
Global Linear and Local Superlinear Convergence of IRLS for Non-Smooth Robust Regression

Liangzu Peng
Mathematical Institute for Data Science
Johns Hopkins University
lpeng25@jhu.edu

Christian Kümmerle
Department of Computer Science
University of North Carolina at Charlotte
kummerle@uncc.edu

René Vidal
Mathematical Institute for Data Science
Johns Hopkins University
rvidal@jhu.edu

Abstract

We advance both the theory and practice of robust ℓ_p -quasinorm regression for $p \in (0, 1]$ by using novel variants of *iteratively reweighted least-squares* (IRLS) to solve the underlying non-smooth problem. In the convex case, $p = 1$, we prove that this IRLS variant converges globally at a linear rate under a mild, deterministic condition on the feature matrix called the *stable range space property*. In the non-convex case, $p \in (0, 1)$, we prove that under a similar condition, IRLS converges locally to the global minimizer at a superlinear rate of order $2 - p$; the rate becomes quadratic as $p \rightarrow 0$. We showcase the proposed methods in three applications: real phase retrieval, regression without correspondences, and robust face restoration. The results show that (1) IRLS can handle a larger number of outliers than other methods, (2) it is faster than competing methods at the same level of accuracy, (3) it restores a sparsely corrupted face image with satisfactory visual quality. <https://github.com/liangzu/IRLS-NeurIPS2022>

1 Introduction

Given a feature matrix $A \in \mathbb{R}^{m \times n}$ with $m \gg n$ and a response vector $y \in \mathbb{R}^m$, the problem

$$\min_{x \in \mathbb{R}^n} \|Ax - y\|_p \quad (1)$$

of ℓ_p -regression has a variety of different applications depending on the choice of $\|v\|_p = (\sum_i |v_i|^p)^{1/p}$. While $p = 2$ corresponds to standard linear regression, the choice of $p > 2$ arises naturally in semi-supervised learning on graphs [1, 2, 3], and a lot of activity has been dedicated recently to the computational complexity analysis for the case $p > 1$ [4, 5, 6].

In this paper, we assume there is a coefficient vector $x^* \in \mathbb{R}^n$ such that the residual $r^* := Ax^* - y$ is k -sparse, in which case a choice of $p \in (0, 1]$ is of interest. Indeed, for the convex and non-smooth case $p = 1$, (1) is known as *least absolute deviation*, which dates back to the time of Boscovich around in the middle of the 18th century [7, 8]. Since then, it has been known intuitively that (1) ($p = 1$) is robust to *large but few measurement errors* (quoting [9]), i.e., that (1) is robust to outliers.

The algorithmic and theoretical understanding of (1) for $p = 1$ has been of long-standing interest to statisticians, which has led to a vast literature spanning from the classic 1964 paper of Huber [10] to recent contributions [11, 12, 13, 14, 15, 16, 17, 18, 19, 20]. That being said, for $p = 1$, there is

Algorithm 1: IRLS for ℓ_p -Regression (IRLS_p)

- 1 Input: $\mathbf{A} = [\mathbf{a}_1, \dots, \mathbf{a}_m]^\top \in \mathbb{R}^{m \times n}$, $\mathbf{y} = [y_1, \dots, y_m]^\top \in \mathbb{R}^m$, $p \in (0, 1]$;
 - 2 Weight initialization $\mathbf{w}^{(0)} \leftarrow [w_1^{(0)}; \dots; w_m^{(0)}] \in \mathbb{R}^m$; // Instead, one can initialize some vector $\mathbf{x}^{(1)}$.
 - 3 For $t \leftarrow 0, 1, \dots$:

$\mathbf{x}^{(t+1)} \leftarrow \operatorname{argmin}_{\mathbf{x} \in \mathbb{R}^n} \sum_{i=1}^m w_i^{(t)} (\mathbf{a}_i^\top \mathbf{x} - y_i)^2 \quad (2)$
Update $\epsilon^{(t+1)}$ suitably based on $\mathbf{x}^{(t+1)}$, $\mathbf{A}$, and $\mathbf{y}$; // See Section 2.2 for details.
 $w_i^{(t+1)} \leftarrow \max \{ |\mathbf{a}_i^\top \mathbf{x}^{(t+1)} - y_i|, \epsilon^{(t+1)} \}^{p-2} \quad \forall i = 1, \dots, m \quad (3)$
-

a close, but somewhat underexplored connection between (1) and the *basis pursuit* problem [21], and more specifically, the theoretical [22, 23, 24] and algorithmic aspects [25, 26, 27] of compressed sensing [28] (see also Section 2.1). For $p \in (0, 1)$, problem (1) is non-smooth and non-convex, and generally less well understood. To the best of our knowledge, the only paper that considered (1) with $p \in (0, 1)$ is [29], where the authors presented a condition (based on *restricted isometry constants*) that guarantees exact recovery of $\mathbf{x}^*$ from (1). Other related works come either from the compressed sensing literature, where an ℓ_p (and noisy) version of basis pursuit has been considered as a sparsity-promoting formulation [30, 31, 32, 26, 33, 34, 35, 36, 37], or from the literature of matrix recovery/completion, where the *Schatten- p* norm has come into play as a non-convex surrogate for rank minimization [38, 39, 40, 41, 42, 43]. A key message from these works is that ℓ_p or Schatten- p minimization with $p \in (0, 1)$ offers better information-theoretic properties (e.g., requires fewer samples for exact recovery) than minimization with $p = 1$.

Here, we study an *iteratively reweighted least-squares* method (IRLS_p) to solve (1) with $p \in (0, 1]$. As listed in Algorithm 1, IRLS_p alternates between solving a weighted least-squares problem (2) and updating the weights (3); see Section 2.2 for more elaboration on IRLS. The simplicity of this idea (with its impressive performance) justifies its popularity in many machine learning [44, 45, 46, 47] and computer vision [48, 49, 50, 51, 52] applications. Based on recent advances on IRLS [26, 27], we make several contributions for understanding (1) and IRLS_p. We state our contributions next.

The Stable Range Space Property (Section 3.1). We put forward the use of the (stable) *range space property* (RSP) for studying the robust regression problem (1). The stable RSP was proposed by [53] in a different context to analyze a compressed sensing algorithm for solving a weighted basis pursuit problem at each iteration. In analogy to the *nullspace property* in compressed sensing [24, 28], we show that the RSP is a necessary and sufficient condition for guaranteeing that the ℓ_p -regression problem in (1) admits $\mathbf{x}^*$ as its unique solution (Proposition 2). Moreover, we show that if $\mathbf{A} \in \mathbb{R}^{m \times n}$ has i.i.d. $\mathcal{N}(0, 1)$ entries and if m is large enough, then the (stable) RSP holds with high probability (Proposition 3). This justifies its use as the core assumption in our analysis.

Global Linear Convergence (Section 3.2). We prove in Theorem 1 that, under a stable RSP assumption and with $\epsilon^{(t)}$ suitably updated, the IRLS_p Algorithm 1 with $p = 1$ converges linearly to the ground-truth $\mathbf{x}^*$ from any initial weight $\mathbf{w}^{(0)}$ (or equivalently from any initial point $\mathbf{x}^{(1)}$). Note that while (accelerated) first-order methods (e.g., (sub-)gradient descent and proximal algorithms) can also solve the convex and non-smooth problem (1) with $p = 1$, they exhibit, in general, (global) sub-linear rates at best [54, 55]. To our knowledge, [17] is the only paper that claims global linear convergence of a different IRLS variant for (1) with $p = 1$; we compare our results with the ones of [17] in Section 3.2. On the other hand, Theorem 1 is inspired by the IRLS method of [27] for basis pursuit; based on [27], we suitably modify their proof strategy, and thus obtain a faster linear rate.

Local Superlinear Convergence (Section 3.3). We prove in Theorem 2 that, under a stable RSP assumption and with $\epsilon^{(t)}$ suitably updated, the IRLS_p Algorithm 1 with $p \in (0, 1)$ converges to $\mathbf{x}^*$ superlinearly, provided that $\mathbf{x}^{(1)}$ falls into a certain neighborhood of $\mathbf{x}^*$. To the best of our knowledge, no similar result exists for ℓ_p -regression (1). While Theorem 2 is inspired by the IRLS algorithm of [26] for ℓ_p -basis pursuit, their IRLS method does not work well for small p (Section 2.2), and unlike in our work, their radius of local convergence diminishes greatly as $p \rightarrow 0$ or $m \rightarrow \infty$.

Applications (Section 4). We illustrate the performance of IRLS_p for *real phase retrieval* [56, 57], *linear regression without correspondences* [58, 59, 60, 61], and *face restoration from sparsely corrupted measurements* [62, 63, 64]. For real phase retrieval (Section 4.1), we show that $\text{IRLS}_{0,1}$ needs only $m = 2n - 1$ measurements to recover $\mathbf{x}^*$ up to sign, with an additional assumption. Theoretically, even brute-force can fail to identify $\pm \mathbf{x}^*$ for fewer than $2n - 1$ measurements [65]. Empirically, many methods, including *Kaczmarz* [66, 67, 57], *PhaseLamp* [68], *truncated Wirtinger flow* [69], and *coordinate descent* [70], fail with $m = 2n - 1$ Gaussian measurements (Figure 2a). For linear regression without correspondences (Section 4.2), we show in Figures 2b-2c that, $\text{IRLS}_{0,1}$ is uniformly faster (20-100x) and more accurate than *PDLP* [71] (merged into *Google or-tools*) and the commercial solver *Gurobi* [72], both of which solve (1) with $p = 1$ as a linear program, and also than *subgradient descent* implemented by Beck & Guttman-Beck [73]. For face restoration (Section 4.3), we present both quantitative and qualitative results on the Extended Yale B dataset [74].

