

A Non-Convex Approach to Blind Calibration for Linear Random Sensing Models

Valerio Cambarelli*, Laurent Jacques*

ICTEAM/ELEN, ISPGROUP, Université catholique de Louvain, Belgium.

Abstract— Performing blind calibration is highly important in modern sensing strategies, particularly when calibration aided by multiple, accurately designed training signals is infeasible or resource-consuming. We here address it as a naturally-formulated non-convex problem for a linear model with sub-Gaussian random sensing vectors in which both the sensor gains and the signal are unknown. A sample complexity bound is derived to assess that solving the problem by projected gradient descent with a suitable initialisation provably converges to the global optimum. These findings are supported by numerical evidence on the phase transition of blind calibration and by an imaging example.

1 Introduction

The problem of acquiring an unknown signal under sensing model errors is crucial for modern sensing strategies such as Compressed Sensing (CS), in which such errors are inevitable in physical implementations, have direct impact on the attainable quality [1], and may be to some extent exploited for security purposes [2]. However, if the signals and model errors are stationary, the use of random sensing operators in CS also suggests that repeating the acquisition (*i.e.*, taking more *snapshots*) under new draws of the sensing model could suffice to diversify the measurements and learn both unknown quantities. We address the case of sensing a single unstructured vector $\mathbf{x} \in \mathbb{R}^n$ by collecting p snapshots of m random projections, *i.e.*,

$$\mathbf{y}_l = \bar{\mathbf{d}} \mathbf{A}_l \mathbf{x}, \quad \bar{\mathbf{d}} := \text{diag}(\mathbf{d}) \in \mathbb{R}^{m \times m}, \quad l \in [p] := \{1, \dots, p\}, \quad (1)$$

where $\mathbf{y}_l = (y_{1,l}, \dots, y_{m,l})^\top \in \mathbb{R}^m$ is the l -th snapshot; $\mathbf{d} = (d_1, \dots, d_m)^\top \in \mathbb{R}_+^m$ is an unknown, positive and bounded gain vector that is identical throughout the p snapshots; the random matrices $\mathbf{A}_l \in \mathbb{R}^{m \times n}$ are independent and identically distributed (i.i.d.) and each $\mathbf{A}_l$ has i.i.d. rows, the i -th row $\mathbf{a}_{i,l}^\top \in \mathbb{R}^n$ being a centred isotropic sub-Gaussian random vector (*i.e.*, $\mathbb{E}[\mathbf{a}_{i,l}] = \mathbf{0}_n, \mathbb{E}[\mathbf{a}_{i,l} \mathbf{a}_{i,l}^\top] = \mathbf{I}_n$).

This essentially bilinear model, referred to as *blind calibration* [3, 4] is motivated by compressive imaging applications, where unknown $\mathbf{d}$ are associated to physical gains and attenuations in a sensor array, while the random projections of an image $\mathbf{x}$ are obtained by light modulation (*e.g.*, by random convolution [5–8]). Similarly, model errors concern radar signal processing applications (*e.g.*, as discussed in [9, 10]). In such contexts, the knowledge of $\mathbf{d}$ could critically improve the accuracy of the sensing model. This is typically achieved by solving convex or alternating minimisation problems [3, 4, 11] that use multiple training signals $\mathbf{x}_l$ instead of randomising the sensing operator. More recently, *lifting* approaches [9, 10] have been proposed to jointly recover $(\mathbf{x}, \mathbf{d})$ in (1) (as well as more general *blind deconvolution* models, [12–14]). Their main limitation is in that a semidefinite program is solved to recover a very large rank-one matrix $\mathbf{x} \mathbf{d}^\top$, an approach that rapidly becomes computationally inefficient as m and n exceed a few hundreds.

*LJ and VC are funded by the Belgian FNRS. Part of this study is funded by the project ALTERSENSE (MIS-FNRS).

