

Compressive Phase Retrieval via Reweighted Amplitude Flow

Liang Zhang, *Student Member, IEEE*, Gang Wang, *Student Member, IEEE*,
Georgios B. Giannakis, *Fellow, IEEE*, and Jie Chen, *Senior Member, IEEE*

Abstract—The problem of reconstructing a sparse signal vector from magnitude-only measurements (a.k.a., compressive phase retrieval), emerges naturally in diverse applications, but it is NP-hard in general. Building on recent advances in nonconvex optimization, this paper puts forth a new algorithm that is termed *compressive reweighted amplitude flow* and abbreviated as CRAF, for compressive phase retrieval. Specifically, CRAF operates in two stages. The first stage seeks a sparse initial guess via a new spectral procedure. In the second stage, CRAF implements a few hard thresholding based iterations using reweighted gradients. When there are sufficient measurements, CRAF provably recovers the underlying signal vector exactly with high probability under suitable conditions. Moreover, its sample complexity coincides with that of the state-of-the-art procedures. Finally, substantial simulated tests showcase remarkable performance of the new spectral initialization, as well as improved exact recovery relative to competing alternatives.

Index terms— Nonconvex optimization, model-based hard thresholding, iteratively reweighting, linear convergence to the global optimum

I. INTRODUCTION

Phase retrieval (PR) refers to the task of reconstructing a signal vector from its phaseless measured linearly transformed entries. It emerges naturally in a wide range of engineering and physics applications such as X-ray crystallography, astronomy, and coherent diffraction imaging [1], [2]. In these setups, the physical sensors can only record the density (the number of photons) of the light waves, but not their phase. This missing phase information renders general phase retrieval ill-posed. In fact, it has been established that reconstructing a discrete, finite-duration signal vector from its Fourier transform magnitudes is generally NP-complete [3]. To obtain useful solutions, additional assumptions have to be made, which include (block) sparsity of underlying signal vectors [4], [5], [6], [7], non-negativity [4], [1], and random Gaussian measurements [8], [9], [10], [11], [12], [13], [14].

The work of L. Zhang, G. Wang, and G. B. Giannakis in this paper was partially supported by NSF grants 1423316, 1442686, 1508993, and 1509040. The work of J. Chen was partially supported by the National Natural Science Foundation of China grants U1509215, 61621063, and the Program for Changjiang Scholars and Innovative Research Team in University (IRT1208). L. Zhang, G. Wang, and G. B. Giannakis are with the Digital Technology Center and the Department of Electrical and Computer Engineering at the University of Minnesota, Minneapolis, MN 55455, USA. G. Wang is also with the State Key Laboratory of Intelligent Control and Decision of Complex Systems, Beijing 100081, P. R. China. J. Chen is with the School of Automation and the State Key Laboratory of Intelligent Control and Decision of Complex Systems, Beijing Institute of Technology, Beijing 100081, P. R. China. E-mails: {zhan3523, gangwang, georgios}@umn.edu; chenjie@bit.edu.cn.

A number of phase retrieval approaches have been developed so far, a sample of which are reviewed next. Alternating projection methods were advocated in [15], [16]. By means of matrix-lifting and upon dropping the nonconvex rank constraint, convex semidefinite programs (SDP) were formulated [17], [18]. Minimizing the least-squares or least-absolute-value loss, several iterative solvers were pursued, namely those abbreviated as AltMinPhase [19], Wirtinger flow (WF) [9], [10], [13], [20], [7], amplitude flow [11], [21], [12], [22], and composite optimization [23]. Convex phase retrieval approaches without matrix lifting can be found in [24], [25]. We also recently developed a reweighted amplitude flow (RAF) algorithm which benchmarks the numerical performance of phase retrieval of signal vectors from Gaussian random measurements [12].

The aforementioned phase retrieval approaches do not exploit possible structural information of the underlying signal vector, and they require for exact recovery that the number of measurements be on the order of the dimension of the vector [9], [12]. This number in large-scale high-resolution imaging applications is on the order of millions, rendering such algorithms inefficient. The signal vectors or their feature maps in many practical setups however, are naturally sparse or admit an (approximately) sparse representation after certain known and deterministic linear transformations have been applied [1]. This prior information can be critical in reducing the number of measurements required by general phase retrieval approaches, and has prompted the development of various (block) sparse phase retrieval solvers. To obtain sparse solutions, the ℓ_1 -regularized PhaseLift was solved in [26]. Targeting nonconvex compressive phase retrieval formulations, a greedy algorithm was devised [5], and the soft-thresholded Wirtinger flow (TWF) [7] as well as the sparse truncated amplitude flow (SPARTA) [6] was developed; see also [27] for the (block) compressive phase retrieval with alternating minimization (CoPRAM).

Building upon and going well beyond our precursors in [12], [6], this paper puts forth a new algorithm termed *compressive reweighted amplitude flow* (CRAF) for (block)-sparse phase retrieval. Generalizing [6], while further accounting for the structured sparsity pattern, the amplitude-based (block)-sparse phase retrieval problem is formulated, and it is solved in two stages, namely the initialization and the refinement stages. To enhance the initialization performance, a new sparse spectral initialization is developed, which judiciously assigns a *negative or positive weight* to each sample. As such, the mean of the resultant initialization matrix features a larger

gap between the first and the second eigenvalues, hence yielding improved performance as will be demonstrated in the numerical tests. The second stage of CRAF successively refines the initialization by means of (model-based) hard thresholding iterations using *reweighted* gradients. From the theoretical side, CRAF provably recovers the true signal vector at a linear rate under suitable conditions. Finally, numerical tests showcase the CRAF's improved recovery, and robustness to unknown sparsity relative to competing approaches.

The remainder of this paper is structured as follows. Section II outlines the (block)-sparse phase retrieval problem. Section III describes the algorithm, and establishes its convergence. Simulated tests are presented in Section IV, and the proofs of the main theorems are given in Section V. Section VI concludes the paper.

Regarding notation, lower- (upper-) case boldface letters stand for column vectors (matrices). Sets are represented by calligraphic letters, e.g., $\mathcal{S}$, with the exception of $\mathcal{T}$ as superscript denoting matrix or vector transposition. The cardinality of set $\mathcal{S}$ is given by $|\mathcal{S}|$. Symbol $\|\cdot\|_2$ is reserved for the Euclidean norm, whereas $\|\cdot\|_0$ for the ℓ_0 (pseudo)-norm counting the number of nonzero entries in a vector. Operator $\lceil \cdot \rceil$ returns the smallest integer greater than or equal to the given scalar. The Gauss error function $\text{erf}(x)$ is defined as $\text{erf}(x) := (1/\sqrt{\pi}) \int_{-x}^x e^{-\tilde{x}^2} d\tilde{x}$. For a positive integer m , $[m]$ denotes the index set $\{1, 2, \dots, m\}$. Finally, the ordered eigenvalues of matrix $\mathbf{X} \in \mathbb{R}^{n \times n}$ are given as $\lambda_1(\mathbf{X}) \geq \lambda_2(\mathbf{X}) \geq \dots \geq \lambda_n(\mathbf{X})$.

II. COMPRESSIVE PHASE RETRIEVAL

The compressive phase retrieval aims at recovering a sparse signal vector from a few magnitude-only measurements [5], [6], [7]. Mathematically, it can be described as follows: Given a small set of phaseless linear measurements

$$\psi_i = |\langle \mathbf{a}_i, \mathbf{x} \rangle|, \quad 1 \leq i \leq m \quad (1)$$

in which $\{\psi_i\}_{i=1}^m$ are the observed magnitudes, and $\{\mathbf{a}_i \in \mathbb{R}^n\}_{i=1}^m$ the known sampling vectors, the goal is to recover a (kB) -sparse solution $\mathbf{x} \in \mathbb{R}^n$, namely $\|\mathbf{x}\|_0 \leq kB$ with kB being the known sparsity level. To accommodate also the block-sparse signal vectors, the following terminology is useful. Suppose without loss of generality that $\mathbf{x}$ is split into N_B blocks $\{\mathbf{x}_b\}_{b=1}^{N_B}$, namely one can write $\mathbf{x} := [\mathbf{x}_1^T \cdots \mathbf{x}_{N_B}^T]^T$ [28]. For notational brevity, let $\mathcal{N}_B := \{1, \dots, N_B\}$ denote the index set of all blocks, and $\mathcal{B}_b$ collect all indices of the entries of $\mathbf{x}$ corresponding to the b -th block. Therefore, $\mathcal{B}_b \subseteq [n]$ for all $b \in \mathcal{N}_b$, where $[n] := \{1, \dots, n\}$ consisting of all indices of $\mathbf{x}$.

Definition 1 (k -block-sparse vectors [29]). *The k -block sparse vectors refer to vectors $\mathbf{x} = [\mathbf{x}_1^T \cdots \mathbf{x}_{N_B}^T]^T$ such that $\mathbf{x}_b = \mathbf{0}$ for all $b \notin \mathcal{S}_B$, where $\mathcal{S}_B$ is a subset of $\mathcal{N}_B$ with cardinality $|\mathcal{S}_B| = k$.*

For simplicity, we consider that each block of the signal vector has equal length, that is, $|\mathcal{B}_b| = B$ for all $b \in \mathcal{N}_b$ with $BN_B = n$. It is clear that when $B = 1$, the block-sparse phase retrieval boils down to the ordinary or unstructured

sparse phase retrieval. Accordingly, we will henceforth focus on developing recovery algorithms for a block-sparse signal vector.

Adopting the least-squares criterion, the task of recovering a k -block sparse vector from m magnitude-only measurements can be cast as [11]

$$\underset{\mathbf{z} \in \mathcal{M}_B^k}{\text{minimize}} \quad \ell(\mathbf{z}) := \frac{1}{2m} \sum_{i=1}^m (\psi_i - |\mathbf{a}_i^T \mathbf{z}|)^2 \quad (2)$$

where $\mathcal{M}_B^k$ denotes the set of all k -block-sparse vectors of dimension n . Because of the nonconvex objective and the combinatorial constraint, the problem in (2) is in general NP-hard, hence computationally intractable.

For analytical concreteness, we focus on the real Gaussian model, which assumes $\mathbf{x} \in \mathbb{R}^n$, and independent and identically distributed (i.i.d.) sensing vectors follow $\mathbf{a}_i \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_n)$ for all $1 \leq i \leq m$. When there are enough measurements, it is reasonable to assume existence of a unique (up to a global sign) k -block-sparse solution $\{\pm \mathbf{x}\}$ to the quadratic system in (2). The critical goal of this paper is to put forth simple and scalable algorithms that can provably reconstruct $\mathbf{x}$ from as few magnitude-only measurements as possible.

III. COMPRESSIVE REWEIGHTED AMPLITUDE FLOW

This section presents the two stages, namely the initialization and the gradient refinement stages of CRAF. To begin, the distance from any estimate $\mathbf{z} \in \mathbb{R}^n$ to the solution set $\{\pm \mathbf{x}\} \subseteq \mathbb{R}^n$ is defined as $\text{dist}(\mathbf{z}, \mathbf{x}) := \min\{\|\mathbf{z} + \mathbf{x}\|_2, \|\mathbf{z} - \mathbf{x}\|_2\}$.

A. Sparse Spectral Initialization

A modified spectral initialization that utilizes the information from all available data samples is delineated first. Relative to existing phase retrieval initializations suggested in [9], [10], [11], [12], enhanced numerical performance is achieved by assigning judicious weights to all sampling vectors. Subsequently, the generalization of the new initialization procedure to compressive phase retrieval settings is justified.

1) *Spectral initialization*: Finding a good initialization is key in enabling strong convergence of iterative nonconvex optimization algorithms. Consider first the general phase retrieval, namely without exploiting the sparse prior information. Similar to past approaches, the new initialization entails estimating the norm $\|\mathbf{x}\|_2$ as well as the directional vector $\mathbf{d} := \mathbf{x}/\|\mathbf{x}\|_2$. Regarding the former, it has been well documented that the term $\hat{r} := \sqrt{(1/m) \sum_{i=1}^m \psi_i^2}$ is an unbiased and tightly concentrated estimate of the norm $r := \|\mathbf{x}\|_2$ when there are enough measurements [9]. The challenge remains to estimate the direction $\mathbf{d}$, namely seek a unit vector $\hat{\mathbf{d}}$ that is maximally correlated with $\mathbf{d}$.

