



Mathematics of Operations Research

Publication details, including instructions for authors and subscription information:
<http://pubsonline.informs.org>

On the Efficiency of Random Permutation for ADMM and Coordinate Descent

Ruoyu Sun, Zhi-Quan Luo, Yinyu Ye

To cite this article:

Ruoyu Sun, Zhi-Quan Luo, Yinyu Ye (2020) On the Efficiency of Random Permutation for ADMM and Coordinate Descent. Mathematics of Operations Research 45(1):233-271. <https://doi.org/10.1287/moor.2019.0990>

Full terms and conditions of use: <https://pubsonline.informs.org/Publications/Librarians-Portal/PubsOnLine-Terms-and-Conditions>

This article may be used only for the purposes of research, teaching, and/or private study. Commercial use or systematic downloading (by robots or other automatic processes) is prohibited without explicit Publisher approval, unless otherwise noted. For more information, contact permissions@informs.org.

The Publisher does not warrant or guarantee the article's accuracy, completeness, merchantability, fitness for a particular purpose, or non-infringement. Descriptions of, or references to, products or publications, or inclusion of an advertisement in this article, neither constitutes nor implies a guarantee, endorsement, or support of claims made of that product, publication, or service.

Copyright 2019, INFORMS

Please scroll down for article—it is on subsequent pages



With 12,500 members from nearly 90 countries, INFORMS is the largest international association of operations research (O.R.) and analytics professionals and students. INFORMS provides unique networking and learning opportunities for individual professionals, and organizations of all types and sizes, to better understand and use O.R. and analytics tools and methods to transform strategic visions and achieve better outcomes.

For more information on INFORMS, its publications, membership, or meetings visit <http://www.informs.org>

On the Efficiency of Random Permutation for ADMM and Coordinate Descent

Ruoyu Sun,^a Zhi-Quan Luo,^{b,c} Yinyu Ye^d

^a Department of Industrial and Enterprise Systems Engineering and Coordinated Science Laboratory, University of Illinois Urbana–Champaign, Champaign, Illinois 61801; ^b Chinese University of Hong Kong, 518172 Shenzhen, China; ^c Shenzhen Research Institute of Big Data, 518172 Shenzhen, China; ^d Department of Management Science and Engineering, Stanford University, Palo Alto, California 94305

Contact: ruoyus@illinois.edu,  <http://orcid.org/0000-0003-2487-5322> (RS); luozq@cuhk.edu.cn (Z-QL); yyye@stanford.edu (YY)

Received: December 26, 2017

Revised: September 17, 2018;

November 16, 2018

Accepted: December 25, 2018

Published Online in Articles in Advance:

November 6, 2019

MSC2010 Subject Classification: Primary: 90C02; secondary: 65F02

OR/MS Subject Classification: Analysis of Algorithms: Computational Complexity

<https://doi.org/10.1287/moor.2019.0990>

Copyright: © 2019 INFORMS

Abstract. Random permutation is observed to be powerful for optimization algorithms: for multiblock ADMM (alternating direction method of multipliers), whereas the classical cyclic version diverges, the randomly permuted version converges in practice; for BCD (block coordinate descent), the randomly permuted version is typically faster than other versions. In this paper we provide strong theoretical evidence that random permutation has positive effects on ADMM and BCD, by analyzing randomly permuted ADMM (RP-ADMM) for solving linear systems of equations, and randomly permuted BCD (RP-BCD) for solving unconstrained quadratic problems. First, we prove that RP-ADMM converges in expectation for solving systems of linear equations. The key technical result is that the spectrum of the expected update matrix of RP-BCD lies in $(-1/3, 1)$, instead of the typical range $(-1, 1)$. Second, we establish expected convergence rates of RP-ADMM for solving linear systems and RP-BCD for solving unconstrained quadratic problems. This expected rate of RP-BCD is $O(n)$ times better than the worst-case rate of cyclic BCD, thus establishing a gap of at least $O(n)$ between RP-BCD and cyclic BCD. To analyze RP-BCD, we propose a conjecture of a new matrix algebraic mean-geometric mean inequality and prove a weaker version of it.

Funding: This research is supported by the leading talents of Guangdong Province program [Grant 00201501]; the Development and Reform Commission of Shenzhen Municipality; Air Force Office of Scientific Research [Grant FA9550-12-1-0396]; National Science Foundation of China [Grant 61731018]; and the National Science Foundation [Grant CCF 1755847].

Keywords: random permutation • ADMM • coordinate descent • convergence • matrix AM-GM inequality • spectral radius

1. Introduction

A simple yet powerful idea for solving large-scale computational problems is to iteratively solve smaller subproblems. The applications of this idea include coordinate descent (CD), projection onto convex sets, and SGD (stochastic gradient descent). They are well suited for large-scale *unconstrained* optimization problem (see, e.g., Wright [44], Shi et al. [35] for recent reviews of CD methods) because it decomposes a large problem into small subproblems. The decomposition idea is crucial for huge problems owing to both the cheap per-iteration cost and small memory requirement. Moreover, this idea is “orthogonal” to other large-scale optimization ideas, such as first-order methods (using only gradient information) and random projection, and thus can be easily combined with other ideas.

This paper is motivated by a natural question: how should we extend the decomposition idea to solve problems *with constraints*? We consider a constrained minimization problem with a convex objective function and linear constraints (this is for motivation; our analysis is for a much simpler version):

$$\begin{aligned} \min_{x_1, \dots, x_n} \quad & f(x_1, x_2, \dots, x_n), \\ \text{s.t.} \quad & A_1 x_1 + \dots + A_n x_n = b, \\ & x_i \in \mathcal{X}_i, \quad i = 1, \dots, n, \end{aligned} \tag{1}$$

where $A_i \in \mathbb{R}^{N \times d_i}$, $b \in \mathbb{R}^{N \times 1}$, $\mathcal{X}_i \subseteq \mathbb{R}^{d_i}$ is a closed convex set, $i = 1, \dots, n$, and $f : \mathbb{R}^{d_1 + d_2 + \dots + d_n} \rightarrow \mathbb{R}$ is a closed convex function. Many machine learning and engineering problems can be cast into linearly constrained optimization problems with two blocks (see Boyd et al. [3] for many examples) or more than two blocks (e.g., linear programming, robust principal component analysis, composite regularizers for structured sparsity; see Chen et al. [6] and Wang et al. [43] for more examples).

To apply the decomposition idea to a constrained problem, one possible way is to form the augmented Lagrangian function and perform coordinate descent for the primal problem and a gradient step for the dual problem (i.e., combining block coordinate descent [BCD] with the augmented Lagrangian method) to obtain the so-called alternating direction method of multipliers (ADMM). ADMM was originally proposed in Glowinski and Marroco [12] (see also Chan and Glowinski [5], Gabay and Mercier [11]) to solve problem (1) when there are only two blocks (i.e., $n = 2$) and the objective function is separable. It is natural and computationally beneficial to extend the original ADMM directly to solve the general n -block problem (1) via the following procedure:

$$\begin{cases} x_1^{k+1} = \arg \min_{x_1 \in \mathcal{X}_1} \mathcal{L}(x_1, x_2^k, \dots, x_n^k; \mu^k), \\ \vdots \\ x_n^{k+1} = \arg \min_{x_n \in \mathcal{X}_n} \mathcal{L}(x_1^{k+1}, \dots, x_{n-1}^{k+1}, x_n; \mu^k), \\ \mu^{k+1} = \mu^k - \beta(A_1 x_1^{k+1} + \dots + A_n x_n^{k+1} - b), \end{cases} \quad (2)$$

where the augmented Lagrangian function

$$\mathcal{L}(x_1, \dots, x_n; \mu) = f(x_1, \dots, x_n) - \mu^T \left(\sum_i A_i x_i - b \right) + \frac{\beta}{2} \left\| \sum_i A_i x_i - b \right\|^2. \quad (3)$$

The convergence of the direct extension of ADMM to a multiblock case had been an open question, until a counter-example was recently given in Chen et al. [6]. More specifically, Chen et al. [6] showed that even for the simplest scenario in which the objective function is 0 and the number of blocks is 3, ADMM can be divergent for a certain choice of $A = [A_1, A_2, A_3]$. There are several proposals to overcome the drawback (see, e.g., [4, 8, 9, 10, 14, 15, 16, 17, 18, 19, 20, 21, 25, 26, 27, 28, 29, 37]), but they either need to restrict the range of original problems being solved, add additional cost in each step of computation, or limit the stepsize in updating the Lagrange multipliers. These solutions typically slow down the performance of ADMM for solving most practical problems. Moreover, it is not clear how to compare the convergence speed of these algorithms because they typically contain different parameters. One may ask whether a “minimal” modification of cyclic multiblock ADMM (2) can lead to convergence, and whether we can provide some convergence speed analysis that is easy to interpret.

One of the simplest modifications of (2) is to add randomness to the update order. Randomness has been very useful in the analysis of BCD methods and SGD methods. In particular, a recent work by Sun and Ye [41] showed that randomized CD (R-CD) can be up to $O(n^2)$ times faster than cyclic CD (C-CD) for quadratic minimization in the worst case, where n is the number of variables.¹ Another example is the comparison of incremental aggregated gradient in Blatt et al. [2] and its randomized version, stochastic average gradient (Schmidt et al. [33]); it turns out that the introduction of randomness leads to better iteration complexity bounds. There is also some recent study on a randomly permuted version of pure SGD (Gürbüzbalaban et al. [13]). These examples show that randomization may improve the algorithm in theory and in practice.

It is important to note that the iteration complexity bounds for randomized algorithms are usually established for independent randomization (sampling with replacement), whereas in practice, random permutation (sampling without replacement) has been reported to exhibit faster convergence (e.g., Shalev-Shwartz and Zhang [34], Recht and Ré [32], Sun [38]). Interestingly, our simulation shows that for solving linear system of equations, randomly permuted ADMM (RP-ADMM) always converges, but independently randomized versions of ADMM can be divergent even for Gaussian data. Therefore, we focus on the analysis of RP-ADMM in this paper.

Random permutation is known to be notoriously difficult to analyze. Even for unconstrained quadratic minimization, the convergence rate of randomly permuted BCD (RP-BCD) is poorly understood. Many existing works treated cyclic BCD and RP-BCD together (Beck and Tetruashvili [1], Sun and Hong [39], Tseng [42]), and thus the best known convergence rate of RP-BCD for general convex problems are in fact the same as that of C-BCD (Sun and Hong [39]). However, in light of a recent study that established an up to $O(n^2)$ gap between cyclic CD and R-CD [41], it is unlikely that RP-CD has the same convergence rate as C-CD, because that would imply RP-CD could be $O(n^2)$ times slower than R-CD. For the special example that demonstrates the gap between C-CD and R-CD, it was shown recently that RP-CD is faster than R-CD² (Wright and Lee [45]). However, the general quadratic case seems to be quite difficult, probably owing to its close connection to a matrix AM-GM (algebraic mean-geometric mean) inequality (Recht and Ré [31]), the difficulty of which is essentially to prove an inequality in noncommutative algebra.

1.1. Summary of Contributions

We consider two extremes of a general RP-ADMM: (i) the objective is zero (i.e., RP-ADMM for solving a linear system); (ii) the constraint is zero and the objective is a quadratic function (i.e., RP-BCD for solving quadratic minimization). Owing to the lack of understanding of random permutation for quadratic minimization as discussed previously, we restrict to the two cases in this paper.

The first result of this paper is the expected convergence of RP-ADMM for solving linear systems. More specifically, when the objective function is zero and the constraint is a nonsingular square linear system of equations, the expected output of randomly permuted ADMM converges to the unique primal-dual optimal solution. A major technical result in this proof is that the eigenvalues of the expected iteration matrix of RP-BCD for quadratic problems lie in $(-1/3, 1)$, instead of the typical range $(-1, 1)$.

The second result is about the expected convergence rate of RP-ADMM for solving linear systems and RP-BCD for solving quadratic problems. We show that RP-BCD for a convex quadratic minimization problem with equal diagonal entries has expected iteration complexity $O(n \frac{\lambda_{\text{avg}}}{\lambda_{\text{min}}} \log(1/\epsilon))$, where λ_{avg} and λ_{min} are the average eigenvalue and the minimum eigenvalue of the coefficient matrix, and one “iteration” here means a cycle of updating all blocks. This improves an existing bound of $O(n^2 \frac{\lambda_{\text{avg}}}{\lambda_{\text{min}}} \log(1/\epsilon))$ for RP-BCD by a factor of n . Built on this result, we further show that RP-ADMM for solving linear systems achieves the same expected iteration complexity bound $O(n \frac{\lambda_{\text{avg}}}{\lambda_{\text{min}}} \log(1/\epsilon))$.

Technically, we provide a simple and clean proof of the expected convergence, by applying a classic result on the eigenvalues of Jordan product. For proving the expected convergence rate, we propose a new variant of the matrix AM-GM inequality conjecture and prove a weaker version of this conjecture.

Our result shows that random permutation may be a good answer to the question of “how to apply the decomposition idea to solve constrained problems.” Because multiblock BCD is widely used for large-scale unconstrained problems, we expect multiblock RP-ADMM to be a good candidate for large-scale linearly constrained problems. Our result provides one of the few direct analyses of random permutation in optimization algorithms and offers an explanation of the mysterious gap between RP-ADMM and cyclic ADMM. As reflected by the proof, the intuition is that random permutation provides “three-level symmetrization” that adjusts the spectrum of the update matrix. On the basis of the analysis for RP-ADMM, we are able to improve the best known complexity of RP-BCD for equally diagonal quadratic problems by a factor of n , when expressing the complexity only in terms of the quantity $\frac{\lambda_{\text{avg}}}{\lambda_{\text{min}}}$.

1.2. Related Works

This paper is a stronger version of a previous technical report by Sun et al. [40], which was not published. Another related work is the paper by Chen et al. [7], which modifies the proof of [40] to make it work with a quadratic objective function.

We highlight a few novel contributions of the present paper (neither in the original technical report [40] nor in the paper [7]).

- i. The present paper provides a much simpler proof for the result of expected convergence.
- ii. The present paper provides the first convergence rate analysis of RP-ADMM. See Theorem 4 and the proof in Section 4.5, Section 7.2, and Section 7.1.
- iii. The present paper provides an improved convergence rate analysis of RP-BCD. See Theorem 3 and the proof in Section 4.3 and Section 7.3.
- iv. The present paper introduces a theory-motivated algorithm, Bernoulli-ADMM, which reduces the sampling time yet still achieves the expected convergence. To our knowledge, this update order has not appeared before even in other algorithm setups. See Section 2.5 and Proposition 1.

Besides the technical contributions, we want to emphasize that the present paper is not just adding new results to our previous technical report [40] but actually completes a missing step of the story. From a mathematical point of view, the most striking consequence of our original proof is that the spectral radius of RP-BCD lies in a smaller region $(-1/3, 1)$. It is natural to think that this fundamental fact should have an impact on the analysis of original RP-BCD. The present paper fills this gap by showing that this result can help build an $O(n)$ gap between the (expected) convergence rate of RP-BCD and cyclic BCD. A general message is that, on one hand, to understand constrained optimization we have to understand unconstrained optimization (analyzing ADMM reduces to analyzing BCD); on the other hand, analyzing constrained optimization helps improve the understanding of unconstrained optimization (the analysis of ADMM leads to progress in BCD). We find this interaction

between unconstrained optimization (BCD) and constrained optimization (ADMM) fascinating. The whole story is only revealed in the present paper but not in the previous technical report [40] or [7].

Besides the above unique aspects, the present paper inherits some interesting numerical findings from the technical report [40] that do not appear in [7]. We find that cyclic ADMM diverges with probability 1 for many random distributions of data, thus showing that the seemingly surprising divergence behavior reported in Chen et al. [6] is quite common. However, it is easy to miss this finding if one uses the Gaussian distribution to generate data. Another interesting finding is that the independently randomized version of ADMM diverges with probability 1 for Gaussian data but not for the counter-example in [6], preventing us from analyzing the independently randomized version. Without these findings, the motivation of studying RP-ADMM would be less clear. See Section 2.4 and Section 8.