Quantity	Finite-sample ($p < \infty$)	Expectation ($\mathbb{E}_{\mathbf{a}_{i,l}}, p \rightarrow \infty$)
Objective: $f(\xi, \gamma)$	$\frac{1}{2mp} \sum_{l=1}^p \ \gamma \mathbf{A}_l \xi - \mathbf{d} \mathbf{A}_l \mathbf{x}\ _2^2$	$\frac{1}{2m} [\ \xi\ _2^2 \ \gamma\ _2^2 + \ \mathbf{x}\ _2^2 \ \mathbf{d}\ _2^2 - 2(\gamma^\top \mathbf{d})(\xi^\top \mathbf{x})]$
Signal gradient: $\nabla_\xi f(\xi, \gamma)$	$\frac{1}{mp} \sum_{l=1}^p \mathbf{A}_l^\top \gamma (\gamma \mathbf{A}_l \xi - \mathbf{d} \mathbf{A}_l \mathbf{x})$	$\frac{1}{m} [\ \gamma\ _2^2 \xi - (\gamma^\top \mathbf{d}) \mathbf{x}]$
Gain gradient: $\nabla_\gamma f(\xi, \gamma)$	$\frac{1}{mp} \sum_{l=1}^p \bar{\mathbf{A}}_l \xi (\gamma \mathbf{A}_l \xi - \mathbf{d} \mathbf{A}_l \mathbf{x})$	$\frac{1}{m} [\ \xi\ _2^2 \gamma - (\xi^\top \mathbf{d}) \mathbf{d}]$
Hessian: $\mathcal{H}f(\xi, \gamma)$	$\frac{1}{mp} \sum_{l=1}^p \begin{bmatrix} \mathbf{A}_l^\top \gamma^2 \mathbf{A}_l & \mathbf{A}_l^\top \gamma \mathbf{A}_l \xi - \mathbf{d} \mathbf{A}_l \mathbf{x} \\ \gamma \mathbf{A}_l \xi - \mathbf{d} \mathbf{A}_l \mathbf{x} & \bar{\mathbf{A}}_l \xi \end{bmatrix}$	$\frac{1}{m} \begin{bmatrix} \ \gamma\ _2^2 \mathbf{I}_n & \gamma \mathbf{d}^\top - \mathbf{x} \mathbf{d}^\top \\ \gamma \mathbf{d}^\top - \mathbf{x} \mathbf{d}^\top & \ \xi\ _2^2 \mathbf{I}_m \end{bmatrix}$
Initialisation: (ξ_0, γ_0)	$\left(\frac{1}{mp} \sum_{l=1}^p (\mathbf{A}_l)^\top \bar{\mathbf{d}} \mathbf{A}_l \mathbf{x}, \mathbf{1}_m \right)$	$\left(\frac{\ \mathbf{d}\ _2}{m} \mathbf{x}, \mathbf{1}_m \right)$

Table 1: Finite-sample and expected values of the objective function; its gradient components and Hessian matrix; the initialisation point (ξ_0, γ_0) . $\bar{\cdot}$ abbreviates the $\text{diag}(\cdot)$ operator.

2 A Non-Convex Approach

Inspired by recent results on fast, provably convergent non-convex approaches to phase retrieval [15–18] we argue that the blind calibration problem of recovering the two unstructured vectors $(\mathbf{x}, \mathbf{d})$ in (1) can be solved exactly by a non-convex problem described hereafter. Since no *a priori* structure is assumed on $(\mathbf{x}, \mathbf{d})$ we operate in the case $mp \geq n + m$ and solve

$$(\hat{\mathbf{x}}, \hat{\mathbf{d}}) = \underset{\substack{\xi \in \mathbb{R}^n, \gamma \in \Pi_+^m}}{\text{argmin}} \frac{1}{2mp} \sum_{l=1}^p \|\gamma \mathbf{A}_l \xi - \mathbf{y}_l\|_2^2, \quad (2)$$