2 Background

2.1 Connection of Robust Regression and Compressed Sensing

The ℓ_p -regression problem (1) has a natural correspondence to the sparse recovery problem

$$\min_{\mathbf{r} \in \mathbb{R}^m} \|\mathbf{r}\|_p \quad \text{s.t.} \quad \mathbf{D}\mathbf{r} = \mathbf{z} \quad (4)$$

where $\mathbf{D}$ is a $(m - n) \times m$ matrix, $\mathbf{z} \in \mathbb{R}^{m-n}$ a vector and the objective $\|\mathbf{r}\|_p$ penalizes coefficient vectors with too many non-zero coordinates. For $p = 1$, (4) is also called *basis pursuit* [75]. Problems (1) and (4) are known to be related in the following sense (as implicitly stated in [21, 29]):

Proposition 1. *Suppose that the range space of $\mathbf{A}$ is equal to the nullspace of $\mathbf{D}$, and that $\mathbf{z} = -\mathbf{D}\mathbf{y}$. If $\mathbf{x}_1$ globally minimizes (1), then $\mathbf{r}_1 := \mathbf{A}\mathbf{x}_1 - \mathbf{y}$ globally minimizes (4). On the other hand, if $\mathbf{r}_2$ globally minimizes (4), then there exists some $\mathbf{x}_2$ with $\mathbf{r}_2 = \mathbf{A}\mathbf{x}_2 - \mathbf{y}$ that globally minimizes (1).*

Proposition 1 sheds light on how we “transfer”, with new insights, from the analysis of [26, 27] for (4) to results for (1). In what follows, we treat [26, 27] in the context of robust regression and highlight the contribution of our work relative to [26, 27] and other existing results, whenever possible.

2.2 Iteratively Reweighted Least-Squares: The Basics and New Insights

We first discuss two different updating rules for the *smoothing parameter* $\epsilon^{(t)}$ of Algorithm 1: the *fixed* rule and the *dynamic* rule. We then consider the weight updating strategy (3). Along the way we use synthetic experiments to illustrate ideas; see Appendix ?? for the experimental setup. In doing so, we intend to provide a review of the state-of-the-art on the variants of IRLS.

Fixed Smoothing Parameter. Instead of updating $\epsilon^{(t)}$ at each iteration, most works on IRLS use a fixed and small positive number, e.g., $\epsilon := \epsilon^{(t)} = 0.001$ for each t [44, 76, 45, 47, 49, 51, 50, 77]. The intuition is to avoid division by the potentially very small residual $|\mathbf{a}_i^\top \mathbf{x}^{(t)} - y_i|$. But this common practice of fixing ϵ comes with at least three issues. First, IRLS with a fixed $\epsilon > 0$ converges only to an “ ϵ -approximate” point, not exactly the global minimizer ([78, 79, 44, 76], Figure 1a). Second, even if setting ϵ small (e.g., $\epsilon = 10^{-15}$) could lead to an accurate enough solution for $p = 1$ (Figure 1a), it can fail for $p < 1$ (Figure 1b). Finally, it makes obtaining a global linear convergence rate guarantee difficult: We are not aware of any theory about a global linear rate of IRLS with fixed ϵ .

Dynamic Smoothing Parameter. Researchers have used different insights to reach the consensus of dynamically updating $\epsilon^{(t)}$ [26, 17, 80, 27, 81]. The first insight is that convergence to global minimizers ensues if $\epsilon^{(t)}$ is suitably decreased to 0 [26]. With $\beta \in (0, 1)$, one such decreasing rule is

$$\epsilon^{(t+1)} \leftarrow \beta \epsilon^{(t)} \text{ if certain conditions are satisfied, or keep } \epsilon^{(t+1)} \leftarrow \epsilon^{(t)} \text{ otherwise [17, 81].} \quad (5)$$

Figure 1a depicts the performance of [17] with an arbitrary choice $\beta = 0.5$ and $\epsilon^{(0)} = 1$. Also shown in Figure 1a are the IRLS methods with update rules of [26] and [27], which we discuss next.

The basis pursuit paper [26] also proposed a dynamic updating rule for $\epsilon^{(t)}$. In our context (1), it sets

$$\epsilon^{(0)} \leftarrow \infty, \quad \epsilon^{(t+1)} \leftarrow \min \{ \epsilon^{(t)}, [\mathbf{r}^{(t+1)}]_{\alpha+1}/m \}, \quad \mathbf{r}^{(t+1)} := \mathbf{A}\mathbf{x}^{(t+1)} - \mathbf{y} \in \mathbb{R}^m, \quad (6)$$

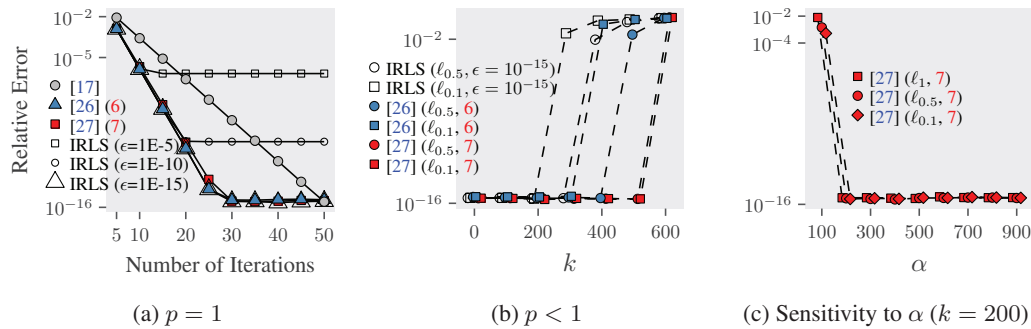


Figure 1: Figure 1a: The relative error $\|\mathbf{x}^{(t)} - \mathbf{x}^*\|_2 / \|\mathbf{x}^*\|_2$ at each iteration t of different IRLS variants for ℓ_1 -regression ($k = 200$). Figure 1b: The relative error of IRLS_p Algorithm 1 and that of [26] for ℓ_p -regression (50 iterations). Figure 1c: The sensitivity of (6) and (7) to mis-specification of α (50 iterations). In Figure 1, we set $m = 1000, n = 10$, and results are averaged over 20 trials.

where the hyper-parameter¹ α is a non-negative integer, and $[\mathbf{r}^{(t+1)}]_{\alpha+1}$ is the $(\alpha + 1)$ -th largest element of the residual $\mathbf{r}^{(t+1)}$ in absolute values. Clearly, (6) creates a non-increasing sequence of smoothing parameters $\epsilon^{(t)}$, and the decay rate of $\epsilon^{(t)}$ is adaptive to the data (in this case to the residual $\mathbf{r}^{(t+1)}$). While at first glance it is not clear how fast $\epsilon^{(t)}$ decays, [26] showed that the decay rate of $\epsilon^{(t)}$ in (6) is locally linear for $p = 1, \alpha = k$ under mild conditions, in accordance with the decay of the objective values.

The final dynamic update rule for the parameter $\epsilon^{(t)}$ is from [27]¹ and improves upon (6) [26] via:

$$\begin{aligned} \epsilon^{(0)} &\leftarrow \infty, \quad \epsilon^{(t+1)} \leftarrow \min \{ \epsilon^{(t)}, \sigma^{(t+1)} / m \} \\ \text{where } \sigma^{(t+1)} &\leftarrow \min \left\{ \|\mathbf{r}^{(t+1)} - \mathbf{z}\|_1 : \mathbf{z} \in \mathbb{R}^m \text{ is } \alpha\text{-sparse} \right\} \end{aligned} \quad (7)$$

Observe that (7) computes the ℓ_1 -norm $\sigma^{(t+1)}$ of the best α -term approximation of the residual $\mathbf{r}^{(t+1)}$, which is in general larger than $[\mathbf{r}^{(t+1)}]_{\alpha+1}$ of (6), while both update rules (6) and (7) rely on the hyper-parameter α . Ideally, if $\alpha = k$, and if IRLS with (6) or (7) converges to $\mathbf{x}^*$, then both $[\mathbf{r}^{(t+1)}]_{\alpha+1}$ and $\sigma^{(t+1)}$ will approach 0 for large enough t . While update rules (6) and (7) perform similarly for ℓ_1 -regression (Figure 1a), (6) fails more easily for small p than (7) (Figure 1b). Finally, we note that overestimating the sparsity level k by setting $\alpha > k$ deteriorates the performance of (7) only slightly (Figure 1c).

Summary. In Section 2.2 we delivered two take-away messages: (1) IRLS with a fixed $\epsilon^{(t)}$ finds some approximate solution, sometimes good enough, (2) dynamically updating $\epsilon^{(t)}$ as per (5)-(7) leads to local [26] or global [17, 27] linear convergence guarantees. In view of Figure 1, we will next consider the IRLS_p Algorithm 1 with update rule (7) for $\epsilon^{(t)}$, and analyze its convergence rates.

3 Convergence Theory of IRLS for Robust Regression

In Section 3.1 we introduce the stable range space property and justify it as our core assumption. Under this assumption, we prove the global linear convergence (Theorem 1) and local superlinear convergence (Theorem 2) of the IRLS_p Algorithm 1 in Sections 3.2 and 3.3, respectively.

