

Packing LPs are Hard to Solve Accurately, Assuming Linear Equations are Hard

Rasmus Kyng ^{*}
Department of Computer Science
ETH Zurich

Di Wang [†]
Google

Peng Zhang [‡]
Department of Computer Science
Yale University

Abstract

We study the complexity of approximately solving packing linear programs. In the Real RAM model, it is known how to solve packing LPs with N non-zeros in time $\tilde{O}(N/\epsilon)$. We investigate whether the ϵ dependence in the running time can be improved.

Our first main result relates the difficulty of this problem to hardness assumptions for solving dense linear equations. We show that, in the Real RAM model, unless linear equations in matrices $n \times n$ with condition number $O(n^{10})$ can be solved to ϵ accuracy faster than $\tilde{O}(n^{2.01} \log(1/\epsilon))$, no algorithm $(1-\epsilon)$ -approximately solves a $O(n) \times O(n)$ packing LPs (where $N = O(n^2)$) in time $\tilde{O}(n^2 \epsilon^{-0.0003})$. It would be surprising to solve linear equations in the Real RAM model this fast, as we currently cannot solve them faster than $\tilde{O}(n^\omega)$, where ω denotes the exponent in the running time for matrix multiplication in the Real RAM model (and equivalently matrix inversion). The current best bound on this exponent is roughly $\omega \leq 2.372$. Note, however, that a fast solver for linear equations does not directly imply faster matrix multiplication. But, our reduction shows that if fast and accurate packing LP solvers exist, then either linear equations can be solved much faster than matrix multiplication or the matrix multiplication constant is very close to 2.

Instantiating the same reduction with different parameters, we show that unless linear equations in matrices with condition number $O(n^{1.5})$ can be solved to ϵ accuracy faster than $\tilde{O}(n^{2.372} \log(1/\epsilon))$, no algorithm $(1-\epsilon)$ -approximately solves packing LPs in time $\tilde{O}(n^2 \epsilon^{-0.067})$. Thus smaller improvements in the exponent for ϵ in the running time of Packing LP solvers also imply improvements in the current state-of-the-art for solving linear equations.

Our second main result relates the difficulty of approximately solving packing linear programs to hardness assumptions for solving sparse linear equations: In the Real RAM model, unless well-conditioned sparse systems of linear equations can be solved faster than $\tilde{O}((\text{no. non-zeros of matrix}) \sqrt{\text{condition number of matrix}})$, no algorithm $(1-\epsilon)$ -approximately solves packing LPs with N non-zeros in time $\tilde{O}(N \epsilon^{-0.165})$. This running time of $\tilde{O}((\text{no. non-zeros of matrix}) \sqrt{\text{condition number of matrix}})$ is obtained by the classical Conjugate Gradient algorithm

by a standard analysis. Our reduction implies that if sufficiently good packing LP solvers exist, then this long-standing best-known bound on the running time for solving well-conditioned systems of linear equations is sub-optimal¹. While we prove results in the Real RAM model, our condition number assumptions ensure that our results can be translated to fixed point arithmetic with $(\log n)^{O(1)}$ bits per number.

1 Introduction

Packing and covering linear programs (LPs) are LPs formulated with non-negative coefficients, constants and variables. Formally, a packing LP in its generic form can be written as

$$(1.1) \quad \max_{x \geq 0} \{c^\top x : Ax \leq b\},$$

where $c \in \mathbb{R}_{\geq 0}^n$, $b \in \mathbb{R}_{\geq 0}^m$ and $A \in \mathbb{R}_{\geq 0}^{m \times n}$. With the same matrix and vectors A, b, c , a covering LP can be written as

$$\min_{y \geq 0} \{b^\top y : A^\top y \geq c\}.$$

Packing LPs and covering LPs are dual to each other, and thus they have the same optimal value. In general, taken non-negative A, b, c as input, a packing / covering LP solver returns both a packing LP solution and a covering LP solution as a primal-dual pair [LN93, KY14, AZO15b]. Thus packing LPs and covering LPs are usually considered equivalent from algorithm design or hardness perspectives. In this work we will focus on packing LPs. We remark that by a slightly modification of our reduction, the same lower bounds hold for covering LPs as well.²

A packing LP is always feasible and has optimum at least 0 because of the all 0 solution. Let OPT denote

¹This bound is the best known for solving sparse linear equations when only a bound on the condition number is known. In other regimes, better results are known.

²We note that packing and covering LPs are also called positive LPs in the literature, although the usage of positive LPs is also used to refer to the more general class of LPs where both packing (i.e. $\leq$) constraints and covering (i.e. $\geq$) constraints appear simultaneously.

^{*}email: kyng@inf.ethz.ch.

[†]email: di.wang@cc.gatech.edu.

[‡]email: peng.zhang@yale.edu. Supported in part by NSF Grant CCF-1562041, Office of Naval Research Award N00014-16-2374, and a Simons Investigator award from the Simons Foundation to Daniel Spielman.

the optimal value of a packing LP of the form (1.1). Given an error parameter $\epsilon > 0$, we say $\mathbf{x}$ is a $(1 - \epsilon)$ -approximate solution to this packing LP iff $\mathbf{Ax} \leq \mathbf{b}$ and $\mathbf{c}^\top \mathbf{x} \geq (1 - \epsilon)\text{OPT}$.

Throughout this paper, we are interested in nearly-linear time solvers for $(1 - \epsilon)$ -approximately solving packing LPs in the Real RAM model³ (see [PS85] for an introduction to the Real RAM model). Specifically, the run time is $\tilde{O}(N/\epsilon^c)$, where N is the number of non-zeros of the input and $c > 0$ is an absolute constant. We use $\tilde{O}(\cdot)$ to hide $\text{polylog}(N/\epsilon)$ factors. Although algorithms for solving general LPs, such as interior point method and ellipsoid method, can be applied to packing LPs with a $\text{polylog}(1/\epsilon)$ dependence on the run time, these methods usually have large dependence on input size and thus are not suitable for large-scale instances.

Small dependency on $1/\epsilon$ is meaningful in two aspects. First, the dependence on $1/\epsilon$ is a natural measure of the efficiency of iterative methods, where a $1/\epsilon^c$ dependence indicates that to get one more bit of accuracy, the algorithm needs to do 2^c times more work. Second, as packing and covering LPs are often used as subroutines in algorithm design (e.g., bipartite matching, set cover [LN93], scheduling [PST95], multicommodity flow [GK07, Mad10, AHK12], etc.), the dependence on $1/\epsilon$ determines the complexity of the overall algorithm. Sometimes, ϵ is chosen to get the best trade-off among multiple parts of the overall algorithm. One example is bipartite matching [Wan17]. Given a bipartite graph with n vertices and m edges, computing its maximum matching can be formulated as a packing LP with $O(m)$ non-zeros. Suppose one can $(1 - \epsilon)$ -approximately solve this packing LP in time $O(m/\epsilon^c)$ for some constant c . To turn this approximate solution to a maximum matching, one needs to compute $O(\epsilon n)$ augmenting paths in time $O(m \cdot \epsilon n)$. Setting $\epsilon = n^{-1/(1+c)}$, the total run time is $O(mn^{c/(1+c)})$. Faster approximate solvers for packing LPs would imply better run time for this algorithm.

Designing nearly linear time approximate solvers for packing and covering LPs was initiated by the seminal work of Luby and Nisan [LN93], which proposed an algorithm with run time $\tilde{O}(N/\epsilon^4)$. The run time dependence on $1/\epsilon$ was then improved by subsequent works including [You01, You14, KY14, AZO15b, AZO15a, WRM16, MRWZ16, CQ18], etc. The best known sequential algorithm has expected run time $\tilde{O}(N/\epsilon)$ [AZO15a, WRM16].

Therefore, a natural question is: Can we design a

$(1 - \epsilon)$ -approximate solver for packing LPs with run time $\tilde{O}(N/\epsilon^{o(1)})$? In this paper, we give a negative answer, conditioned on that linear equations cannot be solved fast.

Solving linear equations is important in numerical linear algebra, and is a central tool in computer science, statistics, economics, physics, and engineering. It is used as a subroutine of interior point method when applied to solving a general LP. Using interior point method, solving a polynomially well-conditioned LP can be reduced to solving $\tilde{O}(\sqrt{\text{rank}(\mathbf{A})})$ linear equations of the same size as the input [LS14], where $\mathbf{A}$ is the LP coefficient matrix. Improvements in linear equation solvers would directly imply improvements for LP solvers in some regimes, and hence would imply improvements for a large class of convex optimization problems.

Our results. We give a conditional lower bound, by relating solving a packing LP to solving a system of linear equations $\mathbf{Ax} = \mathbf{b}$. Specially, given an $n \times n$ matrix $\mathbf{A}$ with full rank, an n -dimensional column vector $\mathbf{b}$, an error parameter $\epsilon \geq 0$, the goal is to compute a vector $\mathbf{x}$ such that $\|\mathbf{Ax} - \mathbf{b}\|_2 \leq \epsilon \|\mathbf{b}\|_2$.

We introduce (several variants of) the Linear Equation Time Hypothesis LTH: Firstly, $\text{LTH}_{1.5}^{2.372}$ is the assumption that an arbitrary linear equation instance $\mathbf{Ax} = \mathbf{b}$, where $\mathbf{A}$ is an $n \times n$ matrix with condition number of $\mathbf{A}^\top \mathbf{A}$ at most $O(n^{2.15})$ cannot be solved to error $\epsilon \leq n^{-15}$ faster than time $\tilde{O}(n^{2.372})$ in the Real RAM model⁴ $\text{LTH}_{1.5}^{2.372}$ is a major open question in numerical linear algebra and scientific computing since falsifying it implies we can solve linear equations faster than current best known bounds on matrix multiplication. The exponent ω is defined as the number s.t. matrix multiplication can be computed in time bounded by $O(n^\omega)$ in the Real RAM model, and the current best bound is $\omega \leq 2.372 \dots$. Again, in the Real RAM model, is known that one can exactly solve $\mathbf{Ax} = \mathbf{b}$ by simply inverting $\mathbf{A}$ and then multiplying $\mathbf{A}^{-1}\mathbf{b}$. We show that in the Real RAM model no algorithm can solve Packing LPs with N non-zero entries in time $\tilde{O}(N/\epsilon^{0.067})$ unless $\text{LTH}_{1.5}^{2.372}$ is false.

More generally, we define LTH_k^γ as the assumption that an arbitrary linear equation instance $\mathbf{Ax} = \mathbf{b}$, where $\mathbf{A}$ is an $n \times n$ matrix with condition number of $\mathbf{A}^\top \mathbf{A}$ at most $O(n^{2k})$, cannot be solved in time $\tilde{O}(n^\gamma)$ in

³Although we prove results in the Real RAM model, our condition number assumptions ensure that our results can be translated to fixed point arithmetic with $(\log n)^{O(1)}$ bits per number.

⁴Note that in the Real RAM model, using matrix multiplication-based matrix inversion, one can in fact solve a non-singular linear equation exactly in time $\tilde{O}(n^\omega)$. However, even in fixed point arithmetic, given our assumptions about $\mathbf{A}$ and the desired accuracy of the solution, the linear equation can be solved to the desired accuracy in $\tilde{O}(n^\omega)$ arithmetic operations on numbers with $(\log n)^{O(1)}$ bits each [DDH07].

the Real RAM model. We show that in the Real RAM model, no algorithm can solve Packing LPs with N non-zero entries in time $\tilde{O}(N/\epsilon^\alpha)$ unless $\text{LTH}_k^{2+\alpha+3\alpha k}$ is false. By instantiating our result for different parameter values, we can state our first main result informally as

THEOREM 1.1. *[Informal Statement] In the Real RAM model,*

- *unless linear equations in matrices with condition number $O(n^{10})$ can be solved to ϵ accuracy faster than $\tilde{O}(n^{2.01} \log(1/\epsilon))$, no algorithm $(1 - \epsilon)$ -approximately solves packing LPs in time $\tilde{O}(N\epsilon^{-0.0003})$.*
- *unless linear equations in matrices with condition number $O(n^{1.5})$ can be solved to ϵ accuracy faster than $\tilde{O}(n^{2.372} \log(1/\epsilon))$, no algorithm $(1 - \epsilon)$ -approximately solves packing LPs in time $\tilde{O}(N\epsilon^{-0.067})$.*

An arguably even more important question in numerical linear algebra and scientific computing is how quickly sparse linear systems can be solved. We define SLTH_k^γ as the assumption that an arbitrary linear equation instance $\mathbf{A}^\top \mathbf{A} \mathbf{x} = \mathbf{A}^\top \mathbf{b}$ cannot be solved to high accuracy faster than time $\tilde{O}(\text{nnz}(\mathbf{A})^\gamma)$, when $\mathbf{A}$ is an $m \times n$ matrix s.t. $\mathbf{A}^\top \mathbf{A}$ has condition number n^{2k} . Conjugate Gradient⁵ rules out SLTH_k^{1+k} , but falsifying $\text{SLTH}_k^{1+0.99k}$ would give the first improvement over Conjugate Gradient for this problem since 1952 [HS52].

We show that no algorithm can solve packing LPs with N non-zero entries in time $\tilde{O}(N/\epsilon^\alpha)$ unless $\text{SLTH}_k^{1+1.5\alpha+3\alpha k}$ is false. This has interesting consequences for sparse matrices, e.g. $\text{nnz}(\mathbf{A}) \approx n \log n$ with small condition number, e.g. at most $n^{1.5}$. Both Conjugate Gradient and matrix-inversion by matrix multiplication (even with $\omega = 2$) are consistent with $\text{SLTH}_{1.5}^{1.99}$, and no algorithm can solve Packing LPs with N non-zero entries in time $\tilde{O}(N/\epsilon^{0.165})$ unless $\text{SLTH}_{1.5}^{1.99}$ is false. We can summarize this second main result informally as

THEOREM 1.2. *[Informal Statement] In the Real RAM model, unless linear equations $\mathbf{A}^\top \mathbf{A} \mathbf{x} = \mathbf{A}^\top \mathbf{b}$ can be solved to high accuracy faster than time $\tilde{O}(\text{nnz}(\mathbf{A}) \sqrt{\text{condition number of } \mathbf{A}^\top \mathbf{A}})$ when $\mathbf{A}^\top \mathbf{A}$ has condition number $\Theta(n^3)$, then no algorithm $(1 - \epsilon)$ -approximately solves packing LPs in time $\tilde{O}(N\epsilon^{-0.165})$.*

⁵To understand the condition number dependence, see Section 3. Running time of Conjugate Gradient scales like $\sqrt{\text{condition number of } \mathbf{A}^\top \mathbf{A}}$, which by our definition of condition number becomes scaling like condition number of $\mathbf{A}$.

Under the assumptions of the informal theorem statement, the Conjugate Gradient algorithm has a worst case running time of $\tilde{O}(\text{nnz}(\mathbf{A}) \sqrt{\text{condition number of } \mathbf{A}^\top \mathbf{A}})$, so the informal theorem also tells us that Conjugate Gradient is suboptimal in this regime if sufficiently good packing LP solvers exist.

While we prove results in the Real RAM model, our condition number assumptions ensure that our results can be translated to fixed point arithmetic with $(\log n)^{O(1)}$ bits per number under appropriate assumptions (see Remark 3.1).

Instead of exploring consequences of having access to a packing LP solver with N non-zero entries in time $\tilde{O}(N/\epsilon^\alpha)$ for small constant α , we could also look at consequences of having a $\tilde{O}(N \log(1/\epsilon))$ solver. Using our reductions, one can then show that general LPs, with polynomially bounded condition numbers and bit complexity, can be solved in $\tilde{O}(N \log(1/\epsilon))$ time (see Appendix D).