given $\{\mathbf{y}_l\}_{l=1}^p, \{\mathbf{A}_l\}_{l=1}^p$ with $\Pi_+^m = \{\gamma \in \mathbb{R}_+^m, \mathbf{1}_m^\top \gamma = m\}$ being the scaled probability simplex. The geometry of this problem is well understood by observing the objective $f(\xi, \gamma)$ with its gradient and Hessian matrix, computed in Table 1 for finite and asymptotic p . All finite-sample expressions therein are unbiased estimates of their expectation w.r.t. $\mathbf{a}_{i,l}$. There, we evince that all points in $\{(\xi, \gamma) \in \mathbb{R}^n \times \mathbb{R}^m : \xi = \frac{1}{\alpha} \mathbf{x}, \gamma = \alpha \mathbf{d}, \alpha \in \mathbb{R} \setminus \{0\}\}$ are global minimisers of $f(\xi, \gamma)$. Moreover, $f(\xi, \gamma)$ is easily shown to be generally non-convex¹ as there exist plenty of counterexamples (ξ', γ') for which $\mathcal{H}f(\xi', \gamma') \not\geq 0$.

By applying the constraint $\gamma \in \Pi_+^m$, one minimiser $(\mathbf{x}^*, \mathbf{d}^*) = \left(\frac{\|\mathbf{d}\|_1}{m} \mathbf{x}, \frac{m}{\|\mathbf{d}\|_1} \mathbf{d} \right)$ remains for $mp \geq n + m$, which is exact up to $\alpha = m/\|\mathbf{d}\|_1$. Assuming that $\mathbf{d} \in \mathbb{R}_+^m$ in (1) is bounded amounts to letting² $\mathbf{d}^* := \mathbf{1}_m + \omega \in \mathcal{C}_\rho \subset \Pi_+^m$, $\omega \in \mathbf{1}_m^\perp \cap \rho \mathbb{B}_\infty^m$ for a maximum deviation $\rho = \|\mathbf{d}^* - \mathbf{1}_m\|_\infty < 1$ where $\mathcal{C}_\rho := \mathbf{1}_m + \mathbf{1}_m^\perp \cap \rho \mathbb{B}_\infty^m$. The test vector is also cast as $\gamma := \mathbf{1}_m + \varepsilon \in \mathcal{C}_\rho$. While the problem clearly remains non-convex, by defining a *distance* $\Delta(\xi, \gamma) := \|\xi - \mathbf{x}^*\|_2^2 + \frac{\|\mathbf{x}^*\|_2^2}{m} \|\gamma - \mathbf{d}^*\|_2^2$, and a *neighbourhood* of the global minimiser $(\mathbf{x}^*, \mathbf{d}^*)$ as $\mathcal{D}_{\kappa, \rho} := \{(\xi, \gamma) \in \mathbb{R}^n \times \mathcal{C}_\rho : \Delta(\xi, \gamma) \leq \kappa^2 \|\mathbf{x}^*\|_2^2\}$ for $\kappa, \rho \in [0, 1]$, it is easily shown that (2) is locally convex in expectation on $\mathcal{D}_{\kappa, \rho}$ for sufficiently small κ . Thus, we require an *initialisation* point (ξ_0, γ_0) to lie in a small $\mathcal{D}_{\kappa, \rho}$ for which a notion of local convexity holds for finite p . Similarly to [15] we see that initialising ξ_0 as in Table 1 grants $\mathbb{E}[\xi_0] \equiv \mathbf{x}^*$, *i.e.*, asymptotically in p this initialisation yields the exact solution. We also let $\gamma_0 = \mathbf{1}_m$ ($\varepsilon_0 = \mathbf{0}_m$) and run projected gradient descent to solve (2), as summarised in Algorithm 1. Note that, since we assume ρ to be small (*i.e.*, $\mathbf{d}^*$ is far from the vertices of Π_+^m) and the optimisation starts on

¹It is *biconvex* [10], *i.e.*, convex once either ξ or γ are fixed.