3.1 The Stable Range Space Property

Definition 1 (Stable Range Space Property). A matrix $\mathbf{A} \in \mathbb{R}^{m \times n}$ is said to satisfy the *range space property (RSP)* of order k , or the k -RSP for short, if the following holds for any vector $\mathbf{d}$ in the range

¹In [26] and [27], α is the sparsity level k , but k might be unknown in practice, so we treat it as a hyper-parameter. That being said, in the experiments we set $\alpha = k$ by default, for simplicity and in light of Figure 1c.

space of $\mathbf{A}$ and any set $S \subset \{1, \dots, m\}$ of cardinality at most k :

$$\sum_{i \in S} |d_i| < \sum_{i \in S^c} |d_i| \quad (8)$$

The *stable RSP* of $\mathbf{A}$ is defined in the same way as the RSP except that we now require

$$\sum_{i \in S} |d_i| \leq \eta \sum_{i \in S^c} |d_i| \quad (9)$$

for some $\eta \in (0, 1)$. We write “ (k, η) -stable RSP” to emphasize the parameters k and η .

By definition, the stable RSP implies the RSP. Note that the (k, η) -stable RSP was proposed in [53] to analyze a reweighted ℓ_1 -minimization algorithm for compressed sensing. With the notation of Proposition 1, we can see that checking (8) for all $\mathbf{d}$ in the range space of $\mathbf{A}$ is equivalent to checking it for all $\mathbf{d}$ in the nullspace of $\mathbf{D}$, the latter being the well-known *nullspace property* (NSP) [24, 28]. In other words, the (stable) RSP of $\mathbf{A}$ is equivalent to the (stable) NSP of $\mathbf{D}$.

In Sections 3.2 and 3.3, we will use the (k, η) -stable RSP as an assumption for analysis. Arguably, this is a weak assumption to make as the (stable) RSP is very close to a sufficient and necessary condition for exact recovery of the coefficient vector via ℓ_p -robust regression (1):

Proposition 2 (Exact Recovery $\Leftrightarrow$ RSP). *Let $p \in (0, 1]$. For all $\mathbf{x}^*$ and $\mathbf{y}$ such that $\mathbf{A}\mathbf{x}^* - \mathbf{y}$ is k -sparse $\mathbf{x}^*$ is the unique solution to (1) if and only if $\mathbf{A} \in \mathbb{R}^{m \times n}$ satisfies the k -RSP.*

It follows from [82, Theorem 5] and Proposition 1 that checking whether a given matrix $\mathbf{A}$ satisfies the stable RSP is *co-NP-complete* [83]. On the other hand, we show that random Gaussian matrices of size $m \times n$ with sufficiently large m satisfy the stable RSP with high probability, which further justifies the usage of the stable RSP as an assumption in our analysis of IRLS_p :

Proposition 3 (Gaussian $\Rightarrow$ RSP). *Suppose $m - n \geq 2k$. Let $\mathbf{A} \in \mathbb{R}^{m \times n}$ be a matrix with i.i.d. $\mathcal{N}(0, 1)$ entries. Let $\delta \in (0, 1)$ and $\eta \in (0, 1]$ be fixed constants. If it holds that*

$$\frac{(m - n)^2}{m - n + 1} \geq 2k \ln(em/k) \cdot \left(1.67 + \eta^{-1} + \frac{\sqrt{18 \ln(2.5\delta^{-1})}}{\sqrt{2k \ln(em/k)}} \right)^2, \quad (10)$$

then $\mathbf{A}$ satisfies the (k, η) -stable RSP with probability at least $1 - \delta$.

The assumption $m - n \geq 2k$ of Proposition 3 is necessary, as it is not hard to prove that, if $m - n < 2k$ and $\mathbf{A}\mathbf{x}^* - \mathbf{y}$ is k -sparse, then the ℓ_0 -minimization problem $\min_{\mathbf{x} \in \mathbb{R}^n} \|\mathbf{A}\mathbf{x} - \mathbf{y}\|_0$ has multiple solutions. With this assumption, (10) roughly becomes $m - n \geq ck \log(em/k)$ for some constant c if $m - n$ is large enough; thus, ignoring logarithmic and constant factors, condition (10) becomes $m - n \geq \Theta(k)$, which is nearly optimal, as the condition $m - n \geq 2k$ is necessary.

Proposition 2 follows directly from [28, Theorem 4.9], so we omit its proof. The reader might also find Proposition 3 corresponds to [28, Corollary 9.34]. This corollary assumes $\mathbf{D}$ of (4) has i.i.d. $\mathcal{N}(0, 1)$ entries, from which is not immediate what the distribution of $\mathbf{A}$ is, thus it does not imply Proposition 3 directly; this is why we provide a complete proof for Proposition 3 in Appendix ??.

3.2 Global Linear Convergence of IRLS_1

With the background provided in Section 2, we are now ready to state the first main result:

Theorem 1 (Global Linear Convergence). *Suppose $\mathbf{A} \in \mathbb{R}^{m \times n}$ obeys the (k, η) -stable RSP with $\eta \in (0, 3/4)$ (cf. Definition 1). Let $\mathbf{A}\mathbf{x}^* - \mathbf{y}$ be k -sparse. If the smoothing parameter $\epsilon^{(t)}$ is updated as per (7) with $\alpha = k$ during the execution of the IRLS_1 Algorithm 1, then it holds for any $t \geq 1$ that*

$$\|\mathbf{A}\mathbf{x}^{(t+1)} - \mathbf{y}\|_1 - \|\mathbf{A}\mathbf{x}^* - \mathbf{y}\|_1 \leq \left(1 - \frac{(3 - 4\eta)^2}{294\eta m} \right)^t \cdot 3 \cdot \|\mathbf{A}\mathbf{x}^{(1)} - \mathbf{y}\|_1, \quad (11)$$

meaning that the IRLS_1 Algorithm 1 converges linearly and globally in objective value.

Discussion. The condition $\eta \in (0, 3/4)$ is better understood via Proposition 3, which asserts that this condition holds with high probability if $\mathbf{A}$ has i.i.d. Gaussian entries and if m is large enough

(compared to n, k). Moreover, if the (k, η) -stable RSP holds, Proposition 2 implies that $\mathbf{x}^*$ is a global minimizer of (1) with $p = 1$. The strength of Theorem 1 is in that it guarantees linear convergence for *any initialization, starting at the first iteration*. However, the price to pay is that the convergence rate seems conservative with constant $1/294$, and depends on the number m of samples: (11) predicts that IRLS_1 converges with accuracy δ in $O(m \log(1/\delta))$ iterations and m can be large. We argue that such dependency is not an artifact of our analysis, because similar dependencies have been established for IRLS variants for ℓ_p -regression (1) with $p > 2$ [2, Theorem 3.1], for interior point methods in convex optimization with m inequality constraints ([84, Chapter 11], [55, Chapter 5]), and for SGD in smooth & strongly convex optimization with a sum of m function components [85, Theorem 2.1]). Moreover, an adversarial initialization for which the rate empirically depends on m was pointed out for IRLS for basis pursuit in [27]. That being said, we conjecture that the linear dependency on m in (11) can be improved to $\sqrt{m} \log m$ for some different choice of $\epsilon^{(t)}$; see [2] and [84, Sections 11.5 and 11.8.2] for why this conjecture makes sense. Furthermore, Figure 1a empirically depicts that, with the least-squares initialization, i.e., with $w_i^{(0)} = 1$ for all i , IRLS_1 typically converges to $\mathbf{x}^*$ in 30 iterations. This hints to the possibility of alleviating the dependency on m by analyzing the least-squares initialization, a challenging task which we leave to future work.

Comparison to [17]. To our knowledge, [17] is the only paper that claimed the global linear convergence of an IRLS variant with $\epsilon^{(t)}$ updated as per (5) and with $p = 1$. In particular, [17] sets $\epsilon^{(t+1)} \leftarrow \beta \epsilon^{(t)}$ whenever $\|\mathbf{x}^{(t+1)} - \mathbf{x}^{(t)}\|_2 \leq 2\beta \epsilon^{(t)}$; see Figure 1a for the convergence of IRLS under this rule. The update rule $\epsilon^{(t+1)} \leftarrow \beta \epsilon^{(t)}$ of [17] has the advantage of being faster to compute than (7). In spite of this, and of their beautiful proof idea, their result has several disadvantages compared to Theorem 1. First, their proof involves sophisticated probabilistic arguments, and is tailored towards a feature matrix that satisfies strong concentration properties. Our analysis, on the other hand, uses the (k, η) -stable RSP, which is close to a necessary and sufficient condition for the success of ℓ_1 -regression (cf. Proposition 2) and expected to hold for a much larger range of matrices than just (sub-)Gaussian matrices (see [86, 87] for related results in the context of basis pursuit). Second, some statements in their proof are inaccurate (e.g., their Lemma 5 does not hold for $k < n$, the last paragraph in their proving Lemma 8 is not rigorous). Third, their update rule incurs two hyper-parameters, $\epsilon^{(0)}$ and β , while Theorem 1 is based on the decreasing rule (7) that is adaptive to data, and the only hyper-parameter α is easier to set (Figure 1c).

Connection to [27]. A related global linear convergence was established for the IRLS method for basis pursuit in [27, Theorem 3.2]. Our proof is inspired by [27], but also improves on its strategy. For example, their Theorem 3.2 involves two parameters for the *stable nullspace property*, while we only have a single RSP parameter η , simplifying matters. Also, we obtain a constant of $1/294$ in Theorem 1, which is better than the value of $1/768$ for the respective constant in [27, Theorem 3.2].

It is important to note that the proofs of [17, 27] and ours heavily rely on certain rules for decreasing the smoothing parameter $\epsilon^{(t)}$; all these proofs of global linear convergence for IRLS would break down if $\epsilon^{(t)}$ were fixed. While the main theorems of both [17] and [27] are limited to the $p = 1$ case, we present our Theorem 2 for $p \in (0, 1)$ next.