Among different initialization strategies, the procedure proposed in [11] proves successful in achieving excellent numerical performance in estimating $\mathbf{d}$; see also [23] for robustified alternatives. However, the truncation therein discards the useful information carried over in a non-negligible portion

of samples. To exploit all the data samples, the new spectral initialization obtains the wanted approximation vector as

$$\hat{\mathbf{d}} := \arg \max_{\|\mathbf{z}\|_2=1} \mathbf{z}^T \left(\frac{\lambda^-}{|\mathcal{I}^-|} \sum_{i \in \mathcal{I}^-} \mathbf{a}_i \mathbf{a}_i^T + \frac{\lambda^+}{|\mathcal{I}^+|} \sum_{i \in \mathcal{I}^+} \mathbf{a}_i \mathbf{a}_i^T \right) \mathbf{z} \quad (3)$$

where $\lambda^- < 0$ and $\lambda^+ > 0$ are preselected coefficients, and the index sets $\mathcal{I}^- := \{i \in [m] : \psi_i^2 \leq \hat{r}^2/2\}$, and $\mathcal{I}^+ := \{i \in [m] : \psi_i^2 \geq \hat{r}^2/2\}$. It is worth pointing out that the judiciously devised index sets satisfy $\mathcal{I} = \mathcal{I}^- \cup \mathcal{I}^+$, suggesting full use of the available data samples. With $\hat{r}$ and $\hat{\mathbf{d}}$ at hand, the initial estimate of $\mathbf{x}$ can be obtained conveniently as $\mathbf{z}^0 := \hat{r}\hat{\mathbf{d}}$.

Intuitively, the initialization strategy in (3) can be justified as follows. Leveraging the rotational invariance of $\mathbf{a} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$, we have for any thresholds $\tau_1, \tau_2 \in [0, 1]$:

$$\begin{aligned} \mathbb{E}[\mathbf{a}\mathbf{a}^T | \langle \mathbf{a}, \mathbf{d} \rangle^2 \leq \tau_1] \\ = \mathbf{I}_n - \mathbf{d}\mathbf{d}^T + \mathbb{E}[\langle \mathbf{a}, \mathbf{d} \rangle^2 | \langle \mathbf{a}, \mathbf{d} \rangle^2 \leq \tau_1] \mathbf{d}\mathbf{d}^T \end{aligned} \quad (4)$$

$$\begin{aligned} \mathbb{E}[\mathbf{a}\mathbf{a}^T | \langle \mathbf{a}, \mathbf{d} \rangle^2 \geq \tau_2] \\ = \mathbf{I}_n - \mathbf{d}\mathbf{d}^T + \mathbb{E}[\langle \mathbf{a}, \mathbf{d} \rangle^2 | \langle \mathbf{a}, \mathbf{d} \rangle^2 \geq \tau_2] \mathbf{d}\mathbf{d}^T. \end{aligned} \quad (5)$$

It has been proved in [23, Lemma 3.2] that

$$\mathbb{E}[\langle \mathbf{a}, \mathbf{d} \rangle^2 | \langle \mathbf{a}, \mathbf{d} \rangle^2 \leq \tau_1] \leq \tau_1/3.$$

Therefore, the smallest eigenvalue of $\mathbb{E}[\mathbf{a}\mathbf{a}^T | \langle \mathbf{a}, \mathbf{d} \rangle^2 \leq \tau_1]$ satisfies

$$\lambda_n(\mathbb{E}[\mathbf{a}\mathbf{a}^T | \langle \mathbf{a}, \mathbf{d} \rangle^2 \leq \tau_1]) \leq \tau_1/3$$

whereas all other eigenvalues are

$$\lambda_i(\mathbb{E}[\mathbf{a}\mathbf{a}^T | \langle \mathbf{a}, \mathbf{d} \rangle^2 \leq \tau_1]) = 1, \quad 1 \leq i \leq n-1.$$

Similarly, one can establish the following lower bound for the second term $\mathbb{E}[\langle \mathbf{a}, \mathbf{d} \rangle^2 | \langle \mathbf{a}, \mathbf{d} \rangle^2 \geq \tau_2]$ in (5).

Lemma 1. Consider any nonzero signal vector $\mathbf{d} \in \mathbb{R}^n$ with $\|\mathbf{d}\|_2 = 1$. If $\mathbf{a} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$, then for any $\tau \geq 0$, it holds that

$$\mathbb{E}[\langle \mathbf{a}, \mathbf{d} \rangle^2 | \langle \mathbf{a}, \mathbf{d} \rangle^2 \geq \tau] \geq \frac{6 - \tau \operatorname{erf}(\sqrt{\tau})}{6 - 3 \operatorname{erf}(\sqrt{\tau})}. \quad (6)$$

Proof. With $\tilde{a} := \langle \mathbf{a}, \mathbf{d} \rangle$, it holds that $\tilde{a} \sim \mathcal{N}(0, 1)$. Let $\tilde{a}'$ be a random variable with the same density as $\tilde{a}$, and $p(\tilde{a})$ and $p(\tilde{a}')$ denote the density of $\tilde{a}$ and $\tilde{a}'$, respectively. It follows that

$$\begin{aligned} \mathbb{E}[\langle \mathbf{a}, \mathbf{d} \rangle^2 | \langle \mathbf{a}, \mathbf{d} \rangle^2 \geq \tau] &= \mathbb{E}[\tilde{a}^2 | \tilde{a}^2 \geq \tau] \\ &= \mathbb{E}[\tilde{a}^2 | |\tilde{a}| \geq \sqrt{\tau}] = \frac{\int_{\sqrt{\tau}}^{\infty} \tilde{a}^2 p(\tilde{a}) d\tilde{a}}{\int_{\sqrt{\tau}}^{\infty} p(\tilde{a}') d\tilde{a}'} \\ &= \frac{1 - \int_{-\sqrt{\tau}}^{\sqrt{\tau}} \tilde{a}^2 p(\tilde{a}) d\tilde{a}}{1 - \int_{-\sqrt{\tau}}^{\sqrt{\tau}} p(\tilde{a}') d\tilde{a}'} = \frac{2 - \mathbb{E}[\tilde{a}^2 | \tilde{a}^2 \leq \tau] \operatorname{erf}(\sqrt{\tau})}{2 - \operatorname{erf}(\sqrt{\tau})} \\ &\geq \frac{6 - \tau \operatorname{erf}(\sqrt{\tau})}{6 - 3 \operatorname{erf}(\sqrt{\tau})} \end{aligned}$$

where the last inequality relies on [23, Lemma 3.2]. $\square$

To help understanding the assertion of Lemma 1, taking $\tau = 0.5$ as an example, we find $\mathbb{E}[\langle \mathbf{a}, \mathbf{d} \rangle^2 | \langle \mathbf{a}, \mathbf{d} \rangle^2 \geq 0.5] \geq 1.42$ by substituting the inequality $\operatorname{erf}(\sqrt{0.5}) \geq 0.68$ into (6). Hence, it holds that $\lambda_1(\mathbb{E}[\langle \mathbf{a}, \mathbf{d} \rangle^2 | \langle \mathbf{a}, \mathbf{d} \rangle^2 \geq 0.5]) \geq 1.42$, and

$\lambda_i(\mathbb{E}[\mathbf{a}\mathbf{a}^T | \langle \mathbf{a}, \mathbf{x} \rangle^2 \geq 0.5]) = 1, \quad 1 \leq i \leq n-1$. Subsequently, it can be deduced that for $\lambda^- < 0$ and $\lambda^+ > 0$, the largest eigenvalue of $\lambda^- \mathbb{E}[\mathbf{a}\mathbf{a}^T | \langle \mathbf{a}, \mathbf{d} \rangle^2 \leq 0.5] + \lambda^+ \mathbb{E}[\mathbf{a}\mathbf{a}^T | \langle \mathbf{a}, \mathbf{d} \rangle^2 \geq 0.5]$ is greater than or equal to $1.42\lambda^+ + 0.166\lambda^-$, and all other eigenvalues equal $\lambda^+ + \lambda^-$. Therefore, once the sample mean matrix $\frac{\lambda^-}{|\mathcal{I}^-|} \sum_{i \in \mathcal{I}^-} \mathbf{a}_i \mathbf{a}_i^T + \frac{\lambda^+}{|\mathcal{I}^+|} \sum_{i \in \mathcal{I}^+} \mathbf{a}_i \mathbf{a}_i^T$ is sufficiently close to its mean, it would be possible to estimate $\mathbf{d}$ with high accuracy based on the matrix perturbation lemma in [30, Corollary 1], which is also included as Lemma 2 in the Appendix for completeness. The aforementioned arguments speak for the effectiveness of the proposed initialization, whereas the next theorem quantifies rigorously the initialization estimation error $\operatorname{dist}(\mathbf{z}^0, \mathbf{x})$.

Theorem 1. Let $\mathbf{z}_0 = \hat{r}\hat{\mathbf{d}}$ with $\hat{\mathbf{d}}$ obtained from (3). For any given constant $\delta_0 \in (0, 1)$, there exists numerical constants $c_0 > 0$ and C_0 such that the following holds

$$\operatorname{dist}(\mathbf{z}_0, \mathbf{x}) \leq \delta_0 \|\mathbf{x}\|_2$$

with probability at least $1 - 10 \exp(-c_0 m)$ when $m \geq C_0 n$.

For readability, the proof of Theorem 1 is deferred to Section V-A. Although the suggested initialization assumes a specific threshold $\hat{r}^2/2$ to split samples into $\mathcal{I}^-$ and $\mathcal{I}^+$, it is straightforward to incorporate two different thresholds $0 \leq \tilde{\tau}_1 \leq \tilde{\tau}_2 \leq 1$ such that $\mathcal{I}^- := \{i \in [m] : \psi_i^2 \leq \tilde{\tau}_1 \hat{r}^2\}$ and $\mathcal{I}^+ := \{i \in [m] : \psi_i^2 \geq \tilde{\tau}_2 \hat{r}^2\}$. By appropriately selecting $\tilde{\tau}_1$ and $\tilde{\tau}_2$, the initialization performance can be further boosted. It is worth pointing out that the weak recovery performance of similar procedures has been studied in [31], which only provides guarantee for the case of $n \rightarrow \infty$.

2) *Support recovery:* The initialization procedure in (3) is developed for general signal vectors $\mathbf{x}$, without leveraging the structural information that is present in diverse applications. When the vector is sparse, the required number of data samples to yield an accurate initialization can be reduced [6]. Next, we demonstrate how to obtain a sparse initialization based on the procedure discussed in Section III-A. Similar to [6], [27], obtaining a sparse initialization entails first estimating the (block)-support of the underlying (block)-sparse signal vectors.

Specifically, define random variables $Z_{i,j} := \psi_i^2 a_{i,j}^2, \forall j \in [n]$. According to [6, Eq. (16)], the following holds

$$\begin{aligned} \mathbb{E}\left[\sum_{j \in \mathcal{B}_b} Z_{i,j}^2\right] &= \mathbb{E}\left[\sum_{j \in \mathcal{B}_b} (\mathbf{a}_i^T \mathbf{x})^4 a_{i,j}^4\right] \\ &= 9B \|\mathbf{x}\|_2^4 + 24 \sum_{j \in \mathcal{B}_b} x_j^4 + 72 \|\mathbf{x}_b\|^2 \|\mathbf{x}\|_2^2. \end{aligned} \quad (7)$$

If $b \in \mathcal{S}_B$, then $\mathbf{x}_b \neq \mathbf{0}$, yielding $\mathbb{E}\left[\sum_{j \in \mathcal{B}_b} Z_{i,j}^2\right] > 9B \|\mathbf{x}\|_2^4 + 72 \|\mathbf{x}_b\|^2 \|\mathbf{x}\|_2^2$ in (7). On the contrary, if $b \notin \mathcal{S}_B$, one has $\mathbf{x}_b = \mathbf{0}$, yielding $\mathbb{E}\left[\sum_{j \in \mathcal{B}_b} Z_{i,j}^2\right] = 9B \|\mathbf{x}\|_2^4$. It is evident that there is a separation of at least $72 \|\mathbf{x}_b\|^2 \|\mathbf{x}\|_2^2$ in the expected values of $\sum_{j \in \mathcal{B}_b} Z_{i,j}^2$ for $b \in \mathcal{S}_B$ and $b \notin \mathcal{S}_B$. As long as the gap $72 \|\mathbf{x}_b\|^2 \|\mathbf{x}\|_2^2$ is large enough, the (block)-support set $\mathcal{S}_B$ can be recovered exactly in this way.

