

# Solving Tall Dense Linear Programs in Nearly Linear Time

Jan van den Brand

KTH Royal Institute of Technology, Sweden  
janvdb@kth.se

Aaron Sidford

Stanford University, U.S.A.  
sidford@stanford.edu

Yin Tat Lee

University of Washington and MSR Redmond, U.S.A.  
yintat@uw.edu

Zhao Song

Princeton University and  
Institute for Advanced Study, U.S.A.  
zhaos@ias.edu

## ABSTRACT

In this paper we provide an  $\tilde{O}(nd + d^3)$  time randomized algorithm for solving linear programs with  $d$  variables and  $n$  constraints with high probability. To obtain this result we provide a robust, primal-dual  $\tilde{O}(\sqrt{d})$ -iteration interior point method inspired by the methods of Lee and Sidford (2014, 2019) and show how to efficiently implement this method using new data-structures based on heavy-hitters, the Johnson–Lindenstrauss lemma, and inverse maintenance. Interestingly, we obtain this running time without using fast matrix multiplication and consequently, barring a major advance in linear system solving, our running time is near optimal for solving dense linear programs among algorithms that do not use fast matrix multiplication.

## CCS CONCEPTS

• Theory of computation → Linear programming.

## KEYWORDS

linear program, interior point method, nearly linear time

## ACM Reference Format:

Jan van den Brand, Yin Tat Lee, Aaron Sidford, and Zhao Song. 2020. Solving Tall Dense Linear Programs in Nearly Linear Time. In *Proceedings of the 52nd Annual ACM SIGACT Symposium on Theory of Computing (STOC '20)*, June 22–26, 2020, Chicago, IL, USA. ACM, New York, NY, USA, 14 pages. <https://doi.org/10.1145/3357713.3384309>

The full version of this paper is available as [65] at <https://arxiv.org/pdf/2002.02304.pdf>.

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from [permissions@acm.org](mailto:permissions@acm.org).  
STOC '20, June 22–26, 2020, Chicago, IL, USA

© 2020 Copyright held by the owner/author(s). Publication rights licensed to ACM.  
ACM ISBN 978-1-4503-6979-4/20/06...\$15.00  
<https://doi.org/10.1145/3357713.3384309>

## 1 INTRODUCTION

Given  $A \in \mathbb{R}^{n \times d}$ ,  $b \in \mathbb{R}^d$ , and  $c \in \mathbb{R}^n$  solving a linear program (P) and its dual (D):

$$(P) = \min_{x \in \mathbb{R}_{\geq 0}^n : A^T x = b} c^T x \text{ and } (D) = \max_{y \in \mathbb{R}^d : Ay \geq c} b^T y. \quad (1)$$

is one of the most fundamental and well-studied problems in computer science and optimization. Developing faster algorithms for (1) has been the subject of decades of extensive research and the pursuit of faster linear programming methods has lead to numerous algorithmic advances and the advent of fundamental optimization techniques, e.g. simplex methods [18], ellipsoid methods [30], and interior-point methods (IPMs) [29].

The current fastest algorithms for solving (1) are the IPMs of Lee and Sidford [34] and Cohen, Lee, and Song [12]. These results build on a long line of work on IPMs [3, 29, 47, 52, 61, 63, 64], fast matrix multiplication [14, 15, 21, 22, 58–60, 66], and linear system solvers [9, 11, 39, 42, 44]. The first, [34] combined an  $\tilde{O}(\sqrt{d})$  iteration IPM from [33, 35] with new techniques for *inverse maintenance*, i.e. maintaining an approximate inverse of a slowly changing matrix, to solve (1) in time  $\tilde{O}(\text{nnz}(A)\sqrt{d} + d^{2.5})$  with high probability. For sufficiently large values of  $\text{nnz}(A)$  or  $n$ , this is the fastest known running time for solving linear programming up to polylogarithmic factors.

The second, [12], developed a stable and robust version of the IPM of [52] (using techniques from [33, 35]) and combined it with novel randomization, data structure, and rectangular matrix multiplication [22] techniques to solve (1) in time  $\tilde{O}(\max\{n, d\}^\omega)$  with high probability where  $\omega < 2.373$  is the current best known matrix multiplication constant [21, 66]. When  $n = \tilde{\Theta}(d)$ , this running time matches that of the best known linear system solvers for solving  $n \times n$  linear systems and therefore is the best possible barring a major linear system solving advance.

Though, these results constitute substantial advances in algorithmic techniques for linear programming, the running times of [34] and [12] are incomparable and neither yield a nearly linear running time when  $n$  grows polynomially with  $d$ , i.e. when  $d = n^\delta$  for any  $\delta > 0$  neither [34] or [12] yields a nearly linear running time. Consequently, it has remained a fundamental open problem to determine whether or not it is possible to solve high-dimensional linear programs to high-precision in nearly linear time for any polynomial ratio of  $n, d$ , and  $\text{nnz}(A)$ . Though achieving such a nearly linear runtime is known in the simpler setting of linear regression

[9–11, 39, 42, 44], achieving analogous results for solving such tall linear programs has been elusive.

In this paper we provide the first such nearly linear time algorithm for linear programming. We provide an algorithm which solves (1) with high probability in time  $\tilde{O}(nd + d^3)$ . Whenever  $A$  is dense, i.e.  $\text{nnz}(A) = \Omega(nd)$ , and sufficiently tall, i.e.  $n = \Omega(d^2)$ , this constitutes a nearly linear running time. In contrast to previous state-of-the-art IPMs for linear programming [12, 34] our algorithm doesn't need to use fast matrix multiplication and thereby matches the best known running time for regression on dense matrices which does not use fast matrix multiplication.

To achieve this result, we introduce several techniques that we believe are independent interest. First, we consider the IPM of [35] and develop an efficiently implementable robust primal-dual variant of it in the style of [5, 12, 37] which only requires solving  $\tilde{O}(\sqrt{d})$  linear systems. Prior to [33, 35], obtaining such an efficient  $\tilde{O}(\sqrt{d})$ -iteration IPM was a major open problem in the theory of IPMs. We believe our primal dual method and its analysis is simpler than that of [35]; by developing a primal-dual method we eliminate the need for explicit  $\ell_p$ -Lewis weight computation for  $p \neq 2$  as in [35] and instead work with the simpler  $\ell_2$ -Lewis weights or leverage scores. Further, we show that this primal dual method is highly stable and can be implemented efficiently given only multiplicative estimates of the variables, Hessians, and leverage scores.

With this IPM in hand, the problem of achieving our desired running time reduces to implementing this IPM efficiently. This problem is that of maintaining multiplicative approximations to vectors (the current iterates), leverage scores (a measure of importance of the rows under local rescaling), and the inverse of matrix (the system one needs to solve to take a step of the IPM) under small perturbations. While variants of vector maintenance have been considered recently [5, 12, 37] and inverse maintenance is well-studied historically [5, 12, 29, 33–35, 37, 48, 49, 61], none of these methods can be immediately applied in our setting where we cannot afford to pay too much in terms of  $n$  each iteration.

Our second contribution is to show that these data-structure problems can be solved efficiently. A key technique we use to overcome these issues is heavy-hitters sketching. We show that it is possible to apply a heavy hitters sketch (in particular [26, 50]) to the iterates of the method such that we can efficiently find changes in the coordinates. This involves carefully sketching groups of updates and dynamically modifying the induced data-structure. These sketches only work against non-adaptive adversaries, and therefore care is need to ensure that the sketch is used only to save time and not affect the progression of the overall IPM in a way that break this non-adaptive assumption. To achieve this we use the sketches to propose short-lists of possible changes which we then filter to ensure that the output of the datastructure is deterministic (up to a low probability failure event). Coupling this technique with known Johnson-Lindenstrauss sketches yields our leverage score maintenance data structure and adapting and simplifying previous inverse maintenance techniques yields our inverse maintenance data structure. We believe this technique of sketching the central path is powerful and may find further applications.

## 1.1 Our Results

The main result of this paper is the following theorem for solving (1). This algorithm's running time is nearly linear whenever the LP is sufficiently dense and tall, i.e.  $\text{nnz}(A) = \tilde{\Omega}(nd)$  and  $n = \tilde{\Omega}(d^2)$ . This is the first polynomial time algorithm with a nearly linear running time for high-dimensional instances (i.e. when  $d$  can grow polynomially with  $n$ ). This algorithm does not use fast matrix multiplication (FMM) and consequently its running time matches the best known running time for checking whether there even exists  $x$  such that  $A^\top x = b$  for dense  $A$  without using FMM.

**THEOREM 1.1 (MAIN RESULT).** *There is an algorithm (Algorithm 2) which given any linear program of the form (1) for non-degenerate<sup>1</sup>  $A \in \mathbb{R}^{n \times d}$  and  $\delta \in [0, 1]$  computes a point  $x \in \mathbb{R}_{\geq 0}^n$  such that*

$$c^\top x \leq \min_{A^\top x = b, x \geq 0} c^\top x + \delta \cdot \|c\|_2 \cdot R \text{ and } \|A^\top x - b\|_2 \leq \delta \cdot (\|A\|_F \cdot R + \|b\|_2)$$

where  $R$  is the diameter of the polytope in  $\ell_2$  norm, i.e.  $\|x\|_2 \leq R$  for all  $x \in \mathbb{R}_{\geq 0}^n$  with  $A^\top x = b$ . Further, the expected running time of the method is  $O((nd + d^3) \log^{O(1)} n \log \frac{n}{\delta})$ .

*Remark.* See [33, 52] on the discussion on converting such an approximate solution to an exact solution. For integral  $A, b, c$ , it suffices to pick  $\delta = 2^{-O(L)}$  to get an exact solution where  $L = \log(1 + d_{\max} + \|c\|_\infty + \|b\|_\infty)$  is the bit complexity and  $d_{\max}$  is the largest absolute value of the determinant of a square sub-matrix of  $A$ . For many combinatorial problems  $L = O(\log(n + \|b\|_\infty + \|c\|_\infty))$ .

Beyond this result (Section 4) we believe our robust primal-dual  $\tilde{O}(\sqrt{d})$ -iteration IPM and our data structures for maintaining multiplicative approximations to vectors, leverage scores, and the inverse of matrices are of independent interest.

## 1.2 Previous Work

Linear programming has been the subject of extensive research for decades and it is impossible to completely cover to this impressive line of work in this short introduction. Here we cover results particularly relevant to our approach. For more detailed coverage of prior-work see, e.g. [47, 68].

**IPMs:** The first proof of a polynomial time IPM was due to Karmarkar in [29]. After multiple running time improvements [29, 48, 49, 61] the current fastest IPMs are the aforementioned results of [35] and [12]. Beyond these results, excitingly the work of [12] was recently extended to obtain comparable running time improvements for solving arbitrary empirical risk minimization problems (ERM) [37] and was recently simplified and de-randomized by van den Brand [5]. These works consider variants of the vector maintenance problem and our work is inspired in part by them. Our IPM leverages the barrier from [35] in a new way, that enables the application of robustness techniques from [5, 12, 35] and new techniques for handling approximately feasible points.

**Heavy Hitters and Sketching:** Sketching is a well-studied problem with a broad range of applications. Johnson-Lindenstrauss

<sup>1</sup>We assume throughout that  $A$  is non-degenerate meaning it has full-column rank and no-zero rows. This assumption can be avoided by preprocessing  $A$  to check for zero rows and adding a tiny amount of noise to the matrix to make it full-column rank. There are other natural ways to remove this assumption, see e.g. [35].

sketches were used extensively in previous IPMs [35], but only for the restricted application of computing leverage scores. Further sketching techniques have been used for the purpose of dimension reduction and sampling in other optimization contexts, e.g. solving linear systems [9–11, 39, 42, 44] and certain forms of  $\ell_p$ -regression [13]. Our methods make use of  $\ell_2$ -heavy hitters sketches [8, 16, 17, 26, 31, 45, 50], in particular the  $\ell_2$ -sketches of [26, 50], to decrease iteration costs. We are unaware of these sketches being used previously to obtain provable improvements for solving offline (as opposed to online, streaming, or dynamic optimization problems) variants of linear programming previously. We believe this is due in part to the difficulty of using these methods with non-oblivious adversaries and consequently we hope the techniques we use to overcome these issues may be of further use. Also, note that the definition of  $\ell_2$ -heavy hitters problem that we use is equivalent to  $\ell_\infty/\ell_2$ -sparse recovery problem. Though the  $\ell_2/\ell_2$ -sparse recovery [7, 19, 23, 25, 43, 51, 55] is a more standard task in compressed sensing, we are unaware of how to efficiently use its guarantees for our applications.

**Leverage Scores and Lewis Weights:** Leverage scores [11, 39, 40, 57] (and more broadly Lewis weights [4, 13, 35, 38]) are fundamental notions of importance of rows of a matrix with numerous applications. In this paper we introduce a natural online problem for maintaining multiplicative approximations to leverage scores and we show how to solve this problem efficiently. Though we are unaware of this problem being studied previously, we know that in the special case of leverage scores induced by graph problems, which are known as effective resistances, there are dynamic algorithms for maintaining them, e.g. [20]. However, these algorithms seem tailored to graph structure and it is unclear how to apply them in our setting. Further, there are streaming algorithms for variants of this problem [2, 27, 28], however their running time is too large for our purposes.

**Inverse Maintenance:** This problem has been studied extensively and are a key component in obtaining efficient linear programming running times with previous IPMs [5, 12, 29, 32–35, 37, 48, 49, 54, 55, 61] and other optimization methods [1, 36]. Outside the area of optimization, this problem is also known as Dynamic Matrix Inverse [6, 53]. Our method for solving inverse maintenance is closely related to these results with modifications needed to fully take advantage of our leverage score maintenance, obtain  $\text{poly}(d)$  (as opposed to  $\text{poly}(n)$  runtimes), ensure that randomness used to make the algorithm succeed doesn't affect the input to the data structure, and ultimately produce solvers that are correct in expectation.

## 2 OVERVIEW OF APPROACH

The proof of Theorem 1.1 is split into two steps. In the full version [65] we provide a new robust, primal-dual,  $\tilde{O}(\sqrt{d})$  iteration IPM inspired by the LS-barriers of [33, 35] and central path robustness techniques of [5, 12, 33, 35, 37]. Then we develop new data structures based on heavy hitters, the Johnson-Lindenstrauss lemma, and inverse maintenance. These data structure then allow us to efficiently implement this new IPM.

Here we give an outline for the new IPM and how to implement this IPM efficiently using the new data structures. Section 4 shows how to use these result to obtain the fast linear program solver.

### 2.1 A Robust Primal-Dual $\tilde{O}(\sqrt{d})$ IPM

Here we provide an overview of our approach to deriving our robust, primal-dual,  $\tilde{O}(\sqrt{d})$  iteration IPM. We assume throughout this section (and the bulk of the paper) that  $A$  is *non-degenerate*, meaning it is full-column rank and no non-zero rows. By standard techniques, in Section 4.5 we efficiently reduce solving (1) in general to solving an instance of where this assumption holds.

We first give a quick introduction to primal-dual path IPMs, review the central path from [35] and from it derive the central path that our method is based on, and explain our new IPM.

*Primal-Dual Path Following IPMs:* Our algorithm for solving the linear programs given by (1) is rooted in classic primal-dual path-following IPMs. Primal-dual IPMs, maintain a *primal feasible point*,  $x \in \mathbb{R}_{\geq 0}^n$  with  $A^\top x = b$ , a *dual feasible point*  $y \in \mathbb{R}^d$  with  $s = Ay - c \geq 0$ , and attempt to decrease the *duality gap*

$$\text{gap}(x, y) \stackrel{\text{def}}{=} c^\top x - b^\top y = (Ay + s)^\top x - (A^\top x)^\top y = s^\top x.$$

Note that  $\text{gap}(x, y)$  upper bounds the error of  $x$  and  $y$ , i.e.  $\text{gap}(x, y) \leq \epsilon$  implies that  $x$  and  $y$  are each optimal up to an additive  $\epsilon$  in objective function, and therefore to solve a linear program it suffices to decrease the duality gap for primal and dual feasible points.