²The orthogonal complement of $\mathbf{1}_m$ is $\mathbf{1}_m^\perp = \{\mathbf{v} \in \mathbb{R}^m : \mathbf{1}_m^\top \mathbf{v} = 0\}$.

```

1: Initialise  $\xi_0 := \frac{1}{mp} \sum_{l=1}^p (\mathbf{A}_l)^\top \mathbf{y}_l$ ,  $\gamma_0 := \mathbf{1}_m$ ,  $k := 0$ .
2: while stop criteria not met do
3:    $\mu_\xi := \operatorname{argmin}_{v \in \mathbb{R}} f(\xi_k - v \nabla_\xi f(\xi_k, \gamma_k), \gamma_k)$  {Line search in  $\xi$ }
4:    $\mu_\gamma := \operatorname{argmin}_{v \in \mathbb{R}} f(\xi_k, \gamma_k - v \nabla_\gamma f(\xi_k, \gamma_k))$  {Line search in  $\gamma$ }
5:    $\xi_{k+1} := \xi_k - \mu_\xi \nabla_\xi f(\xi_k, \gamma_k)$ 
6:    $\gamma_{k+1} := \gamma_k - \mu_\gamma \nabla_\gamma f(\xi_k, \gamma_k)$ 
7:    $\gamma_{k+1} := P_{C_\rho} \gamma_{k+1}$  {Projection on  $C_\rho$ }
8:    $k := k + 1$ 
9: end while

```

Algorithm 1: Non-Convex Blind Calibration by Projected Gradient Descent.

$\gamma_0 = \mathbf{1}_m$, we are allowed to first project the gradient step as $\nabla_\gamma^\perp f(\xi, \gamma) := P_{\mathbf{1}_m}^\perp \nabla_\gamma f(\xi, \gamma)$ with $P_{\mathbf{1}_m}^\perp = \mathbf{I}_m - \frac{1}{m} \mathbf{1}_m \mathbf{1}_m^\top$, then project it on C_ρ (i.e., P_{C_ρ} ; this step is useful for the proofs and not required in all the experiments). The line searches in 3:-4: of Algorithm 1 are solved in closed-form and introduced to improve the convergence rate. Below, we obtain the conditions that ensure convergence of this descent algorithm to the exact solution $(\mathbf{x}^*, \mathbf{d}^*)$ for fixed step sizes μ_ξ, μ_γ .

3 Recovery Guarantees

Having assumed that all mp sensing vectors $\mathbf{a}_{i,l}$ are i.i.d. sub-Gaussian, we establish the sample complexity required for the initialisation to lie in $\mathcal{D}_{\kappa,\rho}$ with very high probability. This and all other statements rely on Bernstein-type concentration inequalities [19] and some simple geometry of (2) in $\mathcal{D}_{\kappa,\rho}$.

Proposition 1 (Initialisation Proximity). *Let (ξ_0, γ_0) be as in Algorithm 1. For any $\epsilon \in (0, 1)$, $t > 1$, provided $n \gtrsim t \log(mp)$ and $mp \gtrsim \epsilon^{-2}(n+m) \log(n\epsilon^{-1})$ we have that, with probability exceeding $1 - Ce^{-c\epsilon^2 mp} - (mp)^{-t}$ for some $C, c > 0$, $\|\xi_0 - \mathbf{x}^*\|_2 \leq \epsilon \|\mathbf{x}^*\|_2$. Since $\gamma_0 = \mathbf{1}_m$ we also have $\|\gamma_0 - \mathbf{d}^*\|_\infty \leq \rho < 1$. Thus $(\xi_0, \gamma_0) \in \mathcal{D}_{\kappa,\rho}$ with the same probability and $\kappa := \sqrt{\epsilon^2 + \rho^2}$.*

With analogue tools we are able to state when a single gradient descent step from any (ξ, γ) reduces, with very high probability, the distance w.r.t. the global minimum. This requires establishing a *regularity condition* on $\nabla f(\xi, \gamma)$ holding uniformly for all $(\xi, \gamma) \in \mathcal{D}_{\kappa,\rho}$ (similarly to [15, Condition 7.9]).

