}_{\mathbf{a}_0} \mathbf{P}_I \mathbf{C}_{\mathbf{a}_0}^* \mathbf{C}_{\mathbf{x}_0} \boldsymbol{\iota} (\mathbf{a}^{(k)} - \mathbf{a}_0).
 \end{aligned}
 \tag{5.26}$$

Now since $\mathbf{C}_{\mathbf{x}_0}^* \mathbf{P}_I \mathbf{C}_{\mathbf{x}_0} \approx n\theta \mathbf{e}_0 \mathbf{e}_0^*$, this suggests that $(n\theta)^{-1} \boldsymbol{\iota}^* \mathbf{C}_{\mathbf{x}_0}^* \mathbf{C}_{\mathbf{a}_0} \mathbf{P}_I \mathbf{C}_{\mathbf{a}_0}^* \mathbf{C}_{\mathbf{x}_0} \boldsymbol{\iota}$ approximates a contraction mapping with fixed point $\mathbf{a}_0$, as follows:

$$\begin{aligned}
 (n\theta)^{-1} \boldsymbol{\iota}^* \mathbf{C}_{\mathbf{x}_0}^* \mathbf{C}_{\mathbf{a}_0} \mathbf{P}_I \mathbf{C}_{\mathbf{a}_0}^* \mathbf{C}_{\mathbf{x}_0} \boldsymbol{\iota} &\approx \boldsymbol{\iota}^* \mathbf{C}_{\mathbf{a}_0} \mathbf{e}_0 \mathbf{e}_0^* \mathbf{C}_{\mathbf{a}_0}^* \boldsymbol{\iota} \\
 &\approx \mathbf{a}_0 \mathbf{a}_0^*.
 \end{aligned}
 \tag{5.27}$$

Hence, if we can ensure all of the above approximation is sufficiently and increasingly accurate as the iterate proceeds, the alternating minimization essentially is a power method which finds the leading eigenvector of matrix $\mathbf{a}_0 \mathbf{a}_0^*$ —and the solution to this algorithm is apparently $\mathbf{a}_0$. Indeed, we prove that the iterates produced by this sequence of operations converge to the ground truth at a linear rate, as long as it is initialized sufficiently nearby.

Theorem 5.2 (linear rate convergence of alternating minimization). *Suppose $\mathbf{y} = \mathbf{a}_0 * \mathbf{x}_0$, where $\mathbf{a}_0$ is μ -shift coherent and $\mathbf{x}_0 \sim \text{BG}(\theta)$. Then there exist some constants C, c, c_μ such that if $(\mu + 1/p) \kappa_I^2 < c_\mu$ and $n > C\theta^{-2} p^2 \log n$, then with probability at least $1 - c/n$, for any starting point $\mathbf{a}^{(0)}$ and $\lambda^{(0)}, I^{(0)}$ such that*

$$\|\mathbf{a}^{(0)} - \mathbf{a}_0\|_2 \leq \mu + 1/p, \quad \lambda^{(0)} = 5\kappa_I(\mu + 1/p), \quad I^{(0)} = \text{supp}(\mathbf{C}_{\mathbf{a}^{(0)}}^* \mathbf{y}),
 \tag{5.28}$$

and for $k = 1, 2, \dots$,

$$\mathbf{x}^{(k+1)} \leftarrow \underset{\mathbf{x}}{\text{argmin}} \frac{1}{2} \|\mathbf{a}^{(k)} * \mathbf{x} - \mathbf{y}\|_2^2 + \lambda^{(k)} \sum_{i \notin I^{(k)}} |\mathbf{x}_i|,
 \tag{5.29}$$

$$\mathbf{a}^{(k+1)} \leftarrow \mathbf{P}_{\mathbb{S}^{p-1}} \left[\underset{\mathbf{a}}{\text{argmin}} \frac{1}{2} \|\mathbf{a} * \mathbf{x}^{(k+1)} - \mathbf{y}\|_2^2 \right],
 \tag{5.30}$$

$$\lambda^{(k+1)} \leftarrow \frac{1}{2} \lambda^{(k)}, \quad I^{(k+1)} \leftarrow \text{supp}(\mathbf{x}^{(k+1)}),
 \tag{5.31}$$

then

$$\|\mathbf{a}^{(k+1)} - \mathbf{a}_0\|_2 \leq (\mu + 1/p) 2^{-k}
 \tag{5.32}$$

for every $k = 0, 1, 2, \dots$.