3.3 Local Superlinear Convergence of IRLS_p

The ℓ_p -regression problem (1) with $p \in (0, 1)$ is more challenging than the case $p = 1$ due to the lack of convexity. However, here we show that at least locally, IRLS_p converges with a superlinear rate of order $2 - p$. Moreover, since it is valid to run IRLS_p even with $p = 0$, we carefully design the proof such that the following result holds not only for $p \in (0, 1]$, but also for $p = 0$, in which case IRLS_p can be interpreted as an algorithm minimizing a sum-of-logarithm objective (see Appendix ??).

Theorem 2 (Local Superlinear Convergence). *Run IRLS_p with $p \in [0, 1]$ and update $\epsilon^{(t)}$ by (7) with $\alpha = k$. Assume $\mathbf{A}$ satisfies the (k, η) -stable RSP (Definition 1). Let $c \in (0, 1)$ be a sufficiently small constant such that $2c^{1-p}\eta(\eta + 1) < (1 - c)^{2-p}$. Let $\mathbf{A}\mathbf{x}^* - \mathbf{y}$ be k -sparse with support S^* . Define*

$$\mu := 2\eta(\eta + 1)(1 - c)^{p-2} \cdot \min_{i \in S^*} |\mathbf{a}_i^\top \mathbf{x}^* - y_i|^{p-1}. \quad (12)$$

If the initialization $\mathbf{x}^{(1)}$ is in a neighborhood of $\mathbf{x}^$, in the sense that*

$$\|\mathbf{A}\mathbf{x}^{(1)} - \mathbf{A}\mathbf{x}^*\|_1 \leq c \cdot \min_{i \in S^*} |\mathbf{a}_i^\top \mathbf{x}^* - y_i|, \quad (13)$$

then IRLS_p achieves the following superlinear convergence rate of order $2 - p$ for every $t \geq 1$:

$$\|\mathbf{Ax}^{(t+1)} - \mathbf{Ax}^*\|_1 \leq \mu \cdot \left(\|\mathbf{Ax}^{(t)} - \mathbf{Ax}^*\|_1\right)^{2-p} < \|\mathbf{Ax}^{(t)} - \mathbf{Ax}^*\|_1 \quad (14)$$

In particular, for $p \in [0, 1)$, the superlinear rate (14) implies

$$\|\mathbf{Ax}^{(t+1)} - \mathbf{Ax}^*\|_1 \leq \left(\mu^{1/(1-p)} \cdot \|\mathbf{Ax}^{(1)} - \mathbf{Ax}^*\|_1\right)^{(2-p)^t - 1} \cdot \|\mathbf{Ax}^{(1)} - \mathbf{Ax}^*\|_1. \quad (15)$$

Example 1 (Quadratic versus linear rates). We illustrate the practical benefit of a superlinear (quadratic) convergence rate with the following simple calculation. Suppose that the inequality in (13) is barely fulfilled such that in the case of $p = 0$, $\mathbf{x}^{(1)}$ satisfies $\mu \cdot \|\mathbf{Ax}^{(1)} - \mathbf{Ax}^*\|_1 \leq 0.9999999$. Then (15) implies that after 30 iterations, the residual error is far below numerical precision already since $\|\mathbf{Ax}^{(31)} - \mathbf{Ax}^*\|_1 \leq 0.9999999^{2^{30}-1} \cdot \|\mathbf{Ax}^{(1)} - \mathbf{Ax}^*\|_1 \approx 10^{-47} \cdot \|\mathbf{Ax}^{(1)} - \mathbf{Ax}^*\|_1$. On the other hand, for $p = 1$, a linear convergence factor μ of $\mu \leq 0.9999999$ is only able to guarantee a decay of the order $\mu^{30} \leq 0.9999999^{30} \approx 0.999997$ after 30 iterations.

Discussion. Even though IRLS_p with $p \in (0, 1)$ is tailored to ℓ_p -regression, we measure the progress of IRLS_p in ℓ_1 -norm, e.g., we provide upper bounds on $\|\mathbf{Ax}^{(t+1)} - \mathbf{Ax}^*\|_1$; this is because doing so avoids the use of Hölder's inequality, which allows us to give tighter results that improve on [26] (see the next paragraph). The second point that deserves discussion is the local convergence neighborhood defined by the right-hand side of (13). Note that, we typically assume that every outlier sample $(\mathbf{a}_i, y_i)$ (with $i \in S^*$) would result in a relatively large residual at $\mathbf{x}^*$, say $|\mathbf{a}_i^\top \mathbf{x}^* - y_i| \gg 0$. Hence the minimization term of the right-hand side in (13) is well-behaved. Next, one might ask how to obtain such an initialization $\mathbf{x}^{(1)}$ that satisfies (13). One might run IRLS_1 to produce such an initialization; note though that the upper bound of (13) can not be computed, so one could not detect when to switch from IRLS_1 to IRLS_p ($p \in [0, 1)$). Or alternatively, one could run IRLS_p with a least-squares initialization and count on empirical global (super)linear convergence of IRLS_p ; the latter is what we did in the experiments and is what we recommend. A final and important remark is that ℓ_p -regression (1) with $p \in (0, 1)$ is in general NP-hard [28, Exercise 2.10], but this does not contradict the theoretical local superlinear convergence of Theorem 2 and the empirical global convergence in Figures 1b and 2b, and this does not mean that IRLS_p solves an NP-hard problem in polynomial time. The catch is that we operate under the assumption that $\mathbf{x}^*$ leads to a k -sparse residual and is a unique global minimizer of (1) (cf. Proposition 2). With this assumption and small enough k/m , ℓ_p -regression is tractable and can be solved via IRLS_p to high accuracy (Figure 1b).

Connection to [26]. [26, Theorem 7.9] proves the local superlinear convergence of IRLS with update rule (6) for ℓ_p -basis pursuit (4), which motivates Theorem 2. However, [26, Theorem 7.9] requires c of (13) to satisfy $c = O(1/m^{2/p-1})$, which means that, as $p \rightarrow 0$ or $m \rightarrow \infty$, the eligible value of c quickly becomes vanishingly small, and thus the radius of local convergence (13) diminishes greatly. In contrast, in our condition $2c^{1-p}\eta(\eta+1) < (1-c)^{2-p}$, c has no direct dependency on m and is well-behaved even if $p \rightarrow 0$. The above difference is due partly to our improved proof strategy, and partly to the different update rules of the smoothing parameter $\epsilon^{(t)}$. Figure 1b showed that rule (6) of [26] does not work well for small p , in contrast to rule (7) that we use; a possible explanation for this phenomenon is that rule (6) decreases $\epsilon^{(t)}$ too fast, yielding a poor initialization for the next iteration of weighted least-squares (similarly to the interior point method, cf. [84]). Finally, [26, Theorem 7.9] does not hold for $p = 0$, for which Theorem 2 still holds and suggests a local quadratic rate.

4 Applications and Experiments

We now explore the performance of IRLS_p in three different applications, *real phase retrieval* (Section 4.1), *linear regression without correspondences* (Section 4.2), and *face restoration* (Section 4.3). Note that the first two applications are special examples of the recent algebraic-geometric framework called *homomorphic sensing* [88], [89, Section 4], [90, Section 1.2]. In Section 4.4, we examine the behavior of IRLS_p for different values of p and under noise.

4.1 Real Phase Retrieval

In real phase retrieval [56, 57], we are given a measurement matrix $\mathbf{A} = [\mathbf{a}_1, \dots, \mathbf{a}_m]^\top$ and $\mathbf{y} = [y_1, \dots, y_m]^\top$, where $y_i := |\mathbf{a}_i^\top \mathbf{x}^*|$, and we need to find either $\mathbf{x}^*$ or $-\mathbf{x}^*$. This problem

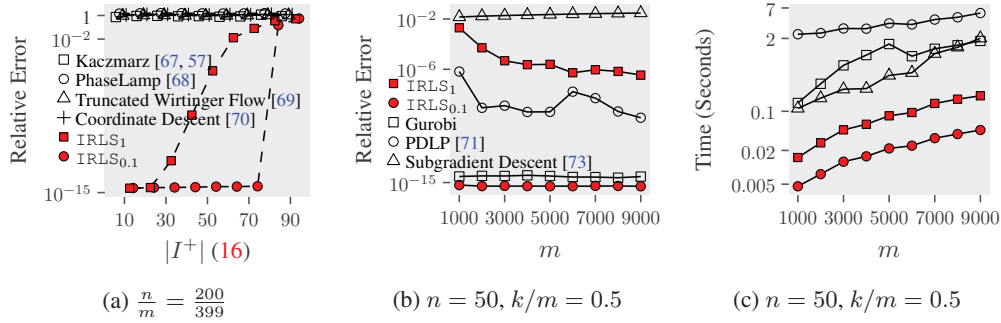


Figure 2: Figure 2a: Relative error $\min\{\|\hat{\mathbf{x}} - \mathbf{x}^*\|_2, \|\hat{\mathbf{x}} + \mathbf{x}^*\|_2\}/\|\mathbf{x}^*\|_2$ of the methods that produce estimates $\hat{\mathbf{x}}$ for real phase retrieval. Figures 2b-2c: Relative errors $\|\hat{\mathbf{x}} - \mathbf{x}^*\|_2/\|\mathbf{x}^*\|_2$ and running times for linear regression without correspondences. IRLS_p run for at most 50 iterations. 50 trials.

is a relative of the complex phase retrieval problem [65], where all data $\mathbf{y}$ and $\mathbf{A}$, as well as the ground-truth $\mathbf{x}^*$, are complex-valued, and which has applications in X-ray crystallography [91].