Our reduction. We will first reduce a linear system instance to a linear program instance with non-negative variables, by standard techniques. Given $\mathbf{A} \in \mathbb{R}^{n \times n}$ and $\mathbf{b} \in \mathbb{R}^n$, solving $\min_{\mathbf{x} \in \mathbb{R}^n} \|\mathbf{A} \mathbf{x} - \mathbf{b}\|_2$ is equivalent to solving the two equations: $\mathbf{A}^\top \mathbf{z} = \mathbf{A}^\top \mathbf{b}$, $\mathbf{A} \mathbf{x} = \mathbf{z}$, where $\mathbf{x} \in \mathbb{R}^n, \mathbf{z} \in \mathbb{R}^n$ are variables. Note these two linear equations are always feasible. Let $\mathbf{x}^+, \mathbf{x}^-, \mathbf{z}^+, \mathbf{z}^- \in \mathbb{R}_{\geq 0}^n$. We replace $\mathbf{x}$ by $\mathbf{x}^+ - \mathbf{x}^-$ and replace $\mathbf{z}$ by $\mathbf{z}^+ - \mathbf{z}^-$, and write an equality constraint as two inequality constraints:

$$\begin{aligned} \mathbf{A}^\top (\mathbf{z}^+ - \mathbf{z}^-) &\leq \mathbf{A}^\top \mathbf{b} \\ -\mathbf{A}^\top (\mathbf{z}^+ - \mathbf{z}^-) &\leq -\mathbf{A}^\top \mathbf{b} \\ \mathbf{A} (\mathbf{x}^+ - \mathbf{x}^-) - (\mathbf{z}^+ - \mathbf{z}^-) &\leq \mathbf{0} \\ -\mathbf{A} (\mathbf{x}^+ - \mathbf{x}^-) + (\mathbf{z}^+ - \mathbf{z}^-) &\leq \mathbf{0} \end{aligned}$$

For notation simplicity, we will write the above linear inequalities in its matrix form: $\mathbf{C} \mathbf{y} \leq \mathbf{p}$, where $\mathbf{y} \in \mathbb{R}_{\geq 0}^{4n}$ contains non-negative variables. Wlog, we can scale $\mathbf{C}, \mathbf{p}$ so that all entries of $\mathbf{C}$ are between -1 and 1, and $\|\mathbf{p}\|_2 = 1$.

The key of our reduction is to reduce the linear program $\mathbf{C} \mathbf{y} \leq \mathbf{p}$ with non-negative variables to a packing LP which has both non-negative variables and non-negative coefficients. We will convert entries of $\mathbf{C}, \mathbf{p}$ to non-negative entries by a bounding box constraint. Specifically, we introduce a new non-negative variable $y_0 \in \mathbb{R}_{\geq 0}$, and a new constraint

$$(1.2) \quad \mathbf{1}^\top \mathbf{y} + y_0 = U,$$

where $\mathbf{1}$ is the all-one vector, and U is a sufficiently large number so that the linear program $\mathbf{C} \mathbf{y} \leq \mathbf{p}, \mathbf{1}^\top \mathbf{y} \leq U$

is feasible. One can use binary search to find a value $U \geq 1$ which satisfies the above conditions and is upper bounded by $\text{poly}(n, \kappa)$, where κ is the condition number of $\mathbf{A}$. We then add the equality (1.2) to each row of $\mathbf{C}\mathbf{y} \leq \mathbf{p}$, and we get

$$(\mathbf{C} + \mathbf{1}\mathbf{1}^\top)\mathbf{y} + y_0\mathbf{1} \leq \mathbf{p} + U\mathbf{1}.$$

Since the entries of $\mathbf{C}, \mathbf{p}$ are between -1 and 1, the above inequality has non-negative coefficients. Finally, we turn the equality constraint (1.2) to an inequality constraint and set the objective to maximizing $\mathbf{1}^\top \mathbf{y} + y_0$:

$$(1.3) \quad \begin{aligned} \max \quad & \mathbf{1}^\top \mathbf{y} + y_0 \\ \text{s.t.} \quad & (\mathbf{C} + \mathbf{1}\mathbf{1}^\top)\mathbf{y} + y_0\mathbf{1} \leq \mathbf{p} + U\mathbf{1} \\ & \mathbf{1}^\top \mathbf{y} + y_0 \leq U \\ & y_0 \geq 0, \mathbf{y} \geq \mathbf{0} \end{aligned}$$

The above LP has size $O(n)$ and the reduction can be implemented in time $O(n)$. In addition, its optimal value is U and its optimal solution $\mathbf{y}$ is a feasible solution to the linear system $\mathbf{C}\mathbf{y} = \mathbf{p}$.

Suppose we are given a $(1 - \epsilon)$ -approximate solution $\mathbf{y}$ of the packing LP (1.3). We can turn it to an approximate solution of $\min_{\mathbf{x}} \|\mathbf{A}\mathbf{x} - \mathbf{b}\|_2$ with error $\epsilon' = \text{poly}(\epsilon, n, \kappa)$. This already implies that if one can solve a packing LP with error ϵ in time $O(n^2 \log(1/\epsilon))$, then one can solve a linear system with error ϵ' in time $O(n^2 \log(n\kappa/\epsilon'))$, which falsifies the hypothesis LTH_{10}^2 .

We further improve the error dependence between the linear system instance and the packing LP instance, by using iterative refinement for linear equations. Specifically, one could solve a system of linear equations with error ϵ' , by iteratively solving a sequence of $O(\log(1/\epsilon'))$ linear systems with error 0.1. To compute an approximate solution of $\min_{\mathbf{x}} \|\mathbf{A}\mathbf{x} - \mathbf{b}\|_2$ with error ϵ' , instead of solving one packing LP with a small error, we will solve $O(\log(1/\epsilon'))$ packing LPs with a relatively larger error independent of ϵ' . By a careful calculation, we prove Theorem 1.1.

We prove Theorem 1.2 by slightly modifying the above reduction to preserve the sparsity of $\mathbf{C}$: for each row of $\mathbf{C}\mathbf{y} \leq \mathbf{p}$, we add a bounding box constraint containing only the variables with non-zero coefficients in this row.

Discussion of our reduction. One should take a moment to reflect on the elements of our construction. Linear equation solvers allow for “accuracy amplification” via iterative refinement. This turns a constant error algorithm into a high accuracy algorithm, provided one is careful to adopt the right notion of error⁶. This

in turn allows us to get a type of hardness amplification when reducing systems of linear equations to packing LPs. The most important aspect of our approach is that we are able to turn crude packing LP solvers into crude linear equation solvers of the form that allow of this “accuracy amplification”, and hence we are able to derive surprisingly strong consequences from the existence of packing LP solvers with running time scaling as $\text{poly}(1/\epsilon)$.

This should be contrasted with other approaches to hardness amplification in fine-grained complexity, such as the Distributed PCP framework of [ARW17], which was used for showing hardness of approximation for subquadratic time algorithms for problems including finding a maximum inner product pair among a set of vectors. It is not clear how to apply this in a setting where algorithms with highly subquadratic running times in the input size are known to exist. We hope that our approach of showing hardness using a reduction from *constant error* linear equation solving will prove fruitful for other questions in fine-grained hardness of approximation.

1.1 Previous Works Previous works on packing and covering LPs can be divided into width-dependent solvers and width-independent solvers. In this context, width is the product of OPT and the largest entry in $\mathbf{A}$. The dependence of width usually comes from multiplicative weights update (e.g., see [PST95, AHK12]). In general, width-dependent solvers can be very slow when the width is large.

The line of width-independent nearly linear time solvers for packing and covering LPs was initiated by the seminal work of Luby and Nisan [LN93], which gave a $\tilde{O}(N/\epsilon^4)$ time algorithm for a packing or covering LP with N non-zeros. Their algorithm can be made parallel with $\tilde{O}(1/\epsilon^4)$ time and $\tilde{O}(N/\epsilon^4)$ total work. Subsequent work along this line of research mostly focus on getting better dependence on $1/\epsilon$ (for example, refer to [You01, You14, KY14, AZO15b, AZO15a, WRM16, MRWZ16, CQ18], etc). Currently, the fastest parallel algorithm for packing and covering LPs runs in $\tilde{O}(1/\epsilon^2)$ time and $\tilde{O}(N/\epsilon^2)$ work ([AZO15b, MRWZ16]). The fastest sequential algorithm has expected running time $\tilde{O}(N/\epsilon)$ and with at least a constant probability solves a packing or covering LP with multiplicative error ϵ

build from a packing LP solver is not guaranteed to act as a linear operator, we make sure it achieves multiplicative error strictly less than 1 to before we use it for iterative refinement. This is in contrast to the more common preconditioning of linear systems by linear operators which allows for “accuracy amplification” even when the preconditioning, corresponding to our crude solver, has much higher error.

⁶This well-known phenomenon is discussed in Appendix B. Note a subtle point: because the crude linear equation solver we

([AZO15a, WRM16]).

From the complexity perspective, the main focuses of previous works are on fast parallel solvers. Trevisan and Xhafa [TX98] showed that, *exactly* solving packing LPs is P-complete, that is, there is no fast parallel exact solvers for packing LP unless $P = NC$. However, their result does not rule out any fast algorithms for *approximately* solving packing LPs. Both exact and approximate solvers for general LPs were proved to be P-complete [DLR79, Ser91, Meg92].

Exactly solvers for the worst cases of many combinatorial problems, geometric problems can be slow. Itai [Ita78] showed that LPs are polynomially equivalent to the following problems: Linear equalities with / without bounded coefficients, Homologous flow, and 2-commodity flow problem. Later, Dobkin and Reiss [DR80] gave more LP-complete problems in computational geometry, such as determining intersection of hyperplanes, finding extreme point, and so on. Allowing some error usually significantly improves running time.

When solving general LPs in the Real RAM model, different results are optimal in different regimes. All known polynomial time algorithms require bounds on bit complexity of the input. By [CLS19], if $\omega \leq 2 + 1/6$ (and another condition on the so-called dual exponent is satisfied), then polynomially well-conditioned LPs can be solved in $\tilde{O}(n^{2+1/6})$ time using matrix inverse maintenance-based techniques, and using currently known matrix multiplication time, they get a bound of $\tilde{O}(n^{2.372\dots})$, matching the current bounds on ω . Solving a polynomially well-conditioned LP can be reduced to solving $\tilde{O}(\sqrt{\text{rank}(\mathbf{A})})$ linear equations of the same size as the input [LS14], where $\mathbf{A}$ is the LP coefficient matrix.

Fine-grained complexity [WW10] has motivated the study of conditional lower bounds related to matrix multiplication, and linear equation solving. For example, [MNS⁺17] showed that, fast algorithms for high accuracy spectrum approximation problems, such as logarithm of matrix determinant, trace of matrix inverse, and trace of matrix exponential, would imply triangle detection algorithms for general graphs running in faster than the state of art matrix multiplication time. In addition, [KZ17] proved that, if one can solve some structured linear equations fast, then one can solve general linear equations fast.

1.2 Organization of the Remaining Paper In Section 2, we give some notations and discuss approximately solving linear equations. In Section 3, we formally define Linear Equation Time Hypotheses, and state our main results for lower bounds of packing LPs. In Section 4, we prove the main results.

2 Preliminaries

2.1 Notations We use $\|\cdot\|_2$ to denote the Euclidean norm on vectors and the spectral norm on matrices. When $\mathbf{M}$ is an $n \times n$ positive semidefinite matrix, we define the $\mathbf{M}$ -norm on a vector $\mathbf{x} \in \mathbb{R}^n$ by $\|\mathbf{x}\|_{\mathbf{M}} \stackrel{\text{def}}{=} \sqrt{\mathbf{x}^\top \mathbf{M} \mathbf{x}}$. Let $\text{nnz}(\mathbf{A})$ denote the number of non-zero entries in a matrix $\mathbf{A}$. We define $\|\mathbf{A}\|_{\max} = \max_{i,j} |\mathbf{A}_{ij}|$.

Given a matrix $\mathbf{A} \in \mathbb{R}^{m \times n}$ and a vector $\mathbf{c} \in \mathbb{R}^m$ for some m, n , we call the tuple $(\mathbf{A}, \mathbf{c})$ a linear system. Given matrix $\mathbf{A} \in \mathbb{R}^{m \times n}$, let $\text{im}(\mathbf{A})$ denote the image of $\mathbf{A}$, and let $\Pi_{\mathbf{A}} \stackrel{\text{def}}{=} \mathbf{A}(\mathbf{A}\mathbf{A}^\top)^\dagger \mathbf{A}^\top$, i.e. the orthogonal projection onto $\text{im}(\mathbf{A})$. Note that $\Pi_{\mathbf{A}} = \Pi_{\mathbf{A}}^\top$ and $\Pi_{\mathbf{A}} = \Pi_{\mathbf{A}}^2$. We define the maximum singular value $\sigma_{\max}(\mathbf{A})$ in the usual way as $\sigma_{\max}(\mathbf{A}) = \max_{\mathbf{x} \in \mathbb{R}^n, \mathbf{x} \neq \mathbf{0}} \sqrt{\frac{\mathbf{x}^\top \mathbf{A}^\top \mathbf{A} \mathbf{x}}{\mathbf{x}^\top \mathbf{x}}}$. We define the minimum singular value as $\sigma_{\min}(\mathbf{A}) = \min_{\mathbf{x} \in \mathbb{R}^n, \mathbf{x} \neq \mathbf{0}} \sqrt{\frac{\mathbf{x}^\top \mathbf{A}^\top \mathbf{A} \mathbf{x}}{\mathbf{x}^\top \mathbf{x}}}$.

We define the condition number of $\mathbf{A}$ as $\kappa(\mathbf{A}) = \frac{\sigma_{\max}(\mathbf{A})}{\sigma_{\min}(\mathbf{A})}$.

2.2 Approximately Solving A Linear Equation

In this section we formally define the notions of approximate solutions to linear systems that we work with throughout this paper. Suppose $\mathbf{A} \in \mathbb{R}^{m \times n}$ with $m \geq n$ and $\text{rank}(\mathbf{A}) = n$. In general, the linear equation $\mathbf{A}\mathbf{x} = \mathbf{b}$ may not have a solution. A reasonable generalization of solving the linear equation is then to solve the problem

$$(2.4) \quad \arg \min_{\mathbf{x} \in \mathbb{R}^n} \|\mathbf{A}\mathbf{x} - \mathbf{b}\|_2^2$$

Defining a useful notion of an *approximate* solution to the above optimization problem requires some care. We use the following definition. [Linear Equation Approximation Problem, LEA] Given linear system $(\mathbf{A}, \mathbf{b})$, where $\mathbf{A} \in \mathbb{R}^{m \times n}$, and $\mathbf{b} \in \mathbb{R}^m$, and given a scalar $0 \leq \epsilon \leq 1$, we refer to the LEA problem for the triple $(\mathbf{A}, \mathbf{b}, \epsilon)$ as the problem of finding $\mathbf{x} \in \mathbb{R}^n$ s.t.

$$\|\mathbf{A}\mathbf{x} - \Pi_{\mathbf{A}}\mathbf{b}\|_2^2 \leq \epsilon \|\Pi_{\mathbf{A}}\mathbf{b}\|_2^2,$$

and we say that such an $\mathbf{x}$ is a solution to the LEA instance $(\mathbf{A}, \mathbf{b}, \epsilon)$.

This definition of LEA is widely used in the literature on solving systems of linear equations. When $\mathbf{b}$ is in the image of $\mathbf{A}$, this definition is equivalent to obtaining a solution that is within a $1 + \epsilon$ multiplicative factor of the optimum value of Problem (2.4), but these notions are not equivalent when $\mathbf{b}$ is not contained in the image of $\mathbf{A}$. In Appendix A, we give more background on this notion of approximately solving linear equations for readers unfamiliar with the subject. Below, we state an important fact from this appendix, which should give

the reader an impression of why the LEA definition is natural. In particular, the definition is equivalent to several other convenient notions of error.