1.3. Notation and Organization

1.3.1. Notation. For a matrix X , we denote $X(i, j)$ as the (i, j) -th entry of X , $\text{eig}(X)$ as the set of eigenvalues of X , $\rho(X)$ as the spectral radius of X (i.e., the maximum modulus of the eigenvalues of X), $\|X\|$ as the spectral norm of X , and X^T as the transpose of X . When X is block partitioned, we use $X[i, j]$ to denote the (i, j) -th block of X . When X is a real symmetric matrix, let $\lambda_{\max}(X)$ and $\lambda_{\min}(X)$ denote the maximum and minimum eigenvalue of X , respectively. For two real symmetric matrices X_1 and X_2 , $X_1 > X_2$ (resp. $X_1 \geq X_2$) means $X_1 - X_2$ is positive definite (resp. positive semidefinite). We use I_m to denote the identity matrix with dimension m , and we will simply use I when it is clear from the context what the dimension is. For square matrices $U_i \in \mathbb{R}^{u_i \times u_i}$, $i = 1, \dots, k$, we denote $\text{Diag}(U_1, U_2, \dots, U_k)$ as the block-diagonal matrix with U_i being the i -th diagonal block.

1.3.2. Organization. In Section 2, we present three versions of randomized ADMM, with an emphasis on RP-ADMM. In Section 3, we present our main results Theorems 1 and 2 and their proofs. The subsequent sections are devoted to the proofs of the two technical results Lemmas 1 and 2, which are used in the proof of Theorem 2. In particular, the proof of Lemma 1 is given in Section 5, and the proof of Lemma 2 is given in Section 6.

2. Algorithms

In this section, we will present both randomly permuted and independently randomized versions of ADMM for solving (1) and specialize RP-ADMM for solving a square system of equations. We also present a rather novel algorithm, Bernoulli-randomized ADMM (motivated by our proof).

2.1. Randomly Permuted ADMM

In this subsection, we first propose RP-ADMM for solving the general optimization problem (1), then we present the update equation of RP-ADMM for solving a linear system of equations.

Define Γ as

$$\Gamma \triangleq \{\sigma \mid \sigma \text{ is a permutation of } \{1, \dots, n\}\}. \quad (4)$$

At each round, we draw a permutation σ of $\{1, \dots, n\}$ uniformly at random from Γ , and update the primal variables in the order of the permutation, followed by updating the dual variables in a usual way. Obviously, all primal and dual variables are updated exactly once at each round. See Algorithm 1 for the details of RP-ADMM. Note that with a little abuse of notation, the function $\mathcal{L}(x_{\sigma(1)}, x_{\sigma(2)}, \dots, x_{\sigma(n)}; \mu)$ in this algorithm should be understood as $\mathcal{L}(x_1, x_2, \dots, x_n; \mu)$. For example, when $n = 3$ and $\sigma = (231)$, $\mathcal{L}(x_{\sigma(1)}, x_{\sigma(2)}, x_{\sigma(3)}; \mu) = \mathcal{L}(x_2, x_3, x_1; \mu)$ should be understood as $\mathcal{L}(x_1, x_2, x_3; \mu)$.

Algorithm 1 (n -Block Randomly Permuted ADMM (RP-ADMM))

Initialization: $x_i^0 \in \mathbb{R}^{d_i \times 1}$, $i = 1, \dots, n$; $\mu^0 \in \mathbb{R}^{N \times 1}$.

Round k ($k = 0, 1, 2, \dots$):

1. Primal update.

Pick a permutation σ of $\{1, \dots, n\}$ uniformly at random.

For $i = 1, \dots, n$, compute $x_{\sigma(i)}^{k+1}$ by

$$x_{\sigma(i)}^{k+1} = \arg \min_{x_{\sigma(i)} \in \mathcal{D}_{\sigma(i)}} \mathcal{L}(x_{\sigma(1)}^{k+1}, \dots, x_{\sigma(i-1)}^{k+1}, x_{\sigma(i)}, x_{\sigma(i+1)}^k, \dots, x_{\sigma(n)}^k; \mu^k). \quad (5)$$

2. Dual update. Update the dual variable by

$$\mu^{k+1} = \mu^k - \beta \left(\sum_{i=1}^n A_i x_i^{k+1} - b \right). \quad (6)$$

2.1.1. Optimization Formulation of Solving a Linear System of Equations. Consider a special case of (1) in which $f_i = 0$, $\mathcal{X}_i = \mathbb{R}^{d_i}$, $\forall i$, and $N = \sum_i d_i$ (i.e., the constraint is a square system of equations). Then problem (1) becomes

$$\begin{aligned} \min_{x \in \mathbb{R}^N} \quad & 0, \\ \text{s.t.} \quad & A_1 x_1 + \cdots + A_n x_n = b, \end{aligned} \quad (7)$$

where $A_i \in \mathbb{R}^{N \times d_i}$, $x_i \in \mathbb{R}^{d_i \times 1}$, $b \in \mathbb{R}^{N \times 1}$. Solving this feasibility problem (with 0 being the objective function) is equivalent to solving a linear system of equations

$$Ax = b, \quad (8)$$

where $A = [A_1, \dots, A_n] \in \mathbb{R}^{N \times N}$, $x = [x_1^T, \dots, x_n^T]^T \in \mathbb{R}^{N \times 1}$, $b \in \mathbb{R}^{N \times 1}$.

Throughout this paper, we assume A is nonsingular. Then the unique solution to (8) is $x = A^{-1}b$, and problem (7) has a unique primal-dual optimal solution $(x, \mu) = (A^{-1}b, 0)$. The augmented Lagrangian function (3) for the optimization problem (7) becomes

$$\mathcal{L}(x, \mu) = -\mu^T(Ax - b) + \frac{\beta}{2}\|Ax - b\|^2. \quad (9)$$

Throughout this paper, we assume $\beta = 1$; note that our algorithms and results can be extended to any $\beta > 0$ by simply scaling μ .

2.1.2. Example of Three-Block ADMM. Before presenting the update equation of general RP-ADMM for solving (7), we consider a simple case $N = n = 3$, $d_i = 1$, $\forall i$ and $\sigma = (123)$ and let $a_i = A_i \in \mathbb{R}^{3 \times 1}$. The update Equations (5) and (6) can be rewritten as

$$\begin{aligned} -a_1^T \mu^k + a_1^T(a_1 x_1^{k+1} + a_2 x_2^k + a_3 x_3^k - b) &= 0, \\ -a_2^T \mu^k + a_2^T(a_1 x_1^{k+1} + a_2 x_2^{k+1} + a_3 x_3^k - b) &= 0, \\ -a_3^T \mu^k + a_3^T(a_1 x_1^{k+1} + a_2 x_2^{k+1} + a_3 x_3^{k+1} - b) &= 0, \\ (a_1 x_1^{k+1})(+a_2 x_2^{k+1} + a_3 x_3^{k+1} - b) + \mu^{k+1} - \mu^k &= 0. \end{aligned}$$

Denote $y^k = [x_1^k; x_2^k; x_3^k; (\mu^k)^T] \in \mathbb{R}^{6 \times 1}$, then the above update equation becomes

$$\begin{bmatrix} a_1^T a_1 & 0 & 0 & 0 \\ a_2^T a_1 & a_2^T a_2 & 0 & 0 \\ a_3^T a_1 & a_3^T a_2 & a_3^T a_3 & 0 \\ a_1 & a_2 & a_3 & I_{3 \times 3} \end{bmatrix} y^{k+1} = \begin{bmatrix} 0 & -a_1^T a_2 & -a_1^T a_3 & a_1^T \\ 0 & 0 & -a_2^T a_3 & a_2^T \\ 0 & 0 & 0 & a_3^T \\ 0 & 0 & 0 & I_{3 \times 3} \end{bmatrix} y^k + \begin{bmatrix} A^T b \\ b \end{bmatrix}. \quad (10)$$

Define

$$L \triangleq \begin{bmatrix} a_1^T a_1 & 0 & 0 \\ a_2^T a_1 & a_2^T a_2 & 0 \\ a_3^T a_1 & a_3^T a_2 & a_3^T a_3 \end{bmatrix}, \quad R \triangleq \begin{bmatrix} 0 & -a_1^T a_2 & -a_1^T a_3 \\ 0 & 0 & -a_2^T a_3 \\ 0 & 0 & 0 \end{bmatrix}. \quad (11)$$

The relation between L and R is

$$L - R = A^T A.$$

Define

$$\bar{L} \triangleq \begin{bmatrix} L & 0 \\ A & I_{3 \times 3} \end{bmatrix}, \quad \bar{R} \triangleq \begin{bmatrix} R & A^T \\ 0 & I_{3 \times 3} \end{bmatrix}, \quad \bar{b} = \begin{bmatrix} A^T b \\ b \end{bmatrix} \quad (12)$$

then the update Equation (10) becomes $\bar{L} y^{k+1} = \bar{R} y^k + \bar{b}$, that is,

$$y^{k+1} = (\bar{L})^{-1} \bar{R} y^k + \bar{L}^{-1} \bar{b}. \quad (13)$$

As a side remark, Chen et al. [6] provide a specific example of $A \in \mathbb{R}^{3 \times 3}$ so that $\rho((\bar{L})^{-1}\bar{R}) > 1$, which implies the divergence of the above iteration if the update order $\sigma = (123)$ is used all the time. This counterexample disproves the convergence of cyclic three-block ADMM.

2.1.3. General Update Equation of RP-ADMM. In general, for the optimization problem (7), the primal update (5) becomes

$$-A_{\sigma(i)}^T \mu^k + A_{\sigma(i)}^T \left(\sum_{j=1}^i A_{\sigma(j)} x_{\sigma(j)}^{k+1} + \sum_{l=i+1}^n A_{\sigma(l)} x_{\sigma(l)}^k - b \right) = 0, \quad i = 1, \dots, n. \quad (14)$$

Replacing $\sigma(i), \sigma(j), \sigma(l)$ by i, j, l , we can rewrite the above equation as

$$-A_i^T \mu^k + A_i^T \left(\sum_{\sigma^{-1}(j) \leq \sigma^{-1}(i)} A_j x_j^{k+1} + \sum_{\sigma^{-1}(l) > \sigma^{-1}(i)} A_l x_l^k - b \right) = 0, \quad i = 1, \dots, n, \quad (15)$$

where σ^{-1} denotes the inverse mapping of a permutation σ , that is, $\sigma(i) = t \Leftrightarrow i = \sigma^{-1}(t)$. Denote the output of Algorithm 1 after round $(k-1)$ as

$$y^k \triangleq [x^k; \mu^k] = [x_1^k; \dots; x_n^k; \mu^k] \in \mathbb{R}^{2N \times 1}. \quad (16)$$

The update equations of Algorithm 1 for solving (7) (i.e., (15) and (6)) can be written in the matrix form as (when the permutation is σ and $\beta = 1$)

$$y^{k+1} = \bar{L}_\sigma^{-1} \bar{R}_\sigma y^k + \bar{L}_\sigma^{-1} \bar{b}, \quad (17)$$

where $\bar{L}_\sigma, \bar{R}_\sigma, L_\sigma, R_\sigma, \bar{b}$ are defined by

$$\bar{L}_\sigma \triangleq \begin{bmatrix} L_\sigma & 0 \\ A & I_{N \times N} \end{bmatrix}, \quad \bar{R}_\sigma \triangleq \begin{bmatrix} R_\sigma & A^T \\ 0 & I_{N \times N} \end{bmatrix}, \quad \bar{b} = \begin{bmatrix} A^T b \\ b \end{bmatrix}, \quad (18)$$

in which $L_\sigma \in \mathbb{R}^{N \times N}$ has $n \times n$ blocks and the (i, j) -th block is defined as

$$L_\sigma[i, j] \triangleq \begin{cases} A_i^T A_j & \sigma^{-1}(j) \leq \sigma^{-1}(i), \\ 0 & \text{otherwise,} \end{cases} \quad (19)$$

and R_σ is defined as

$$R_\sigma \triangleq L_\sigma - A^T A. \quad (20)$$

Another expression of L_σ , equivalent to (19), is the following:

$$L_\sigma[\sigma(i), \sigma(j)] \triangleq \begin{cases} A_{\sigma(i)}^T A_{\sigma(j)} & j \leq i, \\ 0 & j > i. \end{cases} \quad (21)$$

To illustrate the above expression of L_σ , we consider the n -coordinate case that $d_i = 1, \forall i$. In this case, each block x_i is a single coordinate, and each A_i is a vector. Denote $a_i \triangleq A_i \in \mathbb{R}^{N \times 1}$. Let $L_\sigma(k, l)$ denote the (k, l) -th entry of the matrix L_σ , then the definition (21) becomes

$$L_\sigma(\sigma(i), \sigma(j)) \triangleq \begin{cases} a_{\sigma(i)}^T a_{\sigma(j)} & j \leq i, \\ 0 & j > i. \end{cases} \quad (22)$$

A user-friendly rule for writing L_σ is described as follows (use $\sigma = (231)$ as an example). Start from a zero matrix. First, find all reverse pairs of σ ; here, we say (i, j) is a reverse pair if i appears after j in σ . For the permutation (231), all the reverse pairs are (1, 3), (3, 2), and (1, 2). Second, in the positions corresponding to the reverse pairs, write down the corresponding entries of $A^T A$ (i.e., $a_1^T a_3, a_3^T a_2$, and $a_1^T a_2$, respectively). At last, write $a_i^T a_i$ in the diagonal positions. Using this rule, we can write down the expression of $L_{(231)}$ as

$$L_{(231)} = \begin{bmatrix} a_1^T a_1 & a_1^T a_2 & a_1^T a_3 \\ 0 & a_2^T a_2 & 0 \\ 0 & a_3^T a_2 & a_3^T a_3 \end{bmatrix}.$$

A user-friendly rule to quickly check the correctness of an expression of L_σ is the following (still take $\sigma = (231)$ as an example). According to the order of the permutation (231), the second row, the third row, and the first

row should have a strictly decreasing number of zeros (two zeros, one zero, and no zero). In contrast, the second column, the third column, and the first column should have a strictly increasing number of zeros.

For the general case that $d_i \geq 1, \forall i$, we can write down the block partitioned L_σ in a similar way. For example, when $n = 3$ and $\sigma = (231)$, we have

$$L_{(231)} = \begin{bmatrix} A_1^T A_1 & A_1^T A_2 & A_1^T A_3 \\ 0 & A_2^T A_2 & 0 \\ 0 & A_3^T A_2 & A_3^T A_3 \end{bmatrix}.$$

2.2. Randomly Permuted BCD

RP-ADMM is a generalization of RP-BCD. In fact, when the constraint does not exist, RP-ADMM reduces to RP-BCD. In this subsection, we present RP-BCD for solving convex quadratic problems. Note that RP-ADMM for solving linear systems and RP-BCD for solving quadratic problems are two extremes of general RP-ADMM: in the former case the objective function is zero, and in the latter case the constraint is zero. Interestingly, the two extreme cases are related because the expected iteration matrix of RP-BCD appears as a component of the expected iteration matrix of RP-ADMM. We will show later that their eigenvalues are closely related.

Consider a special case of (1) in which $f(x) = \frac{1}{2} \|Ax - b\|^2$, $\mathcal{X}_i = \mathbb{R}^{d_i}, \forall i$, and there is no constraint. With abuse of notation, we use A to denote the coefficient matrix, whereas in the original formulation A denotes the constraint matrix. We “recycle” the notation A so that we can build a connection with RP-ADMM for solving linear systems later. Assume $N = \sum_i d_i$. Then problem (1) becomes a least-squares problem

$$\min_{x \in \mathbb{R}^N} \frac{1}{2} \|Ax - b\|^2 = \frac{1}{2} \|A_1 x_1 + \cdots + A_n x_n - b\|^2, \quad (23)$$

where $A_i \in \mathbb{R}^{N \times d_i}, x_i \in \mathbb{R}^{d_i \times 1}, b \in \mathbb{R}^{N \times 1}$. Similar to Section 2.1.1, we assume A is nonsingular. Then the unique solution to (8) is $x = A^{-1}b$.

In the augmented Lagrangian function given in (9), if we delete the first term, which depends on the dual variable μ , we obtain the quadratic function $\frac{1}{2} \|Ax - b\|^2$. Thus if we eliminate the dual variable μ in the update equations of RP-ADMM, we will obtain the update equations for RP-BCD. Suppose x^k is the iterate after the k -th epoch (i.e., go through all coordinates once), and σ is the order used in the k -th iteration; then, as a simpler version of (17), we have

$$x^{k+1} = L_\sigma^{-1} R_\sigma x^k + L_\sigma^{-1} b, \quad (24)$$

where L_σ and R_σ are defined as in (21) and (20), and σ is a random permutation.

