

High-Probability Polynomial-Time Complexity of Restarted PDHG for Linear Programming

Zikai Xiong

Department of Industrial Engineering and Management Sciences, Northwestern University, 2145 Sheridan Road, Evanston, IL 60208, USA.
zikai.xiong@northwestern.edu.

Abstract. The restarted primal-dual hybrid gradient method (rPDHG) is a first-order method recently known for its computational effectiveness in solving linear programming (LP) problems. Despite its impressive practical performance, the theoretical iteration bounds for rPDHG can be exponentially poor. To investigate this gap from a probabilistic perspective, we show that rPDHG achieves polynomial-time complexity in a high-probability sense, under assumptions on the probability distribution from which the data instance is generated. We consider both Gaussian and more general sub-Gaussian distribution models. For standard-form LP instances with m constraints and n variables, our bounds take a particularly simple form when m is not too close to n : rPDHG iterates settle on the optimal basis in $\tilde{O}\left(\frac{n^{2.5}m^{0.5}}{\delta}\right)$ iterations, followed by $O\left(\frac{n^{0.5}m^{0.5}}{\delta} \ln\left(\frac{1}{\varepsilon}\right)\right)$ iterations to compute an ε -optimal solution, with probability at least $1 - \delta$ for δ that is not exponentially small. The Stage-I bound further improves to $\tilde{O}\left(\frac{n^{2.5}}{\delta}\right)$ in the Gaussian distribution model. Experimental results confirm the tail behavior and the polynomial-time dependence on problem dimensions of the iteration counts. As an application of our probabilistic analysis, we explore how the disparity among the components of the optimal solution bears on the performance of rPDHG, and we provide guidelines for generating challenging LP instances.

Key words: Linear optimization, First-order method, Probabilistic analysis, High-probability complexity

MSC2020 subject classification: Primary: 90C05; Secondary: 90C60

1. Introduction Linear programming (LP) has been one of the most fundamental areas of optimization since the 1950s, with applications spanning many domains. For over 70 years, researchers and practitioners have extensively studied and refined LP solution methods, primarily through two major algorithmic frameworks: the simplex method, introduced by Dantzig in 1947, and the interior-point method, developed by Karmarkar in 1984. The profound impact of these methods on both academia and industry cannot be overstated.

While the two classic methods (simplex and interior-point methods) are successful in solving many LP problems, both of them rely on frequent matrix factorizations, whose computational cost grows superlinearly with problem dimensions. Moreover, these matrix factorizations cannot efficiently utilize modern computational architectures, especially graphics processing units (GPUs). These limitations have motivated the development of the restarted primal-dual hybrid gradient method (rPDHG), a first-order method that eliminates the need for matrix factorizations. By utilizing primarily matrix-vector products for gradient computations, rPDHG achieves better scalability by exploiting problem sparsity and leveraging parallel architectures. It directly addresses the saddle-point formulation (see Applegate et al. [8]), and also solves conic linear programs (see Xiong and Freund [58]) and convex quadratic programs (see Huang et al. [28], Lu and Yang [34]).

The rPDHG method has demonstrated performance comparable to, and sometimes exceeding, that of classical simplex and interior-point methods on many LP instances. Its implementations span both CPUs (PDLP by Applegate et al. [7]) and GPUs (cuPDLP and cuPDLP-C by Lu and Yang [32] and Lu et al. [35]). Recognizing its effectiveness, leading commercial solvers including COPT 7.1 (see Ge et al. [20]), Xpress 9.4 (see Biele and Gade [9]), and Gurobi 13.0 (see Gurobi Optimization, LLC [22]) have integrated rPDHG as a third foundational LP algorithm alongside simplex and interior-point methods. The method has also been adopted by other solvers such as Google OR-Tools (see Applegate et al. [7]), HiGHS (see HiGHS Development Team [24]), and NVIDIA cuOpt (see Blin [10]).

A persistent challenge in LP algorithm development has been the “embarrassing gap” between the practical experience of these LP algorithms’ efficiency and their worst-case iteration bounds. The efficiency is typically measured by the iteration count (to solve the problem) as a function of problem dimensions. For simplex methods, empirical evidence suggests it is usually a low-degree polynomial or even a linear function of the

dimensions (see, for example, Shamir [44]), but worst-case analyses reveal this function could potentially be exponentially large (see, for example, Klee and Minty [30]). For interior-point methods, worst-case iteration bounds are usually a polynomial of the dimension times the “bit-length” L of the problem, but the practical performance is usually much better than the worst-case analyses would suggest. This gap between theory and practice has motivated extensive research into average-case complexity analysis for both classic LP methods, which we review shortly after.

A similar gap exists between the theoretical bound and empirical performance for rPDHG. While the method demonstrates strong practical performance on most LP problems (see, e.g., Applegate et al. [7], Lu and Yang [32]), its theoretical worst-case iteration bounds can be exponentially poor. Recent work has made significant progress in understanding rPDHG’s convergence behavior. Applegate et al. [8] establish the first linear convergence rate of rPDHG using the global Hoffman constant of the matrix of the KKT system that defines the optimal solutions. Xiong and Freund [57, 59] provide a tighter computational guarantee for rPDHG using two natural properties of the LP problem: LP sharpness and the “limiting error ratio.” Furthermore, Xiong [56] gives a closed-form iteration bound for LPs with unique optima and demonstrates that rPDHG has a two-stage performance. In Stage I, the iterates settle on the optimal basis (and thus identify this basis). Stage II then exhibits faster local convergence to the optimal solution. However, all of the existing computational guarantees at least linearly depend on certain condition measures that can be exponentially large in the worst case, which leaves unexplained the strong practical performance of rPDHG observed in many applications.

To help address this gap between theory and practice from a probabilistic perspective, this paper primarily aims at answering the following question:

Q1. Can we prove that rPDHG is efficient with high probability under a probabilistic input model?

Here, “efficient” means having polynomial-time complexity. An affirmative answer to this question would provide a theoretical explanation for rPDHG’s strong performance under the probabilistic model and, more importantly, offer insights into its behavior beyond the worst case. We approach this question through probabilistic analysis and give an affirmative answer in this paper.

Probabilistic analysis has successfully narrowed similar theory-practice gaps for classic LP methods. By assuming the probability distributions according to which the problem data was generated (or equivalently, the distributions on the input data), researchers can provide bounds on the expected number of iterations. For the simplex method, breakthrough results in the 1980s established various polynomial iteration bounds under various probabilistic assumptions (see, for example, Adler et al. [1, 2], Adler and Megiddo [3], Borgwardt [12, 13, 14, 15], Haimovich [23], Megiddo [36], Smale [45, 46]). Similarly, for interior-point methods, Todd [51] proposes a probabilistic model. Various versions of this model lead to polynomial expected iteration bounds in terms of problem dimensions, independent of the bit-length L (see, for example, Anstreicher et al. [5, 6], Huang [26, 27], Ji and Potra [29], Ye [61]). More recently, another probabilistic analysis approach called smoothed analysis provides a new framework that bridges worst-case and average-case analyses. Polynomial-time complexity bounds have also been established under this framework for both simplex and interior-point methods for LP instances (see, for example, Dadush and Huiberts [17], Dunagan et al. [19], Spielman and Teng [47, 48]).

Probabilistic analysis may also provide new tools for revealing new insights into rPDHG. Some practical observations of rPDHG, such as the two-stage performance and sensitivity under perturbations, have been partially explained by the worst-case iteration bounds established by Xiong [56] and others. Probabilistic analysis may provide complementary perspectives and clearer insights into these behaviors of rPDHG.

Despite the above, it should be admitted that real-life LP problems often differ significantly from the probabilistic models that people assume for probabilistic analyses. For example, almost all the existing probabilistic models are based on either the Gaussian distribution or the sign-invariant distribution (see, e.g., Shamir [44] and Todd [51]). These distributional assumptions rarely match the data generation processes of practical LP instances. Consequently, extending these models to broader distribution classes remains an important research direction (see the survey by Shamir [44]).

Therefore, beyond the main question Q1, this paper also addresses two additional questions:

- Q2. Can the probabilistic analysis be extended beyond Gaussian distributions to more general distribution families?
- Q3. What novel insights into rPDHG's performance can be gained through probabilistic analysis?

To address the above three questions, we study a probabilistic model that is built on the classic probabilistic model proposed by Todd [51]. Our model considers standard-form LP instances where the components of the constraint matrix follow specified Gaussian or, more generally, sub-Gaussian distributions, while the right-hand side and the objective vectors are constructed based on random primal and dual solutions. This approach builds on Todd [51]'s framework, which has been dominant in probabilistic analyses of interior-point methods (see, e.g., Anstreicher et al. [5, 6] and Ye [61]). Compared with existing probabilistic models, our work extends the distribution of the input data from Gaussian to general sub-Gaussian distributions, which include Bernoulli and bounded distributions. This extension substantially broadens the class of random LP instances covered by our analysis. Our analysis leverages recent advances in nonasymptotic random matrix theory developed over the past two decades (see, e.g., the ICM 2010 lecture by Rudelson and Vershynin [42]).

1.1. Outline and contributions In this paper, we consider LP instances in standard form:

$$\min_{x \in \mathbb{R}^n} \quad c^\top x \quad \text{s.t. } Ax = b, \ x \geq 0 \quad (1.1)$$

where the constraint matrix $A \in \mathbb{R}^{m \times n}$ and $m \leq n$. Thus, the problem has m linear equality constraints and n decision variables. We also denote by $d := n - m$ their difference. Any LP instance can be reformulated equivalently in the standard form (1.1). The dual to the problem (1.1) can be expressed as:

$$\max_{y \in \mathbb{R}^m, s \in \mathbb{R}^n} \quad b^\top y \quad \text{s.t. } A^\top y + s = c, \ s \geq 0 \quad (1.2)$$

with $s \in \mathbb{R}^n$ denoting the slack. We let $x^\star$ be any optimal solution of (1.1) and let $s^\star$ be any optimal dual slack of (1.2). The rest of the paper is organized as follows.

Section 2 introduces the rPDHG algorithm for solving linear programs.

Section 3 presents our main result: the high-probability polynomial-time complexity of rPDHG. We begin by introducing our probabilistic model with sub-Gaussian input data in Section 3.1. Building on rPDHG's two-stage behavior, in which Stage I settles on the optimal basis and Stage II achieves fast local convergence, Section 3.2 presents high-probability iteration bounds for both stages. Section 3.3 then presents improved bounds under Gaussian input data. Table 1 contains a preview of these results in the case where m is not too close to n . Here ε denotes the target error tolerance of the desired solution. These results establish high-probability polynomial-time complexity guarantees for rPDHG (addressing Q1), and extend the probabilistic LP complexity analysis to sub-Gaussian input data (addressing Q2).

Sections 4 and 5 present the proofs of our results in Sections 3.2 and 3.3, respectively.

Section 6 presents the experimental results. They confirm the tail behavior of the iteration counts for both stages (the linear dependence on $\frac{1}{\delta}$), and also validate the polynomial dependence on problem dimensions in the high-probability iteration counts.

Section 7 provides new insights using probabilistic analysis (addressing Q3). We investigate how the disparity ratio $\phi := \frac{\frac{1}{n} \cdot \sum_{i=1}^n (x_i^\star + s_i^\star)}{\min_{1 \leq i \leq n} (x_i^\star + s_i^\star)}$ among the optimal solution components of $x^\star, s^\star$ influences rPDHG's performance, by deriving a high-probability iteration bound conditioned on unique primal and dual optima that grows linearly in this disparity ratio ϕ . This result yields practical insights for generating challenging test LP instances, which we validate experimentally.

1.2. Related work In addition to the above worst-case analysis of rPDHG (Applegate et al. [8], Xiong [56], Xiong and Freund [59]), Hinder [25] proves that rPDHG has a polynomial iteration bound for total unimodular LPs. Xiong and Freund [58] provide computational guarantees of rPDHG for general conic linear programs based on geometric measures of the primal-dual (sub)level set. Lu and Yang [31, 33] study the vanilla PDHG and the Halpern restarted PDHG using trajectory-based analysis, and demonstrate the two-stage convergence characterized by Hoffman constants of a reduced linear system defined by the limiting

TABLE 1. High-probability iteration bounds for computing an ε -optimal solution using rPDHG (in the case where m is not too close to n).

	Distribution	Stage I	Stage II
Section 3.2 (Theorem 3.1)	sub-Gaussian	$\tilde{O}\left(\frac{n^{2.5}m^{0.5}}{\delta}\right)$	$O\left(\frac{n^{0.5}m^{0.5}}{\delta} \cdot \ln\left(\frac{1}{\varepsilon}\right)\right)$
Section 3.3 (Theorem 3.2)	Gaussian	$\tilde{O}\left(\frac{n^{2.5}}{\delta}\right)$	$O\left(\frac{n^{0.5}m^{0.5}}{\delta} \cdot \ln\left(\frac{1}{\varepsilon}\right)\right)$
Section 7 (Theorem 7.1)	sub-Gaussian	$\tilde{O}\left(\frac{n^{1.5}m^{0.5}\phi}{\delta}\right)$	$O\left(\frac{n^{0.5}m^{0.5}}{\delta} \cdot \ln\left(\frac{1}{\varepsilon}\right)\right)$

Note. These iteration bounds hold with probability at least $1 - \delta$ when δ is not exponentially small in terms of m and $n - m$. In the third row, the probabilities are conditional on the primal and dual problems having unique optimal solutions, and $\phi := \frac{\frac{1}{n} \cdot \sum_{i=1}^n (x_i^* + s_i^*)}{\min_{1 \leq i \leq n} (x_i^* + s_i^*)}$ denotes the disparity ratio among the components of these unique optima. Here $O(\cdot)$ denotes bounds up to distribution-related constant factors, and $\tilde{O}(\cdot)$ hides additional logarithmic factors.

iterate. These deterministic worst-case analyses apply broadly, but these iteration bounds are parameterized by instance-specific constants that can be exponentially large in the worst case. The focus of this paper is complementary: we provide a probabilistic explanation of rPDHG's favorable behavior on a broad class of random instances, offering a further step toward understanding the gap between its worst-case bounds and empirical performance.

Probabilistic analyses of LP algorithms beyond simplex and interior-point methods are very limited. Blum and Dunagan [11] show that the perceptron algorithm, a simple greedy approach, achieves high-probability polynomial smoothed complexity $\tilde{O}(\frac{d^3 m^2}{\delta^2})$ for finding feasible solutions with probability at least $1 - \delta$ for LP instances with m constraints and d variables.

There is some recent work on average-case analyses of first-order methods for unconstrained quadratic optimization problems, such as Paquette et al. [38, 39], Pedregosa and Scieur [40], Scieur and Pedregosa [43]. More recently, Anagnostides and Sandholm [4] study the high-probability performance of some first-order methods for zero-sum matrix games to eliminate condition number dependence.

1.3. Notation For any positive integer n , let $[n]$ denote the set $\{1, 2, \dots, n\}$. For a matrix $A \in \mathbb{R}^{m \times n}$, we denote by $A_{\cdot, i}$ its i -th column ($i \in [n]$) and by $A_{j, \cdot}$ its j -th row ($j \in [m]$). For any subset Θ of $[n]$, A_Θ denotes the submatrix formed by columns indexed by Θ . For a vector v , $\|v\|$ denotes the Euclidean norm, and $\|v\|_1$ denotes the ℓ_1 norm. For a matrix A , $\|A\|$ denotes the spectral norm and $\|A\|_F$ denotes the Frobenius norm. For scalars a and b , we use $a \wedge b$ and $a \vee b$ to denote their minimum and maximum, respectively. We denote the singular values of $A \in \mathbb{R}^{m \times n}$ with $m \leq n$ by $\sigma_1(A), \sigma_2(A), \dots, \sigma_m(A)$, where $\sigma_1(A) \geq \sigma_2(A) \geq \dots \geq \sigma_m(A) \geq 0$. The notation 0_m denotes the m -dimensional zero vector. Throughout this paper, $O(\cdot)$ denotes upper bounds, and $\Omega(\cdot)$ denotes lower bounds, both up to absolute constant factors if not specified. Similarly, $\tilde{O}(\cdot)$ and $\tilde{\Omega}(\cdot)$ denote upper and lower bounds respectively, while hiding additional logarithmic factors. For an event E , we use E^c to denote the complement of E .

2. Restarted PDHG for Linear Programming Recall that this paper studies the standard-form LP problem (1.1). When rows of A are linearly independent, a basis of A contains m columns. We denote by $d := n - m$ the number of nonbasic columns, which is also the difference between n and m . The corresponding dual problem of (1.1) is (1.2). A more symmetric form of the dual problem can be obtained by eliminating y . Let Q be any matrix so that the null space of Q is equal to the image space of $A^\top$, and let $\hat{x}$ be any feasible primal solution. Then (1.2) can be equivalently reformulated in terms of s alone as follows:

$$\min_{s \in \mathbb{R}^n} \hat{x}^\top s \quad \text{s.t. } Qs = Qc, s \geq 0. \quad (2.1)$$

This reformulation of the dual was first proposed by Todd and Ye [52]. The optimal slack s^* of (1.2) is identical to the optimal solution s^* of (2.1). With s^* , any (y, s^*) such that $A^\top y + s^* = c$ forms an optimal solution for (1.2). Similarly, if y^* is an optimal dual solution, then $(y^*, c - A^\top y^*)$ is optimal for (1.2) and $c - A^\top y^*$ is optimal for (2.1).

The optimality conditions state that x^* and s^* are optimal if and only if they are feasible for (1.1) and (2.1) and satisfy the complementary slackness condition $(x^*)^\top s^* = 0$. Since $\hat{x}^\top s = \hat{x}^\top (c - A^\top y) = \hat{x}^\top c - b^\top y$ for any feasible solution $\hat{x}$ of (1.1) and feasible (y, s) of (1.2), the objective vector $\hat{x}$ of (2.1) can be replaced by any primal feasible solution $\hat{x}$ without altering the optimal solutions x^* and s^* . Likewise, x^* and s^* remain optimal if the objective vector c of (1.1) is replaced with any $\hat{c}$ satisfying $Q\hat{c} = Qc$ (namely, $\hat{c} = c + A^\top y_0$ for some $y_0 \in \mathbb{R}^m$). In later sections, we frequently assume c has been replaced by $\bar{c} := c + A^\top \bar{y}$ during presolving, where $\bar{y} := \arg \min_y \|c + A^\top y\|$. This substitution simplifies the analysis without affecting the optimal solutions (x^*, s^*) .

Furthermore, the optimal solutions of primal problem (1.1) and dual problem (1.2) form the saddle point of the Lagrangian $L(x, y)$ that is defined as follows:

$$\min_{x \in \mathbb{R}_+^n} \max_{y \in \mathbb{R}^m} L(x, y) := c^\top x + b^\top y - (Ax)^\top y. \quad (2.2)$$

Any x^* and $(y^*, c - A^\top y^*)$ are optimal solutions of (1.1) and (1.2) if and only if (x^*, y^*) is a saddle point of (2.2), and vice versa. The formulation (2.2) is the problem that the primal-dual hybrid gradient method directly addresses.

2.1. Restarted Primal-dual hybrid gradient method (rPDHG) A single iteration of PDHG, denoted by $(x^+, y^+) \leftarrow \text{ONEPDHG}(x, y)$, is defined as follows:

$$\begin{cases} x^+ \leftarrow P_{\mathbb{R}_+^n}(x - \tau(c - A^\top y)) \\ y^+ \leftarrow y + \sigma(b - A(2x^+ - x)) \end{cases} \quad (2.3)$$

where $P_{\mathbb{R}_+^n}(\cdot)$ denotes the projection onto $\mathbb{R}_+^n$, which means taking the positive part of a vector, and τ and σ are the primal and dual step sizes, respectively. The ONEPDHG iteration involves only two matrix-vector multiplications, one projection onto $\mathbb{R}_+^n$ and several vector-vector products. These operations avoid the computationally expensive matrix factorizations. Therefore, PDHG is well-suited for solving large LP instances and exploiting parallel implementation on modern computational architectures such as GPUs.

Furthermore, Applegate et al. [8] introduce the use of restarts to accelerate the convergence rate of PDHG. Here, the “restarts” mean that the method occasionally restarts from the average iterate of the previous consecutive many iterates. We will refer to this scheme as “rPDHG” for the restarted PDHG. Compared with the vanilla version, rPDHG can achieve linear convergence on LP instances and has shown strong practical performance. The rPDHG, together with some heuristic techniques, is the base method used in most current state-of-the-art first-order LP solvers.

In this paper, we organize PDHG iterations into outer loops, and restart at the averaged iterate of the current inner loop. Following Applegate et al. [8], we use the β -restart criterion based on the *normalized duality gap* (see, e.g., (4a) of Applegate et al. [8]): an outer loop restarts once the normalized duality gap of the averaged iterate decreases to a β fraction of its value at the start of that outer loop. Throughout this paper (both theory and experiments), we take β as a fixed absolute constant (we use $\beta = 1/e$). See Algorithm 1 in Appendix A for the complete algorithm framework considered in this paper. Throughout the remainder of the paper, rPDHG refers to Algorithm 1 in Appendix A.