Primal-dual path-following IPMs carefully trade-off decreasing the duality gap (which corresponds to objective function progress) with staying away from the inequality constraints, i.e.  $x \geq 0$  and  $s \geq 0$  (in order to ensure it is easier to make progress). Formally, they consider a (*weighted*) *central path* defined as the unique set of  $(x_\mu, y_\mu, s_\mu) \in \mathbb{R}_{\geq 0}^n \times \mathbb{R}^d \times \mathbb{R}_{\geq 0}^n$  for  $\mu > 0$  that satisfy

$$\begin{aligned} X_\mu S_\mu 1 &= \mu \cdot \tau_{\text{weight}}(x_\mu, s_\mu), \\ A^\top x_\mu &= b, \\ Ay_\mu + s_\mu &= c, \end{aligned} \tag{2}$$

where  $X_\mu \stackrel{\text{def}}{=} \text{Diag}(x_\mu)$ ,  $S_\mu \stackrel{\text{def}}{=} \text{Diag}(s_\mu)$ , and  $\tau_{\text{weight}} : \mathbb{R}_{\geq 0}^n \times \mathbb{R}_{\geq 0}^n \rightarrow \mathbb{R}$  is a *weight function*. These methods maintain primal and dual feasible points and take Newton steps on the above non-linear inequalities to maintain feasible points that are more central, i.e. have (2) closer to holding, for decreasing  $\mu$ . Since many properties of the central path can be defined in terms of just  $x_\mu$  and  $s_\mu$  we often describe methods using only these quantities and we adopt the following notation.

**Definition 2.1 (Feasible Point).** We say that  $(x, s) \in \mathbb{R}_{\geq 0}^n \times \mathbb{R}_{\geq 0}^n$  is a *feasible point* if there exists  $y \in \mathbb{R}^d$  with  $A^\top x = b$  and  $Ay + s = c$ .

Now, perhaps the most widely-used and simple weight function is  $\tau_{\text{weight}}(x, s) \leftarrow \tau_{\text{std}}(x, s) \stackrel{\text{def}}{=} 1$ , i.e. the all-ones vector. Central path points for this weight function are the solution to the following

$$\begin{aligned} x_\mu &= \underset{x \in \mathbb{R}_{\geq 0}^n : A^\top x = b}{\text{argmin}} \quad c^\top x - \mu \sum_{i \in [n]} \log(x_i) \text{ and} \\ (y_\mu, s_\mu) &= \underset{(y, s) \in \mathbb{R}^d \times \mathbb{R}_{\geq 0}^n : A^\top y + s = c}{\text{argmax}} \quad b^\top y - \mu \sum_{i \in [n]} \log(s_i), \end{aligned}$$

i.e. optimization problems trading off the objective function with logarithmic barriers on the inequality constraints. There are numerous methods for following this central path [24, 41, 47, 52, 69]. By starting with nearly-central feasible points for large  $\mu$  and iteratively finding nearly-central feasible points for small  $\mu$ , they

can compute  $\epsilon$ -approximate solutions to (1) in  $\tilde{O}(\sqrt{n})$ -iterations. For these methods, each iteration (or Newton step) consists of nearly-linear time computation and solving one linear system in the matrix  $\mathbf{A}^\top \mathbf{X} \mathbf{S}^{-1} \mathbf{A} \in \mathbb{R}^{d \times d}$  for  $\mathbf{X} = \text{Diag}(\mathbf{x}) \in \mathbb{R}^{n \times n}$  and  $\mathbf{S} = \text{Diag}(\mathbf{s}) \in \mathbb{R}^{n \times n}$ .

The first such  $\tilde{O}(\sqrt{n})$ -iteration IPM with this iteration cost was established by Renegar in 1988 [52] and no faster-converging method with the same iteration cost was developed until the work of Lee and Sidford in 2014 [33, 35]. It had been known since seminal work of Nesterov and Nemirovski in 1994 [47] that there is a weight function that yields a  $\tilde{O}(\sqrt{d})$ -iteration primal-dual path-following IPM, however obtaining any  $\tilde{O}(\sqrt{d})$ -iteration IPM that can be implemented with iterations nearly as efficient as those of [52] was a long-standing open problem.

*The Lewis Weight Barrier and Beyond* The work of [33, 35] addressed this open problem in IPM theory and provided an efficient  $\tilde{O}(\sqrt{d})$ -iteration IPM by providing new weight functions and new tools for following the central path they induce. In particular, [35] introduced the *Lewis weight barrier* which induces the central path in which for some  $p > 0$

$$\tau_{\text{weight}}(\mathbf{x}, \mathbf{s}) \leftarrow \tau_{\text{ls}}(\mathbf{x}, \mathbf{s}) \stackrel{\text{def}}{=} \sigma^{(p)}(\mathbf{S}^{-1} \mathbf{A})$$

where for any  $\mathbf{B} \in \mathbb{R}^{n \times d}$ ,  $\sigma^{(p)}(\mathbf{B})$  are the  $\ell_p$ -Lewis weights of the rows of  $\mathbf{B}$  [38], a fundamental and natural measure of importance of rows with respect to the  $\ell_p$ -norm [4, 13, 35]. In the case of non-degenerate  $\mathbf{B}$ , they are defined recursively as the vector  $\mathbf{w} \in \mathbb{R}_{>0}^n$  which satisfies

$$\mathbf{w} = \text{diag}(\mathbf{W}^{(1/2)-(1/p)} \mathbf{B} (\mathbf{B}^\top \mathbf{W}^{1-(2/p)} \mathbf{B})^{-1} \mathbf{B}^\top \mathbf{W}^{(1/2)-(1/p)}),$$

where  $\mathbf{W} = \text{Diag}(\mathbf{w})$ . In the special case when  $p = 2$ ,

$$\mathbf{w} = \sigma(\mathbf{B}) \stackrel{\text{def}}{=} \text{diag}(\mathbf{B} (\mathbf{B}^\top \mathbf{B})^\dagger \mathbf{B}^\top)$$

is known as the *leverage scores* of the rows of  $\mathbf{A}$  and is a fundamental object for dimension reduction and solving linear systems [11, 39, 40, 46, 56, 57, 67].

The work of [35] formally showed that there is an  $\tilde{O}(\sqrt{d})$ -iteration primal-dual IPM which uses  $\tau_{\text{weight}}(\mathbf{x}, \mathbf{s}) \leftarrow \tau_{\text{ls}}(\mathbf{x}, \mathbf{s})$  when  $p = \Omega(\log n)$ . This choice of  $p$  is motivated by drew geometric connections between Lewis weights and ellipsoidal approximations of polytopes [35]. Further, [35] showed how to modify this IPM to have iterations of comparable cost to the methods which use  $\tau_{\text{std}}(\mathbf{x}, \mathbf{s})$ . This was achieved by leveraging and modifying efficient algorithms for approximately computing Lewis weights and leverage scores and developing techniques for dealing with the noise such approximate computation induces.

To obtain the results of this paper we further simplify [35] and provide more robust methods for following related central paths. Our first observation is that the central path induced by  $\tau_{\text{ls}}$  can be re-written more concisely in terms of only leverage scores. Note that (2) and the definition of Lewis weights imply that there is  $\mathbf{w}_\mu$  with

$$\mathbf{X}_\mu \mathbf{S}_\mu \mathbf{1} = \mu \cdot \mathbf{w}_\mu \text{ and } \mathbf{w}_\mu = \sigma(\mathbf{W}_\mu^{(1/2)-(1/p)} \mathbf{S}_\mu^{-1} \mathbf{A}).$$

for  $\mathbf{W}_\mu \stackrel{\text{def}}{=} \text{Diag}(\mathbf{w}_\mu)$ . Substituting the first equation into the second, gives more compactly that for  $\alpha = 1/p$

$$\mathbf{X}_\mu \mathbf{S}_\mu \mathbf{1} = \mu \cdot \text{Diag}(\sigma(\mathbf{S}_\mu^{-1/2-\alpha} \mathbf{X}_\mu^{1/2-\alpha} \mathbf{A})).$$

Consequently, rather than defining centrality in terms of Lewis weights, we would get the same central path as the one induced by  $\tau_{\text{ls}}(\mathbf{x}, \mathbf{s})$  by letting  $\tau_{\text{weight}}(\mathbf{x}, \mathbf{s}) \leftarrow \sigma(\mathbf{S}_\mu^{1/2-\alpha} \mathbf{X}_\mu^{1/2-\alpha} \mathbf{A})$ , i.e. defining it in terms of leverage scores. In other words, the optimality of the central path conditions forces  $\mathbf{X}_\mu \mathbf{S}_\mu \mathbf{1}$  to be a type of Lewis weight if we carefully define centrality in terms of leverage scores. Though  $\tau_{\text{weight}}(\mathbf{x}, \mathbf{s}) \leftarrow \sigma(\mathbf{S}^{1/2-\alpha} \mathbf{X}^{1/2-\alpha} \mathbf{A})$  induces the same central path as  $\tau_{\text{ls}}(\mathbf{x}, \mathbf{s})$ , these weight functions can be different outside the central path and thereby lead to slightly different algorithms if only approximate centrality is preserved. Further, with the current state-of-the-art theory, leverage scores are simpler to compute and approximate, as Lewis weight computation is often reduced to leverage score computation [13, 35].

Formally, in this paper we consider the following regularized-variant of this centrality measure:

*Definition 2.2 (Weight Function).* Throughout this paper we let  $\alpha \stackrel{\text{def}}{=} 1/(4 \log(4n/d))$  and for all  $\mathbf{x}, \mathbf{s} \in \mathbb{R}_{>0}^n$  let

$$\tau_{\text{reg}}(\mathbf{x}, \mathbf{s}) \stackrel{\text{def}}{=} \sigma(\mathbf{S}^{-1/2-\alpha} \mathbf{X}^{1/2-\alpha} \mathbf{A}) + \frac{d}{n} \mathbf{1}$$

where  $\mathbf{X} = \text{Diag}(\mathbf{x})$  and  $\mathbf{S} = \text{Diag}(\mathbf{s})$ .

This centrality measure is the same as the Lewis weight barrier except that we add a multiple of the all-ones vector,  $\frac{d}{n} \mathbf{1}$ , to simplify our analysis. Further, this allows us to pick  $\alpha \stackrel{\text{def}}{=} 1/(4 \log(4n/d))$  as opposed to  $\alpha = 1/\Omega(\log n)$  due to the extra stability it provides. Since we use this weight function throughout the paper we overload notation and let  $\tau(\mathbf{x}, \mathbf{s}) \stackrel{\text{def}}{=} \tau_{\text{reg}}(\mathbf{x}, \mathbf{s})$  and  $\tau(\mathbf{B}) \stackrel{\text{def}}{=} \sigma(\mathbf{B}) + \frac{d}{n} \mathbf{1}$ .

*Our Robust Primal-Dual Method:* We obtain our results by proving that there is an efficient primal-dual path-following IPM based on  $\tau_{\text{reg}}(\mathbf{x}, \mathbf{s})$ . We believe that our analysis is slightly simpler than [35] due to its specification in terms of leverage scores, rather than the more general Lewis weights, but remark that the core ingredients of its analysis are similar. Formally, we provide Newton-method type steps that allow us to control centrality with respect to this measure and increase  $\mu$  fast enough that this yields an  $\tilde{O}(\sqrt{d})$ -iteration method.

Beyond providing a simplified  $\tilde{O}(\sqrt{d})$ -iteration IPM, we leverage this analysis to provide a robust method. Critical to the development of recent IPMs is that it is possible to design efficient primal-dual  $\tilde{O}(\sqrt{n})$ -iteration IPMs that take steps using only crude, multiplicative approximations to  $\mathbf{x}$  and  $\mathbf{s}$  [5, 12, 37]. These papers consider the standard central path but measure centrality using potential functions introduced in [33, 35]. This robustness allows these papers to efficiently implement steps by only needing to change smaller amounts of coordinates.

Similarly, we show how to apply these approximate centrality measurement techniques to the central path induced by  $\tau_{\text{reg}}(\mathbf{x}, \mathbf{s})$ . We show that it suffices to maintain multiplicative approximations to the current iterate  $(\mathbf{x}, \mathbf{s} \in \mathbb{R}_{>0}^n)$ , the regularized leverage scores  $(\sigma(\mathbf{S}^{-1/2-\alpha} \mathbf{X}^{1/2-\alpha} \mathbf{A}) + \frac{d}{n} \mathbf{1})$ , and the inverse of the local Hessian  $((\mathbf{A}^\top \mathbf{S}^{-1} \mathbf{X} \mathbf{A})^{-1})$  to maintain approximate centrality with respect to  $\tau_{\text{reg}}(\mathbf{x}, \mathbf{s})$ . Interestingly, to do this we slightly modify the type of steps we take in our IPM. Rather than taking standard Newton steps, we slightly change the steps sizes on  $\mathbf{x}$  and  $\mathbf{s}$  to account for the effect of  $\tau_{\text{reg}}(\mathbf{x}, \mathbf{s})$ . Further, approximation of the Hessian causes

the  $x$  iterates to be infeasible, but in the full version we discuss how to modify the steps to control this infeasibility and still prove the desired theorem.

The main guarantees of this new IPM are given by Theorem 4.1 (Section 4.1) and are proven in the full version. This theorem formalizes the above discussion and quantifies how much the iterates, leverage scores, and Hessian can change in each iteration. These bounds are key to obtaining an efficient method that can maintain multiplicative approximations to these quantities.

## 2.2 Heavy Hitters, Congestion Detection, and Sketching the Central Path

It has long been known that the changes to the central path are essentially sparse or low-rank on average. In Renegar’s IPM [52], the multiplicative change in  $x$  and  $s$  per iteration are bounded in  $\ell_2$ . Consequently, the matrices for which linear systems are solved in Renegar’s method do not change too quickly. In fact, this phenomenon holds for IPMs more broadly, and from the earliest work on polynomial time IPMs [29], to the most recent fastest methods [5, 33–35, 37], and varied work in between [48, 49, 62] this bounded change in IPM iterates has been leveraged to obtain faster linear programming algorithms.

Our IPM also enjoys a variety of stability properties on its iterates. We show that the multiplicative change in the iterates are bounded in a norm induced by  $\tau_{\text{reg}}(x, s)$  and therefore are also bounded in  $\ell_2$ . This is known to imply that changes to  $A^\top S^{-1}XA \in \mathbb{R}^{d \times d}$  can be bounded over the course of the algorithm and we further show that this implies that the changes in the weight function,  $\tau_{\text{reg}}(x, s)$ , can also be bounded. These facts, combined with the robustness properties of our IPM imply that to obtain an efficient linear programming algorithm it suffices to maintain multiplicative approximations to the following three quantities (1) the vectors  $x, s \in \mathbb{R}_{\geq 0}^n$ , (2) the regularized leverage scores  $\tau_{\text{reg}}(x, s)$ , and (3) the Hessian inverse  $(A^\top S^{-1}XA)^{-1} \in \mathbb{R}^{d \times d}$  under bounds on how quickly these quantities change.

We treat each of these problems as a self contained data-structure problem, the first we call the *vector maintenance problem*, the second we call the *leverage score maintenance problem*, and the third has been previously studied (albeit different variants) and is called the *inverse maintenance problem*. For each problem we build efficient solutions by combining techniques from the sketching literature (e.g. heavy hitters sketches and Johnson-Lindenstrauss sketches) and careful potential functions and tools for dealing with sparse and low rank approximations. (Further work is also needed to maintain the gradient of a potential used to measure proximity to the central path and this is discussed in the full version.)

In the remainder of this overview, we briefly survey how we solve each of these problems. A common issue that needs to be addressed in solving each problem is that of hiding randomness and dealing with adversarial input. Each data-structure uses sampling and sketching techniques to improve running times. While these techniques are powerful and succeed with high probability, they only work against an *oblivious adversary*, i.e. one which provides input that does not depend on the randomness of the data structure. However, the output of our data-structures are used to take steps along the central path and provide the next input, so care needs to be taken to argue

that the output of the data structure, as used by the method, doesn’t somehow leak information about the randomness of the sketches and samples into the next input. A key contribution of our work is showing how to overcome this issue in each case.

**Vector Maintenance:** In the vector maintenance problem, we receive two online sequences of vectors  $h^{(1)}, h^{(2)}, \dots \in \mathbb{R}^d$  and  $g^{(1)}, g^{(2)}, \dots \in \mathbb{R}^n$  and must maintain the sum  $y^{(t)} \stackrel{\text{def}}{=} \sum_{i \in [t]} G^{(i)} A h^{(i)}$  for a fixed matrix  $A \in \mathbb{R}^{n \times d}$ . The naive way of solving this would just compute  $G^{(t)} A h^{(t)}$  in iteration  $t$  and add it to the previous result. Unfortunately, this takes  $O(\text{nnz}(A))$  time per iteration, which is too slow for our purposes. Luckily we do not have to maintain this sum exactly. Motivated by the robustness of our IPM, it is enough to maintain a multiplicative approximation  $\tilde{y}^{(t)} \approx_\epsilon y^{(t)}$  of the sum. An exact definition of this problem, together with our upper bounds, can be found in Section 4.2 and the formal proof of these results can be found in the full version.