2.3. Residual Trick for Efficient Implementation of ADMM and BCD

We note here that when $d_i = 1, \forall i$ (in this case BCD becomes CD), per-epoch computation time of ADMM and CD (no matter what order) is $O(n^2)$; or in other words, per-coordinate-update time is $O(n)$. For instance, updating x^{k+1} by (24) in RP-BCD or updating y^{k+1} by (17) in RP-ADMM only takes time $O(n^2)$. As mentioned in section 3.1 of Nesterov [30], the trick is to keep track of the residual. For both efficient practical implementation and calculation of computation complexity, one should use this residual trick, but for the ease of theoretical analysis we use the matrix update forms (17) and (24) in this paper; there is no contradiction because our theory only depends on the value of x^k but not the specific procedure to compute x^k .

For completeness, we briefly explain how this trick works in our settings. Suppose $d_i = 1, \forall i$, and we use CD methods to solve (23) with a certain update order (could be any order, such as cyclic, randomized, or randomly permuted). Suppose the coordinate i is picked, then x_i is updated by

$$x_i^+ = \frac{1}{A_i^T A_i} [A_i^T (b - A_{-i} x_{-i})], \quad (25)$$

where A_{-i} contains all columns of A except A_i , x_{-i} contains all elements of x except x_i and represents the current values, and x_i^+ represents the new value. A straightforward implementation of (25) requires multiplying x_{-i} by

A_{-i} , which takes $O(n^2)$ operations. With the residual trick (e.g., Nesterov [30]), we introduce the residual $r = Ax - b$ and replace (25) by

$$\begin{aligned} x_i^+ &= x_i - \frac{1}{A_i^T A_i} A_i^T r, \\ r^+ &= r + A_i(x_i^+ - x_i). \end{aligned}$$

Now the calculation of x_i^T and r^+ takes time $O(n)$, and thus one epoch of BCD takes time $O(n^2)$. The same trick can be applied to the primal update of ADMM; with this trick, the dual update (6) can be rewritten as $\mu^+ = \mu - \beta r$, which takes time $O(n)$, and thus one epoch of ADMM takes time $O(n^2)$.

Finally, when $d_i > 1$, similar update equations can still be used except a minor difference that $\frac{1}{A_i^T A_i}$ should be replaced by $(A_i^T A_i)^{-1}$. In a special case that $d_i = d, \forall i$ and $N = dn$, each iteration of BCD takes time $O(Nd + d^3)$ and each epoch takes time $O(N^2 + Nd^2)$. This cost can be reduced if we use BCGD (i.e., not solving the subproblem exactly but updating each block of variables by a gradient step). To not make the paper more complicated, we will not discuss the inexact versions of BCD and ADMM in this paper.

2.4. Two Versions of Independently Randomized ADMM

In this subsection, we present two other versions of randomized ADMM which can be divergent according to simulations. The failure of these versions makes us focus on analyzing RP-ADMM in this paper. These versions can be viewed as natural extensions of R-BCD (randomized BCD) (Leventhal and Lewis [24], Nesterov [30]).

In the first algorithm, called primal-dual randomized ADMM (PD-RADMM), the whole dual variable is viewed as the $(n+1)$ -th block. In particular, at each iteration, the algorithm draws one index i from $\{1, \dots, n, n+1\}$, then performs the following update: if $i \leq n$, update the i -th block of the primal variable; if $i = n+1$, update the whole dual variable. The details are given in Algorithm 2. We have tested PD-RADMM for the counter-example given in Chen et al. [6], and found that PD-RADMM always diverges (for random initial points).

A variant of PD-RADMM has been proposed in Hong et al. [20], with two differences: first, instead of minimizing the augmented Lagrangian $\mathcal{L}$, that algorithm minimizes a strongly convex upper bound of $\mathcal{L}$; second, that algorithm uses a diminishing dual stepsize. With these two modifications, Hong et al. [20] show that each limit point of the sequence generated by their algorithm is a primal-dual optimum with probability 1. Note that [20] also proves the same convergence result for the cyclic version of multiblock ADMM with these two modifications, thus it does not show the benefit of randomization.

Algorithm 2 (Primal-Dual Randomized ADMM (PD-RADMM))

```

Iteration  $t$  ( $t = 0, 1, 2, \dots$ ):
  Pick  $i \in \{1, \dots, n, n+1\}$  uniformly at random;
  If  $1 \leq i \leq n$ :
     $x_i^{t+1} = \arg \min_{x_i \in \mathcal{X}_i} \mathcal{L}(x_1^t, \dots, x_{i-1}^t, x_i, x_{i+1}^t, \dots, x_n^t; \mu^t)$ ,
     $x_j^{t+1} = x_j^t, \forall j \in \{1, \dots, n\} \setminus \{i\}$ ,
     $\mu^{t+1} = \mu^t$ .
  Else if  $i = n+1$ :
     $\mu^{t+1} = \mu^t - \beta(\sum_{i=1}^n A_i x_i^{t+1} - b)$ ,
     $x_j^{t+1} = x_j^t, \forall j \in \{1, \dots, n\}$ .
End

```

In the second algorithm, called primal randomized ADMM (P-RADMM), we only perform randomization for the primal variables. In particular, at each round, we first draw n independent random variables $j_1, \dots, j_n$ from the uniform distribution of $\{1, \dots, n\}$ and update $x_{j_1}, \dots, x_{j_n}$ sequentially, then update the dual variable in the usual way. The details are given in Algorithm 3. This algorithm looks quite similar to RP-ADMM because they both update n primal blocks at each round; the difference is that RP-ADMM samples *without replacement*, whereas this algorithm P-RADMM samples *with replacement*. In other words, RP-ADMM updates each block exactly once at each round, whereas P-RADMM may update one block more than one time or does not update one block at each round.

We have tested P-RADMM in various settings. For the counter-example given in Chen et al. [6], we found that P-RADMM does converge. However, if $n \geq 30$ and A is a Gaussian random matrix (each entry is drawn

independent and identically distributed from $\mathcal{N}(0, 1)$, then P-RADMM diverges in almost all cases we have tested. This phenomenon is rather strange because for random Gaussian matrices A the cyclic ADMM actually converges (according to simulations). An implication is that randomized versions do not always outperform their deterministic counterparts in terms of convergence.

Because both Algorithms 2 and 3 can diverge in certain cases, we will not further study them in this paper. In the rest of the paper, we will focus on RP-ADMM (i.e., Algorithm 1).

Algorithm 3 (Primal Randomized ADMM (P-RADMM))

Round k ($k = 0, 1, 2, \dots$):

1. Primal update.

Pick $l_1, \dots, l_n$ independently from the uniform distribution of $\{1, \dots, n\}$.

For $i = 1, \dots, n$:

$t = kn + i - 1$,

$x_{l_i}^{t+1} = \arg \min_{x_{l_i} \in \mathcal{X}_{l_i}} \mathcal{L}(x_1^t, \dots, x_{l_i-1}^t, x_{l_i}, x_{l_i+1}^t, \dots, x_n^t; \mu^t)$,

$x_j^{t+1} = x_j^t, \forall j \in \{1, \dots, n\} \setminus \{l_i\}$,

$\mu^{t+1} = \mu^t$.

End.

2. Dual update.

$\mu^{(k+1)n} = \mu^{kn} - \beta \left(\sum_{i=1}^n A_i x_i^{(k+1)n} - b \right)$.

2.5. Bernoulli-Randomized ADMM

To implement randomly permuted ADMM, one needs to sample from all blocks without replacement. To save the sampling time, we propose another algorithm, which we call Bernoulli-randomized ADMM. This algorithm is motivated by the proof of Theorem 1. This updating scheme can be applied to other algorithms, such as SGD and coordinate descent methods.

The new update order combines the well-known double-sweep order and Bernoulli-randomization. The original double-sweep order is $(1, 2, \dots, n-1, n, n-1, n-2, \dots, 1)$, meaning that $x_1, x_2, \dots, x_{n-1}, x_n, x_{n-1}, x_{n-2}, \dots, x_1$ are updated sequentially in each “cycle.” It combines the normal cyclic order $(1, 2, \dots, n)$ and a reverse order $(n, n-1, \dots, 1)$. We propose the following updating scheme: add a check box to each block, and in each cycle we perform the following operations.

1. Phase I: go through the blocks $x_1, x_2, \dots, x_n$ one by one sequentially as follows: for each block x_i , flip a fair coin and:

- if the outcome is “heads,” update the block x_i and check the check box;
- if the outcome is “tails,” do nothing about x_i and uncheck the check box.

		1	2	3	4	5
Phase I	begin	skip →	skip →	3 →	skip →	5
						↓
Phase II	end	1	← 2	← skip	← 4	← skip

2. Phase II: go through the blocks $x_n, x_{n-1}, \dots, x_1$ in the reverse order, and update x_i if the box is unchecked.

Note that in each cycle we go through each block twice but update each block exactly once so that the number of totally updated blocks remains n . For example, when $n = 5$, (35421) is a possible update order, as shown in the following diagram. Similarly, (13542) is also a possible update order.

But (13524) and (35412) are not possible. The set of all possible update orders is given by

$$\Gamma_{\text{BR}} \triangleq \{\sigma \in \Gamma \mid \exists i \in \{1, \dots, n-1\} \text{ such that } \sigma(1) < \sigma(2) < \dots < \sigma(i) \text{ and } \sigma(i+1) > \dots > \sigma(n)\},$$

where Γ is the set of permutations of $\{1, 2, \dots, n\}$ as defined in (4). In other words, a sequence from Γ_{BR} is a concatenation of an increasing sequence and a decreasing sequence. Note that the permutation $(1, 2, \dots, n)$ is in Γ_{BR} because it can be viewed as the concatenation of an increasing sequence $(1, 2, \dots, n-1)$ and a “decreasing sequence” (n) , and we can let $i = n-1$ in the above definition to cover this case. Similarly, the permutation $(n, n-1, \dots, 1)$ is also in Γ_{BR} because $i = 1$ will cover this case.

The algorithm Bernoulli-randomized ADMM (BR-ADMM) is formally described below. We skip the epoch index k because otherwise the notation would be cumbersome.

Algorithm 4 (n -Block Bernoulli-Randomized ADMM (BR-ADMM))

Initialization: $x_i^0 \in \mathbb{R}^{d_i \times 1}, i = 1, \dots, n; \mu^0 \in \mathbb{R}^{N \times 1}$.

Round k ($k = 0, 1, 2, \dots$):

1. Primal update.

Set $c_i = 0, i = 1, \dots, n$.

Phase I.

For $i = 1, 2, \dots, n$:

Draw a random variable $\xi \sim \text{Bernoulli}(1/2)$, that is, $\Pr(\xi = 1) = \Pr(\xi = 0) = 1/2$.

If $\xi = 1$: set $c_i = 1$ and update x_i by

$$x_i \leftarrow \arg \min_{x_i \in \mathcal{X}_i} \mathcal{L}(x_1, \dots, x_{i-1}, x_i, x_{i+1}, \dots, x_n; \mu). \quad (26)$$

Phase II.

For $i = n, n-1, \dots, 1$: if $c_i = 0$, update x_i by (26).

2. Dual update. Update the dual variable by

$$\mu \leftarrow \mu - \beta \left(\sum_{i=1}^n A_i x_i - b \right). \quad (27)$$

For solving linear systems of equations, the update formula is the same as (17), the update formula of RP-ADMM. The difference is that for RP-ADMM σ can be an arbitrary permutation, whereas for BR-ADMM there is some restriction on σ : it has to be a permutation in Γ_{BR} .

3. Main Results

3.1. Expected Convergence of RP-ADMM

Let σ_i denote the permutation used in round i of Algorithm 1, which is a uniform random variable drawn from the set of permutations Γ . After round k , Algorithm 1 generates a random output y^{k+1} , which depends on the observed draw of the random variable

$$\xi_k = (\sigma_0, \sigma_1, \dots, \sigma_k). \quad (28)$$

We will show that the expected iterate (the iterate y^k is defined in (16))

$$\phi^k = E_{\xi_{k-1}}(y^k) \quad (29)$$

converges to the primal-dual solution of the problem (7). Although the expected convergence does not necessarily imply the convergence in a particular realization, it serves as an evidence of convergence. Our proof seems much different from and more difficult than previous proofs for other randomized methods, because random permutation and spectral radius of nonsymmetric matrices are difficult objects to deal with—not many existing mathematical tools are available to help.³ Note that the extension of this result to the nonsquare full column-rank case is simple.⁴

Theorem 1. Assume the coefficient matrix $A = [A_1, \dots, A_n]$ of the constraint in (7) is a nonsingular square matrix. Suppose Algorithm 1 is used to solve problem (7); then the expected output converges to the unique primal-dual optimal solution to (7), that is,

$$\{\phi^k\}_{k \rightarrow \infty} \rightarrow \begin{bmatrix} A^{-1}b \\ 0 \end{bmatrix}. \quad (30)$$

Because the update matrix does not depend on previous iterates, we claim (and prove in Section 4.1) that Theorem 1 holds if the expected update matrix has a spectral radius less than 1 (i.e., if the following Theorem 2 holds).

Theorem 2. Suppose $A = [A_1, \dots, A_n] \in \mathbb{R}^{N \times N}$ is nonsingular, and $\bar{L}_\sigma^{-1}, \bar{R}_\sigma$ are defined by (18) for any permutation σ . Define

$$M \triangleq E_\sigma(\bar{L}_\sigma^{-1} \bar{R}_\sigma) = \frac{1}{n!} \sum_{\sigma \in \Gamma} (\bar{L}_\sigma^{-1} \bar{R}_\sigma), \quad (31)$$

where the expectation is taken over the uniform random distribution over Γ , the set of permutations of $\{1, 2, \dots, n\}$. Then the spectral radius of M is smaller than 1, that is,

$$\rho(M) < 1. \quad (32)$$

Remark 1. For the counterexample in Chen et al. [6] in which $A = [1, 1, 1; 1, 1, 2; 1, 2, 2]$, it is easy to verify that $\rho(M_\sigma) > 1.02$ for any permutation σ of $(1, 2, 3)$. Interestingly, Theorem 2 shows that even if each M_σ is “bad” (with spectral radius larger than 1), the average of them is always “good” (with spectral radius smaller than 1).

Theorem 2 is just a linear algebra result and can be understood even without knowing the details of the algorithm. However, the proof of Theorem 2 is rather nontrivial. This proof will be provided in Section 4.2, and the technical results used in this proof will be proved in Sections 5 and 6.

The convergence rate of RP-ADMM for solving linear systems of equations is closely related to the convergence rate of RP-BCD (randomly permuted BCD) for solving quadratic problems. We will discuss their relation and how our results in this paper improve our understanding for RP-BCD.

A similar convergence result holds for BR-ADMM proposed in Section 2.5, as presented below. The proof is a simple modification of the proof of Theorem 1 and can be found in Section 6.4.

Proposition 1. Assume the coefficient matrix $A = [A_1, \dots, A_n]$ of the constraint in (7) is a nonsingular square matrix. Suppose Algorithm 4 is used to solve problem (7); then the expected output converges to the unique primal-dual optimal solution to (7).

3.2. Expected Convergence Rate of RP-ADMM and RP-BCD

There is a close relation between RP-ADMM for solving linear systems and RP-CD for solving quadratic problems (see Lemma 2). Thus it is not surprising that we need to understand RP-BCD before understanding RP-ADMM. We will first present an expected convergence rate of RP-BCD (in terms of the expected iterates) for solving quadratic problems, which improves the best existing convergence rate (one type of rates, to be precise) by a factor of n .⁵ The result is proved via establishing a weak version of matrix AM-GM inequality. This result also establishes a large gap of $O(n)$ between RP-BCD and C-BCD (cyclic BCD). Second, built upon the result for RP-BCD, we establish a convergence rate of RP-ADMM that is similar to RP-BCD and also n times better than that of C-BCD.

The first result is about the expected convergence rate of RP-BCD for the case $A_i^T A_i = I$. This assumption is made so that the expression is simple, and the case for general A_i is given in the next result.