Other restart triggers have also been used in practice. For example, Applegate et al. [8] discuss fixed-frequency restarts in addition to the normalized duality gap restart, and Lu and Yang [31] consider alternative triggers based on progress measures. These restart schemes have similar theoretical iteration bounds and empirical performance. No single trigger is consistently better than the others. We use the restart scheme based on the normalized duality gap in this paper, as it is widely used in implementations such as Applegate et al. [7], Lu and Yang [32].

Note that rPDHG contains nested loops, but for notational simplicity, we use a single superscript indexed by the number of ONEPDHG steps performed. We denote the primal and dual solutions of rPDHG after k steps of ONEPDHG by x^k and y^k . The corresponding slack $c - A^\top y^k$ is denoted by s^k .

3. High-Probability Computational Guarantees of rPDHG In this section, we describe the probabilistic model used in our analysis and present the polynomial-time complexity of rPDHG when applied to this model in a high-probability sense.

3.1. Probabilistic linear programming model Our probabilistic model builds on the classic model proposed by Todd [51], in which the constraint matrix is drawn from a specific probability distribution, and the right-hand side and the objective vector are computed based on given random primal and dual solutions. Different versions of this model have been extensively studied in the average-case complexity analyses of interior-point methods (see, for example, Anstreicher et al. [5, 6], Huang [26, 27], Ji and Potra [29], Ye [61]).

Compared with other probabilistic models, a significant difference lies in the distribution of the constraint matrix. It is usually assumed that each component of the constraint matrix obeys a Gaussian distribution (see, e.g., Anstreicher et al. [6] and Todd [51]). Our probabilistic model weakens this assumption to the more general case of sub-Gaussian distributions. Below, we provide the definition of sub-Gaussian random variables along with some typical examples (see proofs and additional examples in Wainwright [54]):

DEFINITION 3.1 (SUB-GAUSSIAN RANDOM VARIABLE). A random variable X is sub-Gaussian if there exists $\sigma > 0$ such that $\mathbb{E} \left[e^{\lambda(X - \mathbb{E}[X])} \right] \leq e^{\frac{\sigma^2 \lambda^2}{2}}$ for all $\lambda \in \mathbb{R}$. The parameter σ is referred to as the sub-Gaussian parameter.

EXAMPLE 3.1 (GAUSSIAN RANDOM VARIABLE). Let X be a Gaussian random variable with mean μ and variance σ^2 . Then X is sub-Gaussian with the parameter σ .

EXAMPLE 3.2 (BOUNDED RANDOM VARIABLE). Let X be a random variable supported on the bounded interval $[a, b]$. Then X is sub-Gaussian with a parameter σ that is less than or equal to $\frac{b-a}{2}$.

The family of sub-Gaussian distributions contains many commonly used distributions. In particular, it contains any bounded distribution, such as the Bernoulli distribution and the uniform distribution on a closed interval.

Our probabilistic model involves the notions of sub-Gaussian matrices and nonnegative absolutely continuous sub-Gaussian vectors defined as follows.

DEFINITION 3.2 (SUB-GAUSSIAN MATRIX). A matrix A is a sub-Gaussian matrix if its elements are independent and identically distributed (i.i.d.), each obeying a mean-zero, unit-variance sub-Gaussian distribution. The sub-Gaussian parameter of each component is denoted by $\sigma_{A_{ij}}$, and the sub-Gaussian parameter of A is defined to be $\sigma_A := \max_{ij} \sigma_{A_{ij}}$.

DEFINITION 3.3 (NONNEGATIVE ABSOLUTELY CONTINUOUS SUB-GAUSSIAN VECTOR). A vector u is a nonnegative absolutely continuous sub-Gaussian vector if its components are independent (and potentially different) nonnegative sub-Gaussian random variables whose probability density functions are bounded above by one. We denote the maximum of the means and the maximum of the sub-Gaussian parameters over all components of u by μ_u and σ_u , respectively.

With the above definitions, we now describe the probabilistic model considered in this paper.

DEFINITION 3.4 (PROBABILISTIC MODEL). Instances of the probabilistic model are as follows:

$$\min_{x \in \mathbb{R}^n} \hat{s}^\top x, \quad \text{s.t. } Ax = A\hat{x}, x \geq 0. \quad (3.1)$$

The constraint matrix A is a sub-Gaussian matrix in $\mathbb{R}^{m \times n}$ as defined in Definition 3.2, where $m < n$. The right-hand side and the objective vectors are computed based on A and the primal and dual solutions $\hat{x}$ and $\hat{s}$, where $\hat{x}$ and $\hat{s}$ are generated as follows:

$$\hat{x} = \begin{pmatrix} u^1 \\ 0_d \end{pmatrix} \text{ and } \hat{s} = \begin{pmatrix} 0_m \\ u^2 \end{pmatrix} \quad (3.2)$$

where $u^1 \in \mathbb{R}_+^m$ and $u^2 \in \mathbb{R}_+^d$, and the vector $u := (u^1, u^2)$ is a nonnegative absolutely continuous sub-Gaussian vector in $\mathbb{R}^n$ as defined in Definition 3.3. The random matrix A and the random vector u are independent. (The objective vector $\hat{s}$ can be replaced by any c of the form $c = \hat{s} + A^\top y$ since the optimal solution sets (X^*, S^*) remain unchanged. Optionally, $\hat{s}$ can be replaced by $\bar{c}$, the objective vector with the smallest norm, defined as $\bar{c} := \hat{s} + A^\top \hat{y}$ where $\hat{y} = \arg \min_{y \in \mathbb{R}^m} \|\hat{s} + A^\top y\|$, and so satisfies $A\bar{c} = 0$.)

In the probabilistic model in Definition 3.4, the components of A have unit variance and the probability densities of the components of u are at most one. If these requirements do not hold, they can be satisfied by scalar rescaling of the data instances.

Instances of the probabilistic model have the option of using the objective vector $\bar{c}$ instead of $\hat{s}$. As discussed above, this $\bar{c}$ satisfies $A\bar{c} = 0$ and does not change the optimal solution sets (X^*, S^*) because $\bar{c} = \hat{s} + A^\top \hat{y}$. Keeping $\hat{s}$ would introduce an additional minor term involving the component in the image space of $A^\top$ in the complexity bound (see Remark 3.5 in Xiong and Freund [59]), which makes the bound look unnecessarily complicated. Furthermore, computing $\bar{c}$ (e.g., using the conjugate gradient method to compute $\hat{y}$) is usually much less expensive than solving the LP itself, and $A\bar{c} = 0$ is often assumed in rPDHG analyses such as Xiong [56], Xiong and Freund [58]. To enhance clarity and to align with these previously established results used in our analysis, we presume throughout this paper that the objective vector is set to $\bar{c}$ during a presolving step before applying rPDHG.

When applying rPDHG to instances of this model, we assume that rPDHG regards them as regular standard-form LP instances without knowing any prior information about the distribution of the input data.

The model in Definition 3.4 is analogous to the model of Anstreicher et al. [6], specifically the TDMV1 model. It is an important case of Todd [51]’s classic probabilistic model and has been studied by Anstreicher et al. [5], Ye [61, 62] and others for analyzing the average-case performance of interior-point methods. Our model and those of Anstreicher et al. [6], Todd [51] all assume that the constraint matrix and a pair of primal-dual feasible (and optimal) solutions are sampled from specific probability distributions, after which the right-hand side b and the objective vector c of the random LP instance are computed by $b = A\hat{x}$ and $c = \hat{s}$. One advantage of using this model is that any instance of this model is feasible and has a bounded optimal solution that is randomly distributed.

Distribution of the constraint matrix. Our model allows the constraint matrix to follow a general sub-Gaussian distribution, whereas the probabilistic models most commonly studied in the LP literature assume Gaussian or other symmetry-based distributions. Generalizing the distributional assumptions is important because it broadens the range of random LP instances for which the analysis can offer insight into rPDHG’s observed behavior. For example, in the probabilistic analysis of simplex methods, Borgwardt [14], Smale [45] assume that the polytopes come from a special spherically symmetric distribution. Later, Adler et al. [1], Adler and Megiddo [3], Todd [50] assume that the constraint matrix is drawn from a sign-invariant distribution. Spielman and Teng [48] study Gaussian perturbations of input data to the simplex method. On the other hand, most probabilistic analyses of interior-point methods have been built on Todd [51]’s probabilistic model, which assumes that the constraint matrix is a Gaussian matrix. Examples include Anstreicher et al. [5, 6], Mizuno et al. [37], Ye [61].

We will show later in Section 3.2 that, for our model with general sub-Gaussian distributions, rPDHG has polynomial-time complexity in a high-probability sense. Given the popularity of Gaussian distributions in the literature on probabilistic analysis, later in Section 3.3, we will show that rPDHG has even better high-probability polynomial-time complexity when the constraint matrix is Gaussian.

Distribution of solutions $\hat{x}$ and $\hat{s}$. In our probabilistic model, $\hat{x}$ and $\hat{s}$ are optimal primal-dual solutions because they are feasible primal and dual solutions and satisfy the complementary slackness condition. Because rPDHG is invariant under permutations of variables, without loss of generality, in the model we assign the indices of the possible nonzero components of $\hat{x}$ and $\hat{s}$ to $\{1, \dots, m\}$ and $\{m+1, \dots, n\}$. We use B and N to denote the submatrices of A corresponding to the column indices $\{1, \dots, m\}$ and $\{m+1, \dots, n\}$, respectively, so that A is represented as (B, N) . When B is full-rank and $(\hat{x}, \hat{s})$ satisfies the strict complementary slackness condition, the optimal solution sets X^* , $\mathcal{Y}^*$ and S^* are all singletons. It is formally stated as follows (proof in Appendix B):

LEMMA 3.1. *For an instance of the probabilistic model, X^* , $\mathcal{Y}^*$ and S^* are all singletons with $X^* = \{\hat{x}\}$ and $S^* = \{\hat{s}\}$ if and only if B is full-rank and $(\hat{x}, \hat{s})$ satisfies the strict complementary slackness condition, namely, $u = (u^1, u^2) > 0$.*

Furthermore, once m is sufficiently large, it is highly probable that the instance of the probabilistic model has unique primal and dual optimal solutions.

LEMMA 3.2. *There exists a constant $\nu > 0$ that depends only (and at most polynomially) on the sub-Gaussian parameter of A , such that the probability of B being full-rank is at least $1 - e^{-\nu m}$.*

The proof of Lemma 3.2 is provided in Section 4.1.

The distributions of $\hat{x}$ and $\hat{s}$ also differ between our probabilistic model and that of Anstreicher et al. [6]. On the one hand, the model of Anstreicher et al. [6] permits varying numbers of nonzeros in $\hat{x}$ and $\hat{s}$ while maintaining strict complementarity. This may result in multiple optimal solutions with high probability. On the other hand, the components of u in their model are i.i.d. from a folded Gaussian distribution (the distribution of the absolute value of a Gaussian random variable). In contrast, our model fixes the numbers of nonzeros in $\hat{x}$ and $\hat{s}$ to m and d , respectively, while allowing the components of u to follow potentially different sub-Gaussian distributions.

We acknowledge that for a very large number of LP instances occurring in practice, the optimal solution is not unique. But the unique optimum property is generic in theory, especially in probabilistic models, because most randomly generated LP problems (unless specially designed) are nondegenerate (see, for example, Adler et al. [1], Borgwardt [15], Spielman and Teng [47], Todd [50], Ye [61]).

The flexibility in the distribution of u in our model provides tools for various aims of analyses. When u is a random vector of the folded Gaussian distribution, the probabilistic model is for the average-case analysis. When u is a given fixed vector with random perturbations, the probabilistic model is for smoothed analysis on the dependence of the optimal solution. Furthermore, we will show later in Section 7 that the performance of rPDHG is highly influenced by the distribution of the optimal solution, and the unique optimum property of our model is helpful in generating artificial LP instances of various difficulty levels.

It should be mentioned that Todd [51] also provides a model that allows control over the degree of degeneracy in $\hat{x}$ and $\hat{s}$ (TDMV2 of Anstreicher et al. [6]). But Todd and Anstreicher et al. [6] later pointed out a subtle error in the analysis of this model. Due to this error, several papers, including Anstreicher et al. [5] and Ye [62], that had claimed to analyze the average-case performance of interior-point methods on this model actually applied only to the TDMV1 model of Anstreicher et al. [6]. This subtle error also affected several papers that studied the version of Todd [51]’s probabilistic model with $\hat{x} = e$ and $\hat{s} = e$. These errors were later corrected by Huang [26] and Ji and Potra [29].

Next, in Section 3.2, we present our main results on the high-probability polynomial-time complexity of rPDHG for instances of the probabilistic model with sub-Gaussian input data. In Section 3.3, we then consider instances with Gaussian input data.

3.2. High-probability performance guarantees for LP instances with sub-Gaussian input data

In this subsection, we analyze the performance of rPDHG on instances of the probabilistic model defined in Definition 3.4. Our focus is on the dependence of the iteration bounds on the dimensions m and n . The bounds may also contain some constants that depend only on the parameters of the distributions of A and u , namely σ_A , μ_u and σ_u in Definitions 3.2 and 3.3. Below we define an ε -optimal solution for the LP primal and dual problems (1.1) and (1.2).

DEFINITION 3.5 (ε -OPTIMAL SOLUTIONS). A primal-dual pair (x, s) is an ε -optimal solution if there exists a primal-dual optimal pair (x^*, s^*) such that $\|(x, s) - (x^*, s^*)\| \leq \varepsilon$.

It is often observed in practice (and shown in theory) that rPDHG exhibits a “two-stage phenomenon,” in which its iterations can be divided into Stage I and Stage II. In Stage I, the iterates eventually reach the point where their positive components correspond to the optimal basis. In Stage II, rPDHG exhibits fast local convergence to an optimal solution (see, e.g., Lu and Yang [31, 33], Xiong [56]). In other words, there exists T_{basis} such that, for all $t \geq T_{basis}$, $x_i^t > 0$ if and only if i is in the optimal basis, after which rPDHG exhibits faster local linear convergence to an optimum. We let T_{local} denote the number of additional iterations beyond T_{basis} required to compute an ε -optimal solution (x^t, s^t) . The following theorem presents our high-probability bounds on T_{basis} and T_{local} .

THEOREM 3.1. *Suppose that rPDHG is applied to an instance of the probabilistic model in Definition 3.4 (with objective vector $\bar{c}$). Let $c_0 := \mu_u + 2\sigma_u$. There exist constants $C_0, C_1, C_2 > 0$ that depend only (and at most polynomially) on σ_A , for which the following high-probability iteration bounds hold:*

1. (Optimal basis identification)

$$\Pr \left[T_{basis} \leq \frac{m^{0.5}n^{3.5}}{d+1} \cdot \frac{c_0 C_1 \cdot \ln(3/\delta)}{\delta} \cdot \ln \left(\frac{m^{0.5}n^{3.5}}{d+1} \cdot \frac{c_0 C_1 \cdot \ln(3/\delta)}{\delta} \right) \right] \geq 1 - \delta - 6 \left(\frac{1}{e^{C_0}} \right)^m - \left(\frac{1}{2} \right)^{d+1} \quad (3.3)$$

for any $\delta \in (0, 1]$.

2. (Fast local convergence) Let $\varepsilon > 0$ be any given tolerance.

$$\Pr \left[T_{local} \leq m^{0.5}n^{0.5} \cdot \frac{C_2}{\delta} \cdot \max \left\{ 0, \ln \left(\frac{c_0}{\varepsilon} \right) \right\} \right] \geq 1 - \delta - 3 \left(\frac{1}{e^{C_0}} \right)^m \quad (3.4)$$

for any $\delta \in (0, 1]$.

In the above theorem, we have divided the rPDHG iterations into two stages and presented probabilistic bounds for them separately. Technically speaking, once the iterates settle on the optimal basis, the optimal solution (x^*, s^*) could be directly computed by solving two linear systems—one for the primal basic system and the other for the dual basic system. (This is common in finite termination methods for interior-point methods, such as Ye [60], where the critical effort lies in computing projections.) Indeed, one could compute optimal solutions in finite time using rPDHG by first running rPDHG for T_{basis} ONEPDHG iterations and then solving the two associated linear systems. However, this is not a practical approach for at least two reasons. First, determining whether (x^t, s^t) has already settled on the optimal basis can be difficult. Second, rPDHG automatically exhibits fast local convergence after identifying the optimal basis, so computing an ε -optimal solution requires at most $T_{basis} + T_{local}$ iterations. The Stage-II bound in (3.4) depends only logarithmically on the target accuracy and on c_0 . This result does not conflict with the global linear convergence guarantee established by Applegate et al. [8]. Actually, rPDHG converges linearly in Stage I. We discuss this further in Remark 4.2 in Section 4.

The inequalities (3.3) and (3.4) are high-probability upper bounds for T_{basis} and T_{local} . The right-hand sides of (3.3) and (3.4) both contain $1 - \delta$ as well as some additional terms that decrease exponentially in the dimensions m and d . In other words, if m and d are sufficiently large – essentially larger than $O(\ln \frac{1}{\delta})$ – these additional terms are negligible, and so (3.3) and (3.4) describe the upper bounds on T_{basis} and T_{local} with a probability that is roughly equal to $1 - \delta$. Similar requirements on the dimensions being sufficiently large are common in other probabilistic analyses of linear programming, such as Huang [27], Mizuno et al. [37], Ye [61] for interior-point methods, and Adler et al. [1], Borgwardt [12], Shamir [44] for simplex methods.

The constant c_0 depends only on μ_u and σ_u , whereas C_0, C_1 , and C_2 depend only on σ_A . See Remark 4.1 in Section 4 for further discussion of these constants. The term $\frac{m^{0.5}n^{3.5}}{d+1}$ is at most as large as $\frac{m^{0.5}n^{3.5}}{2}$. When d (recall $d := n - m$) is not too small compared to n in the sense that $d + 1 \geq n/C_3$ for some absolute constant $C_3 > 1$, then the term $\frac{m^{0.5}n^{3.5}}{d+1}$ is $O(m^{0.5}n^{2.5})$.

The corollary below summarizes the high-probability iteration bounds when (i) m is not too close to n , and (ii) δ is not exponentially small in d and m :

COROLLARY 3.1. *Let $C_3 > 1$ be a given absolute constant. Suppose that rPDHG is applied to an instance of the probabilistic model (with objective vector $\bar{c}$). When d satisfies $\frac{n}{d+1} \leq C_3$, it holds with probability at least $1 - \delta$ that rPDHG computes an ε -optimal solution within at most*

$$\tilde{O} \left(\frac{n^{2.5}m^{0.5}}{\delta} \right) + O \left(\frac{n^{0.5}m^{0.5}}{\delta} \cdot \ln \left(\frac{1}{\varepsilon} \right) \right)$$

iterations for all $\varepsilon \in (0, e^{-1}]$ and any $\delta \in (0, 1]$ satisfying $\delta > 11 \cdot \max \left\{ \left(\frac{1}{e^{C_0}} \right)^m, \left(\frac{1}{2} \right)^{d+1} \right\}$. Moreover, T_{basis} is at most $\tilde{O} \left(\frac{n^{2.5}m^{0.5}}{\delta} \right)$. Here $\tilde{O}(\cdot)$ omits factors of an absolute constant, c_0, C_1, C_3 and logarithmic terms that involve c_0, C_1, C_3, m, n and $\frac{1}{\delta}$, and $O(\cdot)$ additionally omits factors of C_2 and $\ln c_0$.

Corollary 3.1 shows that under a mild assumption on δ and on the dimensions m and d (the interval for δ is nonempty only when $m > \frac{\ln(11)}{C_0}$ and $d > \log_2(11) - 1 \approx 2.4594$), rPDHG settles on the optimal basis within $\tilde{O}\left(\frac{n^{2.5}m^{0.5}}{\delta}\right)$ iterations, and it computes an ε -optimal solution in an additional $O\left(\frac{n^{0.5}m^{0.5}}{\delta} \cdot \ln \frac{1}{\varepsilon}\right)$ iterations, with probability at least $1 - \delta$. To the best of our knowledge, the above result is the first high-probability polynomial-time complexity bound for rPDHG under a probabilistic input model. It provides a probabilistic explanation for why rPDHG can avoid its exponentially poor worst-case complexity on a broad class of random LP instances, thereby making progress toward understanding the gap between its worst-case theory and observed performance. Moreover, it shows that the dependence of the high-probability complexity on n in Stage I (settling on the optimal basis) is higher than in Stage II (local convergence). This observation also aligns with the two-stage phenomenon reported in worst-case complexity and practical experimental results (see Xiong [56] and Lu and Yang [31, 33]).