To estimate the (block)-support $\mathcal{S}_B$ in practice, compute first the so-called block marginals

$$\zeta_b := \sum_{j \in \mathcal{B}_b} \left(\frac{1}{m} \sum_{i=1}^m \psi_i^2 |a_{i,j}|^2 \right)^2, \quad \forall b \in \mathcal{N}_b$$

which serves as an empirical estimate of $\mathbb{E}[\sum_{j \in \mathcal{B}_b} Z_{i,j}^2]$. As explained earlier, the larger ζ_b is, the more likely is for the block to be nonzero, namely $\|\mathbf{x}_b\|_2 > 0$ [27]. Upon collecting $\{\zeta_b\}_{b=1}^{N_B}$, one can pick the indices associated with the k -largest values in $\{\zeta_b\}_{b=1}^{N_B}$, which form the estimated block-support set $\hat{\mathcal{S}}_B$. Subsequently, an estimate of the support of $\mathbf{x}$ denoted as $\hat{\mathcal{S}}$ can be determined as $\hat{\mathcal{S}} := \{i \in \mathcal{B}_b \mid \forall b \in \hat{\mathcal{S}}_B\}$.

The support estimation procedure is summarized in Steps 2-4 of Algorithm 1. Appealing to [27, Theorem 5.1] (also included as Lemma 3 for completeness in the Appendix), Steps 2-4 recover the support of $\mathbf{x}$ exactly with probability at least $1 - \frac{6}{m}$ provided that $m \geq C_0' k^2 B \log(mn)$ for some positive constant C_0' and the minimum block

$$x_{\min}^B := \min_{b \in \mathcal{S}_B} \|\mathbf{x}_b\|_2^2$$

is on the order of $(1/k)\|\mathbf{x}\|_2^2$, namely, $x_{\min}^B = (C_0''/k)\|\mathbf{x}\|_2^2$ for some number $C_0'' > 0$.

If the support has been exactly recovered, that is, $\hat{\mathcal{S}} = \mathcal{S}$, one can rewrite $\psi_i = |\mathbf{a}_i^\top \mathbf{x}| = |\mathbf{a}_{i,\hat{\mathcal{S}}}^\top \mathbf{x}_{\hat{\mathcal{S}}}|$ for all $i \in [m]$, where $\mathbf{a}_{i,\hat{\mathcal{S}}} \in \mathbb{R}^{k_B}$ contains entries of $\mathbf{a}_i$ whose indices belong to $\hat{\mathcal{S}}$; and likewise for $\mathbf{x}_{\hat{\mathcal{S}}} \in \mathbb{R}^k$. Then, the proposed initialization in (3) can be applied to the dimensionality-reduced data $\{(\mathbf{a}_{i,\hat{\mathcal{S}}}, \psi_i)\}_{i=1}^m$ to obtain

$$\hat{\mathbf{d}}_{\hat{\mathcal{S}}} := \max_{\mathbf{z} \in \mathbb{R}^{k_B}} \mathbf{z}^\top \left(\frac{\lambda^-}{|\mathcal{I}^-|} \sum_{i \in \mathcal{I}^-} \mathbf{a}_{i,\hat{\mathcal{S}}} \mathbf{a}_{i,\hat{\mathcal{S}}}^\top + \frac{\lambda^+}{|\mathcal{I}^+|} \sum_{i \in \mathcal{I}^+} \mathbf{a}_{i,\hat{\mathcal{S}}} \mathbf{a}_{i,\hat{\mathcal{S}}}^\top \right) \mathbf{z}.$$

Subsequently, an estimate of the n -dimensional vector $\mathbf{d}$ can be constructed by zero-padding entries of $\hat{\mathbf{d}}_{\hat{\mathcal{S}}}$ whose indices do not belong to $\hat{\mathcal{S}}$.

B. Refinement via Reweighted Gradient Iterations

Upon obtaining an accurate initial point, successive refinements based on reweighted gradient iterations are effected. To account for the block-sparsity structure of the wanted signal vector $\mathbf{x}$, the model-based iterative hard thresholding (M-IHT) [29] is invoked. To start, recall that the generalized gradient of the objective function in (2) is [11]

$$\nabla \ell(\mathbf{z}) := \frac{1}{m} \sum_{i \in [m]} \left(\mathbf{a}_i^\top \mathbf{z} - \psi_i \frac{\mathbf{a}_i^\top \mathbf{z}}{|\mathbf{a}_i^\top \mathbf{z}|} \right) \mathbf{a}_i \quad (8)$$

in which the convention $\mathbf{a}_i^\top \mathbf{z} / |\mathbf{a}_i^\top \mathbf{z}| := 0$ for $|\mathbf{a}_i^\top \mathbf{z}| = 0$ is adopted.

With $t \geq 0$ denoting the iteration count and $\mathbf{z}^0$ being the initial point, the M-IHT algorithm proceeds with the following k -block-sparse hard thresholding, namely

$$\mathbf{z}^{t+1} = \mathcal{H}_k^B \left(\mathbf{z}^t - \frac{\mu}{m} \nabla \ell(\mathbf{z}^t) \right) \quad (9)$$

Algorithm 1 Compressive Reweighted Amplitude Flow (CRAF)

- 1: **Input:** Data $\{(\mathbf{a}_i; \psi_i)\}_{i=1}^m$, block length B , and block sparsity level k ; initialization parameters $\lambda^- = -3$ and $\lambda^+ = 1$; step size $\mu = 1$; and weighting parameters $\{\beta_i = 0.6\}_{i=1}^m$, $\tau_w = 0.1$.
- 2: **For** $b = 1$ **to** N_B , **compute**

$$\zeta_b := \sum_{j \in \mathcal{B}_b} \left(\frac{1}{m} \sum_{i=1}^m \psi_i^2 |a_{i,j}|^2 \right)^2.$$

- 3: **Set** $\hat{\mathcal{S}}_B$ to include indices corresponding to the k -largest instances in $\{\zeta_b\}_{b=1}^{N_B}$.
- 4: **Set** $\hat{\mathcal{S}}$ to comprise indices of $\mathcal{B}_b$ for $b \in \hat{\mathcal{S}}_B$.
- 5: **Compute** the principal eigenvector $\hat{\mathbf{d}}_{\hat{\mathcal{S}}} \in \mathbb{R}^{k_B}$ of

$$\frac{\lambda^-}{|\mathcal{I}^-|} \sum_{i \in \mathcal{I}^-} \mathbf{a}_{i,\hat{\mathcal{S}}} \mathbf{a}_{i,\hat{\mathcal{S}}}^\top + \frac{\lambda^+}{|\mathcal{I}^+|} \sum_{i \in \mathcal{I}^+} \mathbf{a}_{i,\hat{\mathcal{S}}} \mathbf{a}_{i,\hat{\mathcal{S}}}^\top$$

where $\mathcal{I}^- := \{i \in [m] : \psi_i^2 \leq \hat{r}^2/2\}$ and $\mathcal{I}^+ := \{i \in [m] : \psi_i^2 \geq \hat{r}^2/2\}$ with $\hat{r} := \sqrt{\sum_{i=1}^m \psi_i^2 / m}$.

- 6: **Initialize** $\mathbf{z}^0$ as $\hat{r} \tilde{\mathbf{d}}$, where $\tilde{\mathbf{d}} \in \mathbb{R}^n$ is given by augmenting $\hat{\mathbf{d}}_{\hat{\mathcal{S}}}$ in Step 5 with $\tilde{d}_i = 0$ for $i \notin \hat{\mathcal{S}}$.
- 7: **Loop:** **For** $t = 0$ **to** $T - 1$

$$\mathbf{z}^{t+1} = \mathcal{H}_k^B \left(\mathbf{z}^t - \frac{\mu}{m} \sum_{i \in [m]} w_i^t \left(\mathbf{a}_i^\top \mathbf{z}^t - \psi_i \frac{\mathbf{a}_i^\top \mathbf{z}^t}{|\mathbf{a}_i^\top \mathbf{z}^t|} \right) \mathbf{a}_i \right)$$

where $w_i^t := \max \left\{ \tau_w, \frac{|\mathbf{a}_i^\top \mathbf{z}^t|}{|\mathbf{a}_i^\top \mathbf{z}^t| + \psi_i \beta_i} \right\}$.

- 8: **Output:** $\mathbf{z}^T$.
-

where $\mu > 0$ is the preselected step size, and the block-sparse hard thresholding operator $\mathcal{H}_k^B(\bar{\mathbf{u}}) : \mathbb{R}^n \rightarrow \mathbb{R}^n$ converts an n -dimensional vector $\bar{\mathbf{u}} := [\bar{\mathbf{u}}_1^\top \dots \bar{\mathbf{u}}_{N_B}^\top]^\top$ into a k -block-sparse one $\mathbf{u} := [\mathbf{u}_1^\top \dots \mathbf{u}_{N_B}^\top]^\top$ such that

$$\mathbf{u}_b = \begin{cases} \bar{\mathbf{u}}_b, & \text{if } b \in \mathcal{U}_B \\ \mathbf{0}, & \text{if } b \notin \mathcal{U}_B \end{cases}$$

where $\mathcal{U}_B$ comprises indices corresponding to the k -largest entities in $\{\|\bar{\mathbf{u}}_b\|_2\}_{b=1}^{N_B}$.

Unfortunately, the negative gradient $-\nabla \ell(\mathbf{z})$ may not drag the iterate sequence $\{\mathbf{z}^t\}$ to the global optimum $\mathbf{x}$ because the estimated sign $\mathbf{a}_i^\top \mathbf{z} / |\mathbf{a}_i^\top \mathbf{z}|$ in $\nabla \ell(\mathbf{z})$ may not coincide with the true one $\mathbf{a}_i^\top \mathbf{x} / |\mathbf{a}_i^\top \mathbf{x}|$ [11]. As a consequence, the update in (9) may not always reduce the distance of the iterate to the global optimum. To alleviate the negative influence of the erroneously estimated signs, SPARTA implements the following truncated gradient $\nabla \ell_{\text{tr}}(\mathbf{z}^t)$ [6]

$$\nabla \ell_{\text{tr}}(\mathbf{z}^t) := \frac{1}{m} \sum_{i \in \mathcal{I}^t} \left(\mathbf{a}_i^\top \mathbf{z}^t - \psi_i \frac{\mathbf{a}_i^\top \mathbf{z}^t}{|\mathbf{a}_i^\top \mathbf{z}^t|} \right) \mathbf{a}_i \quad (10)$$

where

$$\mathcal{I}^t := \left\{ 1 \leq i \leq m \mid \left| \frac{\mathbf{a}_i^\top \mathbf{z}^t}{|\mathbf{a}_i^\top \mathbf{x}|} \right| \geq \tau_g \right\}$$

for some preselected truncation parameter. It is clear that $\nabla \ell_{\text{tr}}(\mathbf{z})$ is based on data samples whose associated $|\mathbf{a}_i^\top \mathbf{x}|$ is of

relatively large sizes. The reason for this gradient truncation is that gradients (summands in (10)) of large $|\mathbf{a}_i^\top \mathbf{z}|/|\mathbf{a}_i^\top \mathbf{x}|$ provably point toward the global optimum $\mathbf{x}$ with high probability [11]. However, as pointed out in [12], the truncation operation may reject meaningful samples, which hampers the efficacy of $\nabla \ell_{\text{tr}}$ especially when the sample size is limited.