Theorem 3 (Rate of RP-BCD for Quadratic Functions with Identity Diagonal Blocks). Assume the coefficient matrix $A = [A_1, \dots, A_n]$ is a nonsingular square matrix, and $A_i^T A_i = I, \forall i$. Suppose RP-BCD is used to solve problem (23), where x^k denotes the variable after k epochs (each epoch represents one cycle of updating all coordinates). Denote the unique optimal solution as $x^* = A^{-1}b$. Then

$$\|E(x^k) - x^*\| \leq \max \left\{ 1 - \frac{1}{n} \lambda_{\min}(AA^T), \frac{1}{3} \right\}^k \|x^0 - x^*\|. \quad (33)$$

To put this convergence rate result in the context, we consider the simple case that each $d_i = 1$ (i.e., each block consists of a single coordinate). In this case, every diagonal entry of $A^T A$ is 1; thus, the average eigenvalue of $A^T A$ is 1. Throughout the paper, we consider the total computation complexity⁶; note that we assume the residual trick as described in Section 2.3 is always used for all methods.

Our Theorem 3 provides an expected computational complexity upper bound $O(n^3 \kappa_{CD} \log \frac{1}{\epsilon})$ for RP-CD, because each epoch takes $O(n^2)$ time and it requires $O(n^2 \log \frac{1}{\epsilon})$ epochs to achieve error ϵ according to (33). It is known that the computational complexity of R-CD (randomized coordinate descent) to achieve relative accuracy ϵ ⁷ is $O(n^2 \kappa_{CD} \log \frac{1}{\epsilon})$, where $\kappa_{CD} = \lambda_{\text{avg}}(A^T A) / \lambda_{\min}(A^T A) = 1 / \lambda_{\min}(A^T A)$ is the ratio of the average eigenvalue over the minimum eigenvalue. It was recently shown that in terms of κ_{CD} and n only, the worst-case complexity of C-CD (cyclic CD) is $O(n^4 \kappa_{CD} \log \frac{1}{\epsilon})$, which is n^2 times worse than R-CD and n times worse than GD. This shows a large gap between C-CD and R-CD in the worst case.

It was widely conjectured that RP-CD is at least as fast as R-CD, but this conjecture is considered to be rather difficult to prove. For a special class of matrices, recent works (Lee and Wright [23], Wright and Lee [45]) validated the conjecture. However, to our knowledge, even for a general quadratic function with equal

diagonal entries 1, the previously best known convergence rate of RP-CD is almost the same as C-CD (see Sun and Hong [39], Sun and Ye [41]), which can be n^2 times worse than that of R-CD. Our Theorem 3 provides an expected computational complexity upper bound $O(n^3 \kappa_{\text{CD}} \log \frac{1}{\epsilon})$ for RP-CD, which is n times faster than C-CD and n times slower than R-CD. This improves the best existing rate by a factor of n .⁸ We summarize the comparison of the complexity for C-CD, R-CD, and RP-CD in Table 1.

The following proposition generalizes Theorem 3 to the non-identity-diagonal case (i.e., $A_i^T A_i$ does not need to be an identity matrix).

Proposition 2 (Rate of RP-BCD for Quadratic Functions, with Nonidentity Blocks). *Assume the coefficient matrix $A = [A_1, \dots, A_n]$ is a nonsingular square matrix. Suppose RP-BCD is used to solve problem (23). Denote $D = \text{diag}(A_1^T A_1, \dots, A_n^T A_n)$ as a block-diagonal matrix, and the norm $\|z\|_D = \sqrt{z^T D z}$. Then*

$$\|E(x^k) - x^*\|_D \leq \max \left\{ 1 - \frac{1}{n} \lambda_{\min}(D^{1/2} A^T A D^{-1/2}), \frac{1}{3} \right\}^k \|x^0 - x^*\|_D. \quad (34)$$

The proof of Proposition 2 is given in Section 4.4. One can easily transform the quantity $\lambda_{\min}(D^{1/2} A^T A D^{-1/2})$ to a certain quantity that only depends on the eigenvalues of $A_i^T A_i$ and $A^T A$. However, as noted in [41], it is far from clear how tight the transformation is, thus we skip the transformation here. In fact, it is related to some open question on the so-called Jacobi preconditioning. We refer the interested readers to [41] for a detailed discussion of the subtle issues in the non-identity-diagonal case.

At last, we present a result on the expected convergence rate of RP-ADMM for solving linear systems, under the assumption that $A_i^T A_i = I$, $\forall i$. Very similar to Proposition 2, we can also generalize this result to the non-identity-diagonal case (i.e., $A_i^T A_i \neq I$), but to save space we skip the generalization here. The proof of Theorem 4 is given in Section 4.5.

Theorem 4 (Expected Convergence Rate of RP-ADMM for Linear Systems). *Assume the coefficient matrix $A = [A_1, \dots, A_n]$ of the constraint in (7) is a nonsingular square matrix and $A_i^T A_i = I_{d_i}$. Suppose Algorithm 1 is used to solve problem (7). Denote $y^* = [A_0^{-1} b]$ as the unique primal-dual optimal solution to the problem (7), then*

$$\|E(y^k) - y^*\| \leq \left(1 - \frac{1}{2n} \lambda_{\min}(A A^T) \right)^k \|y^0 - y^*\|. \quad (35)$$

This result implies that similar to RP-CD for solving quadratic problems, the complexity of RP-ADMM in terms of the expected iterates for solving linear systems is also at most

$$T_{\text{RP-ADMM}} = O(n^3 \kappa_{\text{CD}} \log(1/\epsilon)).$$

In light of the fact that C-CD has been shown to only achieve a rate $O(n^4 \kappa_{\text{CD}} \log(1/\epsilon))$ [41], the rate of RP-ADMM we obtain is already quite good. Nevertheless, we conjecture that this complexity upper bound can be improved to $O(n^2 \kappa_{\text{CD}} \log(1/\epsilon))$, the same as the conjectured complexity for RP-CD. However, an improved rate of RP-ADMM leads to an improved rate of RP-BCD (this should be clear via the comparison of (50) and (57)), thus proving this conjecture is an even more difficult problem than the long-standing open question on RP-CD.

3.3. Matrix AM-GM Inequality

To analyze the convergence rate of randomly permuted algorithms, one major technical challenge is matrix AM-GM (algebraic mean-geometric mean) inequality. The following conjecture of matrix AM-GM inequality was proposed in Recht and Ré [31]: for any positive semidefinite matrix $A_1, \dots, A_n \in \mathbb{R}^{n \times n}$,

$$\left\| \frac{1}{n!} \sum_{\sigma=(\sigma_1, \dots, \sigma_n) \in \Gamma} A_{\sigma_n} A_{\sigma_{n-1}} \dots A_{\sigma_1} \right\| \leq \left\| \left(\frac{1}{n} \sum_i A_i \right)^n \right\|. \quad (36)$$

The original version is more general: the number of matrices does not need to be the same as the dimension of the matrix. For simplicity, we just present a simpler version here.

Table 1. Worst-Case Computation Complexity Comparison, Using Only κ_{CD} as Parameter, for Equal-Diagonal Quadratic Case (Ignore $O(\log \frac{1}{\epsilon})$ Factor) and Considering the Error in the Expected Iterates for RP-CD

	GD	C-CD	R-CD	RP-CD (Theorem 3)	RP-CD (conjectured)
Computation complexity	$n^3 \kappa_{\text{CD}}$	$n^4 \kappa_{\text{CD}}$	$n^2 \kappa_{\text{CD}}$	$n^3 \kappa_{\text{CD}}$	$n^2 \kappa_{\text{CD}}$

The matrix AM-GM inequality is a generalization of the well-known AM-GM inequality: for nonnegative numbers $a_1, \dots, a_n$, the geometric mean $(a_1 a_2 \dots a_n)^{1/n}$ is no more than the algebraic mean $\frac{1}{n} \sum_{i=1}^n a_i$. When extending this inequality to matrix domain, the noncommutative nature of matrix multiplication makes the problem rather difficult to prove.

We observe that we only need to prove a matrix AM-GM inequality for projection matrices. We conjecture that the following matrix AM-GM inequality holds.

Conjecture 1 (Matrix AM-GM Inequality for Projection Matrices). *Suppose $P_i \in \mathbb{R}^{N \times N}$, $i = 1, \dots, n$ are projection matrices, then*

$$\frac{1}{n!} \sum_{\sigma=(\sigma_1, \dots, \sigma_n) \in \Gamma} P_{\sigma_n} P_{\sigma_{n-1}} \dots P_{\sigma_1} \leq \left(\frac{1}{n} \sum_i P_i \right)^n. \quad (37)$$

Compared with (36), our conjecture makes a stronger claim on the relation, but it only applies to projection matrices. We have found examples to show that (37) does not hold for general positive semidefinite matrices, but it holds for projection matrices in all of our experiments.

We are not able to prove the new conjecture—that would solve the open question of the best convergence rate of RP-CD for quadratic problem. Nevertheless, inspired by the new conjecture, we prove a weaker version (see Lemma 3), which can lead to an improved convergence rate estimate for RP-CD.

4. Proof of Main Results

4.1. Proof of Theorem 1

Denote σ_k as the permutation used in round k , and define ξ_k as in (28). Rewrite the update Equation (17) below (replacing σ by σ_k):

$$y^{k+1} = \bar{L}_{\sigma_k}^{-1} \bar{R}_{\sigma_k} y^k + \bar{L}_{\sigma_k}^{-1} \bar{b}. \quad (38)$$

We first prove (30) for the case $b = 0$. By (18) we have $\bar{b} = 0$, then (38) is simplified to $y^{k+1} = \bar{L}_{\sigma_k}^{-1} \bar{R}_{\sigma_k} y^k$. Taking the expectation of both sides of this equation in ξ_k (see its definition in (28)), and noting that y^k is independent of σ_k , we get

$$\phi^{k+1} = E_{\xi_k} \left(\bar{L}_{\sigma_k}^{-1} \bar{R}_{\sigma_k} y^k \right) = E_{\sigma_k} \left(E_{\xi_{k-1}} \left(\bar{L}_{\sigma_k}^{-1} \bar{R}_{\sigma_k} y^k \right) \right) = E_{\sigma_k} \left(\bar{L}_{\sigma_k}^{-1} \bar{R}_{\sigma_k} \phi^k \right) = M \phi^k.$$

Because the spectral radius of M is less than 1 by Theorem 2, we have that $\{\phi^k\} \rightarrow 0$, that is, (30).

We then prove (30) for general b . Let $y^* = [A^{-1}b; 0]$ denote the optimal solution. Then it is easy to verify that

$$y^* = \bar{L}_{\sigma_k}^{-1} \bar{R}_{\sigma_k} y^* + \bar{L}_{\sigma_k}^{-1} \bar{b}$$

for all $\sigma_k \in \Gamma$ (i.e., the optimal solution is the fixed point of the update equation for any order). Computing the difference between this equation and (38) and letting $\hat{y}^k = y^k - y^*$, we get $\hat{y}^{k+1} = \bar{L}_{\sigma_k}^{-1} \bar{R}_{\sigma_k} \hat{y}^k$. According to the proof for the case $b = 0$, we have $E(\hat{y}^k) \rightarrow 0$, which implies $E(y^k) \rightarrow y^*$.

4.2. Proof of Theorem 2

The difficulty of proving Theorem 2 (bounding the spectral radius of M defined in (31)) is two-fold. First, M is a nonsymmetric matrix, and there are very few tools to bound the spectral radius of a nonsymmetric matrix. In fact, spectral radius is neither subadditive nor submultiplicative (see, e.g., Kittaneh [22]). Note that the spectral norm of M can be much larger than 1 (there are examples that $\|M\| > 2$), thus we cannot bound the spectral radius simply by the spectral norm. Second, although it is possible to explicitly write each entry of M as a function of the entries of $A^T A$, these functions are very complicated (n -th order polynomials), and it is not clear how to utilize this explicit expression.

The proof outline of Theorem 2 and the main techniques are described below. In Step 0, we provide an expression of the expected update matrix M . In Step 1, we establish the relationship between the eigenvalues of M and the eigenvalues of a simple symmetric matrix AQA^T , where Q is defined in (39). As a consequence, the spectral radius of M is smaller than one if and only if (iff) the eigenvalues of AQA^T lie in the region $(0, 4/3)$. This step partially resolves the first difficulty (i.e., how to deal with the spectral radius of a nonsymmetric matrix). In Step 2, we show that the eigenvalues of AQA^T do lie in $(0, 4/3)$ using mathematical induction. The induction analysis circumvents the second difficulty (i.e., how to utilize the relation between M and A).

Step 0. Compute the Expression of the Expected Update Matrix M .

Define

$$Q \triangleq E_\sigma(L_\sigma^{-1}) = \frac{1}{n!} \sum_{\sigma \in \Gamma} L_\sigma^{-1}. \quad (39)$$

It is easy to prove that Q defined by (39) is symmetric. In fact, note that $L_\sigma^T = L_{\bar{\sigma}}, \forall \sigma \in \Gamma$, where $\bar{\sigma}$ is a reverse permutation of σ satisfying $\bar{\sigma}(i) = \sigma(n+1-i), \forall i$, thus $Q = \frac{1}{n!} \sum_{\sigma} Q_\sigma = (\frac{1}{n!} \sum_{\sigma} Q_{\bar{\sigma}})^T = Q^T$, where the last step is because the sum of all $Q_{\bar{\sigma}}$ is the same as the sum of all Q_σ .

Denote

$$M_\sigma \triangleq \bar{L}_\sigma^{-1} \bar{R}_\sigma = \bar{L}_\sigma^{-1} \begin{bmatrix} R_\sigma & A^T \\ 0 & I \end{bmatrix}. \quad (40)$$

Substituting the expression of $\bar{L}_\sigma^{-1}$ into the above relation, and replacing R_σ by $L_\sigma - A^T A$, we obtain

$$M_\sigma = \begin{bmatrix} L_\sigma^{-1} & 0 \\ -AL_\sigma^{-1} & I \end{bmatrix} \begin{bmatrix} L_\sigma - A^T A & A^T \\ 0 & I \end{bmatrix} = \begin{bmatrix} I - L_\sigma^{-1} A^T A & L_\sigma^{-1} A^T \\ -A + AL_\sigma^{-1} A^T A & I - AL_\sigma^{-1} A^T \end{bmatrix}. \quad (41)$$

Because M_σ is linear in L_σ^{-1} , we have

$$\begin{aligned} M = E_\sigma(M_\sigma) &= \begin{bmatrix} I - E_\sigma(L_\sigma^{-1}) A^T A & E_\sigma(L_\sigma^{-1}) A^T \\ -A + AE_\sigma(L_\sigma^{-1}) A^T A & I - AE_\sigma(L_\sigma^{-1}) A^T \end{bmatrix} \\ &= \begin{bmatrix} I - QA^T A & QA^T \\ -A + AQA^T A & I - AQA^T \end{bmatrix}. \end{aligned} \quad (42)$$

Step 1. Relate M to a Simple Symmetric Matrix.

The main result of Step 1 is given below, and the proof of this result is relegated to Section 5.

Lemma 1. Suppose $A \in \mathbb{R}^{N \times N}$ is nonsingular and $Q \in \mathbb{R}^{N \times N}$ is an arbitrary matrix. Define $M \in \mathbb{R}^{2N \times 2N}$ as

$$M = \begin{bmatrix} I - QA^T A & QA^T \\ -A + AQA^T A & I - AQA^T \end{bmatrix}. \quad (43)$$

Then

$$\lambda \in \text{eig}(M) \iff \frac{(1-\lambda)^2}{1-2\lambda} \in \text{eig}(QA^T A). \quad (44)$$

Furthermore, when Q is symmetric, we have

$$\rho(M) < 1 \iff \text{eig}(QA^T A) \subseteq \left(0, \frac{4}{3}\right). \quad (45)$$

Remark 2. For our problem, the matrix Q as defined by (39) is symmetric (see the argument after Equation (39)), thus the relation (45) indeed holds according to Lemma 1. For a general nonsymmetric Q , (45) does not need to hold, but the first conclusion (44) still holds.

Step 2. Bound the Eigenvalues of $QA^T A$.

The main result of Step 2 is summarized in the following Lemma 2. The proof of Lemma 2 is given in Section 6.