Furthermore, to the best of our knowledge, this is the first probabilistic iteration complexity bound for an LP algorithm that allows the constraint matrix to have general, not necessarily Gaussian, sub-Gaussian entries. The Gaussian matrix is relatively easy to analyze because it has nice symmetry in the sense of geometry, its range is the unique orthogonally invariant distribution, and the corresponding LP instance is nondegenerate with probability 1 (see Todd [51]). The more general sub-Gaussian random matrix model is harder to analyze than it may look. Even the behavior of the extreme singular values of sub-Gaussian matrices was not well understood until about 15 years ago (see the invited lecture at ICM 2010 by Rudelson and Vershynin [42]). Our analysis approach relies on the nonasymptotic result of Rudelson and Vershynin [41] on the smallest singular value of a sub-Gaussian matrix. Proofs of our results above are presented in Section 4.

It should be noted that real-world LP instances often differ from our model by having multiple optimal solutions or input data that are not well-approximated by sub-Gaussian distributions (e.g., heavy-tailed or highly structured data). In such cases, we do not yet have high-probability performance guarantees. The two-stage characterization used in Theorem 3.1 is not well-defined when the optimal basis is not unique, and our probabilistic analysis relies on the accessible iteration bound of Xiong [56], which is for LP instances with unique optimal solutions. A possible extension to general degenerate instances is to start from the global convergence analyses that apply to general LPs, such as [8] and [59], and then study the typical size of the condition measures used in these results under the probabilistic model. Moreover, sub-Gaussianity enables us to invoke nonasymptotic random matrix singular value estimates. Extending the analysis to heavy-tailed or structured inputs would require different tools (e.g., Cook [16], Dumitriu and Zhu [18]). These are left for future investigation.

3.3. High-probability performance guarantees for LP instances with Gaussian input data In this subsection, we analyze the performance of rPDHG under the probabilistic model wherein the constraint matrix A is a Gaussian matrix, which we now define.

DEFINITION 3.6 (GAUSSIAN MATRIX). A matrix A is called a Gaussian matrix if its elements are i.i.d., each obeying the mean-zero Gaussian distribution with unit variance.

A Gaussian matrix A is a special case of a sub-Gaussian matrix in Definition 3.2, with its sub-Gaussian parameter σ_A equal to 1 (see Wainwright [54]). In this section, we show that the high-probability iteration bound of rPDHG, particularly for Stage I (identifying the optimal basis), can be further improved in this classical Gaussian setting. Similar to Theorem 3.1, the theorem below presents high-probability bounds on T_{basis} and T_{local} when the matrix A is Gaussian.

THEOREM 3.2. *Suppose that rPDHG is applied to an instance of the probabilistic model in Definition 3.4 (with objective vector $\bar{c}$) and suppose that the constraint matrix is a Gaussian matrix. Let $c_0 := \mu_u + 2\sigma_u$. There exist absolute constants $C_0, C_1, C_2 > 0$ for which the following high-probability iteration bounds hold:*

1. (Optimal basis identification)

$$\Pr \left[T_{basis} \leq \frac{n^{3.5}}{d+1} \cdot \frac{c_0 C_1 \cdot \ln(6/\delta) \sqrt{\ln(4n/\delta)}}{\delta} \cdot \ln \left(\frac{n^{3.5}}{d+1} \cdot \frac{c_0 C_1 \cdot \ln(6/\delta) \sqrt{\ln(4n/\delta)}}{\delta} \right) \right] \geq 1 - \delta - (n+5) \left(\frac{1}{e^{C_0}} \right)^{m \wedge d} \quad (3.5)$$

for any $\delta \in (0, 1]$.

2. (Fast local convergence) Let $\varepsilon > 0$ be any given tolerance.

$$\Pr \left[T_{local} \leq m^{0.5} n^{0.5} \cdot \frac{C_2}{\delta} \cdot \max \left\{ 0, \ln \left(\frac{C_0}{\varepsilon} \right) \right\} \right] \geq 1 - \delta - 3 \left(\frac{1}{e^{C_0}} \right)^m \quad (3.6)$$

for any $\delta \in (0, 1]$.

The inequalities (3.5) and (3.6) are high-probability upper bounds on T_{basis} and T_{local} . When d and m are sufficiently large relative to $\frac{n}{\delta}$, specifically when $m \wedge d \geq \Omega(\ln \frac{n}{\delta})$, the right-hand side of each inequality approaches $1 - \delta$. However, in the extreme case wherein m or d is too small compared with n , these bounds become trivial for all $\delta \in (0, 1]$ since their right-hand sides become nonpositive.

Notice in Theorem 3.2 that C_0, C_1 , and C_2 are absolute constants. Indeed, these constants otherwise depend only on σ_A , and $\sigma_A = 1$ for a Gaussian matrix. See Remark 5.1 in Section 5 for further discussion of these constants. The term $\frac{n^{3.5}}{d+1}$ is at most as large as $\frac{n^{3.5}}{2}$. When d is not too small compared to n , $\frac{n^{3.5}}{d+1}$ is $O(n^{2.5})$. Similarly to Corollary 3.1, the corollary below summarizes the high-probability iteration bounds when (i) m is not too close to n , and (ii) δ is not exponentially small:

COROLLARY 3.2. *Let $C_3 > 0$ be any given absolute constant. Suppose that rPDHG is applied to an instance of the probabilistic model (with objective vector $\bar{c}$), and suppose that the constraint matrix A is a Gaussian matrix. If d satisfies $\frac{n}{d+1} \leq C_3$, then it holds with probability at least $1 - \delta$ that rPDHG computes an ε -optimal solution within at most*

$$\tilde{O} \left(\frac{n^{2.5}}{\delta} \right) + O \left(\frac{n^{0.5} m^{0.5}}{\delta} \cdot \ln \left(\frac{1}{\varepsilon} \right) \right)$$

iterations, for all $\varepsilon \in (0, e^{-1}]$ and any $\delta \in (0, 1]$ satisfying $\delta > 2(n+5)e^{-C_0(m \wedge d)}$. Moreover, T_{basis} is at most $\tilde{O} \left(\frac{n^{2.5}}{\delta} \right)$. Here $\tilde{O}(\cdot)$ omits factors of constants $c_0 := \mu_u + 2\sigma_u$, an absolute constant and logarithmic terms that involve an absolute constant, m, n and $\frac{1}{\delta}$, and the $O(\cdot)$ additionally omits factors of C_2 and $\ln c_0$.

Corollary 3.2 shows that under a mild assumption on δ and the dimensions m and $d = n - m$ (the interval for δ is nonempty only when $m \wedge d > \frac{1}{C_0} \ln(2(n+5))$), rPDHG settles on the optimal basis within $\tilde{O} \left(\frac{n^{2.5}}{\delta} \right)$ iterations, and it computes an ε -optimal solution in an additional $O \left(\frac{n^{0.5} m^{0.5}}{\delta} \cdot \ln \frac{1}{\varepsilon} \right)$ iterations, with probability at least $1 - \delta$. Comparing Corollary 3.2 with Corollary 3.1 (for sub-Gaussian matrices), the dependence on $m^{0.5}$ is eliminated in the Stage-I iteration bound while the Stage-II iteration bound remains almost identical.

Now we compare this high-probability iteration bound with the probabilistic bounds of interior-point methods and simplex methods. In general, the interior-point method has far less dependence on the dimension in the number of iterations, but the per-iteration complexity is significantly higher than that of rPDHG because it needs to solve a normal equation of a dense normal matrix. The model of Anstreicher et al. [6] is the most similar one to ours, and Anstreicher et al. [6] prove that the expected number of iterations of an interior-point method is at most $O(n \cdot \ln(n))$. They also remark that their proof can be easily modified for other interior-point methods and derive an $O(n^2 \ln(n))$ expected iteration bound for the interior-point methods of Wright [55], Zhang [64], Zhang and Zhang [65] and others. Huang [27] shows that the expected and high-probability numbers of iterations are both bounded above by $O(n^{1.5})$ in another model of Todd [51]. Ye [61] proves that the number of iterations is at most $O(\sqrt{n} \cdot \ln \frac{1}{\varepsilon})$ with high probability under a different model.

As for simplex methods, direct comparison of the probabilistic bounds between rPDHG and simplex methods is challenging because different models and forms of LP are used in probabilistic analyses. Nevertheless, our high-probability iteration bound for rPDHG demonstrates comparable polynomial-time complexity to those established for simplex methods, with the additional advantage that rPDHG requires only two matrix-vector multiplications per iteration. To make the comparison as fair as possible, we consider problems where the number of constraints is of the same order as the number of variables, which we refer to

as the dimension of the problem. Adler et al. [2], Adler and Megiddo [3], Todd [50] prove that the expected number of steps of several equivalent simplex methods on a certain probabilistic model is bounded by a quadratic function of the dimension. Adler and Megiddo [3] also prove that the dependence is tight. Spielman and Teng [48] use the smoothed analysis framework to study models with Gaussian perturbations and derive a polynomial bound on the expected number of simplex pivots. Dadush and Huiberts [17] significantly simplify the analysis and establish a tighter polynomial bound with dimension exponent at most 3.5.

We note that Theorem 3.2 does not imply a bound on the expected number of iterations of rPDHG. Indeed, the tails of (3.5) and (3.6) do not decay to zero. In probabilistic analyses of linear programming, high-probability complexity bounds are also established by Blum and Dunagan [11], Ye [61] and others. We are unable to prove bounds on the expected number of iterations for rPDHG, which is in contrast to the case of interior-point methods and simplex methods where such results are proven in Anstreicher et al. [6], Dadush and Huiberts [17], Todd [50] and others.

4. Proof of Theorem 3.1 In this section, we prove Theorem 3.1. Section 4.1 introduces several useful helper lemmas frequently used in the proofs, including concentration inequalities for sub-Gaussian random variables, bounds on the extreme singular values of random matrices, and a tail bound for the product of two independent random variables. Section 4.2 recalls the worst-case iteration bound of rPDHG. Section 4.3 contains the detailed proof of Theorem 3.1.

4.1. Lemmas of random variables and random matrices The Hoeffding bound provides a tail bound for the sum of independent sub-Gaussian random variables (see Wainwright [54]).

LEMMA 4.1 (Hoeffding bound). *Suppose that the sub-Gaussian variables $\{X_i\}_{i=1}^n$ are independent, and each X_i has the sub-Gaussian parameter σ_i . Then for all $t \geq 0$, we have $\Pr \left[\sum_{i=1}^n (X_i - \mathbb{E}[X_i]) \geq t \right] \leq \exp \left(-\frac{t^2}{2 \sum_{i=1}^n \sigma_i^2} \right)$.*

For a random matrix as defined in Definition 3.2, the following result from Rudelson and Vershynin [41] provides a tail bound for the smallest singular value of the random matrix.

LEMMA 4.2 (Theorem 1.1 of Rudelson and Vershynin [41]). *Let A be a random matrix as defined in Definition 3.2, where $A \in \mathbb{R}^{m \times n}$ and $n \geq m$. For every $\varepsilon > 0$, we have*

$$\Pr \left[\sigma_m(A) \leq \varepsilon \left(\sqrt{n} - \sqrt{m-1} \right) \right] \leq (C_{rv1} \varepsilon)^{d+1} + e^{-C_{rv2} n} \quad (4.1)$$

where $C_{rv1}, C_{rv2} > 0$ depend only (and at most polynomially) on the sub-Gaussian parameter σ_A .

It is worth noting that the smallest singular value may not always be strictly greater than zero, which implies that the matrix A may not be full-rank. However, the above lemma shows that as n increases, the probability of A being full-rank becomes very high. A direct application of this result is the proof of Lemma 3.2.

Proof of Lemma 3.2. According to Lemma 3.1, the probability of Condition 1 being true is equal to $\Pr [\sigma_m(B) > 0]$. Since B is a random matrix as defined in Definition 3.2 with $B \in \mathbb{R}^{m \times m}$, Lemma 4.2 implies that $\Pr [\sigma_m(B) = 0] = \Pr [\sigma_m(B) \leq 0] \leq \lim_{\varepsilon \downarrow 0} \Pr \left[\sigma_m(B) \leq \varepsilon (\sqrt{m} - \sqrt{m-1}) \right] \leq \lim_{\varepsilon \downarrow 0} \left((C_{rv1} \varepsilon)^1 + e^{-C_{rv2} m} \right) = e^{-C_{rv2} m}$. This completes the proof. $\square$

The largest singular value of random matrices is known to be upper bounded by $O(\sqrt{n})$ with high probability. It is formally stated in the following lemma.

LEMMA 4.3. *Let A be a random matrix as defined in Definition 3.2, where $A \in \mathbb{R}^{m \times n}$ and $n \geq m$. Then,*

$$\Pr [\sigma_1(A) \geq 10\sigma_A \sqrt{n}] \leq e^{-6n}.$$

Next, we prove a tail bound of the product of two independent random variables with known heavy tails.

LEMMA 4.4. *Let X and Y be two independent nonnegative random variables, and suppose that there exist $C_1, C_2, \delta_1, \delta_2 > 0$ such that for all $T > 0$:*

$$\Pr[X \geq T] \leq \frac{C_1}{T} + \delta_1 \quad \text{and} \quad \Pr[Y \geq T] \leq \frac{C_2}{T} + \delta_2. \quad (4.2)$$

Then for any $\delta \in (0, 1]$, the following inequality holds:

$$\Pr\left[XY \geq \frac{6C_1C_2 \cdot \ln(3/\delta)}{\delta}\right] \leq \delta + \delta_1 + 2\delta_2. \quad (4.3)$$

Proofs of Lemmas 4.3 and 4.4 are provided in Appendix C.

4.2. Worst-case iteration bounds of rPDHG According to Lemmas 3.1 and 3.2, for instances of the probabilistic model, the probability of having full-rank B and unique optimal solution $\mathcal{X}^* = \{\hat{x}\}$ and $\mathcal{S}^* = \{\hat{s}\}$ is at least $1 - e^{-\nu m}$. It happens with high probability when m is sufficiently large. Therefore, in this subsection, we recall a theoretical iteration bound under the following condition of a unique optimum.

CONDITION 1. *The problem (1.1) has a unique optimal solution x^* , and the dual problem (1.2) has a unique optimal solution (y^*, s^*) .*

Recently, Xiong [56] proves an accessible iteration bound for rPDHG applied to LPs under Condition 1. This new iteration bound is in closed form of the optimal solution and optimal basis. Suppose that the optimal basis of x^* is $\{1, 2, \dots, m\}$. We still let B and N denote the submatrices consisting of the columns indexed by $\{1, 2, \dots, m\}$ and $\{m+1, m+2, \dots, n\}$. By Lemma 3.1, B must be invertible. Then the iteration bound relies on the key quantity Φ defined as follows:

$$\Phi := (\|x^* + s^*\|_1) \cdot \max \left\{ \max_{1 \leq j \leq d} \frac{\sqrt{\|(B^{-1}N)_{\cdot,j}\|^2 + 1}}{s_{m+j}^*}, \max_{1 \leq i \leq m} \frac{\sqrt{\|(B^{-1}N)_{i,\cdot}\|^2 + 1}}{x_i^*} \right\}. \quad (4.4)$$

If Condition 1 does not hold, for notational simplicity, we let $\Phi := \infty$. Then we have the following worst-case iteration bounds for T_{basis} and T_{local} :

THEOREM 4.1 (Theorem 4.1 of Xiong [56]). *Suppose that $Ac = 0$ and rPDHG is applied to solve the LP instance (1.1). The following iteration bounds hold:*

1. (Optimal basis identification) *There exists an absolute constant $\check{c}_1 > 1$ such that:*

$$T_{\text{basis}} \leq \check{c}_1 \cdot \kappa \Phi \cdot \ln(\kappa \Phi). \quad (4.5)$$

2. (Fast local convergence) *There exists an absolute constant $\check{c}_2 > 0$ such that:*

$$T_{\text{local}} \leq \check{c}_2 \cdot \|B^{-1}\| \|A\| \cdot \max \left\{ 0, \ln \left(\frac{\min_{1 \leq i \leq n} \{x_i^* + s_i^*\}}{\varepsilon} \right) \right\}. \quad (4.6)$$

Besides the “complex” expression of Φ in (4.4), Φ has the following simpler upper bound:

LEMMA 4.5 (Proposition 3.1 of Xiong [56]). *When Condition 1 holds, $\Phi \leq \frac{\|x^* + s^*\|_1}{\min_{1 \leq i \leq n} \{x_i^* + s_i^*\}} \cdot \|B^{-1}A\|$.*

Furthermore, it is actually indicated by Xiong [56] that Φ is equal to $\frac{\|x^* + s^*\|_1}{\zeta_p \wedge \zeta_d}$ in which ζ_p and ζ_d are equivalent to three types of condition measures for primal and dual problems respectively. They are (i) stability under data perturbations, (ii) proximity to multiple optima, and (iii) the LP sharpness of the instance.

4.3. Proof of Theorem 3.1 In this section, we prove Theorem 3.1. Let B be the $m \times m$ submatrix formed by the first m columns of A . We denote the condition number of A by κ , defined as $\kappa := \frac{\sigma_1(A)}{\sigma_m(A)}$. In addition, we let φ denote the quantity defined as follows:

$$\varphi := \frac{\|\hat{x} + \hat{s}\|_1}{\min_{1 \leq i \leq n} (\hat{x}_i + \hat{s}_i)}. \quad (4.7)$$

We set $\varphi := \infty$ when $\min_{1 \leq i \leq n} (\hat{x}_i + \hat{s}_i) = 0$, $\kappa := \infty$ when $\sigma_m(A) = 0$, and $\|B^{-1}\| := \infty$ when B lacks full rank.

We then establish the following upper bound for $\kappa\Phi$:

LEMMA 4.6. *For the random LP defined in Definition 3.4, if $\hat{x}$ and $\hat{s}$ satisfy the strict complementary slackness condition, the following inequality holds:*

$$\kappa\Phi \leq \kappa \cdot \|B^{-1}\| \|A\| \cdot \varphi. \quad (4.8)$$

Proof. When B is full-rank, Lemma 3.1 implies that Condition 1 holds with $X^* = \{\hat{x}\}$ and $S^* = \{\hat{s}\}$. In this case, by Lemma 4.5, (4.8) holds. When B is not full-rank, Lemma 3.1 implies that Condition 1 does not hold and by definition, $\Phi = \infty$. Simultaneously, $\|B^{-1}\| = \infty$ so inequality (4.8) still holds in this case. $\square$

Roadmap of the proof of Theorem 3.1. The above Lemma 4.6 indicates that in order to analyze the tail behavior of $\kappa\Phi$, we can study the tails of its components: κ , $\|B^{-1}\| \|A\|$ and φ . Therefore, below we study them in Steps 1 to 3, after which we will combine these results to derive a tail bound for $\kappa \cdot \|B^{-1}\| \|A\| \cdot \varphi$ (Step 4) and complete the proof of Theorem 3.1 and Corollary 3.1 (Step 5).

Step 1. Tail bound of κ . Before showing the result, we define some constants that will be frequently used in the following analysis:

$$c_1 := C_{rv2} \wedge 6, \quad c_2 := 20\sigma_A C_{rv1}, \quad c_3 := 2c_2^2. \quad (4.9)$$

Note that here c_1, c_2 and c_3 depend only on σ_A , the sub-Gaussian parameter of A , and are independent of the distribution of $(\hat{x}, \hat{s})$ and the values of m and n , because C_{rv1} and C_{rv2} are the constants in Lemma 4.2 that depend only on σ_A .

LEMMA 4.7. *The following inequality holds for κ :*

$$\Pr \left[\kappa \geq 2c_2 \cdot \frac{n}{d+1} \right] \leq \left(\frac{1}{2} \right)^{d+1} + 2e^{-c_1 n} \quad (4.10)$$

Proof. We apply Lemma 4.2 and Lemma 4.3 to bound the largest and smallest singular values of A . For any $\varepsilon > 0$, let E_1 denote the event $\sigma_m(A) > \varepsilon(\sqrt{n} - \sqrt{m-1})$ and let E_2 denote the event $\sigma_1(A) < 10\sigma_A\sqrt{n}$. In the event $E_1 \cap E_2$, it holds that $\kappa = \frac{\sigma_1(A)}{\sigma_m(A)} < \frac{10\sigma_A\sqrt{n}}{\varepsilon(\sqrt{n} - \sqrt{m-1})} = \frac{10\sigma_A}{\varepsilon} \cdot \frac{\sqrt{n}(\sqrt{n} + \sqrt{m-1})}{d+1} \leq \frac{20\sigma_A}{\varepsilon} \cdot \frac{n}{d+1}$, and thus

$$\Pr \left[\kappa \geq \frac{20\sigma_A}{\varepsilon} \cdot \frac{n}{d+1} \right] \leq 1 - \Pr[E_1 \cap E_2] \leq \Pr[E_1^c] + \Pr[E_2^c] \leq (C_{rv1}\varepsilon)^{d+1} + e^{-C_{rv2}n} + e^{-6n} \quad (4.11)$$

for any $\varepsilon > 0$. Here the last inequality is due to Lemmas 4.2 and 4.3. Therefore, setting $\varepsilon = \frac{1}{2C_{rv1}}$ yields the desired inequality (4.10) for c_1 and c_2 defined in (4.9). $\square$

Step 2. Tail bound of $\|B^{-1}\| \|A\|$. Next, we analyze the tail bound of $\|B^{-1}\| \|A\|$.