Here we show that real phase retrieval can be solved via ℓ_p -regression (1). Consider the index sets

$$I^+ := \{i : \mathbf{a}_i^\top \mathbf{x}^* > 0\}, \quad I^- := \{i : \mathbf{a}_i^\top \mathbf{x}^* < 0\}. \quad (16)$$

We might assume $\mathbf{a}_i^\top \mathbf{x}^* \neq 0$ for every i without loss of generality. Then we know that either $\mathbf{A}\mathbf{x}^* - \mathbf{y}$ has its ℓ_0 norm equal to $m - |I^-|$, or $\mathbf{A}(-\mathbf{x}^*) - \mathbf{y}$ has its ℓ_0 norm equal to $m - |I^+|$. With such sparsity patterns on these two residuals, we can minimize (1), which serves as a (non-convex) relaxation of ℓ_0 -minimization, to recover $\mathbf{x}^*$ or $-\mathbf{x}^*$, whichever corresponds to a sparser residual. Here, we essentially treat one of the two clusters defined by (16) as inliers, and the other as outliers. This robust regression point of view on real phase retrieval seems to be known by experts [92, Section 1.2], but we have not found a paper that actually proposes to solve (1) for real phase retrieval.

Here we argue with Figure 2a that, solving (1) (by IRLS_{0.1}) for real phase retrieval can be beneficial. Indeed, theoretically, IRLS_{0.1} converges to $\pm \mathbf{x}^*$ locally but superlinearly (as a corollary of Theorem 2). Empirically, Figure 2a shows that IRLS_{0.1} succeeds in recovering $\pm \mathbf{x}^*$ by using only $m = 2n - 1$ samples, which is exactly the theoretical minimum for real phase retrieval [65, Proposition 2.5]. Figure 2a also shows that multiple other state-of-the-art methods² [57, 67, 68, 69, 70] fail to recover $\pm \mathbf{x}^*$ in this extreme situation, even though some of them have *nearly optimal* sample complexity, which requires roughly $O(n)$ samples up to a logarithmic factor to succeed (*caution*: nearly optimal $\neq$ optimal). As suggested by Proposition 3, for $|I^+|$ fixed, minimizing (1) with a Gaussian matrix $\mathbf{A}$ enjoys $m = O(n)$ sample complexity, up to also a logarithmic factor.

Despite these advantages, we also emphasize that solving (1) (via IRLS_{0.1}) for real phase retrieval has an inherent limit: The performance depends on the number $|I^+|$ of positive signs and the number m of samples; we might call $|I^+|/m$ the inlier rate or outlier rate. At $n/m = 200/399$, we see that IRLS_{0.1} succeeds if $|I^+| \leq 70$, and if $|I^+| \geq 399 - 70$ by symmetry, but it fails if $|I^+|$ and $|I^-|$ get closer (Figure 2a). For IRLS_{0.1} to successfully handle the case $|I^+| = |I^-|$, we need more samples, empirically, say $m \geq 5n$. It is an interesting future direction to design an algorithm that can handle the case of minimum samples ($m = 2n - 1$) and balanced data ($|I^+| = |I^-|$).

4.2 Linear Regression without Correspondences

Compared to (real) phase retrieval, the problem of *linear regression without correspondences* is an also important, but less developed subject; see [58, 94, 59, 60, 61, 95, 96, 97] for recent advances. In this problem we are given $\mathbf{y} \in \mathbb{R}^m$ and $\mathbf{A} \in \mathbb{R}^{m \times n}$, and we need to solve the equations

$$\mathbf{y} = \mathbf{\Pi} \mathbf{A} \mathbf{x} \quad (17)$$

for some unknown permutation matrix $\mathbf{\Pi} \in \mathbb{R}^{m \times m}$ and vector $\mathbf{x} \in \mathbb{R}^n$, under the assumption that there exists a solution $(\mathbf{\Pi}^*, \mathbf{x}^*)$. Solving (17) is in general NP-hard for $n > 1$ [94, 98]. However, if

²Reviewing these algorithms is beyond the scope of this paper. Some of them are designed for complex phase retrieval but is applicable to the real case. We used the implementations from PhasePack [93].

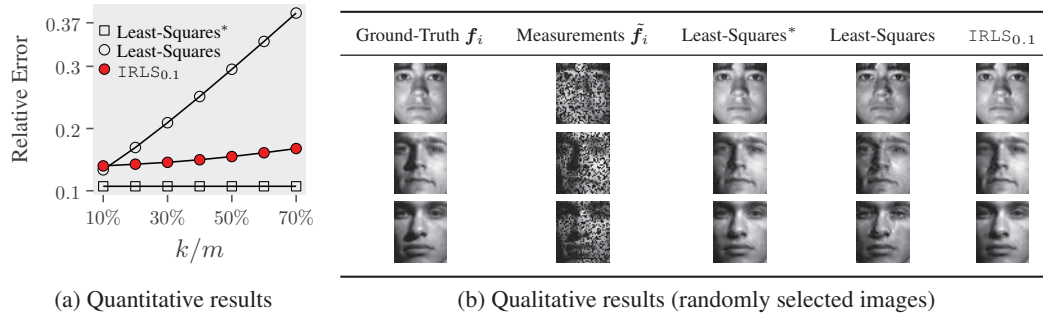


Figure 3: Recover a face image of the Extended Yale B dataset [74] that is corrupted by *salt & pepper* noise [63]. Figure 3a: relative error as a function of the amount of random salt & pepper noise or sparsity level, averaged over 20 trials, all faces, and all individuals. Figure 3b: qualitative results.

we further assume that Π^* permutes at most k rows of A , then we know that the residual $Ax^* - y$ is k -sparse. This is an insight of [59], where it was proposed to estimate x^* by solving (1) with $p = 1$.

We show that IRLS _{p} with $p \in (0, 1)$ is more suitable than several baselines for solving this robust regression problem reduced from linear regression without correspondences. As baselines, we use the PDLF [71] and Gurobi [72] solvers to solve (1) with $p = 1$ as a linear program, and we also use a subgradient descent method implemented in the FOM toolbox of Beck & Guttman-Beck [73]. Figure 2b shows that IRLS_{0.1} is the most accurate, and Figure 2c shows that IRLS_{0.1} is 30 times faster than Gurobi and subgradient descent, and is more than 100 times faster than PDLF. Interestingly, IRLS₁ is not very competitive in terms accuracy (Figure 2b). Moreover, IRLS₁ is slower than IRLS_{0.1}, as IRLS_{0.1} converges faster and thus terminates earlier than IRLS₁; we terminate IRLS _{p} whenever the objective stops decreasing (up to a tolerance 10^{-15}). Finally, see also [45, 48, 49, 50, 51, 52] where their IRLS variants outperform a different set of baselines (for different problems).

4.3 Face Restoration from Sparsely Corrupted Measurements

Here we consider a simple face restoration experiment on the Extended Yale B dataset [74], downsampled as per [99]. This dataset contains the face images of 38 individuals under about $n \approx 60$ different illuminations. Let $F = [f_1, \dots, f_{n+1}]$ be the matrix of faces of the same individual, where each $f_i \in \mathbb{R}^m$ is a (vectorized) face image with $m = 2,016$ pixels. We assume that F is approximately low-rank (which is true under the *Lambertian reflectance* assumption [100]), thus each f_i can be approximately represented as a linear combination of other faces of the individual, i.e., $f_i \approx F_i x_i^*$ for some $x_i^* \in \mathbb{R}^n$, where $F_i \in \mathbb{R}^{m \times n}$ is the same as F , except with the i -th column excluded.

In our experiments of Figure 3, we corrupt f_i by random salt & pepper noise [63] and obtain the noisy face $\tilde{f}_i$ as our measurement; so $F_i x_i^* - \tilde{f}_i$ is approximately sparse. We consider three methods to estimate x_i^* and thus to find an estimate $\hat{f}_i$ of f_i : (1) Least-Squares that solves (1) with $A = F_i$, $y = \tilde{f}_i$, and $p = 2$, (2) Least-Squares* that solves (1) with $A = F_i$, $y = \tilde{f}_i$, and $p = 2$, used as a golden baseline, and (3) IRLS_{0.1} which solves (1) with $A = F_i$ and $y = \tilde{f}_i$. Figure 3a shows that IRLS_{0.1} is more robust to salt & pepper noise, whose amount is k/m , than Least-Squares, and has lower error $\|\hat{f}_i - f_i\|_2 / \|f_i\|_2$ (recall though that Least-Squares is statistically optimal for Gaussian noise). Figure 3b shows some images randomly chosen from the dataset; observe that Least-Squares tends to “average” the illumination in all images F_i , while IRLS_{0.1} delivers faithful restoration.

4.4 The Choice of p and Performance of IRLS _{p} Under Noise

Here we outline two experiments that illustrate two different aspects of IRLS _{p} empirically.

In the first experiment, visualized in Figure 4a, we compare the decay of the relative errors of the iterates of IRLS _{p} for different choices of p in the case that the residual is exactly k -sparse. We observe that IRLS₁ exhibits a linear error decay on the one hand and a superlinear decay for IRLS _{p} with $p = 0.5$ and $p = 0.1$ that accelerates as p decreases towards 0, confirming the rates predicted by Theorem 1 and Theorem 2. ($p = 1, 0.5, 0.1$ are rather arbitrary choices in our experiments.)