Proposition 2 (Regularity condition in $\mathcal{D}_{\kappa,\rho}$). *For any $\delta \in (0, 1)$, $\rho \in [0, 1)$, $t > 1$, provided $\rho < \frac{1-2\delta}{9}$, $n \gtrsim t \log(mp)$, $p \gtrsim \delta^{-2} \log m$ and $\sqrt{mp} \gtrsim \delta^{-2}(n+m) \log(\frac{n}{\delta})$, with probability exceeding $1 - C[me^{-c\delta^2 p} + e^{-c\delta^2 \sqrt{mp}} + (mp)^{-t}]$ for some $C, c > 0$, we have that for all $(\xi, \gamma) \in \mathcal{D}_{\kappa,\rho}$,*

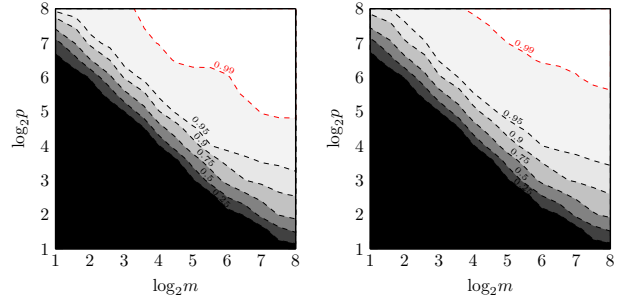
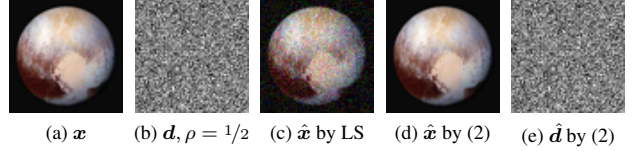
$$\begin{aligned} \left\langle \nabla^\perp f(\xi, \gamma), \begin{bmatrix} \xi - \mathbf{x}^* \\ \gamma - \mathbf{d}^* \end{bmatrix} \right\rangle &\geq \frac{1}{2} \eta \Delta(\xi, \gamma) \quad (\text{Bounded curvature}) \\ \|\nabla^\perp f(\xi, \gamma)\|_2^2 &\leq L^2 \Delta(\xi, \gamma) \quad (\text{Lipschitz gradient}) \end{aligned}$$

for $\eta := 2(1 - 9\rho - 2\delta) > 0$, $L := 4\sqrt{2}[1 + \rho + (1 + \kappa)\|\mathbf{x}^*\|_2]$.

From Proposition 2 it follows that the neighbourhood $\mathcal{D}_{\kappa,\rho}$ is a basin of attraction to the solution $(\mathbf{x}^*, \mathbf{d}^*)$, i.e., we can show that the update from any $(\xi_k, \gamma_k) \in \mathcal{D}_{\kappa,\rho}$ to $(\xi_{k+1}, \gamma_{k+1})$ is so that the distance $\Delta(\xi, \gamma)$ decreases under an upper bound on μ_ξ, μ_γ specified in the following result.

Theorem 1 (Provable Convergence to the Exact Solution). *Under the conditions of Proposition 1, 2 we have that, with probability exceeding $1 - C[me^{-c\delta^2 p} + e^{-c\delta^2 \sqrt{mp}} + e^{-c\epsilon^2 mp} + (mp)^{-t}]$ for some $C, c > 0$, Algorithm 1 with $\mu_\xi := \mu, \mu_\gamma := \mu \frac{m}{\|\mathbf{x}^*\|_2^2}$ has error decay*

$$\Delta(\xi_k, \gamma_k) \leq (1 - \eta\mu + \frac{L^2}{\tau} \mu^2)^k (\epsilon^2 + \rho^2) \|\mathbf{x}^*\|_2^2, (\xi_k, \gamma_k) \in \mathcal{D}_{\kappa,\rho}$$

Figure 1: Empirical phase transition of (2) for $n = 2^8$, $\rho = 10^{-3}$ (left) and $\rho = 10^{-0.5}$ (right). The probability value P_ζ is reported on each level set.Figure 2: A practical, high-dimensional example of blind calibration against unknown, unstructured gains $\mathbf{d}$ with $m = n$, $p = 4$ snapshots.

at any iteration $k > 0$ provided $\mu \in (0, \tau\eta/L^2)$, $\tau := \min\{1, \|\mathbf{x}^*\|_2^2/m\}$. Hence, $\Delta(\xi_k, \gamma_k) \xrightarrow[k \rightarrow \infty]{} 0$.