LEMMA 4.8. *For any $t > 0$, the following inequality holds:*

$$\Pr \left[\|B^{-1}\| \|A\| \geq t \right] \leq \frac{c_2 \sqrt{mn}}{t} + 2e^{-c_1 m}. \quad (4.12)$$

Proof. The proof follows a similar structure to that of Lemma 4.7. We use Lemmas 4.2 and 4.3 to bound the largest singular value of A (equal to $\|A\|$) and the smallest singular value of B (equal to $1/\|B^{-1}\|$). For any $\varepsilon > 0$, let E_1 denote the event $\sigma_m(B) > \varepsilon(\sqrt{m} - \sqrt{m-1})$ and let E_2 denote the event $\sigma_1(A) < 10\sigma_A\sqrt{n}$. In the event $E_1 \cap E_2$, it holds that $\|B^{-1}\| \|A\| = \frac{\sigma_1(A)}{\sigma_m(B)} < \frac{10\sigma_A\sqrt{n}}{\varepsilon(\sqrt{m}-\sqrt{m-1})} = \frac{10\sigma_A}{\varepsilon} \cdot \sqrt{n}(\sqrt{m} + \sqrt{m-1}) \leq \frac{20\sigma_A}{\varepsilon} \cdot \sqrt{mn}$. and thus

$$\begin{aligned} \Pr \left[\|B^{-1}\| \|A\| \geq \frac{20\sigma_A}{\varepsilon} \cdot \sqrt{mn} \right] &\leq 1 - \Pr[E_1 \cap E_2] \leq \Pr[E_1^c] + \Pr[E_2^c] \\ &\leq C_{rv1}\varepsilon + e^{-C_{rv2}m} + e^{-6n} \leq C_{rv1}\varepsilon + e^{-C_{rv2}m} + e^{-6m} \end{aligned} \quad (4.13)$$

where the third inequality uses Lemmas 4.2 and 4.3, and the last inequality is due to $m \leq n$. Finally, replacing ε with $\frac{20\sigma_A}{t} \cdot \sqrt{mn}$ and using the definitions of c_1 and c_2 from (4.9), we conclude the inequality (4.12). The result is valid for any $\varepsilon > 0$ and the corresponding $t > 0$. $\square$

Step 3. Tail bound of φ . Before showing the result, we define two constants $c_4, c_5 > 0$ that will be frequently used:

$$c_4 := \mu_u + 2\sigma_u \quad \text{and} \quad c_5 := C_{rv2} \wedge 2. \quad (4.14)$$

Recall that here μ_u and σ_u denote the maxima of the means and the sub-Gaussian parameters of u 's components, respectively.

LEMMA 4.9. *For all $t > 0$, the following tail bound holds:*

$$\Pr[\varphi \geq t] \leq \frac{c_4 n^2}{t} + e^{-2n}. \quad (4.15)$$

Furthermore, the following probabilistic bound holds:

$$\Pr \left[\min_{1 \leq i \leq n} (\hat{x}_i + \hat{s}_i) \leq c_4 \right] \geq 1 - e^{-2n}. \quad (4.16)$$

Proof. Recall that the vector $u := (u^1, u^2)$ equals $\hat{x} + \hat{s}$ by definition, and

$$\min_{1 \leq i \leq n} (\hat{x}_i + \hat{s}_i) = \min_{1 \leq i \leq n} u_i, \quad \varphi = \frac{\|\hat{x} + \hat{s}\|_1}{\min_{1 \leq i \leq n} (\hat{x}_i + \hat{s}_i)} = \frac{\sum_{i=1}^n u_i}{\min_{1 \leq i \leq n} u_i}. \quad (4.17)$$

In the remainder of the proof we will mainly work on u .

We first derive the probabilistic upper bound of $\sum_{i=1}^n u_i$. Suppose that for each u_i , its sub-Gaussian parameter is σ_i . Using Lemma 4.1, it holds for all $t > 0$ that

$$\Pr \left[\sum_{i=1}^n u_i \geq n\mu_u + t \right] \leq \Pr \left[\sum_{i=1}^n u_i \geq \sum_{i=1}^n \mathbb{E}[u_i] + t \right] \leq \exp \left\{ -\frac{t^2}{2 \sum_{i=1}^n \sigma_i^2} \right\} \leq \exp \left\{ -\frac{t^2}{2n\sigma_u^2} \right\}. \quad (4.18)$$

Here the first inequality is due to $\mathbb{E}[u_i] \leq \mu_u$ for all $i = 1, 2, \dots, n$, the second inequality uses Lemma 4.1, and the last inequality is due to $\sigma_u \geq \sigma_i$ for all $i = 1, 2, \dots, n$. Taking $t = 2n\sigma_u$ in (4.18) gives $\Pr[\sum_{i=1}^n u_i \geq nc_4] \leq e^{-2n}$. Since $\min_{1 \leq i \leq n} u_i \leq n^{-1} \sum_{i=1}^n u_i$, this proves (4.16).

Next we analyze the probabilistic lower bound of $\min_{1 \leq i \leq n} u_i$. Because the probability density of each u_i is bounded by 1, the union bound gives, for any $t > 0$,

$$\Pr \left[\min_{1 \leq i \leq n} u_i \leq t \right] \leq \sum_{i=1}^n \Pr[u_i \leq t] \leq nt. \quad (4.19)$$

Let E_1 denote the event $\sum_{i=1}^n u_i \geq n\mu_u + 2n\sigma_u$. As shown above, $\Pr[E_1] \leq e^{-2n}$. For all $\delta > 0$ we use $E_{2,\delta}$ to denote the event $\min_{1 \leq i \leq n} u_i \leq \frac{\delta}{n}$. By (4.19) we have $\Pr[E_{2,\delta}] \leq \delta$. Therefore, for all $\delta > 0$, it holds that

$$\Pr \left[\frac{\sum_{i=1}^n u_i}{\min_{1 \leq i \leq n} u_i} \geq \frac{n^2(\mu_u + 2\sigma_u)}{\delta} \right] \leq \Pr[E_1 \text{ or } E_{2,\delta}] \leq \Pr[E_1] + \Pr[E_{2,\delta}] \leq e^{-2n} + \delta. \quad (4.20)$$

Replacing δ in (4.20) with $\frac{c_4 n^2}{t}$ and noting that $\mu_u + 2\sigma_u = c_4$ and $\frac{\sum_{i=1}^n u_i}{\min_{1 \leq i \leq n} u_i} = \varphi$, we can conclude (4.15). $\square$

Step 4. Tail bound of $\kappa\Phi$. With the tail bounds of κ and $\|B^{-1}\| \|A\|$ provided in Lemmas 4.7 and 4.8, we can derive a tail bound on their product. Using this tail bound, we then have the following tail bound of $\kappa\Phi$.

LEMMA 4.10. *For any $\delta \in (0, 1]$, the following bound holds:*

$$\Pr \left[\kappa\Phi \geq \frac{\ln(3/\delta)}{\delta} \cdot 6c_3c_4 \cdot \frac{n^{3.5}m^{0.5}}{d+1} \right] \leq \delta + \left(\frac{1}{2}\right)^{d+1} + 6e^{-c_5m}. \quad (4.21)$$

Proof. First of all, we prove the following tail bound of $\kappa\|B^{-1}\| \|A\|$. For any $t > 0$, we claim the following:

$$\Pr \left[\kappa\|B^{-1}\| \|A\| \geq t \right] \leq \frac{1}{t} \cdot \frac{c_3n^{1.5}m^{0.5}}{d+1} + \left(\frac{1}{2}\right)^{d+1} + 4e^{-c_1m}. \quad (4.22)$$

Let t_0 be an arbitrary positive scalar, and we use E_1 and E_2 to denote the events $\kappa \geq \alpha_1 := 2c_2 \cdot \frac{n}{d+1}$ and $\|B^{-1}\| \|A\| \geq \alpha_2 := t_0 \cdot c_2\sqrt{mn}$, respectively. Then we have

$$\begin{aligned} \Pr \left[\kappa\|B^{-1}\| \|A\| \geq t_0 \cdot 2c_2^2 \cdot \frac{n^{1.5}m^{0.5}}{d+1} \right] &= \Pr \left[\kappa\|B^{-1}\| \|A\| \geq \alpha_1\alpha_2 \right] \\ &\leq \Pr \left[\kappa \geq \alpha_1 \text{ or } \|B^{-1}\| \|A\| \geq \alpha_2 \right] = \Pr[E_1 \cup E_2] \leq \Pr[E_1] + \Pr[E_2] \\ &\leq \left(\frac{1}{2}\right)^{d+1} + 2e^{-c_1n} + \frac{1}{t_0} + 2e^{-c_1m} \leq \frac{1}{t_0} + \left(\frac{1}{2}\right)^{d+1} + 4e^{-c_1m} \end{aligned} \quad (4.23)$$

where the second inequality uses the union bound, and the third inequality uses Lemmas 4.7 and 4.8 on $\Pr[E_1]$ and $\Pr[E_2]$, and the last inequality is due to $m \leq n$. The above inequality holds for any $t_0 > 0$. Note that in the above (4.23), $2c_2^2$ is equal to c_3 . Let $t_0 = t \cdot \left(c_3 \cdot \frac{n^{1.5}m^{0.5}}{d+1}\right)^{-1}$ and then (4.23) simplifies to (4.22). This proves the tail bound claim (4.22).

With the tail bound of $\kappa\|B^{-1}\| \|A\|$ and the tail bound of φ from Lemma 4.9, we can now derive a tail bound of $\kappa\|B^{-1}\| \|A\| \cdot \varphi$ using Lemma 4.4. In the rest of the proof we use Z to denote $\kappa\|B^{-1}\| \|A\|$.

From the definition of the probabilistic model, Z and φ are independent. Then we can use Lemma 4.4 to conclude the high-probability bound of $Z\varphi$. Specifically, for any $\delta \in (0, 1]$:

$$\Pr \left[Z\varphi \geq \frac{\ln(3/\delta)}{\delta} \cdot 6c_4n^2 \cdot \frac{c_3n^{1.5}m^{0.5}}{d+1} \right] \leq \delta + \left(\frac{1}{2}\right)^{d+1} + 4e^{-c_1m} + 2e^{-2n} \leq \delta + \left(\frac{1}{2}\right)^{d+1} + 6e^{-c_5m}. \quad (4.24)$$

Here the last inequality holds by noting that $c_5 := C_{rv2} \wedge 2$ is no larger than either 2 or the c_1 defined in (4.9), and $n \geq m$.

Finally, Lemma 4.6 states that $\kappa\Phi \leq Z\varphi$ when $(\hat{x}, \hat{s})$ satisfies strict complementary slackness, which holds almost surely. Therefore, for any $T > 0$, $\Pr[\kappa\Phi \geq T] \leq \Pr[Z\varphi \geq T]$. Substituting this relationship into (4.24) completes the proof. $\square$

Step 5. Proof of Theorem 3.1 and Corollary 3.1. Finally, we finish the proof.

Proof of Theorem 3.1. Let E_1 denote the event $\kappa\Phi \leq \alpha_1 := \frac{\ln(3/\delta)}{\delta} \cdot 6c_3c_4 \cdot \frac{n^{3.5}m^{0.5}}{d+1}$. Using (4.5) of Theorem 4.1, $T_{basis} \leq \check{c}_1\kappa\Phi \cdot \ln(\kappa\Phi)$ for an absolute constant $\check{c}_1 > 1$, and we have:

$$\Pr[T_{basis} \leq \check{c}_1\alpha_1 \ln(\check{c}_1\alpha_1)] \geq \Pr[\check{c}_1\kappa\Phi \cdot \ln(\kappa\Phi) \leq \check{c}_1\alpha_1 \ln(\alpha_1)] = \Pr[E_1] \geq 1 - \left(\delta + \left(\frac{1}{2}\right)^{d+1} + 6e^{-c_5m}\right)$$

where the first inequality uses $\check{c}_1 > 1$ and the last inequality follows from Lemma 4.10. This inequality is exactly (3.3) if letting c_0 be c_4 , C_1 be $6\check{c}_1c_3$ and C_0 be c_5 . Furthermore, C_1 and C_0 depend only (and at most polynomially) on σ_A .

Now, let E_2 denote the event $\|B^{-1}\| \|A\| \leq \alpha_2 := c_2\sqrt{mn}/\delta$ and let E_3 denote the event $\min_{1 \leq i \leq n} u_i \leq c_4$. By the convention following (4.7), E_2 implies that B is full-rank. Since $u > 0$ almost surely, Lemma 3.1

then gives $x^\star = \hat{x}$ and $s^\star = \hat{s}$ on E_2 . Thus, (4.6) gives $T_{local} \leq \check{c}_2 \|B^{-1}\| \|A\| \max\{0, \ln(\min_{1 \leq i \leq n} u_i / \varepsilon)\}$. On $E_2 \cap E_3$, this implies

$$T_{local} \leq \check{c}_2 \alpha_2 \max\left\{0, \ln\left(\frac{c_4}{\varepsilon}\right)\right\}.$$

Furthermore, Lemma 4.8, (4.16), and the union bound give

$$\Pr[E_2 \cap E_3] \geq 1 - \delta - 2e^{-c_1 m} - e^{-2n} \geq 1 - \delta - 3e^{-c_5 m}, \quad (4.25)$$

where the last inequality uses $c_5 \leq c_1$, $c_5 \leq 2$, and $m \leq n$. Therefore, the preceding bound and (4.25) yield (3.4) with $c_0 := c_4$, $C_2 := \check{c}_2 c_2$, and $C_0 := c_5$. $\square$

REMARK 4.1. From the above proof, we may conclude that the constants C_0, C_1, C_2 in Theorem 3.1 can be explicitly represented as $C_0 := (C_{rv2} \wedge 2)$, $C_1 := 4800 \check{c}_1 \sigma_A^2 C_{rv1}^2$, and $C_2 := 20 \check{c}_2 \sigma_A C_{rv1}$. Here (C_{rv1}, C_{rv2}) are the constants from Theorem 1.1 of Rudelson and Vershynin [41] (used in Lemma 4.2), and $\check{c}_1, \check{c}_2$ are the absolute constants in Theorem 4.1 of Xiong [56] (used in Theorem 4.1). Here (C_{rv1}, C_{rv2}) depend only (and at most polynomially) on the sub-Gaussian parameter σ_A of A .

REMARK 4.2. The main idea of the proof of (3.3) is to establish the high-probability upper bound for $\kappa\Phi$ in Lemma 4.10. Theorem 3.1 of Xiong [56] shows a global linear convergence iteration bound $O(\kappa\Phi \cdot \log(\frac{\kappa\Phi(\|x^\star\| + \|s^\star\|)}{\varepsilon}))$ for reaching an ε -optimal solution. Since this complexity bound also heavily relies on $\kappa\Phi$, it is possible to use Lemma 4.10 to derive a high-probability bound on the global linear convergence of rPDHG. The bound in (3.4) is consistent with fast local linear convergence because it depends only logarithmically on c_0/ε .

Proof of Corollary 3.1. Fix $\delta \in (0, 1]$, set $\delta_0 := \delta/22$, and apply Theorem 3.1 with δ_0 . The condition in the corollary implies $2\delta_0 > \max\{e^{-C_0 m}, 2^{-(d+1)}\}$. Therefore, the right-hand sides of (3.3) and (3.4) are lower bounded by $1 - 15\delta_0$ and $1 - 7\delta_0$, respectively. By the union bound, both events occur with probability at least $1 - 22\delta_0 = 1 - \delta$. For $\varepsilon \in (0, e^{-1}]$, $\max\{0, \ln(c_0/\varepsilon)\} \leq (1 + \max\{0, \ln c_0\}) \ln(1/\varepsilon)$. The factor $1 + \max\{0, \ln c_0\}$ is absorbed into the $O(\cdot)$ notation, and the fixed rescaling from δ to δ_0 is absorbed into both order terms. $\square$

5. Proof of Theorem 3.2 In this section, we prove Theorem 3.2. Section 5.1 introduces a few lemmas on random matrices, such as probabilistic lower bounds for the intermediate singular values of a random matrix and the Hanson-Wright inequality. Section 5.2 contains the detailed proof of Theorem 3.2.

5.1. Useful helper lemmas on random matrices In this subsection, we introduce some useful lemmas. The first lemma below shows that the k -th largest singular value of a random matrix in $\mathbb{R}^{n \times n}$ grows linearly as $\Omega\left(\frac{n+1-k}{\sqrt{n}}\right)$ with high probability for all $k = 1, 2, \dots, n$.

LEMMA 5.1. *Let $W \in \mathbb{R}^{n \times n}$ be a random matrix as in Definition 3.2. For every $\delta > 0$, we have*

$$\Pr\left[\sigma_k(W) \geq \frac{\delta(n-k+1)}{4C_{rv1}\sqrt{n}} \text{ for all } k = 1, 2, \dots, n\right] \geq 1 - \delta - ne^{-C_{rv2}n}, \quad (5.1)$$

where C_{rv1} and C_{rv2} are constants from Lemma 4.2.

The proof of the above lemma uses the min-max principle for singular values and Lemma 4.2 on submatrices of the random matrix. The proof is deferred to Appendix D.

The second lemma is the Hanson-Wright inequality, which provides a tail bound on the quadratic form of a random vector.

LEMMA 5.2 (Hanson-Wright inequality, Theorem 6.2.1 of Vershynin [53]). *Let $v \in \mathbb{R}^n$ be a random vector with i.i.d. components, each following a sub-Gaussian distribution with mean zero, unit variance, and parameter σ . Let M be a matrix in $\mathbb{R}^{n \times n}$. Then for all $t \geq 0$, we have*

$$\Pr\left[|v^\top M v - \mathbb{E}[v^\top M v]| \geq t\right] \leq 2 \cdot \exp\left[-C_{hw} \cdot \min\left(\frac{t^2}{\sigma^4 \|M\|_F^2}, \frac{t}{\sigma^2 \|M\|}\right)\right], \quad (5.2)$$

where C_{hw} is a positive absolute constant.

The Hanson-Wright inequality directly yields a tail bound for $\|W^{-1}v\|^2$, where v is a random vector and W is a matrix in $\mathbb{R}^{n \times n}$ with a guaranteed lower bound on all its singular values. The proof is deferred to Appendix D.

LEMMA 5.3. *Suppose $W \in \mathbb{R}^{n \times n}$ satisfies the condition that there exists a constant $c_0 > 0$ such that for all $k = 1, 2, \dots, n$ it holds that*

$$\sigma_k(W) \geq \frac{c_0(n-k+1)}{\sqrt{n}}. \quad (5.3)$$

Let $v \in \mathbb{R}^n$ be a random vector with i.i.d. components that follow a zero-mean, unit-variance, sub-Gaussian distribution of parameter σ . Then for all $\gamma \geq 0$, the following inequality holds:

$$\Pr \left[\|W^{-1}v\|^2 \geq \frac{2n(1+\sigma^2\gamma)}{c_0^2} \right] \leq 2 \cdot \exp(-2C_{hw} \min\{\gamma^2, \gamma\}). \quad (5.4)$$

Here C_{hw} is the constant in Lemma 5.2.

5.2. Proof of Theorem 3.2 In the proof of Theorem 3.1, we used Lemma 4.5 to derive an upper bound for $\kappa\Phi$. However, directly using the expression of Φ in Definition 4.4, we can derive a tighter upper bound for $\kappa\Phi$ in the case of the probabilistic model with Gaussian constraint matrices. The proof of Theorem 3.2 is organized into four steps.

Roadmap of the proof of Theorem 3.2. In Step 1, we show that $\kappa\Phi$ can be bounded as $\kappa\Phi \leq \kappa \cdot \varphi \cdot (Z_p \vee Z_d)$, where Z_p and Z_d are newly defined random variables. We have already established probabilistic upper bounds for κ and φ in Section 4, so in Step 2, we derive probabilistic upper bounds for $Z_p \vee Z_d$. In Step 3, we combine these results to obtain a new high-probability bound for $\kappa\Phi$. Finally, we complete the proof of Theorem 3.2 and Corollary 3.2 in Step 4.

Step 1. A new upper bound on $\kappa\Phi$. We start with some useful properties of Gaussian matrices. Gaussian matrices are orthogonally invariant. Specifically, suppose that A is an $m \times n$ Gaussian matrix where $m < n$. The null space of A has the same distribution as the row space of another $d \times n$ Gaussian matrix Q that depends on A . This property yields useful features of the probabilistic model, such as symmetry between the primal and dual problems (see more in Todd [51]). In other words, for an instance of the probabilistic model with a Gaussian constraint matrix, the dual problem (2.1) is itself an instance of the probabilistic model with a Gaussian constraint matrix.