Proof. See subsection SM8.3. ■

Remark 5.3. The estimates $\mathbf{x}^{(k)}$ also converges to the ground truth $\mathbf{x}_0$ at a linear rate.

6. Experiments. We demonstrate that the tradeoffs between the motif length p_0 and sparsity rate θ produce a transition region for successful SaS deconvolution under generic choices of $\mathbf{a}_0$ and $\mathbf{x}_0$. For fixed values of $\theta \in [10^{-3}, 10^{-2}]$ and $p_0 \in [10^3, 10^4]$, we draw 50 instances of synthetic data by choosing $\mathbf{a}_0 \sim \text{Unif}(\mathbb{S}^{p_0-1})$ and $\mathbf{x}_0 \in \mathbb{R}^n$ with $\mathbf{x}_0 \sim_{\text{i.i.d.}} \text{BG}(\theta)$, where $n = 5 \times 10^5$. Note that choosing $\mathbf{a}_0$ this way implies $\mu(\mathbf{a}_0) \approx \frac{1}{\sqrt{p_0}}$.

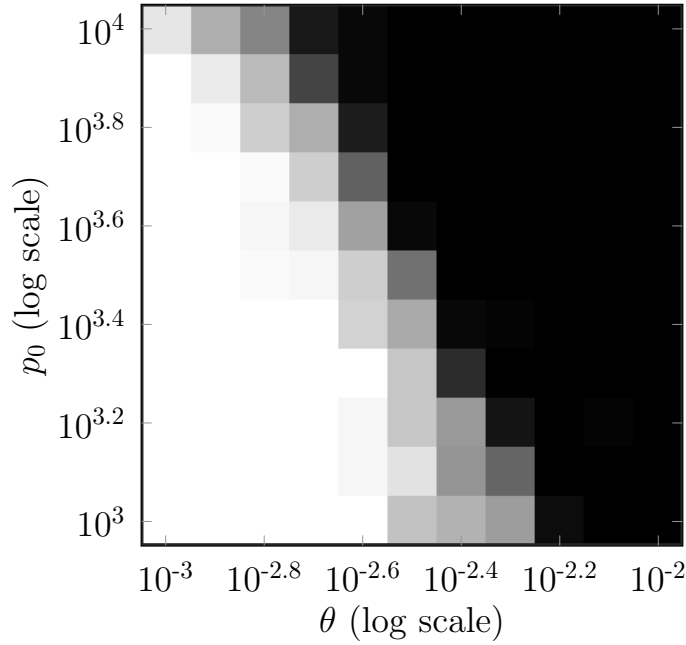


Figure 12. Success probability of SaS deconvolution under generic $\mathbf{a}_0, \mathbf{x}_0$ with varying kernel length p_0 and sparsity rate θ . When sparsity rate decreases sufficiently with respect to kernel length, successful recovery becomes very likely (brighter), and vice versa (darker). A transition line is shown with slope $\frac{\log p_0}{\log \theta} \approx -2$, implying *Algorithm 6.1* works with high probability when $\theta \lesssim \frac{1}{\sqrt{p_0}}$ in the generic case.

For each instance, we recover $\mathbf{a}_0$ and $\mathbf{x}_0$ from $\mathbf{y} = \mathbf{a}_0 * \mathbf{x}_0$ by minimizing problem (2.5). For ease of computation, we modify *Algorithm 3.1* by replacing curvilinear search with the *accelerated Riemannian gradient descent* method (*Algorithm 6.1*), which is an adaptation of accelerated gradient descent [5] to the sphere. In particular, we apply momentum and increment by the Riemannian gradient via the exponential and logarithmic operators

$$(6.1) \quad \text{Exp}_{\mathbf{a}}(\mathbf{u}) := \cos(\|\mathbf{u}\|_2) \cdot \mathbf{a} + \sin(\|\mathbf{u}\|_2) \cdot \frac{\mathbf{u}}{\|\mathbf{u}\|_2},$$