Let us mention finally that if additive measurement noise $\mathbf{N} := (\mathbf{n}_1, \dots, \mathbf{n}_p) \in \mathbb{R}^{m \times p}$ corrupts the p snapshots of (1), then it can be shown that $\Delta(\xi_k, \gamma_k) \rightarrow \sigma^2$ as $k \rightarrow \infty$, with $\sigma^2 \gtrsim \frac{1}{mp} \|\mathbf{N}\|_F^2$. Thus, Algorithm 1 is robust to noise, its solution degrading gracefully with σ^2 .

4 Numerical Experiments

Empirical Phase Transition: To characterise the phase transition of (2) solved by Algorithm 1 we ran some extensive simulations by generating 256 random instances of (2) for each configuration $\log_2 n \times \log_2 m \times \log_2 p \times \log_{10} \rho \in \{1, \dots, 8\}^3 \times \{-3, \dots, 0\}$, drawing $\mathbf{d}^* = \mathbf{1}_m + \omega$ with $\omega \in \mathbb{S}_m^{m-1}$ drawn uniformly at random on the sphere $\rho \mathbb{S}_m^{m-1}$. Then we evaluated $P_\zeta := \mathbb{P} \left[\max \left\{ \frac{\|\hat{\mathbf{d}} - \mathbf{d}^*\|_2}{\|\mathbf{d}^*\|_2}, \frac{\|\hat{\mathbf{x}} - \mathbf{x}^*\|_2}{\|\mathbf{x}^*\|_2} \right\} < \zeta \right]$ on the trials with $\zeta = 10^{-3}$ chosen according to the stop criterion $f(\xi, \gamma) < 10^{-7}$. Of this large dataset we only report the case $n = 2^8$ in Fig. 1, highlighting the contour levels of P_ζ for $\rho = \{10^{-3}, 10^{-0.5}\}$ as a function of $\log_2 m$ and $\log_2 p$.

Blind Calibration of an Imaging System: To test (2) in a realistic context, we assume that $\mathbf{x}$ is a $n = 64 \times 64$ pixel image acquired by an imaging system following (1), in which its $m = 64 \times 64$ pixel sensor array suffers from large *pixel response non-uniformity* [20]. This is simulated by generating $\mathbf{d}$ as before with $\rho = 1/2$. We capture $p = 4$ snapshots with $\mathbf{a}_{i,l} \sim \mathcal{N}(\mathbf{0}_n, \mathbf{I}_n)$. By running Algorithm 1, the recovered $(\hat{\mathbf{x}}, \hat{\mathbf{d}}) \equiv (\mathbf{x}, \mathbf{d})$ attains $\max \left\{ \frac{\|\hat{\mathbf{d}} - \mathbf{d}^*\|_2}{\|\mathbf{d}^*\|_2}, \frac{\|\hat{\mathbf{x}} - \mathbf{x}^*\|_2}{\|\mathbf{x}^*\|_2} \right\} \approx -147.38$ dB. Instead, by fixing $\gamma = \mathbf{1}_m$ and solving (2) only in ξ , i.e., finding a least squares (LS) $\hat{\mathbf{x}}$, we obtain $\frac{\|\hat{\mathbf{x}} - \mathbf{x}^*\|_2}{\|\mathbf{x}^*\|_2} \approx -5.50$ dB.

5 Conclusion

We presented and solved a non-convex formulation of the blind calibration problem for a random linear model. The obtained sample complexity could be further improved by exploiting, e.g., a sparse model for $\mathbf{x}, \mathbf{d}$; in absence of such priors, the exact solution is achieved at a rate $\sqrt{mp} \gtrsim (n+m) \log(n)$.