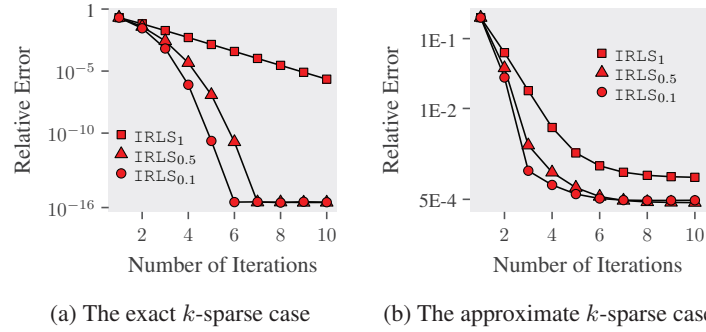


Figure 4: Empirical convergence rates of IRLS_p for different values of p ($m = 1000, n = 10, k = 200$, averaged over 20 trials). Figure 4a: $\mathbf{Ax}^* - \mathbf{y}$ is exactly k -sparse. Figure 4b: $\mathbf{Ax}^* - \mathbf{y}$ is approximately k -sparse with 1% noise (see Section ?? for data generation).

In another experiment, visualized in Figure 4b, we explore the robustness of IRLS_p if the model assumption of k -sparse residuals is not satisfied, but only approximately. In particular, the input vector $\mathbf{y} \in \mathbb{R}^m$ of IRLS_p is now such that its restriction on S^* with $|S^*| = k$ has i.i.d. $\mathcal{N}(0, 1)$ entries (corresponding to the sparse corruptions), but its other coordinates are distributed as independent Gaussian variables with mean $\mathbf{a}_i^\top \mathbf{x}^*$ and variance 0.01, where $i \in (S^*)^c$ (corresponding to dense noise); see also the code provided in Appendix ?? . We observe that also for this case of only approximately k -sparse residuals, IRLS_p performs well, with faster and more accurate convergence for $p = 0.1$ and $p = 0.5$ compared to $p = 1$. We note that, strictly speaking, Theorem 1 and Theorem 2 do not apply directly to this case of approximately k -sparse residuals, but it is not hard to generalize it to this case; see [27, Theorem A.1] for a similar result for IRLS applied to basis pursuit.

5 Discussion and Future Work

In this work, we established novel global linear and local superlinear rates of IRLS_p for ℓ_p -regression (1) under the assumption of the stable range space property. Furthermore, we explored several applications for which IRLS_p exhibits state-of-the-art results.

There are several directions that deserve further investigation. Theoretically, the Gaussian assumption of Proposition 3 might be weakened; similar results would hold for sub-Gaussian distributions [86, 87]. Note also that, in experiments (Figures 1 and 2), IRLS_p works well even when $m - n \rightarrow 2k$. This leaves the question of whether the constant or logarithmic factors of condition (10) are too stringent for guaranteeing the stable RSP to hold; can this be improved? Finally, it is tempting to carry out a global rate analysis of IRLS_p with $p \in (0, 1)$ for ℓ_p -regression; can one prove global (linear) rates under the stable RSP or other assumption?

Algorithmically, designing inexact solvers for the inner (weighted) least-squares problem (2) based on Krylov-subspace methods [101, 102] is expected to further accelerate our current IRLS implementation (our current implementation is attached in Appendix ??). Also, our empirical experiments suggest that IRLS_p (in particular $p \in (0, 1)$) is worth being adopted and applied to many other outlier-robust estimation tasks beyond robust regression [45, 48, 49, 50, 51, 52, 103, 104, 105, 106, 107, 108].

Acknowledgments and Disclosure of Funding

This work was supported by grants NSF 1704458, NSF 1934979, NSF-IIS-1837991, ONR MURI 503405-78051, and the NSF-Simons Collaboration on the Mathematical Foundations of Deep Learning (NSF grant 2031985).

References

- [1] A. El Alaoui, X. Cheng, A. Ramdas, M. J. Wainwright, and M. I. Jordan, “Asymptotic behavior of ℓ_p -based Laplacian regularization in semi-supervised learning,” in *Conference on Learning Theory*,

pp. 879–906, 2016.

- [2] D. Adil, R. Peng, and S. Sachdeva, “Fast, provably convergent IRLS algorithm for p -norm linear regression,” *Advances in Neural Information Processing Systems*, vol. 32, 2019.
- [3] D. Slepcev and M. Thorpe, “Analysis of p -Laplacian regularization in semisupervised learning,” *SIAM Journal on Mathematical Analysis*, vol. 51, no. 3, pp. 2085–2120, 2019.
- [4] S. Bubeck, M. B. Cohen, Y. T. Lee, and Y. Li, “An homotopy method for ℓ_p regression provably beyond self-concordance and in input-sparsity time,” in *ACM Symposium on Theory of Computing*, pp. 1130–1137, 2018.
- [5] D. Adil, R. Kyng, R. Peng, and S. Sachdeva, “Iterative refinement for ℓ_p -norm regression,” in *ACM-SIAM Symposium on Discrete Algorithms*, pp. 1405–1424, 2019.
- [6] A. Jambulapati, Y. P. Liu, and A. Sidford, “Improved iteration complexities for overconstrained p -norm regression,” tech. rep., arXiv:2111.01848v2 [cs.DS], 2021.
- [7] R. J. Boscovich, “De litteraria expeditione per pontificiam ditionem, et synopsis amplioris operis, ac habentur plura ejus ex exemplaria etiam sensorum impressa,” *Bonomiensi Scientiarum et Artum Instituto Atque Academia Commentarii*, vol. 4, pp. 353–396, 1757.
- [8] R. L. Plackett, “Studies in the History of Probability and Statistics. XXIX: The discovery of the method of least squares,” *Biometrika*, vol. 59, pp. 239–251, 08 1972.
- [9] J. Wright and Y. Ma, *High-Dimensional Data Analysis with Low-Dimensional Models: Principles, Computation, and Applications*. Cambridge University Press, 2020.
- [10] P. J. Huber, “Robust estimation of a location parameter,” *The Annals of Mathematical Statistics*, vol. 35, no. 1, pp. 73–101, 1964.
- [11] E. Candes, M. Rudelson, T. Tao, and R. Vershynin, “Error correction via linear programming,” in *IEEE Symposium on Foundations of Computer Science*, pp. 668–681, 2005.
- [12] K. Bhatia, P. Jain, and P. Kar, “Robust regression via hard thresholding,” *Advances in Neural Information Processing Systems*, vol. 28, 2015.
- [13] K. Bhatia, P. Jain, P. Kamalaruban, and P. Kar, “Consistent robust regression,” *Advances in Neural Information Processing Systems*, vol. 30, 2017.
- [14] J. Yang, Y.-L. Chow, C. Ré, and M. W. Mahoney, “Weighted SGD for ℓ_p regression with randomized preconditioning,” *The Journal of Machine Learning Research*, vol. 18, no. 1, pp. 7811–7853, 2017.
- [15] D. Durfee, K. A. Lai, and S. Sawlani, “ ℓ_1 regression using lewis weights preconditioning and stochastic gradient descent,” in *Conference On Learning Theory*, pp. 1626–1656, 2018.
- [16] A. S. Suggala, K. Bhatia, P. Ravikumar, and P. Jain, “Adaptive hard thresholding for near-optimal consistent robust regression,” in *Conference on Learning Theory*, pp. 2892–2897, 2019.
- [17] B. Mukhoty, G. Gopakumar, P. Jain, and P. Kar, “Globally-convergent iteratively reweighted least squares for robust regression problems,” in *International Conference on Artificial Intelligence and Statistics*, pp. 313–322, 2019.
- [18] S. Pesme and N. Flammarion, “Online robust regression via sgd on the ℓ_1 loss,” *Advances in Neural Information Processing Systems*, vol. 33, pp. 2540–2552, 2020.
- [19] X. Chen and M. Derezhinski, “Query complexity of least absolute deviation regression via robust uniform convergence,” in *Conference on Learning Theory*, pp. 1144–1179, 2021.
- [20] A. Parulekar, A. Parulekar, and E. Price, “ ℓ_1 regression with Lewis weights subsampling,” tech. rep., arXiv:2105.09433 [cs.LG], 2021.
- [21] E. J. Candes and T. Tao, “Decoding by linear programming,” *IEEE Transactions on Information Theory*, vol. 51, no. 12, pp. 4203–4215, 2005.
- [22] E. J. Candes, J. K. Romberg, and T. Tao, “Stable signal recovery from incomplete and inaccurate measurements,” *Communications on Pure and Applied Mathematics*, vol. 59, no. 8, pp. 1207–1223, 2006.
- [23] D. L. Donoho, “For most large underdetermined systems of linear equations the minimal ℓ_1 -norm solution is also the sparsest solution,” *Communications on Pure and Applied Mathematics*, vol. 59, no. 6, pp. 797–829, 2006.
- [24] A. Cohen, W. Dahmen, and R. DeVore, “Compressed sensing and best k -term approximation,” *Journal of the American Mathematical Society*, vol. 22, no. 1, pp. 211–231, 2009.
- [25] T. Blumensath and M. E. Davies, “Iterative hard thresholding for compressed sensing,” *Applied and Computational Harmonic Analysis*, vol. 27, no. 3, pp. 265–274, 2009.
- [26] I. Daubechies, R. DeVore, M. Fornasier, and C. S. Güntürk, “Iteratively reweighted least squares minimization for sparse recovery,” *Communications on Pure and Applied Mathematics*, vol. 63, no. 1, pp. 1–38, 2010.