$$(6.2) \quad \text{Log}_{\mathbf{a}}(\mathbf{b}) := \arccos(\langle \mathbf{a}, \mathbf{b} \rangle) \cdot \frac{P_{\mathbf{a}^\perp}(\mathbf{b} - \mathbf{a})}{\|P_{\mathbf{a}^\perp}(\mathbf{b} - \mathbf{a})\|_2},$$

derived from [1]. Here $\text{Exp}_{\mathbf{a}} : \mathbf{a}^\perp \rightarrow \mathbb{S}^{p-1}$ takes a tangent vector of $\mathbf{a}$ and produces a new point on the sphere, whereas $\text{Log}_{\mathbf{a}} : \mathbb{S}^{p-1} \rightarrow \mathbf{a}^\perp$ takes a point $\mathbf{b} \in \mathbb{S}^{p-1}$ and returns the tangent vector which points from $\mathbf{a}$ to $\mathbf{b}$.

For each recovery instance, we say the local minimizer $\mathbf{a}_{\min}$ generated from *Algorithm 6.1* is sufficiently close to a solution of the SaS deconvolution problem if

$$(6.3) \quad \text{success}(\mathbf{a}_{\min}; \mathbf{a}_0) := \{ \max_{\ell} |\langle s_{\ell}[\mathbf{a}_0], \mathbf{a}_{\min} \rangle| > 0.95 \}.$$

The result is shown in *Figure 12*. Our source code can be accessed via the following address:

https://github.com/sbdsphere/sbd_experiments.git

Algorithm 6.1. SaS deconvolution with accelerated Riemannian gradient descent.

Input: Observation $\mathbf{y}$, sparsity penalty $\lambda = 0.5/\sqrt{p_0\theta}$, momentum parameter $\eta \in [0, 1)$.

Initialize $\mathbf{a}^{(0)} \leftarrow -\mathbf{P}_{\mathbb{S}^{p-1}} \nabla \varphi_\rho(\mathbf{P}_{\mathbb{S}^{p-1}} [\mathbf{0}^{p_0-1}; [\mathbf{y}_0, \dots, \mathbf{y}_{p_0-1}]; \mathbf{0}^{p_0-1}])$,

for $k = 1, 2, \dots, K$ **do**

 Get momentum: $\mathbf{w} \leftarrow \text{Exp}_{\mathbf{a}^{(k)}}(\eta \cdot \text{Log}_{\mathbf{a}^{(k-1)}}(\mathbf{a}^{(k)}))$.

 Get negative gradient direction: $\mathbf{g} \leftarrow -\text{grad}[\varphi_\rho](\mathbf{w})$.

 Armijo step $\mathbf{a}^{(k+1)} \leftarrow \text{Exp}_{\mathbf{w}}(t\mathbf{g})$, choosing $t \in (0, 1)$ s.t. $\varphi_\rho(\mathbf{a}^{(k+1)}) - \varphi_\rho(\mathbf{w}) < -t \|\mathbf{g}\|_2^2$.

end for

Output: Return $\mathbf{a}^{(K)}$.

7. Discussion. In this section, we close by discussing the most important limitations of our results when $\mathbf{a}_0$ is coherent, regarding scenarios when the signal setting breaches our assumption, especially when $\mathbf{x}_0$ is either highly sparse or nonsymmetric, and highlighting corresponding directions for future work.

The main drawback of our proposed method is that it does not succeed when the target motif $\mathbf{a}_0$ has shift coherence very close to 1. For instance, a common scenario in image blind deconvolution involves deblurring an image with a smooth, low-pass point spread function (e.g., Gaussian blur). Both our analysis and numerical experiments show that in this situation minimizing φ_ρ does not find the generating signal pairs $(\mathbf{a}_0, \mathbf{x}_0)$ consistently—the minimizer of φ_ρ is often spurious and is not close to any particular shift of $\mathbf{a}_0$. We do not suggest minimizing φ_ρ in this situation. On the other hand, minimizing the Bilinear Lasso objective φ_{lasso} over the sphere often succeeds even if the true signal pair $(\mathbf{a}_0, \mathbf{x}_0)$ is coherent and dense.