FACT 2.1. *The vector $\mathbf{x}$ being a solution to $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon)$ is equivalent to*

$$\|\mathbf{x} - \mathbf{x}^*\|_{\mathbf{A}^T \mathbf{A}}^2 \leq \epsilon \|\mathbf{x}^*\|_{\mathbf{A}^T \mathbf{A}}^2 = \epsilon \left\| \mathbf{A}^T \mathbf{b} \right\|_{(\mathbf{A}^T \mathbf{A})^{-1}}^2$$

Our notion of LEA is useful to work with because an algorithm with this guarantee can self-amplify its accuracy

LEMMA 2.1. *If we have a solver for $\text{LEA}(\mathbf{A}, \mathbf{b}, 0.1)$ that works for arbitrary $\mathbf{b}$, then we can obtain a solver for $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon)$ by iterating it $O(\log(1/\epsilon))$ times.*

This Lemma is well-known, but for completeness we prove it in Appendix B.

3 Linear Equation Hardness Assumptions and Consequences

In this section we cover the state of the art for solving systems of linear equations, and define natural hardness assumptions for solving systems of linear equations. We then state our main results on reductions from linear equation problems to solving Packing LPs and derive some important consequences.

Linear Equations in the Real RAM model.

Since [Str69], it has been known that the inverse a square, full rank, $n \times n$ matrix $\mathbf{A}$ can be computed using $\tilde{O}(n^\omega)$ time in the Real RAM model where ω is the *matrix multiplication constant*: the exponent of the fastest algorithm for multiplying two $n \times n$ matrices, again in the Real RAM model. Currently the best known upper bound on the exponent is $\omega < 2.372\dots$ [LG14]. In general, to output the product of two $n \times n$ matrices requires specifying n^2 entries of the output, so $\omega \geq 2$. It is, however, known that matrix multiplication requires $n^2 \log n$ operations in an arithmetic circuit model [Raz02].

In the Real RAM model, it is also known that if one can compute the inverse of a square $n \times n$ matrix in time $O(n^c)$, then one can compute the product of two square $n \times n$ matrices in $\tilde{O}(n^c)$ time [Mat]. Thus, up to logarithmic factors, running times for matrix multiplication and matrix inversion must be the same in this model.

In the Real RAM model, in the general case, given a matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ and a vector $\mathbf{b} \in \mathbb{R}^n$, the fastest known algorithm for solving a system of linear equations $\mathbf{A}\mathbf{x} = \mathbf{b}$, i.e. finding $\mathbf{x}$ that satisfies the equation is based on simply computing $\mathbf{A}^{-1}$ and applying this matrix to $\mathbf{b}$ to get $\mathbf{x} = \mathbf{A}^{-1}\mathbf{b}$. Consequently, the best

known running time in this model for solving general linear equations is $\tilde{O}(n^\omega)$. However, it is not known whether a faster algorithm could exist, and this is a major open question in numerical linear algebra and scientific computing, and resolving it would have wide impact across theoretical computer science.

Linear Equations in fixed point arithmetic.

In fixed point arithmetic, when we operate on integer inputs, an $(\mathbf{A}, \mathbf{b}, \epsilon)$ instance of LEA can be solved in $\tilde{O}(n^\omega)$ arithmetic operations on numbers represented with $(\log(\|\mathbf{A}\| + \|\mathbf{b}\|)/\epsilon) \log n)^{O(1)}$ bits each, if the input matrix has $n^{O(1)}$ condition number [DDH07]. A crucial point is that the polynomially bounded condition number ensures that the output and intermediate calculations can be approximately represented using few bits.

Linear Equations in Real RAM with polynomially conditioned inputs. Our analysis is based on inputs with polynomially bounded condition numbers, but we assume the Real RAM model of computation for simplicity.

REMARK 3.1. *Because we study $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon)$ instances with polynomially bounded condition numbers, and because we rely on simple, numerically stable reductions it would be straightforward to extend our analysis to the fixed point arithmetic model assuming integer inputs and $(\|\mathbf{A}\| + \|\mathbf{b}\|)/\epsilon = \text{poly}(n)$. Then using $\log^{O(1)} n$ bits per number for our computations should suffice.*

The current, long-standing, state of the art for solving linear equations in both the Real RAM model and in fixed point arithmetic on inputs with polynomially bounded condition number motivates our Linear Equation Time Hypothesis, which formalizes the hypothesis that even moderately well-conditioned linear equations cannot be solved quickly. We state a parameterized family of hypotheses, and then discuss the significance of the hypotheses for different parameter values.

DEFINITION 3.1. (LINEAR EQUATION TIME HYPOTHESIS: LTH_k^γ) *In the Real RAM model, in the worst case, when $\mathbf{A}$ is a full rank $n \times n$ matrix and has condition number most $\kappa(\mathbf{A}) \leq n^k$, the problem $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon)$ with $\epsilon = n^{-10k}$ cannot be solved in $\tilde{O}(n^\gamma)$ time.*

Note that the particular exponent for the bound on ϵ is not important. Also note that in general, solving linear equations requires reading the whole input, so the running time must be at least $\Omega(n^2)$. Hence, LTH_k^γ is true for $\gamma < 2$, at least for $k = \Omega(1)$.

For sparse linear equations in matrices with bounded condition number, a different class of algorithms give a better running time. The best known

general result in this case is based on using Conjugate Gradient Descent (or Chebyshev Iterations, or an accelerated first order method such as Nesterov's algorithm) [Saa03], and these solve $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon)$ in time⁷ $\text{nnz}(\mathbf{A})\kappa(\mathbf{A})\log(1/\epsilon)$ in the Real RAM model, or if we adapt a different formulation of the LEA for positive definite matrices, the time would be $\text{nnz}(\mathbf{A})\sqrt{\kappa(\mathbf{A})}\log(1/\epsilon)$, (when appropriately phrased for $\mathbf{A}$ being square and positive definite).

Conjugate gradient can also be shown in the Real RAM model to converge to an exact solution in n iterations, giving a running time of order $n \cdot \text{nnz}(\mathbf{A})$. However, this is highly misleading, because achieving this type of behavior in a *floating* point arithmetic model requires about n bits per number, as opposed to $(\log(\|\mathbf{A}\| + \|\mathbf{b}\|/\epsilon) \log n)^{O(1)}$ bits in *floating* point arithmetic to achieve the condition number dependent behavior [MMS18].

For our formulation of LEA, no better algorithm than CG is known, and improving on the $\text{nnz}(\mathbf{A})\kappa(\mathbf{A})\log(1/\epsilon)$ running time is a major open problem. This motivates the following parameterized family of hardness assumptions.

DEFINITION 3.2. (SPARSE LINEAR EQUATION TIME HYPOTHESIS: SLTH_k^γ) *In the Real RAM model, in the worst case, when $\mathbf{A}$ is an $m \times n$ matrix with $m \geq n$, $\text{rank}(\mathbf{A}) = n$, and $\kappa(\mathbf{A}) \leq (\text{nnz}(\mathbf{A}))^k$, $\text{LEA}(\mathbf{A}, \gamma, \epsilon)$ with $\epsilon = n^{-10k}$ cannot be solved in time $\tilde{O}(\text{nnz}(\mathbf{A})^\gamma)$.*

Note that Conjugate Gradient shows that SLTH_k^{1+k} is false. However, there is essentially no evidence against $\text{SLTH}_k^{1+0.99k}$ and falsifying this hardness assumption would constitute major progress in solving sparse systems of linear equations. Note that if we instead formulated LEA for Positive Definite matrices and with a different notion of error, Conjugate Gradient would falsify $\text{SLTH}_k^{1+0.5k}$ and falsifying $\text{SLTH}_k^{1+0.49k}$ would be the right bar for making progress.

We state our main results for lower bounds of both dense and sparse packing LP instances in the following. The statements about running times all refer to the Real RAM model.

THEOREM 3.1. (DENSE REDUCTION) *If we can $(1-\epsilon)$ -approximately solve a packing LP of size $O(n) \times O(n)$ in time $\tilde{O}(n^\beta/\epsilon^\alpha)$, then given a condition number upper bound $K_A \geq \kappa(\mathbf{A})$, we can solve $\text{LEA}(\mathbf{A}, \mathbf{b}, 0.1)$ for an $n \times n$ matrix in time*

$$\tilde{O}(n^\beta (nK_A^3)^\alpha).$$

⁷Note that applying CG to $\mathbf{A}^\top \mathbf{A}$ gives a condition number dependence of the form $\sqrt{\kappa(\mathbf{A})^2} = \kappa(\mathbf{A})$.

THEOREM 3.2. (SPARSE REDUCTION) *If we can $(1-\epsilon)$ -approximately solve a packing LP that has N non-zeros in time $\tilde{O}(N/\epsilon^\alpha)$, for any constant α , then given a condition number upper bound $K_A \geq \kappa(\mathbf{A})$, we can solve $\text{LEA}(\mathbf{A}, \mathbf{b}, 0.1)$ for an $m \times n$ matrix $\mathbf{A}$ in time*

$$\tilde{O}(\text{nnz}(\mathbf{A}) (m^{3/2} K_A^3)^\alpha).$$

REMARK 3.2. *One could replace $\text{LEA}(\mathbf{A}, \mathbf{b}, 0.1)$ in Theorem 3.1 and 3.2 by $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon')$ for any $\epsilon' = 1/\text{poly}(n)$. According to Lemma 2.1, the run times of solving $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon')$ and $\text{LEA}(\mathbf{A}, \mathbf{b}, 0.1)$ only differ by $O(\log n)$, which will be hidden in $\tilde{O}(\cdot)$.*

By connecting these two theorems with our Linear Equation Time Hypotheses defined as above, we get the following corollaries.

COROLLARY 3.1. *If we can $(1-\epsilon)$ -approximately solve a packing LP of size $n \times n$ in time $O(n^\beta/\epsilon^\alpha)$, then for any constant $k > 0$, then $\text{LTH}_k^{\beta+\alpha+3\alpha k}$ is false.*

Proof. By Theorem 3.1, we get an algorithm for $\text{LEA}(\mathbf{A}, \mathbf{b}, 0.1)$ with $K_A = n^k$, which by Theorem 2.1, we can turn into an algorithm for $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon)$ with $\epsilon = n^{-10k}$ using $\log(1/\epsilon) = O(k \log n)$ iterations of the $\text{LEA}(\mathbf{A}, \mathbf{b}, 0.1)$ solver with different $\mathbf{b}$. Altogether, the running time is $\tilde{O}(k \log n \cdot n^\beta (nK_A^3)^\alpha)$, which for any constant k is

$$\tilde{O}(n^{\beta+\alpha+3k\alpha}).$$

This falsifies $\text{LTH}_k^{\beta+\alpha+3\alpha k}$. $\square$

Restricting our attention to matrices $\mathbf{A}$ with condition number of $\mathbf{A}^\top \mathbf{A}$ at most $O(n^{10})$, we can state our central corollary of this as:

THEOREM 3.3. [Informal Statement] *In the Real RAM model,*

- *unless linear equations in matrices with condition number $O(n^{10})$ can be solved to ϵ accuracy faster than $\tilde{O}(n^{2.01} \log(1/\epsilon))$, no algorithm $(1-\epsilon)$ -approximately solves packing LPs in time $\tilde{O}(N\epsilon^{-0.0003})$.*
- *unless linear equations in matrices with condition number $O(n^{1.5})$ can be solved to ϵ accuracy faster than $\tilde{O}(n^{2.372} \log(1/\epsilon))$, no algorithm $(1-\epsilon)$ -approximately solves packing LPs in time $\tilde{O}(N\epsilon^{-0.067})$.*

Proof. From our $O(N\epsilon^{-0.0003})$ time packing LP solver, by Theorem 3.1 and Lemma 2.1, noting that $N \leq n^2$ and $2 + 0.0003 \cdot (3 \cdot 10 + 1) \leq 2.01$ immediately

get our $\tilde{O}(n^{2.01} \log(1/\epsilon))$ solver for $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon)$ with $K_{\mathbf{A}} = O(n^{10})$. The second statement follows similarly by checking the appropriate parameters. $\square$

COROLLARY 3.2. *If we can $(1 - \epsilon)$ -approximately solve a packing LP that has N non-zeros in time $\tilde{O}(N/\epsilon^\alpha)$, then $\text{SLTH}_k^{1+1.5\alpha+3k\alpha}$ is false.*

Proof. The proof is similar to that of Corollary 3.1, but it uses the reduction from Theorem 3.2 instead, and that given our assumptions on $\mathbf{A}$,

$$\text{nnz}(\mathbf{A}) \left(m^{3/2} K_{\mathbf{A}}^3 \right)^\alpha \leq (\text{nnz}(\mathbf{A})^{1+1.5\alpha+3k\alpha})$$

$\square$

This has interesting consequences for sparse matrices, e.g. $\text{nnz}(\mathbf{A}) \approx n \log n$ with small condition number, e.g. at most $n^{1.5}$. Both Conjugate Gradient and matrix-inversion by matrix multiplication (even with $\omega = 2$) are consistent with $\text{SLTH}_{1.5}^{1.99}$, and no algorithm can solve Packing LPs with N non-zero entries in time $\tilde{O}(N/\epsilon^{-0.165})$ unless $\text{SLTH}_{1.5}^{1.99}$ is false. We can summarize this consequence as

THEOREM 3.4. *[Informal Statement] In the Real RAM model, unless linear equations $\mathbf{A}^\top \mathbf{A} \mathbf{x} = \mathbf{A}^\top \mathbf{b}$ can be solved to high accuracy faster than time $\tilde{O}(\text{nnz}(\mathbf{A}) \sqrt{\text{condition number of } \mathbf{A}^\top \mathbf{A}})$ when $\mathbf{A}^\top \mathbf{A}$ has condition number $\Theta(n^3)$, then no algorithm $(1 - \epsilon)$ -approximately solves packing LPs in time $\tilde{O}(N\epsilon^{-0.165})$.*

Proof. Similar to the proof of our first main theorem, we combine Lemma 2.1 and Corollary 3.2, and note that for $\alpha = 0.165$ and $k = 1.5$, we have $1 + 1.5\alpha + 3k\alpha \leq 1.99$, so a packing LP solver with the stated parameters would outperform Conjugate Gradient. $\square$

REMARK 3.3. *Suppose that there is likely to be very large gap between $\kappa(\mathbf{A})$ and a known upper bound $K_{\mathbf{A}}$. Then we might want to get reductions that depend on $\kappa(\mathbf{A})$ instead of $K_{\mathbf{A}}$. We can essentially do this, and reduce the dependence on $K_{\mathbf{A}}$ to logarithmic. For example, in the proof of Theorem 3.1, we constructed a solver for $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon)$ with running time $\tilde{O}(n^\beta (nK_{\mathbf{A}}^3)^\alpha \log(1/\epsilon))$. and we can turn this into a solver with running time $\tilde{O}(n^\beta (n\kappa(\mathbf{A})^3)^\alpha \log(K_{\mathbf{A}}/\epsilon))$. The same type of improvement can be obtained in the sparse case. We explain these modified reductions in Appendix C.*

4 Reduction

In this section we discuss how to reduce solving linear systems to solving packing LPs. We introduce two

reductions from linear equation problems to packing LP problems. Using the two reductions, we will prove Theorems 3.1 and 3.2 later in this section.

Normalization. Given a linear equation approximation problem instance $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon)$ as defined in Definition 2.2, where $\mathbf{A} \in \mathbb{R}^{m \times n}$, $\mathbf{b} \in \mathbb{R}^m$, $0 \leq \epsilon \leq 1$. We will assume that $m \geq n$, $\sigma_{\min}(\mathbf{A}) > 0$, and $\kappa(\mathbf{A}) \leq K_{\mathbf{A}}$ for some known $K_{\mathbf{A}}$. Furthermore, we will assume $\mathbf{A}^\top \mathbf{b} \neq 0$, otherwise $\mathbf{x} = 0$ is a solution to $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon)$, in which case $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon)$ is solved in time $O(mn)$.