We now outline how vector maintenance problem can be solved. For simplicity assume we already have an accurate approximation  $\tilde{y}^{(t-1)} \approx_{\epsilon/2} y^{(t-1)}$ . Then we only have to change the entries  $i$  of  $\tilde{y}^{(t-1)}$  where  $y_i^{(t)} \not\approx_{\epsilon/2} y_i^{(t-1)}$ , because for all other  $i$  we have  $\tilde{y}_i^{(t-1)} \approx_\epsilon y_i^{(t)}$ . This means we must detect the large entries of  $y^{(t)} - y^{(t-1)} = G^{(t)} A h^{(t)}$ , which can be done via heavy hitter techniques. From past research on heavy hitters, we know one can construct a small, sparse, random matrix  $\Phi \in \mathbb{R}^{k \times n}$  with  $k \ll n$ , such that known the much smaller vector  $x \in \mathbb{R}^k$  with  $x = \Phi y$  allows for a quick reconstruction of the large entries of  $y$ . Thus for our task, we maintain the product  $\Phi G^{(t)} A$ , which can be done quickly if  $g^{(t)}$  does not change in too many entries compared to  $g^{(t-1)}$ . Then we can reconstruct the large entries of  $G^{(t)} A h^{(t)}$  by computing  $(\Phi G^{(t)} A) h^{(t)}$ . Note that this product can be computed much faster than  $G^{(t)} A h^{(t)}$ , because  $\Phi G^{(t)} A$  is a  $k \times d$  matrix and  $k \ll n$ .

One issue we must overcome, is that the output of our data-structure must not leak any information about  $\Phi$ . This matrix is randomly constructed and the large entries of  $y$  can only be reconstructed from  $x = \Phi y$ , if the vector  $y$  is independent from the randomness in  $\Phi$ . Thus if the output of the data-structure depends on  $\Phi$  and the next future input  $h^{(t)}$  depends on the previous output, then the required independence can no longer be guaranteed. We overcome this problem by computing any entry  $y_i^{(t)}$  exactly, whenever the heavy hitter techniques detect a large change in said entry. As we now know the value of said entry exactly, we can compare it to the previous result and verify, if the entry did indeed change by some large amount. This allows us to define the output of our data-structure in a deterministic way, e.g. maintain  $\tilde{y}_i^{(t)} \approx_\epsilon y_i^{(t)}$  by setting  $\tilde{y}_i^{(t)} = y_i^{(t)}$  whenever  $\tilde{y}_i^{(t-1)} \not\approx_\epsilon y_i^{(t)}$ . (The exact deterministic definition we use is slightly more complex, but this is the high-level idea.)

So far we only explained how we can detect large changes in  $y^{(t)}$  that occur within a single iteration. However, it could be that some entry changes slowly over several iterations. It is easy to extend our heavy hitter technique to also detect these slower changes. The idea is that we not just detect changes within one iteration (i.e.  $y_i^{(t)} \not\approx_{\epsilon/2} y_i^{(t-1)}$ ) but also changes within any power of two (i.e.  $y_i^{(t)} \not\approx_{\epsilon/2} y_i^{(t-2^i)}$  for  $i = 1, \dots, \log t$ ). To make sure that this does not

become too slow, we only check every  $2^i$  iterations if there was a large change within the past  $2^i$  iterations. One can prove that this is enough to also detect slowly changing entries.

**Leverage Score Maintenance:** In the leverage score maintenance problem, we must maintain an approximation of the regularized leverage scores of a matrix of the form  $G^{1/2}A \in \mathbb{R}^{n \times d}$ , for a slowly changing diagonal matrix  $G \in \mathbb{R}^{n \times n}$ . That means we are interested in a data-structure that maintains a vector  $\tilde{\tau} \in \mathbb{R}^n$  with  $\tilde{\tau}_i \approx_{\epsilon} \tau_i := (G^{1/2}A(A^TGA)^{-1}A^TG^{1/2})_{i,i} + d/n$ . The high-level idea for how to solve this task is identical to the previously outlined vector-maintenance: we try to detect for which  $i$  the value  $\tau_i$  changed significantly, and then update  $\tilde{\tau}_i$  so that the vector stays a valid approximation. We show that detecting these indices  $i$  can be reduced to the previous vector maintenance. Here we simply outline this reduction, a formal description of our result can be found in Section 4.3 and its proof and analysis is in the full version.

Consider the case where we change from  $G$  to some  $G'$ . We want to detect indices  $i$  where the  $i$ th leverage score changed significantly. Define  $M \stackrel{\text{def}}{=} (A^TGA)^{-1}A^TG^{1/2} - (A^TG'A)^{-1}A^TG'^{1/2} \in \mathbb{R}^{d \times n}$  and note that

$$(G^{1/2}A(A^TGA)^{-1}A^TG^{1/2})_{i,i} = \|e_i^T G^{1/2}A(A^TGA)^{-1}A^TG^{1/2}\|_2^2,$$

so a large change in the  $i$ th leverage score, when changing  $G \in \mathbb{R}^{n \times n}$  to  $G' \in \mathbb{R}^{n \times n}$ , results in a large norm of  $e_i^T G^{1/2}AM$ , if  $G_{i,i}$  and  $G'_{i,i}$  are roughly the same. (As  $G$  is slowly changing there are not too many  $i$  where  $G_{i,i}$  and  $G'_{i,i}$  significantly. So we can just compute the  $i$ th leverage score exactly for these entries to check if the  $i$ th score changed significantly.) By multiplying this term with a Johnson-Lindenstrauss matrix  $J$  the task of detecting a large leverage score change, becomes the task of detecting rows of  $G^{1/2}AMJ$  for which the row-norm changed significantly. Given that  $J$  only needs some  $O(\log n)$  columns to yield a good approximation of these norms, we know that any large change in the  $i$ th leverage score, must result in some index  $j$  where  $|(G^{1/2}AMJ)_{i,j}|$  must be large. Detecting these large entries can be done in the same way as in the vector-maintenance problem by considering each column of  $MJ$  as a vector.

To make sure that no information about the random matrix  $J$  is leaked, we use the same technique previously outlined in the vector-maintenance paragraph. That is, after detecting a set  $I \subset [n]$  of indices  $i$  for which the leverage score might have changed significantly, we compute the  $i$ th leverage score to verify the large change and set  $\tilde{\tau}_i$  to be this computed leverage score, if the change was large enough. Unlike the vector case however, the  $i$ th leverage score is not computed in a deterministic way (as this would be prohibitively expensive). Instead we use another random Johnson-Lindenstrauss matrix  $J'$ , so the output  $\tilde{\tau}$  is actually defined w.r.t the input and this new matrix  $J'$ . By using a fresh independent Johnson-Lindenstrauss  $J'$  to verify changes to leverage scores, this data structure works against an adaptive adversary.

**Inverse Maintenance:** In the inverse maintenance problem, we maintain a spectral sparsifier of  $A^TWA \in \mathbb{R}^{d \times d}$  and its inverse, for a slowly changing diagonal matrix  $W \in \mathbb{R}^{n \times n}$ . Using the leverage score data-structure, we can assume an approximation to the leverage scores of  $A^TWA$  is given. Hence, we can sample  $\tilde{O}(d)$  many

rows of  $A$  to form a spectral sparsifier. This allows us to get a spectral sparsifier and its inverse in  $\tilde{O}(d^\omega)$ . To speed up the runtime, we follow the idea in [34], which resamples the row only if the leverage score changed too much. This makes sure the sampled matrix is slowly changing in  $\ell_0$  sense and hence we can try to apply the lazy low-rank update idea in [12] to update the inverse in  $\tilde{O}(d^2)$  time. Unfortunately, this algorithm works only against oblivious adversary and the sampled matrix is changing too fast in  $\ell_2$  sense, which is required for the leverage score maintenance.

Fortunately, we note that we do not need a sparsifier that satisfies both conditions at all time. Therefore, our final data structure has two ways to output the inverse of the sparsifier. The sparsifier that works only against oblivious adversary is used in implementing the Newton steps and the sparsifier that is slowly changing is used to compute the sketch  $MJ$  used in the leverage score maintenance problem mentioned above.

For the Newton steps, we only need to make sure the sparsifier does not leak the randomness since the input  $W$  depends on the Newton step. Since we only need to solve linear systems of the form  $A^TWAx = b \in \mathbb{R}^d$ , we can handle this problem by adding an appropriate noise to the output  $x$ . This makes sure the randomness we use in this data structure does not leak when the output  $x$  is used. This idea is also used [34], but extra care is needed to remove the  $\text{nnz}(A)$  per step cost in their algorithm.

For computing the sketch  $MJ$ , we do not need to worry about leakage of randomness, but only need to make sure it is slowly changing in  $\ell_2$  sense. Instead of using [12] as a black-box, we show how to combine the idea of resampling in [34] and the low-rank update in [12]. This gives us an alternate smoother scheme that satisfies the requirement for slowly changing.

### 3 PRELIMINARIES

Here we discuss varied notation we use throughout the paper. We adopt similar notation to [33] and some of the explanations here are copied directly from this work.

**Matrices:** We call a matrix  $A$  *non-degenerate* if it has full column-rank and no zero rows. We call symmetric matrix  $B \in \mathbb{R}^{n \times n}$  positive semidefinite (PSD) if  $x^TBx \geq 0$  for all  $x \in \mathbb{R}^n$  and positive definite (PD) if  $x^TBx > 0$  for all  $x \in \mathbb{R}^n$ .

**Matrix Operations:** For symmetric matrices  $A, B \in \mathbb{R}^{n \times n}$  we write  $A \leq B$  to indicate that  $x^TAx \leq x^TBx$  for all  $x \in \mathbb{R}^n$  and define  $<, \leq$ , and  $\geq$  analogously. For  $A, B \in \mathbb{R}^{n \times m}$ , we let  $A \circ B$  denote the Schur product, i.e.  $[A \circ B]_{i,j} \stackrel{\text{def}}{=} A_{i,j} \cdot B_{i,j}$  for all  $i \in [n]$  and  $j \in [m]$ . We use  $\text{nnz}(A)$  to denote the number of nonzero entries in  $A$ .

**Diagonals:** For  $A \in \mathbb{R}^{n \times n}$  we define  $\text{diag}(A) \in \mathbb{R}^n$  with  $\text{diag}(A)_i = A_{i,i}$  for all  $i \in [n]$  and for  $x \in \mathbb{R}^n$  we define  $\text{Diag}(x) \in \mathbb{R}^{n \times n}$  as the diagonal matrix with  $\text{diag}(\text{Diag}(x)(x)) = x$ . We often use upper case to denote a vectors associated diagonal matrix and in particular let  $X \stackrel{\text{def}}{=} \text{Diag}(x)$ ,  $S \stackrel{\text{def}}{=} \text{Diag}(s)$ ,  $W \stackrel{\text{def}}{=} \text{Diag}(w)$ ,  $T \stackrel{\text{def}}{=} \text{Diag}(t)$ ,  $\bar{X} \stackrel{\text{def}}{=} \text{Diag}(\bar{x})$ ,  $\bar{S} \stackrel{\text{def}}{=} \text{Diag}(\bar{s})$ ,  $\bar{W} \stackrel{\text{def}}{=} \text{Diag}(\bar{w})$ ,  $X_t \stackrel{\text{def}}{=} \text{Diag}(x_t)$ ,  $S_t \stackrel{\text{def}}{=} \text{Diag}(s_t)$ ,  $W_t \stackrel{\text{def}}{=} \text{Diag}(w_t)$ , and  $T_t \stackrel{\text{def}}{=} \text{Diag}(t_t)$ .

**Fundamental Matrices:** For any non-degenerate matrix  $A \in \mathbb{R}^{n \times d}$  we let  $P(A) \stackrel{\text{def}}{=} A(A^TA)^{-1}A^T$  denote the orthogonal projection matrix onto  $A$ 's image. Further, we let  $\sigma(A) \stackrel{\text{def}}{=} \text{diag}(P(A))$

denote  $A$ 's *leverage scores* and we let  $\tau(A) \stackrel{\text{def}}{=} \sigma(A) + \frac{d}{n} \mathbf{1}$  denote its regularized leverage scores. Further, we define  $\Sigma(A) \stackrel{\text{def}}{=} \text{Diag}(\sigma(A))$ ,  $T(A) \stackrel{\text{def}}{=} \text{Diag}(\tau(A))$ ,  $P^{(2)}(A) \stackrel{\text{def}}{=} P(A) \circ P(A)$  (where  $\circ$  denotes entry-wise product), and  $\Lambda(A) \stackrel{\text{def}}{=} \Sigma(A) - P^{(2)}(A)$ .

**Approximations:** We use  $x \approx_\epsilon y$  to denote that  $\exp(-\epsilon)y \leq x \leq \exp(\epsilon)y$  and  $A \approx_\epsilon B$  to denote that  $\exp(-\epsilon)B \leq A \leq \exp(\epsilon)B$ .

**Norms:** For PD  $A \in \mathbb{R}^{n \times n}$  we let  $\|\cdot\|_A$  denote the norm where  $\|x\|_A^2 \stackrel{\text{def}}{=} x^\top A x$  for all  $x \in \mathbb{R}^n$ . For positive  $w \in \mathbb{R}_{>0}^n$  we let  $\|\cdot\|_w$  denote the norm where  $\|x\|_w^2 \stackrel{\text{def}}{=} \sum_{i \in [n]} w_i x_i^2$  for all  $x \in \mathbb{R}^n$ . For any norm  $\|\cdot\|$  and matrix  $M$ , its induced operator norm of  $M$  is defined by  $\|M\| = \sup_{\|x\|=1} \|Mx\|$ .

**Time and Probability:** We use  $\tilde{O}(\cdot)$  to hide factors polylogarithmic in  $n$  and  $d$ . We say an algorithm has a property “with high probability (w.h.p.) in  $n$ ” if it holds with probability at least  $1 - 1/O(\text{poly}(n))$  for any polynomial by choice of the constants in the runtime of the algorithm.

**Misc:** We let  $[z] \stackrel{\text{def}}{=} \{1, 2, \dots, z\}$ . We let  $\mathbf{1}_n, \mathbf{0}_n \in \mathbb{R}^n$  denote the all-one and all-zero vectors,  $\mathbf{0}_n, \mathbf{I}_n \in \mathbb{R}^{n \times n}$  denote the all zero and identity matrices, and drop subscripts when the dimensions are clear. We let  $\mathbf{1}_i$  denote the indicator vector for coordinate  $i$ , i.e. the  $i$ -th basis vector.

## 4 LINEAR PROGRAMMING ALGORITHM

Here we prove the main result of this paper that there is a  $\tilde{O}(nd + d^3)$  time algorithm for solving linear programs. This theorem is restated below for convenience:

**THEOREM 1.1 (MAIN RESULT).** *There is an algorithm (Algorithm 2) which given any linear program of the form (1) for non-degenerate<sup>2</sup>  $A \in \mathbb{R}^{n \times d}$  and  $\delta \in [0, 1]$  computes a point  $x \in \mathbb{R}_{\geq 0}^n$  such that*

$$c^\top x \leq \min_{A^\top x = b, x \geq 0} c^\top x + \delta \cdot \|c\|_2 \cdot R \text{ and } \|A^\top x - b\|_2 \leq \delta \cdot (\|A\|_F \cdot R + \|b\|_2)$$

where  $R$  is the diameter of the polytope in  $\ell_2$  norm, i.e.  $\|x\|_2 \leq R$  for all  $x \in \mathbb{R}_{\geq 0}^n$  with  $A^\top x = b$ . Further, the expected running time of the method is  $O((nd + d^3) \log^{O(1)} n \log \frac{n}{\delta})$ .

The proof for Theorem 1.1 uses four intermediate results, which we formally state in the next Sections 4.1, to 4.4. Each of these intermediate results is self-contained and analyzed in the full version. In this section we show how these results can be combined to obtain Theorem 1.1. The first result is a new, improved IPM as outlined in Section 2.1. The exact statement is given in Section 4.1. Here we give a rough summary to motivate the other three results used by our linear programming algorithm.

Our IPM is robust in the sense, that it makes progress, even if we maintain the primal dual solution pair  $(x, s)$  only approximately. Additionally, the linear system that is solved in each iteration, allows for spectral approximations. More accurately, it is enough to maintain a spectral approximation of an inverse of a matrix of the form  $A^\top W A$  for some diagonal matrix  $W$ . For this robust IPM we

<sup>2</sup>We assume throughout that  $A$  is non-degenerate meaning it has full-column rank and no-zero rows. This assumption can be avoided by preprocessing  $A$  to check for zero rows and adding a tiny amount of noise to the matrix to make it full-column rank. There are other natural ways to remove this assumption, see e.g. [35].

also require to compute approximate leverage scores, which allows the IPM to converge in just  $\tilde{O}(\sqrt{d})$  iterations. These properties of our IPM motivate three new data-structures, all of which are proven and analyzed in the full version:

(i) In Section 4.2, we present a data-structure that can maintain an approximation of the primal dual solution pair  $(x, s)$  efficiently. More formally, this data-structure maintains an approximation of the sum  $\sum_{k \in [t]} W^{(k)} A h^{(k)}$  for diagonal matrices  $W^{(k)} \in \mathbb{R}^{n \times n}$  and vectors  $h^{(k)} \in \mathbb{R}^d$ .