Lemma 2. Suppose $A = [A_1, \dots, A_n] \in \mathbb{R}^{N \times N}$ is nonsingular. Define Q as

$$Q \triangleq E_\sigma(L_\sigma^{-1}) = \frac{1}{n!} \sum_{\sigma \in \Gamma} L_\sigma^{-1}, \quad (46)$$

in which L_σ is defined by (21) and Γ is defined by (4). Then all eigenvalues of $QA^T A$ lie in $(0, 4/3)$, that is

$$\text{eig}(QA^T A) \subseteq \left(0, \frac{4}{3}\right). \quad (47)$$

Remark 3. The upper bound $\frac{4}{3}$ in (47) is probably tight, because we have found numerical examples with $\text{eig}(QA^TA) > 1.3333$. Now the expected convergence of RP-ADMM seems to be a pleasant coincidence: Lemma 1 shows that to prove the expected convergence we need to prove $\sup_A \text{eig}(QA^TA)$, a quantity that can be defined without knowing ADMM, is bounded by $4/3$; Lemma 2 and numerical experiments show that this quantity happens to be exactly $4/3$, so that RP-ADMM can converge (in expectation).

Theorem 2 follows immediately from Lemmas 1 and 2.

4.3. Proof of Theorem 3

We first describe the outline of the proof. The expected update matrix of RP-BCD is $I - QA^TA$, and the eigenvalues of this matrix lie in $(-1, 1)$. The expected convergence speed of RP-BCD depends on the distance between the eigenvalues and the two extremes -1 and 1 . Lemma 2 shows that the distance to -1 is at least $1/3$, which is a constant. We will show that the distance to 1 is at least $\lambda_{\min}(A^TA)/n$, by proving a weaker version of matrix AM-GM inequality. Combining the two results, we obtain the expected convergence speed of RP-BCD.

The formal proof is presented below.

According to (24), we have $x^{k+1} - x^* = (I - L_\sigma^{-1}A^TA)(x^k - x^*)$, where σ is the randomly picked permutation at the k -th epoch. Therefore, the expected update formula of RP-BCD for solving the least squares problem is

$$E(x^{k+1}) - x^* = (I - QA^TA)(E(x^k) - x^*). \quad (48)$$

It implies

$$\|E(x^{k+1}) - x^*\| \leq \rho(I - QA^TA)\|E(x^k) - x^*\|. \quad (49)$$

Suppose the eigenvalues of QA^TA are $\eta_1 \geq \eta_2 \geq \dots \geq \eta_n$; then, according to Lemma 2,

$$4/3 > \eta_1 > \dots > \eta_n > 0.$$

The eigenvalues of $I - QA^TA$ are

$$-\frac{1}{3} < 1 - \eta_1 \leq \dots \leq 1 - \eta_n < 1,$$

thus the spectral radius of $I - QA^TA$ is

$$\rho(I - QA^TA) = \max\{1 - \eta_n, |1 - \eta_1|\} \leq \max\left\{1 - \eta_n, \frac{1}{3}\right\} = \max\left\{\lambda_{\max}(I - QA^TA), \frac{1}{3}\right\}. \quad (50)$$

An interesting phenomenon occurs here. The spectral radius is either $1 - \eta_n$ or $|1 - \eta_1|$. In the latter case, $\rho(I - QA^TA) = |1 - \eta_1| \leq 1/3$, implying that $\|E(x^k) - x^*\| \leq \frac{1}{3^k} \|E(x^0) - x^*\|$, or equivalently, the relative error $\|E(x^k) - x^*\|/\|E(x^0) - x^*\|$ achieves ϵ in $\log 3 \log(1/\epsilon)$ epochs. We do not even need to compute η_1 because it will only affect the convergence speed when the speed is already very fast. From a theoretical perspective, the improvement from $\log 3$ to $\log(1/(1 - |1 - \eta_1|))$ is just an improvement in the constant. Therefore, it is reasonable to ignore η_1 and focus on the estimate of $1 - \eta_n$.

To estimate the maximum eigenvalue of $I - QA^TA$ (or equivalently, that of $I - AQA^T$), we first provide a useful identity that connects $I - AQA^T$ and projection matrices $P_i = I - A_iA_i^T$.

Claim 1. Suppose $A = [A_1, \dots, A_n]$ is a nonsingular square matrix, and $A_i^TA_i = I, \forall i$. For a permutation $\sigma = (\sigma_1, \dots, \sigma_n) \in \Gamma$, L_σ is defined as in (21), and $Q_\sigma = L_\sigma^{-1}$. Denote $P_i = I - A_iA_i^T, i = 1, \dots, n$. Then we have

$$I - AQ_\sigma A^T = P_{\sigma_n} P_{\sigma_{n-1}} \dots P_{\sigma_1}, \quad (51a)$$

$$I - AQA^T = \frac{1}{n!} \sum_{\sigma=(\sigma_1, \dots, \sigma_n) \in \Gamma} P_{\sigma_n} P_{\sigma_{n-1}} \dots P_{\sigma_1}. \quad (51b)$$

The proof of Claim 1 is given at the end of this subsection. Claim 1 states that $I - AQA^T$ is exactly equal to $\frac{1}{n!} \sum_{\sigma=(\sigma_1, \dots, \sigma_n) \in \Gamma} P_{\sigma_n} P_{\sigma_{n-1}} \dots P_{\sigma_1}$, thus we only need to estimate the maximal eigenvalue of the latter expression. This is achieved by the following lemma (the proof is given in Section 7.3).

Lemma 3 (Weak Matrix AM-GM Inequality). *Suppose $P_i \in \mathbb{R}^{N \times N}$, $i = 1, \dots, n$ are projection matrices, then*

$$\frac{1}{n!} \sum_{\sigma=(\sigma_1, \dots, \sigma_n) \in \Gamma} P_{\sigma_n} P_{\sigma_{n-1}} \dots P_{\sigma_1} \leq \frac{1}{n} \sum_i P_i. \quad (52)$$

The above Lemma 3 and Claim 4.1 immediately lead to the following corollary.

Corollary 1. *Suppose $A = [A_1, \dots, A_n]$ is a nonsingular square matrix, and $A_i^T A_i = I$, $\forall i$. Suppose $P_i = I - A_i A_i^T$, $\forall i$. L_σ is defined as in (21), and $Q = E_\sigma(L_\sigma^{-1})$. Then*

$$I - AQA^T \leq \frac{1}{n} \sum_i P_i. \quad (53)$$

Note that $\frac{1}{n} \sum_i P_i = \frac{1}{n} (nI - \sum_i A_i A_i^T) = I - \frac{1}{n} AA^T$, thus (53) implies

$$I - AQA^T \leq I - \frac{1}{n} AA^T,$$

which implies

$$\lambda_{\max}(I - AQA^T) \leq 1 - \frac{1}{n} \lambda_{\min}(AA^T). \quad (54)$$

Substituting into (50), we get

$$\rho(I - QA^T A) \leq \max \left\{ \lambda_{\max}(I - QA^T A), \frac{1}{3} \right\} \leq \max \left\{ 1 - \frac{1}{n} \lambda_{\min}(AA^T), 1/3 \right\}.$$

Substituting this relation into (49), we obtain

$$\|E(x^{k+1}) - x^*\| \leq \max \left\{ 1 - \frac{1}{n} \lambda_{\min}(AA^T), \frac{1}{3} \right\} \|E(x^k) - x^*\|. \quad \text{Q.E.D.}$$

Remark 4. There is a coefficient $1/n$ in front of $\lambda_{\min}(AA^T)$ in (54), and this is why the complexity of RP-CD we establish is n times worse than the conjectured one in Table 1. If Conjecture 1 holds, then this factor of $1/n$ would be removed and the conjectured (expected) complexity of RP-CD in Table 1 would hold.

4.3.1. Proof of Claim 1. We prove (51a) by induction on n . Without loss of generality, we can assume $\sigma = (1, 2, \dots, n)$, then

$$L_\sigma = \begin{bmatrix} A_1^T A_1 & 0 & \dots & 0 \\ A_2^T A_1 & A_2^T A_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ A_n^T A_1 & A_n^T A_2 & \dots & A_n^T A_n \end{bmatrix}.$$

In this case, (51a) becomes

$$I - AL_\sigma^{-1} A^T = P_n P_{n-1} \dots P_1.$$

The expression obviously holds for $n = 1$. Suppose the expression holds for $n - 1$ (i.e., for $\hat{A} = [A_1, \dots, A_{n-1}]$); we have

$$\hat{Z} \triangleq I - \hat{A} \hat{L}_\sigma^{-1} \hat{A}^T = P_{n-1} \dots P_2 P_1, \quad (55)$$

where $\hat{\sigma} = (1, 2, \dots, n-1)$ is a permutation of $n-1$ elements and $\hat{L}_{\hat{\sigma}}$ is the counterpart of L_{σ} for $n-1$ blocks defined as

$$\hat{L}_{\hat{\sigma}} = \begin{bmatrix} A_1^T A_1 & 0 & \dots & 0 \\ A_2^T A_1 & A_2^T A_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ A_{n-1}^T A_1 & A_{n-1}^T A_2 & \dots & A_{n-1}^T A_{n-1} \end{bmatrix}.$$

The two matrices L_{σ} and $\hat{L}_{\hat{\sigma}}$ are related by

$$L_{\sigma} = \begin{bmatrix} \hat{L}_{\hat{\sigma}} & 0 \\ A_n^T \hat{A} & I \end{bmatrix},$$

which implies

$$L_{\sigma}^{-1} = \begin{bmatrix} \hat{L}_{\hat{\sigma}}^{-1} & 0 \\ -A_n^T \hat{A} \hat{L}_{\hat{\sigma}}^{-1} & I \end{bmatrix}.$$

Therefore, we have

$$\begin{aligned} AL_{\sigma}^{-1}A^T &= [\hat{A}, A_n] \begin{bmatrix} \hat{L}_{\hat{\sigma}}^{-1} & 0 \\ -A_n^T \hat{A} \hat{L}_{\hat{\sigma}}^{-1} & I \end{bmatrix} [\hat{A}, A_n]^T = \hat{A} \hat{L}_{\hat{\sigma}}^{-1} \hat{A}^T - A_n A_n^T \hat{A} \hat{L}_{\hat{\sigma}}^{-1} \hat{A}^T + A_n A_n^T \\ &= \hat{Z} - A_n A_n^T \hat{Z} + A_n A_n^T \\ &= I - (I - A_n A_n^T)(I - \hat{Z}) \\ &= I - P_n P_{n-1} \dots P_1, \end{aligned}$$

where in the last step we use the induction hypothesis (55). Thus we have proved (51a). Summing up (51a) for all possible permutations σ and dividing by $n!$, we obtain (51b). Q.E.D.

4.4. Proof of Proposition 2

According to (48), the (expected) update equation of RP-BCD is given by $E(x^{k+1}) - x^* = (I - QA^T A)(E(x^k) - x^*) = Z(E(x^k) - x^*)$, where $Z = I - QA^T A = I - E(L_{\sigma}^{-1}A^T A)$.

Consider a new coefficient matrix $\tilde{A} = [\tilde{A}_1, \dots, \tilde{A}_n]$, where $\tilde{A}_i = A_i(A_i^T A_i)^{-\frac{1}{2}}$. Clearly $\tilde{A}_i^T \tilde{A}_i = I_{d_i}$. Denote the corresponding matrices as $\tilde{L}_{\sigma}, \tilde{Z}$. Define $\Lambda \triangleq \text{Diag}((A_1^T A_1)^{\frac{1}{2}}, \dots, (A_n^T A_n)^{\frac{1}{2}}) = D^{1/2}$. When $\sigma = (1, 2, \dots, n)$, we have

$$L_{\sigma} = \begin{bmatrix} A_1^T A_1 & 0 & \dots & 0 \\ A_2^T A_1 & A_2^T A_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ A_n^T A_1 & A_n^T A_2 & \dots & A_n^T A_n \end{bmatrix}, \quad \tilde{L}_{\sigma} = \begin{bmatrix} \tilde{A}_1^T \tilde{A}_1 & 0 & \dots & 0 \\ \tilde{A}_2^T \tilde{A}_1 & \tilde{A}_2^T \tilde{A}_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ \tilde{A}_n^T \tilde{A}_1 & \tilde{A}_n^T \tilde{A}_2 & \dots & \tilde{A}_n^T \tilde{A}_n \end{bmatrix} = \Lambda^{-1} L_{\sigma} \Lambda^{-1}.$$

It is not hard to verify that the above relation $\tilde{L}_{\sigma} = \Lambda^{-1} L_{\sigma} \Lambda^{-1}$ is true for any σ . Similarly, we have $\tilde{A}^T \tilde{A} = \Lambda^{-1} A^T A \Lambda^{-1}$, thus

$$\tilde{L}_{\sigma}^{-1} \tilde{A}^T \tilde{A} = \Lambda L_{\sigma}^{-1} \Lambda^{-1} A^T A \Lambda^{-1} = \Lambda L_{\sigma}^{-1} A^T A \Lambda^{-1}.$$

This implies

$$\tilde{Z} = E(I - \tilde{L}_{\sigma}^{-1} \tilde{A}^T \tilde{A}) = \Lambda(I - E(L_{\sigma}^{-1} A^T A)) \Lambda^{-1} = \Lambda Z \Lambda^{-1}.$$

Consider a sequence $\tilde{x}^k = \Lambda x^k$ and define $\tilde{x}^* = \Lambda x^*$. Then from the original update equation we have $\Lambda^{-1}(E(\tilde{x}^{k+1}) - \tilde{x}^*) = Z \Lambda^{-1}(E(\tilde{x}^k) - \tilde{x}^*)$, that is,

$$E(\tilde{x}^{k+1}) - \tilde{x}^* = \Lambda Z \Lambda^{-1}(E(\tilde{x}^k) - \tilde{x}^*) = \tilde{Z}(E(\tilde{x}^k) - \tilde{x}^*).$$

According to Theorem 3, we have

$$\|E(\tilde{x}^k) - \tilde{x}^*\| \leq \left\{ 1 - \frac{1}{n} \lambda_{\min}(\tilde{A}^T \tilde{A}), \frac{1}{3} \right\}^k \|\tilde{x}^0 - \tilde{x}^*\|. \quad (56)$$

Note that $\|E(\tilde{x}^k) - \tilde{x}^*\| = \|\Lambda(E(x^k) - x^*)\| = \sqrt{(E(x^k) - x^*)^T \Lambda^2 E(x^k) - x^*} = \|E(x^k) - x^*\|_D$, and $\tilde{A}^T \tilde{A} = \Lambda^{-1} A^T A \Lambda^{-1} = D^{-1/2} A^T \tilde{A} D^{-1/2}$. Substituting into (56), we obtain the desired inequality. Q.E.D.

4.5. Proof of Theorem 4

Now we consider the expected convergence rate of RP-ADMM. The difference with the analysis for RP-BCD is that here we need to consider the distance between the eigenvalues of $I - AQA^T$ with $-1/3$, whereas for RP-BCD what matters is the distance between the eigenvalues of $I - AQA^T$ and -1 , which is at least $2/3$ and thus can be ignored.

Claim 2. Suppose the minimum and maximum eigenvalues of $QA^T A$ are $0 < \tau_{\min} \leq \tau_{\max} < 4/3$. Then

$$\rho(M) = \max \left\{ \sqrt{(1 - \tau_{\min})_+}, (\tau_{\max} - 1)_+, \sqrt{\tau_{\max}(\tau_{\max} - 1)_+} \right\},$$

where $z_+ = \max\{z, 0\}$. Furthermore, we have

$$\rho(M) \leq \max \left\{ 1 - \frac{3}{4}(4 - 3\tau_{\max}), 1 - \frac{1}{2}\tau_{\min} \right\}. \quad (57)$$

The proof of Claim 2 is given in Section 7.1. The next lemma provides a universal estimate of the maximum eigenvalues of $QA^T A$.