We denote $\mathbf{p} \stackrel{\text{def}}{=} \mathbf{A}^\top \mathbf{b}$, and WLOG we assume $\mathbf{A}, \mathbf{b}$ are normalized so that

$$\|\mathbf{A}\|_{\max} \stackrel{\text{def}}{=} \max_{i,j} |\mathbf{A}_{i,j}| = 1, \quad \|\mathbf{p}\|_2 = 1.$$

This normalization can be done in $\text{nnz}(\mathbf{A})$ time, and given any ϵ -approximate solution to the normalized linear system, we can scale the solution to get an ϵ -approximate solution to the original linear system. Moreover, the condition number $\kappa(\mathbf{A})$ is not changed by the normalization. Furthermore, as $\|\mathbf{A}\|_{\max} = 1$, we must have $\sigma_{\max}(\mathbf{A}) \geq 1$, which gives the following claim.

CLAIM 4.1. $\sigma_{\min}^{-1}(\mathbf{A}) \leq \kappa(\mathbf{A})$.

We will present two reductions from linear systems to packing LPs. Both reductions use the same high-level idea, but one reduction always constructs a packing LP that is dense, while the other preserves the sparsity of the given linear system instance. Formally, given a linear system $\mathbf{A} \mathbf{x} = \mathbf{b}$, where $\mathbf{A} \in \mathbb{R}^{m \times n}$ with $m \geq n$ and $\mathbf{A}$ has $\text{nnz}(\mathbf{A})$ non-zeros and $\kappa(\mathbf{A}) \leq K_{\mathbf{A}}$ for some known $K_{\mathbf{A}} = \text{poly}(n)$; and a scalar parameter U which we will discuss shortly, we have a function $\text{PLP}_{\text{dense}}(\mathbf{A}, \mathbf{b}, U)$ that constructs a dense packing LP of size $O(m) \times O(m)$, and a function $\text{PLP}_{\text{sparse}}(\mathbf{A}, \mathbf{b}, U)$ that constructs a packing LP of size $O(m) \times O(m)$ and $O(\text{nnz}(\mathbf{A}))$ non-zeros.

In both the dense reduction and sparse reduction, we first rewrite solving the linear system as finding a feasible solution to a LP which we denote by $\text{LP}(\mathbf{A}, \mathbf{b})$ as follows. Recall $\mathbf{p} \stackrel{\text{def}}{=} \mathbf{A}^\top \mathbf{b}$, and solving $\mathbf{A} \mathbf{x} = \mathbf{b}$ can be rewritten as $\mathbf{A}^\top \mathbf{A} \mathbf{x} = \mathbf{p}$, which in turn can be written as finding $\mathbf{x} \in \mathbb{R}^n, \mathbf{z} \in \mathbb{R}^m$ such that

$$(4.5) \quad \begin{aligned} \mathbf{A}^\top \mathbf{z} &= \mathbf{p} \\ \mathbf{A} \mathbf{x} &= \mathbf{z} \end{aligned}$$

Note that even when $\mathbf{A} \mathbf{x} = \mathbf{b}$ does not have an exact solution, the linear system (4.5) always has an exact solution, and we can bound the norm of the exact solution.

CLAIM 4.2. *There exists $\mathbf{x}^*, \mathbf{z}^*$ satisfying the linear system (4.5), and*

$$\|\mathbf{z}^*\|_1 \leq \sqrt{m} \|\Pi_A \mathbf{b}\|_2, \quad \|\mathbf{x}^*\|_1 \leq \sqrt{n} \sigma_{\min}^{-1}(\mathbf{A}) \|\Pi_A \mathbf{b}\|_2,$$

moreover, we know

$$\|\Pi_A \mathbf{b}\|_2 \in [\sigma_{\max}^{-1}(\mathbf{A}) \|\mathbf{p}\|_2, \sigma_{\min}^{-1}(\mathbf{A}) \|\mathbf{p}\|_2].$$

Proof. Consider the solution where $\mathbf{z}^* = \Pi_A \mathbf{b} = \mathbf{A}(\mathbf{A}^T \mathbf{A})^{-1} \mathbf{p}$ and $\mathbf{x}^* = (\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T \mathbf{b}$. It is straightforward to check $\mathbf{z}^*, \mathbf{x}^*$ satisfy (4.5). Moreover,

$$\|\mathbf{z}^*\|_1 \leq \sqrt{m} \|\mathbf{z}^*\|_2 = \sqrt{m} \|\Pi_A \mathbf{b}\|_2,$$

and

$$\|\mathbf{x}^*\|_1 \leq \sqrt{n} \|\mathbf{x}^*\|_2 \leq \sqrt{n} \sigma_{\min}^{-1}(\mathbf{A}) \|\Pi_A \mathbf{b}\|_2$$

where the last inequality follows from $\mathbf{A} \mathbf{x}^* = \Pi_A \mathbf{b}$.

Consider the SVD of $\mathbf{A} = \sum_i \sigma_i \mathbf{u}_i \mathbf{v}_i^T$. Then $\|\Pi_A \mathbf{b}\|_2^2 = \sum_i (\mathbf{u}_i^T \mathbf{b})^2$ and $\|\mathbf{p}\|_2^2 = \|\mathbf{A}^T \mathbf{b}\|_2^2 = \sum_i \sigma_i^2 (\mathbf{u}_i^T \mathbf{b})^2$. Thus,

$$\sigma_{\max}^{-1}(\mathbf{A}) \leq \frac{\|\Pi_A \mathbf{b}\|_2}{\|\mathbf{p}\|_2} \leq \sigma_{\min}^{-1}(\mathbf{A}).$$

□

In both of our reductions $\text{PLP}_{\text{dense}}(\mathbf{A}, \mathbf{b}, U)$ and $\text{PLP}_{\text{sparse}}(\mathbf{A}, \mathbf{b}, U)$, the parameter U is supposed to be a tight upper-bound of $\|\mathbf{z}^*\|_1 + \|\mathbf{x}^*\|_1$. Claim 4.2 implies that it suffices to pick $U = 2K_A^2 \sqrt{m}$.

From the linear system (4.5), we construct a feasibility linear program $\text{LP}(\mathbf{A}, \mathbf{b})$ as follows. We want all variables to be non-negative, and we can do this in a fairly standard way by creating two non-negative variables per $\mathbf{x}_i$ and $\mathbf{z}_i$:

$$(4.6) \quad \mathbf{x}_j = \mathbf{x}_j^{(+)} - \mathbf{x}_j^{(-)}$$

$$(4.7) \quad \mathbf{z}_i = \mathbf{z}_i^{(+)} - \mathbf{z}_i^{(-)}$$

Let $\mathbf{x}^{(+)}, \mathbf{x}^{(-)}, \mathbf{z}^{(+)},$ and $\mathbf{z}^{(-)}$ be vectors whose entries are $\mathbf{x}_j^{(+)}, \mathbf{x}_j^{(-)}, \mathbf{z}_i^{(+)}, \mathbf{z}_i^{(-)}$ respectively. $\text{LP}(\mathbf{A}, \mathbf{b})$ is the feasibility LP of finding a solution satisfying

$$\begin{aligned} \mathbf{A}^T \mathbf{z}^{(+)} - \mathbf{A}^T \mathbf{z}^{(-)} &\leq \mathbf{p} \\ -\mathbf{A}^T \mathbf{z}^{(+)} + \mathbf{A}^T \mathbf{z}^{(-)} &\leq -\mathbf{p} \\ \mathbf{A} \mathbf{x}^{(+)} - \mathbf{A} \mathbf{x}^{(-)} - (\mathbf{z}^{(+)} - \mathbf{z}^{(-)}) &\leq 0 \\ -(\mathbf{A} \mathbf{x}^{(+)} - \mathbf{A} \mathbf{x}^{(-)}) + (\mathbf{z}^{(+)} - \mathbf{z}^{(-)}) &\leq 0 \\ \mathbf{x}^{(+)}, \mathbf{x}^{(-)}, \mathbf{z}^{(+)}, \mathbf{z}^{(-)} &\geq 0 \end{aligned}$$

The two reductions $\text{PLP}_{\text{dense}}(\mathbf{A}, \mathbf{b}, U)$ and $\text{PLP}_{\text{sparse}}(\mathbf{A}, \mathbf{b}, U)$ proceed differently to write $\text{LP}(\mathbf{A}, \mathbf{b})$ into packing LPs, and we present the remaining steps of the two reduction and the analysis in the rest of the section.

4.1 $\text{PLP}_{\text{dense}}(\mathbf{A}, \mathbf{b}, U)$ To turn $\text{LP}(\mathbf{A}, \mathbf{b})$ into a packing LP, we add one non-negative slack variable $s \geq 0$, and denote

$$\alpha_{\text{sum}} \stackrel{\text{def}}{=} s + \sum_{1 \leq j \leq n} (\mathbf{x}_j^{(+)} + \mathbf{x}_j^{(-)}) + \sum_{1 \leq i \leq m} (\mathbf{z}_i^{(+)} + \mathbf{z}_i^{(-)}).$$

Note α_{sum} is merely a short-hand for the sum of all the variables instead of a new variable. For each existing constraint of $\text{LP}(\mathbf{A}, \mathbf{b})$ except the positivity constraints on single variables, we add α_{sum} and U to the LHS and RHS of the constraint respectively, and we also add

$$(4.8) \quad \alpha_{\text{sum}} \leq U$$

as an additional constraint. This completes the construction of $\text{PLP}_{\text{dense}}(\mathbf{A}, \mathbf{b}, U)$'s constraints, and if written explicitly are Equation (4.9).

The objective of $\text{PLP}_{\text{dense}}(\mathbf{A}, \mathbf{b}, U)$ is to maximize α_{sum} , that is

$$\max s + \sum_{1 \leq j \leq n} (\mathbf{x}_j^{(+)} + \mathbf{x}_j^{(-)}) + \sum_{1 \leq i \leq m} (\mathbf{z}_i^{(+)} + \mathbf{z}_i^{(-)}).$$

LEMMA 4.1. *Given $U = 2K_A^2 \sqrt{m} > 1$ is an upper bound of the ℓ_1 norm of some solution of the linear system (4.5), $\text{PLP}_{\text{dense}}(\mathbf{A}, \mathbf{b}, U)$ is a packing LP of size $O(m) \times O(m)$, and $\text{PLP}_{\text{dense}}(\mathbf{A}, \mathbf{b}, U)$ has optimum at U .*

Proof. The size of $\text{PLP}_{\text{dense}}(\mathbf{A}, \mathbf{b}, U)$ is obvious. For it to be a packing LP, we need to show

1. All variables are non-negative: This holds by construction.
2. All coefficients are non-negative: This is true since $\|\mathbf{A}\|_{\max} \leq 1$, so by adding 1 to every coefficient, all coefficients become non-negative.
3. All constants on the RHS of the constraints are non-negative: This is true since $\|\mathbf{p}\|_2 = 1$, so $\|\mathbf{p}\|_{\infty} \leq 1$, and we add $U \geq 1$ to the RHS of each constraint.
4. All coefficients in the objective function are non-negative: This is true since all coefficients are 1 in the objective.

If $U \geq \|\mathbf{x}^*\|_1 + \|\mathbf{z}^*\|_1$ for some $\mathbf{x}^*$ and $\mathbf{z}^*$ satisfying the linear system (4.5), consider the following solution to $\text{PLP}_{\text{dense}}(\mathbf{A}, \mathbf{b}, U)$ where we set $\mathbf{x}^{(+)}$ and $\mathbf{x}^{(-)}$ based on the signs of $\mathbf{x}^*$

$$(\mathbf{x}_i^{(+)}, \mathbf{x}_i^{(-)}) := \begin{cases} (\mathbf{x}_i^*, 0) & \text{if } \mathbf{x}_i^* \geq 0, \\ (0, -\mathbf{x}_i^*) & \text{if } \mathbf{x}_i^* \leq 0, \end{cases}$$

$$\begin{aligned}
(4.9) \quad & s + \sum_{1 \leq i \leq m} (1 + \mathbf{A}_{ij}) \mathbf{z}_i^{(+)} + (1 - \mathbf{A}_{ij}) \mathbf{z}_i^{(-)} + \sum_{1 \leq j' \leq n} \mathbf{x}_{j'}^{(+)} + \mathbf{x}_{j'}^{(-)} \leq \mathbf{p}_j + U \quad \forall 1 \leq j \leq n \\
& s + \sum_{1 \leq i \leq m} (1 - \mathbf{A}_{ij}) \mathbf{z}_i^{(+)} + (1 + \mathbf{A}_{ij}) \mathbf{z}_i^{(-)} + \sum_{1 \leq j' \leq n} \mathbf{x}_{j'}^{(+)} + \mathbf{x}_{j'}^{(-)} \leq -\mathbf{p}_j + U \quad \forall 1 \leq j \leq n \\
& s + \sum_{1 \leq j \leq n} (1 + \mathbf{A}_{ij}) \mathbf{x}_j^{(+)} + (1 - \mathbf{A}_{ij}) \mathbf{x}_j^{(-)} + 2\mathbf{z}_i^{(-)} + \sum_{1 \leq i' \leq n, i' \neq i} \mathbf{z}_{i'}^{(+)} + \mathbf{z}_{i'}^{(-)} \leq U \quad \forall 1 \leq i \leq m \\
& s + \sum_{1 \leq j \leq n} (1 - \mathbf{A}_{ij}) \mathbf{x}_j^{(+)} + (1 + \mathbf{A}_{ij}) \mathbf{x}_j^{(-)} + 2\mathbf{z}_i^{(+)} + \sum_{1 \leq i' \leq n, i' \neq i} \mathbf{z}_{i'}^{(+)} + \mathbf{z}_{i'}^{(-)} \leq U \quad \forall 1 \leq i \leq m \\
& s + \sum_{1 \leq j \leq n} (\mathbf{x}_j^{(+)} + \mathbf{x}_j^{(-)}) + \sum_{1 \leq i \leq m} (\mathbf{z}_i^{(+)} + \mathbf{z}_i^{(-)}) \leq U \\
& s, \mathbf{x}^{(+)}, \mathbf{x}^{(-)}, \mathbf{z}^{(+)}, \mathbf{z}^{(-)} \geq 0
\end{aligned}$$

and similarly for $\mathbf{z}^{(+)}$ and $\mathbf{z}^{(-)}$:

$$(\mathbf{z}_i^{(+)}, \mathbf{z}_i^{(-)}) := \begin{cases} (\mathbf{z}_i^*, 0) & \text{if } \mathbf{z}_i^* \geq 0, \\ (0, -\mathbf{z}_i^*) & \text{if } \mathbf{z}_i^* \leq 0, \end{cases}$$

For the added slack variable s , as $U \geq \|\mathbf{x}^*\|_1 + \|\mathbf{z}^*\|_1$, we can set

$$\begin{aligned}
s &\stackrel{\text{def}}{=} U - \|\mathbf{x}^*\|_1 + \|\mathbf{z}^*\|_1 \\
&= U - \left(\sum_{1 \leq j \leq n} (\mathbf{x}_j^{(+)} + \mathbf{x}_j^{(-)}) + \sum_{1 \leq i \leq m} (\mathbf{z}_i^{(+)} + \mathbf{z}_i^{(-)}) \right).
\end{aligned}$$

It is straightforward to see this solution is feasible, and gives objective value U , which must be the optimal due to the constraint (4.8). $\square$

In the following lemma, we translate the error bound between the original linear system and $\text{PLP}_{\text{dense}}(\mathbf{A}, \mathbf{b}, U)$.