(ii) In Section 4.3 we present a data-structure that can maintain approximate leverage scores, required by the IPM.

(iii) The last requirement of the IPM is for us to maintain a spectral approximation of  $(A^\top W A)^{-1}$  for some diagonal matrix  $W$ . We present a data-structure that can maintain this inverse approximately in  $\tilde{O}(d^{\omega - \frac{1}{2}} + d^2)$  amortized time per step, when the matrix  $W$  changes slowly w.r.t  $\ell_2$ -norm and if we have estimates of the leverage scores of  $W^{1/2} A$ . The exact result is stated in Section 4.4.

With this we have all tools available for proving the main result Theorem 1.1. In Section 4.5 we show how to combine all these tools to obtain the fast linear programming algorithm.

### 4.1 Interior Point Method

In the full version we derive and analyze the core subroutine of our primal-dual robust  $\tilde{O}(\sqrt{d})$ -iteration IPM (Algorithm 1). This subroutine takes an approximate central point for parameter  $\mu^{(\text{init})}$  and in  $\tilde{O}(\sqrt{d} \log(\mu^{(\text{target})}/\mu^{(\text{init})}))$  iterations outputs an approximate central path point for any given parameter  $\mu^{(\text{target})}$ . Here we simply state the routine (Algorithm 1) and the main theorem regarding its performance (Theorem 4.1). We defer the full motivation of the method and its analysis (i.e. the proof of Theorem 4.1) to the full version. In the remainder of this section we argue how with the appropriate data-structures, this theorem implies our main result.

**THEOREM 4.1.** *There exists a constant  $\zeta > 0$  such that, given  $x^{(\text{init})}, s^{(\text{init})} \in \mathbb{R}_{>0}^n$ ,  $\mu^{(\text{init})} > 0$ ,  $\mu^{(\text{target})} > 0$ , and  $\epsilon \in (0, \alpha/16000)$  with  $x^{(\text{init})} s^{(\text{init})} \approx_{2\epsilon} \mu^{(\text{init})} \cdot \tau(x^{(\text{init})}, s^{(\text{init})})$  and*

$$\frac{1}{\mu^{(\text{init})}} \|Ax^{(\text{init})} - b\|_{(A^\top X^{(\text{init})} S^{(\text{init}) - 1} A)^{-1}} \leq \frac{\zeta \epsilon^2}{\log^6 n}$$

Algorithm 1 outputs  $(x^{(\text{final})}, s^{(\text{final})})$  such that

$$x^{(\text{final})} s^{(\text{final})} \approx_\epsilon \mu^{(\text{target})} \cdot \tau(x^{(\text{final})}, s^{(\text{final})}) \quad \text{and}$$

$$\frac{1}{\mu^{(\text{target})}} \|Ax^{(\text{final})} - b\|_{(A^\top X^{(\text{final})} S^{(\text{final}) - 1} A)^{-1}} \leq \frac{\zeta \epsilon^2}{\sqrt{d} \log^6 n}$$

$$\text{in } O\left(\sqrt{d} \log(n) \cdot \left(\frac{1}{\epsilon \alpha} \cdot \log\left(\frac{\mu^{(\text{target})}}{\mu^{(\text{init})}}\right) + \frac{1}{\alpha^3}\right)\right) \text{ iterations.}$$

Furthermore, throughout Algorithm 1, we have

- $xs \approx_{4\epsilon} \mu \cdot \tau(x, s)$  for some  $\mu$  where  $(x, s)$  is immediate points in the algorithms
- $\|X^{-1} \delta_x\|_{\tau+\infty} \leq \frac{\epsilon}{2}$ ,  $\|S^{-1} \delta_s\|_{\tau+\infty} \leq \frac{\epsilon}{2}$ ,  $\|\text{diag}(\tau)^{-1} \delta_\tau\|_{\tau+\infty} \leq 2\epsilon$  where  $\delta_x$ ,  $\delta_s$  and  $\delta_\tau$  is the change of  $x$ ,  $s$  and  $\tau$  in one iteration and  $\|x\|_{\tau+\infty} \stackrel{\text{def}}{=} \|x\|_\infty + C_{\text{norm}} \|x\|_\tau$  for  $C_{\text{norm}} \stackrel{\text{def}}{=} \frac{10}{\alpha}$ .

**Remark.** We will take  $\epsilon = 1/\text{poly} \log(n)$ .

**Algorithm 1:** Path Following (Theorem 4.1)

---

```

1 procedure CENTERING( $x^{(\text{init})} \in \mathbb{R}_{>0}^n, s^{(\text{init})} \in \mathbb{R}_{>0}^n, \mu^{(\text{init})} > 0, \mu^{(\text{target})} > 0, \epsilon > 0$ )
2    $\alpha \leftarrow 1/(4 \log(4n/d)), \lambda \leftarrow \frac{2}{\epsilon} \log(\frac{2^{16}n\sqrt{d}}{\alpha^2}),$ 
3    $\gamma \leftarrow \min(\frac{\epsilon}{4}, \frac{\alpha}{50\lambda})$ 
4    $\mu \leftarrow \mu^{(\text{init})}, x \leftarrow x^{(\text{init})}, s \leftarrow s^{(\text{init})}$ 
5   while true do
6     Pick any  $\bar{x}, \bar{s} \in \mathbb{R}_{>0}^n$  such that  $\bar{x} \approx_\epsilon x, \bar{s} \approx_\epsilon s$ 
7     Find  $\bar{v}$  such that  $\|\bar{v} - v\|_\infty \leq \gamma$  where
8      $v = \mu \cdot \tau(x, s)/w$ 
9     Let  $\Phi(v) \stackrel{\text{def}}{=} \sum_{i=1}^n \exp(\lambda(v_i - 1)) + \exp(-\lambda(v_i - 1))$ 
10    for all  $v \in \mathbb{R}^n$ 
11    if  $\mu = \mu^{(\text{target})}$  and  $\Phi(\bar{v}) \leq \frac{2^{16}n\sqrt{d}}{\alpha^2}$  then break;
12     $h = \gamma \nabla \Phi(\bar{v})^{b(\bar{v})}$  for  $\bar{\tau} \approx_\epsilon \tau(x, s)$ 
13    Pick any  $H \in \mathbb{R}^{d \times d}$  with
14     $H \approx_{(c\epsilon)/(d^{1/4} \log^3 n)} A^\top \bar{S}^{-1} \bar{X} A$  and
15     $\mathbb{E}[H] = A^\top \bar{S}^{-1} \bar{X} A$  for small constant  $c > 0$ .
16    Let  $Q = \bar{S}^{-1/2} \bar{X}^{-1/2} A H^{-1} A^\top \bar{X}^{1/2} \bar{S}^{-1/2}, \bar{W} = \bar{X} \bar{S}$ 
17     $x \leftarrow x + (1 + 2\alpha) \bar{X} \bar{W}^{-1/2} (I - Q) \bar{W}^{1/2} h$ 
18     $s \leftarrow s + (1 - 2\alpha) \bar{S} \bar{W}^{-1/2} Q \bar{W}^{1/2} h$ 
19    Pick any  $\bar{x}^{(\text{new})}, \bar{s}^{(\text{new})}, \bar{\tau}^{(\text{new})} \in \mathbb{R}_{>0}^n$  with
20     $\bar{x}^{(\text{new})} \approx_\epsilon x, \bar{s}^{(\text{new})} \approx_\epsilon s, \bar{\tau}^{(\text{new})} \approx_1 \tau(x, s)$ 
21     $\delta_\lambda \leftarrow \text{MAINTAINFEASIBILITY}(x, s, \bar{\tau}^{(\text{new})})$ 
22     $x \leftarrow x + \bar{X}^{(\text{new})} (\bar{S}^{(\text{new})})^{-1} A \delta_\lambda$ 
23    if  $\mu > \mu^{(\text{target})}$  then
24       $\mu \leftarrow \max\{\mu^{(\text{target})}, (1 - \frac{\gamma\alpha}{2^{15}\sqrt{d}})\mu\};$ 
25    else if  $\mu < \mu^{(\text{target})}$  then
26       $\mu \leftarrow \min\{\mu^{(\text{target})}, (1 + \frac{\gamma\alpha}{2^{15}\sqrt{d}})\mu\};$ 
27  return  $(x, s)$ 

```

---

Note that the complexity of Theorem 4.1 depends on the cost of implementing Lines 5, 6, 12, and 13 of Algorithm 1, as well as the cost of the function MAINTAINFEASIBILITY. Here Lines 5, 12 and 13 ask us to maintain an approximation of the primal dual solution pair  $(x, s)$ . A data-structure for this task is presented in Section 4.2. Additionally, to compute Lines 12 and 13, we must have access to an approximate inverse of  $A^\top \bar{S}^{-1} \bar{X} A$  (see Line 10). The task of maintaining this inverse will be performed by the data-structure presented in Section 4.4. At last, consider Line 6. To implement this line, we must have an approximation of the leverage scores  $\tau(x, s)$ . In Section 4.3, we present a data-structure that can efficiently maintain such an approximation.

To help us analyze the cost of function MAINTAINFEASIBILITY we prove the following in the full version.

**THEOREM 4.2 (MAINTAIN FEASIBILITY).** *The additional amortized cost of calling MAINTAINFEASIBILITY in Line 15 of Algorithm 1 is  $\tilde{O}(nd^{0.5} + d^{2.5}/\epsilon^2)$  per call, plus the cost of querying  $\tilde{O}(n/\sqrt{d} + d^{1.5}/\epsilon^2)$  entries of  $x$  and  $s$  (assuming  $x, s$  are given implicitly, e.g. via some data structure).*

**4.2 Vector Data Structure**

Consider an online sequence of  $n \times n$  diagonal matrices  $G^{(1)}, \dots, G^{(T)} \in \mathbb{R}^{n \times n}$  and vectors  $h^{(1)}, \dots, h^{(T)} \in \mathbb{R}^d, \delta^{(1)}, \dots, \delta^{(T)} \in \mathbb{R}^n$  and define  $y^{(t+1)} := \sum_{k=1}^t G^{(k)} A h^{(k)} + \delta^{(k)}$ . In this subsection we describe a data-structure that can efficiently maintain an approximation  $\tilde{y}^{(t)} \approx_\epsilon y^{(t)}$ , when the relative changes  $\|(Y^{(k)})^{-1} G^{(k)} A h^{(k)}\|_2$  and  $\|(Y^{(k)})^{-1} \delta^{(k)}\|_2$  are small. This is motivated by the following requirement of our IPM: we must maintain a multiplicative approximation of a sequence of vectors  $x^{(t)}, s^{(t)} \in \mathbb{R}^n$  (see Line 5 of Algorithm 1), where  $x^{(k+1)} = x^{(k)} + \delta_x^{(k)}, s^{(k+1)} = s^{(k)} + \delta_s^{(k)}$  and the terms  $\delta_x^{(k)}$  and  $\delta_s^{(k)}$  are roughly of the form (see Lines 12 and 13 of Algorithm 1):

$$\delta_x^{(k)} = (1 + 2\alpha) \bar{X}^{(k)} (\bar{W}^{(k)})^{-1/2} (I - Q^{(k)}) (\bar{W}^{(k)})^{1/2} v^{(k)},$$

$$\delta_s^{(k)} = (1 - 2\alpha) \bar{S}^{(k)} (\bar{W}^{(k)})^{-1/2} Q^{(k)} (\bar{W}^{(k)})^{1/2} v^{(k)}.$$

To maintain an approximation of  $x^{(t)}$ , we can then use the data-structure for maintaining an approximation of  $y^{(t)}$  by choosing

$$G^{(k)} = (1 + 2\alpha) \bar{X}^{(k)} (\bar{W}^{(k)})^{-1/2},$$

$$h^{(k)} = -Q^{(k)} (\bar{W}^{(k)})^{1/2} v^{(k)}, \delta^{(k)} = (\bar{W}^{(k)})^{1/2} v^{(k)}.$$

Likewise, we can maintain an approximation of  $s^{(t)}$  by a slightly different choice of parameters. The exact results is the following Theorem 4.3:

**THEOREM 4.3 (VECTOR MAINTENANCE).** *There exists a Monte-Carlo data-structure, that works against an adaptive adversary, with the following procedures:*

- **INITIALIZE**( $A, g, x^{(0)}, \epsilon$ ): Given matrix  $A \in \mathbb{R}^{n \times d}$ , scaling  $g \in \mathbb{R}^n$ , initial vector  $x^{(0)}$ , and target accuracy  $\epsilon \in (0, 1/10)$ , the data-structure preprocesses in  $O(\text{nnz}(A) \log^5 n)$  time.
- **SCALE**( $i, u$ ): Given  $i \in [n]$  and  $u \in \mathbb{R}$  sets  $g_i = u$  in  $O(d \log^5 n)$  amortized time.
- **QUERY**( $h^{(t)}, \delta^{(t)}$ ): Let  $g^{(t)} \in \mathbb{R}^n$  be the scale vector  $g \in \mathbb{R}^n$  during  $t$ -th call to **QUERY** and let  $h^{(t)} \in \mathbb{R}^d, \delta^{(t)} \in \mathbb{R}^n$  be the vectors given during that query. Define

$$x^{(t)} = x^{(0)} + \sum_{k \in [t]} G^{(k)} A h^{(k)} + \sum_{k \in [t]} \delta^{(k)}.$$

Then, w.h.p. in  $n$  the data-structure outputs a vector  $y \in \mathbb{R}^n$  such that  $y \approx_\epsilon x^{(t)}$ . Furthermore, the total cost over  $T$  steps is

$$O(T(n \log n + \sum_{k \in [T]} (\|(X^{(k)})^{-1} G^{(k)} A h^{(k)}\|_2^2 + \|(X^{(k)})^{-1} \delta^{(k)}\|_2^2) \cdot \epsilon^{-2} \cdot d \log^6 n)).$$

- **COMPUTEEXACT**( $i$ ): Output  $x_i^{(t)} \in \mathbb{R}^n$  exactly in amortized time  $O(d \log n)$ .

**4.3 Leverage Score Maintenance**

The IPM of Theorem 4.1, requires approximate leverage scores of some matrix of the form  $GA$ , where  $G$  is a diagonal matrix (see Line 6 of Algorithm 1, where  $G = (\bar{X}/\bar{S})^{1/2}$ ). Here the matrix  $G$  changes slowly from one iteration of the IPM to the next one, which

allows us to create a data-structure that can maintain the scores more efficiently than recomputing them from scratch every time  $\mathbf{G}$  changes.

**THEOREM 4.4 (LEVERAGE SCORE MAINTENANCE).** *There exists a Monte-Carlo data-structure, that works against an adaptive adversary, with the following procedures:*

- **INITIALIZE**( $\mathbf{A}, g, \epsilon$ ): Given matrix  $\mathbf{A} \in \mathbb{R}^{n \times d}$ , scaling  $g \in \mathbb{R}^n$  and target accuracy  $\epsilon > 0$ , the data-structure preprocesses in  $O(nd\epsilon^{-2} \log^4 n)$  time.
- **SCALE**( $i, u$ ): Given  $i \in [n]$  and  $u \in \mathbb{R}$  sets  $g_i = u$  in time  $O(d\epsilon^{-2} \log^5 n)$ .
- **QUERY**( $\Psi^{(t)}, \Psi_{(\text{safe})}^{(t)}$ ): Let  $g^{(t)}$  be the vector  $g$  during  $t$ -th call to **QUERY** and define  $\mathbf{H}^{(t)} = \mathbf{A}^\top (\mathbf{G}^{(t)})^2 \mathbf{A}$ . Given random input-matrices  $\Psi^{(t)} \in \mathbb{R}^{d \times d}$  and  $\Psi_{(\text{safe})}^{(t)} \in \mathbb{R}^{d \times d}$  such that  $\Psi^{(t)} \approx_{\epsilon/(24 \log n)} (\mathbf{H}^{(t)})^{-1}$ ,  $\Psi_{(\text{safe})}^{(t)} \approx_{\epsilon/(24 \log n)} (\mathbf{H}^{(t)})^{-1}$ .

and any randomness used to generate  $\Psi_{(\text{safe})}^{(t)}$  is independent of the randomness used to generate  $\Psi^{(t)}$ , w.h.p. in  $n$  the data-structure outputs a vector  $\tilde{\tau} \in \mathbb{R}^n$  independent of  $\Psi^{(1)}, \dots, \Psi^{(t)}$  such that  $\tilde{\tau}_i \approx_{\epsilon} \tau_i(\mathbf{G}^{(t)} \mathbf{A})$  for all  $i \in [n]$ .