LEMMA 5.4 (Theorem 2.4 of Todd [51]). *For an instance of the probabilistic model in Definition 3.4 with a Gaussian constraint matrix $A \in \mathbb{R}^{m \times n}$, the dual problem (2.1) is also an instance of the probabilistic model with a Gaussian constraint matrix $Q \in \mathbb{R}^{d \times n}$.*

According to Lemma 3.1, instances of the probabilistic model with Gaussian constraint matrices satisfy Condition 1 almost surely with $\mathcal{X}^* = \{\hat{x}\}$ and $\mathcal{S}^* = \{\hat{s}\}$, because Gaussian matrices are full-rank almost surely. Consequently, the optimal basis of the primal problem is almost surely $\{1, 2, \dots, m\}$, while the optimal basis of its dual problem (2.1) is almost surely $\{m+1, m+2, \dots, n\}$. Let Θ denote the primal optimal basis $\{1, 2, \dots, m\}$, and let $\bar{\Theta}$ denote its complement $\{m+1, m+2, \dots, n\}$. Then A_Θ is exactly B and $A_{\bar{\Theta}}$ is N . For the primal problem, the simplex tableau $A_\Theta^{-1}A_{\bar{\Theta}}$ at the optimal solution is $B^{-1}N$. For the dual problem, it is $Q_{\bar{\Theta}}^{-1}Q_\Theta$, where $Q_{\bar{\Theta}}^{-1}Q_\Theta$ is equal to $-(B^{-1}N)^\top$, as stated below:

LEMMA 5.5 (Lemma 3.6 of Xiong [56]). *The matrix $Q_{\bar{\Theta}}^{-1}Q_\Theta$ is equal to $-(B^{-1}N)^\top$.*

Using Lemma 5.5, the terms $\|(B^{-1}N)_{\cdot,j}\|^2$ in Definition 4.4 of Φ can be converted into $\|(Q_{\bar{\Theta}}^{-1}Q_{\Theta})_{\cdot,i}\|^2$. Given that Condition 1 holds almost surely with $\mathcal{X}^* = \{\hat{x}\}$ and $\mathcal{S}^* = \{\hat{s}\}$, we can rewrite Φ as follows:

$$\begin{aligned} \Phi &= (\|\hat{x} + \hat{s}\|_1) \cdot \max \left\{ \max_{1 \leq j \leq d} \frac{\sqrt{\|(B^{-1}N)_{\cdot,j}\|^2 + 1}}{\hat{s}_{j+m}}, \max_{1 \leq i \leq m} \frac{\sqrt{\|(Q_{\bar{\Theta}}^{-1}Q_{\Theta})_{\cdot,i}\|^2 + 1}}{\hat{x}_i} \right\} \\ &\leq \underbrace{\frac{\|\hat{x} + \hat{s}\|_1}{\min_{1 \leq i \leq n} (\hat{x}_i + \hat{s}_i)}}_{\text{Denoted by } \varphi} \cdot \max \left\{ \underbrace{\max_{1 \leq j \leq d} \sqrt{\|(B^{-1}N)_{\cdot,j}\|^2 + 1}}_{\text{Denoted by } Z_p}, \underbrace{\max_{1 \leq i \leq m} \sqrt{\|(Q_{\bar{\Theta}}^{-1}Q_{\Theta})_{\cdot,i}\|^2 + 1}}_{\text{Denoted by } Z_d} \right\} \end{aligned} \quad (5.5)$$

in which the first term of the product is denoted by φ , and the two terms in the bracket are denoted by Z_p and Z_d , respectively. Therefore, almost surely, we have

$$\kappa\Phi \leq \kappa \cdot \varphi \cdot (Z_p \vee Z_d). \quad (5.6)$$

Tail bounds for κ and φ have already been established in Lemma 4.7 and Lemma 4.9, respectively. Thus, the primary focus will be on deriving a tail bound for $Z_p \vee Z_d$.

Step 2. Tail bound of $Z_p \vee Z_d$. Although Z_p and Z_d are not independent, they have a symmetric structure, and the matrices B and $Q_{\bar{\Theta}}$ are both Gaussian random matrices, which facilitates our analysis. To simplify notation, we define the following parameters for any $\gamma \geq 0$:

$$c_6 := \sqrt{256 \cdot C_{rv1}^2 + 1} \quad \text{and} \quad \delta_\gamma := 2n \cdot \exp(-2C_{hw} \min\{\gamma^2, \gamma\}) + n \cdot e^{-C_{rv2} \cdot (m \wedge d)}. \quad (5.7)$$

Here c_6 is a fixed absolute constant as C_{rv1} depends only on σ_A , which is equal to 1 for Gaussian matrices. Additionally, δ_γ decreases exponentially with γ , m , and d , so for sufficiently large γ , m , and d , δ_γ becomes negligible compared to δ .

LEMMA 5.6. *For all $\delta > 0$ and $\gamma \geq 1$, the following bound holds:*

$$\Pr \left[Z_p \vee Z_d \geq \frac{c_6 \sqrt{n\gamma}}{\delta} \right] \leq \delta + \delta_\gamma \quad (5.8)$$

for c_6 and δ_γ defined in (5.7).

Proof. When $\delta > 1$, (5.8) is trivial. Later we consider the case $\delta \in (0, 1]$.

We first study the upper bound for Z_p . The same reasoning will apply symmetrically to Z_d . For the matrix B , let E_δ denote the event $\sigma_k(B) \geq \frac{\delta}{4C_{rv1}} \cdot \frac{m-k+1}{\sqrt{m}}$ for all $k = 1, 2, \dots, m$. By Lemma 5.1, we have

$$\Pr[E_\delta] \geq 1 - \delta - me^{-C_{rv2}m}. \quad (5.9)$$

Next, Lemma 5.3 implies that for each $j \in \{1, \dots, d\}$,

$$\Pr \left[\|B^{-1}N_{\cdot,j}\|^2 \geq \frac{2m(1+\sigma^2\gamma)}{\frac{\delta^2}{16C_{rv1}^2}} \middle| E_\delta \right] \leq 2 \cdot \exp(-2C_{hw} \min\{\gamma^2, \gamma\}). \quad (5.10)$$

Here $\sigma = 1$ because $N_{\cdot,j}$ is a standard Gaussian vector (unit variance entries). Using the union bound over all $j \in \{1, \dots, d\}$, it follows that:

$$\Pr \left[\max_{1 \leq j \leq d} \|B^{-1}N_{\cdot,j}\|^2 \geq \frac{32mC_{rv1}^2(1+\gamma)}{\delta^2} \middle| E_\delta \right] \leq 2d \cdot \exp(-2C_{hw} \min\{\gamma^2, \gamma\}). \quad (5.11)$$

Combining (5.9) and (5.11), we have:

$$\Pr \left[\max_{1 \leq j \leq d} \|B^{-1}N_{\cdot,j}\|^2 \geq \frac{32mC_{rv1}^2(1+\gamma)}{\delta^2} \right] \leq 2d \cdot \exp(-2C_{hw} \min\{\gamma^2, \gamma\}) + \delta + me^{-C_{rv2}m}. \quad (5.12)$$

Similarly, for $\max_{1 \leq i \leq m} \|(Q_{\bar{\Theta}}^{-1}Q_{\Theta})_{\cdot,i}\|^2$, we have:

$$\Pr \left[\max_{1 \leq i \leq m} \|(Q_{\bar{\Theta}}^{-1}Q_{\Theta})_{\cdot,i}\|^2 \geq \frac{32dC_{rv1}^2(1+\gamma)}{\delta^2} \right] \leq 2m \cdot \exp(-2C_{hw} \min\{\gamma^2, \gamma\}) + \delta + de^{-C_{rv2}d}. \quad (5.13)$$

Combining (5.12) and (5.13), and noting that $n = m + d$, we have:

$$\Pr \left[Z_p \vee Z_d \geq \sqrt{\frac{32nC_{rv1}^2(1+\gamma)}{\delta^2} + 1} \right] \leq 2n \cdot \exp(-2C_{hw} \min\{\gamma^2, \gamma\}) + 2\delta + n \cdot e^{-C_{rv2} \cdot (m \wedge d)}. \quad (5.14)$$

Replacing δ in (5.14) with $\delta/2$ proves:

$$\Pr \left[Z_p \vee Z_d \geq \sqrt{\frac{128n \cdot C_{rv1}^2(1+\gamma)}{\delta^2} + 1} \right] \leq \delta + 2n \cdot \exp(-2C_{hw} \min\{\gamma^2, \gamma\}) + n \cdot e^{-C_{rv2} \cdot (m \wedge d)}. \quad (5.15)$$

To show (5.8), note that for $\gamma \geq 1$ and $\delta \in (0, 1]$,

$$\sqrt{\frac{128n \cdot C_{rv1}^2(1+\gamma)}{\delta^2} + 1} \leq \sqrt{\frac{256C_{rv1}^2 \cdot n\gamma}{\delta^2} + 1} \leq \sqrt{\frac{(256C_{rv1}^2 + 1) \cdot n\gamma}{\delta^2}} = \frac{c_6\sqrt{n\gamma}}{\delta}.$$

Substituting this inequality into (5.15) yields (5.8). $\square$

Step 3. Tail bound of $\kappa\Phi$. We now combine the tail bounds established in Lemma 4.7, Lemma 4.9, and Lemma 5.6 to derive a tail bound for $\kappa \cdot \varphi \cdot (Z_p \vee Z_d)$, which directly provides a probabilistic upper bound for $\kappa\Phi$. To simplify the notation, we define two absolute constants:

$$c_7 := C_{rv2} \wedge \ln(2) \quad \text{and} \quad c_8 := \max \left\{ 1, \frac{1}{\sqrt{2C_{hw}}} \right\}. \quad (5.16)$$

Here c_7 is an absolute constant because C_{rv2} only depends on the sub-Gaussian parameter σ_A , which is always equal to 1 for Gaussian matrices.

LEMMA 5.7. *For all $\delta \in (0, 1]$, the following inequality holds:*

$$\Pr \left[\kappa\Phi \geq 24c_2c_4c_6c_8 \cdot \frac{n^{3.5}}{d+1} \cdot \frac{\ln\left(\frac{6}{\delta}\right) \sqrt{\ln\left(\frac{4n}{\delta}\right)}}{\delta} \right] \leq \delta + (n+5) \left(\frac{1}{e^{c_7}} \right)^{m \wedge d}. \quad (5.17)$$

Proof. First of all, we study the tail bound of $\kappa \cdot (Z_p \vee Z_d)$. According to Lemma 4.7, using the union bound, we obtain the following probabilistic upper bound of $\kappa \cdot (Z_p \vee Z_d)$:

$$\Pr \left[\kappa \cdot (Z_p \vee Z_d) \geq \frac{c_6\sqrt{n\gamma}}{\delta} \cdot 2c_2 \cdot \frac{n}{d+1} \right] \leq \delta + \underbrace{\delta_\gamma + \left(\frac{1}{2}\right)^{d+1} + 2 \left(\frac{1}{e^{c_1}}\right)^n}_{\text{Denoted by } \delta_0} \quad (5.18)$$

for all $\gamma \geq 1$ and $\delta > 0$. Replacing $\frac{c_6\sqrt{n\gamma}}{\delta} \cdot 2c_2 \cdot \frac{n}{d+1}$ by t yields the following tail bound:

$$\Pr \left[\kappa \cdot (Z_p \vee Z_d) \geq t \right] \leq \frac{2c_2c_6\sqrt{n\gamma} \cdot \frac{n}{d+1}}{t} + \delta_\gamma + \delta_0. \quad (5.19)$$

Next, observe that $\kappa \cdot (Z_p \vee Z_d)$ and φ are independent, and φ already has a tail bound in Lemma 4.9. Let $X = \kappa \cdot (Z_p \vee Z_d)$ and $Y = \varphi$. Using Lemma 4.4, we can derive the probabilistic upper bound for $XY = \kappa \cdot \varphi \cdot (Z_p \vee Z_d)$:

$$\Pr \left[\kappa \cdot \varphi \cdot (Z_p \vee Z_d) \geq \frac{6 \cdot 2c_2c_6\sqrt{n\gamma} \cdot \frac{n}{d+1} \cdot c_4n^2 \cdot \ln(3/\delta)}{\delta} \right] \leq \delta + \delta_\gamma + \delta_0 + 2 \cdot e^{-2n} \quad (5.20)$$

for all $\delta \in (0, 1]$ and $\gamma \geq 1$. According to (5.6), $\kappa\Phi \leq \kappa \cdot \varphi \cdot (Z_p \vee Z_d)$ so

$$\Pr \left[\kappa\Phi \geq 12c_2c_4c_6 \cdot \frac{n^{3.5}}{d+1} \cdot \frac{\sqrt{\gamma} \cdot \ln(3/\delta)}{\delta} \right] \leq \delta + \delta_\gamma + \delta_0 + 2 \cdot e^{-2n}. \quad (5.21)$$

Then we simplify the right-hand side of (5.21). Substituting the values of δ_γ and δ_0 , we have:

$$\begin{aligned} \delta + \delta_\gamma + \delta_0 + 2 \cdot e^{-2n} &= \delta + 2n \cdot \exp(-2C_{hw} \min\{\gamma^2, \gamma\}) + n \cdot \left(\frac{1}{e^{C_{rv2}}} \right)^{m \wedge d} + \left(\frac{1}{2} \right)^{d+1} \\ &+ 2 \left(\frac{1}{e^{C_{rv2} \wedge 6}} \right)^n + 2 \cdot \left(\frac{1}{e^2} \right)^n \leq \delta + 2n \cdot \exp(-2C_{hw} \min\{\gamma^2, \gamma\}) + (n+5) \left(\frac{1}{e^{c_7}} \right)^{m \wedge d} \end{aligned} \quad (5.22)$$

where the last inequality follows from $c_7 = C_{rv2} \wedge \ln(2)$, as defined in (5.16). We now choose:

$$\gamma = \max \left\{ 1, \frac{1}{2C_{hw}} \cdot \ln \left(\frac{2n}{\delta} \right) \right\}. \quad (5.23)$$

In this way, since $\gamma \geq 1$ we have $\min\{\gamma^2, \gamma\} = \gamma$, and hence $2n \cdot \exp(-2C_{hw} \min\{\gamma^2, \gamma\}) = 2n \cdot \exp(-2C_{hw}\gamma) \leq \delta$. From (5.22) we conclude that $\delta + \delta_\gamma + \delta_0 + 2 \cdot e^{-2n}$ is further upper bounded by $2\delta + (n+5) \left(\frac{1}{e^{c_7}} \right)^{m \wedge d}$. Substituting the value of γ into (5.21) yields

$$\Pr \left[\kappa\Phi \geq 12c_2c_4c_6 \cdot \frac{n^{3.5}}{d+1} \cdot \frac{\ln \left(\frac{3}{\delta} \right) \cdot \max \left\{ 1, \sqrt{\frac{1}{2C_{hw}} \cdot \ln \left(\frac{2n}{\delta} \right)} \right\}}{\delta} \right] \leq 2\delta + (n+5) \left(\frac{1}{e^{c_7}} \right)^{m \wedge d}. \quad (5.24)$$

Finally, since $\max \left\{ 1, \sqrt{\frac{1}{2C_{hw}} \cdot \ln \left(\frac{2n}{\delta} \right)} \right\} \leq c_8 \sqrt{\ln \left(\frac{2n}{\delta} \right)}$ for $\delta \leq \frac{1}{2}$ (as per (5.16)), replacing δ with $\frac{\delta}{2}$ in the above inequality yields the desired probabilistic bound (5.17). $\square$

Step 4. Proof of Theorem 3.2 and Corollary 3.2. Finally, we finish the proof.

Proof of Theorem 3.2. Let E denote the event $\kappa\Phi \leq \alpha := 24c_2c_4c_6c_8 \cdot \frac{n^{3.5}}{d+1} \cdot \frac{\ln \left(\frac{6}{\delta} \right) \sqrt{\ln \left(\frac{4n}{\delta} \right)}}{\delta}$. By (4.5) of Theorem 4.1, $T_{basis} \leq \check{c}_1 \kappa\Phi \cdot \ln(\kappa\Phi)$ for an absolute constant $\check{c}_1 > 1$, and thus we have:

$$\Pr[T_{basis} \leq \check{c}_1 \alpha \cdot \ln(\check{c}_1 \alpha)] \geq \Pr[\check{c}_1 \kappa\Phi \cdot \ln(\kappa\Phi) \leq \check{c}_1 \alpha \cdot \ln(\alpha)] \geq \Pr[E] \geq 1 - \delta - (n+5) \left(\frac{1}{e^{c_7}} \right)^{m \wedge d}. \quad (5.25)$$

where the first inequality uses $\check{c}_1 > 1$ and the last inequality is due to Lemma 5.7. The inequality (5.25) yields (3.5) by setting $c_0 = c_4$, $C_0 = c_7$, and $C_1 = 24\check{c}_1c_2c_6c_8$. Here c_2, c_6, c_7 , and c_8 are all absolute constants because $\sigma_A = 1$ for a Gaussian matrix A , so C_0 and C_1 are absolute constants as well. This proves (3.5) with $C_0 = c_7$.

As for T_{local} , the proof is identical to that of Theorem 3.1 because a Gaussian matrix is also a sub-Gaussian matrix. We may take $C_0 = c_7$ because $c_1 = C_{rv2} \wedge 6 \geq c_7 = C_{rv2} \wedge \ln(2)$ and $1 - \delta - 2e^{-c_1m} - e^{-2n} \geq 1 - \delta - 3e^{-c_7m}$. Thus, C_0 and $C_2 = \check{c}_2c_2$ are absolute because $\sigma_A = 1$. This completes the proof of Theorem 3.2. $\square$

REMARK 5.1. From the above proof, (C_0, C_1, C_2) may be explicitly given by $C_0 := (C_{rv2} \wedge \ln 2)$, $C_1 := 24\check{c}_1(20\sigma_A C_{rv1})\sqrt{256C_{rv1}^2 + 1} \cdot \max\left\{1, (2C_{hw})^{-1/2}\right\}$, and $C_2 := 20\check{c}_2\sigma_A C_{rv1}$. Here (C_{rv1}, C_{rv2}) are the constants from Theorem 1.1 of Rudelson and Vershynin [41], $\check{c}_1, \check{c}_2$ are the absolute constants in Theorem 4.1 in Xiong [56], and C_{hw} is the absolute constant in the Hanson-Wright inequality as stated, e.g., in Theorem 6.2.1 of Vershynin [53] (used in Lemma 5.2). Note that (C_{rv1}, C_{rv2}) depend only on σ_A but for the Gaussian matrix A defined in Definition 3.6, σ_A is equal to 1. Thus, C_0, C_1 , and C_2 are absolute constants.

Proof of Corollary 3.2. Fix $\delta \in (0, 1]$, set $\delta_0 := \delta/8$, and apply Theorem 3.2 with δ_0 . The condition in the corollary, $m \wedge d \leq m$, and $n + 5 \geq 6$ imply $(n + 5)e^{-C_0(m \wedge d)} < 4\delta_0$ and $3e^{-C_0 m} < 2\delta_0$. Thus, the right-hand sides of (3.5) and (3.6) are lower bounded by $1 - 5\delta_0$ and $1 - 3\delta_0$, respectively. By the union bound, both events occur with probability at least $1 - 8\delta_0 = 1 - \delta$. For $\varepsilon \in (0, e^{-1}]$, $\max\{0, \ln(c_0/\varepsilon)\} \leq (1 + \max\{0, \ln c_0\}) \ln(1/\varepsilon)$. The factor $1 + \max\{0, \ln c_0\}$ is absorbed into the $O(\cdot)$ notation, and the fixed rescaling from δ to δ_0 is absorbed into both order terms. $\square$

6. Experimental Confirmation of the High-Probability Polynomial-Time Complexity This section presents the experimental results of rPDHG applied to randomly generated LP instances to validate our high-probability complexity analysis. Section 6.1 confirms the tail behavior of the iteration counts in both stages. Section 6.2 demonstrates the polynomial dependence of the number of iterations on the dimension n .