References

- [1] M. A. Herman and T. Strohmer, "General deviants: An analysis of perturbations in compressed sensing," *IEEE Journal of Selected Topics in Signal Processing*, vol. 4, no. 2, pp. 342–349, 2010.
- [2] V. Cambareri, M. Mangia, F. Pareschi, R. Rovatti, and G. Setti, "Low-complexity multiclass encryption by compressed sensing," *IEEE Transactions on Signal Processing*, vol. 63, no. 9, pp. 2183–2195, 2015.
- [3] L. Balzano and R. Nowak, "Blind calibration of networks of sensors: Theory and algorithms," in *Networked Sensing Information and Control*. Springer, 2008, pp. 9–37.
- [4] C. Bilen, G. Puy, R. Gribonval, and L. Daudet, "Convex Optimization Approaches for Blind Sensor Calibration Using Sparsity," *IEEE Transactions on Signal Processing*, vol. 62, no. 18, pp. 4847–4856, Sep. 2014.
- [5] J. Romberg, "Compressive sensing by random convolution," *SIAM Journal on Imaging Sciences*, vol. 2, no. 4, pp. 1098–1128, 2009.
- [6] S. Bahmani and J. Romberg, "Compressive deconvolution in random mask imaging," *IEEE Transactions on Computational Imaging*, vol. 1, no. 4, pp. 236–246, 2015.
- [7] T. Bjorklund and E. Magli, "A parallel compressive imaging architecture for one-shot acquisition," in *2013 IEEE Picture Coding Symposium (PCS)*. IEEE, 2013, pp. 65–68.
- [8] K. Degraux, V. Cambareri, B. Geelen, L. Jacques, G. Lafruit, and G. Setti, "Compressive Hyperspectral Imaging by Out-of-Focus Modulations and Fabry-Pérot Spectral Filters," in *International Traveling Workshop on Interactions between Sparse models and Technology (iTWIST)*, 2014.
- [9] B. Friedlander and T. Strohmer, "Bilinear compressed sensing for array self-calibration," in *2014 48th Asilomar Conference on Signals, Systems and Computers*, Nov. 2014, pp. 363–367.
- [10] S. Ling and T. Strohmer, "Self-calibration and biconvex compressive sensing," *Inverse Problems*, vol. 31, no. 11, p. 115002, 2015.
- [11] J. Lipor and L. Balzano, "Robust blind calibration via total least squares," in *2014 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 2014, pp. 4244–4248.
- [12] A. Ahmed, B. Recht, and J. Romberg, "Blind deconvolution using convex programming," *IEEE Transactions on Information Theory*, vol. 60, no. 3, pp. 1711–1732, 2014.
- [13] S. Ling and T. Strohmer, "Blind deconvolution meets blind demixing: Algorithms and performance bounds," *arXiv preprint arXiv:1512.07730*, 2015.
- [14] A. Ahmed, A. Cosse, and L. Demanet, "A convex approach to blind deconvolution with diverse inputs," in *2015 IEEE 6th International Workshop on Computational Advances in Multi-Sensor Adaptive Processing (CAMSAP)*, Dec 2015, pp. 5–8.
- [15] E. Candès, X. Li, and M. Soltanolkotabi, "Phase Retrieval via Wirtinger Flow: Theory and Algorithms," *IEEE Transactions on Information Theory*, vol. 61, no. 4, pp. 1985–2007, Apr. 2015.
- [16] C. D. White, S. Sanghavi, and R. Ward, "The local convexity of solving systems of quadratic equations," *arXiv:1506.07868 [math, stat]*, Jun. 2015, arXiv: 1506.07868.
- [17] Y. Chen and E. Candès, "Solving random quadratic systems of equations is nearly as easy as solving linear systems," in *Advances in Neural Information Processing Systems*, 2015, pp. 739–747.
- [18] J. Sun, Q. Qu, and J. Wright, "A geometric analysis of phase retrieval," in *2016 IEEE International Symposium on Information Theory (ISIT)*, July 2016, pp. 2379–2383.
- [19] R. Vershynin, "Introduction to the non-asymptotic analysis of random matrices," in *Compressed Sensing: Theory and Applications*. Cambridge University Press, 2012, pp. 210–268.
- [20] M. M. Hayat, S. N. Torres, E. Armstrong, S. C. Cain, and B. Yasuda, "Statistical algorithm for nonuniformity correction in focal-plane arrays," *Applied Optics*, vol. 38, no. 5, pp. 772–780, 1999.