Lemma 4. The maximum eigenvalues of $QA^T A$ is at most $\frac{4}{3} - \frac{4}{9n+1}$, that is,

$$\tau_{\max} = \lambda_{\max}(QA^T A) \leq \frac{4}{3} - \frac{4}{9n+1}. \quad (58)$$

The proof of Lemma 4 is given in Section 7.2

According to (54), which is established in the proof of the expected convergence rate of RP-BCD, we have

$$\tau_{\min} = \lambda_{\min}(QA^T A) \geq \frac{1}{n} \lambda_{\min}(A^T A). \quad (59)$$

Substituting the bounds (58) and (59) into (57), we obtain

$$\rho(M) \leq \max \left\{ 1 - \frac{3}{4}(4 - 3\tau_{\max}), 1 - \frac{1}{2}\tau_{\min} \right\} = \max \left\{ 1 - \frac{1}{n+1}, 1 - \frac{1}{2n} \lambda_{\min}(A^T A) \right\}. \quad (60)$$

Because $\lambda_{\min}(A^T A) \leq 1$, $\frac{1}{2n} \leq \frac{1}{n+1}$, this bound can be simplified to

$$\rho(M) \leq 1 - \frac{1}{2n} \lambda_{\min}(A^T A). \quad \text{Q.E.D.}$$

Remark 5. The eigenvalues of $QA^T A$ lie in the region $(0, 4/3)$, which guarantees the expected convergence of RP-ADMM. To obtain the expected convergence rate, we need to know the distance of the spectrum to the two extremes 0 and $4/3$. We conjecture that the bound can be improved to $\rho(M) \leq 1 - \frac{1}{2} \lambda_{\min}(A^T A)$. This requires more effort than the conjecture of RP-CD: besides showing $\tau_{\min} \geq O(\lambda_{\min}(A^T A))$, we also need to show $\tau_{\max} \leq \frac{4}{3} - O(\lambda_{\min}(A^T A))$. This is left as future work.

5. Proof of Lemma 1

The proof of Lemma 1 relies on two simple techniques. The first technique, as elaborated in Step 1 below, is to factorize M and rearrange the factors. The second technique, as elaborated in the Step 2 below, is to reduce the dimension by eliminating a variable from the eigenvalue equation.

Step 1. Factorizing M and Rearranging the Order of Multiplication.

The following observation is crucial: the matrix M defined by (43) can be factorized as

$$M = \begin{bmatrix} I & 0 \\ -A & I \end{bmatrix} \begin{bmatrix} QA^T & I \\ I & A \end{bmatrix} \begin{bmatrix} -A & I \\ I & 0 \end{bmatrix}.$$

Switching the order of the products by moving the first component to the last, we get a new matrix

$$M' \triangleq \begin{bmatrix} QA^T & I \\ I & A \end{bmatrix} \begin{bmatrix} -A & I \\ I & 0 \end{bmatrix} \begin{bmatrix} I & 0 \\ -A & I \end{bmatrix} = \begin{bmatrix} QA^T & I \\ I & A \end{bmatrix} \begin{bmatrix} -2A & I \\ I & 0 \end{bmatrix} = \begin{bmatrix} I - 2QA^T A & QA^T \\ -A & I \end{bmatrix}. \quad (61)$$

Note that $\text{eig}(XY) = \text{eig}(YX)$ for any two square matrices, thus

$$\text{eig}(M) = \text{eig}(M').$$

To prove (44), we only need to prove

$$\lambda \in \text{eig}(M') \iff \frac{(1-\lambda)^2}{1-2\lambda} \in \text{eig}(QA^T A). \quad (62)$$

Step 2. Relate the Eigenvalues of M' to the Eigenvalues of $QA^T A$ (i.e., Prove (62)).

This step is simple because we only use the definition of eigenvalues. However, note that, without Step 1, just applying the definition of eigenvalues of the original matrix M may not lead to a simple relationship as (62).

We first prove one direction of (62):

$$\lambda \in \text{eig}(M') \implies \frac{(1-\lambda)^2}{1-2\lambda} \in \text{eig}(QA^T A). \quad (63)$$

Suppose $v \in \mathbb{C}^{2N \times 1} \setminus \{0\}$ is an eigenvector of M' corresponding to the eigenvalue λ , that is,

$$M'v = \lambda v.$$

Partition v as $v = \begin{bmatrix} v_1 \\ v_0 \end{bmatrix}$, where $v_1, v_0 \in \mathbb{C}^{N \times 1}$. Using the expression of M' in (61), we can write the above equation as

$$\begin{bmatrix} I - 2QA^T A & QA^T \\ -A & I \end{bmatrix} \begin{bmatrix} v_1 \\ v_0 \end{bmatrix} = \lambda \begin{bmatrix} v_1 \\ v_0 \end{bmatrix},$$

which implies

$$(I - 2QA^T A)v_1 + QA^T v_0 = \lambda v_1, \quad (64a)$$

$$-Av_1 + v_0 = \lambda v_0. \quad (64b)$$

We claim that (63) holds when $v_1 = 0$. In fact, in this case we must have $v_0 \neq 0$ (otherwise $v = 0$ cannot be an eigenvector). By (64b) we have $\lambda v_0 = v_0$, thus $\lambda = 1$. By (64a) we have $0 = QA^T v_0 = QA^T A(A^{-1}v_0)$, which implies $\frac{(1-\lambda)^2}{1-2\lambda} = 0 \in \text{eig}(QA^T A)$, therefore (63) holds in this case.

We then prove (63) for the case

$$v_1 \neq 0. \quad (65)$$

The Equation (64b) implies $(1-\lambda)v_0 = Av_1$. Multiplying both sides of (64a) by $(1-\lambda)$ and invoking this equation, we get

$$(1-\lambda)(I - 2QA^T A)v_1 + QA^T Av_1 = (1-\lambda)\lambda v_1.$$

This relation can be simplified to

$$(1-2\lambda)QA^T Av_1 = (1-\lambda)^2 v_1. \quad (66)$$

We must have $\lambda \neq \frac{1}{2}$; otherwise, the above relation implies $v_1 = 0$, which contradicts (65). Then (66) becomes

$$QA^T Av_1 = \frac{(1-\lambda)^2}{1-2\lambda} v_1. \quad (67)$$

Therefore, $\frac{(1-\lambda)^2}{1-2\lambda}$ is an eigenvalue of $QA^T A$, with the corresponding eigenvector $v_1 \neq 0$, which finishes the proof of (63).

The other direction⁹

$$\lambda \in \text{eig}(M) \iff \frac{(1-\lambda)^2}{1-2\lambda} \in \text{eig}(QA^T A) \quad (68)$$

is easy to prove. Suppose $\frac{(1-\lambda)^2}{1-2\lambda} \in \text{eig}(QA^T A)$. We consider two cases.

Case 1. $\frac{(1-\lambda)^2}{1-2\lambda} = 0$.

In this case $\lambda = 1$. Because $0 = \frac{(1-\lambda)^2}{1-2\lambda} \in \text{eig}(QA^T A)$, there exists $v_0 \in \mathbb{C}^N \setminus \{0\}$ such that $QA^T A v_0 = 0$, and let $v_1 = (0, \dots, 0)^T \in \mathbb{C}^{N \times 1}$; then v_0, v_1 and $\lambda = 1$ satisfy (64). Thus, $v = \begin{bmatrix} v_1 \\ v_0 \end{bmatrix} \in \mathbb{C}^{2N} \setminus \{0\}$ satisfies $Mv = \lambda v$, which implies $\lambda = 1 \in \text{eig}(M)$.

Case 2. $\frac{(1-\lambda)^2}{1-2\lambda} \neq 0$, then $\lambda \neq 1$.

Let v_1 be the eigenvector corresponding to $\frac{(1-\lambda)^2}{1-2\lambda}$ (i.e., pick v_1 that satisfies (67)), and define $v_0 = v_1/(1-\lambda)$. It is easy to verify that $v = \begin{bmatrix} v_1 \\ v_0 \end{bmatrix}$ satisfies $Mv = \lambda v$, which implies $\lambda \in \text{eig}(M)$.

Step 3. When Q is symmetric, prove (45) by simple algebraic computation.

Because Q is symmetric, we know that $\text{eig}(QA^T A) = \text{eig}(AQA^T) \subseteq \mathbb{R}$. Suppose $\tau \in \mathbb{R}$ is an eigenvalue of $QA^T A$; then any λ satisfying $\frac{(1-\lambda)^2}{1-2\lambda} = \tau$ is an eigenvalue of M . This relation can be rewritten as $\lambda^2 + 2(\tau - 1)\lambda + (1 - \tau) = 0$, which, as a real-coefficient quadratic equation in λ , has two roots:

$$\lambda_1 = 1 - \tau + \sqrt{\tau(\tau - 1)}, \quad \lambda_2 = 1 - \tau - \sqrt{\tau(\tau - 1)}. \quad (69)$$

Note that when $\tau(\tau - 1) < 0$, the expression $\sqrt{\tau(\tau - 1)}$ denotes a complex number $i\sqrt{\tau(1 - \tau)}$, where i is the imaginary unit. To prove (45), we only need to prove

$$\max\{|\lambda_1|, |\lambda_2|\} < 1 \iff 0 < \tau < \frac{4}{3}. \quad (70)$$

Consider three cases.

Case 1. $\tau < 0$.

Then $\tau(\tau - 1) = |\tau|(|\tau| + 1) > 0$. In this case, $\lambda_1 = 1 + |\tau| + \sqrt{|\tau|(|\tau| + 1)} > 1$.

Case 2. $0 < \tau < 1$.

Then $\tau(\tau - 1) < 0$, and (69) can be rewritten as

$$\lambda_{1,2} = 1 - \tau \pm i\sqrt{\tau(1 - \tau)},$$

which implies $|\lambda_1| = |\lambda_2| = \sqrt{(1 - \tau)^2 + \tau(1 - \tau)} = \sqrt{1 - \tau} < 1$.

Case 3. $\tau > 1$.

Then $\tau(\tau - 1) > 0$. According to (69), it is easy to verify $\lambda_1 > 0 > \lambda_2$ and

$$|\lambda_2| = \tau - 1 + \sqrt{\tau(\tau - 1)} > 1 - \tau + \sqrt{\tau(\tau - 1)} = |\lambda_1|.$$

Then we have

$$\max\{|\lambda_1|, |\lambda_2|\} < 1 \iff |\lambda_2| = \tau - 1 + \sqrt{\tau(\tau - 1)} < 1 \iff 1 < \tau < \frac{4}{3}.$$

Combining the conclusions of the three cases immediately leads to (70).

6. Proof of Lemma 2

This section is devoted to the proof of Lemma 2. We first give a proof overview in Section 6.1. The formal proof of Lemma 2 is given in Section 6.2. The proofs of the technical results involved in the proof are given in the subsequent subsections.

Without loss of generality, we can assume

$$A_i^T A_i = I_{d_i \times d_i}, \quad i = 1, \dots, n.$$

To show this, let us write M_σ, M as $M_\sigma(A_1, \dots, A_n)$ and $M(A_1, \dots, A_n)$ respectively, that is, functions of the coefficient matrix $(A_1, \dots, A_n)$. Define $\tilde{A}_i = A_i(A_i^T A_i)^{-\frac{1}{2}}$ and

$$D \triangleq \text{Diag}\left((A_1^T A_1)^{-\frac{1}{2}}, \dots, (A_n^T A_n)^{-\frac{1}{2}}, I_{N \times N}\right).$$

It is easy to verify that $M_\sigma(A_1, \dots, A_n) = D^{-1} M_\sigma(\tilde{A}_1, \dots, \tilde{A}_n) D$, which implies

$$M(A_1, \dots, A_n) = D^{-1} M(\tilde{A}_1, \dots, \tilde{A}_n) D.$$

Thus $\rho(M(A_1, \dots, A_n)) = \rho(M(\tilde{A}_1, \dots, \tilde{A}_n))$. In other words, normalizing A_i to $\tilde{A}_i$, which satisfies $\tilde{A}_i^T \tilde{A}_i = I_{d_i \times d_i}$, does not change the spectral radius of M .

6.1. Proof Overview

In the proof overview, we discuss a few issues one may encounter when proving the result, and how we resolve these issues.

The simulations show that $\|QA^T A\| < \frac{4}{3} \ll \|Q\| \|A^T A\|$, thus, we cannot relax $\|QA^T A\|$ to the product of $\|Q\|$ and $\|A^T A\|$ and have to treat $QA^T A$ as a single subject. However, each entry of $QA^T A$ is a complicated function (in fact, a high order polynomial) of the entries of $A^T A$. In other words, Q is like a black box. To open the “black box,” we use a simple expression of $Z = I - AQA^T$ proved in Claim 1, that is, $Z = E_\sigma(P_{\sigma_1}, \dots, P_{\sigma_n})$, where $P_i = I - A_i A_i^T$ is directly related to A_i . The problem becomes how to connect the eigenvalues of $E_\sigma(P_{\sigma_1}, \dots, P_{\sigma_n})$ with those of $AA^T = \sum_i A_i A_i^T = n - \sum_i P_i$.

Although this is a clear linear algebra problem, it is not easy to obtain a lower bound of $E_\sigma(P_{\sigma_1}, \dots, P_{\sigma_n})$. In fact, even though we know the eigenvalues of $Z = E_\sigma(P_{\sigma_1}, \dots, P_{\sigma_n})$ are lower bounded by -1 because RP-CD converges, it is not clear how to prove this lower bound directly from a linear algebra perspective.

In our solution, we apply two tricks. The first trick is to view $E_\sigma(P_{\sigma_1}, \dots, P_{\sigma_n})$ as an induction formula that connects it and its lower dimensional analogs. This is based on a simple observation that any permutation $(\sigma_1 \sigma_2 \dots \sigma_n)$ can be written as the concatenation of $(\sigma_1 \sigma_2 \dots \sigma_{n-1})$ and σ_n , thus the expression of $Z = E_\sigma(P_{\sigma_1}, \dots, P_{\sigma_n})$ can be decomposed accordingly. We then reduce the problem to bounding the eigenvalues of a Jordan product $P_n \hat{Z} + \hat{Z} P_n$, where P_n is a projection matrix and $\hat{Z}$ is the lower dimensional analog of Z . The second trick is to apply a formula on the eigenvalues of Jordan product developed by Strang in 1962 [36]. Somewhat surprisingly, his formula exactly leads to the desired lower bound of $-1/3$.

6.2. Proof of Lemma 2

The proof can be divided into three steps: first provide an alternative expression of AQA^T , then prove an induction formula, and finally apply Strang’s formula to perform mathematical induction. This subsection contains the major part of the proof, and the intermediate technical results will be proved in later subsections.

Step 0. Expression of $I - AQA^T$.

As proved in Claim 1, we have a simple expression of the update matrix $I - AQA^T$:

$$I - AQA^T = \frac{1}{n!} \sum_{\sigma=(\sigma_1, \dots, \sigma_n) \in \Gamma} P_{\sigma_n} P_{\sigma_{n-1}} \dots P_{\sigma_1}.$$

Step 1. Induction Formula.

For any $k \in [n]$, define

$$\Gamma_k \triangleq \{\sigma' \mid \sigma' \text{ is a permutation of } [n] \setminus \{k\}\}. \quad (71)$$

For any $\sigma' \in \Gamma_k$, we define $L_{\sigma'} \in \mathbb{R}^{(N-d_k) \times (N-d_k)}$ as a $(n-1) \times (n-1)$ block-partitioned matrix, with the $(\sigma'(i), \sigma'(j))$ -th block being

$$L_{\sigma'}[\sigma'(i), \sigma'(j)] \triangleq \begin{cases} A_{\sigma'(i)}^T A_{\sigma'(j)} & i \geq j, \\ 0 & i < j. \end{cases} \quad (72)$$

We then define $\hat{Q}_k \in \mathbb{R}^{(N-d_k) \times (N-d_k)}$ by

$$\hat{Q}_k \triangleq \frac{1}{|\Gamma_k|} \sum_{\sigma' \in \Gamma_k} L_{\sigma'}^{-1}, \quad k = 1, \dots, n. \quad (73)$$

Define W_k as the k -th block-column of $A^T A$ excluding the block $A_k^T A_k$, that is,

$$W_k = [A_k^T A_1, \dots, A_k^T A_{k-1}, A_k^T A_{k+1}, \dots, A_k^T A_n]^T, \quad \forall k \in [n]. \quad (74)$$

On the basis of the expression of $I - AQA^T$ presented before, we build a connection between the update matrix $I - AQA^T$ and its lower dimensional analogs. The proof of Proposition 3 is given in Section 6.3.