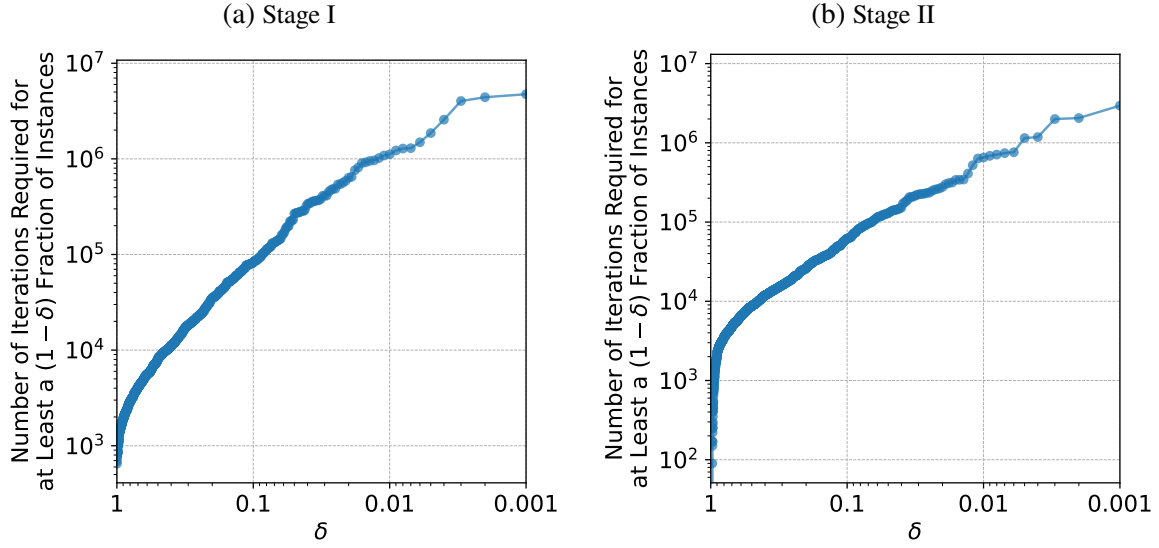
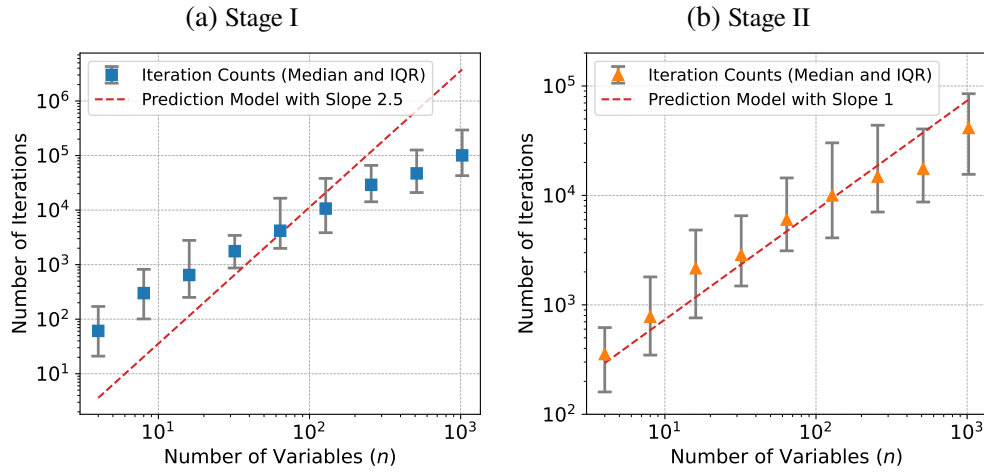
We implement rPDHG for the standard-form problem (1.1), following Algorithm 1 in Appendix A. In our numerical experiments, we regard an LP instance as successfully solved if rPDHG computes a primal-dual solution pair (x^k, y^k) within Euclidean distance 10^{-4} of the optimal solution, namely, $\|(x^k, y^k) - (x^*, y^*)\| \leq 10^{-4}$. Compared with the commonly used KKT residuals, Euclidean distance is a more straightforward measure of optimality, and the probabilistic model with Gaussian input data ensures easy access to (x^*, y^*) . It is also the same criterion used in the experiments of Xiong [56]. We manually classify the iterations into two stages. Stage I comprises all iterations until the support set of x^k settles down to the optimal basis, while Stage II consists of all subsequent iterations. These stages correspond to the optimal basis identification and fast local convergence stages analyzed in Section 3.

6.1. Tail behavior of the iteration count This subsection presents experimental confirmation of the tail behavior of the Stage-I and Stage-II iteration counts. We apply rPDHG to 1,000 LP instances generated according to Definition 3.4 for $n = 100$ and $m = 50$. The constraint matrix A is a Gaussian matrix. The vectors $\hat{x}$ and $\hat{s}$ are constructed with i.i.d. nonzero components, each drawn from the folded Gaussian distribution (absolute value of a Gaussian random variable of zero mean and unit variance). In all experiments, we use the objective vector $\bar{c}$ so that $A\bar{c} = 0$. Figure 1 shows the number of iterations required to complete Stages I and II of at least a $(1 - \delta)$ fraction of instances, for different values of δ in $(0, 1)$.

Theorems 3.1 and 3.2 theoretically predict that the number of iterations required to complete Stages I and II with a $(1 - \delta)$ success rate should grow with $\frac{1}{\delta}$ (at rates $\tilde{O}(\frac{1}{\delta})$ and $O(\frac{1}{\delta})$, respectively), except for exponentially small δ . If the theory is exact, in Figures 1a and 1b (log-log plots with reversed horizontal axes), the slopes should be roughly equal to 1, meaning that the number of iterations required has a reciprocal relationship with δ . Indeed, Figure 1 shows that both slopes are approximately 1, particularly for δ between 0.1 and 0.01. These observations confirm our theoretical prediction for the dependence on δ and the tail behavior of the iteration counts.

6.2. Polynomial dependence of the iteration count on the problem dimension This subsection examines how iteration counts of the two stages scale with the problem dimension (number of variables) in practice. We generate LP instances according to Definition 3.4 with $m = n/2$ for n in $\{4, 8, 16, 32, \dots\}$. As in Section 6.1, the constraint matrix A is a Gaussian matrix, and the vector u has i.i.d. components, each obeying the folded Gaussian distribution. For each value of n , we generate 100 LP instances and compute the first quartile, median, and third quartile of both Stage-I and Stage-II iteration counts. Figure 2 shows the relation between these statistics of the iteration counts and the number of variables n .

Theorem 3.2 predicts that in most cases, the Stage-I and Stage-II iteration counts should scale as $\tilde{O}(n^{2.5})$ and $O(n)$, respectively. If the theory is exact, since Figures 2a and 2b are log-log plots, the slopes should

FIGURE 1. Number of iterations required to complete Stages I and II of at least a $(1 - \delta)$ fraction of instancesFIGURE 2. The median values and interquartile range (IQR) of Stage-I and Stage-II iteration counts for various dimensions n (for LP instances with Gaussian input data).

be approximately 2.5 and 1, respectively, corresponding to polynomial dependencies of degree 2.5 and 1 on n . Based on these insights, we fit prediction models to the median iteration counts using slopes of 2.5 and 1 for Stages I and II. The predicted iteration counts using these models are shown by the dashed lines in the two figures for Stage I and Stage II, respectively. Figure 2a shows that the Stage-I iteration count grows more slowly than the theoretical bound $\tilde{O}(n^{2.5})$ in Theorem 3.2, as it does not grow as fast as the dashed line. Instead, the practical iteration count dependence on n is smaller than our model by roughly a factor proportional to n . For Stage II, Figure 2b shows that the iteration count aligns well with the $O(n)$ dependence predicted by Theorems 3.1 and 3.2. Overall, the slower growth of Stage-II iterations relative to Stage-I iterations validates the insight that local linear convergence is faster, particularly for larger problem dimensions.

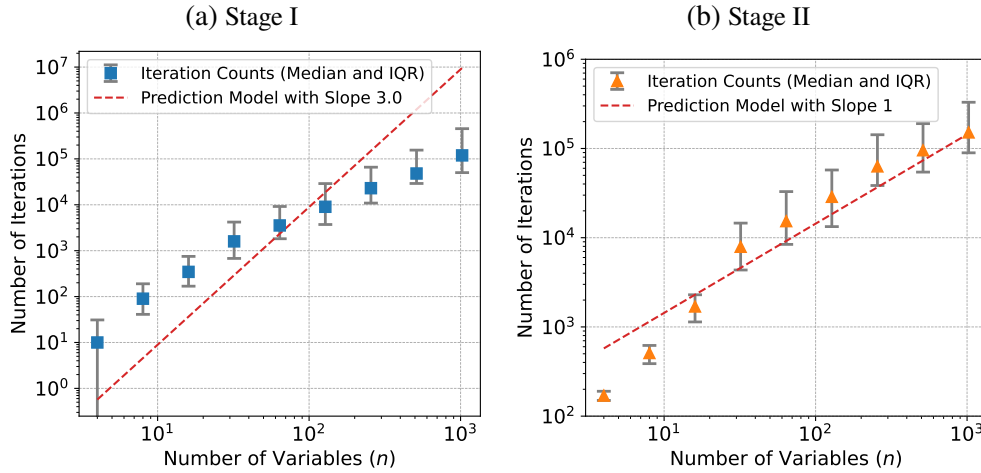
A similar gap between observed practical performance and probabilistic analysis is common in the average-case analysis literature on classic methods as well. For instance, Shamir [44] summarizes many reported observations of simplex method performance on real-life LP problems and concludes that the

number of iterations for real-life problems is usually observed to be no higher than $O(m)$ for LP instances with m linear constraints, but the average-case analyses of Adler and Megiddo [3], Todd [50] and Adler et al. [1] all prove a quadratic dependence of the expected iteration count on the problem dimension in theory. Similarly, the interior-point method iteration count often shows logarithmic growth in the problem dimension n in practice, despite the polynomial dependence on n proven in the probabilistic analyses by Anstreicher et al. [6], Dunagan et al. [19], Ye [61] and others. Our high-probability Stage-I iteration bound (shown by Theorem 3.2) also overestimates the empirical scaling by a factor of $O(n)$, while our Stage-II bound accurately predicts the observed dependence on n . In addition to revealing that the complexity of rPDHG is polynomial in dimension with high probability, an interesting research question is how to further shrink the gap between theory and practice in terms of the polynomial dependence on n . We conjecture that the Stage-I iteration count might actually be bounded by $\tilde{O}(\frac{n^{1.5}}{\delta})$ with probability at least $1 - \delta$ for δ that is not exponentially small. Clearly, this issue merits more research in the future.

From Figure 2a, a small prefactor (about 1.906) multiplying $n^{2.5}$ already upper-bounds the median Stage-I iteration counts over the tested range. Similarly, Figure 2b suggests that a modest prefactor (about 135.0) multiplying n already upper-bounds the median Stage-II iteration counts. These observations empirically suggest that the constants $c_0 C_1$ and C_2 need not be large for the bounds in Theorem 3.2 to hold with probability at least one half.

Furthermore, we conducted the same experiment for LP instances with sub-Gaussian input data. The only difference is that the entries of A are i.i.d. Rademacher random variables (taking values 1 and -1 with probability $1/2$). The experimental results are shown in Figure 3. Theorem 3.1 predicts that in most cases, the Stage-I and Stage-II iteration counts should scale as $\tilde{O}(n^3)$ and $O(n)$, respectively. Based on these insights, we also fit prediction models to the median iteration counts using slopes of 3 and 1 for Stages I and II. The insights from this sub-Gaussian input data case (including the gap with the theoretical predictions) are no different from those in the experiment with Gaussian input data.

FIGURE 3. The median values and interquartile range (IQR) of Stage-I and Stage-II iteration counts for various dimensions n (for LP instances with sub-Gaussian input data).



7. How the Disparity among Optimal Solution Components Affects the Performance of rPDHG

In this section, we use probabilistic analysis to investigate how the disparity among the optimal solution components affects the performance of rPDHG. Let E_{uniq} denote the event that Condition 1 holds, namely, that the primal and dual problems have unique optimal solutions. On E_{uniq} , let x^* and s^* denote the unique primal optimal solution and dual optimal slack, respectively, and define their disparity ratio as follows:

$$\phi := \frac{\frac{1}{n} \cdot \sum_{i=1}^n (x_i^* + s_i^*)}{\min_{1 \leq i \leq n} (x_i^* + s_i^*)}. \quad (7.1)$$

Note that $\phi \geq 1$, with equality holding if and only if $x^\star + s^\star$ is a scalar multiple of the all-ones vector e . The quantities in the conditional bounds below are evaluated on E_{uniq} . Their values outside this event are immaterial. We conduct probabilistic analysis of the iteration bound conditioned on E_{uniq} . Our result shows that, among instances with unique primal and dual optima, T_{basis} depends linearly on the realized disparity ratio ϕ with high conditional probability, while T_{local} is nearly independent of ϕ .

THEOREM 7.1. *Suppose that rPDHG is applied to an instance of the probabilistic model in Definition 3.4 (with objective vector $\bar{c}$). There exist constants $C_0, C_1, C_2 > 0$ that depend only (and at most polynomially) on σ_A for which the following high-probability iteration bounds hold:*

1. (Optimal basis identification)

$$\Pr \left[T_{\text{basis}} \leq \frac{m^{0.5} n^{2.5}}{d+1} \cdot \frac{C_1 \phi}{\delta} \cdot \ln \left(\frac{m^{0.5} n^{2.5}}{d+1} \cdot \frac{C_1 \phi}{\delta} \right) \middle| E_{\text{uniq}} \right] \geq 1 - \delta - 4 \left(\frac{1}{e^{C_0}} \right)^m - \left(\frac{1}{2} \right)^{d+1} \quad (7.2)$$

for any $\delta \in (0, 1]$.

2. (Fast local convergence) Let $\varepsilon > 0$ be any given tolerance.

$$\Pr \left[T_{\text{local}} \leq m^{0.5} n^{0.5} \cdot \frac{C_2}{\delta} \cdot \max \left\{ 0, \ln \left(\frac{\min_{1 \leq i \leq n} (x_i^\star + s_i^\star)}{\varepsilon} \right) \right\} \middle| E_{\text{uniq}} \right] \geq 1 - \delta - 2 \left(\frac{1}{e^{C_0}} \right)^m \quad (7.3)$$

for any $\delta \in (0, 1]$.

The above theorem characterizes the high-probability performance of rPDHG conditioned on the primal and dual optima being unique. The quantities ϕ and $\min_i (x_i^\star + s_i^\star)$ are their realized values on E_{uniq} . Compared with Theorem 3.1, the constant c_0 is absent, and the disparity ratio ϕ is the only quantity related to the optimal solution in (7.2).

The quantity $\frac{n^{2.5}}{d+1}$ is at most as large as $\frac{n^{2.5}}{2}$. When d (recall $d = n - m$) is not too small compared to n in the sense that $d+1 \geq \frac{n}{C_3}$ for some absolute constant $C_3 > 1$, the quantity $\frac{n^{2.5}}{d+1}$ is $O(n^{1.5})$. The high-probability iteration bound for Stage I (optimal basis identification) depends linearly on ϕ , while that for Stage II (local convergence) depends only logarithmically on $\min_{1 \leq i \leq n} (x_i^\star + s_i^\star)$. Roughly speaking, the greater the disparity among the components of $x^\star + s^\star$, the larger ϕ becomes, making an LP instance more challenging for Stage I of rPDHG according to our theory. For example, when $x^\star + s^\star = e$ (the all-ones vector) and m is not too close to n , the Stage-I iteration count is bounded by $\tilde{O}(\frac{n^{1.5} m^{0.5}}{\delta})$ with probability at least $1 - O(\delta)$ for non-exponentially small δ . This bound on T_{basis} is almost n times smaller than the bound of T_{basis} for the random LP instances in Theorem 3.1. Conversely, highly imbalanced components in $x^\star + s^\star$ can lead to larger bounds on T_{basis} than those in Theorem 3.1.

It should be noted that metrics similar to the disparity defined above (such as the ratios between the largest and smallest components, or the smallest positive component) also influence the complexity of other methods, especially interior-point methods, for solving LPs. The proximity measures $\frac{\max_i x_i s_i}{\min_i x_i s_i}$ quantify the distance from the central path and enter the design and analysis of interior-point methods (see Güler and Ye [21], Terlaky [49]). The minimum positive coordinate scale over strictly complementary solutions is used to determine when the iterates of interior-point methods can identify the optimal face (see, e.g., Terlaky [49], Ye [60]). A quantity involving the ratio between the norm and the smallest positive component of a pair of primal-dual strictly complementary optimal solutions directly appears in the finite termination analysis of interior-point methods (see, e.g., Anstreicher et al. [6]). For these methods, such metrics typically enter the iteration complexity only logarithmically, unlike that for rPDHG. Beyond interior-point methods, the strongly polynomial complexity analyses of the simplex method and policy iteration methods for discounted Markov decision problems rely on boundedness of ratios such as $\frac{\|x^\star\|_1}{\min_{i: x_i^\star > 0} x_i^\star}$ (see, e.g., Ye [63]). Additionally, Lu and Yang [33] demonstrate that PDHG (without restarts) exhibits faster local linear convergence within a neighborhood whose size relates to $\min_{1 \leq i \leq n} \{x_i^\star + s_i^\star\}$. In these two settings, such metrics have a larger effect because they appear outside a logarithm.

7.1. Application: Generating difficult LP instances Our analysis suggests that we can generate LP instances with controlled difficulty levels for rPDHG by manipulating the distributions of $\hat{x}$ and $\hat{s}$ in our probabilistic model. By Lemma 3.1, $\hat{x}$ and $\hat{s}$ are almost surely the unique primal and dual optimal solutions if the constraint matrix is a Gaussian matrix. The larger the ratio $\frac{\frac{1}{n} \sum_{i=1}^n (\hat{x}_i + \hat{s}_i)}{\min_{1 \leq i \leq n} (\hat{x}_i + \hat{s}_i)}$, the more difficult the LP instance is likely to be. Such instances can serve as useful benchmarks for evaluating the performance of rPDHG and other first-order LP methods. Below we present an approach to generate difficult random LP instances.

Following Definition 3.4, we consider random instances with a Gaussian constraint matrix $A \in \mathbb{R}^{m \times n}$, where $n = 2m$. The only difference from Definition 3.4 is that we define $\hat{x} = \hat{x}^l$ and $\hat{s} = \hat{s}^l$ for $l \geq 0$ to control ϕ :

$$\hat{x}^l := \begin{pmatrix} u^l \\ 0_m \end{pmatrix} \quad \text{and} \quad \hat{s}^l := \begin{pmatrix} 0_m \\ u^l \end{pmatrix} \quad \text{where } u^l := \left[\underbrace{4^{-l}, 4^{-l}, \dots, 4^{-l}}_{\lfloor m/2 \rfloor \text{ copies}}, \underbrace{1, 1, \dots, 1}_{m - \lfloor m/2 \rfloor \text{ copies}} \right]. \quad (7.4)$$

Here $\lfloor m/2 \rfloor$ denotes the largest integer no greater than $m/2$. These instances have full row rank almost surely, and $\hat{x}^l$ and $\hat{s}^l$ satisfy strict complementary slackness, ensuring that they are the unique pair of optimal solutions. Therefore, the smallest component of $x^* + s^*$ is

$$\min_{1 \leq i \leq n} (x_i^* + s_i^*) = \min_{1 \leq i \leq n} (\hat{x}_i^l + \hat{s}_i^l) = 4^{-l}, \quad (7.5)$$

and the corresponding disparity ratio ϕ_l is equal to

$$\phi_l = \frac{1}{2m} \cdot \frac{2\lfloor m/2 \rfloor \cdot 4^{-l} + 2(m - \lfloor m/2 \rfloor) \cdot 1}{4^{-l}} = \frac{\lfloor m/2 \rfloor}{m} + \left(1 - \frac{\lfloor m/2 \rfloor}{m}\right) \cdot 4^l. \quad (7.6)$$

Note that ϕ_l is approximately equal to 2^{2l-1} when l and m are large enough. In this way, the parameter l controls the magnitude of ϕ_l . Instances with large values of l have large ϕ_l , unbalanced $x^* + s^*$ and small $\min_{1 \leq i \leq n} (x_i^* + s_i^*)$.

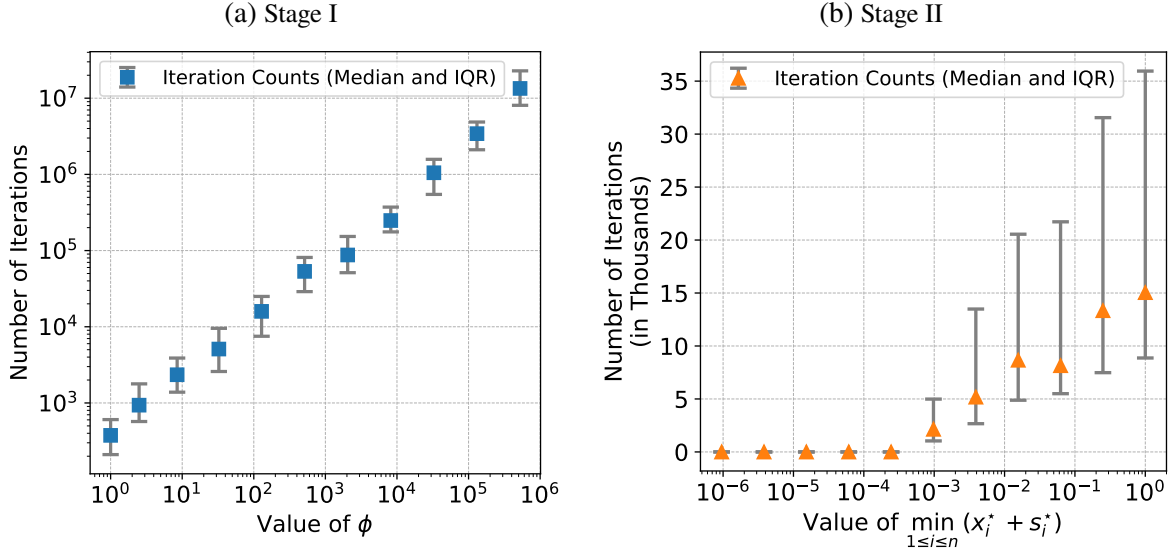
Now we confirm the difficulty of these instances via experiments. We follow the same experimental setup of rPDHG as in Section 6, and we still consider an instance solved when a primal-dual solution pair is within Euclidean distance 10^{-4} of the optimal solution. We set $m = 50$ and $n = 100$. For each $l \in \{0, 1, \dots, 10\}$, we generate 100 LP instances and compute the first quartile, median, and third quartile of both Stage-I and Stage-II iteration counts of rPDHG. Figure 4 shows the relation between these statistics of the Stage-I iteration count and the value of ϕ , and the relation between these statistics of the Stage-II iteration count and the value of $\min_{1 \leq i \leq n} (x_i^* + s_i^*)$, for each family of instances (grouped by l).