An alternative to the truncation trimming procedure is to introduce different weights for different gradients [12], which helps fusing useful information from all gradient directions. Specifically, the ensuing reweighted gradient used in [12] proves successful in phase retrieval of general signal vectors

$$\nabla \ell_{\text{rw}}(\mathbf{z}^t) := \frac{1}{m} \sum_{i \in [m]} w_i^t \left(\mathbf{a}_i^\top \mathbf{z}^t - \psi_i \frac{\mathbf{a}_i^\top \mathbf{z}^t}{|\mathbf{a}_i^\top \mathbf{z}^t|} \right) \mathbf{a}_i \quad (11)$$

where the weights are given by

$$w_i^t := \max \left\{ \tau_w, \frac{|\mathbf{a}_i^\top \mathbf{z}^t|}{|\mathbf{a}_i^\top \mathbf{z}^t| + \psi_i \beta_i} \right\}, \quad \forall i \in [m] \quad (12)$$

for certain preselected parameters $\tau_w > 0$ and $\beta_i > 0$ for all $i \in [m]$. Evidently, it holds that $\tau_w \leq w_i^t \leq 1$ for all $i \in [m]$, and the larger the ratio $|\mathbf{a}_i^\top \mathbf{z}^t|/|\mathbf{a}_i^\top \mathbf{x}|$, the larger the weight w_i^t . In this sense, w_i^t reflects the confidence in the i -th negative gradient pointing toward the global optimum $\mathbf{x}$.

In the context of phase retrieval of block-sparse vectors, it is thus reasonable to implement the M-IHT based iteration using reweighted gradients, namely

$$\mathbf{z}^{t+1} := \mathcal{H}_k^B(\mathbf{z}^t - \mu \nabla \ell_{\text{rw}}(\mathbf{z}^t)). \quad (13)$$

The proposed block-sparse phase retrieval solver is summarized in Algorithm 1, whose exact recovery is established in the following theorem.

Theorem 2. *Let $\mathbf{x} \in \mathbb{R}^n$ be any k -block-sparse ($kB \ll n$) vector with $\mathbf{x}_{\min}^B := (C_0^B/k)\|\mathbf{x}\|_2^2$. Consider noiseless measurements $\{\psi_i = |\mathbf{a}_i^\top \mathbf{x}|\}_{i=1}^m$ from the real Gaussian model. If $m \geq C_1 k^2 B \log(mn)$, there exists a constant learning rate $\mu > 0$, such that the successive estimates $\mathbf{z}^t$ in Algorithm 1 obey*

$$\|\mathbf{z}^t - \mathbf{x}\|_2 \leq \delta_0 \rho^t \|\mathbf{x}\|_2, \quad t = 0, 1, \dots \quad (14)$$

with probability at least $1 - c_2 \exp(-c_1 m) - 6/m$. Here, $0 < \delta_0 < 1$, $0 < \rho < 1$, $\mu, c_1 > 0$, $c_2 > 0$, C_0^B , and C_1 are certain numerical constants.

The proof of Theorem 2 is provided in Section V-B. Regarding its implication, a couple of observations come in order. To start, as soon as $m \geq C_1 k^2 B \log(mn)$, CRAF recovers exactly k -block-sparse vectors $\mathbf{x}$ of non-negligible blocks. This sample complexity is consistent with the Block CoPRAM method in [27]. Furthermore, CRAF converges exponentially fast. Expressed differently, it takes CRAF at most $T := \mathcal{O}(\log(1/\epsilon))$ iterations to reach a solution of ϵ -relative accuracy.

IV. NUMERICAL TESTS

This section demonstrates the efficacy of the proposed initialization and the CRAF algorithm relative to the state-of-the-art approaches for sparse phase retrieval, including

SPARTA [6] and CoPRAM [27]. In all experiments, the support $\mathcal{S}$ of the true signal vectors $\mathbf{x} \in \mathbb{R}^{3,000}$ was randomly chosen. The nonzero entries were generated using $\mathbf{x}_{\mathcal{S}} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$. The obtained $\mathbf{x}$ was subsequently normalized such that $\|\mathbf{x}\|_2 = 1$. The sampling vectors were generated using $\mathbf{a}_i \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$, $1 \leq i \leq m$. For SPARTA, its suggested parameters were used. The parameters of CRAF were set as $\lambda^- = -3$, $\lambda^+ = 1$, $\{\beta_i = 0.6\}_{i=1}^m$, $\tau_w = 0.1$, and $\mu = 1$. For all simulated algorithms, the maximum iterations were fixed to $T = 1,000$, and all reported results are averaged over 100 Monte Carlo simulations.

The first experiment evaluates the performance of our initialization relative to that in SPARTA [6] and CoPRAM [27] for block length $B = 1$. Figure 1 depicts the average relative error of the three initialization schemes with the sparsity level k varying from 25 to 35, and m/k fixed to 30. Clearly, the new initialization outperforms the other two with large margins.

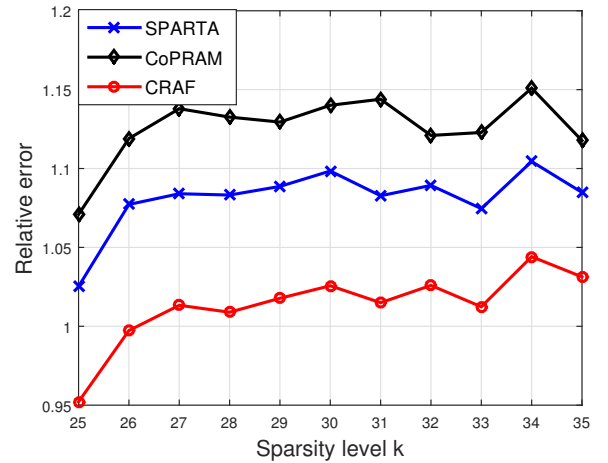


Fig. 1: Average relative error for $B = 1$.

The second experiment examines the empirical success rates of CRAF, SPARTA, and CoPRAM for solving the ordinary compressive phase retrieval with $B = 1$. Each of the 100 Monte Carlo trials is declared a success if the relative error $\text{dist}(\mathbf{z}^T, \mathbf{x})/\|\mathbf{x}\|_2$ is less than 10^{-5} . The empirical success rates of CRAF, SPARTA, and CoPRAM are presented in Fig. 2 with m increasing from 400 to 1,800. Notably, the curves showcase improved exact recovery performance of CRAF relative to its competing alternatives. Since in certain applications, the sparsity level k may not be accurately known, it is desirable to have the compressive phase retrieval algorithms remain operational for unknown or inexact k values. Let $\hat{k}$ be an estimate of the sparsity level k . The recovery performance of CRAF is tested with $\hat{k}$ set as the upper limit of the theoretically affordable sparsity level, namely $\sqrt{3,000} \approx 55$. From Fig. 3, it is clear that CRAF offers the best numerical performance for unknown k . A careful comparison between Figs. 2 and 3 demonstrates that CRAF is more robust to unknown k values than CoPRAM.

The next experiment tests the performance of CRAF relative to that of Block CoPRAM and SPARTA for 20-block-sparse phase retrieval with block length $B = 2$. The empirical success

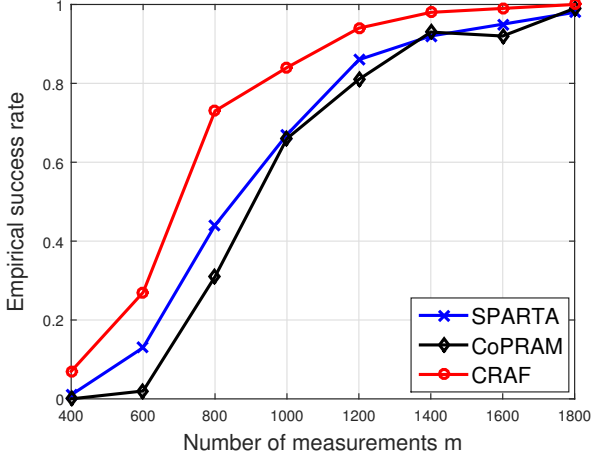


Fig. 2: Empirical success rate versus m for $B = 1$, $k = 30$.

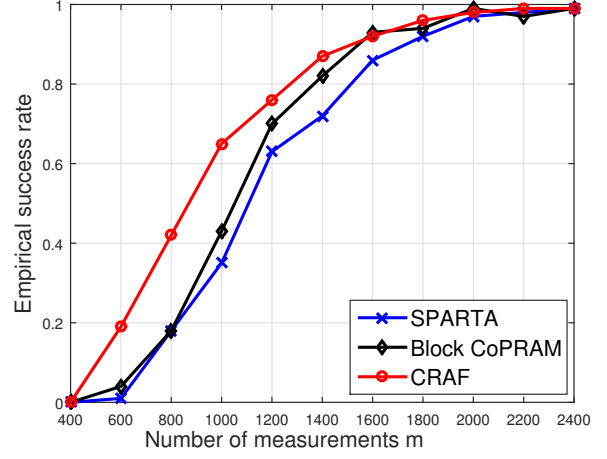


Fig. 4: Empirical success rate versus m for $B = 2$, $k = 20$.

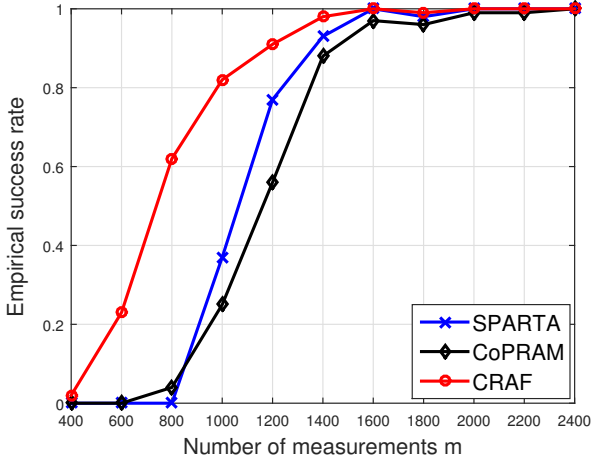


Fig. 3: Empirical success rate versus m for $B = 1$, $k = 30$, and $\hat{k} = 55$.

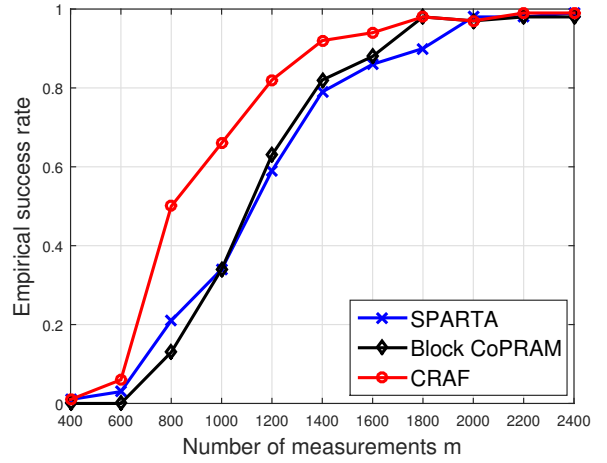


Fig. 5: Empirical success rate versus m for $B = 2$, $k = 20$, and $\hat{k} = 28$.

rates for the three schemes from 100 independent trials with k known and unknown are reported in Figs. 4 and 5, respectively. In both cases, CRAF yields the best recovery performance.

The last experiment validates the robustness of CRAF with respect to noisy measurements of the following form:

$$\psi_i = |\mathbf{a}_i^T \mathbf{x}| + \eta_i, \quad 1 \leq i \leq m$$

where $\{\eta_i\}$ are independently sampled from $\mathcal{N}(0, \sigma^2)$. In this experiment, $k = 30$, $B = 1$, and $m = 1,600$ were simulated. Figure 6 depicts the relative errors of the three approaches versus varying σ^2 from 0.1 to 0.6, from which it is clear that CRAF offers the most accurate estimates for all noise levels. In other words, CRAF achieves improved robustness relative to SPARTA and CoPRAM.

V. PROOFS

The proofs of Theorem 1 and 2 are presented next. The proof of Theorem 1 is mainly based on that of [23, Prop. 2], whereas the proof of Theorem 2 builds upon [6, Thm. 1].