The total cost of  $T$  calls to **QUERY** is

$$O((\sum_{t \in [T]} \|\mathbf{G}^{(t)} \mathbf{A} \Psi^{(t)} \mathbf{A}^\top \mathbf{G}^{(t)} - \mathbf{G}^{(t-1)} \mathbf{A} \Psi^{(t-1)} \mathbf{A}^\top \mathbf{G}^{(t-1)}\|_F)^2 \cdot \epsilon^{-4} n \log^7 n + T(T_\Psi + \epsilon^{-2} d^2 \log^3 n))$$

where  $T_\Psi$  is the time required to multiply a vector with  $\Psi^{(t)}$  (i.e. in case it is given implicitly via a data structure).

#### 4.4 Inverse Maintenance

For the IPM we must approximately maintain the inverse  $(\mathbf{A}^\top \mathbf{W} \mathbf{A})^{-1}$  where  $\mathbf{A} \in \mathbb{R}^{n \times d}$  undergoes changes to the diagonal matrix  $\mathbf{W}$  (see Line 10 of Algorithm 1 where  $\mathbf{W} = \mathbf{S}^{-1} \bar{\mathbf{X}}$ ). By using estimates of the leverage scores of  $\mathbf{W}^{1/2} \mathbf{A}$  (as maintained by Theorem 4.4, Section 4.3), we are able to maintain the inverse in amortized  $\tilde{O}(d^{\omega-\frac{1}{2}} + d^2)$  time per step, even for  $n \gg d$ . The exact result is stated as Theorem 4.5.

**THEOREM 4.5 (INVERSE MAINTENANCE).** *Given a full rank matrix  $\mathbf{A} \in \mathbb{R}^{n \times d}$  with  $n \geq d$  and error tolerance  $\epsilon \in (0, 1/10)$ , there is a data structure that approximately solves a sequence of linear systems  $\mathbf{A}^\top \bar{\mathbf{W}} \mathbf{A} \mathbf{y} = \mathbf{b} \in \mathbb{R}^d$  for positive diagonal matrices  $\bar{\mathbf{W}} \in \mathbb{R}^{n \times n}$  through the following operation:*

- **INITIALIZE**( $\mathbf{A}, \tilde{\mathbf{w}}, \tilde{\tau}, \epsilon$ ): Given matrix  $\mathbf{A} \in \mathbb{R}^{n \times d}$ , scaling  $\tilde{\mathbf{w}} \in \mathbb{R}_{>0}^n$ , shifted leverage score estimates  $\tilde{\tau} \in \mathbb{R}_{>0}^n$ , and accuracy  $\epsilon \in (0, 1/10)$ , the data-structure preprocesses in  $O(d^\omega)$  time.
- **UPDATE**( $\tilde{\mathbf{w}}, \tilde{\tau}$ ): Output a matrix  $\Psi \in \mathbb{R}^{d \times d}$  and vector  $\tilde{\mathbf{w}}^{(\text{alg})}$  where  $\Psi^{-1}$  is close to  $\mathbf{A}^\top \bar{\mathbf{W}} \mathbf{A} \in \mathbb{R}^{d \times d}$  and  $\tilde{\mathbf{w}}^{(\text{alg})}$  is close to  $\tilde{\mathbf{w}}$ .
- **SOLVE**( $\mathbf{b}, \tilde{\mathbf{w}}, \delta$ ): Input  $\tilde{\mathbf{w}} \approx \tilde{\mathbf{w}}$  and  $\delta > 0$ , output  $\mathbf{y} = \Psi \mathbf{b} \in \mathbb{R}^d$  for some random matrix  $\Psi^{-1} \in \mathbb{R}^{d \times d}$  that is close to  $\mathbf{A}^\top \bar{\mathbf{W}} \mathbf{A} \in \mathbb{R}^{d \times d}$ .

Let  $\tau(\mathbf{w}) \stackrel{\text{def}}{=} \tau(\mathbf{W} \mathbf{A})$ . Suppose that all estimate shifted leverage scores  $\tilde{\tau}_i \in (1 \pm \frac{1}{16 \log d}) \tau(\mathbf{w})_i$  for  $i \in [n]$ <sup>3</sup> and that there is a sequence

<sup>3</sup>Recall that  $\tau(\mathbf{w})_i = (\sqrt{\mathbf{W}} \mathbf{A} (\mathbf{A}^\top \mathbf{W} \mathbf{A})^{-1} \mathbf{A}^\top \sqrt{\mathbf{W}})_i$ ,  $\forall i \in [n]$

$\mathbf{w}^{(0)}, \mathbf{w}^{(1)}, \dots, \mathbf{w}^{(K)} \in \mathbb{R}^n$  such that the  $\mathbf{w}^{(k)}$  satisfy

$$\frac{1}{\epsilon^2} \|\mathbf{W}^{(k)}\|^{-1} (\mathbf{w}^{(k+1)} - \mathbf{w}^{(k)})\|_{\tau(\mathbf{w}^{(k)})}^2 + \|(\mathbf{T}(\mathbf{w}^{(k)}))^{-1} (\tau(\mathbf{w}^{(k+1)}) - \tau(\mathbf{w}^{(k)}))\|_{\tau(\mathbf{w}^{(k)})}^2 \leq \frac{1}{80} \quad (3)$$

for  $k = 0, 1, \dots, K-1$  with  $K = n^{O(1)}$ . Further assume that the update sequence  $\tilde{\mathbf{w}}^{(0)}, \tilde{\mathbf{w}}^{(1)}, \dots, \tilde{\mathbf{w}}^{(K)} \in \mathbb{R}^n$  satisfies  $\tilde{\mathbf{w}}^{(k)} \approx_{\epsilon/(16 \log d)} \mathbf{w}^{(k)}$  for all  $k$  and the  $\tilde{\mathbf{w}}^{(k)}$  are independent to the output of **UPDATE** and **SOLVE**. (The input  $\mathbf{b}^{(k)}$  can depend on the previous output of the data structure.) Then, we have the following:

- The amortized time per call of **UPDATE** is  $O(\epsilon^{-2} \cdot (d^{\omega-\frac{1}{2}} + d^2) \cdot \log^{3/2}(n))$ .
- The time per call of **SOLVE** is  $O(\delta^{-2} \cdot d^2 \cdot \log^2(n/\delta))$ .
- **UPDATE** outputs some  $\Psi, \tilde{\mathbf{w}}^{(\text{alg})}$  where  $\Psi^{-1} = \mathbf{A}^\top \tilde{\mathbf{W}}^{(\text{alg})} \mathbf{A} \approx_{\epsilon} \mathbf{A}^\top \bar{\mathbf{W}} \mathbf{A}$  with probability  $1 - 1/\text{poly}(n)$  and  $\tilde{\mathbf{w}}^{(\text{alg})} \approx_{\epsilon} \tilde{\mathbf{w}}$ .
- **SOLVE** outputs some  $\mathbf{y} = \Psi \mathbf{b}$  where  $\Psi^{-1} \approx_{\delta} \mathbf{A}^\top \bar{\mathbf{W}} \mathbf{A}$  with probability  $1 - 1/\text{poly}(n)$  and  $\mathbb{E}[\Psi \mathbf{b}] = (\mathbf{A}^\top \bar{\mathbf{W}} \mathbf{A})^{-1} \mathbf{b}$ .

In general, Theorem 4.5 does not work against adaptive adversaries, i.e. the input  $\mathbf{w}$  and  $\tilde{\tau}$  to the **UPDATE** procedure is not allowed to depend on previous outputs. In the full version we show, that this algorithm can be improved such that the input  $\mathbf{w}$  and  $\tilde{\tau}$  is allowed to depend on the output of **SOLVE**. However, the input is still not allowed to depend on the output of **UPDATE**.

**LEMMA 4.6.** *Theorem 4.5 holds even if the input  $\mathbf{w}$  and  $\tilde{\tau}$  of the algorithm depends on the output of **SOLVE**. Furthermore, we have*

$$\mathbb{E} \left[ \sum_{k \in [K-1]} \left\| \sqrt{\tilde{\mathbf{W}}^{(\text{alg})}(k+1)} \mathbf{A} \Psi^{(k+1)} \mathbf{A}^\top \sqrt{\tilde{\mathbf{W}}^{(\text{alg})}(k+1)} - \sqrt{\tilde{\mathbf{W}}^{(\text{alg})}(k)} \mathbf{A} \Psi^{(k)} \mathbf{A}^\top \sqrt{\tilde{\mathbf{W}}^{(\text{alg})}(k)} \right\|_F \right] \leq 16K \log^{5/2} n$$

where  $\Psi^{(k)} \in \mathbb{R}^{d \times d}$ ,  $\tilde{\mathbf{w}}^{(\text{alg})}(k)$  is the output of the  $k$ -th step of **UPDATE**( $\tilde{\mathbf{w}}^{(k)}, \tilde{\tau}^{(k)}$ ).

#### 4.5 Linear Programming Algorithm

Here we show how to combine the tools from Section 4.1 to 4.4 to obtain a linear program solver that runs in  $\tilde{O}(nd + d^3)$  time. First, we give a brief summary of our linear programming algorithm, Algorithm 2. The algorithm consists of two phases. In the first phase we construct a good initial feasible solution, and in the second we move along the central path towards the optimal solution.

The construction of the initial point works as follows: via a simple transformation (stated below as Theorem 4.7), we obtain a feasible solution pair  $(x, s)$  where both  $x$  and  $s$  are close to the all 1 vector and hence good enough as a point close to the central path of the standard log barrier function. However, we need to find a point such that  $xs \approx_{\epsilon} \mu \cdot (\sigma(\mathbf{S}^{-1/2-\alpha} \mathbf{X}^{1/2-\alpha} \mathbf{A}) + \frac{d}{n} \mathbf{1})$ . By picking  $x = 1$  and  $\mu = 1$ , the initial slack  $s$  needs to satisfy  $s \approx_{\epsilon} \sigma(\mathbf{S}^{-1/2-\alpha} \mathbf{A}) + \frac{d}{n} \mathbf{1}$ , which (up to the additive  $\frac{d}{n} \mathbf{1}$ ) is exactly the condition for  $\ell_p$  Lewis weight with  $p = \frac{1}{1+\alpha}$ . Cohen and Peng showed that such a vector  $s$  can be found efficiently as long as  $p \in (0, 4]$  [13]. We note that such  $s$  might not satisfy  $\mathbf{A} \mathbf{y} + s = c$ . That is why we define  $c^{(\text{tmp})} := \mathbf{A} \mathbf{y} + s$  for  $y = 0$ , so that  $(x, s)$  is a feasible solution pair for the cost vector  $c^{(\text{tmp})}$ . In the first phase of Algorithm 2, we move the point  $(x, s)$

along the central path of the temporary cost vector  $c^{(\text{tmp})}$  and bring the points to a location where we can switch the cost  $c^{(\text{tmp})}$  to  $c$  without violating the centrality conditions. This is how we obtain our feasible starting point for the cost vector  $c$ . In the subsequent second phase, we move along the path for the cost  $c$  until it is close to the optimal solution.

Moving along these two paths of the first and second phase is performed via the IPM of Algorithm 1 (Section 4.1, Theorem 4.1). Note that Algorithm 1 does not specify in Line 5 how to obtain the approximate solution pair  $(\bar{x}, \bar{s})$ , so we must implement this step on our own. Likewise, we must specify how to efficiently compute the steps in Lines 12 and 13. These implementations can be found in the second part of Algorithm 2. The high-level idea is to use the data-structures presented in Section 4.2 to 4.4.

To illustrate, consider Line 13 of Algorithm 1, which computes

$$s \leftarrow s + (1 - 2\alpha)\bar{S}\bar{W}^{-1/2}\bar{Q}\bar{W}^{1/2}h$$

for  $\bar{W} = \bar{X}\bar{S}$  and  $\bar{Q} = \bar{S}^{-1/2}\bar{X}^{-1/2}\bar{A}\bar{H}^{-1}\bar{A}^\top\bar{X}^{-1/2}\bar{S}^{-1/2}$  for  $\bar{H} \approx_\epsilon \bar{A}^\top\bar{S}^{-1}\bar{X}\bar{A}$ . We split this task into three parts: (i) compute  $r := \bar{A}^\top\bar{X}^{-1/2}\bar{S}^{-1/2}\bar{W}^{1/2}h$ , (ii) compute  $v := \bar{H}^{-1}r$ , and (iii) compute  $(1 - 2\alpha)\bar{S}\bar{W}^{-1/2}\bar{X}^{-1/2}\bar{A}v = (1 - 2\alpha)\bar{A}v$ . Part (i), the vector  $r$ , can be maintained efficiently, because we maintain the approximate solutions  $\bar{x}, \bar{s}$  (thus also  $\bar{w}$ ) and vector  $h$ , (via an Algorithm in the full version) in such a way, that per iteration only few entries change on average. Part (ii) is solved by the inverse maintenance data-structure of Section 4.4 (Theorem 4.5). The last part (iii) is solved implicitly by the data-structure of Section 4.2 (Theorem 4.3), which is also used to obtain the approximate solutions  $\bar{x}, \bar{s}$  in Line 5 of Algorithm 1. We additionally run the data-structure of Section 4.3 (Theorem 4.4) in parallel, to maintain an approximation of the leverage scores, which allows us to find the approximation  $\bar{v}$  required in Line 6. These modifications to Algorithm 1 are given in the second part of Algorithm 2.

The following theorem shows how to reduce solving any bounded linear program to solving a linear program with a non-degenerate constrain matrix and an explicit initial primal and dual interior point.

**THEOREM 4.7 (INITIAL POINT).** *Consider some linear program  $\min_{\bar{A}^\top \bar{x} = b, \bar{x} \geq 0} \bar{c}^\top \bar{x}$  with  $n$  variables and  $d$  constraints. Assume that*

1. *Diameter of the polytope:  $\|\bar{x}\|_2 \leq R$  for all  $\bar{x} \geq 0$  with  $\bar{A}^\top \bar{x} = b$*
2. *Lipschitz constant of the linear program:  $\|\bar{c}\|_2 \leq L$ .*
3. *The constraint matrix  $\bar{A}$  is non-degenerate.*

*For any  $\delta \in (0, 1]$ , the modified linear program  $\min_{\bar{A}^\top \bar{x} = \bar{b}, \bar{x} \geq 0} \bar{c}^\top \bar{x}$  with*

$$\bar{A} = \begin{bmatrix} \bar{A} & 1_n \|\bar{A}\|_F \\ 0 & 1 \|\bar{A}\|_F \\ \frac{1}{R} \bar{b}^\top - 1_n^\top \bar{A} & 0 \end{bmatrix} \in \mathbb{R}^{(n+2) \times (d+1)},$$

$$\bar{b} = \begin{bmatrix} \frac{1}{R} \bar{b} \\ (n+1) \|\bar{A}\|_F \end{bmatrix} \in \mathbb{R}^{d+1} \text{ and } \bar{c} = \begin{bmatrix} \frac{\delta}{L} \cdot \bar{c} \\ 0 \\ 1 \end{bmatrix} \in \mathbb{R}^{n+2}$$

*satisfies the following:*

$$1. \bar{x} = \begin{bmatrix} 1_n \\ 1 \\ 1 \end{bmatrix}, \bar{y} = \begin{bmatrix} 0_d \\ -1 \end{bmatrix} \text{ and } \bar{s} = \begin{bmatrix} 1_n + \frac{\delta}{L} \cdot \bar{c} \\ 1 \\ 1 \end{bmatrix} \text{ are feasible.}$$