Theorem 7.1 predicts that the Stage-I iteration count grows linearly with the value of ϕ . If the theory is exact, in the log-log plot in Figure 4a, the slopes of the three statistics should all be equal to 1. Figure 4a indeed shows that the three statistics of the Stage-I iteration count all have a slope of approximately 1, confirming the linear dependence of the high-probability upper bound on T_{basis} in Theorem 7.1 on ϕ . This also validates our approach of controlling the value of $\frac{\frac{1}{n} \sum_{i=1}^n (\hat{x}_i + \hat{s}_i)}{\min_{1 \leq i \leq n} (\hat{x}_i + \hat{s}_i)}$ to generate challenging LP instances for rPDHG (and perhaps even other first-order LP methods). When $l = 10$, $\phi_l \geq 5 \times 10^5$ and the corresponding iteration count is likely to be higher than 10^7 . These instances are significantly more difficult than the LP instances studied in Section 6.2 for the high-probability performance of rPDHG. See the data points in Figure 2 along the horizontal axis of n near 10^2 for comparison.

Additionally, Theorem 7.1 predicts that the Stage-II iteration count grows at most logarithmically with the value of $\min_{1 \leq i \leq n} (x_i^* + s_i^*)$ in a high-probability sense. This relationship can indeed be observed from Figure 4b, particularly for $\min_{1 \leq i \leq n} (x_i^* + s_i^*) \geq 10^{-3}$. Theorem 7.1 also indicates that T_{local} is likely to be 0 when $\min_{1 \leq i \leq n} (x_i^* + s_i^*)$ is smaller than or equal to the tolerance ε . Indeed, Figure 4b shows that when $\min_{1 \leq i \leq n} (x_i^* + s_i^*) \leq 2 \cdot 10^{-4}$, all quartiles of the Stage-II iteration count become zero, suggesting that a sufficiently accurate solution is found even before the iterates settle on the optimal basis. However, in our instances, a smaller value of $\min_{1 \leq i \leq n} (x_i^* + s_i^*)$ does not indicate an easier problem, as it leads to larger values of l and the corresponding ϕ_l (see (7.5) and (7.6)).

In the remainder of this section, we prove Theorem 7.1.

FIGURE 4. The median values and the interquartile range (IQR) of Stage-I and Stage-II iteration counts for the random LP instances with various values of ϕ and $\min_{1 \leq i \leq n} (x_i^* + s_i^*)$.



7.2. Proof of Theorem 7.1 We first derive a conditional high-probability upper bound on $\kappa\Phi$.

LEMMA 7.1. *For the random LP defined in Definition 3.4, for any $\delta \in (0, 1]$, it holds that*

$$\Pr \left[\kappa\Phi \leq \frac{c_3 n^{2.5} m^{0.5} \phi}{\delta(d+1)} \mid E_{\text{uniq}} \right] \geq 1 - \delta - 4e^{-c_1 m} - \left(\frac{1}{2}\right)^{d+1}.$$

Proof. For a fixed $\delta \in (0, 1]$, let $G := \left\{ \kappa \|B^{-1}\| \|A\| \leq \frac{c_3 n^{1.5} m^{0.5}}{\delta(d+1)} \right\} \cap \{u > 0\}$. Since $u > 0$ almost surely, the tail bound (4.22) gives

$$\Pr[G] \geq 1 - \delta - 4e^{-c_1 m} - \left(\frac{1}{2}\right)^{d+1}.$$

Recall the convention following (4.7) that $\|B^{-1}\| := \infty$ when B is not full-rank. Thus, on G , B is full-rank and $(\hat{x}, \hat{s})$ satisfies strict complementary slackness. Lemma 3.1 then implies that $G \subseteq E_{\text{uniq}}$ and that $x^* = \hat{x}$ and $s^* = \hat{s}$ on G . Hence, by (4.7) and (7.1), $\varphi = n\phi$ on G . The definition of G and Lemma 4.6 therefore yield

$$\kappa\Phi \leq \kappa \|B^{-1}\| \|A\| \varphi \leq \frac{c_3 n^{2.5} m^{0.5} \phi}{\delta(d+1)}$$

on G . Moreover, Lemmas 3.1 and 3.2 imply $\Pr[E_{\text{uniq}}] \geq 1 - e^{-\nu m} > 0$. Consequently,

$$\Pr \left[\kappa\Phi \leq \frac{c_3 n^{2.5} m^{0.5} \phi}{\delta(d+1)} \mid E_{\text{uniq}} \right] \geq \frac{\Pr[G]}{\Pr[E_{\text{uniq}}]} \geq \Pr[G] \geq 1 - \delta - 4e^{-c_1 m} - \left(\frac{1}{2}\right)^{d+1}. \quad (7.7)$$

where the second inequality uses $0 < \Pr[E_{\text{uniq}}] \leq 1$. $\square$

Proof of Theorem 7.1. By (4.5), $T_{\text{basis}} \leq \check{c}_1 \kappa\Phi \ln(\kappa\Phi)$ for an absolute constant $\check{c}_1 > 1$. Furthermore, $\kappa\Phi \geq 1$ by (4.4) and $\kappa \geq 1$. Since $z \mapsto z \ln z$ is increasing on $[1, \infty)$, Lemma 7.1 proves (7.2) with $C_0 = c_1$ and $C_1 = \check{c}_1 c_3$.

For the local convergence bound, Lemma 4.8 gives $\Pr[\|B^{-1}\| \|A\| \leq c_2 \sqrt{mn}/\delta] \geq 1 - \delta - 2e^{-c_1 m}$. This event implies that B is full-rank, and its intersection with $\{u > 0\}$ is therefore contained in E_{uniq} . On this

intersection, (4.6) gives the bound inside (7.3) with $C_2 = \check{c}_2 c_2$. Thus, the conditional probability in (7.3) is at least

$$\frac{\Pr \left[\|B^{-1}\| \|A\| \leq \frac{c_2 \sqrt{mn}}{\delta}, u > 0 \right]}{\Pr[E_{\text{uniq}}]} \geq 1 - \delta - 2e^{-c_1 m}.$$

Taking $C_0 = c_1$, the constants C_0, C_1, C_2 depend only (and at most polynomially) on σ_A . $\square$

Appendix A: Restarted PDHG Method for Linear Programming Algorithm 1 presents the framework of rPDHG. It requires no matrix factorizations. Line 4 of Algorithm 1 is an iteration of the vanilla PDHG

Algorithm 1: rPDHG: restarted-PDHG

```

1 Input: Initial iterate  $(x^{0,0}, y^{0,0}) = (0, 0)$ ,  $\ell \leftarrow 0$ ,  $k \leftarrow 0$ , step sizes  $\tau, \sigma$ ;
2 repeat
3   repeat
4     conduct one step of PDHG:  $(x^{\ell,k+1}, y^{\ell,k+1}) \leftarrow \text{ONEPDHG}(x^{\ell,k}, y^{\ell,k})$ ;
5     compute the average iterate:  $(\bar{x}^{\ell,k+1}, \bar{y}^{\ell,k+1}) \leftarrow \frac{1}{k+1} \sum_{i=1}^{k+1} (x^{\ell,i}, y^{\ell,i})$ ;
6      $k \leftarrow k + 1$ ;
7   until satisfying the restart condition;
8   restart the outer loop:  $(x^{\ell+1,0}, y^{\ell+1,0}) \leftarrow (\bar{x}^{\ell,k}, \bar{y}^{\ell,k})$ ,  $\ell \leftarrow \ell + 1$ ,  $k \leftarrow 0$ ;
9 until  $(x^{\ell,0}, y^{\ell,0})$  satisfies some convergence condition;
10 Output:  $(x^{\ell,0}, y^{\ell,0})$ 

```

defined in (2.3). Line 5 can be computed efficiently by updating $(\bar{x}^{\ell,k}, \bar{y}^{\ell,k})$ incrementally. The restart condition in Line 7 follows the β -restart criterion of Applegate et al. [8], which triggers when the “normalized duality gap” of $(\bar{x}^{\ell,k}, \bar{y}^{\ell,k})$ reduces to a β fraction of $(x^{\ell,0}, y^{\ell,0})$ ’s normalized duality gap. For the normalized duality gap’s precise definition, see (4a) of Applegate et al. [8] for general primal-dual first-order methods and Definition 2.1 of Xiong and Freund [58] for PDHG in conic linear optimization. The normalized duality gap measures the violations of feasibility and optimality, and $(x^{\ell,0}, y^{\ell,0})$ becomes approximately optimal when its normalized duality gap is sufficiently small.

In practice, the restart condition is checked periodically (every few hundred iterations), and the normalized duality gap can be easily approximated by a separable norm variant within an absolute constant factor. See Section 6 of Applegate et al. [8] and Appendix A of Xiong and Freund [58] for its efficient approximation. This approach is implemented in Applegate et al. [8], Lu and Yang [32] and our experiments in Section 6.

Following Theorem 3.1 of Xiong [56], we set the step sizes as: $\tau = \frac{\lambda_{\min}}{2\lambda_{\max}}$ and $\sigma = \frac{1}{2\lambda_{\min}\lambda_{\max}}$, where $\lambda_{\max}$ and $\lambda_{\min}$ denote the largest and smallest nonzero singular values of A , respectively. The restart parameter β is set to $\frac{1}{e}$, where e is the base of the natural logarithm. The largest singular value $\lambda_{\max}$ can be efficiently computed via power iterations. Using an imprecise estimate of $\lambda_{\min}$ is equivalent to applying rPDHG to a scalar-rescaled LP problem, so the method still converges to optimal solutions (Applegate et al. [8], Xiong [56]). rPDHG with the above parameter setup has been studied in analyses of Applegate et al. [8], Xiong [56], Xiong and Freund [57, 58, 59].

Appendix B: Proof of Lemma 3.1

Proof of Lemma 3.1. First, observe that $\hat{x}$ and $\hat{s}$ are the optimal primal-dual solutions, so $\hat{x} \in \mathcal{X}^*$ and $\hat{s} \in \mathcal{S}^*$.

For sufficiency, suppose B is full-rank. Then $u^1 \in \mathbb{R}^m$, $u^2 \in \mathbb{R}^d$, $u^1 > 0$ and $u^2 > 0$ imply that $\hat{x}$ and $\hat{s}$ are both nondegenerate optimal solutions. Furthermore, since rows of A are linearly independent, the instance has unique primal and dual optimal solutions: $\mathcal{X}^* = \{\hat{x}\}$ and $\mathcal{S}^* = \{\hat{s}\}$. Moreover, the corresponding y satisfying $A^\top y + \hat{s} = c$ is also unique. Thus, $\mathcal{X}^*$, $\mathcal{Y}^*$ and $\mathcal{S}^*$ are all singletons with $\mathcal{X}^* = \{\hat{x}\}$ and $\mathcal{S}^* = \{\hat{s}\}$.

For necessity, we consider two cases. If $(\hat{x}, \hat{s})$ does not satisfy strict complementary slackness, at least a pair of strictly complementary optimal solutions exists. Alternatively, if B is not full-rank but $u^1 > 0$ and $u^2 > 0$, then any solution $\tilde{x} = \begin{pmatrix} \tilde{x}_{[m]} \\ 0 \end{pmatrix}$ satisfying $B\tilde{x}_{[m]} = Bu^1$ and $\tilde{x}_{[m]} \geq 0$ is also optimal, where $[m]$ denotes $\{1, 2, \dots, m\}$. Multiple such $\tilde{x}_{[m]}$ exist because B is not full-rank and $u^1 > 0$. Therefore, in either case, the optimal solution sets $\mathcal{X}^*$ and $\mathcal{S}^*$ cannot both be singletons with $\mathcal{X}^* = \{\hat{x}\}$ and $\mathcal{S}^* = \{\hat{s}\}$. $\square$

Appendix C: Proof of Section 4.1

C.1. Proof of Lemma 4.3 We actually have the following more general result.

LEMMA C.1. *Let A be a random matrix as defined in Definition 3.2, where $A \in \mathbb{R}^{m \times n}$ and $n \geq m$. Then, for all $t \geq 10\sigma_A$, we have:*

$$\Pr[\sigma_1(A) \geq t\sqrt{n}] \leq \exp\left(-\frac{t^2 n}{16\sigma_A^2}\right). \quad (\text{C.1})$$

Proof. We apply the argument of Proposition 2.3 of Rudelson and Vershynin [41] to $A^\top$, so their dimensions (N, n) correspond to our (n, m) . First, replace both 1/2-nets in their proof by 1/4-nets. The volume comparison in the proof of their Proposition 2.1 gives cardinality bounds 9^n and 9^m , respectively.

Next, by applying Lemma 4.1 to both tails, the constants c_1 and C_1 in the second displayed equation of their proof can be taken as $1/(2\sigma_A^2)$ and 2, respectively, because the squared coefficients in the bilinear form sum to one for any fixed pair of unit vectors.

Finally, with the modified nets, Proposition 2.2 therein gives an approximation factor $(1 - 1/4)^{-2} = 16/9 < 2$. Thus, in the final union-bound step, we apply the fixed-pair tail bound at threshold $t\sqrt{n}/2$, rather than $t\sqrt{n}$. This replaces the last displayed estimate of their proof by

$$\Pr[\sigma_1(A) \geq t\sqrt{n}] \leq 2 \cdot 9^n \cdot 9^m \exp\left(-\frac{t^2 n}{8\sigma_A^2}\right) \leq \exp\left(-\frac{t^2 n}{16\sigma_A^2}\right), \quad (\text{C.2})$$

where the last inequality follows from $\ln 2 + (m+n) \ln 9 \leq n \ln 162 < \frac{25n}{4} \leq \frac{t^2 n}{16\sigma_A^2}$ for $m \leq n$ and $t \geq 10\sigma_A$. This proves the lemma. $\square$

Now we can prove Lemma 4.3.

Proof of Lemma 4.3. Setting $t = 10\sigma_A$ in (C.1) gives $\Pr[\sigma_1(A) \geq 10\sigma_A\sqrt{n}] \leq e^{-25n/4} \leq e^{-6n}$. $\square$

C.2. Proof of Lemma 4.4 We first present a technical result before proving Lemma 4.4.

LEMMA C.2. *For any $y \in (0, \frac{1}{e}]$, if there exists $x \in (0, \frac{1}{e}]$ such that $x \ln(\frac{1}{x}) = y$, then $x \geq \frac{y}{2 \ln(1/y)}$.*

Proof. Define $f(x) := x \cdot \ln(1/x)$ and $x_0 := \frac{y}{2 \ln(1/y)}$. Then we have $f(x_0) = \frac{y}{2 \ln(1/y)} \cdot \ln(2 \ln(1/y)/y) = \frac{y}{2 \ln(1/y)} \cdot [\ln(2 \ln(1/y)) + \ln(1/y)] = \frac{y}{2} + \frac{y}{2} \cdot \frac{\ln(2 \ln(1/y))}{\ln(1/y)}$. Note that as $y \in (0, 1/e]$, $2 \ln(1/y) - 1/y \leq \max_{t \geq e} [2 \ln(t) - t] = 2 - e < 0$, so $\frac{\ln(2 \ln(1/y))}{\ln(1/y)} \leq 1$ and thus $f(x_0) \leq \frac{y}{2} + \frac{y}{2} = y$. Finally, since $f(x) := x \cdot \ln(1/x)$ is monotonically increasing on $(0, 1/e]$, we have $x_0 \leq x$. This proves the statement. $\square$

Now we prove Lemma 4.4.

Proof of Lemma 4.4. First of all, we prove that for all $T > C_1 C_2$:

$$\Pr[XY \geq T] \leq \frac{2C_1 C_2}{T} + \frac{C_1 C_2}{T} \cdot \ln\left(\frac{T}{C_1 C_2}\right) + \delta_1 + 2\delta_2. \quad (\text{C.3})$$

For a random variable η , let F_η denote the distribution function of η . The probability $\Pr[XY \geq T]$ has the following upper bound:

$$\Pr[XY \geq T] = \int_0^\infty \Pr\left[X \geq \frac{T}{y}\right] dF_Y(y) \stackrel{(4.2)}{\leq} \underbrace{\delta_1 + \int_0^\infty \min\left\{\frac{C_1 y}{T}, 1\right\} \cdot dF_Y(y)}_{\text{Denoted by I}}, \quad (\text{C.4})$$

in which

$$\mathbf{I} = \underbrace{\int_0^{\frac{T}{C_1}} \frac{C_1 y}{T} \cdot dF_Y(y)}_{\text{Denoted by } \mathbf{II}} + \int_{\frac{T}{C_1}}^{\infty} dF_Y(y) = \mathbf{II} + \Pr \left[Y \geq \frac{T}{C_1} \right] \stackrel{(4.2)}{\leq} \mathbf{II} + \frac{C_1 C_2}{T} + \delta_2. \quad (\text{C.5})$$

The term $\mathbf{II}$ is $\frac{C_1}{T} \int_0^{\frac{T}{C_1}} y \cdot dF_Y(y)$, and $\frac{T}{C_1} \cdot \mathbf{II}$ has the following upper bound:

$$\begin{aligned} \frac{T}{C_1} \cdot \mathbf{II} &= \int_0^{\frac{T}{C_1}} y \cdot dF_Y(y) = y \cdot F_Y(y) \Big|_0^{\frac{T}{C_1}} - \int_0^{\frac{T}{C_1}} F_Y(y) \cdot dy \\ &= \frac{T}{C_1} \cdot \Pr \left[Y \leq \frac{T}{C_1} \right] - \int_0^{\frac{T}{C_1}} (1 - \Pr[Y > y]) dy \leq \frac{T}{C_1} - \frac{T}{C_1} + \int_0^{\frac{T}{C_1}} \Pr[Y > y] dy \\ &\stackrel{(4.2)}{\leq} \int_0^{\frac{T}{C_1}} \min \left\{ \frac{C_2}{y} + \delta_2, 1 \right\} dy \leq \frac{T}{C_1} \cdot \delta_2 + \int_0^{\frac{T}{C_1}} \min \left\{ \frac{C_2}{y}, 1 \right\} dy \\ &\leq \frac{T}{C_1} \cdot \delta_2 + \int_0^{C_2} 1 dy + \int_{C_2}^{\frac{T}{C_1}} \frac{C_2}{y} dy = \frac{T}{C_1} \cdot \delta_2 + C_2 + C_2 \ln \left(\frac{T}{C_1 C_2} \right). \end{aligned} \quad (\text{C.6})$$

In other words,

$$\mathbf{II} \leq \frac{C_1}{T} \left(\frac{T}{C_1} \cdot \delta_2 + C_2 + C_2 \ln \left(\frac{T}{C_1 C_2} \right) \right) = \delta_2 + \frac{C_1 C_2}{T} + \frac{C_1 C_2}{T} \cdot \ln \left(\frac{T}{C_1 C_2} \right). \quad (\text{C.7})$$

Finally, substituting the upper bound of $\mathbf{II}$ in (C.7) back into (C.5) and (C.4) yields (C.3).