A. Proof of Theorem 1

For ease of presentation, some notation is established first. To begin, let

$$\mathbf{M}^- := \frac{1}{|\mathcal{I}^-|} \sum_{i \in \mathcal{I}^-} \mathbf{a}_i \mathbf{a}_i^T, \text{ and } \mathbf{M}^+ := \frac{1}{|\mathcal{I}^+|} \sum_{i \in \mathcal{I}^+} \mathbf{a}_i \mathbf{d}_i^T$$

denote the first and second parts of the matrix used in the new initialization procedure (3). Upon defining

$$\begin{aligned} \phi(\tau^-) &:= \mathbb{E}[\langle \mathbf{a}, \mathbf{d} \rangle^2 | \langle \mathbf{a}, \mathbf{d} \rangle^2 \leq \tau^-] - 1 \\ \phi(\tau^+) &:= \mathbb{E}[\langle \mathbf{a}, \mathbf{d} \rangle^2 | \langle \mathbf{a}, \mathbf{d} \rangle^2 \geq \tau^+] - 1 \end{aligned}$$

it can be verified that

$$\mathbb{E}[\mathbf{M}^-] = \mathbf{I}_n + \phi(\tau^-) \mathbf{d} \mathbf{d}^T, \text{ and } \mathbb{E}[\mathbf{M}^+] = \mathbf{I}_n + \phi(\tau^+) \mathbf{d} \mathbf{d}^T.$$

Without loss of generality, one can then write

$$\begin{aligned} \mathbf{M}^- &= \mathbf{I}_n + \phi^-(\epsilon) \mathbf{d} \mathbf{d}^T + \mathbf{\Delta}^- \\ \mathbf{M}^+ &= \mathbf{I}_n + \phi^+(\epsilon) \mathbf{d} \mathbf{d}^T + \mathbf{\Delta}^+ \end{aligned}$$

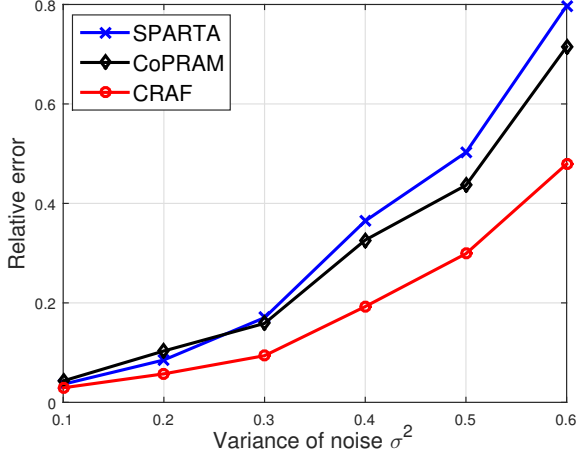


Fig. 6: Average relative error of signal recovery versus the variance of noise σ^2 for $B = 1$, $k = 30$, and $m = 1600$.

where Δ^- (Δ^+) describes the difference between M^- (M^+) and their mean $\mathbb{E}[M^-] = \mathbf{I}_n + \phi^-(\epsilon)\mathbf{d}\mathbf{d}^T$ ($\mathbb{E}[M^+] = \mathbf{I}_n + \phi^+(\epsilon)\mathbf{d}\mathbf{d}^T$). Appealing to a standard eigenvector perturbation result [30, Corollary 1] (also included as Lemma 2 in the Appendix), to bound $\text{dist}(\hat{\mathbf{d}}, \mathbf{d})$, it suffices to bound $\|\lambda^-\Delta^- + \lambda^+\Delta^+\|_2$.

It has been established in [23, Proposition 2] that for arbitrarily small $\delta^- \in (0, 1)$, and some absolute constants $c^- > 0$, $C^- < \infty$ dependent on δ^- , the following holds with probability at least $1 - 5\exp(-c^-m)$

$$\|\Delta^-\|_2 \leq \delta^- \quad (15)$$

whenever $m \geq C^-n$.

The task now remains to bound $\|\Delta^+\|_2$. To account for the estimation error in $\hat{r}$, the following two index sets surrounding $\mathcal{I}^+$ are introduced [23, Proposition 2]

$$\begin{aligned} \mathcal{I}_{-\epsilon}^+ &:= \left\{ i \in [m] \mid \langle \mathbf{a}_i, \mathbf{d} \rangle^2 > \frac{1-\epsilon}{2} \right\} \\ \mathcal{I}_{+\epsilon}^+ &:= \left\{ i \in [m] \mid \langle \mathbf{a}_i, \mathbf{d} \rangle^2 \geq \frac{1+\epsilon}{2} \right\} \end{aligned}$$

for some numerical constant $\epsilon \in (0, 1)$. For convenience, one can set $\kappa := \exp((\epsilon - 1)/4)/\sqrt{\pi(1 - \epsilon)}$ which upper bounds $P(\langle \mathbf{a}, \mathbf{d} \rangle^2 \in [\frac{1-\epsilon}{2}, \frac{1+\epsilon}{2}])/\epsilon$, and

$$p_0(\epsilon) := P\left(\langle \mathbf{a}, \mathbf{d} \rangle^2 \geq \frac{1+\epsilon}{2}\right).$$

It can be readily checked that

$$\begin{aligned} p_0(\epsilon) &= P\left(\chi_1^2 \geq \frac{1+\epsilon}{2}\right) = 1 - P\left(\chi_1^2 \leq \frac{1+\epsilon}{2}\right) \\ &\geq 1 - \sqrt{\frac{1+\epsilon}{2} \exp\left(\frac{1-\epsilon}{2}\right)} \end{aligned}$$

leveraging the tail bound of the χ_1^2 distribution.

Subsequently, five events denoted as $\{\mathcal{E}_i\}_{i=1}^5$ occurring with high probability are introduced in (16). The constants $p \geq 1$ and $q \geq 1$ in (16) satisfy $1/p + 1/q = 1$. On the event $\mathcal{E}_1$, it

can be verified that $\hat{r}^2 \in [(1 - \epsilon)r^2, (1 + \epsilon)r^2]$, hence $\mathcal{I}_{+\epsilon}^+ \subset \mathcal{I}^+ \subset \mathcal{I}_{-\epsilon}^+$. When all five events $\{\mathcal{E}_i\}_{i=1}^5$ occur, $\|\Delta^+\|_2$ can be bounded. In detail, Δ^+ can be rewritten as in (17), implying that $\Delta^+ = \Delta_1^+ + \Delta_2^+ + \Delta_3^+$. The three terms $\|\Delta_1^+\|_2$, $\|\Delta_2^+\|_2$, and $\|\Delta_3^+\|_2$ are bounded next. Note that on the event $\mathcal{E}_5$, it holds that $\|\Delta_1^+\|_2 \leq \epsilon$. To bound $\|\Delta_2^+\|_2$, observe that on the event $\mathcal{E}_3$ and $\mathcal{E}_4$, the following are true

$$\begin{aligned} \|\Delta_2^+\|_2 &= \frac{m}{|\mathcal{I}_{+\epsilon}^+|} \left\| \frac{1}{m} \sum_{i \in \mathcal{I}^+ \setminus \mathcal{I}_{+\epsilon}^+} \mathbf{a}_i \mathbf{a}_i^T \right\|_2 \\ &\leq \frac{m}{|\mathcal{I}_{+\epsilon}^+|} \left\| \frac{1}{m} \sum_{i \in \mathcal{I}_{-\epsilon}^+ \setminus \mathcal{I}_{+\epsilon}^+} \mathbf{a}_i \mathbf{a}_i^T \right\|_2 \\ &\leq \frac{2}{p_0(\epsilon)} \cdot 4q(\kappa\epsilon)^{1/p} \end{aligned}$$

where the last inequality arises from the definitions of $\mathcal{E}_3$ and $\mathcal{E}_4$. Regarding $\|\Delta_3^+\|_2$, the next holds true

$$\begin{aligned} \|\Delta_3^+\|_2 &= m \frac{|\mathcal{I}^+| - |\mathcal{I}_{+\epsilon}^+|}{|\mathcal{I}_{+\epsilon}^+| |\mathcal{I}^+|} \left\| \frac{1}{m} \sum_{i \in \mathcal{I}^+} \mathbf{a}_i \mathbf{a}_i^T \right\|_2 \\ &\leq m \frac{|\mathcal{I}^+| - |\mathcal{I}_{+\epsilon}^+|}{|\mathcal{I}_{+\epsilon}^+| |\mathcal{I}^+|} \left\| \frac{1}{m} \sum_{i=1}^m \mathbf{a}_i \mathbf{a}_i^T \right\|_2 \\ &\leq \frac{8\epsilon(1 + \epsilon)\kappa}{p_0(\epsilon)^2}. \end{aligned}$$

To sum, the following is true

$$\|\Delta^+\|_2 = \|\Delta_1^+\|_2 + \|\Delta_2^+\|_2 + \|\Delta_3^+\|_2 \quad (18)$$

$$\leq \epsilon + \frac{2}{p_0(\epsilon)} \cdot 4q(\kappa\epsilon)^{1/p} + \frac{8\epsilon(1 + \epsilon)\kappa}{p_0(\epsilon)^2}. \quad (19)$$

Taking $p = 1 + \frac{1}{\log \frac{1}{\kappa\epsilon}}$ and $q = 1 + \log \frac{1}{\kappa\epsilon}$, one has $\|\Delta^+\|_2 \leq \delta^+$, with

$$\delta^+ := \epsilon + \frac{8\epsilon(1 - \log \kappa\epsilon)\kappa\epsilon}{p_0(\epsilon)} + \frac{8\epsilon(1 + \epsilon)\kappa}{p_0(\epsilon)^2}.$$

Since $p_0(\epsilon)$ and κ are bounded away from 0 for sufficiently small $\epsilon > 0$, δ^+ approaches 0 as ϵ approaches 0. Based on the established bounds on $\|\Delta^-\|_2$ in (15) and $\|\Delta^+\|_2$ in (18), one has

$$\|\lambda^-\Delta^- + \lambda^+\Delta^+\|_2 \leq \lambda^+\delta^+ + \lambda^-\delta^-.$$

From Lemma 2 in the Appendix, the next can be deduced

$$\text{dist}^2(\hat{\mathbf{d}}, \mathbf{d}) \leq 2 - 2|\langle \hat{\mathbf{d}}, \mathbf{d} \rangle| \leq \left(\frac{2(\lambda^+\delta^+ + \lambda^-\delta^-)}{\lambda^+\phi(\tau^+) + \lambda^-\phi(\tau^-)} \right)^2$$

implying

$$\text{dist}(\hat{\mathbf{d}}, \mathbf{d}) \leq \frac{2(\lambda^+\delta^+ + \lambda^-\delta^-)}{\lambda^+\phi(\tau^+) + \lambda^-\phi(\tau^-)}. \quad (20)$$

Combining $\mathcal{E}_1$ and the bound in (20) gives rise to

$$\begin{aligned} \text{dist}(\mathbf{z}^0, \mathbf{x}) &\leq \hat{r} \text{dist}(\hat{\mathbf{d}}, \mathbf{d}) + |r - \hat{r}| \\ &\leq \left(\sqrt{1 + \epsilon} \text{dist}(\hat{\mathbf{d}}, \mathbf{d}) + \sqrt{1 + \epsilon} - 1 \right) \|\mathbf{x}\|_2. \end{aligned}$$