- [27] C. Kümmerle, C. Mayrink Verdun, and D. Stöger, “Iteratively reweighted least squares for basis pursuit with global linear convergence rate,” *Advances in Neural Information Processing Systems*, 2021.
- [28] S. Foucart and H. Rauhut, *A Mathematical Introduction to Compressive Sensing*. Springer New York, 2013.
- [29] R. Chartrand, “Nonconvex compressed sensing and error correction,” in *IEEE International Conference on Acoustics, Speech and Signal Processing*, vol. 3, pp. III–889, 2007.
- [30] R. Chartrand, “Exact reconstruction of sparse signals via nonconvex minimization,” *IEEE Signal Processing Letters*, vol. 14, no. 10, pp. 707–710, 2007.
- [31] R. Chartrand and V. Staneva, “Restricted isometry properties and nonconvex compressive sensing,” *Inverse Problems*, vol. 24, no. 3, p. 035020, 2008.
- [32] S. Foucart and M.-J. Lai, “Sparsest solutions of underdetermined linear systems via ℓ_q -minimization for $0 < q \leq 1$,” *Applied and Computational Harmonic Analysis*, vol. 26, no. 3, pp. 395–407, 2009.
- [33] M. Wang, W. Xu, and A. Tang, “On the performance of sparse recovery via ℓ_p -minimization ($0 \leq p \leq 1$),” *IEEE Transactions on Information Theory*, vol. 57, no. 11, pp. 7255–7278, 2011.
- [34] Q. Sun, “Recovery of sparsest signals via ℓ_q -minimization,” *Applied and Computational Harmonic Analysis*, vol. 32, no. 3, pp. 329–341, 2012.
- [35] D. Ba, B. Babadi, P. L. Purdon, and E. N. Brown, “Convergence and stability of iteratively re-weighted least squares algorithms,” *IEEE Transactions on Signal Processing*, vol. 62, no. 1, pp. 183–195, 2013.
- [36] Z. Zhou, Q. Zhang, and A. M.-C. So, “ $\ell_{1,p}$ -norm regularization: Error bounds and convergence rate analysis of first-order methods,” in *International Conference on Machine Learning*, vol. 37, pp. 1501–1510, 2015.
- [37] L. Zheng, A. Maleki, H. Weng, X. Wang, and T. Long, “Does ℓ_p -minimization outperform ℓ_1 -minimization?,” *IEEE Transactions on Information Theory*, vol. 63, no. 11, pp. 6896–6935, 2017.
- [38] G. Marjanovic and V. Solo, “On ℓ_q optimization and matrix completion,” *IEEE Transactions on Signal Processing*, vol. 60, no. 11, pp. 5714–5724, 2012.
- [39] K. Mohan and M. Fazel, “Iterative reweighted algorithms for matrix rank minimization,” *The Journal of Machine Learning Research*, vol. 13, no. 1, pp. 3441–3473, 2012.
- [40] F. Nie, H. Huang, and C. Ding, “Low-rank matrix recovery via efficient Schatten p -norm minimization,” in *AAAI Conference on Artificial Intelligence*, 2012.
- [41] C. Kümmerle and J. Sigl, “Harmonic mean iteratively reweighted least squares for low-rank matrix recovery,” *The Journal of Machine Learning Research*, vol. 19, no. 1, pp. 1815–1863, 2018.
- [42] P. Giampouras, R. Vidal, A. Rontogiannis, and B. Haeffele, “A novel variational form of the Schatten- p quasi-norm,” *Advances in Neural Information Processing Systems*, vol. 33, pp. 21453–21463, 2020.
- [43] C. Kümmerle and C. Mayrink Verdun, “A scalable second order method for ill-conditioned matrix completion from few samples,” in *International Conference on Machine Learning*, pp. 5872–5883, 2021.
- [44] G. Lerman, M. B. McCoy, J. A. Tropp, and T. Zhang, “Robust computation of linear models by convex relaxation,” *Foundations of Computational Mathematics*, vol. 15, no. 2, pp. 363–410, 2015.
- [45] M. C. Tsakiris and R. Vidal, “Dual principal component pursuit,” *Journal of Machine Learning Research*, vol. 19, no. 18, pp. 1–50, 2018.
- [46] Z. Zhu, Y. Wang, D. Robinson, D. Naiman, R. Vidal, and M. C. Tsakiris, “Dual principal component pursuit: Improved analysis and efficient algorithms,” in *Advances in Neural Information Processing Systems*, 2018.
- [47] T. Ding, Z. Zhu, T. Ding, Y. Yang, R. Vidal, M. C. Tsakiris, and D. Robinson, “Noisy dual principal component pursuit,” in *International Conference on Machine Learning*, pp. 1617–1625, 2019.
- [48] K. Aftab and R. Hartley, “Convergence of iteratively re-weighted least squares to robust M-estimators,” in *IEEE Winter Conference on Applications of Computer Vision*, pp. 480–487, 2015.
- [49] W. Dong, X.-j. Wu, and J. Kittler, “Sparse subspace clustering via smoothed ℓ_p minimization,” *Pattern Recognition Letters*, vol. 125, pp. 206–211, 2019.
- [50] D. Iwata, M. Waechter, W.-Y. Lin, and Y. Matsushita, “An analysis of sketched IRLS for accelerated sparse residual regression,” in *European Conference on Computer Vision*, pp. 609–626, 2020.
- [51] T. Ding, Y. Yang, Z. Zhu, D. P. Robinson, R. Vidal, L. Kneip, and M. C. Tsakiris, “Robust homography estimation via dual principal component pursuit,” in *IEEE Conference on Computer Vision and Pattern Recognition*, pp. 6080–6089, 2020.
- [52] H. Yang, P. Antonante, V. Tzoumas, and L. Carlone, “Graduated non-convexity for robust spatial perception: From non-minimal solvers to global outlier rejection,” *IEEE Robotics and Automation Letters*, vol. 5, no. 2, pp. 1127–1134, 2020.

- [53] Y.-B. Zhao and D. Li, “Reweighted ℓ_1 -minimization for sparse solutions to underdetermined linear systems,” *SIAM Journal on Optimization*, vol. 22, no. 3, pp. 1065–1088, 2012.
- [54] A. Beck, *First-Order Methods in Optimization*. Society for Industrial and Applied Mathematics, 2017.
- [55] Y. Nesterov, *Lectures on Convex Optimization*. Springer, 2018.
- [56] Y. Chen, Y. Chi, J. Fan, and C. Ma, “Gradient descent with random initialization: fast global convergence for nonconvex phase retrieval,” *Mathematical Programming*, vol. 176, no. 1, pp. 5–37, 2019.
- [57] Y. S. Tan and R. Vershynin, “Phase retrieval via randomized Kaczmarz: Theoretical guarantees,” *Information and Inference: A Journal of the IMA*, vol. 8, no. 1, pp. 97–123, 2019.
- [58] J. Unnikrishnan, S. Haghighatshoar, and M. Vetterli, “Unlabeled sensing with random linear measurements,” *IEEE Transactions on Information Theory*, vol. 64, no. 5, pp. 3237–3253, 2018.
- [59] M. Slawski and E. Ben-David, “Linear regression with sparsely permuted data,” *Electronic Journal of Statistics*, vol. 13, no. 1, pp. 1–36, 2019.
- [60] M. C. Tsakiris, L. Peng, A. Conca, L. Kneip, Y. Shi, and H. Choi, “An algebraic-geometric approach for linear regression without correspondences,” *IEEE Transactions on Information Theory*, vol. 66, no. 8, pp. 5130–5144, 2020.
- [61] L. Peng and M. C. Tsakiris, “Linear regression without correspondences via concave minimization,” *IEEE Signal Processing Letters*, vol. 27, pp. 1580–1584, 2020.
- [62] J. Wright, A. Y. Yang, A. Ganesh, S. S. Sastry, and Y. Ma, “Robust face recognition via sparse representation,” *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 31, no. 2, pp. 210–227, 2008.
- [63] C. Solomon and T. Breckon, *Fundamentals of Digital Image Processing: A practical approach with examples in Matlab*. John Wiley & Sons, 2011.
- [64] Y. Yao, L. Peng, and M. Tsakiris, “Unlabeled principal component analysis,” *Advances in Neural Information Processing Systems*, 2021.
- [65] R. Balan, P. Casazza, and D. Edidin, “On signal reconstruction without phase,” *Applied and Computational Harmonic Analysis*, vol. 20, no. 3, pp. 345 – 356, 2006.
- [66] T. Strohmer and R. Vershynin, “A randomized Kaczmarz algorithm with exponential convergence,” *Journal of Fourier Analysis and Applications*, vol. 15, no. 2, pp. 262–278, 2009.
- [67] K. Wei, “Solving systems of phaseless equations via Kaczmarz methods: A proof of concept study,” *Inverse Problems*, vol. 31, no. 12, p. 125008, 2015.
- [68] O. Dhifallah, C. Thrampoulidis, and Y. M. Lu, “Phase retrieval via linear programming: Fundamental limits and algorithmic improvements,” in *Allerton Conference on Communication, Control, and Computing*, pp. 1071–1077, 2017.
- [69] Y. Chen and E. Candes, “Solving random quadratic systems of equations is nearly as easy as solving linear systems,” *Advances in Neural Information Processing Systems*, vol. 28, 2015.
- [70] W.-J. Zeng and H.-C. So, “Coordinate descent algorithms for phase retrieval,” tech. rep., arXiv:1706.03474 [cs.IT], 2017.
- [71] D. Applegate, M. Díaz, O. Hinder, H. Lu, M. Lubin, B. O’Donoghue, and W. Schudy, “Practical large-scale linear programming using primal-dual hybrid gradient,” *Advances in Neural Information Processing Systems*, vol. 34, 2021.
- [72] “Gurobi 9.5.1.” <https://www.gurobi.com/>, Date Accessed: May 10, 2022.
- [73] A. Beck and N. Guttman-Beck, “FOM – a MATLAB toolbox of first-order methods for solving convex optimization problems,” *Optimization Methods and Software*, vol. 34, no. 1, pp. 172–193, 2019.
- [74] A. S. Georgiades, P. N. Belhumeur, and D. J. Kriegman, “From few to many: Illumination cone models for face recognition under variable lighting and pose,” *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 23, no. 6, pp. 643–660, 2001.
- [75] S. S. Chen, D. L. Donoho, and M. A. Saunders, “Atomic decomposition by basis pursuit,” *SIAM Review*, vol. 43, no. 1, pp. 129–159, 2001.
- [76] G. Lerman and T. Maunu, “Fast, robust and non-convex subspace recovery,” *Information and Inference: A Journal of the IMA*, vol. 7, no. 2, pp. 277–336, 2018.
- [77] Q. Qu, Z. Zhu, X. Li, M. C. Tsakiris, J. Wright, and R. Vidal, “Finding the sparsest vectors in a subspace: Theory, algorithms, and applications,” tech. rep., arXiv:2001.06970 [cs.LG], 2020.
- [78] T. F. Chan and P. Mulet, “On the convergence of the lagged diffusivity fixed point method in total variation image restoration,” *SIAM Journal on Numerical Analysis*, vol. 36, no. 2, pp. 354–367, 1999.