## Algorithm 2: LP Algorithm (based on Algorithm 1)

```

1 global variables
2    $\mathbf{A} \in \mathbb{R}^{n \times d}, \mu > 0$ 
3    $D^{-1};$  // Theorem 4.5
4    $D^{(x)}, D^{(s)};$  // Theorem 4.3
5    $D_{\text{LEVERAGE}};$  // Theorem 4.4
6    $D_{\text{GRADIENT}};$  // See appendix of full version
7    $\bar{\tau} \in \mathbb{R}^n;$  //  $\bar{\tau}_{\gamma/8} \approx (\sigma(S^{-1/2-\alpha} X^{1/2-\alpha} A) + \frac{d}{n} 1)$ 
   /* Same parameters as in Theorem 4.1 */
8    $\alpha \stackrel{\text{def}}{=} 1/(4 \log(4n/d)), \epsilon \stackrel{\text{def}}{=} \frac{\alpha}{16000}$ 
9    $\lambda \stackrel{\text{def}}{=} \frac{2}{\epsilon} \log(\frac{2^{16} n \sqrt{d}}{\alpha^2}), \gamma \stackrel{\text{def}}{=} \min(\epsilon/4, \frac{\alpha}{50\lambda})$ 
10 procedure SOLVE( $\mathbf{A} \in \mathbb{R}^{n \times d}, b \in \mathbb{R}^d, c \in \mathbb{R}^d, \delta > 0$ )
11   Modify the LP and obtain an initial  $x, y$  and  $s$  by
     Lemma 4.7 to accuracy  $\delta/8n^2$ 
   /* For notational simplicity, we use  $\mathbf{A}, b, c, n, d$ 
     for the modified LP induced by Lemma 4.7
     in the remainder of the code. */
12   Find  $s$  with  $s \approx_\epsilon \sigma(S^{-1/2-\alpha} A) + \frac{d}{n} 1$  via Theorem 4.8
13    $\tau \leftarrow s, \mu \leftarrow 1$  // Since  $x = 1$  (Lemma 4.7),  $xs \approx_\epsilon \mu \tau$ 
14    $D^{-1}.$ INITIALIZE( $\mathbf{A}, S^{-1-2\alpha} x^{1-2\alpha}, \tau, \gamma/512 \log n$ )
15    $D_{\text{LEVERAGE}}.$ INITIALIZE( $\mathbf{A}, S^{-1-2\alpha} x^{1-2\alpha}, \gamma/8$ )
16    $D^{(x)}.$ INITIALIZE( $\mathbf{A}, S^{-1} x, \gamma/8$ )
17    $D^{(s)}.$ INITIALIZE( $\mathbf{A}, 1, s, \gamma/8$ )
18    $D_{\text{GRADIENT}}.$ INITIALIZE( $\mathbf{A}, \mu \tau / (xs), \tau, x, \gamma$ )
19    $c^{(\text{tmp})} \leftarrow s;$  //  $c^{(\text{tmp})} = \mathbf{A}y + s$  for  $y = 0$ 
20    $x, s, \bar{\tau}, \mu \leftarrow \text{CENTERING}(x, s, \tau, \mu, \Theta(n^2 \sqrt{d}/(\gamma \alpha^2)))$ 
21    $s^{(\text{new})} \leftarrow s + c - c^{(\text{tmp})}$ 
22    $D^{(s)}.$ INITIALIZE( $\mathbf{A}, 1, s^{(\text{new})}, \gamma/8$ )
23    $x, s, \bar{\tau}, \mu \leftarrow \text{CENTERING}(x, s, \bar{\tau}, \mu, \delta^2/(8^3 n^4 d))$ 
24   Reduce  $\frac{1}{\mu} \|\mathbf{A}^\top x - b\|_{(\mathbf{A}^\top \bar{X} \bar{S}^{-1} \bar{A})^{-1}}$  to some small  $O(\delta/n^2)$ 
25   Return an approximate solution of the original linear
     program according to Lemma 4.7

```

2. Let  $(\bar{x}, \bar{y}, \bar{s})$  be primal dual vectors of the modified LP and  $\Phi_b \stackrel{\text{def}}{=} \frac{1}{\mu} \cdot \|\bar{A}^\top \bar{x} - \bar{b}\|_{(\bar{A}^\top \bar{X} \bar{S}^{-1} \bar{A})^{-1}}^2$ , then  $\|\bar{x}\|_\infty \leq (1 + O(\Phi_b)) \cdot O(n)$ .

3. Let  $(\bar{x}, \bar{y}, \bar{s})$  be primal dual vectors of the modified LP with  $\bar{x} \cdot \bar{s} \approx_{0.5} \mu \cdot \tau(\bar{x}, \bar{s})$  for  $\mu < \delta^2/(8d)$  and small enough  $\Phi_b := \frac{1}{\mu} \cdot \|\bar{A}^\top \bar{x} - \bar{b}\|_{(\bar{A}^\top \bar{X} \bar{S}^{-1} \bar{A})^{-1}}^2 = O(1)$  (i.e.  $\bar{x}$  does not have to be feasible). The vector  $\hat{x} \stackrel{\text{def}}{=} R \cdot \bar{x}_{1:n}$  where  $\bar{x}_{1:n}$  is the first  $n$  coordinates of  $\bar{x}$  is an approximate solution to the original linear program in the following sense

$$c^\top \hat{x} \leq \min_{\bar{A}^\top \bar{x} = b, \bar{x} \geq 0} c^\top \bar{x} + O(nLR) \cdot (\sqrt{\Phi_b} + \delta),$$

$$\|\mathbf{A}^\top \hat{x} - b\|_2 \leq O(n^2) \cdot (\|\bar{A}\|_F R + \|\bar{b}\|_2) \cdot (\sqrt{\Phi_b} + \delta),$$

$$\hat{x} \geq 0.$$

As outlined before, the initial points given in Theorem 4.7 do not satisfy  $xs \approx_\epsilon \mu \cdot (\sigma(S^{-1/2-\alpha} X^{1/2-\alpha} A) + (d/n)1)$ . To satisfy this condition we pick  $x = 1$  and  $\mu = 1$ , and pick the initial slack vector  $s$  to satisfy  $s \approx_\epsilon \sigma(S^{-1/2-\alpha} A) + \frac{d}{n} 1$ , which is exactly the condition

**Algorithm 3:** LP Algorithm (based on Algorithm 1), Continuation of Algorithm 2.

---

```

26 procedure CENTERING( $x^{(\text{init})} \in \mathbb{R}_{>0}^n, s^{(\text{init})} \in \mathbb{R}_{>0}^n, \tau^{(\text{init})} >$ 
     $0, \mu^{(\text{init})} > 0, \mu^{(\text{target})} > 0$ )
27    $\mu \leftarrow \mu^{(\text{init})}, \bar{\tau} \leftarrow \tau^{(\text{init})}, \tau^{(\text{tmp})} \leftarrow \tau^{(\text{init})}, \bar{x} \leftarrow x^{(\text{init})},$ 
     $x^{(\text{tmp})} \leftarrow x^{(\text{init})}, \bar{s} \leftarrow s^{(\text{init})}, s^{(\text{tmp})} \leftarrow s^{(\text{init})}$ 
28   while TRUE do
29      $\bar{x}_i = x_i^{(\text{tmp})}$  for all  $i$  such that  $\bar{x}_i \not\approx_{\gamma/8} x_i^{(\text{tmp})}$ 
30      $\bar{s}_i = s_i^{(\text{tmp})}$  for all  $i$  such that  $\bar{s}_i \not\approx_{\gamma/8} s_i^{(\text{tmp})}$ 
31      $\bar{\tau}_i = \tau_i^{(\text{tmp})}$  for all  $i$  such that  $\bar{\tau}_i \not\approx_{\gamma/8} \tau_i^{(\text{tmp})}$ 
32      $\bar{w} \leftarrow \bar{X}\bar{s}, \bar{v} \leftarrow \mu \cdot \bar{W}^{-1}\bar{\tau}$ 
33     if  $\mu = \mu^{(\text{target})}$  and  $\Phi(\bar{v}) \leq \frac{2^{16}n\sqrt{d}}{\alpha^2}$  then break;
34      $D_{\text{GRADIENT}}.\text{UPDATE}(i, \bar{v}_i, \bar{\tau}_i, \bar{x}_i)$  for  $i$  where  $\bar{v}_i, \bar{\tau}_i$  or  $\bar{x}_i$  changed
35      $h, r \leftarrow D_{\text{GRADIENT}}.\text{QUERY}()$ 
36      $\Psi^{(\alpha)}, g \leftarrow D^{-1}.\text{UPDATE}(\bar{S}^{-1-2\alpha}\bar{x}^{1-2\alpha}, \bar{\tau})$ 
37      $D_{\text{LEVERAGE}}.\text{UPDATE}(i, \sqrt{g})$  for  $i$  where  $g_i$  changed
38      $\Psi_{(\text{safe})}^{(\alpha)}(b) \stackrel{\text{def}}{=} D^{-1}.\text{SOLVE}(b, \bar{S}^{-1-2\alpha}\bar{x}^{1-2\alpha}, \frac{\gamma}{512\log n})$ 
39      $\Psi_{(\text{safe})}^{(\alpha)}(b) \stackrel{\text{def}}{=} D^{-1}.\text{SOLVE}(b, \bar{S}^{-1}\bar{x}, (c\epsilon)/(d^{1/4}\log^3 n))$ 
40      $D^{(x)}.\text{SCALE}(\bar{x}/\bar{s})$  // Only scale coordinates
        where  $\bar{x}$  or  $\bar{s}$  changed.
41      $x^{(\text{tmp})} \leftarrow D^{(x)}.\text{QUERY}((1+2\alpha)\Psi_{(\text{safe})}^{(\alpha)}r, (1+2\alpha)\bar{X}h)$ 
42      $s^{(\text{tmp})} \leftarrow D^{(s)}.\text{QUERY}((1-2\alpha)\Psi_{(\text{safe})}^{(\alpha)}r, 0n)$ 
43      $\tau^{(\text{tmp})} \leftarrow D_{\text{LEVERAGE}}.\text{QUERY}(\Psi^{(\alpha)}, \Psi_{(\text{safe})}^{(\alpha)})$ 
44      $\delta_\lambda \leftarrow \text{MAINTAINFEASIBILITY}(D^{(x)}, D^{(s)}, \mu)$ 
45      $x^{(\text{tmp})} \leftarrow D^{(x)}.\text{QUERY}(\delta_\lambda, 0)$ 
46     if  $\mu > \mu^{(\text{target})}$  then
         $\mu \leftarrow \max(\mu^{(\text{target})}, (1 - \frac{\gamma\alpha}{2^{15}\sqrt{d}})\mu);$ 
47     else if  $\mu < \mu^{(\text{target})}$  then
         $\mu \leftarrow \min(\mu^{(\text{target})}, (1 + \frac{\gamma\alpha}{2^{15}\sqrt{d}})\mu);$ 
48      $x \leftarrow D^{(x)}.\text{COMPUTEEXACT}(1, \dots, n)$ 
49      $s \leftarrow D^{(s)}.\text{COMPUTEEXACT}(1, \dots, n)$ 
50   return  $(x, s, \tau^{(\text{tmp})}, \mu)$ 

```

---

for  $\ell_p$  Lewis weight with  $p = \frac{1}{1+\alpha}$ . The following theorem shows that such a vector  $s$  can be found efficiently as long as  $p \in (0, 4)$ . This vector  $s$  might not be a valid slack vector, so as outlined before, the algorithm runs in two phases: first using a cost vector  $c^{(\text{tmp})}$  for which the initial  $s$  is feasible, and then switching to the correct  $c$ .

**THEOREM 4.8 ([13]).** *Given  $p \in (0, 4)$ ,  $\eta > 0$  and non-degenerate  $A \in \mathbb{R}^{n \times d}$  w.h.p. in  $n$ , we can compute  $w \in \mathbb{R}_{>0}^n$  with  $w \approx_\epsilon \sigma(W^{\frac{1}{2}-\frac{1}{p}}A) + \eta 1$  in  $\tilde{O}((\text{nnz}(A) + d^\omega)\text{poly}(1/\epsilon))$  time.*

**PROOF.** Our proof is similar to [13], which proved a variant when  $\eta = 0$ . Consider the map  $T(w) \stackrel{\text{def}}{=} (W^{\frac{2}{p}-1}(\sigma(W^{\frac{1}{2}-\frac{1}{p}}A) + \eta 1))^{p/2}$ . Further, fix any positive vectors  $v, w \in \mathbb{R}^n$  such that  $v \approx_\alpha w$ . We

have that  $A^\top V^{1-\frac{1}{p}}A \approx_{|1-\frac{2}{p}|\alpha} A^\top W^{1-\frac{2}{p}}A$  and hence

$$a_i^\top (A^\top V^{1-\frac{2}{p}}A)^{-1}a_i \approx_{|1-\frac{2}{p}|\alpha} a_i^\top (A^\top W^{1-\frac{2}{p}}A)^{-1}a_i. \quad (4)$$

Note that

$$\begin{aligned} T(v)_i^{2/p} &= \frac{v_i^{1-\frac{2}{p}} a_i^\top (A^\top V^{1-\frac{2}{p}}A)^{-1}a_i + \eta}{v_i^{1-\frac{2}{p}}} \\ &= a_i^\top (A^\top V^{1-\frac{2}{p}}A)^{-1}a_i + \eta v_i^{\frac{2}{p}-1}. \end{aligned}$$

Using (4), we have

$$\begin{aligned} T(v)_i^{2/p} &\leq e^{|1-\frac{2}{p}|\alpha} a_i^\top (A^\top W^{1-\frac{2}{p}}A)^{-1}a_i + \eta v_i^{\frac{2}{p}-1} \\ &\leq e^{|1-\frac{2}{p}|\alpha} (a_i^\top (A^\top W^{1-\frac{2}{p}}A)^{-1}a_i + \eta w_i^{\frac{2}{p}-1}) \\ &= e^{|1-\frac{2}{p}|\alpha} T(w)_i^{2/p}. \end{aligned}$$

Similarly, we have  $T(v)_i^{2/p} \geq e^{-|1-\frac{2}{p}|\alpha} T(w)_i^{2/p}$ . Taking  $p/2$  power of both sides, we have

$$e^{-|1-\frac{2}{p}|\alpha} T(w)_i \leq T(v)_i \leq e^{|1-\frac{2}{p}|\alpha} T(w)_i.$$

Hence,  $T(v) \approx_{|p/2-1|\alpha} T(w)$ .

Consequently, let  $w_0 = \eta 1$  and consider the algorithm  $w_{k+1} = T(w_k)$ . Since  $\eta > 0$  we have that  $w_0 = w_0 \eta^{-p/2} \eta^{p/2} \leq T(w_0)$  and  $T(w_0) \leq w_0 \eta^{-p/2} (1 + \eta)^{p/2} \leq w_0 \exp(p\eta^{-1}/2)$ . Consequently,  $T(w_0) \approx_{p\eta^{-1}/2} w_0$  and after  $k$  steps we have that

$$T(w_k) \approx_{\exp(|p/2-1|k p \eta^{-1}/2)} w_k.$$

Since for  $p \in (0, 4)$ , we have that  $|p/2-1| < 1$  we see that after  $O(\log(\eta^{-1}/\epsilon))$  steps we have  $w_k \approx_\epsilon T(w_k)$ . Further, since  $w_k \geq \eta 1$  implies that  $T(w_k) \geq \eta 1$  we have that  $w_k \in \mathbb{R}_{>0}^n$  as desired. To implement the steps, one can check, by the same proof, that it suffices to get a  $\approx_{O(\epsilon)}$  multiplicative approximation to  $T(w)$  in each step, which can be done in  $\tilde{O}((\text{nnz}(A) + d^\omega)/\epsilon^2)$  per step by standard leverage score estimation techniques.  $\square$

Now, we first prove the correctness of the Algorithm 2.

**LEMMA 4.9.** *Algorithm 2 outputs  $x$  such that w.h.p. in  $n$*

$$\begin{aligned} c^\top x &\leq \min_{A^\top x=b, x \geq 0} c^\top x + LR \cdot \delta, \\ \|A^\top x - b\|_2 &\leq \delta \cdot \left( \|A\|_F \cdot R + \|b\|_2 \right), \\ x &\geq 0. \end{aligned}$$

**PROOF.** We define the primal dual point  $(x, s)$  via the formula of lines 12 and 13. We start by showing that our implementation of Algorithm 1 in Algorithm 2 satisfies all required conditions, i.e. that we can apply Theorem 4.1. Throughout this proof, states hold only w.h.p. and therefore the restatement of this is often omitted for brevity.

*Invariant:*  $\bar{x} \approx_{\gamma/4} x, \bar{s} \approx_{\gamma/4} s, \bar{\tau} \approx_{\gamma/4} \sigma(\bar{S}^{-1/2-\alpha}\bar{X}^{1/2-\alpha}A) + \frac{d}{n}1$  and  $\Psi_{(\text{safe})} \approx_\epsilon (A^\top \bar{S}^{-1}\bar{X}A)^{-1}$ : We first show that throughout CENTERING, we have the invariant above, assuming the input parameter  $\tau_i^{(\text{init})}$  satisfied  $\tau_i^{(\text{init})} \approx_{\gamma/4} \sigma(S^{-1/2-\alpha}X^{1/2-\alpha}A) + \frac{d}{n}1$ . Theorem 4.3 shows that  $x^{(\text{tmp})} \approx_{\gamma/8} x$  and  $s^{(\text{tmp})} \approx_{\gamma/8} s$  (Line 41 and 42).

Then, by the update rule (Line 29 and 30), we have that  $x^{(\text{tmp})} \approx_{\gamma/8} \bar{x}$  and  $s^{(\text{tmp})} \approx_{\gamma/8} \bar{s}$ . Hence, we have the desired approximation for  $\bar{x}$  and  $\bar{s}$ . Next, Theorem 4.5 shows that

- $\Psi^{(\alpha)} \approx_{\gamma/(512 \log n)} (\mathbf{A}^\top \mathbf{S}^{-1-2\alpha} \mathbf{X}^{1-2\alpha} \mathbf{A})^{-1}$  (Line 14) and
- $\Psi_{(\text{safe})}^{(\alpha)} \approx_{\gamma/(512 \log n)} (\mathbf{A}^\top \mathbf{S}^{-1-2\alpha} \mathbf{X}^{1-2\alpha} \mathbf{A})^{-1}$  (Line 38).