We are particularly interested in (C.3) when $T \geq 3C_1 C_2$, as otherwise the bound is trivial. We use δ_0 to denote $\frac{C_1 C_2}{T}$ (which takes values in $(0, 1/3]$), and (C.3) reduces to $\Pr \left[XY \geq \frac{C_1 C_2}{\delta_0} \right] \leq 3\delta_0 \cdot \ln \left(\frac{1}{\delta_0} \right) + \delta_1 + 2\delta_2$. We let $\hat{\delta} := \delta_0 \cdot \ln(1/\delta_0)$, and $\hat{\delta}$ may take any value in $(0, \frac{\ln(3)}{3}]$. Using Lemma C.2, δ_0 is at least as large as $\frac{\hat{\delta}}{2 \ln(1/\hat{\delta})}$. Therefore,

$$\Pr \left[XY \geq \frac{2C_1 C_2 \cdot \ln(1/\hat{\delta})}{\hat{\delta}} \right] \leq \Pr \left[XY \geq \frac{C_1 C_2}{\delta_0} \right] \leq 3\hat{\delta} + \delta_1 + 2\delta_2. \quad (\text{C.8})$$

Finally, we replace $3\hat{\delta}$ by δ (which may take any value in $(0, 1]$) and (C.8) then becomes the desired inequality (4.3). $\square$

Appendix D: Proof of Section 5.1

D.1. Proof of Lemma 5.1

Proof of Lemma 5.1. According to the min-max principle for singular values, the k -th largest singular value $\sigma_k(W)$ has the following equivalent representation:

$$\sigma_k(W) = \max_{S: \dim(S)=k} \min_{x \in S, \|x\|=1} \|Wx\|. \quad (\text{D.1})$$

Let $S_k := \{x \in \mathbb{R}^n : x_{k+1} = x_{k+2} = \dots = x_n = 0\}$ denote the k -dimensional subspace of vectors with zeros in the last $n - k$ components. Since $\dim(S_k) = k$, (D.1) implies that $\sigma_k(W)$ has the following lower bound:

$$\sigma_k(W) \geq \min_{x \in S_k, \|x\|=1} \|Wx\| = \min_{y \in \mathbb{R}^k, \|y\|=1} \|W_{[k]}y\| = \sigma_k(W_{[k]}) \quad (\text{D.2})$$

where $W_{[k]}$ denotes the first k columns of W . Thus, $\sigma_k(W)$ is bounded below by the smallest singular value of $W_{[k]}$.

Since $W_{[k]}$ is also a sub-Gaussian matrix, it further holds that

$$(C_{rv1}\delta_0)^{n-k+1} + e^{-C_{rv2}n} \geq \Pr\left[\sigma_k(W_{[k]}) < \delta_0(\sqrt{n} - \sqrt{k-1})\right] \geq \Pr\left[\sigma_k(W_{[k]}) < \delta_0 \cdot \frac{n-k+1}{2\sqrt{n}}\right] \quad (D.3)$$

for all $\delta_0 > 0$, where the first inequality follows from Lemma 4.2 applied to $W_{[k]}$ and the second inequality is due to $\sqrt{n} - \sqrt{k-1} = \frac{n-k+1}{\sqrt{n}+\sqrt{k-1}} \geq \frac{n-k+1}{2\sqrt{n}}$. Therefore, this yields the following bound on all singular values for all $\delta_0 > 0$:

$$\begin{aligned} & \Pr\left[\sigma_k(W) \geq \frac{\delta_0(n-k+1)}{2\sqrt{n}} \text{ for } k=1, 2, \dots, n\right] \\ & \geq 1 - \sum_{k=1}^n \Pr\left[\sigma_k(W) < \frac{\delta_0(n-k+1)}{2\sqrt{n}}\right] \stackrel{(D.2)}{\geq} 1 - \sum_{k=1}^n \Pr\left[\sigma_k(W_{[k]}) < \frac{\delta_0(n-k+1)}{2\sqrt{n}}\right] \\ & \stackrel{(D.3)}{\geq} 1 - \sum_{k=1}^n (C_{rv1}\delta_0)^{n-k+1} - n \cdot e^{-C_{rv2}n} \geq 1 - 2C_{rv1}\delta_0 - n \cdot e^{-C_{rv2}n}. \end{aligned} \quad (D.4)$$

The final inequality holds because for $\delta_0 \in (0, \frac{1}{2C_{rv1}})$, $\sum_{k=1}^n (C_{rv1}\delta_0)^k \leq \sum_{k=1}^{\infty} (C_{rv1}\delta_0)^k = \frac{C_{rv1}\delta_0}{1-C_{rv1}\delta_0} \leq 2C_{rv1}\delta_0$. Replacing $2C_{rv1}\delta_0$ in (D.4) with δ establishes (5.1) for all $\delta \in (0, 1)$. As for the case $\delta \geq 1$, (5.1) is trivial. $\square$

D.2. Proof of Lemma 5.3

Proof of Lemma 5.3. Let M denote $W^{-\top}W^{-1}$. Then $\|W^{-1}v\|^2 = v^\top W^{-\top}W^{-1}v = v^\top Mv$. We will first prove the following tail bound of $v^\top Mv$ for all $\gamma \geq 0$:

$$\Pr\left[v^\top Mv \geq \mathbb{E}[v^\top Mv] + \frac{2\gamma\sigma^2n}{c_0^2}\right] \leq 2 \cdot \exp\left(-2C_{hw} \min\{\gamma^2, \gamma\}\right). \quad (D.5)$$

To establish (D.5), we apply Lemma 5.2, which requires bounds on $\|M\|$ and $\|M\|_F$. We begin by characterizing the singular values of M . For $k=1, 2, \dots, n$:

$$\sigma_k(M) = \sigma_k(W^{-\top}W^{-1}) = \frac{1}{\sigma_{n+1-k}^2(W)} \quad (D.6)$$

Then we have the following upper bounds for $\|M\|_F^2$ and $\|M\|$:

$$\|M\| = \sigma_1(M) \stackrel{(D.6)}{=} \frac{1}{\sigma_n^2(W)} \stackrel{(5.3)}{\leq} \frac{n}{c_0^2} \quad (D.7)$$

$$\|M\|_F^2 = \sum_{i=1}^n \sigma_i^2(M) \stackrel{(D.6)}{=} \sum_{i=1}^n \frac{1}{\sigma_i^4(W)} \stackrel{(5.3)}{\leq} \sum_{i=1}^n \frac{n^2}{c_0^4 i^4} \leq \frac{n^2}{c_0^4} \sum_{i=1}^{\infty} \frac{1}{i^4} = \frac{n^2}{c_0^4} \cdot \frac{\pi^4}{90} \leq \frac{2n^2}{c_0^4}. \quad (D.8)$$

where the last equality uses the definition of the Riemann zeta function and its value at 4. With the above upper bounds of $\|M\|_F^2$ and $\|M\|$, now we can use Lemma 5.2:

$$\Pr\left[v^\top Mv \geq \mathbb{E}[v^\top Mv] + t\right] \leq 2 \cdot \exp\left[-C_{hw} \cdot \min\left(\frac{c_0^4 t^2}{2\sigma^4 n^2}, \frac{c_0^2 t}{\sigma^2 n}\right)\right]. \quad (D.9)$$

For any $\gamma \geq 0$, set $t = \frac{2\gamma\sigma^2n}{c_0^2}$. Substituting this choice into (D.9) yields:

$$\begin{aligned} \Pr\left[v^\top Mv \geq \mathbb{E}[v^\top Mv] + \frac{2\gamma\sigma^2n}{c_0^2}\right] & \leq 2 \cdot \exp\left[-C_{hw} \cdot \min\left(\frac{c_0^4}{2\sigma^4 n^2} \cdot \frac{4\gamma^2\sigma^4 n^2}{c_0^4}, \frac{c_0^2}{\sigma^2 n} \cdot \frac{2\gamma\sigma^2n}{c_0^2}\right)\right] \\ & = 2 \cdot \exp\left[-2C_{hw} \cdot \min\{\gamma^2, \gamma\}\right]. \end{aligned} \quad (D.10)$$

This proves (D.5).

We now use (D.5) to derive the tail bound of $v^\top M v$. We start from the following expression of $E[v^\top M v]$:

$$E[v^\top M v] = E[\|W^{-1} v\|^2] = \sum_{i=1}^n E[(W^{-1} v)_i^2] = \sum_{i=1}^n \|(W^{-1})_{i,\cdot}\|^2 = \|W^{-1}\|_F^2 = \sum_{i=1}^n \frac{1}{\sigma_i^2(W)} \quad (D.11)$$

where the third equality is because components of v are independent and have zero mean and unit variance. Furthermore, $E[v^\top M v]$ is upper bounded as follows:

$$E[v^\top M v] \stackrel{(D.11)}{=} \sum_{i=1}^n \frac{1}{\sigma_i^2(W)} \stackrel{(5.3)}{\leq} \sum_{i=1}^n \frac{n}{c_0^2 i^2} \leq \frac{n}{c_0^2} \sum_{i=1}^{\infty} \frac{1}{i^2} = \frac{n}{c_0^2} \cdot \frac{\pi^2}{6} \leq \frac{2n}{c_0^2}, \quad (D.12)$$

where the last equality uses the definition of the Riemann zeta function and its value at 2. Therefore, since $\|W^{-1} v\|^2 = v^\top M v$, for all $t > 0$:

$$\Pr \left[\|W^{-1} v\|^2 \geq \frac{2n}{c_0^2} + t \right] \leq \Pr [v^\top M v \geq E[v^\top M v] + t]. \quad (D.13)$$

Finally, letting $t = \frac{2\gamma\sigma^2 n}{c_0^2}$ and applying (D.5) to this inequality completes the proof. $\square$

Acknowledgments

The author thanks Yinyu Ye, Moïse Blanchard and Haihao Lu for the helpful discussions. The author is also grateful to Robert M. Freund for the valuable suggestions while preparing this manuscript. The author also thanks the three anonymous referees for helpful comments and suggestions that improved the quality of the manuscript. Part of this work was performed at the Massachusetts Institute of Technology and was supported by AFOSR Grant No. FA9550-22-1-0356. Another part of this work was performed at Georgia Institute of Technology and was supported by ONR Award #N00014-25-1-2088.

References

- [1] Adler I, Karp R, Shamir R (1986) A family of simplex variants solving an $m \times d$ linear program in expected number of pivot steps depending on d only. *Mathematics of Operations Research* 11(4):570–590.
- [2] Adler I, Karp RM, Shamir R (1987) A simplex variant solving an $m \times d$ linear program in $O(\min(m^2, d^2))$ expected number of pivot steps. *Journal of Complexity* 3(4):372–387.
- [3] Adler I, Megiddo N (1985) A simplex algorithm whose average number of steps is bounded between two quadratic functions of the smaller dimension. *Journal of the ACM* 32(4):871–895.
- [4] Anagnostides I, Sandholm T (2024) Convergence of $\log(1/\epsilon)$ for gradient-based algorithms in zero-sum games without the condition number: A smoothed analysis. *Proceedings of the Advances in Neural Information Processing Systems* 37, volume 37, 125839–125878.
- [5] Anstreicher KM, Ji J, Potra FA, Ye Y (1993) Average performance of a self-dual interior-point algorithm for linear programming. Pardalos PM, ed., *Complexity in Numerical Optimization*, 1–15 (World Scientific).
- [6] Anstreicher KM, Ji J, Potra FA, Ye Y (1999) Probabilistic analysis of an infeasible-interior-point algorithm for linear programming. *Mathematics of Operations Research* 24(1):176–192.
- [7] Applegate D, Díaz M, Hinder O, Lu H, Lubin M, O’Donoghue B, Schudy W (2021) Practical large-scale linear programming using primal-dual hybrid gradient. *Proceedings of the Advances in Neural Information Processing Systems* 34, volume 34, 20243–20257.
- [8] Applegate D, Hinder O, Lu H, Lubin M (2023) Faster first-order primal-dual methods for linear programming using restarts and sharpness. *Mathematical Programming* 201(1-2):133–184.
- [9] Biele A, Gade D (2024) FICO® Xpress solver 9.4. <https://community.fico.com/s/blog-post/a5QQi00000019II5MAM/fico4824>. Accessed: August 1, 2026.

- [10] Blin N (2024) Accelerate large linear programming problems with NVIDIA cuOpt. <https://developer.nvidia.com/blog/accelerate-large-linear-programming-problems-with-nvidia-cuopt/>. Accessed: August 1, 2026.
- [11] Blum A, Dunagan J (2002) Smoothed analysis of the perceptron algorithm for linear programming. *Proceedings of the 13th Annual ACM-SIAM Symposium on Discrete Algorithms*, 905–914.
- [12] Borgwardt KH (1977) *Untersuchungen zur Asymptotik der mittleren Schrittzahl von Simplexverfahren in der Linearen Optimierung*. Ph.D. thesis, Universität Kaiserslautern, Kaiserslautern.
- [13] Borgwardt KH (1978) Zum Rechenaufwand von Simplexverfahren. *Operations Research Verfahren* 31:83–97.
- [14] Borgwardt KH (1982) The average number of pivot steps required by the simplex-method is polynomial. *Zeitschrift für Operations Research* 26:157–177.
- [15] Borgwardt KH (1982) Some distribution-independent results about the asymptotic order of the average number of pivot steps of the simplex method. *Mathematics of Operations Research* 7(3):441–462.
- [16] Cook N (2018) Lower bounds for the smallest singular value of structured random matrices. *The Annals of Probability* 46(6):3442–3500.
- [17] Dadush D, Huiberts S (2018) A friendly smoothed analysis of the simplex method. *Proceedings of the 50th Annual ACM SIGACT Symposium on Theory of Computing*, 390–403.
- [18] Dumitriu I, Zhu Y (2024) Extreme singular values of inhomogeneous sparse random rectangular matrices. *Bernoulli* 30(4):2904–2931.
- [19] Dunagan J, Spielman DA, Teng SH (2011) Smoothed analysis of condition numbers and complexity implications for linear programming. *Mathematical Programming* 126(2):315–350.
- [20] Ge D, Hu H, Huangfu Q, Liu J, Liu T, Lu H, Yang J, Ye Y, Zhang C (2024) cuPDL-C. <https://github.com/COPT-Public/cuPDL-C>. Accessed: August 1, 2026.
- [21] Güler O, Ye Y (1993) Convergence behavior of interior-point algorithms. *Mathematical Programming* 60(1-3):215–228.
- [22] Gurobi Optimization, LLC (2025) Gurobi releases version 13.0 with improved performance and new solving capabilities. <https://www.gurobi.com/company/newsroom/gurobi-releases-version-13-0-with-improved-performance-and-new-solving-capabilities>. Accessed: August 1, 2026.
- [23] Haimovich M (1983) The simplex algorithm is very good! On the expected number of pivot steps and related properties of random linear programs. Draft, Columbia University, 415 Uris Hall, New York.
- [24] HiGHS Development Team (2024) HiGHS v1.7.0 release notes. <https://github.com/ERGO-Code/HiGHS/releases/tag/v1.7.0>. Accessed: August 1, 2026.
- [25] Hinder O (2024) Worst-case analysis of restarted primal-dual hybrid gradient on totally unimodular linear programs. *Operations Research Letters* 57:107199.
- [26] Huang S (2000) Average number of iterations of some polynomial interior-point algorithms for linear programming. *Science in China A: Mathematics* 43(8):829–835.
- [27] Huang SM (2005) Expected number of iterations of interior-point algorithms for linear programming. *Journal of Computational Mathematics* 23(1):93–100.
- [28] Huang Y, Zhang W, Li H, Ge D, Liu H, Ye Y (2025) A restarted primal-dual hybrid conjugate gradient method for large-scale quadratic programming. *INFORMS Journal on Computing*.
- [29] Ji J, Potra FA (2008) On the probabilistic complexity of finding an approximate solution for linear programming. *Journal of Complexity* 24(2):214–227.
- [30] Klee V, Minty GJ (1972) How good is the simplex algorithm? Shisha O, ed., *Inequalities III*, 159–175 (Academic Press).
- [31] Lu H, Yang J (2024) Restarted Halpern PDHG for linear programming. *arXiv preprint arXiv:2407.16144*.
- [32] Lu H, Yang J (2025) cuPDL.jl: A GPU implementation of restarted primal-dual hybrid gradient for linear programming in Julia. *Operations Research* 73(6):3440–3452.

- [33] Lu H, Yang J (2025) On the geometry and refined rate of primal–dual hybrid gradient for linear programming. *Mathematical Programming* 212:349–387.
- [34] Lu H, Yang J (2026) A practical and optimal first-order method for large-scale convex quadratic programming. *Mathematical Programming* 215(1-2):771–808.
- [35] Lu H, Yang J, Hu H, Huangfu Q, Liu J, Liu T, Ye Y, Zhang C, Ge D (2023) cuPDLP-C: A strengthened implementation of cuPDLP for linear programming by C language. *arXiv preprint arXiv:2312.14832*.
- [36] Megiddo N (1986) Improved asymptotic analysis of the average number of steps performed by the self-dual simplex algorithm. *Mathematical Programming* 35:140–172.
- [37] Mizuno S, Todd MJ, Ye Y (1993) On adaptive-step primal-dual interior-point algorithms for linear programming. *Mathematics of Operations Research* 18(4):964–981.
- [38] Paquette C, Lee K, Pedregosa F, Paquette E (2021) SGD in the large: Average-case analysis, asymptotics, and stepsize criticality. *Proceedings of Thirty Fourth Conference on Learning Theory*, volume 134, 3548–3626 (PMLR).
- [39] Paquette C, van Merriënboer B, Paquette E, Pedregosa F (2023) Halting time is predictable for large models: A universality property and average-case analysis. *Foundations of Computational Mathematics* 23(2):597–673.
- [40] Pedregosa F, Scieur D (2020) Acceleration through spectral density estimation. *Proceedings of the 37th International Conference on Machine Learning*, volume 119, 7553–7562 (PMLR).
- [41] Rudelson M, Vershynin R (2009) Smallest singular value of a random rectangular matrix. *Communications on Pure and Applied Mathematics* 62(12):1707–1739.
- [42] Rudelson M, Vershynin R (2010) Non-asymptotic theory of random matrices: extreme singular values. *Proceedings of the International Congress of Mathematicians 2010*, volume III, 1576–1602 (World Scientific).
- [43] Scieur D, Pedregosa F (2020) Universal average-case optimality of Polyak momentum. *Proceedings of the 37th International Conference on Machine Learning*, 8565–8572 (PMLR).
- [44] Shamir R (1987) The efficiency of the simplex method: A survey. *Management Science* 33(3):301–334.
- [45] Smale S (1983) On the average number of steps of the simplex method of linear programming. *Mathematical Programming* 27(3):241–262.
- [46] Smale S (1983) The problem of the average speed of the simplex method. Bachem A, Korte B, Grötschel M, eds., *Mathematical Programming: The State of the Art*, 530–539 (Springer).
- [47] Spielman DA, Teng SH (2003) Smoothed analysis of termination of linear programming algorithms. *Mathematical Programming* 97:375–404.
- [48] Spielman DA, Teng SH (2004) Smoothed analysis of algorithms: Why the simplex algorithm usually takes polynomial time. *Journal of the ACM* 51(3):385–463.
- [49] Terlaky T (2009) Twenty-five years of interior point methods. *Decision Technologies and Applications*, 1–33 (INFORMS).
- [50] Todd MJ (1986) Polynomial expected behavior of a pivoting algorithm for linear complementarity and linear programming problems. *Mathematical Programming* 35(2):173–192.
- [51] Todd MJ (1991) Probabilistic models for linear programming. *Mathematics of Operations Research* 16(4):671–693.
- [52] Todd MJ, Ye Y (1990) A centered projective algorithm for linear programming. *Mathematics of Operations Research* 15(3):508–529.
- [53] Vershynin R (2018) *High-dimensional probability: An introduction with applications in data science* (Cambridge University Press), 1st edition.
- [54] Wainwright MJ (2019) *High-dimensional statistics: A non-asymptotic viewpoint* (Cambridge University Press).
- [55] Wright SJ (1994) An infeasible-interior-point algorithm for linear complementarity problems. *Mathematical Programming* 67(1-3):29–51.
- [56] Xiong Z (2026) Accessible complexity bounds for restarted PDHG on linear programs with a unique optimizer. *Mathematics of Operations Research*.

- [57] Xiong Z, Freund RM (2023) On the relation between LP sharpness and limiting error ratio and complexity implications for restarted PDHG. *arXiv preprint arXiv:2312.13773* .
- [58] Xiong Z, Freund RM (2024) The role of level-set geometry on the performance of PDHG for conic linear optimization. *arXiv preprint arXiv:2406.01942* .
- [59] Xiong Z, Freund RM (2026) Computational guarantees for restarted PDHG for LP based on “limiting error ratios” and LP sharpness. *Mathematical Programming* .
- [60] Ye Y (1992) On the finite convergence of interior-point algorithms for linear programming. *Mathematical Programming* 57(1-3):325–335.
- [61] Ye Y (1994) Toward probabilistic analysis of interior-point algorithms for linear programming. *Mathematics of Operations Research* 19(1):38–52.
- [62] Ye Y (1997) *Interior point algorithms: Theory and analysis* (Wiley-Interscience).
- [63] Ye Y (2011) The simplex and policy-iteration methods are strongly polynomial for the Markov decision problem with a fixed discount rate. *Mathematics of Operations Research* 36(4):593–603.
- [64] Zhang Y (1994) On the convergence of a class of infeasible interior-point methods for the horizontal linear complementarity problem. *SIAM Journal on Optimization* 4(1):208–227.
- [65] Zhang Y, Zhang D (1994) Superlinear convergence of infeasible-interior-point methods for linear programming. *Mathematical Programming* 66(1-3):361–377.