Proposition 3. Define

$$Z = I - AQA^T, \quad \hat{Z}_k = I - \hat{A}_k \hat{Q}_k \hat{A}_k^T,$$

where Q is defined as in (39), $\hat{A}_k = [A_1, \dots, A_{k-1}, A_{k+1}, \dots, A_n]$, and $\hat{Q}_k$ is defined in (73), and $P_k = I - A_k A_k^T$. Then we have

$$Z = \frac{1}{2n} \sum_{k=1}^n (P_k \hat{Z}_k + \hat{Z}_k P_k). \quad (75)$$

Step 2. Applying Strang's Result on Jordan Product to Perform Mathematical Induction.

It is obvious that the product of two symmetric matrices is not necessarily symmetric, so it is common to encounter the symmetrized product $XY + YX$, which is called Jordan product of two matrices X and Y . Our induction formula basically states that Z is the average of the Jordan product of the lower dimensional analog and P_k .

The eigenvalues of the Jordan product of two matrices have been studied before. The following result is proved in Strang [36].

Lemma 5 (Strang [36, theorem 1]; Eigenvalues of Jordan Product). Suppose two symmetric positive-semidefinite matrices X and Y satisfy

$$\alpha_1 I \leq X \leq \alpha_n I, \quad \beta_1 I \leq Y \leq \beta_n I,$$

then the maximal (resp. minimal) eigenvalue of the Jordan product $XY + YX$ are the largest (resp. smallest) of the set

$$\left\{ 2\alpha_i \beta_j, i, j \in \{1, n\}, \frac{16\alpha_1 \alpha_n \beta_1 \beta_n - (\beta_1 - \beta_n)^2 (\alpha_1 - \alpha_n)^2}{4(\alpha_1 + \alpha_n)(\beta_1 + \beta_n)} \right\}. \quad (76)$$

Let us come back to the proof of Lemma 2. We use mathematical induction to prove Lemma 2. For the basis of the induction ($n = 1$), Lemma 2 holds because $QA^T A = I_{d_1 \times d_1}$. Assume Lemma 2 holds for $n - 1$; we will prove Lemma 2 for n .

Consider one term of (75) $P_k \hat{Z}_k + \hat{Z}_k P_k$. Note that $P_k = I - A_k A_k^T$ is a projection matrix, because we have assumed $A_k^T A_k = I$. Combining with the induction hypothesis, we have

$$0 \leq P_k \leq I, \quad -\frac{1}{3}I < \hat{Z}_k < I.$$

Let $\alpha_1 = 0, \alpha_n = 1, \beta_1 = -1/3, \beta_n = 1$; then the set (76) becomes (keep the repeated values)

$$\{0, 0, -2/3, 2, -2/3\}.$$

Then by Lemma 5 we have

$$-\frac{1}{3}I \leq \frac{1}{2}(P_k \hat{Z}_k + \hat{Z}_k P_k) \leq I.$$

Note that because by the induction hypothesis the eigenvalues of $\hat{Z}_k$ cannot achieve the extreme values of region $(-1/3, 1)$, the eigenvalues of $\frac{1}{2}(P_k \hat{Z}_k + \hat{Z}_k P_k)$ also cannot.¹⁰ So we have

$$-\frac{1}{3}I < \frac{1}{2}(P_k \hat{Z}_k + \hat{Z}_k P_k) < I.$$

Thus, according to (75) we have

$$-\frac{1}{3}I < Z < I.$$

This finishes the induction step. Q.E.D.

Remark 6. Where does the magical number $-1/3$ come from? It is actually the strange and complicated term $\frac{16\alpha_1\alpha_n\beta_1\beta_n-(\beta_1-\beta_n)^2(\alpha_1-\alpha_n^2)}{4(\alpha_1+\alpha_n)(\beta_1+\beta_n)}$ in Strang's result (76), which occurs owing to the special structure of the Jordan product.

6.3. Proof of Proposition 3 (the Induction Formula)

It is easy to build an induction formula from the expression (51b). For example, when $n = 3$, the matrix $\sum_{\sigma} P_{\sigma_1} P_{\sigma_2} P_{\sigma_3}$ can be decomposed as the sum of $P_1(P_2P_3 + P_3P_2) + (P_2P_3 + P_3P_2)P_1$ and two other similar terms (changing the outside part P_1 to P_2, P_3 and the inside part $P_2P_3 + P_3P_2$ correspondingly). The inside part $P_2P_3 + P_3P_2$ only involves two matrices, thus is a lower-dimensional analog. To make this even easier to see, denote $X = P_1, Y = P_2, Z = P_3$, then

$$2 \sum_{\text{permute } X, Y, Z} XYZ = [X(YZ + ZY) + (YZ + ZY)X] + [Y(XZ + ZX) + (XZ + ZX)Y] \\ + [Z(XY + YX) + (XY + YX)Z].$$

A rigorous argument based on the above intuition is given as follows. Applying the Formula (51b) to the matrix $P_1, \dots, P_{k-1}, P_{k+1}, \dots, P_n$, and by the definition $\hat{A}_k = [A_1, \dots, A_{k-1}, A_{k+1}, \dots, A_n]$ and the definition of $\hat{Q}_k$ in (73), we have

$$I - \hat{A}_k \hat{Q}_k \hat{A}_k = \frac{1}{(n-1)!} \sum_{\sigma=(\sigma_1, \dots, \sigma_{n-1}) \in \Gamma_k} P_{\sigma_{n-1}} P_{\sigma_{n-2}} \dots P_{\sigma_1}.$$

We then have

$$2(I - AQA^T) = \frac{2}{n!} \sum_{\sigma=(\sigma_1, \dots, \sigma_n) \in \Gamma} P_{\sigma_n} P_{\sigma_{n-1}} \dots P_{\sigma_1} \\ = \frac{1}{n} \frac{1}{(n-1)!} \sum_{k=1}^n \sum_{\sigma=(\sigma_1, \dots, \sigma_{n-1}) \in \Gamma_k} (P_k P_{\sigma_{n-1}} P_{\sigma_{n-2}} \dots P_{\sigma_1} + P_{\sigma_{n-1}} P_{\sigma_{n-2}} \dots P_{\sigma_1} P_k) \\ = \frac{1}{n} \sum_{k=1}^n (P_k (I - \hat{A}_k \hat{Q}_k \hat{A}_k^T) + (I - \hat{A}_k \hat{Q}_k \hat{A}_k) P_k),$$

which is the desired formula.

6.4. Proof of Proposition 1

We provide the proof of the expected convergence of BR-ADMM here, because this proof is a slightly smaller subset of the proof of Theorem 1. We will just describe the necessary modifications.

We only need to prove a similar version of Theorem 2 (i.e., the spectral radius of the expected update matrix of BR-ADMM is less than 1). Throughout the proof, we need to change the matrix $Q = \frac{1}{|\Gamma|} \sum_{\sigma \in \Gamma} Q_{\sigma}$ to another one defined as

$$Q^{\text{BR}} \triangleq \frac{1}{|\Gamma^{\text{BR}}|} \sum_{\sigma \in \Gamma^{\text{BR}}} Q_{\sigma}, \quad (77)$$

where Γ^{BR} denotes the set of all possible permutations according to the Bernoulli randomization rule. It is easy to see that $|\Gamma^{\text{BR}}| = 2^n$. Other matrices such as M should be changed accordingly.

The proof of Theorem 2 mainly consists of Lemma 1 and Lemma 2. Because Lemma 1 has nothing to do with the specific expression of Q , we only need to prove Lemma 2 for BR-ADMM; that is, the matrix $AQ^{\text{BR}}A^T$ has all eigenvalues in the region $(0, 4/3)$. Following the proof of Lemma 2, we divide the proof into three steps.

Step 0. Expression of $Z^{\text{BR}} \triangleq I - AQ^{\text{BR}}A^T$.

In Claim 1, we have proved the expression (51a) that $I - AQ_{\sigma}A^T = P_{\sigma_n} P_{\sigma_{n-1}} \dots P_{\sigma_1}$ for any permutation σ , which implies

$$Z^{\text{BR}} = I - AQ^{\text{BR}}A^T \stackrel{(77)}{=} \frac{1}{2^n} \sum_{\sigma \in \Gamma^{\text{BR}}} P_{\sigma_n} P_{\sigma_{n-1}} \dots P_{\sigma_1}.$$

Step 1. Induction Formula.

Notice that a characteristic of the Bernoulli randomization rule is that the first block is either updated first or updated last. For instance, when $n = 4$, $(1, 3, 4, 2)$ is a feasible permutation in Γ^{BR} and $(3, 4, 2, 1)$ is also a feasible permutation, but $(3, 1, 4, 2)$ is not feasible. After removing the first block, the rest $n - 1$ blocks form a permutation in $\hat{\Gamma}^{\text{BR}}$, where $\hat{\Gamma}^{\text{BR}}$ is the set of all permutation of $2, 3, \dots, n$ according to the Bernoulli randomization rule. In other words, we have $\Gamma^{\text{BR}} = \{(1, \hat{\sigma}), (\hat{\sigma}, 1), \text{ where } \hat{\sigma} \in \hat{\Gamma}^{\text{BR}}\}$. Thus we have an induction formula

$$Z^{\text{BR}} = \frac{1}{2^n} \sum_{\sigma=(\sigma_1, \dots, \sigma_{n-1}) \in \hat{\Gamma}^{\text{BR}}} (P_1 P_{\sigma_{n-1}} \dots P_{\sigma_1} + P_{\sigma_{n-1}} \dots P_{\sigma_1} P_1) = \frac{1}{2} (P_1 \hat{Z}^{\text{BR}} + \hat{Z}^{\text{BR}} P_1), \quad (78)$$

where $\hat{Z}^{\text{BR}}$ is the lower dimensional analog of Z^{BR} for the rest $n - 1$ blocks (after removing the first block).

Step 2. Applying Mathematical Induction.

This step is almost the same as Step 2 of the proof of Lemma 2. More specifically, combining the induction hypothesis that $\text{eig}(\hat{Z}^{\text{BR}}) \in (-1/3, 1)$, Strang's result Lemma 5 and (78), we obtain the desired result $\text{eig}(Z^{\text{BR}}) \in (-1/3, 1)$. This finishes the proof.

7. Proof of Technical Results for Expected Convergence Rates**7.1. Proof of Claim 2**

Suppose all the distinct eigenvalues of $I - QA^T A$ are $0 < \tau_{N'} < \dots < \tau_1 < 4/3$, where $1 \leq N' \leq N$. Denote $\tau_{\min} = \tau_{N'}$, $\tau_{\max} = \tau_1$. According to Lemma 1, the expected update matrix of RP-ADMM M has $2N'$ distinct eigenvalues $\lambda_{k,1}, \lambda_{k,2}$ given by

$$\lambda_{k,1} = 1 - \tau_k + \sqrt{\tau_k(\tau_k - 1)}, \quad \lambda_{k,2} = 1 - \tau_k - \sqrt{\tau_k(\tau_k - 1)}, \quad k = 1, \dots, N'.$$

Suppose the integer $m \in [1, N' + 1]$ satisfies $\tau_m \leq 1 < \tau_{m-1}$. When $m = 1$, every $\tau_k \leq 1$; when $m = N' + 1$, every $\tau_k > 1$.

For $N' \geq k \geq m$, that is, $\tau_k \leq 1$, we have $\tau_k(\tau_k - 1) \leq 0$; thus, the two corresponding eigenvalues of M are

$$\lambda_{k,1} = 1 - \tau \pm i\sqrt{\tau(1 - \tau)}, \quad \lambda_{k,2} = 1 - \tau \pm i\sqrt{\tau(1 - \tau)},$$

which implies $|\lambda_{k,1}| = |\lambda_{k,2}| = \sqrt{(1 - \tau_k)^2 + \tau_k(1 - \tau_k)} = \sqrt{1 - \tau_k}$. Thus, $\rho_1 = \max_{N' \geq k \geq m} \{|\lambda_{k,1}|, |\lambda_{k,2}|\} = \sqrt{1 - \tau_{N'}} = \sqrt{1 - \tau_{\min}}$ if such k exists; when such k does not exist (i.e., $\tau_k > 1 \forall k$), we denote $\rho_1 = 0$, which equals $\sqrt{(1 - \tau_{\min})_+}$. In summary, we have $\rho_1 = \sqrt{(1 - \tau_{\min})_+}$.

For $m - 1 \geq k \geq 1$ (i.e., $\tau_k > 1$), we have $\tau_k(\tau_k - 1) > 0$. It is easy to verify $\lambda_{k,1} > 0 > \lambda_{k,2}$ and

$$|\lambda_{k,2}| = \tau_k - 1 + \sqrt{\tau_k(\tau_k - 1)} > 1 - \tau_k + \sqrt{\tau_k(\tau_k - 1)} = |\lambda_{k,1}|.$$

Denote $\rho_2 = \max_{m-1 \geq k \geq 1} \{|\lambda_{k,1}|, |\lambda_{k,2}|\}$, then

$$\rho_2 = \max_{m-1 \geq k \geq 1} \{|\lambda_{k,2}|\} = \max_{m-1 \geq k \geq 1} \left\{ \tau_k - 1 + \sqrt{\tau_k(\tau_k - 1)} \right\} = \tau_{\max} - 1 + \sqrt{\tau_{\max}(\tau_{\max} - 1)}$$

if such k exists; when such k does not exist (i.e., $\tau_k \leq 1 \forall k$), we denote $\rho_2 = 0$, which equals $(\tau_{\max} - 1)_+ + \sqrt{\tau_{\max}((\tau_{\max} - 1)_+)}$.

Combining the two scenarios, we have $\rho(M) = \max_{N' \geq k \geq 1} \{|\lambda_{k,1}|, |\lambda_{k,2}|\} = \max\{\rho_1, \rho_2\} = \max\{\sqrt{(1 - \tau_{\min})_+}, (\tau_{\max} - 1)_+ + \sqrt{\tau_{\max}((\tau_{\max} - 1)_+)}\}$.

Next, we prove

$$\begin{aligned} (\tau_{\max} - 1)_+ + \sqrt{\tau_{\max}(\tau_{\max} - 1)_+} &\leq \max \left\{ 1 - \frac{3}{4}(4 - 3\tau_{\max}), 0 \right\}, \\ \sqrt{(1 - \tau_{\min})_+} &\leq 1 - \frac{1}{2}\tau_{\min}. \end{aligned} \quad (79)$$

In fact, when $4/3 \geq \tau \geq 1$, we have

$$1 - \left(\tau - 1 + \sqrt{\tau(\tau - 1)} \right) = 2 - \tau - \sqrt{\tau(\tau - 1)} = \frac{(2 - \tau)^2 - \tau(\tau - 1)}{2 - \tau + \sqrt{\tau(\tau - 1)}} = \frac{3 - 4\tau}{2 - \tau + \sqrt{\tau(\tau - 1)}} \geq \frac{3}{4}(3 - 4\tau),$$

thus $\tau - 1 + \sqrt{\tau(\tau - 1)} \leq 1 - \frac{3}{4}(3 - 4\tau)$. When $\tau < 1$, clearly $\tau - 1 + \sqrt{\tau(\tau - 1)} = 0$. Thus $(\tau - 1)_+ + \sqrt{\tau(\tau - 1)}_+ \leq \max\{0, 1 - \frac{3}{4}(3 - 4\tau)\}$. For the second relation, if $0 \leq \tau < 1$ then $\sqrt{1 - \tau} = 1 - \frac{\tau}{1 + \sqrt{1 - \tau}} \leq 1 - \frac{\tau}{2}$; if $1 \leq \tau \leq 4/3$, then $\sqrt{(1 - \tau)_+} = 0 < 1 - \frac{1}{2}\tau$. Thus, $\sqrt{(1 - \tau)_+} \leq 1 - \frac{1}{2}\tau$ holds for any $\tau \in [0, 4/3]$.

Substituting (79) into the expression of $\rho(M)$, we obtain the desired inequality

$$\rho(M) \leq \max \left\{ 1 - \frac{3}{4}(4 - 3\tau_{\max}), 1 - \frac{1}{2}\tau_{\min} \right\}.$$

7.2. Proof of Lemma 4

This is one of the two main lemmas of proving the expected convergence rate of RP-ADMM (the other is the expected convergence rate of RP-CD).