LEMMA 4.2. *Given $\mathbf{A} \in \mathbb{R}^{m \times n}$, $\mathbf{b} \in \mathbb{R}^m$, $0 \leq \epsilon' \leq 1$, where $m \geq n$, and $\mathbf{A}^\top \mathbf{b} \neq 0$, and $\kappa(\mathbf{A})$ is upper bounded by some known number K_A . Let*

$$U \stackrel{\text{def}}{=} 2K_A^2 \sqrt{m}, \quad \epsilon \stackrel{\text{def}}{=} \frac{\epsilon'}{2K_A \sqrt{m} U}$$

Suppose $\mathbf{x}^{(+)}, \mathbf{x}^{(-)}, \mathbf{z}^{(+)}, \mathbf{z}^{(-)}, s$ is a feasible solution of $\text{PLP}_{\text{dense}}(\mathbf{A}, \mathbf{b}, U)$ with objective value at least $(1 - \epsilon)U$, then $\mathbf{x} = \mathbf{x}^{(+)} - \mathbf{x}^{(-)}$ is a solution to $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon')$, that is,

$$\|\mathbf{Ax} - \Pi_A \mathbf{b}\|_2 \leq \epsilon' \|\Pi_A \mathbf{b}\|_2.$$

Proof. By Claim 4.2, U is an upper bound of the ℓ_1 norm of some solution of the linear system (4.5), that is, U satisfies the condition of Lemma 4.1. By Claim 4.2,

$$(4.10) \quad \|\Pi_A \mathbf{b}\|_2 \geq \frac{1}{\sigma_{\max}(\mathbf{A})}$$

Let $\mathbf{x} = \mathbf{x}^{(+)} - \mathbf{x}^{(-)}$, $\mathbf{z} = \mathbf{z}^{(+)} - \mathbf{z}^{(-)}$. We want to bound the error in $\mathbf{Ax} = \mathbf{z}$, $\mathbf{A}^\top \mathbf{z} = \mathbf{p}$. Consider the i -th equality constraint in the linear system: $\mathbf{A}_i \mathbf{x} = \mathbf{z}_i$, we have a corresponding pair of constraints in $\text{PLP}_{\text{dense}}(\mathbf{A}, \mathbf{b}, U)$

$$\begin{aligned}
\alpha_{\text{sum}} - (\mathbf{z}_i^{(+)} - \mathbf{z}_i^{(-)}) + \sum_{1 \leq j \leq n} \mathbf{A}_{ij} (\mathbf{x}_j^{(+)} - \mathbf{x}_j^{(-)}) &\leq U \\
\alpha_{\text{sum}} + (\mathbf{z}_i^{(+)} - \mathbf{z}_i^{(-)}) - \sum_{1 \leq j \leq n} \mathbf{A}_{ij} (\mathbf{x}_j^{(+)} - \mathbf{x}_j^{(-)}) &\leq U
\end{aligned}$$

Since we have a feasible solution, we know from the above two constraints that

$$|\mathbf{A}_i \mathbf{x} - \mathbf{z}_i| \leq U - \alpha_{\text{sum}}.$$

The same argument holds for any linear constraint $(\mathbf{A}^\top)_j \mathbf{z} = \mathbf{p}_j$, so we have

$$\|\mathbf{Ax} - \mathbf{z}\|_\infty, \|\mathbf{A}^\top \mathbf{z} - \mathbf{p}\|_\infty \leq U - \alpha_{\text{sum}},$$

and by the approximation guarantee the RHS is at most $\epsilon \cdot U$.

Now we bound the error of the solution $\mathbf{x}$ in terms of the linear system

$$\begin{aligned}
(4.11) \quad & \|\mathbf{Ax} - \Pi_A \mathbf{b}\|_2 \\
& \leq \|\mathbf{Ax} - \Pi_A \mathbf{z}\|_2 + \|\Pi_A \mathbf{z} - \Pi_A \mathbf{b}\|_2 \\
& \leq \|\mathbf{Ax} - \mathbf{z}\|_2 + \sigma_{\min}^{-1}(A) \|\mathbf{A}^\top \mathbf{z} - \mathbf{A}^\top \mathbf{b}\|_2 \\
& \leq \sqrt{m} \|\mathbf{Ax} - \mathbf{z}\|_\infty + \sqrt{n} \sigma_{\min}^{-1}(A) \|\mathbf{A}^\top \mathbf{z} - \mathbf{A}^\top \mathbf{b}\|_\infty \\
& \leq \epsilon \kappa(\mathbf{A}) (\sqrt{m} + \sqrt{n}) U.
\end{aligned}$$

In the second inequality, we have $\|\mathbf{Ax} - \Pi_A \mathbf{z}\|_2 \leq \|\mathbf{Ax} - \mathbf{z}\|_2$ because $\mathbf{Ax}$ is in the span of columns of

$\mathbf{A}$, so projecting $\mathbf{z}$ to the column span of $\mathbf{A}$ makes the distance to $\mathbf{Ax}$ smaller. Also in the second inequality, we use

$$\begin{aligned} & \|\Pi_A \mathbf{z} - \Pi_A \mathbf{b}\|_2 \\ &= \left\| \mathbf{A}(\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T (\mathbf{z} - \mathbf{b}) \right\| \\ &= \sqrt{(\mathbf{z} - \mathbf{b})^T \mathbf{A}(\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T \mathbf{A}(\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T (\mathbf{z} - \mathbf{b})} \\ &= \sqrt{(\mathbf{z} - \mathbf{b})^T \mathbf{A}(\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T (\mathbf{z} - \mathbf{b})} \\ &\leq \sigma_{\min}^{-1}(\mathbf{A}) \left\| \mathbf{A}^T \mathbf{z} - \mathbf{A}^T \mathbf{b} \right\|_2. \end{aligned}$$

The last inequality is due to Claim 4.1. Plugging Equation (4.10) into Equation (4.11):

$$\begin{aligned} \|\mathbf{Ax} - \Pi_A \mathbf{b}\|_2 &\leq \epsilon \kappa(\mathbf{A}) (\sqrt{m} + \sqrt{n}) U \cdot \|\Pi_A \mathbf{b}\|_2 \\ &\leq 2\epsilon K_A \sqrt{m} U \cdot \|\Pi_A \mathbf{b}\|_2. \end{aligned}$$

By the setting of ϵ , we have $\|\mathbf{Ax} - \Pi_A \mathbf{b}\|_2 \leq \epsilon' \|\Pi_A \mathbf{b}\|_2$. $\square$

Proof. [Proof of Theorem 3.1] Given an arbitrary linear equation approximation problem instance $\text{LEA}(\mathbf{A}, \mathbf{b}, 0.1)$ where $\mathbf{A} \in \mathbb{R}^{n \times n}$, $\mathbf{b} \in \mathbb{R}^n$ and $\kappa(\mathbf{A}) \leq K_A$ for some known K_A . We first check whether $\mathbf{A}^\top \mathbf{b} = 0$. If $\mathbf{A}^\top \mathbf{b} = 0$, then $\mathbf{x} = 0$ is a solution to $\text{LEA}(\mathbf{A}, \mathbf{b}, 0.1)$. Otherwise, we will construct $\text{PLP}_{\text{dense}}(\mathbf{A}, \mathbf{b}, U)$ with $U = K_A^2 \sqrt{n}$ and solve it up to multiplicative error $\epsilon = \frac{1}{100K_A \sqrt{n} U}$. By Lemma 4.2, we will get a solution to $\text{LEA}(\mathbf{A}, \mathbf{b}, 0.1)$.

Suppose we can solve the packing LP $\text{PLP}_{\text{dense}}(\mathbf{A}, \mathbf{b}, U)$, whose size is $O(n) \times O(n)$, up to multiplicative error ϵ in time $\tilde{O}(n^\beta / \epsilon^\alpha)$, then we can solve the $\text{LEA}(\mathbf{A}, \mathbf{b}, 0.1)$ in time $\tilde{O}(n^\beta (K_A^3 n)^\alpha)$. $\square$

4.2 $\text{PLP}_{\text{sparse}}(\mathbf{A}, \mathbf{b}, U)$ The sparsity preserving reduction follows the same approach as the dense reduction. However, instead of adding the same α_{sum} and U to the LHS and RHS of every inequality of $\text{LP}(\mathbf{A}, \mathbf{b})$ to make the coefficients non-negative, we will pin-point only the non-zero coefficients to preserve sparsity. We construct $\text{PLP}_{\text{sparse}}(\mathbf{A}, \mathbf{b}, U)$ starting from $\text{LP}(\mathbf{A}, \mathbf{b})$ as follows.

For each $j \in [1, n]$, look at the pair of constraints in $\text{LP}(\mathbf{A}, \mathbf{b})$ corresponding to the j -th row of $\mathbf{A}^T \mathbf{z} = \mathbf{p}$ in the linear system (4.5):

$$\begin{aligned} (\mathbf{A}^T)_j \mathbf{z}^{(+)} - (\mathbf{A}^T)_j \mathbf{z}^{(-)} &\leq \mathbf{p}_j \\ -(\mathbf{A}^T)_j \mathbf{z}^{(+)} + (\mathbf{A}^T)_j \mathbf{z}^{(-)} &\leq -\mathbf{p}_j \end{aligned}$$

We add a non-negative slack variable $\mathbf{s}_j^{(p)}$ that serves as the slack for *both* of these constraints. Furthermore, we

denote

$$\alpha_j^{(p)} \stackrel{\text{def}}{=} \mathbf{s}_j^{(p)} + \sum_{i: A_{ij} \neq 0} \mathbf{z}_i^{(+)} + \mathbf{z}_i^{(-)},$$

and again note $\alpha_j^{(p)}$ is merely a shorthand rather than a new variable. We add $\alpha_j^{(p)}$ and U to the LHS and RHS respectively of both the constraints above. Since $\|\mathbf{A}\|_{\max} \leq 1$, all the coefficients in these constraints become non-negative. We then add to $\text{PLP}_{\text{sparse}}(\mathbf{A}, \mathbf{b}, U)$ a new inequality

$$\alpha_j^{(p)} \leq U.$$

Similarly, for each $i \in [1, m]$, we consider the pair of constraints corresponding to the i -th row in $\mathbf{Ax} = \mathbf{z}$

$$\begin{aligned} \mathbf{A}_i \mathbf{x}^{(+)} - \mathbf{A}_i \mathbf{x}^{(-)} - (\mathbf{z}_i^{(+)} - \mathbf{z}_i^{(-)}) &\leq 0 \\ -(\mathbf{A}_i \mathbf{x}^{(+)} - \mathbf{A}_i \mathbf{x}^{(-)}) + (\mathbf{z}_i^{(+)} - \mathbf{z}_i^{(-)}) &\leq 0 \end{aligned}$$

We add a non-negative slack variable $\mathbf{s}_i^{(0)}$ for both of these constraints, and denote

$$\alpha_i^{(0)} := \mathbf{s}_i^{(0)} + \mathbf{z}_i^{(+)} + \mathbf{z}_i^{(-)} + \sum_{1 \leq j \leq n: A_{ij} \neq 0} \mathbf{x}_j^{(+)} + \mathbf{x}_j^{(-)}.$$

We add $\alpha_i^{(0)}$ and U to the LHS and RHS of the pair of constraints, and add $\alpha_i^{(0)} \leq U$ as a new constraint.

We can write the constraints of $\text{PLP}_{\text{sparse}}(\mathbf{A}, \mathbf{b}, U)$ explicitly as Equation (4.12).

Finally, the objective of $\text{PLP}_{\text{sparse}}(\mathbf{A}, \mathbf{b}, U)$ is to maximize $\sum_{1 \leq j \leq n} \alpha_j^{(p)} + \sum_{1 \leq i \leq m} \alpha_i^{(0)}$.

LEMMA 4.3. *Given $U = 2K_A^2 \sqrt{m} > 1$ is an upper bound of the ℓ_1 norm of some solution of the linear system (4.5), $\text{PLP}_{\text{sparse}}(\mathbf{A}, \mathbf{b}, U)$ constructs a packing LP with $O(\text{nnz}(\mathbf{A}))$ non-zeros and $\text{PLP}_{\text{sparse}}(\mathbf{A}, \mathbf{b}, U)$ has optimal value $(m+n)U$.*

The proof is a straightforward but tedious adaptation of the proof of Lemma 4.1, so we omit it.

Now we translate the error bound between the original linear system and $\text{PLP}_{\text{sparse}}(\mathbf{A}, \mathbf{b}, U)$.

LEMMA 4.4. *Given $\mathbf{A} \in \mathbb{R}^{m \times n}$, $\mathbf{b} \in \mathbb{R}^m$, $0 \leq \epsilon' \leq 1$, where $m \geq n$, $\mathbf{A}^\top \mathbf{b} \neq 0$ and $\kappa(\mathbf{A}) \leq K_A$ for some known number K_A . Let*

$$U \stackrel{\text{def}}{=} 2K_A^2 \sqrt{m}, \quad \epsilon \stackrel{\text{def}}{=} \frac{\epsilon'}{2K_A m U}.$$

Suppose $\mathbf{x}^{(+)}, \mathbf{x}^{(-)}, \mathbf{z}^{(+)}, \mathbf{z}^{(-)}, \mathbf{s}^{(p)}, \mathbf{s}^{(0)}$ is a feasible solution to $\text{PLP}_{\text{sparse}}(\mathbf{A}, \mathbf{b}, U)$ with objective value at least $(1 - \epsilon) \cdot (m+n)U$. Then, $\mathbf{x} = \mathbf{x}^{(+)} - \mathbf{x}^{(-)}$ is a solution to $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon')$, that is,

$$\|\mathbf{Ax} - \Pi_A \mathbf{b}\|_2 \leq \epsilon' \|\Pi_A \mathbf{b}\|_2.$$

$$\begin{aligned}
(4.12) \quad & s_j^{(p)} + \sum_{1 \leq i \leq m: \mathbf{A}_{ij} \neq 0} (1 + \mathbf{A}_{ij}) \mathbf{z}_i^{(+)} + (1 - \mathbf{A}_{ij}) \mathbf{z}_i^{(-)} \leq \mathbf{p}_j + U \quad \forall 1 \leq j \leq n \\
& s_j^{(p)} + \sum_{1 \leq i \leq m: \mathbf{A}_{ij} \neq 0} (1 - \mathbf{A}_{ij}) \mathbf{z}_i^{(+)} + (1 + \mathbf{A}_{ij}) \mathbf{z}_i^{(-)} \leq -\mathbf{p}_j + U \quad \forall 1 \leq j \leq n \\
& s_j^{(p)} + \sum_{1 \leq i \leq m: \mathbf{A}_{ij} \neq 0} \mathbf{z}_i^{(+)} + \mathbf{z}_i^{(-)} \leq U \quad \forall 1 \leq j \leq n \\
& 2\mathbf{z}_i^{(-)} + s_i^{(0)} + \sum_{1 \leq j \leq n: \mathbf{A}_{ij} \neq 0} (1 + \mathbf{A}_{ij}) \mathbf{x}_j^{(+)} + (1 - \mathbf{A}_{ij}) \mathbf{x}_j^{(-)} \leq U \quad \forall 1 \leq i \leq m \\
& 2\mathbf{z}_i^{(+)} + s_i^{(0)} + \sum_{1 \leq j \leq n: \mathbf{A}_{ij} \neq 0} (1 - \mathbf{A}_{ij}) \mathbf{x}_j^{(+)} + (1 + \mathbf{A}_{ij}) \mathbf{x}_j^{(-)} \leq U \quad \forall 1 \leq i \leq m \\
& \mathbf{z}_i^{(+)} + \mathbf{z}_i^{(-)} + s_i^{(0)} + \sum_{1 \leq j \leq n: \mathbf{A}_{ij} \neq 0} \mathbf{x}_j^{(+)} + \mathbf{x}_j^{(-)} \leq U \quad \forall 1 \leq i \leq m \\
& \mathbf{x}^{(+)}, \mathbf{x}^{(-)}, \mathbf{z}^{(+)}, \mathbf{z}^{(-)}, \mathbf{s}^{(p)}, \mathbf{s}^{(0)} \geq 0
\end{aligned}$$

Proof. Note U is an upper bound of the ℓ_1 norm of some solution of the linear system (4.5), and thus U satisfies the condition of Lemma 4.3. Besides, $\|\Pi_A \mathbf{b}\|_2 \geq \frac{1}{\sigma_{\max}(\mathbf{A})}$.