Letting $\delta_0 := \frac{2\sqrt{1+\epsilon}(\lambda^+\delta^+ + \lambda^-\delta^-)}{\lambda^+\phi(\tau^+) + \lambda^-\phi(\tau^-)} + \sqrt{1 + \epsilon} - 1$, we have $\text{dist}(\mathbf{z}^0, \mathbf{x}) \leq \delta_0 \|\mathbf{x}\|_2$. It is worth stressing that $\lim_{\epsilon \rightarrow 0} \delta_0 =$

$$\begin{aligned}
\mathcal{E}_1 &:= \left\{ \frac{1}{m} \|\mathbf{A}^\mathcal{T} \mathbf{A}\|_2 \in [1 - \epsilon, 1 + \epsilon] \right\}, & \mathcal{E}_2 &:= \{ |\mathcal{I}_{-\epsilon}^+| \leq |\mathcal{I}_{+\epsilon}^+| + 2\epsilon\kappa m \} \\
\mathcal{E}_3 &:= \left\{ |\mathcal{I}_{+\epsilon}^+| \geq \frac{1}{2} m p_0(\epsilon) \right\}, & \mathcal{E}_4 &:= \left\{ \left\| \frac{1}{m} \sum_{i \in \mathcal{I}_{-\epsilon}^+ \setminus \mathcal{I}_{+\epsilon}^+} \mathbf{a}_i \mathbf{a}_i^\mathcal{T} \right\|_2 \leq 4q(\kappa\epsilon)^{\frac{1}{p}} \right\} \\
\mathcal{E}_5 &:= \left\{ \left\| \frac{1}{|\mathcal{I}_{+\epsilon}^+|} \sum_{i \in \mathcal{I}_{+\epsilon}^+} \left[\mathbf{a}_i \mathbf{a}_i^\mathcal{T} - \left(\mathbf{I}_n + \phi^+ \left(\frac{1+\epsilon}{2} \right) \mathbf{x} \mathbf{x}^\mathcal{T} \right) \right] \right\|_2 \leq \epsilon \right\}
\end{aligned} \tag{16}$$

$$\begin{aligned}
\Delta^+ &= \mathbf{M}^+ - \left(\mathbf{I}_n + \phi^+ \left(\frac{1+\epsilon}{2} \right) \mathbf{d} \mathbf{d}^\mathcal{T} \right) \\
&= \underbrace{\frac{1}{|\mathcal{I}_{+\epsilon}^+|} \sum_{i \in \mathcal{I}_{+\epsilon}^+} \left[\mathbf{a}_i \mathbf{a}_i^\mathcal{T} - \left(\mathbf{I}_n + \phi^+ \left(\frac{1+\epsilon}{2} \right) \mathbf{d} \mathbf{d}^\mathcal{T} \right) \right]}_{:= \Delta_1^+} + \underbrace{\frac{1}{|\mathcal{I}_{+\epsilon}^+|} \sum_{i \in \mathcal{I}^+ \setminus \mathcal{I}_{+\epsilon}^+} \mathbf{a}_i \mathbf{a}_i^\mathcal{T}}_{:= \Delta_2^+} - \underbrace{\left(\frac{1}{|\mathcal{I}_{+\epsilon}^+|} - \frac{1}{|\mathcal{I}^+|} \right) \sum_{i \in \mathcal{I}^+} \mathbf{a}_i \mathbf{a}_i^\mathcal{T}}_{:= \Delta_3^+}
\end{aligned} \tag{17}$$

0, suggesting that δ_0 can be brought arbitrarily close to 0 by increasing m .

So far, it has been proved that $\text{dist}(\mathbf{z}^0, \mathbf{x}) \leq \delta_0 \|\mathbf{x}\|_2$ on the events $\{\mathcal{E}_i\}_{i=1}^5$. The next step is to show the five events occur simultaneously with high probability. Recall that it has been shown in [23, Proposition 2] that each of the events $\mathcal{E}_1, \mathcal{E}_2$, and $\mathcal{E}_4$ occurs with probability at least $1 - \exp(-c^- m)$ when $m > C^- n$.

To complete the proof, we first show that

$$P(\mathcal{E}_3) \geq 1 - \exp\left(-\frac{m p_0(\epsilon)^2}{2}\right).$$

To that end, rewrite $|\mathcal{I}_{+\epsilon}^+|$ as

$$|\mathcal{I}_{+\epsilon}^+| = \sum_{i=1}^m \mathbb{1}_{\{\langle \mathbf{a}_i, \mathbf{d} \rangle^2 \geq (1+\epsilon)/2\}}.$$

Since $\{\mathbb{1}_{\{\langle \mathbf{a}_i, \mathbf{d} \rangle^2 \geq (1+\epsilon)/2\}}\}_{i=1}^m$ are i.i.d. Bernoulli random variables with

$$P(\langle \mathbf{a}_i, \mathbf{d} \rangle^2 \geq (1+\epsilon)/2) \geq p_0(\epsilon), \quad \forall i \in [m]$$

the following holds

$$P(|\mathcal{I}_{+\epsilon}^+| \leq \frac{1}{2} m p_0(\epsilon)) \leq \exp\left(-\frac{m p_0(\epsilon)^2}{2}\right)$$

by Hoeffding's inequality [32]. Therefore, $P(\mathcal{E}_3) \geq 1 - \exp\left(-\frac{m p_0(\epsilon)^2}{2}\right)$. Similar to [23, Lemma A.6], it can be shown that for

$$m \geq \log^2 p_0(\epsilon) n / c^+ \epsilon^2 p_0(\epsilon)$$

with some absolute constant $c^+ > 0$,

$$P(\mathcal{E}_5 | \mathcal{E}_3) \geq 1 - \exp\left(-\frac{c^+ \epsilon^2 m p_0(\epsilon)}{\log^2 p_0(\epsilon)}\right).$$

Thus, setting $c_0 = \min\left\{\frac{c^+ \epsilon^2 p_0(\epsilon)}{\log^2 p_0(\epsilon)}, \frac{p_0(\epsilon)^2}{2}, c^-\right\}$ and

$$C_0 = \max\{\log^2 p_0(\epsilon) / c^+ \epsilon^2 p_0(\epsilon), C^-\}$$

confirms the assertion of Theorem 1.

B. Proof of Theorem 2

Some notation used only for this section is introduced first. For all $t \geq 0$, let

$$\mathbf{v}^{t+1} := \mathbf{z}^t - \frac{\mu}{m} \sum_{i=1}^m w_i^t \left(\mathbf{a}_i^\mathcal{T} \mathbf{z}^t - \psi_i \frac{\mathbf{a}_i^\mathcal{T} \mathbf{z}^t}{|\mathbf{a}_i^\mathcal{T} \mathbf{z}^t|} \right) \mathbf{a}_i$$

represent the estimate prior to effecting the hard thresholding operation in (9). The support of $\mathbf{x}$ and $\mathbf{z}^t$ is denoted as $\mathcal{S}$ and $\hat{\mathcal{S}}^t$, respectively. Hence, the support for the reconstruction error $\mathbf{h}^t := \mathbf{x} - \mathbf{z}^t$ defined as Ω^t is given by $\mathcal{S} \cup \hat{\mathcal{S}}^t$. Additionally, let $\Omega^t \setminus \Omega^{t+1}$ be the difference between sets Ω^t and Ω^{t+1} . Evidently, it holds that $|\mathcal{S}| = |\hat{\mathcal{S}}^t| = s$ for $s := kB$, which implies $|\Omega^t| \leq 2s$, $|\Omega^t \setminus \Omega^{t+1}| \leq 2s$, and $|\Omega^t \cup \Omega^{t+1}| \leq 3s$, $\forall t \geq 0$. Vectors with sets as subscript, e.g., $\mathbf{v}_{\Omega^t}$, are formed by zeroing all entries of the vector except for those in the set.

According to the triangle inequality of the vector 2-norm, it holds that

$$\begin{aligned}
\|\mathbf{x} - \mathbf{z}^{t+1}\|_2 &= \|\mathbf{x} - \mathbf{v}_{\Omega^{t+1}}^{t+1} + \mathbf{v}_{\Omega^{t+1}}^{t+1} - \mathbf{z}^{t+1}\|_2 \\
&\leq \|\mathbf{x} - \mathbf{v}_{\Omega^{t+1}}^{t+1}\|_2 + \|\mathbf{z}^{t+1} - \mathbf{v}_{\Omega^{t+1}}^{t+1}\|_2 \\
&\leq 2\|\mathbf{x}_{\Omega^{t+1}} - \mathbf{v}_{\Omega^{t+1}}^{t+1}\|_2.
\end{aligned} \tag{21}$$

The last inequality in (21) comes from $\|\mathbf{z}^{t+1} - \mathbf{v}_{\Omega^{t+1}}^{t+1}\|_2 \leq \|\mathbf{x}_{\Omega^{t+1}} - \mathbf{v}_{\Omega^{t+1}}^{t+1}\|_2$ since $\mathbf{z}^{t+1}$ achieves the minimal distance to $\mathbf{v}_{\Omega^{t+1}}^{t+1}$ among all vectors belonging to $\mathcal{M}_B^k$ and supported on Ω^{t+1} . Substituting the definitions of $\mathbf{h}^t$ and $\mathbf{v}^t$ into (21), one arrives at

$$\begin{aligned}
\frac{1}{2} \|\mathbf{h}^{t+1}\|_2 &\leq \left\| \mathbf{h}_{\Omega^{t+1}}^t - \frac{\mu}{m} \sum_{i=1}^m w_i^t \mathbf{a}_i^\mathcal{T} \mathbf{h}^t \mathbf{a}_{i, \Omega^{t+1}} \right. \\
&\quad \left. - \frac{\mu}{m} \sum_{i=1}^m w_i^t \left(\frac{\mathbf{a}_i^\mathcal{T} \mathbf{z}^t}{|\mathbf{a}_i^\mathcal{T} \mathbf{z}^t|} - \frac{\mathbf{a}_i^\mathcal{T} \mathbf{x}}{|\mathbf{a}_i^\mathcal{T} \mathbf{x}|} \right) \mathbf{a}_i^\mathcal{T} \mathbf{x} \mathbf{a}_{i, \Omega^{t+1}} \right\|_2 \\
&\leq \left\| \mathbf{h}_{\Omega^{t+1}}^t - \frac{\mu}{m} \sum_{i=1}^m w_i^t \mathbf{a}_{i, \Omega^{t+1}} \mathbf{a}_{i, \Omega^{t+1}}^\mathcal{T} \mathbf{h}_{\Omega^{t+1}}^t \right\|_2
\end{aligned}$$

$$\begin{aligned}
& + \left\| \frac{\mu}{m} \sum_{i=1}^m w_i^t \mathbf{a}_{i,\Omega^{t+1}} \mathbf{a}_{i,\Omega^t \setminus \Omega^{t+1}}^\mathcal{T} \mathbf{h}_{\Omega^t \setminus \Omega^{t+1}}^t \right\|_2 \\
& + \left\| \frac{\mu}{m} \sum_{i=1}^m w_i^t \left(\frac{\mathbf{a}_i^\mathcal{T} \mathbf{z}^t}{|\mathbf{a}_i^\mathcal{T} \mathbf{z}^t|} - \frac{\mathbf{a}_i^\mathcal{T} \mathbf{x}}{|\mathbf{a}_i^\mathcal{T} \mathbf{x}|} \right) |\mathbf{a}_i^\mathcal{T} \mathbf{x}| \mathbf{a}_{i,\Omega^{t+1}} \right\|_2. \quad (22)
\end{aligned}$$

Hence, bounding $\|\mathbf{h}^{t+1}\|_2$ suffices to bound the three terms on the right hand side of (22).