In light of the above observations, we view the analysis of the Bilinear Lasso as the most important direction for future theoretical work on SaS deconvolution. The drop quadratic formulation studied here has commonalities with the Bilinear Lasso: both exhibit local minima at signed shifts, and both exhibit negative curvature in symmetry breaking directions. A major difference (and hence major challenge) is that gradient methods for Bilinear Lasso do not retract to a union of subspaces—they retract to a more complicated, nonlinear set.

Our model assumes $\mathbf{x}_0$ to be Bernoulli–Gaussian vectors, which are sparse and symmetric i.i.d. random variables. When $\mathbf{x}_0$ is sparse but nonsymmetric, (e.g., Bernoulli), one can apply our result with a simple symmetrization trick, using the concatenated observation vectors $[\mathbf{y}, -\mathbf{y}]$ as an input to our algorithm.

When $\mathbf{x}_0$ is highly sparse and if $\mathbf{y}$ is noiseless, it is possible to identify a short copy of $\mathbf{a}_0$ via looking for the shortest consecutive nonzero entries within $\mathbf{y}$. When $\theta \ll 1/p_0$, these isolated copies are very common. Once θ exceeds $1/p_0$, or when support $\mathbf{x}_0$ is not Bernoulli random while being more clustered, they become very uncommon. In particular, the probability of an isolated copy is small unless $n \gtrsim \exp(p_0\theta)$. Our proposed approach succeeds when $n \geq \text{poly}(p_0)$.

In applications involving noisy data, optimization approaches often outperform direct inspection, even for samples with isolated copies of $\mathbf{a}_0$. An intuition for this is that optimization methods aggregate information across the sample. One practical avenue for obtaining the best of both worlds is to try to optimize the choice of data segment used for initialization. This can be a potential improvement for our data-driven initialization scheme, both in theory and in practice.

Finally, there are several directions in which our analysis could be improved. Our lower bounds on the length n of the random vector $\mathbf{x}_0$ required for success are clearly suboptimal. We also suspect our sparsity-coherence tradeoff between μ, θ (roughly, $\theta \lesssim 1/(\sqrt{\mu}p_0)$) is suboptimal, even for the φ_p objective. Articulating optimal sparsity-coherence tradeoffs is another interesting direction in this line of work. Extending our current result for cases when $\mathbf{y}$ is affected by noise can also be a natural next step for future work.