Given the $(1 - \epsilon)$ -approx optimal solution to $\text{PLP}_{\text{sparse}}(\mathbf{A}, \mathbf{b}, U)$, we consider the following solution to the linear system

$$\begin{aligned}
\mathbf{x} &:= \mathbf{x}^{(+)} - \mathbf{x}^{(-)} \\
\mathbf{z} &:= \mathbf{z}^{(+)} - \mathbf{z}^{(-)}
\end{aligned}$$

We want to bound the error in $\mathbf{Ax} = \mathbf{z}$, $\mathbf{A}^T \mathbf{z} = \mathbf{p}$.

Consider the i -th equality constraint in the linear system $\mathbf{A}_i \mathbf{x} = \mathbf{z}_i$, we have a corresponding pair of constraints in $\text{PLP}_{\text{sparse}}(\mathbf{A}, \mathbf{b}, U)$:

$$\begin{aligned}
\alpha_i^{(0)} + \mathbf{A}_i \mathbf{x}^{(+)} - \mathbf{A}_i \mathbf{x}^{(-)} - (\mathbf{z}_i^{(+)} - \mathbf{z}_i^{(-)}) &\leq U \\
\alpha_i^{(0)} - (\mathbf{A}_i \mathbf{x}^{(+)} - \mathbf{A}_i \mathbf{x}^{(-)}) + (\mathbf{z}_i^{(+)} - \mathbf{z}_i^{(-)}) &\leq U
\end{aligned}$$

which subtracting away $\alpha_i^{(0)}$ and U from the LHS and RHS respectively gives

$$|(\mathbf{Ax})_i - \mathbf{z}_i| \leq U - \alpha_i^{(0)}.$$

The same argument holds for any pair of linear constraints corresponding to $(\mathbf{A}^T \mathbf{z})_j = \mathbf{p}_j$, and in total we have

$$\begin{aligned}
& \|\mathbf{Ax} - \mathbf{z}\|_1 + \|\mathbf{A}^T \mathbf{z} - \mathbf{p}\|_1 \\
& \leq (m + n)U - \sum_j \alpha_j^{(p)} - \sum_i \alpha_i^{(0)} \\
& \leq \epsilon(m + n)U.
\end{aligned}$$

Now we bound the error of the solution $\mathbf{x}$ in terms of LEA

$$\begin{aligned}
\|\mathbf{Ax} - \Pi_A \mathbf{b}\|_2 &\leq \|\mathbf{Ax} - \Pi_A \mathbf{z}\|_2 + \|\Pi_A \mathbf{z} - \Pi_A \mathbf{b}\|_2 \\
&\leq \|\mathbf{Ax} - \mathbf{z}\|_2 + \sigma_{\min}^{-1}(\mathbf{A}) \|\mathbf{A}^T \mathbf{z} - \mathbf{A}^T \mathbf{b}\|_2 \\
&\leq \|\mathbf{Ax} - \mathbf{z}\|_1 + \sigma_{\min}^{-1}(\mathbf{A}) \|\mathbf{A}^T \mathbf{z} - \mathbf{A}^T \mathbf{b}\|_1 \\
&\leq \epsilon \sigma_{\min}^{-1}(\mathbf{A}) (m + n)U \\
&\leq \epsilon \sigma_{\min}^{-1}(\mathbf{A}) (m + n)U \cdot \|\Pi_A \mathbf{b}\|_2 \\
&\leq 2\epsilon K_A m U \|\Pi_A \mathbf{b}\|_2.
\end{aligned}$$

Setting

$$\epsilon = \frac{\epsilon'}{2K_A m U},$$

we have $\|\mathbf{Ax} - \Pi_A \mathbf{b}\|_2 \leq \epsilon' \|\Pi_A \mathbf{b}\|_2$. $\square$

Proof. [Proof of Theorem 3.2] The proof is similar to the proof of Theorem 3.1. If $\mathbf{A}^\top \mathbf{b} = 0$ then $\mathbf{x} = 0$ is a solution to $\text{LEA}(\mathbf{A}, \mathbf{b}, 0.1)$. Otherwise we solve $\text{PLP}_{\text{sparse}}(\mathbf{A}, \mathbf{b}, U)$. If we can ϵ -approximately solve a packing LP in time $\tilde{O}(N\epsilon^{-\alpha})$, then we can solve $\text{LEA}(\mathbf{A}, \mathbf{b}, 0.1)$ in time $\tilde{O}(nnz(\mathbf{A})(m^{3/2}K_A^3)^\alpha)$. $\square$

Acknowledgements

The authors thank Richard Peng for many helpful discussions related to this project, and we thank anonymous reviewers for their helpful comments.

References

- [AHK12] Sanjeev Arora, Elad Hazan, and Satyen Kale. The multiplicative weights update method: a meta-algorithm and applications. *Theory of Computing*, 8(1):121–164, 2012.

- [ARW17] A. Abboud, A. Rubinfeld, and R. Williams. Distributed pcp theorems for hardness of approximation in p. In *2017 IEEE 58th Annual Symposium on Foundations of Computer Science (FOCS)*, pages 25–36, Oct 2017.
- [AZO15a] Zeyuan Allen-Zhu and Lorenzo Orecchia. Nearly-linear time positive LP solver with faster convergence rate. In *Proceedings of the Forty-Seventh Annual ACM on Symposium on Theory of Computing*, STOC '15, pages 229–236, 2015.
- [AZO15b] Zeyuan Allen-Zhu and Lorenzo Orecchia. Using optimization to break the epsilon barrier: A faster and simpler width-independent algorithm for solving positive linear programs in parallel. In *Proceedings of the Twenty-Sixth Annual ACM-SIAM Symposium on Discrete Algorithms*, SODA '15, pages 1439–1456, 2015.
- [CLS19] Michael B. Cohen, Yin Tat Lee, and Zhao Song. Solving linear programs in the current matrix multiplication time. In *Proceedings of the 51st Annual ACM SIGACT Symposium on Theory of Computing*, STOC 2019, pages 938–942, New York, NY, USA, 2019. ACM.
- [CQ18] Chandra Chekuri and Kent Quanrud. Randomized mwu for positive lps. In *Proceedings of the Twenty-Ninth Annual ACM-SIAM Symposium on Discrete Algorithms*, pages 358–377. SIAM, 2018.
- [DDH07] James Demmel, Ioana Dumitriu, and Olga Holtz. Fast linear algebra is stable. *Numer. Math.*, 108(1):59–91, October 2007.
- [DLR79] David Dobkin, Richard J Lipton, and Steven Reiss. Linear programming is log-space hard for p. *Information Processing Letters*, 8(2):96–97, 1979.
- [DR80] David P. Dobkin and Steven P. Reiss. The complexity of linear programming. *Theor. Comput. Sci.*, 11(1):1–18, 1980.
- [GK07] Naveen Garg and Jochen Koenemann. Faster and simpler algorithms for multicommodity flow and other fractional packing problems. *SIAM Journal on Computing*, 37(2):630–652, 2007.
- [HS52] Magnus Rudolph Hestenes and Eduard Stiefel. *Methods of conjugate gradients for solving linear systems*, volume 49. NBS Washington, DC, 1952.
- [Ita78] Alon Itai. Two-commodity flow. *Journal of the ACM (JACM)*, 25(4):596–611, 1978.
- [KY14] Christos Koufogiannakis and Neal E Young. A nearly linear-time ptas for explicit fractional packing and covering linear programs. *Algorithmica*, 70(4):648–674, 2014.
- [KZ17] Rasmus Kyng and Peng Zhang. Hardness results for structured linear systems. In *Foundations of Computer Science (FOCS), 2017 IEEE 58th Annual Symposium on*, pages 684–695. IEEE, 2017.
- [LG14] François Le Gall. Powers of tensors and fast matrix multiplication. In *Proceedings of the 39th international symposium on symbolic and algebraic computation*, pages 296–303. ACM, 2014.
- [LN93] Michael Luby and Noam Nisan. A parallel approximation algorithm for positive linear programming. In *Proceedings of the Twenty-Fifth Annual ACM Symposium on Theory of Computing, May 16-18, 1993, San Diego, CA, USA*, pages 448–457, 1993.
- [LS14] Yin Tat Lee and Aaron Sidford. Path finding methods for linear programming: Solving linear programs in $\tilde{O}(n^3)$ iterations and faster algorithms for maximum flow. In *Foundations of Computer Science (FOCS), 2014 IEEE 55th Annual Symposium on*, pages 424–433. IEEE, 2014.
- [Mad10] Aleksander Madry. Faster approximation schemes for fractional multicommodity flow problems via dynamic graph algorithms. In *Proceedings of the forty-second ACM symposium on Theory of computing*, pages 121–130. ACM, 2010.
- [Mat] Notes for math 242, linear algebra, lehigh university fall 2008. Available at: <https://www.lehigh.edu/~gi02/m242/08linstras.pdf>.
- [Meg92] Nimrod Megiddo. A note on approximate linear programming. *Inf. Process. Lett.*, 42(1):53–, April 1992.
- [MMS18] Cameron Musco, Christopher Musco, and Aaron Sidford. Stability of the lanczos method for matrix function approximation. In *Proceedings of the Twenty-Ninth Annual ACM-SIAM Symposium on Discrete Algorithms*, pages 1605–1624. Society for Industrial and Applied Mathematics, 2018.
- [MNS⁺17] Cameron Musco, Praneeth Netrapalli, Aaron Sidford, Shashanka Ubaru, and David P Woodruff. Spectrum approximation beyond fast matrix multiplication: Algorithms and hardness. *arXiv preprint arXiv:1704.04163*, 2017.
- [MRWZ16] Michael W. Mahoney, Satish Rao, Di Wang, and Peng Zhang. Approximating the solution to mixed packing and covering lps in parallel $\tilde{O}(\epsilon^{-3})$ time. In *43rd International Colloquium on Automata, Languages, and Programming, ICALP 2016, July 11-15, 2016, Rome, Italy*, pages 52:1–52:14, 2016.
- [PS85] Franco P. Preparata and Michael I. Shamos. *Computational Geometry: An Introduction*. Springer-Verlag, Berlin, Heidelberg, 1985.
- [PST95] Serge A Plotkin, David B Shmoys, and Éva Tardos. Fast approximation algorithms for fractional packing and covering problems. *Mathematics of Operations Research*, 20(2):257–301, 1995.
- [Raz02] Ran Raz. On the complexity of matrix product. In *Proceedings of the thirty-fourth annual ACM symposium on Theory of computing*, pages 144–151. ACM, 2002.
- [Ren95] James Renegar. Incorporating condition measures into the complexity theory of linear programming. *SIAM Journal on Optimization*, 5(3):506–524, 1995.
- [Saa03] Yousef Saad. *Iterative methods for sparse linear systems*, volume 82. siam, 2003.
- [Ser91] Maria Serna. Approximating linear programming is log-space complete for p. *Information Processing Letters*, 37(4):233–236, 1991.
- [ST14] D. Spielman and S. Teng. Nearly linear time algorithms for preconditioning and solving symmetric, diagonally dominant linear systems. *SIAM Journal on*

Matrix Analysis and Applications, 35(3):835–885, 2014. Available at <http://arxiv.org/abs/cs/0607105>.

- [Str69] Volker Strassen. Gaussian elimination is not optimal. *Numerische mathematik*, 13(4):354–356, 1969.
- [TB97] Lloyd N. Trefethen and David Bau. *Numerical linear algebra*. SIAM, 1997.
- [TX98] Luca Trevisan and Fatos Xhafa. The parallel complexity of positive linear programming. *Parallel Processing Letters*, 8(04):527–533, 1998.
- [Wan17] Di Wang. *Fast Approximation Algorithms for Positive Linear Programs*. PhD thesis, UC Berkeley, 2017.
- [WRM16] Di Wang, Satish Rao, and Michael W. Mahoney. Unified acceleration method for packing and covering problems via diameter reduction. In *43rd International Colloquium on Automata, Languages, and Programming, ICALP 2016, July 11-15, 2016, Rome, Italy*, pages 50:1–50:13, 2016.
- [WW10] Virginia Vassilevska Williams and Ryan Williams. Subcubic equivalences between path, matrix and triangle problems. In *Foundations of Computer Science (FOCS), 2010 51st Annual IEEE Symposium on*, pages 645–654. IEEE, 2010.
- [You01] Neal E Young. Sequential and parallel algorithms for mixed packing and covering. In *Foundations of Computer Science, 2001. Proceedings. 42nd IEEE Symposium on*, pages 538–546. IEEE, 2001.
- [You14] Neal E Young. Nearly linear-time approximation schemes for mixed packing/covering and facility-location linear programs. *ArXiv e-prints*, 2014.

A Background on Approximately Solving Linear Equations

In this Appendix, for completeness, we give some background on notions of approximately solving systems of linear equations for readers unfamiliar with the subject. Our presentation is mostly reproduced from [KZ17] by a subset of the authors.

Recall our definition

[Linear Equation Approximation Problem, LEA] Given linear system $(\mathbf{A}, \mathbf{b})$, where $\mathbf{A} \in \mathbb{R}^{m \times n}$, and $\mathbf{b} \in \mathbb{R}^m$, and given a scalar $0 \leq \epsilon \leq 1$, we refer to the LEA problem for the triple $(\mathbf{A}, \mathbf{b}, \epsilon)$ as the problem of finding $\mathbf{x} \in \mathbb{R}^n$ s.t.

$$\|\mathbf{Ax} - \Pi_{\mathbf{A}}\mathbf{b}\|_2^2 \leq \epsilon \|\Pi_{\mathbf{A}}\mathbf{b}\|_2^2,$$

and we say that such an $\mathbf{x}$ is a solution to the LEA instance $(\mathbf{A}, \mathbf{b}, \epsilon)$.

This definition of a LEA instance and solution has several advantages: when $\text{im}(\mathbf{A}) = \mathbb{R}^m$, we get $\Pi_{\mathbf{A}} = \mathbf{I}$, and it reduces to the natural condition $\|\mathbf{Ax} - \mathbf{b}\|_2^2 \leq \epsilon \|\mathbf{b}\|_2^2$, which because $\text{im}(\mathbf{A}) = \mathbb{R}^m$, can be satisfied for any ϵ , and for $\epsilon = 0$ tells us that $\mathbf{Ax} = \mathbf{b}$.