Regarding the first term, the following holds

$$\begin{aligned}
& \left\| \mathbf{h}_{\Omega^{t+1}}^t - \frac{\mu}{m} \sum_{i=1}^m w_i^t \mathbf{a}_{i,\Omega^{t+1}} \mathbf{a}_{i,\Omega^{t+1}}^\mathcal{T} \mathbf{h}_{\Omega^{t+1}}^t \right\|_2 \\
& \leq \left\| \mathbf{I} - \frac{\mu}{m} \sum_{i=1}^m w_i^t \mathbf{a}_{i,\Omega^{t+1}} \mathbf{a}_{i,\Omega^{t+1}}^\mathcal{T} \right\|_2 \|\mathbf{h}_{\Omega^{t+1}}^t\|_2 \\
& \leq \max\{1 - \mu\bar{\lambda}, \mu\bar{\lambda} - 1\} \|\mathbf{h}_{\Omega^{t+1}}^t\|_2 \quad (23)
\end{aligned}$$

in which $\bar{\lambda}$ and $\underline{\lambda} > 0$ denote the largest and smallest eigenvalue of $(1/m) \sum_{i=1}^m w_i^t \mathbf{a}_{i,\Omega^{t+1}} \mathbf{a}_{i,\Omega^{t+1}}^\mathcal{T}$, respectively. Since $\tau_w \leq w_i^t \leq 1$ and $\mathbf{a}_{i,\Omega^{t+1}} \mathbf{a}_{i,\Omega^{t+1}}^\mathcal{T}$, $\forall i \in [m]$ are positive semidefinite, the next is true

$$\begin{aligned}
\tau_w \sum_{i=1}^m \mathbf{a}_{i,\Omega^{t+1}} \mathbf{a}_{i,\Omega^{t+1}}^\mathcal{T} & \leq \sum_{i=1}^m w_i^t \mathbf{a}_{i,\Omega^{t+1}} \mathbf{a}_{i,\Omega^{t+1}}^\mathcal{T} \\
& \leq \sum_{i=1}^m \mathbf{a}_{i,\Omega^{t+1}} \mathbf{a}_{i,\Omega^{t+1}}^\mathcal{T}. \quad (24)
\end{aligned}$$

To estimate the eigenvalues of $(1/m) \sum_{i=1}^m \mathbf{a}_{i,\Omega^{t+1}} \mathbf{a}_{i,\Omega^{t+1}}^\mathcal{T}$, we resort to the restricted isometry property of Gaussian matrices $\mathbf{A} \in \mathbb{R}^{m \times n}$ whose entries are i.i.d. standard Gaussian variables [33]. Specifically, if $\mathcal{K} \subsetneq \{1, \dots, n\}$ with $|\mathcal{K}| \leq 3s$, then for constant $\delta_{3s} \leq \epsilon$, the following holds for all $\mathbf{u} \in \mathbb{R}^m$

$$\sqrt{(1 - \delta_{3s})m} \|\mathbf{u}\|_2 \leq \|\mathbf{A}_{\mathcal{K}}^\mathcal{T} \mathbf{u}\|_2 \leq \sqrt{(1 + \delta_{3s})m} \|\mathbf{u}\|_2$$

with probability at least $1 - e^{-c'_1 m}$, provided that $m \geq C'_1 \epsilon^{-2} (3s) \log(n/(3s))$ for numerical constants $c'_1, C'_1 > 0$ [34, Proposition 3.1]. Hence,

$$\lambda_1 \left(\frac{1}{m} \sum_{i=1}^m \mathbf{a}_{i,\Omega^{t+1}} \mathbf{a}_{i,\Omega^{t+1}}^\mathcal{T} \right) \leq 1 + \delta_{3s} \quad (25)$$

$$\lambda_n \left(\frac{1}{m} \sum_{i=1}^m \mathbf{a}_{i,\Omega^{t+1}} \mathbf{a}_{i,\Omega^{t+1}}^\mathcal{T} \right) \geq 1 - \delta_{3s} \quad (26)$$

due to $|\Omega^{t+1}| \leq 2s$. Substituting (25) and (26) into (24) yields

$$\bar{\lambda} \leq 1 + \delta_{3s}, \quad \text{and} \quad \underline{\lambda} \geq \tau_w(1 - \delta_{3s})$$

which together with (23) suggests that

$$\begin{aligned}
& \left\| \mathbf{h}_{\Omega^{t+1}}^t - \frac{\mu}{m} \sum_{i=1}^m w_i^t \mathbf{a}_{i,\Omega^{t+1}} \mathbf{a}_{i,\Omega^{t+1}}^\mathcal{T} \mathbf{h}_{\Omega^{t+1}}^t \right\|_2 \\
& \leq \max\{1 - \mu\tau_w(1 - \delta_{3s}), \mu(1 + \delta_{3s}) - 1\} \|\mathbf{h}_{\Omega^{t+1}}^t\|_2. \quad (27)
\end{aligned}$$

Consider now the second term in (22). For convenience, define

$$\begin{aligned}
\mathbf{A}_{\Omega^{t+1}}^\mathcal{T} & := [\mathbf{a}_{1,\Omega^{t+1}} \cdots \mathbf{a}_{m,\Omega^{t+1}}] \\
\mathbf{A}_{\Omega^t \setminus \Omega^{t+1}}^\mathcal{T} & := [\mathbf{a}_{1,\Omega^t \setminus \Omega^{t+1}} \cdots \mathbf{a}_{m,\Omega^t \setminus \Omega^{t+1}}] \\
\mathbf{A}_{\Omega^t \cup \Omega^{t+1}}^\mathcal{T} & := [\mathbf{a}_{1,\Omega^t \cup \Omega^{t+1}} \cdots \mathbf{a}_{m,\Omega^t \cup \Omega^{t+1}}]
\end{aligned}$$

and let $\mathbf{W}$ be a diagonal matrix with its i -th diagonal entry being w_i^t . Then, one has

$$\begin{aligned}
& \left\| \frac{\mu}{m} \sum_{i=1}^m w_i^t \mathbf{a}_{i,\Omega^{t+1}} \mathbf{a}_{i,\Omega^t \setminus \Omega^{t+1}}^\mathcal{T} \mathbf{h}_{\Omega^t \setminus \Omega^{t+1}}^t \right\|_2 \\
& \leq \left\| \frac{\mu}{m} \mathbf{A}_{\Omega^{t+1}}^\mathcal{T} \mathbf{W} \mathbf{A}_{\Omega^t \setminus \Omega^{t+1}}^\mathcal{T} \right\|_2 \|\mathbf{h}_{\Omega^t \setminus \Omega^{t+1}}^t\|_2 \\
& \leq \left\| \frac{\mu}{m} \mathbf{A}_{\Omega^t \cup \Omega^{t+1}}^\mathcal{T} \mathbf{W} \mathbf{A}_{\Omega^t \cup \Omega^{t+1}}^\mathcal{T} - \mu \frac{\tau_w + 1}{2} \mathbf{I} \right\|_2 \|\mathbf{h}_{\Omega^t \setminus \Omega^{t+1}}^t\|_2 \\
& \leq \mu \frac{1 - \tau_w + 2\delta_{3s}}{2} \|\mathbf{h}_{\Omega^t \setminus \Omega^{t+1}}^t\|_2 \quad (28)
\end{aligned}$$

where the first inequality is due to the definition of the matrix norm, the second inequality comes from the fact that $\mathbf{A}_{\Omega^{t+1}}^\mathcal{T} \mathbf{W} \mathbf{A}_{\Omega^t \setminus \Omega^{t+1}}^\mathcal{T}$ is a submatrix of $\mathbf{A}_{\Omega^t \cup \Omega^{t+1}}^\mathcal{T} \mathbf{W} \mathbf{A}_{\Omega^t \cup \Omega^{t+1}}^\mathcal{T}$, and the last inequality stems from $\tau_w < 1$.

Finally, for the last term in (22), let $\mathbf{v}^t := [v_1^t \cdots v_m^t]^\mathcal{T}$ with $v_i^t := w_i^t \left(\frac{\mathbf{a}_i^\mathcal{T} \mathbf{z}^t}{|\mathbf{a}_i^\mathcal{T} \mathbf{z}^t|} - \frac{\mathbf{a}_i^\mathcal{T} \mathbf{x}}{|\mathbf{a}_i^\mathcal{T} \mathbf{x}|} \right) |\mathbf{a}_i^\mathcal{T} \mathbf{x}|$, $\forall i \in [m]$. By the definition of the induced matrix 2-norm, it holds that

$$\begin{aligned}
& \left\| \frac{1}{m} \sum_{i=1}^m w_i^t \left(\frac{\mathbf{a}_i^\mathcal{T} \mathbf{z}^t}{|\mathbf{a}_i^\mathcal{T} \mathbf{z}^t|} - \frac{\mathbf{a}_i^\mathcal{T} \mathbf{x}}{|\mathbf{a}_i^\mathcal{T} \mathbf{x}|} \right) |\mathbf{a}_i^\mathcal{T} \mathbf{x}| \mathbf{a}_{i,\Omega^{t+1}} \right\|_2 \\
& = \frac{1}{m} \|\mathbf{A}_{\Omega^{t+1}}^\mathcal{T} \mathbf{v}^t\|_2 \leq \left\| \frac{1}{\sqrt{m}} \mathbf{A}_{\Omega^{t+1}}^\mathcal{T} \right\|_2 \left\| \frac{1}{\sqrt{m}} \mathbf{v}^t \right\|_2 \\
& \leq (1 + \delta_{3s}) \left\| \frac{1}{\sqrt{m}} \mathbf{v}^t \right\|_2. \quad (29)
\end{aligned}$$

Regarding the term $\left\| \frac{1}{\sqrt{m}} \mathbf{v}^t \right\|_2$, the following holds

$$\begin{aligned}
\frac{1}{m} \|\mathbf{v}^t\|_2^2 & = \frac{1}{m} \sum_{i=1}^m w_i^t \left(\frac{\mathbf{a}_i^\mathcal{T} \mathbf{z}^t}{|\mathbf{a}_i^\mathcal{T} \mathbf{z}^t|} - \frac{\mathbf{a}_i^\mathcal{T} \mathbf{x}}{|\mathbf{a}_i^\mathcal{T} \mathbf{x}|} \right)^2 |\mathbf{a}_i^\mathcal{T} \mathbf{x}|^2 \\
& \leq 2 \cdot \frac{1}{m} \sum_{i=1}^m |\text{sgn}(\mathbf{a}_i^\mathcal{T} \mathbf{z}) - \text{sgn}(\mathbf{a}_i^\mathcal{T} \mathbf{x})| |\mathbf{a}_i^\mathcal{T} \mathbf{x}| |\mathbf{a}_i^\mathcal{T} \mathbf{h}| \\
& \leq 4 \frac{\sqrt{1 + \delta_{2s}}}{1 - \rho_0} \left(\delta_{2s} + \sqrt{\frac{21}{20}} \rho_0 \right) \|\mathbf{h}\|_2^2 \quad (30)
\end{aligned}$$

where the first inequality originates from that $|\mathbf{a}_i^\mathcal{T} \mathbf{x}| \leq |\mathbf{a}_i^\mathcal{T} \mathbf{h}|$ when $\text{sgn}(\mathbf{a}_i^\mathcal{T} \mathbf{z}) \neq \text{sgn}(\mathbf{a}_i^\mathcal{T} \mathbf{x})$ and $0 < w_i^t < 1$; the second inequality is obtained by appealing to Lemma 4 in the Appendix, which is adapted from [20, Lemma 7.17]. Taking the result in (30) into (29) gives rise to

$$\begin{aligned}
& \left\| \frac{1}{m} \sum_{i=1}^m \left(\frac{\mathbf{a}_i^\mathcal{T} \mathbf{z}^t}{|\mathbf{a}_i^\mathcal{T} \mathbf{z}^t|} - \frac{\mathbf{a}_i^\mathcal{T} \mathbf{x}}{|\mathbf{a}_i^\mathcal{T} \mathbf{x}|} \right) |\mathbf{a}_i^\mathcal{T} \mathbf{x}| \mathbf{a}_{i,\Omega^{t+1}} \right\|_2 \\
& \leq (1 + \delta_{3s}) \gamma \|\mathbf{h}^t\|_2 \quad (31)
\end{aligned}$$

where $\gamma := 2 \sqrt{\frac{\sqrt{1 + \delta_{2s}}}{1 - \rho_0} \left(\delta_{2s} + \sqrt{\frac{21}{20}} \rho_0 \right)}$.

Plugging the bounds in (27), (28), and (31) into (22) confirms that

$$\begin{aligned}
\|\mathbf{h}^{t+1}\|_2 & \leq 2 \max\{1 - \mu\tau_w(1 - \delta_{3s}), \mu(1 + \delta_{3s}) - 1\} \|\mathbf{h}_{\Omega^{t+1}}^t\|_2 \\
& \quad + \mu(1 - \tau_w + 2\delta_{3s}) \|\mathbf{h}_{\Omega^t \setminus \Omega^{t+1}}^t\|_2 + 2\mu(1 + \delta_{3s}) \gamma \|\mathbf{h}^t\|_2 \\
& \leq 2\sqrt{2} \max\{\max\{1 - \mu\tau_w(1 - \delta_{3s}), \mu(1 + \delta_{3s}) - 1\}, \\
& \quad \mu(1 - \tau_w + 2\delta_{3s})/2\} \|\mathbf{h}^t\|_2 + 2\mu(1 + \delta_{3s}) \gamma \|\mathbf{h}^t\|_2 \\
& \leq 2 \left[\sqrt{2} \max\{1 - \mu\tau_w(1 - \delta_{3s}), \mu(1 + \delta_{3s}) - 1\}, \right. \\
& \quad \left. \mu(1 - \tau_w + 2\delta_{3s})/2\} + \mu(1 + \delta_{3s}) \gamma \right] \|\mathbf{h}^t\|_2
\end{aligned}$$

$$:= \rho \|\mathbf{h}^t\|_2 \quad (32)$$

where the second inequality is due to

$$\|\mathbf{h}_{\Omega^{t+1}}^t\|_2 + \|\mathbf{h}_{\Omega^t \setminus \Omega^{t+1}}^t\|_2 \leq \sqrt{2} \|\mathbf{h}^t\|_2$$

for disjoint sets Ω^{t+1} and $\Omega^t \setminus \Omega^{t+1}$. From (32), it is clear that for proper τ_w and sufficiently small ρ_0 , δ_{2s} , and δ_{3s} , one can select a constant step size $\mu > 0$ such that $\rho < 1$. Theorem V-B can be then directly deduced by combining Theorem 1, Lemma 3, and equation (32).