REFERENCES

- [1] P.-A. ABSIL, R. MAHONY, AND R. SEPULCHRE, *Optimization Algorithms on Matrix Manifolds*, Princeton University Press, 2009.
- [2] A. AHMED, B. RECHT, AND J. ROMBERG, *Blind deconvolution using convex programming*, IEEE Trans. Inform. Theory, 60 (2014), pp. 1711–1732.
- [3] G. AYERS AND J. C. DAINTY, *Iterative blind deconvolution method and its applications*, Opt. Lett., 13 (1988), pp. 547–549.
- [4] S. BAKER AND T. KANADE, *Limits on super-resolution and how to break them*, IEEE Trans. Pattern Anal. Mach. Intell., 24 (2002), pp. 1167–1183.
- [5] A. BECK AND M. TEOULLE, *A fast iterative shrinkage-thresholding algorithm for linear inverse problems*, SIAM J. Imaging Sci., 2 (2009), pp. 183–202, <https://doi.org/10.1137/080716542>.
- [6] A. J. BELL AND T. J. SEJNOWSKI, *An information-maximization approach to blind separation and blind deconvolution*, Neural Comput., 7 (1995), pp. 1129–1159.
- [7] A. BENICHOX, E. VINCENT, AND R. GRIBONVAL, *A fundamental pitfall in blind deconvolution with sparse and shift-invariant priors*, in the 38th IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP-2013), IEEE, 2013, pp. 6108–6112.
- [8] P. BONES, C. PARKER, B. SATHERLEY, AND R. WATSON, *Deconvolution and phase retrieval with use of zero sheets*, J. Opt. Soc. A, 12 (1995), pp. 1842–1857.
- [9] D. BRIERS, D. D. DUNCAN, E. R. HIRST, S. J. KIRKPATRICK, M. LARSSON, W. STEENBERGEN, T. STROMBERG, AND O. B. THOMPSON, *Laser speckle contrast imaging: Theoretical and practical limitations*, J. Biomed. Opt., 18 (2013), 066018.
- [10] P. CAMPISI AND K. EGIAZARIAN, *Blind Image Deconvolution: Theory and Applications*, CRC Press, 2016.
- [11] E. J. CANDÈS, M. B. WAKIN, AND S. P. BOYD, *Enhancing sparsity by reweighted ℓ_1 minimization*, J. Fourier Anal. Appl., 14 (2008), pp. 877–905.
- [12] M. CANNON, *Blind deconvolution of spatially invariant image blurs with phase*, IEEE Trans. Acoustics Speech Signal Process., 24 (1976), pp. 58–63.
- [13] A. S. CARASSO, *Direct blind deconvolution*, SIAM J. Appl. Math., 61 (2001), pp. 1980–2007, <https://doi.org/10.1137/S0036139999362592>.
- [14] T. F. CHAN AND C.-K. WONG, *Total variation blind deconvolution*, IEEE Trans. Image Process., 7 (1998), pp. 370–375.
- [15] S. CHEUNG, Y. LAU, Z. CHEN, J. SUN, Y. ZHANG, J. WRIGHT, AND A. PASUPATHY, *Beyond the Fourier transform: A nonconvex optimization approach to microscopy analysis*, submitted.
- [16] Y. CHI, *Guaranteed blind sparse spikes deconvolution via lifting and convex optimization*, IEEE J. Selected Topics Signal Process., 10 (2016), pp. 782–794.
- [17] S. CHOUDHARY AND U. MITRA, *Fundamental Limits of Blind Deconvolution Part II: Sparsity-Ambiguity Trade-Offs*, preprint, <https://arxiv.org/abs/1503.03184>, 2015.
- [18] W. DONG, L. ZHANG, G. SHI, AND X. WU, *Image deblurring and super-resolution by adaptive sparse domain selection and adaptive regularization*, IEEE Trans. Image Process., 20 (2011), pp. 1838–1857.
- [19] B. EFRON, T. HASTIE, I. JOHNSTONE, AND R. TIBSHIRANI, *Least angle regression*, Ann. Statist., 32 (2004), pp. 407–499.
- [20] C. EKANADHAM, D. TRANCHINA, AND E. P. SIMONCELLI, *A blind sparse deconvolution method for neural spike identification*, in Advances in Neural Information Processing Systems 24, 2011, pp. 1440–1448.
- [21] R. FERGUS, B. SINGH, A. HERTZMANN, S. T. ROWEIS, AND W. T. FREEMAN, *Removing camera shake from a single photograph*, ACM Trans. Graphics, 25 (2006), pp. 787–794.