When $\text{im}(\mathbf{A})$ does not include all of $\mathbb{R}^m$, the vector $\Pi_{\mathbf{A}}\mathbf{b}$ is exactly the projection of $\mathbf{b}$ onto $\text{im}(\mathbf{A})$, and so

a solution can still be obtained for any ϵ . Further, as $(\mathbf{I} - \Pi_{\mathbf{A}})\mathbf{b}$ is orthogonal to $\Pi_{\mathbf{A}}\mathbf{b}$ and $\mathbf{Ax}$, it follows that

$$\|\mathbf{Ax} - \mathbf{b}\|_2^2 = \|(\mathbf{I} - \Pi_{\mathbf{A}})\mathbf{b}\|_2^2 + \|\mathbf{Ax} - \Pi_{\mathbf{A}}\mathbf{b}\|_2^2.$$

Thus, when $\mathbf{x}$ is a solution to the LEA instance $(\mathbf{A}, \mathbf{b}, \epsilon)$, then $\mathbf{x}$ also gives an $\epsilon^2 \|\Pi_{\mathbf{A}}\mathbf{b}\|_2^2$ additive approximation to

$$(A.1) \quad \min_{\mathbf{x} \in \mathbb{R}^n} \|\mathbf{Ax} - \mathbf{b}\|_2^2 = \|(\mathbf{I} - \Pi_{\mathbf{A}})\mathbf{b}\|_2^2.$$

Similarly, an $\mathbf{x}$ which gives an additive $\epsilon^2 \|\Pi_{\mathbf{A}}\mathbf{b}\|_2^2$ approximation to Problem (A.1) is always a solution to the LEA instance $(\mathbf{A}, \mathbf{b}, \epsilon)$. These observations prove the following (well-known) fact:

FACT A.1. *Let $\mathbf{x}^* \in \arg \min_{\mathbf{x} \in \mathbb{R}^n} \|\mathbf{Ax} - \mathbf{b}\|_2^2$, then for every $\mathbf{x}$,*

$$\|\mathbf{Ax} - \mathbf{b}\|_2^2 \leq \|\mathbf{Ax}^* - \mathbf{b}\|_2^2 + \epsilon^2 \|\Pi_{\mathbf{A}}\mathbf{b}\|_2^2$$

if and only if $\mathbf{x}$ is a solution to the LEA instance $(\mathbf{A}, \mathbf{b}, \epsilon)$.

When the linear system $\mathbf{Ax} = \mathbf{b}$ does not have a solution, a natural notion of solution is any minimizer of Problem (A.1). A simple calculation shows that this is equivalent to requiring that $\mathbf{x}$ is a solution to the linear system $\mathbf{A}^\top \mathbf{Ax} = \mathbf{A}^\top \mathbf{b}$, which always has a solution even when $\mathbf{Ax} = \mathbf{b}$ does not. The system $\mathbf{A}^\top \mathbf{Ax} = \mathbf{A}^\top \mathbf{b}$ is referred to as the *normal equation* associated with $\mathbf{Ax} = \mathbf{b}$ (see [TB97]).

FACT A.2. *$\mathbf{x}^* \in \arg \min_{\mathbf{x} \in \mathbb{R}^n} \|\mathbf{Ax} - \mathbf{b}\|_2^2$, if and only if $\mathbf{A}^\top \mathbf{Ax}^* = \mathbf{A}^\top \mathbf{b}$, and this linear system always has a solution.*

This leads to a natural question: Suppose we want to approximately solve the linear system $\mathbf{A}^\top \mathbf{Ax} = \mathbf{A}^\top \mathbf{b}$. Can we choose our notion of approximation to be equivalent to that of a solution to the LEA instance $(\mathbf{A}, \mathbf{b}, \epsilon)$?

A second natural question is whether we can choose a notion of distance between a proposed solution $\mathbf{x}$ and an optimal solution $\mathbf{x}^* \in \arg \min_{\mathbf{x} \in \mathbb{R}^n} \|\mathbf{Ax} - \mathbf{b}\|_2^2$ s.t. this distance being small is equivalent to $\mathbf{x}$ being a solution to the LEA instance $(\mathbf{A}, \mathbf{b}, \epsilon)$? The answer to both questions is yes, as demonstrated by the following facts:

FACT A.3. *Suppose $\mathbf{x}^* \in \arg \min_{\mathbf{x} \in \mathbb{R}^n} \|\mathbf{Ax} - \mathbf{b}\|_2^2$ then*

$$1. \quad \left\| \mathbf{A}^\top \mathbf{Ax} - \mathbf{A}^\top \mathbf{b} \right\|_{(\mathbf{A}^\top \mathbf{A})^{-1}} = \|\mathbf{Ax} - \Pi_{\mathbf{A}}\mathbf{b}\|_2 = \|\mathbf{x} - \mathbf{x}^*\|_{\mathbf{A}^\top \mathbf{A}}.$$

2. The following statements are each equivalent to $\mathbf{x}$ being a solution to the LEA instance $(\mathbf{A}, \mathbf{b}, \epsilon)$:

- (a) $\left\| \mathbf{A}^\top \mathbf{A} \mathbf{x} - \mathbf{A}^\top \mathbf{b} \right\|_{(\mathbf{A}^\top \mathbf{A})^{-1}} \leq \epsilon \left\| \mathbf{A}^\top \mathbf{b} \right\|_{(\mathbf{A}^\top \mathbf{A})^{-1}}$ if and only if $\mathbf{x}$ is a solution to the LEA instance $(\mathbf{A}, \mathbf{b}, \epsilon)$.
- (b) $\left\| \mathbf{x} - \mathbf{x}^* \right\|_{\mathbf{A}^\top \mathbf{A}} \leq \epsilon \left\| \mathbf{x}^* \right\|_{\mathbf{A}^\top \mathbf{A}}$ if and only if $\mathbf{x}$ is a solution to the LEA instance $(\mathbf{A}, \mathbf{b}, \epsilon)$.
- (c) $\left\| \mathbf{x} - \mathbf{x}^* \right\|_{\mathbf{A}^\top \mathbf{A}} \leq \epsilon \left\| \mathbf{A}^\top \mathbf{b} \right\|_{(\mathbf{A}^\top \mathbf{A})^{-1}}$ if and only if $\mathbf{x}$ is a solution to the LEA instance $(\mathbf{A}, \mathbf{b}, \epsilon)$.

Fact A.3 explains connection between our Definition 2.2, and the usual convention for measuring error in the Laplacian solver literature [ST14]. In this setting, we consider a Laplacian matrix $\mathbf{L}$, which can be written as $\mathbf{L} = \mathbf{A}^\top \mathbf{A} \in \mathbb{R}^{n \times n}$, and a vector $\mathbf{b}$ s.t. $\mathbf{\Pi}_{\mathbf{A}^\top \mathbf{A}} \mathbf{b} = \mathbf{b}$. This condition on $\mathbf{b}$ is easy to verify in the case of Laplacians, since for the Laplacian of a connected graph, $\mathbf{\Pi}_{\mathbf{A}^\top \mathbf{A}} = \mathbf{I} - \frac{1}{n} \mathbf{1} \mathbf{1}^\top$. Additionally, it is also equivalent to the condition that there exists $\mathbf{c}$ s.t. $\mathbf{b} = \mathbf{A}^\top \mathbf{c}$. For Laplacians it is possible to compute both $\mathbf{A}$ and a vector $\mathbf{c}$ s.t. $\mathbf{b} = \mathbf{A}^\top \mathbf{c}$ in time linear in $\text{nnz}(\mathbf{L})$. For Laplacian solvers, the approximation error of an approximate solution $\mathbf{x}$ is measured by the ϵ s.t. $\left\| \mathbf{A}^\top \mathbf{A} \mathbf{x} - \mathbf{b} \right\|_{(\mathbf{A}^\top \mathbf{A})^\dagger} \leq \epsilon \left\| \mathbf{b} \right\|_{(\mathbf{A}^\top \mathbf{A})^\dagger}$. By Fact A.3, we see that this is exactly equivalent to $\mathbf{x}$ being a solution to the LEA instance $(\mathbf{A}, \mathbf{c}, \epsilon)$.

B Iterative Refinement

In this Appendix, we prove the well-known Lemma 2.1 for completeness.

LEMMA 2.1. *If we have a solver for $\text{LEA}(\mathbf{A}, \mathbf{b}, 0.1)$ that works for arbitrary $\mathbf{b}$, then we can obtain a solver for $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon)$ by iterating it $O(\log(1/\epsilon))$ times.*

Proof. Consider the linear equation $\mathbf{A}^\top \mathbf{A} \mathbf{x} = \mathbf{A}^\top \mathbf{b}$. Let $\mathbf{M} = \mathbf{A}^\top \mathbf{A}$ and $\mathbf{p} = \mathbf{A}^\top \mathbf{b}$. We basically do iterative refinement involving the matrix $\mathbf{M}$. Note the desired solution is $\mathbf{x}^* = \mathbf{M}^{-1} \mathbf{p}$. (See Fact A.3). We start with defining $\mathbf{x}^{(0)} = \mathbf{0}$ and

$$\mathbf{x}^{(0)} = \mathbf{0} \text{ and } \mathbf{b}^{(0)} = \mathbf{b} \text{ and } \mathbf{p}^{(0)} = \mathbf{A}^\top \mathbf{b}$$

Starting from $t = 0$, we will run an $\text{LEA}(\mathbf{A}, \mathbf{b}^{(t)}, 0.1)$ solver to produce $\mathbf{x}^{(t+1)}$ such that

$$\left\| \mathbf{x}^{(t+1)} - \mathbf{M}^{-1} \mathbf{p}^{(t)} \right\|_{\mathbf{M}} \leq 0.1 \left\| \mathbf{M}^{-1} \mathbf{p}^{(t)} \right\|_{\mathbf{M}}.$$

We then set

$$\mathbf{b}^{(t+1)} = \left(\mathbf{b}^{(t)} - \mathbf{A} \mathbf{x}^{(t+1)} \right)$$

and we let $\mathbf{p}^{(t+1)} = \mathbf{A}^\top \mathbf{b}^{(t+1)}$, which ensures

$$\mathbf{p}^{(t+1)} = \left(\mathbf{p}^{(t)} - \mathbf{M} \mathbf{x}^{(t+1)} \right)$$

By definition we have

$$\mathbf{M}^{-1} \mathbf{p}^{(t+1)} = \mathbf{M} \mathbf{p}^{(t)} - \mathbf{x}^{(t+1)},$$

and thus

$$\left\| \mathbf{M}^{-1} \mathbf{p}^{(t+1)} \right\|_{\mathbf{M}} \leq 0.1 \left\| \mathbf{M}^{-1} \mathbf{p}^{(t)} \right\|_{\mathbf{M}}.$$

For simplicity let's consider a 2-step version. Feeding in $\mathbf{b}^{(1)}$ as input to the solver in turn gives $\mathbf{x}^{(2)}$ such that

$$\begin{aligned} \left\| \mathbf{x}^{(2)} - \mathbf{M}^{-1} \mathbf{p}^{(1)} \right\|_{\mathbf{M}} &\leq 0.1 \left\| \mathbf{M}^{-1} \mathbf{p}^{(1)} \right\|_{\mathbf{M}} \\ &\leq 0.01 \left\| \mathbf{M}^{-1} \mathbf{p}^{(0)} \right\|_{\mathbf{M}}. \end{aligned}$$

Expanding the LHS gives

$$\mathbf{x}^{(2)} - \mathbf{M}^{-1} \mathbf{p}^{(1)} = \mathbf{x}^{(2)} + \mathbf{x}^{(1)} - \mathbf{M}^{-1} \mathbf{p},$$

which means $\mathbf{x}^{(1)} + \mathbf{x}^{(2)}$ is now a solution with error 0.01. Repeating this gives a $10\times$ smaller error after each step. $\square$

C Improved Dependence on Condition Number Upper Bounds

In Section 3, we saw how our reductions from linear equations to packing LPs lead to algorithms to the $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon)$ problem. We saw how different running times for packing LP algorithms would lead to different running times for the $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon)$ problem. These algorithms resulting from our reductions depend polynomially on an upper bound $K_{\mathbf{A}}$ on the condition number of $\mathbf{A}$.

When there might be a large gap between $\kappa(\mathbf{A})$ and a known upper bound $K_{\mathbf{A}}$, it could be desirable to get reductions that depend on $\kappa(\mathbf{A})$ instead of $K_{\mathbf{A}}$. In this appendix, we describe how to do this, reducing the dependence on $K_{\mathbf{A}}$ to logarithmic. The reduction is straightforward, but we include it for completeness. In the proof of Corollary 3.1 we constructed a solver for $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon)$ with running time $\tilde{O}(n^\beta (n K_{\mathbf{A}}^3)^\alpha \log(1/\epsilon))$, assuming a packing LP solver with running time $\tilde{O}(n^\beta / \epsilon^\alpha)$. We can turn this into an $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon)$ solver with running time $\tilde{O}(n^\beta (n \kappa(\mathbf{A})^3)^\alpha \log(K_{\mathbf{A}}/\epsilon))$.

One approach to doing so is to make sure we can certify approximation quality of the solution when a sufficiently good one is found. When run with a too-small (invalid) “bound” on $\kappa(\mathbf{A})$, the algorithms we

construct may fail, but they do not take more time than the previously derived bound. Consequently, if can check whether a solution is valid, we can then search for an adequate bound on $\kappa(\mathbf{A})$ using successive doublings until the algorithm succeeds.

We assume an convenient normalization of $\mathbf{A}$ and $\mathbf{b}$ (see Section 4), namely

1. $\max_{ij} |\mathbf{A}_{ij}| = 1$
2. $\|\mathbf{A}^\top \mathbf{b}\|_2 = 1$

Ultimately, our goal is to find $\mathbf{x}$ s.t.

$$(C.2) \quad \|\mathbf{Ax} - \Pi_{\mathbf{A}} \mathbf{b}\|_2 \leq \epsilon \|\Pi_{\mathbf{A}} \mathbf{b}\|_2$$

We will show that for any ϵ' , if we find a solution with

$$(C.3) \quad \|\mathbf{Ax} - \Pi_{\mathbf{A}} \mathbf{b}\|_2 \leq \epsilon' \|\Pi_{\mathbf{A}} \mathbf{b}\|_2$$

then this implies

$$(C.4) \quad \|\mathbf{A}^\top \mathbf{Ax} - \mathbf{A}^\top \mathbf{b}\|_2 \leq K_{\mathbf{A}} \sqrt{mn} \cdot \epsilon' \|\mathbf{A}^\top \mathbf{b}\|_2$$

We can also show that

$$(C.5) \quad \|\mathbf{A}^\top \mathbf{Ax} - \mathbf{A}^\top \mathbf{b}\|_2 \leq \epsilon'' \|\mathbf{A}^\top \mathbf{b}\|_2$$

implies

$$(C.6) \quad \|\mathbf{Ax} - \Pi_{\mathbf{A}} \mathbf{b}\|_2 \leq K_{\mathbf{A}} \sqrt{mn} \cdot \epsilon'' \|\Pi_{\mathbf{A}} \mathbf{b}\|_2$$

Equation (C.5) has the advantage of being a linear time checkable condition.

Now suppose we the run $\tilde{O}(n^\beta (nB^3)^\alpha \log(1/\epsilon'))$ time algorithm for $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon')$, with $\epsilon' = \epsilon/(K_{\mathbf{A}} \sqrt{mn})^2$. If the algorithm succeeds, then Equation (C.3) is satisfied and this implies Equation (C.4). This latter condition we can check in linear time, or equivalently, we check that Equation (C.5) is satisfied with $\epsilon'' = \epsilon/(K_{\mathbf{A}} \sqrt{mn})$. This then implies Equation (C.6) is satisfied, and hence Equation (C.2), by our choice of ϵ' and ϵ'' .

Thus, whenever if we accept whenever Equation (C.5) is satisfied, we only accept sufficiently accurate solutions. On the other hand, if the $\tilde{O}(n^\beta (nB^3)^\alpha \log(1/\epsilon'))$ time algorithm is run with B being an upper bound on $\kappa(\mathbf{A})$, it will produce a solution satisfying (C.3) and hence Equations (C.5) and (C.2).

Now, finally, we run the $\tilde{O}(n^\beta (nB^3)^\alpha \log(1/\epsilon'))$ time algorithm with successive doublings of B and each time check if Equation (C.5) is satisfied for the output. This must succeed when $B \geq \kappa(\mathbf{A})$, but possibly before, in which case we just get a correct solution sooner.