VI. CONCLUDING REMARKS

This contribution developed a compressive reweighted amplitude flow (CRAF) algorithm for phase retrieval of (block)-sparse signal vectors. CRAF first estimates the support of the underlying signal vector, which is followed by a new spectral procedure to obtain an effective initialization. To strengthen this initial guess, CRAF proceeds with (model-based) hard thresholding iterations relying on reweighted gradients of the amplitude-based least-squares loss function. CRAF provably recovers the true signal vectors exponentially fast when a sufficient number of measurements become available. Judicious numerical tests corroborate the merits of CRAF relative to state-of-the-art solvers of the same kind.

APPENDIX: SUPPORTING LEMMAS

The following lemma which bounds the distance between the principle eigenvectors of two symmetric matrices is adapted from [30, Corollary 1].

Lemma 2. [30, Corollary 1] *Let $\mathbf{Z} := \mathbf{X} + \Delta$ with $\mathbf{X}$ and Δ being symmetric matrices, unit vectors $\mathbf{v}_1$ and $\mathbf{u}_1$ be the principal eigenvectors of $\mathbf{Z}$ and $\mathbf{X}$, and $\theta := \cos^{-1}(\langle \mathbf{u}_1, \mathbf{v}_1 \rangle)$ represent the angles between $\mathbf{u}_1$ and $\mathbf{v}_1$. It then holds that*

$$\sqrt{1 - \langle \mathbf{u}_1, \mathbf{v}_1 \rangle^2} = |\sin \theta| \leq \frac{2\|\Delta\|_2}{\lambda_1(\mathbf{X}) - \lambda_2(\mathbf{X})}. \quad (33)$$

The next lemma adopted from [27] certifies that Steps 2-4 in Algorithm 1 recover the true support of $\mathbf{x}$ with high probability.

Lemma 3 (Support estimate [27]). *Consider any k -block-sparse signal vector $\mathbf{x} \in \mathbb{R}^n$ with support $\mathcal{S}$ and $\mathbf{x}_{\min}^B := \min_{b \in \mathcal{S}_B} \|\mathbf{x}_b\|_2^2$ on the order of $(1/k)\|\mathbf{x}\|_2^2$. Assume $\{\mathbf{a}_i\}_{i=1}^m$ are i.i.d standard Gaussian, that is, $\mathbf{a}_i \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_n)$. There exists an event of probability exceeding $1 - 6/m$ such that, Steps 2-4 in Algorithm 1 recover $\mathcal{S}$ if $m \geq C_0' k^2 B \log(mn)$ for some positive constant C_0' .*

The last lemma that is useful in establishing the convergence of CRAF is proved in [20, Lemma 7.17].

Lemma 4. [20, Lemma 7.17] *Consider m noise-free measurements $\{\psi_i = |\mathbf{a}_i^\top \mathbf{x}|\}_{i=1}^m$ in which $\mathbf{x} \in \mathbb{R}^n$ is s -sparse with support $\mathcal{S}$, and $\{\mathbf{a}_i \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_n)\}_{i=1}^m$ are i.i.d. sensing vectors. Let $\mathbf{z} \in \mathbb{R}^n$ be an s -sparse vector satisfying $\|\mathbf{z} - \mathbf{x}\|_2 \leq \rho_0 \|\mathbf{x}\|_2$. If $\mathbf{h} \in \mathbb{R}^n$ is $(2s)$ -sparse and $m > C_3(2s) \log(n/(2s))$ for some numerical constants C_3 , then the following holds for $\delta_{2s} > 0$*

$$\frac{1}{m} \sum_{i=1}^m |\operatorname{sgn}(\mathbf{a}_i^\top \mathbf{z}) - \operatorname{sgn}(\mathbf{a}_i^\top \mathbf{x})| |\mathbf{a}_i^\top \mathbf{x}| |\mathbf{a}_i^\top \mathbf{h}|$$

$$\leq 2 \frac{\sqrt{1 + \delta_{2s}}}{1 - \rho_0} \left(\delta_{2s} + \sqrt{\frac{21}{20}} \rho_0 \right) \|\mathbf{h}\|_2^2 \quad (34)$$

with probability exceeding $1 - 3e^{-c_3 m}$ for a fixed numerical constant $c_3 > 0$.

REFERENCES

- [1] K. Jaganathan, Y. C. Eldar, and B. Hassibi, "Phase retrieval: An overview of recent developments," *Opt. Compressive Sens.*; also in *arXiv:1510.07713*, 2015.
- [2] E. J. Candès, Y. C. Eldar, T. Strohmer, and V. Voroninski, "Phase retrieval via matrix completion," *SIAM Rev.*, vol. 57, no. 2, pp. 225–251, May 2015.
- [3] H. Sahinoglu and S. D. Cabrera, "On phase retrieval of finite-length sequences using the initial time sample," *IEEE Trans. Circuits and Syst.*, vol. 38, no. 8, pp. 954–958, Aug. 1991.
- [4] J. R. Fienup, "Phase retrieval algorithms: A comparison," *Appl. Opt.*, vol. 21, no. 15, pp. 2758–2769, Aug. 1982.
- [5] Y. Shechtman, A. Beck, and Y. C. Eldar, "GESPAR: Efficient phase retrieval of sparse signals," *IEEE Trans. Signal Process.*, vol. 62, no. 4, pp. 928–938, Feb. 2014.
- [6] G. Wang, L. Zhang, G. B. Giannakis, J. Chen, and M. Akçakaya, "Sparse phase retrieval via truncated amplitude flow," *IEEE Trans. Signal Process.*, 2017 (to appear); see also *arXiv:1611.07641*, 2016.
- [7] T. Cai, X. Li, and Z. Ma, "Optimal rates of convergence for noisy sparse phase retrieval via thresholded Wirtinger flow," *Ann. Stat.*, vol. 44, no. 5, pp. 2221–2251, 2016.
- [8] P. Netrapalli, P. Jain, and S. Sanghavi, "Phase retrieval using alternating minimization," *IEEE Trans. Signal Process.*, vol. 63, no. 18, pp. 4814–4826, Sept. 2015.
- [9] E. J. Candès, X. Li, and M. Soltanolkotabi, "Phase retrieval via Wirtinger flow: Theory and algorithms," *IEEE Trans. Inf. Theory*, vol. 61, no. 4, pp. 1985–2007, Apr. 2015.
- [10] Y. Chen and E. J. Candès, "Solving random quadratic systems of equations is nearly as easy as solving linear systems," *Comm. Pure Appl. Math.*, vol. 70, no. 5, pp. 822–883, Dec. 2017.
- [11] G. Wang, G. B. Giannakis, and Y. C. Eldar, "Solving systems of random quadratic equations via truncated amplitude flow," *IEEE Trans. Inf. Theory*, vol. 63, no. 12, pp. 1–22, Dec. 2017.
- [12] G. Wang, G. B. Giannakis, Y. Saad, and J. Chen, "Solving most systems of random quadratic equations," in *Adv. Neural Inf. Process. Syst.*, Long Beach, CA, USA, December 4-9 2017.
- [13] H. Zhang, Y. Zhou, Y. Liang, and Y. Chi, "Reshaped Wirtinger flow and incremental algorithm for solving quadratic system of equations," *arXiv:1605.07719*, 2016.
- [14] C. Ma, K. Wang, Y. Chi, and Y. Chen, "Implicit regularization in nonconvex statistical estimation: Gradient descent converges linearly for phase retrieval, matrix completion and blind deconvolution," *arXiv:1711.10467*, 2017.
- [15] R. W. Gerchberg and W. O. Saxton, "A practical algorithm for the determination of phase from image and diffraction," *Optik*, vol. 35, pp. 237–246, Nov. 1972.
- [16] J. R. Fienup, "Reconstruction of an object from the modulus of its Fourier transform," *Opt. letters*, vol. 3, no. 1, pp. 27–29, July 1978.
- [17] E. J. Candès, T. Strohmer, and V. Voroninski, "PhaseLift: Exact and stable signal recovery from magnitude measurements via convex programming," *Appl. Comput. Harmon. Anal.*, vol. 66, no. 8, pp. 1241–1274, Nov. 2013.
- [18] I. Waldspurger, A. d'Aspremont, and S. Mallat, "Phase recovery, maxcut and complex semidefinite programming," *Math. Program.*, vol. 149, no. 1, pp. 47–81, 2015.
- [19] P. Netrapalli, P. Jain, and S. Sanghavi, "Phase retrieval using alternating minimization," *IEEE Trans. Signal Process.*, vol. 63, no. 18, pp. 4814–4826, Sept. 2015.
- [20] M. Soltanolkotabi, "Structured signal recovery from quadratic measurements: Breaking sample complexity barriers via nonconvex optimization," *arXiv:1702.06175*, 2017.
- [21] G. Wang, G. B. Giannakis, and J. Chen, "Solving large-scale systems of random quadratic equations via stochastic truncated amplitude flow," in *European Signal Process. Conf.*, Kos Island, August 28-September 3, 2017.
- [22] G. Wang and G. B. Giannakis, "Solving random systems of quadratic equations via truncated generalized gradient flow," in *Adv. Neural Inf. Process. Syst.*, Barcelona, Spain, 2016, pp. 568–576.

- [23] J. C. Duchi and F. Ruan, "Solving (most) of a set of quadratic equalities: Composite optimization for robust phase retrieval," *arXiv:1705.02356*, May 2017.
- [24] T. Goldstein and S. Studer, "PhaseMax: Convex phase retrieval via basis pursuit," *arXiv:1610.07531v1*, 2016.
- [25] P. Hand and V. Voroninski, "Compressed sensing from phaseless gaussian measurements via linear programming in the natural parameter space," *arXiv:1611.05985*, 2016.
- [26] H. Ohlsson, A. Y. Yang, R. Dong, and S. S. Sastry, "CPRL—An extension of compressive sensing to the phase retrieval problem," in *Adv. Neural Inf. Process. Syst.*, Stateline, NV, 2012, pp. 1367–1375.
- [27] G. Jagatap and C. Hedge, "Phase retrieval using structured sparsity: A sample efficient algorithmic framework," *arXiv:1705.06412*, 2017.
- [28] L. Zhang, G. Wang, D. Romero, and G. B. Giannakis, "Randomized block Frank-Wolfe for convergent large-scale learning," *IEEE Trans. Signal Process.*, vol. 65, no. 24, pp. 6448–6461, Dec. 2017.
- [29] R. G. Baraniuk, V. Cevher, M. F. Duarte, and C. Hegde, "Model-based compressive sensing," *IEEE Trans. Inf. Theory*, vol. 56, no. 4, pp. 1982–2001, Apr. 2010.
- [30] Y. Yu, T. Wang, and R. J. Samworth, "A useful variant of the Davis–Kahan theorem for statisticians," *Biometrika*, vol. 102, no. 2, pp. 315–323, Jun. 2015.
- [31] M. Mondelli and A. Montanari, "Fundamental limits of weak recovery with applications to phase retrieval," *arXiv:1708.05932*, Sep. 2017.
- [32] R. Vershynin, "Introduction to the non-asymptotic analysis of random matrices," *arXiv:1011.3027*, 2010.
- [33] E. J. Candès and T. Tao, "Decoding by linear programming," *IEEE Trans. Inf. Theory*, vol. 51, no. 12, pp. 4203–4215, Dec. 2005.
- [34] D. Needell and J. A. Tropp, "CoSaMP: Iterative signal recovery from incomplete and inaccurate samples," *Appl. Comput. Harmon. Anal.*, vol. 26, no. 3, pp. 301–321, May 2009.