- [22] R. GE, F. HUANG, C. JIN, AND Y. YUAN, *Escaping from saddle points—online stochastic gradient for tensor decomposition*, in Conference on Learning Theory, 2015, pp. 797–842.
- [23] D. GOLDFARB, *Curvilinear path steplength algorithms for minimization which use directions of negative curvature*, Math. Programming, 18 (1980), pp. 31–40.
- [24] D. GOLDFARB, C. MU, J. WRIGHT, AND C. ZHOU, *Using negative curvature in solving nonlinear programs*, Comput. Optim. Appl., 68 (2017), pp. 479–502.
- [25] S. HARMEILING, M. HIRSCH, S. SRA, AND B. SCHOLKOPF, *Online blind deconvolution for astronomical imaging*, in the 2009 IEEE International Conference on Computational Photography (ICCP 2009), IEEE, 2009, pp. 1–7.
- [26] R. JOHNSON, P. SCHNITER, T. J. ENDRES, J. D. BEHM, D. R. BROWN, AND R. A. CASAS, *Blind equalization using the constant modulus criterion: A review*, Proc. IEEE, 86 (1998), pp. 1927–1950.
- [27] N. JOSHI, R. SZELISKI, AND D. J. KRIEGMAN, *PSF estimation using sharp edge prediction*, in the 2008 IEEE Conference on Computer Vision and Pattern Recognition, IEEE, 2008, pp. 1–8.
- [28] K. F. KAARESEN AND T. TAXT, *Multichannel blind deconvolution of seismic signals*, Geophys., 63 (1998), pp. 2093–2107.
- [29] M. KECH AND F. KRAHMER, *Optimal injectivity conditions for bilinear inverse problems with applications to identifiability of deconvolution problems*, SIAM J. Appl. Algebra Geom., 1 (2017), pp. 20–37, <https://doi.org/10.1137/16M1067469>.
- [30] D. KRISHNAN, T. TAY, AND R. FERGUS, *Blind deconvolution using a normalized sparsity measure*, in the 2011 IEEE Conference on Computer Vision and Pattern Recognition (CVPR), IEEE, 2011, pp. 233–240.
- [31] D. KUNDUR AND D. HATZINAKOS, *Blind image deconvolution*, IEEE Signal Process. Mag., 13 (1996), pp. 43–64.
- [32] R. LANE AND R. BATES, *Automatic multidimensional deconvolution*, J. Opt. Soc. A, 4 (1987), pp. 180–188.
- [33] R. G. LANE, *Blind deconvolution of speckle images*, J. Opt. Soc. A, 9 (1992), pp. 1508–1514.
- [34] Y. LAU, Q. QU, H.-W. KUO, P. ZHOU, Y. ZHANG, AND J. WRIGHT, *Short-and-sparse deconvolution: A geometric approach*, in the 8th International Conference on Learning Representations, 2020; preprint, <https://arxiv.org/abs/1908.10959>, 2019.
- [35] A. LEVIN, R. FERGUS, F. DURAND, AND W. T. FREEMAN, *Deconvolution Using Natural Image Priors*, Computer Science and Artificial Intelligence Laboratory, Massachusetts Institute of Technology, 2007.
- [36] A. LEVIN, Y. WEISS, F. DURAND, AND W. T. FREEMAN, *Understanding blind deconvolution algorithms*, IEEE Trans. Pattern Anal. Mach. Intell., 33 (2011), pp. 2354–2367.
- [37] M. S. LEWICKI, *A review of methods for spike sorting: The detection and classification of neural action potentials*, Network, 9 (1998), pp. R53–R78.
- [38] Y. LI AND Y. BRESLER, *Global geometry of multichannel sparse blind deconvolution on the sphere*, in Advances in Neural Information Processing Systems 31, Curran Associates, 2018, pp. 1132–1143.
- [39] Y. LI, K. LEE, AND Y. BRESLER, *Identifiability in blind deconvolution with subspace or sparsity constraints*, IEEE Trans. Inform. Theory, 62 (2016), pp. 4266–4275.
- [40] Y. LI, K. LEE, AND Y. BRESLER, *Identifiability and stability in blind deconvolution under minimal assumptions*, IEEE Trans. Inform. Theory, 63 (2017), pp. 4619–4633.
- [41] S. LING AND T. STROHMER, *Self-calibration and biconvex compressive sensing*, Inverse Problems, 31 (2015), 115002.
- [42] S. LING AND T. STROHMER, *Blind deconvolution meets blind demixing: Algorithms and performance bounds*, IEEE Trans. Inform. Theory, 63 (2017), pp. 4497–4520.
- [43] J. MARKHAM AND J.-A. CONCHELLO, *Parametric blind deconvolution: A robust method for the simultaneous estimation of image and blur*, J. Opt. Soc. A, 16 (1999), pp. 2377–2391.
- [44] M. MIYOSHI AND Y. KANEDA, *Inverse filtering of room acoustics*, IEEE Trans. Acoustics Speech Signal Process., 36 (1988), pp. 145–152.
- [45] P. A. NAYLOR AND N. D. GAUBITCH, *Speech Dereverberation*, Springer Science & Business Media, 2010.
- [46] M. R. OSBORNE, B. PRESNELL, AND B. A. TURLACH, *A new approach to variable selection in least squares problems*, IMA J. Numer. Anal., 20 (2000), pp. 389–403.
- [47] D. PERRONE AND P. FAVARO, *Total variation blind deconvolution: The devil is in the details*, in Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2014, pp. 2909–2916.
- [48] E. A. PNEVMATIKAKIS, D. SOUDRY, Y. GAO, T. A. MACHADO, J. MEREL, D. PFAU, T. REARDON,