The running time will be dominated by the last call to the algorithm, and on that last call we must have $B < 2\kappa(\mathbf{A})$. Based on our choice of ϵ' , we now get an overall running time of $\tilde{O}(n^\beta (n\kappa(\mathbf{A})^3)^\alpha \log(K_{\mathbf{A}}/\epsilon))$.

All that remains is to establish the implication from Equation (C.3) to (C.4), and from (C.5) to (C.6). First, we note that with our normalization of $\mathbf{A}$ and $\mathbf{b}$, one can show $1 \leq \sigma_{\max}(\mathbf{A}) \leq \sqrt{mn}$. This further implies that $1/\kappa(\mathbf{A}) \leq \sigma_{\min}(\mathbf{A})$. Combining this with Fact A.3, we get

$$\begin{aligned} \|\mathbf{A}^\top \mathbf{Ax} - \mathbf{A}^\top \mathbf{b}\|_2 &\leq \sigma_{\max} \|\mathbf{A}^\top \mathbf{Ax} - \mathbf{A}^\top \mathbf{b}\|_{(\mathbf{A}^\top \mathbf{A})^{-1}} \\ &\leq \sqrt{mn} \|\mathbf{A}^\top \mathbf{Ax} - \mathbf{A}^\top \mathbf{b}\|_{(\mathbf{A}^\top \mathbf{A})^{-1}} \\ &= \sqrt{mn} \|\mathbf{Ax} - \Pi_{\mathbf{A}} \mathbf{b}\|_2 \end{aligned}$$

and

$$\begin{aligned} \|\mathbf{A}^\top \mathbf{Ax} - \mathbf{A}^\top \mathbf{b}\|_2 &\geq \sigma_{\min} \|\mathbf{A}^\top \mathbf{Ax} - \mathbf{A}^\top \mathbf{b}\|_{(\mathbf{A}^\top \mathbf{A})^{-1}} \\ &\geq \frac{1}{\kappa(\mathbf{A})} \|\mathbf{A}^\top \mathbf{Ax} - \mathbf{A}^\top \mathbf{b}\|_{(\mathbf{A}^\top \mathbf{A})^{-1}} \\ &= \frac{1}{\kappa(\mathbf{A})} \|\mathbf{Ax} - \Pi_{\mathbf{A}} \mathbf{b}\|_2 \end{aligned}$$

Also, by Claim 4.2

$$\|\Pi_{\mathbf{A}} \mathbf{b}\|_2 \in \left[\frac{1}{\sqrt{mn}} \|\mathbf{A}^\top \mathbf{b}\|_2, \kappa(\mathbf{A}) \|\mathbf{A}^\top \mathbf{b}\|_2 \right]$$

Thus, if we assume Equation (C.3), we get

$$\begin{aligned} \|\mathbf{A}^\top \mathbf{Ax} - \mathbf{A}^\top \mathbf{b}\|_2 &\leq \sqrt{mn} \|\mathbf{Ax} - \Pi_{\mathbf{A}} \mathbf{b}\|_2 \\ &\leq \sqrt{mn} \epsilon' \|\Pi_{\mathbf{A}} \mathbf{b}\|_2 \leq \epsilon' \kappa(\mathbf{A}) \sqrt{mn} \|\mathbf{A}^\top \mathbf{b}\|_2, \end{aligned}$$

i.e. we conclude Equation (C.4) holds.

Meanwhile, if we assume Equation (C.5), we get

$$\begin{aligned} \|\mathbf{Ax} - \Pi_{\mathbf{A}} \mathbf{b}\|_2 &\leq \kappa(\mathbf{A}) \|\mathbf{A}^\top \mathbf{Ax} - \mathbf{A}^\top \mathbf{b}\|_2 \\ &\leq \kappa(\mathbf{A}) \epsilon'' \|\mathbf{A}^\top \mathbf{b}\|_2 \\ &\leq \epsilon'' \kappa(\mathbf{A}) \sqrt{mn} \|\Pi_{\mathbf{A}} \mathbf{b}\|_2 \\ &\leq \epsilon'' K_{\mathbf{A}} \sqrt{mn} \|\Pi_{\mathbf{A}} \mathbf{b}\|_2 \end{aligned}$$

so Equation (C.6) holds.

Of course, we can use similar modifications to the algorithm to get an new algorithm for the sparse case of $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon)$ assuming an $\tilde{O}(N/\epsilon^\alpha)$ time algorithm for packing LPs with N non-zeros. This will give an algorithm for $\text{LEA}(\mathbf{A}, \mathbf{b}, \epsilon)$ with running time $\tilde{O}(\text{nnz}(\mathbf{A}) (m^{3/2} \kappa(\mathbf{A})^3)^\alpha \log(K_{\mathbf{A}}/\epsilon))$.

D Reducing a General LP Instance to a Packing LP Instance

In this appendix, we show a reduction from a well-conditioned LP instance to a packing LP instance, by slightly modifying our reduction from Section 4. Let us first introduce some notations and definitions for a general LP.

Given $\mathbf{A} \in \mathbb{R}^{m \times n}$, $\mathbf{b} \in \mathbb{R}^m$, $\mathbf{c} \in \mathbb{R}^n$. We use $L = (\mathbf{A}, \mathbf{b}, \mathbf{c})$ to denote an LP written in the canonical form:

$$(D.7) \quad \begin{aligned} \max \quad & \mathbf{c}^\top \mathbf{x} \\ \text{s.t.} \quad & \mathbf{A}\mathbf{x} \leq \mathbf{b} \\ & \mathbf{x} \geq \mathbf{0} \end{aligned}$$

Let $\text{OPT}(\mathbf{A}, \mathbf{b}, \mathbf{c})$ be the optimal value of LP $(\mathbf{A}, \mathbf{b}, \mathbf{c})$. Let

$$\|L\|_\infty \stackrel{\text{def}}{=} \max\{\|\mathbf{A}\|_{\max}, \|\mathbf{b}\|_\infty, \|\mathbf{c}\|_\infty\}.$$

where $\|\mathbf{A}\|_{\max} = \max_{i,j} |\mathbf{A}_{ij}|$.

Renegar [Ren95] introduced the notion of condition numbers to measure the complexity of solving LPs. Define the primal and dual condition numbers of a LP L as $\kappa_P(L) \stackrel{\text{def}}{=} \sup\{\delta : L + \Delta L \text{ is primal feasible if } \|\Delta L\|_\infty < \delta\}$ and $\kappa_D(L) \stackrel{\text{def}}{=} \sup\{\delta : L + \Delta L \text{ is dual feasible if } \|\Delta L\|_\infty < \delta\}$. Note κ_P and κ_D are not scaling invariant. Renegar proved the following lemma.

LEMMA D.1. ([Ren95]) *Let $L = (\mathbf{A}, \mathbf{b}, \mathbf{c})$ be an LP instance with primal optimal solution $\mathbf{x}^*$ and dual optimal solution $\mathbf{y}^*$. Then*

$$\begin{aligned} \|\mathbf{x}^*\|_1 &\leq \frac{\max\{\|\mathbf{b}\|_\infty, -\text{OPT}(L)\}}{\kappa_P(L)} \\ \|\mathbf{y}^*\|_1 &\leq \frac{\max\{\|\mathbf{c}\|_\infty, \text{OPT}(L)\}}{\kappa_D(L)} \end{aligned}$$

We state our theorem in the following.

THEOREM D.1. *If we can $(1 - \epsilon)$ -approximately solve a packing LP with N non-zeros in time $\tilde{O}(N \log(1/\epsilon))$, then given a feasible LP instance $L = (\mathbf{A}, \mathbf{b}, \mathbf{c})$ that has N non-zeros, and polynomially bounded $\|L\|_\infty$, optimum and condition numbers, we can compute a solution $\mathbf{x} \geq \mathbf{0}$ satisfying*

$$(D.8) \quad \mathbf{c}^\top \mathbf{x} \geq \text{OPT}(\mathbf{A}, \mathbf{b}, \mathbf{c}) - \epsilon', \mathbf{A}\mathbf{x} \leq \mathbf{b} + \epsilon' \mathbf{1}$$

in time $\tilde{O}(N \log(1/\epsilon'))$.

Proof. Let m be the number of rows of $\mathbf{A}$ and n be the number of columns of $\mathbf{A}$.

Let $\mathbf{x}^*$ be a primal optimal solution of L and let $\mathbf{y}^*$ be a dual optimal solution of L . Since L is feasible and bounded, by the strong duality theorem, $\mathbf{x}^*, \mathbf{y}^*$ is a feasible solution to the following LP.

$$(D.9) \quad \begin{aligned} \mathbf{c}^\top \mathbf{x} &= \mathbf{b}^\top \mathbf{y} \\ \mathbf{A}\mathbf{x} &\leq \mathbf{b} \\ \mathbf{A}^\top \mathbf{y} &\geq \mathbf{c} \\ \mathbf{x}, \mathbf{y} &\geq \mathbf{0} \end{aligned}$$

Let U be a real number satisfying $U \geq \|\mathbf{x}^*\|_1 + \|\mathbf{y}^*\|_1$, which can be upper bounded by Lemma D.1. One can use binary search to find U . For simplicity let $K = \|L\|_\infty$.

Similar to our construction of $\text{PLP}_{\text{sparse}}$ in Section 4.2, we introduce a new variable and a new bounding box constraint for each inequality and set the objective properly, and we get the packing LP (D.10). We can check that the packing LP has $O(N)$ non-zeros and the optimum is $(n + m + 1)U \stackrel{\text{def}}{=} U'$.

Let $\mathbf{x}, \mathbf{y}$ be a $(1 - \epsilon)$ -approximate solution to the packing LP (D.10). Then

$$\begin{aligned} \mathbf{A}\mathbf{x} - \mathbf{b} &\leq \epsilon K U' \mathbf{1} \\ \mathbf{c} - \mathbf{A}^\top \mathbf{y} &\leq \epsilon K U' \mathbf{1} \\ |\mathbf{c}^\top \mathbf{x} - \mathbf{b}^\top \mathbf{y}| &\leq \epsilon K U' \end{aligned}$$

From the third inequality above,

$$(D.11) \quad \mathbf{c}^\top \mathbf{x} \geq \mathbf{b}^\top \mathbf{y} - \epsilon K U' \geq \mathbf{b}^\top \tilde{\mathbf{y}}^* - \epsilon K U'$$

where $\tilde{\mathbf{y}}^*$ is an optimal solution of the following perturbed dual:

$$\begin{aligned} \min \quad & \mathbf{b}^\top \mathbf{y} \\ \text{s.t.} \quad & \mathbf{A}^\top \mathbf{y} \geq \mathbf{c} - \epsilon K U' \mathbf{1} \\ & \mathbf{y} \geq \mathbf{0} \end{aligned}$$

Given $L(\mathbf{A}, \mathbf{b}, \mathbf{c})$ is dual feasible and bounded, this perturbed dual linear program is feasible and bounded. Its dual is the following perturbed version of $L(\mathbf{A}, \mathbf{b}, \mathbf{c})$:

$$\begin{aligned} \max \quad & \mathbf{c}^\top \mathbf{x} - \epsilon K U' \mathbf{1}^\top \mathbf{x} \\ \text{s.t.} \quad & \mathbf{A}\mathbf{x} \leq \mathbf{b} \\ & \mathbf{x} \geq \mathbf{0} \end{aligned}$$

Let $\tilde{\mathbf{x}}^*$ be an optimal solution of this perturbed LP. Then by the strong duality theorem and Equation (D.11),

$$\begin{aligned} \mathbf{c}^\top \mathbf{x} &\geq \mathbf{c}^\top \tilde{\mathbf{x}}^* - \epsilon K U' \mathbf{1}^\top \tilde{\mathbf{x}}^* - \epsilon K U' \\ &\geq \mathbf{c}^\top \mathbf{x}^* - \epsilon K U' \mathbf{1}^\top \mathbf{x}^* - \epsilon K U' \\ &\geq \mathbf{c}^\top \mathbf{x}^* - \epsilon \frac{K U' \max\{\|\mathbf{b}\|_\infty, -\text{OPT}(L)\}}{\kappa_P(L)} - \epsilon K U' \end{aligned}$$

$$\begin{aligned}
(D.10) \quad & \max \quad \sum_{1 \leq j \leq n} \theta_j + \sum_{1 \leq i \leq m} \eta_i + \zeta \\
& \text{s.t.} \quad K\gamma + \sum_{1 \leq j \leq n: c_j \neq 0} (K + \mathbf{c}_j) \mathbf{x}_j + \sum_{1 \leq i \leq m: \mathbf{b}_i \neq 0} (K - \mathbf{b}_i) \mathbf{y}_i \leq KU \\
& \quad K\gamma + \sum_{1 \leq j \leq n: c_j \neq 0} (K - \mathbf{c}_j) \mathbf{x}_j + \sum_{1 \leq i \leq m: \mathbf{b}_i \neq 0} (K + \mathbf{b}_i) \mathbf{y}_i \leq KU \\
& \quad \gamma + \sum_{1 \leq j \leq n: c_j \neq 0} \mathbf{x}_j + \sum_{1 \leq i \leq m: \mathbf{b}_i \neq 0} \mathbf{y}_i \leq U \\
& \quad K\beta_i + \sum_{1 \leq j \leq n: \mathbf{A}_{ij} \neq 0} (K + \mathbf{A}_{ij}) \mathbf{x}_j \leq \mathbf{b}_i + KU, \forall 1 \leq i \leq m \\
& \quad \beta_i + \sum_{1 \leq j \leq n: \mathbf{A}_{ij} \neq 0} \mathbf{x}_j \leq U, \forall 1 \leq i \leq m \\
& \quad K\alpha_j + \sum_{1 \leq i \leq m: \mathbf{A}_{ij} \neq 0} (K - \mathbf{A}_{ij}) \mathbf{y}_i \leq -\mathbf{c}_j + KU, \forall 1 \leq j \leq n \\
& \quad \alpha_j + \sum_{1 \leq i \leq m: \mathbf{A}_{ij} \neq 0} \mathbf{y}_i \leq U, \forall 1 \leq j \leq n \\
& \quad \gamma \geq 0, \boldsymbol{\alpha}, \boldsymbol{\beta}, \mathbf{x}, \mathbf{y} \geq \mathbf{0}
\end{aligned}$$

Here, $\theta_j \stackrel{\text{def}}{=} \alpha_j + \sum_{1 \leq i \leq m: \mathbf{A}_{ij} \neq 0} \mathbf{y}_i, \forall 1 \leq j \leq n$, $\eta_i \stackrel{\text{def}}{=} \beta_i + \sum_{1 \leq j \leq n: \mathbf{A}_{ij} \neq 0} \mathbf{x}_j, \forall 1 \leq i \leq m$, and $\zeta \stackrel{\text{def}}{=} \gamma + \sum_{1 \leq j \leq n: c_j \neq 0} \mathbf{x}_j + \sum_{1 \leq i \leq m: \mathbf{b}_i \neq 0} \mathbf{y}_i$. All θ_j, β_i, ζ are shorthands rather than new variables.

The last inequality is due to Lemma D.1.

To make solution $\mathbf{x}$ satisfy Equation (D.8), it suffices to set

$$\epsilon \leq \min \left\{ \frac{\epsilon'}{2KU'}, \frac{\epsilon' \kappa_P(L)}{2KU' \max\{\|\mathbf{b}\|_\infty, -\text{OPT}(L)\}} \right\}$$

If we can $(1 - \epsilon)$ -approximately solve the packing LP (D.10) in time $\tilde{O}(N \log(1/\epsilon))$, then we can compute a solution $\mathbf{x}$ satisfying Equation (D.8) in time

$$\tilde{O} \left(N \log \left(\frac{KU(m+n) \max\{1, |\text{OPT}(L)|\}}{\kappa_P(L) \epsilon'} \right) \right)$$

Given all parameters in the log is polynomially bounded, the run time is $\tilde{O}(N \log(1/\epsilon'))$. $\square$