- [79] A. Beck, “On the convergence of alternating minimization for convex programming with applications to iteratively reweighted least squares and decomposition schemes,” *SIAM Journal on Optimization*, vol. 25, no. 1, pp. 185–209, 2015.
- [80] A. Y. Aravkin, J. V. Burke, and D. He, “IRLS for sparse recovery revisited: Examples of failure and a remedy,” tech. rep., arXiv:1910.07095 [math.ST], 2019.
- [81] X. Yang, J. Wang, and H. Wang, “Towards an efficient approach for the nonconvex ℓ_p -ball projection: Algorithm and analysis,” tech. rep., arXiv:2101.01350v5 [math.OC], 2021.
- [82] A. M. Tillmann and M. E. Pfetsch, “The computational complexity of the restricted isometry property, the nullspace property, and related concepts in compressed sensing,” *IEEE Transactions on Information Theory*, vol. 60, no. 2, pp. 1248–1259, 2013.
- [83] S. Arora and B. Barak, *Computational Complexity: A Modern Approach*. Cambridge University Press, 2009.
- [84] S. Boyd, S. P. Boyd, and L. Vandenberghe, *Convex optimization*. Cambridge University Press, 2004.
- [85] Z. Allen-Zhu, “Katyusha: The first direct acceleration of stochastic gradient methods,” *The Journal of Machine Learning Research*, vol. 18, no. 1, pp. 8194–8244, 2017.
- [86] G. Lecué and S. Mendelson, “Sparse recovery under weak moment assumptions,” *Journal of the European Mathematical Society*, vol. 19, no. 3, pp. 881–904, 2017.
- [87] S. Dirksen, G. Lecué, and H. Rauhut, “On the gap between restricted isometry properties and sparse recovery conditions,” *IEEE Transactions on Information Theory*, vol. 64, no. 8, pp. 5478–5487, 2016.
- [88] M. C. Tsakiris and L. Peng, “Homomorphic sensing,” in *International Conference on Machine Learning*, 2019.
- [89] L. Peng, B. Wang, and M. Tsakiris, “Homomorphic sensing: Sparsity and noise,” in *International Conference on Machine Learning*, pp. 8464–8475, 2021.
- [90] L. Peng and M. C. Tsakiris, “Homomorphic sensing of subspace arrangements,” *Applied and Computational Harmonic Analysis*, vol. 55, pp. 466–485, 2021.
- [91] P. Grohs, S. Koppensteiner, and M. Rathmair, “Phase retrieval: Uniqueness and stability,” *SIAM Review*, vol. 62, no. 2, pp. 301–350, 2020.
- [92] M. Huang and Y. Wang, “Linear convergence of randomized Kaczmarz method for solving complex-valued phaseless equations,” tech. rep., arXiv:2109.11811 [math.NA], 2021.
- [93] R. Chandra, T. Goldstein, and C. Studer, “PhasePack: A phase retrieval library,” in *International conference on Sampling Theory and Applications*, pp. 1–5, 2019.
- [94] A. Pananjady, M. J. Wainwright, and T. A. Courtade, “Linear regression with shuffled data: Statistical and computational limits of permutation recovery,” *IEEE Transactions on Information Theory*, vol. 64, no. 5, pp. 3286–3300, 2018.
- [95] M. Slawski, G. Diao, and E. Ben-David, “A pseudo-likelihood approach to linear regression with partially shuffled data,” *Journal of Computational and Graphical Statistics*, vol. 0, no. 0, pp. 1–31, 2021.
- [96] Y. Xie, Y. Mao, S. Zuo, H. Xu, X. Ye, T. Zhao, and H. Zha, “A hypergradient approach to robust regression without correspondence,” in *International Conference on Learning Representations*, 2021.
- [97] F. Li, K. Fujiwara, F. Okura, and Y. Matsushita, “Generalized shuffled linear regression,” in *IEEE/CVF International Conference on Computer Vision*, pp. 6474–6483, 2021.
- [98] D. Hsu, K. Shi, and X. Sun, “Linear regression without correspondence,” in *Advances in Neural Information Processing Systems*, 2017.
- [99] C. You, D. Robinson, and R. Vidal, “Scalable sparse subspace clustering by orthogonal matching pursuit,” in *IEEE Conference on Computer Vision and Pattern Recognition*, pp. 3918–3927, 2016.
- [100] R. Basri and D. W. Jacobs, “Lambertian reflectance and linear subspaces,” *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 25, no. 2, pp. 218–233, 2003.
- [101] M. Fornasier, S. Peter, H. Rauhut, and S. Worm, “Conjugate gradient acceleration of iteratively reweighted least squares methods,” *Computational Optimization and Applications*, vol. 65, no. 1, pp. 205–259, 2016.
- [102] S. Gazzola, C. Meng, and J. G. Nagy, “Krylov methods for low-rank regularization,” *SIAM Journal on Matrix Analysis and Applications*, vol. 41, no. 4, pp. 1477–1504, 2020.
- [103] A. Chatterjee and V. M. Govindu, “Robust relative rotation averaging,” *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 40, no. 4, pp. 958–972, 2017.
- [104] C. Sidhartha and V. M. Govindu, “It is all in the weights: Robust rotation averaging revisited,” in *International Conference on 3D Vision*, pp. 1134–1143, 2021.

- [105] L. Peng, M. C. Tsakiris, and R. Vidal, “ARCS: Accurate rotation and correspondences search,” in *IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 11153–11163, 2022.
- [106] S. H. Lee and J. Civera, “HARA: A hierarchical approach for robust rotation averaging,” in *IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 15777–15786, 2022.
- [107] L. Peng, M. Fazlyab, and R. Vidal, “Semidefinite relaxations of truncated least-squares in robust rotation search: Tight or not,” in *European Conference on Computer Vision*, pp. 0–0, Springer, 2022.
- [108] L. Carlone, “Estimation contracts for outlier-robust geometric perception,” tech. rep., arXiv:2208.10521 [stat.ML], 2022.

Checklist

1. For all authors...
 - (a) Do the main claims made in the abstract and introduction accurately reflect the paper’s contributions and scope? **[Yes]** Furthermore, we made a lot of efforts in comparing the theoretical results of prior works and position our contribution in the literature.
 - (b) Did you describe the limitations of your work? **[Yes]** See for example at Lines 268-274 on real phase retrieval. We also discussed several directions for improvements or future research at the end of the paper.
 - (c) Did you discuss any potential negative societal impacts of your work? **[No]**
 - (d) Have you read the ethics review guidelines and ensured that your paper conforms to them? **[Yes]**
2. If you are including theoretical results...
 - (a) Did you state the full set of assumptions of all theoretical results? **[Yes]** In particular, we used an entire subsection (Section 3.1) to justify the major assumption, the stable range space property.
 - (b) Did you include complete proofs of all theoretical results? **[Yes]** We omitted the proofs of Proposition 1 and 2 as they follow directly from existing results; see Line 83 and Line 155. All other theoretical results are proved in the appendix.
3. If you ran experiments...
 - (a) Did you include the code, data, and instructions needed to reproduce the main experimental results (either in the supplemental material or as a URL)? **[Yes]** We attached a copy of our code in the supplemental material.
 - (b) Did you specify all the training details (e.g., data splits, hyperparameters, how they were chosen)? **[Yes]** In the appendix, we use an entire section to describe our experimental setup. We did not do that in the main paper in interest of space.
 - (c) Did you report error bars (e.g., with respect to the random seed after running experiments multiple times)? **[No]** We did not report error bars for two reasons. The main reason is that some of our figures are quite dense, and adding error bars on them makes it harder to read. A slightly subjective reason is that we believe our current figures are informative enough, correctly deliver the message, and are fair for other methods.
 - (d) Did you include the total amount of compute and the type of resources used (e.g., type of GPUs, internal cluster, or cloud provider)? **[N/A]**
4. If you are using existing assets (e.g., code, data, models) or curating/releasing new assets...
 - (a) If your work uses existing assets, did you cite the creators? **[Yes]**
 - (b) Did you mention the license of the assets? **[No]** In the main paper we used the **Extended Yale B face dataset**, which is free for research purposes.
 - (c) Did you include any new assets either in the supplemental material or as a URL? **[No]**
 - (d) Did you discuss whether and how consent was obtained from people whose data you’re using/curating? **[Yes]**
 - (e) Did you discuss whether the data you are using/curating contains personally identifiable information or offensive content? **[No]**
5. If you used crowdsourcing or conducted research with human subjects...
 - (a) Did you include the full text of instructions given to participants and screenshots, if applicable? **[N/A]**

- (b) Did you describe any potential participant risks, with links to Institutional Review Board (IRB) approvals, if applicable? [N/A]
- (c) Did you include the estimated hourly wage paid to participants and the total amount spent on participant compensation? [N/A]