Hence, we can apply Theorem 4.4 and get

$$\tau^{(\text{tmp})} \approx_{\gamma/8} \sigma(\bar{\mathbf{S}}^{-1/2-\alpha} \bar{\mathbf{X}}^{1/2-\alpha} \mathbf{A}) + \frac{d}{n} 1$$

(Line 15 and Line 43). Again, by update rule (Line 31), we have the desired approximation for  $\bar{\tau}$ .

Finally,  $\Psi_{(\text{safe})} \approx_{\epsilon} (\mathbf{A}^\top \bar{\mathbf{S}}^{-1} \bar{\mathbf{X}} \mathbf{A})^{-1}$  follows from Theorem 4.5 (Line 39). These invariants also imply  $\|\bar{v} - v\|_\infty \leq \gamma$  and that  $D_{\text{Gradient}}$  maintains  $\nabla \Phi(\bar{v}')^{b(\bar{\tau})}$  for some  $\|\bar{v}' - v\|_\infty \leq \gamma$  (see full version). Thus, in summary, CENTERING of Algorithm 2 behaves like CENTERING of Theorem 4.1.

*Invariant:*  $xs \approx \mu \cdot \tau(x, s)$ : Initially,  $x = 1$  due to the reduction (Lemma 4.7),  $\mu = 1$  and  $s \approx_{\epsilon} \sigma(\mathbf{S}^{-1/2-\alpha} \mathbf{A}) + \frac{d}{n} 1$  (Line 12). Hence,  $xs \approx_{\epsilon} \mu \cdot (\sigma(\mathbf{S}^{-1/2-\alpha} \mathbf{X}^{1/2-\alpha} \mathbf{A}) + \frac{d}{n} 1)$  initially.

We change  $x, s, \mu$  in Line 20 by calling CENTERING. By Theorem 4.1 we then have  $xs \approx_{\epsilon} \mu \tau$  after this call to CENTERING. Next, consider the step where we switch the cost vector on Line 21. Let  $s^{(\text{new})}$  be the vector  $s$  after Line 21 and  $s$  is before Line 21. Since  $\mathbf{A}y + s = c^{(\text{tmp})}$  we have that  $\mathbf{A}y + s^{(\text{new})} = c^{(\text{tmp})} - s + s^{(\text{new})} = c$ , i.e.  $s$  is a valid slack vector for cost  $c$ . Further, we have that, by design

$$\frac{s^{(\text{new})} - s}{s} = \frac{c - c^{(\text{tmp})}}{s}.$$

First, we bound the denominator. We have that  $sx \approx_{1/2} \mu \tau$ , so for all  $i \in [n]$  we have

$$s_i \geq \frac{\mu}{2x_i} \frac{d}{n} \geq \frac{\mu}{\Omega(n^2)}$$

where we used  $\|x\|_\infty \leq O(n)$  by Lemma 4.7 as we ensure  $x$  is sufficiently close to feasible for the modified linear program by Theorem 4.1. For the numerator, we note that  $\|c\|_\infty \leq 1$  for the modified linear program and  $\|c^{(\text{tmp})}\|_\infty = \|s\|_\infty \leq 3$  for the  $s$  computed by Theorem 4.8 in Line 12, again by the definition of  $s$  (Line 12) and the modified  $\mathbf{A}$  and  $y$ . Hence, we have that

$$\left\| \frac{s^{(\text{new})} - s}{s} \right\|_\infty \leq \frac{16n^2}{\mu} \leq \frac{\gamma \alpha^2}{c \cdot \sqrt{d}}.$$

for any constant  $c$  by choosing the constant in the  $O(\cdot)$  in Line 20 appropriately. Thus  $xs \approx_{2\epsilon} \mu \cdot \tau(x, s)$  and  $\bar{\tau} \approx_{\gamma/4} \sigma(\mathbf{S}^{-1/2-\alpha} \mathbf{X}^{1/2-\alpha} \mathbf{A}) + \frac{d}{n} 1$ , so when we call CENTERING in Line 23, we again obtain  $xs \approx_{\epsilon} \mu \cdot \tau(x, s)$ .

*Conclusion:* Before the algorithm ends, we have  $xs \approx_{1/4} \mu \tau$ . In Line 24, we reduce  $\Phi_b := \|\mathbf{A}^\top x - b\|_{(\mathbf{A}^\top \mathbf{X} \mathbf{S}^{-1} \mathbf{A})^{-1}}^2$  to some small enough  $\Phi_b = O(\delta/n^2)$ . As argued in the full version, this does not move the vector  $x$  too much, i.e. we still have  $x \cdot s \approx_{1/2} \mu \cdot \tau(x, s)$ . Hence, by the choice of the new  $\mu$  in Line 23, Lemma 4.7 shows that we can output a point  $\hat{x}$  with the desired properties.  $\square$

Finally, we analyze the cost of the Algorithm 2.

LEMMA 4.10. *Algorithm 2 takes  $O((nd + d^3) \log^{O(1)} n \log(n/\delta))$  time with high probability in  $n$ .*

PROOF. For simplicity, we use  $\tilde{O}(\cdot)$  to suppress all terms that are  $\log^{O(1)} n$ . We first note that all parameters  $\alpha, \epsilon, \lambda, \gamma$  are either  $\log^{O(1)} n$  or  $1/\log^{O(1)} n$ . The cost of Algorithm 2 is dominated by the repeated calls to CENTERING.

*Number of iterations:* By Theorem 4.1 the calls to CENTERING in Lines 20 and 23 perform  $\tilde{O}(\sqrt{d} \log(1/\delta))$  many iterations in total.

*Cost of  $D^{-1}$ :* Note that  $\|\tau^{-1} \delta_\tau\|_{\tau+\infty} = 2\epsilon$ ,  $\|\mathbf{S}^{-1} \delta_s\|_{\tau+\infty} \leq \epsilon/2$ , and  $\|\mathbf{X}^{-1} \delta_x\|_{\tau+\infty} \leq \epsilon/2$  by Theorem 4.1. Hence, the weight of the vector  $w = \mathbf{S}^{-1-2\alpha} \mathbf{X}^{1-2\alpha}$  satisfies  $\|\mathbf{W}^{-1} \delta_w\|_\tau = O(\epsilon)$ . Note however, that we need  $D^{-1} \cdot \text{UPDATE}$  to compute the inverse with accuracy  $O(\gamma/\log n)$  and hence Theorem 4.5 requires the relative movement of  $w$  to be less than  $\gamma/\log n$ . This can be fixed by splitting the step into  $\tilde{O}(1)$  pieces. Now, Theorem 4.5 shows that the total cost is  $\tilde{O}(d^\omega + \sqrt{d} \log(1/\delta)(d^{\omega-\frac{1}{2}} + d^2)) = \tilde{O}((d^\omega + d^{2.5}) \log(1/\delta))$  where the term  $d^\omega$  is the initial cost and  $d^{\omega-\frac{1}{2}} + d^2$  is the amortized cost per step. Note that the solver runs with accuracy  $\tilde{O}(d^{-1/4})$  so the total time for all calls to  $D^{-1} \cdot \text{SOLVE}$  will be  $\tilde{O}(d^3)$ .

*Cost of  $D_{\text{LEVERAGE}}$ :* By Lemma 4.6, the total movement of the projection matrix is

$$\sum_{k \in [K-1]} \left\| \sqrt{\mathbf{W}^{(k+1)}} \mathbf{A} \Psi^{(k+1)} \mathbf{A}^\top \sqrt{\mathbf{W}^{(k+1)}} - \sqrt{\mathbf{W}^{(k)}} \mathbf{A} \Psi^{(k)} \mathbf{A}^\top \sqrt{\mathbf{W}^{(k)}} \right\|_F = \tilde{O}(K)$$

where  $K$  is the number of steps and  $w = \mathbf{S}^{-1-2\alpha} \mathbf{X}^{1-2\alpha}$ . Then, Theorem 4.4 shows that the total cost is  $O(\frac{n}{d} K^2 \cdot d) = O(nK^2) = O(nd \log^2(1/\delta))$ .

*Cost of  $D^{(x)}$  and  $D^{(s)}$ :* By Theorem 4.3, the total cost for  $D^{(x)}$  is bounded by  $\tilde{O}(K^2 \max \|\mathbf{X}^{-1} \delta_x\|_2^2 \cdot d + Kn)$  where  $\max \|\mathbf{X}^{-1} \delta_x\|_2^2$  is the maximum movement in one step in  $\ell_2$ -norm. Note that  $\tau \geq \frac{d}{n}$  and hence  $\|\mathbf{X}^{-1} \delta_x\|_2^2 \leq \frac{n}{d} \|\mathbf{X}^{-1} \delta_x\|_\tau^2 = \tilde{O}(\frac{n}{d})$ . Hence, we have that the cost is  $\tilde{O}(d \log^2(1/\delta) \cdot \frac{n}{d} \cdot d) = \tilde{O}(nd)$ . The bound for  $D^{(s)}$  is the same.

*Cost of maintaining  $\mathbf{A}^\top \bar{\mathbf{X}} h$  ( $D_{\text{Gradient}}$ ):* The cost of maintaining  $\mathbf{A}^\top \bar{\mathbf{X}} h$  is exactly equals to  $d$  times the number of coordinates changes in  $\bar{x}, \bar{s}, \bar{\tau}$ . Note that we change the exact  $x, s, \tau$  by at most around  $(1 \pm \gamma/8)$  multiplicative factor in every step. So, the total number of entry changes performed to  $\bar{x}, \bar{s}, \bar{\tau}$  is bounded by

$$\tilde{O} \left( K^2 \left( \max \|\mathbf{X}^{-1} \delta_x\|_2^2 + \max \|\mathbf{S}^{-1} \delta_s\|_2^2 + \max \|\tau^{-1} \delta_\tau\|_2^2 \right) \right),$$

where  $\max$  refers to the maximum movement in any step in  $\ell_2$ -norm. As we showed before, all the movement terms are bounded by  $\tilde{O}(\frac{n}{d})$  (see the paragraphs regarding  $D^{(x)}$ ,  $D^{(s)}$ , and  $D^{-1}$ ). So in total there are  $\tilde{O}(n \log^2(1/\delta))$  many entry changes we perform to  $\bar{x}, \bar{s}, \bar{\tau}$ . Hence, the total cost of maintenance is again  $\tilde{O}(nd \log^2(1/\delta))$ .

*Cost of implementing MAINTAINFEASIBILITY:* By Theorem 4.2 we pay  $\tilde{O}(nd^{0.5+d^{2.5}})$  amortized time per call to MAINTAINFEASIBILITY. Additionally, we must compute  $\tilde{O}(n/\sqrt{d} + d^{1.5})$  entries of  $x$  in each iteration (i.e. call  $D^{(x)} \cdot \text{COMPUTEEXACT}(i)$ ), which costs  $\tilde{O}(d)$  per call, so we have total cost of  $\tilde{O}(nd + d^3)$  after  $\tilde{O}(\sqrt{d})$  iterations.

*Removing the extra  $\log(1/\delta)$  term:* We note that all the extra  $\log(1/\delta)$  terms are due to running the data structures for  $\sqrt{d} \log(1/\delta)$  steps. However, we can reinitialize the data structures every  $\sqrt{d}$  iterations. This decreases the  $K$  dependence from  $K^2$  to  $K\sqrt{d}$ .

*Independence and Adaptive Adversaries:* Randomized data structures often can not handle inputs that depend on outputs of the previous iteration. For example Theorem 4.5 is such a case, where the input to UPDATE is not allowed to depend on the output of any previous calls to UPDATE. In our Algorithm 2 the input to the data-structures inherently depends on their previous output, so here we want to verify that this does not cause any issues.

The data-structures  $D^{(x)}$  and  $D^{(s)}$  are given by Theorem 4.3, which explicitly states that the data-structure works against adaptive adversaries. This means, that the input is allowed to depend on the output of previous iterations. Likewise,  $D_{\text{LEVERAGE}}$  works against adaptive adversaries by Theorem 4.4 (provided the input  $\Psi_{(\text{safe})}^{(\alpha)}$  are chosen by randomness independent of the randomness chosen for  $\Psi^{(\alpha)}$  which is the case by Lemma 4.6). The only issue is with  $D^{-1}$  (given by Theorem 4.5), where the input  $w, \bar{\tau}$  to UPDATE is not allowed to depend on the output of previous calls to UPDATE (however,  $w$  and  $\bar{\tau}$  are allowed to depend on the output of SOLVE by Lemma 4.6). Also note, that the inputs  $\bar{w}, b, \delta$  to SOLVE are allowed to depend on any previous output of UPDATE and SOLVE, as Theorem 4.5 only has issues with the inputs  $w$  and  $\bar{\tau}$  to UPDATE. So to show that our algorithm works, we are only left with verifying that the input  $w, \bar{\tau}$  to UPDATE does not depend on previous results of UPDATE. Let us prove this by induction.

The result of UPDATE is  $\Psi^{(\alpha)} \leftarrow D^{-1} \cdot \text{UPDATE}(\bar{S}^{-1-2\alpha} \bar{x}^{1-2\alpha}, \bar{\tau})$ , and when executing this line for the very first time, it can obviously not depend on a previous output yet. The matrix  $\Psi^{(\alpha)}$  is only used as input to  $D_{\text{LEVERAGE}} \cdot \text{QUERY}(\Psi^{(\alpha)}, \Psi_{(\text{safe})}^{(\alpha)})$ , but by Theorem 4.4 the output of this procedure does not depend on  $\Psi^{(\alpha)}$ . Hence  $\Psi^{(\alpha)}$  does not affect anything else, which also means the input  $\bar{S}^{-1-2\alpha} \bar{x}^{1-2\alpha}$  and  $\bar{\tau}$  to the next call of  $D^{-1} \cdot \text{UPDATE}$  do not depend on previous  $\Psi^{(\alpha)}$ .  $\square$

## 5 OPEN PROBLEMS

Our main result is a linear program solver that runs in expected  $\tilde{O}(nd + d^3)$  time. For dense constraint matrices  $A \in \mathbb{R}^{n \times d}$  this is optimal among algorithms that do not use fast matrix multiplication, barring a major improvement in solving linear systems. The fastest linear system solvers run in  $\tilde{O}(\text{nnz}(A) + d^\omega)$  time, where  $\text{nnz}(A)$  is the number of nonzero entries in  $A$  and  $\omega$  is the matrix exponent. This leads to two open questions: (i) Can the  $nd$  term in our complexity be improved to  $\text{nnz}(A)$ ? (ii) Can the  $d^3$  term be improved to  $d^\omega$  by exploiting fast matrix multiplication?

A major bottleneck for question (i) is how to detect large entries of the product  $Ah$  for  $A \in \mathbb{R}^{n \times d}$ ,  $h \in \mathbb{R}^d$ . Currently the complexity of our data structure for that problem (see full version) is  $\tilde{O}(\|GAh\|^2 \cdot \epsilon^{-2} \cdot d)$ , which can be interpreted as “ $d$  times the number of entries larger than  $\epsilon$ ”. Note that verifying the answer (that is, for a list of indices  $I \subset [n]$ , check that  $(Ah)_i$  is indeed larger than  $\epsilon$ ) requires the same complexity for dense matrices  $A$ . However, for matrices with  $z \leq d$  entries per row, the verification complexity is just  $z \cdot |I|$ . This suggests that it might be possible to improve the data structure to run in  $\tilde{O}(\|GAh\|^2 \cdot \epsilon^{-2} \cdot z)$  for matrices  $A$  with  $z$  nonzero entries per row. If such a data structure could be found, it would result in a faster linear programming algorithm.

So far question (ii) has only been answered for the case  $d = \Omega(n)$  [5, 12] and no progress has been made for the more general case  $d = o(n)$ , even when allowing for a worse dependency on  $n$  than our  $\tilde{O}(nd + d^3)$  bound. So far the best dependency on  $d$  is  $d^{2.5} \gg d^\omega$  [33, 34], so a first step could be to improve our complexity to  $\tilde{O}(nd + d^{2.5})$  time. Currently the  $d^3$  bottleneck comes from maintaining the feasibility of the primal solution  $x$  (Theorem 4.2) so a first step would be to improve the complexity of maintaining the feasibility.