- Y. MU, C. LACEFIELD, W. YANG ET AL., *Simultaneous denoising, deconvolution, and demixing of calcium imaging data*, *Neuron*, 89 (2016), pp. 285–299.
- [49] S. K. SAHA, *Diffraction-Limited Imaging with Large and Moderate Telescopes*, World Scientific, 2007.
- [50] Y. SATO, *A method of self-recovering equalization for multilevel amplitude-modulation systems*, *IEEE Trans. Commun.*, 23 (1975), pp. 679–682.
- [51] O. SHALVI AND E. WEINSTEIN, *New criteria for blind deconvolution of nonminimum phase systems (channels)*, *IEEE Trans. Inform. Theory*, 36 (1990), pp. 312–321.
- [52] Q. SHAN, J. JIA, AND A. AGARWALA, *High-quality motion deblurring from a single image*, *ACM Trans. Graphics*, 27 (2008), pp. 1–10.
- [53] G. SHTENGEL, J. A. GALBRAITH, C. G. GALBRAITH, J. LIPPINCOTT-SCHWARTZ, J. M. GILLETTE, S. MANLEY, R. SOUGRAT, C. M. WATERMAN, P. KANCHANAWONG, M. W. DAVIDSON, R. D. FETTER, AND H. F. HESS, *Interferometric fluorescent super-resolution microscopy resolves 3D cellular ultrastructure*, *Proc. Natl. Acad. Sci. USA*, 106 (2009), pp. 3125–3130.
- [54] T. G. STOCKHAM, T. M. CANNON, AND R. B. INGEBRETSEN, *Blind deconvolution through digital signal processing*, *Proc. IEEE*, 63 (1975), pp. 678–692.
- [55] J. SUN, Q. QU, AND J. WRIGHT, *Complete dictionary recovery over the sphere II: Recovery by Riemannian trust-region method*, *IEEE Trans. Inform. Theory*, 63 (2017), pp. 885–914.
- [56] P. WALK, P. JUNG, G. E. PFANDER, AND B. HASSIBI, *Blind Deconvolution with Additional Autocorrelations via Convex Programs*, preprint, <https://arxiv.org/abs/1701.04890>, 2017.
- [57] L. WANG AND Y. CHI, *Blind deconvolution from multiple sparse inputs*, *IEEE Signal Process. Lett.*, 23 (2016), pp. 1384–1388.
- [58] Z. WEN AND W. YIN, *A feasible method for optimization with orthogonality constraints*, *Math. Program.*, 142 (2013), pp. 397–434.
- [59] D. WIPF AND H. ZHANG, *Revisiting Bayesian blind deconvolution*, *J. Mach. Learn. Res.*, 15 (2014), pp. 3595–3634.
- [60] L. XU AND J. JIA, *Two-phase kernel estimation for robust motion deblurring*, in the 11th European Conference on Computer Vision, Springer, 2010, pp. 157–170.
- [61] J. YANG, J. WRIGHT, T. S. HUANG, AND Y. MA, *Image super-resolution via sparse representation*, *IEEE Trans. Image Process.*, 19 (2010), pp. 2861–2873.
- [62] Y.-L. YOU AND M. KAVEH, *Anisotropic blind image restoration*, in Proceedings of the 3rd International Conference on Image Processing, Vol. 2, IEEE, 1996, pp. 461–464.
- [63] Y. ZHANG, H.-W. KUO, AND J. WRIGHT, *Structured local optima in sparse blind deconvolution*, *IEEE Trans. Inform. Theory*, 66 (2020), pp. 419–452.
- [64] Y. ZHANG, Y. LAU, H.-W. KUO, S. CHEUNG, A. PASUPATHY, AND J. WRIGHT, *On the global geometry of sphere-constrained sparse blind deconvolution*, in Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, IEEE, 2017, pp. 4894–4902.