## ACKNOWLEDGEMENTS

We thank Sébastien Bubeck, Ofer Dekel, Jerry Li, Ilya Razenshteyn, and Microsoft Research for facilitating conversations between and hosting researchers involved in this collaboration. We thank Vasileios Nakos, Jelani Nelson, and Zhengyu Wang for very helpful discussions about sparse recovery literature. We thank Jonathan Kelner, Richard Peng, and Sam Chiu-wai Wong for helpful conversations. This project has received funding from the European Research Council (ERC) under the European Unions Horizon 2020 research and innovation programme under grant agreement No 715672. This project was supported in part by NSF awards CCF-1749609, CCF-1740551, DMS-1839116, CCF-1844855, and Microsoft Research Faculty Fellowship. This project was supported in part by Special Year on Optimization, Statistics, and Theoretical Machine Learning (being led by Sanjeev Arora) at Institute for Advanced Study.

## REFERENCES

- [1] Deeksha Adil, Rasmus Kyng, Richard Peng, and Sushant Sachdeva. 2019. Iterative Refinement for lp-norm Regression. In *Proceedings of the Thirtieth Annual ACM-SIAM Symposium on Discrete Algorithms (SODA)*. SIAM, <https://arxiv.org/pdf/1901.06764.pdf>, 1405–1424.
- [2] Kook Jin Ahn, Sudipto Guha, and Andrew McGregor. 2013. Spectral Sparsification in Dynamic Graph Streams. In *Approximation, Randomization, and Combinatorial Optimization. Algorithms and Techniques - 16th International Workshop, APPROX 2013, and 17th International Workshop, RANDOM 2013, Berkeley, CA, USA, August 21-23, 2013. Proceedings*. 1–10.
- [3] Kurt M. Anstreicher. 1996. Volumetric path following algorithms for linear programming. *Math. Program.* 76 (1996), 245–263.
- [4] Jean Bourgain, Joram Lindenstrauss, and V Milman. 1989. Approximation of zonoids by zonotopes. *Acta mathematica* 162, 1 (1989), 73–141.
- [5] Jan van den Brand. 2020. A Deterministic Linear Program Solver in Current Matrix Multiplication Time. In *Proceedings of the Thirty-First Annual ACM-SIAM Symposium on Discrete Algorithms (SODA)*. <https://arxiv.org/pdf/1910.11957.pdf>.
- [6] Jan van den Brand, Danupon Nanongkai, and Thatchaphol Saranurak. 2019. Dynamic Matrix Inverse: Improved Algorithms and Matching Conditional Lower Bounds. In *60th Annual IEEE Symposium on Foundations of Computer Science (FOCS)*. <https://arxiv.org/pdf/1905.05067.pdf>.
- [7] Emmanuel J Candes, Justin K Romberg, and Terence Tao. 2006. Stable signal recovery from incomplete and inaccurate measurements. *Communications on pure and applied mathematics* 59, 8 (2006), 1207–1223.
- [8] Moses Charikar, Kevin Chen, and Martin Farach-Colton. 2002. Finding frequent items in data streams. In *International Colloquium on Automata, Languages, and Programming (ICALP)*. Springer, 693–703.
- [9] Kenneth L Clarkson and David P Woodruff. 2013. Low rank approximation and regression in input sparsity time. In *Proceedings of the 45th annual ACM symposium on Symposium on theory of computing (STOC)*. ACM, <https://arxiv.org/pdf/1207.6365.pdf>, 81–90.
- [10] Michael B. Cohen. 2016. Nearly Tight Oblivious Subspace Embeddings by Trace Inequalities. In *Proceedings of the Twenty-Seventh Annual ACM-SIAM Symposium on Discrete Algorithms (SODA)*. 278–287.
- [11] Michael B Cohen, Yin Tat Lee, Cameron Musco, Christopher Musco, Richard Peng, and Aaron Sidford. 2015. Uniform sampling for matrix approximation. In *Proceedings of the 2015 Conference on Innovations in Theoretical Computer Science (ITCS)*. ACM, <https://arxiv.org/pdf/1408.5099.pdf>, 181–190.
- [12] Michael B Cohen, Yin Tat Lee, and Zhao Song. 2019. Solving linear programs in the current matrix multiplication time. In *Proceedings of the 51st Annual ACM SIGACT Symposium on Theory of Computing (STOC)*. ACM, <https://arxiv.org/pdf/1810.07896.pdf>, 938–942.

- [13] Michael B. Cohen and Richard Peng. 2015.  $\ell_p$  Row Sampling by Lewis Weights. In *Proceedings of the Forty-Seventh Annual ACM Symposium on Theory of Computing, (STOC)*. <https://arxiv.org/pdf/1412.0588.pdf>, 183–192.
- [14] Don Coppersmith and Shmuel Winograd. 1982. On the asymptotic complexity of matrix multiplication. *SIAM J. Comput.* 11, 3 (1982), 472–492.
- [15] Don Coppersmith and Shmuel Winograd. 1990. Matrix Multiplication via Arithmetic Progressions. *J. Symb. Comput.* 9, 3 (1990), 251–280.
- [16] Graham Cormode and Marios Hadjieleftheriou. 2009. Finding the frequent items in streams of data. *Commun. ACM* 52, 10 (2009), 97–105.
- [17] Graham Cormode and Shan Muthukrishnan. 2004. An improved data stream summary: the count-min sketch and its applications. In *Proceedings of the Annual Conference on Foundations of Software Technology and Theoretical Computer Science*.
- [18] George B Dantzig. 1951. Maximization of a linear function of variables subject to linear inequalities. New York (1951).
- [19] David L. Donoho. 2006. Compressed sensing. *IEEE Trans. Information Theory* 52, 4 (2006), 1289–1306.
- [20] David Durfee, Yu Gao, Gramoz Goranci, and Richard Peng. 2019. Fully dynamic spectral vertex sparsifiers and applications. In *Proceedings of the 51st Annual ACM SIGACT Symposium on Theory of Computing (STOC)*. <https://arxiv.org/pdf/1906.10530.pdf>, 914–925.
- [21] François Le Gall. 2014. Powers of tensors and fast matrix multiplication. In *ISSAC*.
- [22] François Le Gall and Florent Urrutia. 2018. Improved Rectangular Matrix Multiplication using Powers of the Coppersmith-Winograd Tensor. In *Proceedings of the Twenty-Ninth Annual ACM-SIAM Symposium on Discrete Algorithms (SODA)*. <https://arxiv.org/pdf/1708.05622.pdf>, 1029–1046.
- [23] Anna C Gilbert, Yi Li, Ely Porat, and Martin J Strauss. 2010. Approximate sparse recovery: optimizing time and measurements. *SIAM Journal on Computing* 2012 (A preliminary version of this paper appears in *STOC* 2010) 41, 2 (2010), 436–453.
- [24] Clovis C Gonzaga. 1992. Path-following methods for linear programming. *SIAM review* 34, 2 (1992), 167–224.
- [25] Haitam Al Hassaneh. 2016. *Sparse recovery and Fourier sampling*. Ph.D. Dissertation. Massachusetts Institute of Technology.
- [26] Daniel M Kane, Jelani Nelson, Ely Porat, and David P Woodruff. 2011. Fast moment estimation in data streams in optimal space. In *Proceedings of the forty-third annual ACM symposium on Theory of computing (STOC)*. ACM, <https://arxiv.org/pdf/1007.4191.pdf>, 745–754.
- [27] Michael Kapralov, Yin Tat Lee, Cameron Musco, Christopher Musco, and Aaron Sidford. 2017. Single Pass Spectral Sparsification in Dynamic Streams. In *SIAM J. Comput.*, Vol. 46(1). <https://arxiv.org/pdf/1407.1289.pdf>, 456–477.
- [28] Michael Kapralov, Navid Nouri, Aaron Sidford, and Jakab Tardos. 2020. Dynamic Streaming Spectral Sparsification in Nearly Linear Time and Space. In *Proceedings of the Thirty-First Annual ACM-SIAM Symposium on Discrete Algorithms (SODA)*. <https://arxiv.org/pdf/1903.12150.pdf>.
- [29] Narendra Karmarkar. 1984. A new polynomial-time algorithm for linear programming. In *Proceedings of the sixteenth annual ACM symposium on Theory of computing (STOC)*. ACM, 302–311.
- [30] Leonid G Khachiyan. 1980. Polynomial algorithms in linear programming. *U. S. S. R. Comput. Math. and Math. Phys.* 20, 1 (1980), 53–72.
- [31] Kasper Green Larsen, Jelani Nelson, Huy L Nguyễn, and Mikkel Thorup. 2016. Heavy hitters via cluster-preserving clustering. In *57th Annual IEEE Symposium on Foundations of Computer Science (FOCS)*. IEEE, <https://arxiv.org/pdf/1604.01357>, 61–70.
- [32] Yin Tat Lee. 2016. *Faster algorithms for convex and combinatorial optimization*. Ph.D. Dissertation. Massachusetts Institute of Technology.
- [33] Yin Tat Lee and Aaron Sidford. 2014. Path-Finding Methods for Linear Programming: Solving Linear Programs in  $\tilde{O}(\sqrt{\text{rank}})$  Iterations and Faster Algorithms for Maximum Flow. In *55th Annual IEEE Symposium on Foundations of Computer Science (FOCS)*. <https://arxiv.org/pdf/1312.6677.pdf>, <https://arxiv.org/pdf/1312.6713.pdf>, 424–433.
- [34] Yin Tat Lee and Aaron Sidford. 2015. Efficient Inverse Maintenance and Faster Algorithms for Linear Programming. In *56th Annual IEEE Symposium on Foundations of Computer Science (FOCS)*. <https://arxiv.org/pdf/1503.01752.pdf>, 230–249.
- [35] Yin Tat Lee and Aaron Sidford. 2019. Solving Linear Programs with  $\sqrt{\text{rank}}$  Linear System Solves. In *arXiv preprint*. <http://arxiv.org/pdf/1910.08033.pdf>. Journal submission.
- [36] Yin Tat Lee, Aaron Sidford, and Sam Chiu-wai Wong. 2015. A Faster Cutting Plane Method and its Implications for Combinatorial and Convex Optimization. In *56th Annual IEEE Symposium on Foundations of Computer Science (FOCS)*.
- [37] Yin Tat Lee, Zhao Song, and Qiuyu Zhang. 2019. Solving Empirical Risk Minimization in the Current Matrix Multiplication Time. In *Conference on Learning Theory (COLT)*. <https://arxiv.org/pdf/1905.04447.pdf>, 2140–2157.
- [38] D. Lewis. 1978. Finite dimensional subspaces of  $L_p$ . *Studia Mathematica* 63, 2 (1978), 207–212. <http://eudml.org/doc/218208>
- [39] Mu Li, Gary L. Miller, and Richard Peng. 2013. Iterative Row Sampling. In *54th Annual IEEE Symposium on Foundations of Computer Science (FOCS)*. <https://arxiv.org/pdf/1211.2713.pdf>, 127–136.
- [40] Michael W. Mahoney. 2011. Randomized Algorithms for Matrices and Data. *Foundations and Trends in Machine Learning* 3, 2 (2011), 123–224.
- [41] Sanjay Mehrotra. 1992. On the Implementation of a Primal-Dual Interior Point Method. *SIAM Journal on Optimization* 2, 4 (1992), 575–601.
- [42] Xiangrui Meng and Michael W Mahoney. 2013. Low-distortion subspace embeddings in input-sparsity time and applications to robust linear regression. In *Proceedings of the 45th annual ACM symposium on Symposium on theory of computing (STOC)*. ACM, <https://arxiv.org/pdf/1210.3135.pdf>, 91–100.
- [43] Vasileios Nakos and Zhao Song. 2019. Stronger L2/L2 Compressed Sensing: Without Iterating. In *Proceedings of the 51st Annual ACM Symposium on Theory of Computing (STOC)*. <https://arxiv.org/pdf/1903.02742.pdf>.
- [44] Jelani Nelson and Huy L Nguyễn. 2013. OSNAP: Faster numerical linear algebra algorithms via sparser subspace embeddings. In *54th Annual IEEE Symposium on Foundations of Computer Science (FOCS)*. <https://arxiv.org/pdf/1211.1002.pdf>, 117–126.
- [45] Jelani Nelson, Huy L Nguyễn, and David P Woodruff. 2014. On deterministic sketching and streaming for sparse recovery and norm estimation. In *Linear Algebra and its Applications*, Vol. 441. <https://arxiv.org/pdf/1206.5725.pdf>, 152–167.
- [46] Jelani Nelson and David P. Woodruff. 2014. Personal communication. . (2014).
- [47] Yurii Nesterov and Arkadii Semenov. 1994. *Interior-point polynomial algorithms in convex programming*. Vol. 13. Society for Industrial and Applied Mathematics.
- [48] Yu Nesterov and Arkadii Nemirovskiy. 1989. *Self-concordant functions and polynomial-time methods in convex programming*. USSR Academy of Sciences, Central Economic & Mathematic Institute.
- [49] Yu Nesterov and A Nemirovsky. 1991. Acceleration and parallelization of the path-following interior point method for a linearly constrained convex quadratic problem. *SIAM Journal on Optimization* 1, 4 (1991), 548–564.
- [50] Rasmus Pagh. 2013. Compressed matrix multiplication. *ACM Transactions on Computation Theory (TOCT)* 5, 3 (2013), 9.
- [51] Eric C Price. 2013. *Sparse recovery and Fourier sampling*. Ph.D. Dissertation. Massachusetts Institute of Technology.
- [52] James Renegar. 1988. A polynomial-time algorithm, based on Newton’s method, for linear programming. *Mathematical Programming* 40, 1-3 (1988), 59–93.
- [53] Piotr Sankowski. 2004. Dynamic Transitive Closure via Dynamic Matrix Inverse (Extended Abstract). In *45th Annual IEEE Symposium on Foundations of Computer Science (FOCS)*. 509–517.
- [54] Aaron Daniel Sidford. 2015. *Iterative methods, combinatorial optimization, and linear programming beyond the universal barrier*. Ph.D. Dissertation. Massachusetts Institute of Technology.
- [55] Zhao Song. 2019. *Matrix Theory: Optimization, Concentration and Algorithms*. Ph.D. Dissertation. The University of Texas at Austin.
- [56] Zhao Song, David P Woodruff, and Peilin Zhong. 2019. Relative Error Tensor Low Rank Approximation. In *Proceedings of the Thirtieth Annual ACM-SIAM Symposium on Discrete Algorithms (SODA)*. <https://arxiv.org/pdf/1704.08246.pdf>.
- [57] Daniel A Spielman and Nikhil Srivastava. 2011. Graph sparsification by effective resistances. *SIAM J. Comput.* 40, 6 (2011), 1913–1926.
- [58] Andrew James Stothers. 2010. On the complexity of matrix multiplication. (2010).
- [59] Volker Strassen. 1969. Gaussian elimination is not optimal. *Numerische mathematik* 13, 4 (1969), 354–356.
- [60] Volker Strassen. 1986. The asymptotic spectrum of tensors and the exponent of matrix multiplication. In *27th Annual IEEE Symposium on Foundations of Computer Science (FOCS)*. 49–54.
- [61] Pravin M. Vaidya. 1989. A New Algorithm for Minimizing Convex Functions over Convex Sets (Extended Abstract). In *30th Annual IEEE Symposium on Foundations of Computer Science (FOCS)*. 338–343.
- [62] Pravin M Vaidya. 1989. Speeding-up linear programming using fast matrix multiplication. In *30th Annual IEEE Symposium on Foundations of Computer Science (FOCS)*. 332–337.
- [63] Pravin M. Vaidya. 1990. Reducing the Parallel Complexity of Certain Linear Programming Problems (Extended Abstract). In *31st Annual IEEE Symposium on Foundations of Computer Science (FOCS)*. 583–589.
- [64] Pravin M Vaidya and David S Atkinson. 1993. A technique for bounding the number of iterations in path following algorithms. *Complexity in Numerical Optimization* (1993), 462–489.
- [65] Jan van den Brand, Yin Tat Lee, Aaron Sidford, and Zhao Song. 2020. Solving Tall Dense Linear Programs in Nearly Linear Time. *CoRR* abs/2002.02304 (2020).
- [66] Virginia Vassilevska Williams. 2012. Multiplying matrices faster than Coppersmith-Winograd. In *Proceedings of the forty-fourth annual ACM symposium on Theory of computing (STOC)*. ACM, 887–898.
- [67] David P. Woodruff. 2014. Sketching as a Tool for Numerical Linear Algebra. *Foundations and Trends in Theoretical Computer Science* 10, 1-2 (2014), 1–157.
- [68] Yinyu Ye. 2011. *Interior point algorithms: theory and analysis*. Vol. 44. John Wiley & Sons.
- [69] Yinyu Ye, Michael J Todd, and Shinji Mizuno. 1994. An  $O(\sqrt{nL})$ -iteration homogeneous and self-dual linear programming algorithm. *Mathematics of Operations Research* 19, 1 (1994), 53–67.