The proof outline of Lemma 4 and the main techniques are described below. The previous proof for the expected convergence of RP-ADMM in Section 6 is not strong enough to prove a convergence rate. We have to obtain a more refined estimate of the spectral radius of AQA^T . To do so, we transform the induction formula in Proposition 3 to a “dual” form: instead of AQA^T , we consider a similar matrix QA^TA . We then apply the two simple techniques used in the proof of Lemma 1: factorize and rearrange, and reduce the dimension by eliminating a variable from the eigenvalue equation. We obtain a somewhat complicated inequality relating $\lambda_{\max}(QA^TA)$ and its lower-dimensional analog $\lambda_{\max}(\hat{Q}\hat{A}^T\hat{A})$. Finally, we perform a detailed analysis of the inequality to prove the desired bound.

7.2.1. Step 1. Mathematical Induction and Induction Formula.

Define a sequence $\{\alpha_k\}_{k=1}^{\infty}$ such that

$$\alpha_1 = 1/3, \quad \alpha_{k+1} = h(\alpha_k) \triangleq \frac{\alpha_k}{8} \frac{16 - 3\alpha_k}{2 + 3\alpha_k}. \quad (80)$$

It is easy to verify that $0 < \alpha_{k+1} < \alpha_k \leq 1/3$ for all k . The following claim provides a bound of α_k (the proof will be given in Section 7.2.4).

Claim 3. Suppose the sequence $\{\alpha_k\}_{k=1}^{\infty}$ satisfies (80); then $\alpha_k \geq \frac{4}{9(k+1)}, \forall k \geq 1$.

According to this claim, to prove the desired result $\lambda_{\max}(AQA^T) \leq \frac{4}{3} - \frac{4}{9(k+1)}$, we only need to prove the following result:

$$\text{eig}(AQA^T) \subseteq \left(0, \frac{4}{3} - \alpha_n\right]. \quad (81)$$

We prove this result by mathematical induction. When $n = 1$, because $A^TA = A_1^TA_1 = I$, we have $\lambda_{\min}(AQA^T) = \lambda_{\max}(AQA^T) = 1 = \frac{4}{3} - \alpha_1$.

Suppose the result holds for $n - 1$ (i.e., for a problem with $n - 1$ blocks); the eigenvalues of the corresponding matrix $\hat{A}\hat{Q}\hat{A}^T$ lie in the region $(0, \frac{4}{3} - \alpha_{n-1})$.

Next, we build the induction formula, which is the dual form of the one we derived before. According to (75), we have

$$2(I - AQA^T) = \frac{1}{n} \sum_{k=1}^n (P_k(I - \hat{A}_k\hat{Q}_k\hat{A}_k) + (I - \hat{A}_k\hat{Q}_k\hat{A}_k)P_k),$$

which can be rewritten as

$$AQA^T = \frac{1}{n} \sum_{k=1}^n \left[I - \frac{1}{2}P_k(I - \hat{A}_k\hat{Q}_k\hat{A}_k) - \frac{1}{2}(I - \hat{A}_k\hat{Q}_k\hat{A}_k)P_k \right]. \quad (82)$$

Note that

$$\begin{aligned} I - P_k(I - \hat{A}_k\hat{Q}_k\hat{A}_k) &= I - (I - A_kA_k^T)(I - \hat{A}_k\hat{Q}_k\hat{A}_k^T) \\ &= A_kA_k^T + \hat{A}_k\hat{Q}_k\hat{A}_k^T - A_kA_k^T\hat{A}_k\hat{Q}_k\hat{A}_k^T \\ &= [\hat{A}_k, A_k] \begin{bmatrix} \hat{Q}_k & 0 \\ -A_k^T\hat{A}_k\hat{Q}_k & I \end{bmatrix} [\hat{A}_k, A_k]^T. \end{aligned}$$

Thus the symmetrized version

$$I - \frac{1}{2}P_k(I - \hat{A}_k\hat{Q}_k\hat{A}_k) - \frac{1}{2}(I - \hat{A}_k\hat{Q}_k\hat{A}_k)P_k \quad (83)$$

$$= [\hat{A}_k, A_k] \begin{bmatrix} \hat{Q}_k & -\frac{1}{2}\hat{Q}_k^T\hat{A}_k^TA_k \\ -\frac{1}{2}A_k^T\hat{A}_k\hat{Q}_k & I \end{bmatrix} [\hat{A}_k, A_k]^T \quad (84)$$

$$= \bar{A}_k Q_k \bar{A}_k^T, \quad (85)$$

where in the last step we use the definitions

$$\bar{A}_k \triangleq [\hat{A}_k, A_k], \quad Q_k \triangleq \begin{bmatrix} \hat{Q}_k & -\frac{1}{2}\hat{Q}_k W_k \\ -\frac{1}{2}W_k^T \hat{Q}_k & I_{d_k \times d_k} \end{bmatrix}. \quad (86)$$

Summing up (85) for $k = 1, \dots, n$ and applying (82), we have

$$AQA^T = \frac{1}{n} \sum_{k=1}^n \bar{A}_k Q_k \bar{A}_k^T. \quad (87)$$

Consequently,

$$\frac{1}{n} \sum_{k=1}^n \lambda_{\min}(\bar{A}_k Q_k \bar{A}_k^T) \leq \lambda_{\min}(AQA^T) \leq \lambda_{\max}(AQA^T) \leq \frac{1}{n} \sum_{k=1}^n \lambda_{\max}(\bar{A}_k Q_k \bar{A}_k^T). \quad (88)$$

To prove $\text{eig}(AQA^T) \subseteq (0, \frac{4}{3} - \alpha_n]$, we only need to prove for any $k = 1, \dots, n$,

$$\text{eig}(\bar{A}_k Q_k \bar{A}_k^T) \subseteq \left(0, \frac{4}{3} - \alpha_n\right). \quad (89)$$

Note that $\hat{Q}_k$ only depends on the entries of $\hat{A}_k^T \hat{A}_k \in \mathbb{R}^{(N-d_k) \times (N-d_k)}$, which has $(n-1) \times (n-1)$ blocks; thus by the induction hypothesis, we have

$$\text{eig}(\hat{Q}_k \hat{A}_k^T \hat{A}_k) \subseteq \left(0, \frac{4}{3} - \alpha_{n-1}\right). \quad (90)$$

Proposition 4. Suppose $A = [\hat{A}_n, A_n] \in \mathbb{R}^{N \times N}$ is a nonsingular matrix, where $\hat{A}_n \in \mathbb{R}^{N \times (N-d_n)}$, and $A_n \in \mathbb{R}^{N \times d_n}$ satisfies $A_n^T A_n = I_{d_n \times d_n}$. Suppose $\hat{Q}_n \in \mathbb{R}^{(N-d_n) \times (N-d_n)}$ is symmetric, satisfying

$$\text{eig}(A \hat{Q}_n A^T) \subseteq \left(0, \frac{4}{3} - \alpha_{n-1}\right), \quad (91)$$

where $\{\alpha_k\}$ is defined in (80). Define

$$W_n \triangleq \hat{A}_n^T A_n \in \mathbb{R}^{(N-d_n) \times d_n}, \quad Q_n \triangleq \begin{bmatrix} \hat{Q}_n & -\frac{1}{2}\hat{Q}_n W_n \\ -\frac{1}{2}W_n^T \hat{Q}_n & I_{d_n \times d_n} \end{bmatrix}. \quad (92)$$

Then $\text{eig}(AQ_n A^T) \subseteq (0, \frac{4}{3} - \alpha_n]$.

The proof of Proposition 4 will be divided into two parts and given in Sections 7.2.2 and 7.2.3.

We claim that (89) follows from the induction hypothesis (90) and the expressions of $\bar{A}_k$ and Q_k in (86). In fact, the above proposition directly proves (89) for $k = n$. If we replace $A, \hat{A}_n, A_n, \hat{Q}_n, Q_n$ by $\bar{A}_k, \hat{A}_k, A_k, \hat{Q}_k, Q_k$, respectively, in the following proposition, we will obtain (89) for any k . Finally, as mentioned earlier, the desired result $\text{eig}(AQA^T) \subseteq (0, \frac{4}{3} - \alpha_n]$ in Lemma 2 follows immediately from (89) and (88).

7.2.2. Step 2. Relation Between $\lambda_{\max}(A_n Q_n A_n^T)$ and Its Analog.

In this subsection, we provide a proof of a weaker result $\text{eig}(AQ_n A^T) \subseteq (0, \frac{4}{3})$ under the conditions of Proposition 4; the proof of the desired result $\text{eig}(AQ_n A^T) \subseteq (0, \frac{4}{3} - \alpha_n]$ will be provided in the next subsection.

For simplicity, throughout this proof, we denote

$$W \triangleq W_n \in \mathbb{R}^{(N-d_n) \times d_n}, \hat{Q} \triangleq \hat{Q}_n \in \mathbb{R}^{(N-d_n) \times (N-d_n)}, \hat{A} \triangleq \hat{A}_n \in \mathbb{R}^{N \times (N-d_n)}.$$

According to the assumption of Proposition 4, we have

$$\hat{\lambda} \triangleq \lambda_{\max}(A\hat{Q}A^T) \in \left(0, \frac{4}{3} - \alpha_{n-1}\right]. \quad (93)$$

We first prove

$$0 \leq \Theta \triangleq W^T \hat{Q} W < \frac{4}{3} I. \quad (94)$$

Because $\text{eig}(\hat{Q}\hat{A}^T\hat{A}) \subseteq (0, \infty)$ and $\hat{A}$ is nonsingular, thus $\hat{Q} > 0$. Then we have $\Theta = W^T \hat{Q} W \geq 0$, which proves the first relation of (94). By the definition $W = \hat{A}^T A_n$ we have

$$\begin{aligned} \rho(\Theta) &= \rho(A_n^T \hat{A} \hat{Q} \hat{A}^T A_n) = \max_{v \in \mathbb{R}^{d_n \times 1}, \|v\|=1} v^T A_n^T \hat{A} \hat{Q} \hat{A}^T A_n v \\ &\leq \rho(\hat{A} \hat{Q} \hat{A}^T) \max_{v \in \mathbb{R}^{d_n \times 1}, \|v\|=1} \|A_n v\|^2 = \rho(\hat{A} \hat{Q} \hat{A}^T) \|A_n\|^2 = \rho(\hat{A} \hat{Q} \hat{A}^T) < \frac{4}{3}, \end{aligned} \quad (95)$$

where the last equality is due to the assumption $A_n^T A_n = I$, and the last inequality is due to the assumption (91). By (95) we have $\Theta < \frac{4}{3} I$, thus (94) is proved.

We apply a trick that we have previously used: factorize Q_n and change the order of multiplication. To be specific, Q_n defined in (92) can be factorized as

$$Q_n = \begin{bmatrix} I & 0 \\ -\frac{1}{2}W^T & I \end{bmatrix} \begin{bmatrix} \hat{Q} & 0 \\ 0 & I - \frac{1}{4}W^T \hat{Q} W \end{bmatrix} \begin{bmatrix} I & -\frac{1}{2}W \\ 0 & I \end{bmatrix} = J \begin{bmatrix} \hat{Q} & 0 \\ 0 & C \end{bmatrix} J^T, \quad (96)$$

where

$$J \triangleq \begin{bmatrix} I & 0 \\ -\frac{1}{2}W^T & I \end{bmatrix},$$

I in the upper left block denotes the $(N - d_n)$ -dimensional identity matrix, I in the lower right block denotes the d_n -dim identity matrix, and

$$C \triangleq I - \frac{1}{4}W^T \hat{Q} W \in \mathbb{R}^{d_n \times d_n}. \quad (97)$$

It is easy to prove

$$\text{eig}(AQ_n A^T) \subseteq (0, \infty). \quad (98)$$

In fact, we only need to prove $Q_n > 0$. According to (96), we only need to prove $\begin{bmatrix} \hat{Q} & 0 \\ 0 & C \end{bmatrix} > 0$. This follows from $\hat{Q} > 0$ and the fact $C = I - \frac{1}{4}W^T \hat{Q} W \stackrel{(94)}{>} I - \frac{1}{3}I > 0$. Thus (98) is proved.

It remains to prove

$$\rho(AQ_n A^T) < \frac{4}{3}. \quad (99)$$

Denote $\hat{B} \triangleq \hat{A}^T \hat{A} \in \mathbb{R}^{(N-d_n) \times (N-d_n)}$; then we can write $A^T A$ as

$$A^T A = \begin{bmatrix} \hat{B} & W \\ W^T & I \end{bmatrix}. \quad (100)$$

We simplify the expression of $\rho(AQ_n A^T)$ as follows:

$$\rho(AQ_n A^T) = \rho\left(AJ \begin{bmatrix} \hat{Q} & 0 \\ 0 & C \end{bmatrix} J^T A^T\right) = \rho\left(\begin{bmatrix} \hat{Q} & 0 \\ 0 & C \end{bmatrix} J^T A^T A J\right). \quad (101)$$

By algebraic computation, we have

$$\begin{aligned} J^T A^T A J &= \begin{bmatrix} I & -\frac{1}{2}W \\ 0 & I \end{bmatrix} \begin{bmatrix} \hat{B} & W \\ W^T & I \end{bmatrix} \begin{bmatrix} I & 0 \\ -\frac{1}{2}W^T & I \end{bmatrix} \\ &= \begin{bmatrix} I & -\frac{1}{2}W \\ 0 & I \end{bmatrix} \begin{bmatrix} \hat{B} - \frac{1}{2}WW^T & W \\ \frac{1}{2}W^T & I \end{bmatrix} = \begin{bmatrix} \hat{B} - \frac{3}{4}WW^T & \frac{1}{2}W \\ \frac{1}{2}W^T & I \end{bmatrix}, \end{aligned} \quad (102)$$

thus,

$$Y \triangleq \begin{bmatrix} \hat{Q} & 0 \\ 0 & C \end{bmatrix} J^T A^T A J = \begin{bmatrix} \hat{Q} & 0 \\ 0 & C \end{bmatrix} \begin{bmatrix} \hat{B} - \frac{3}{4}WW^T & \frac{1}{2}W \\ \frac{1}{2}W^T & I \end{bmatrix} = \begin{bmatrix} \hat{Q}\hat{B} - \frac{3}{4}\hat{Q}WW^T & \frac{1}{2}\hat{Q}W \\ \frac{1}{2}CW^T & C \end{bmatrix}. \quad (103)$$

Suppose $\lambda > 0$ is the maximal eigenvalue of Y . According to (101) that $\rho(AQ_n A^T) = \rho(Y)$, we also have $\lambda = \lambda_{\max}(AQ_n A^T)$. To prove (99), we only need to prove

$$\lambda < \frac{4}{3}. \quad (104)$$

Suppose $v \in \mathbb{R}^{N \times 1} \setminus \{0\}$ is the eigenvector corresponding to λ (i.e., $Zv = \lambda v$). Partition v into $v = \begin{bmatrix} v_1 \\ v_0 \end{bmatrix}$, where $v_1 \in \mathbb{R}^{N-d_n}, v_0 \in \mathbb{R}^{d_n}$. According to the expression of Z in (103), $Zv = \lambda v$ implies

$$\left(\hat{Q}\hat{B} - \frac{3}{4}\hat{Q}WW^T \right) v_1 + \frac{1}{2}\hat{Q}Wv_0 = \lambda v_1, \quad (105a)$$

$$\frac{1}{2}CW^T v_1 + Cv_0 = \lambda v_0. \quad (105b)$$

If $\lambda I - C$ is singular (i.e., λ is an eigenvalue of C), then by (94) we have $\frac{2}{3}I < C = 1 - \frac{1}{4}\Theta \leq I$, which implies $\lambda \leq 1$, thus (104) holds. In the following, we assume

$$\lambda I - C \text{ is nonsingular.} \quad (106)$$

An immediate consequence is

$$v_1 \neq 0,$$

because otherwise (105b) implies $Cv_0 = \lambda v_0$, which combined with (106) leads to $v_0 = 0$ and thus $v = 0$, a contradiction.

By (105b) we get

$$v_0 = \frac{1}{2}(\lambda I - C)^{-1}CW^T v_1.$$

Plugging into (105a), we obtain

$$\lambda v_1 = \left(\hat{Q}\hat{B} - \frac{3}{4}\hat{Q}WW^T \right) v_1 + \frac{1}{2}\hat{Q}W \frac{1}{2}(\lambda I - C)^{-1}CW^T v_1 = (\hat{Q}\hat{B} + \hat{Q}W\Phi W^T) v_1, \quad (107)$$

where

$$\begin{aligned} \Phi &\triangleq -\frac{3}{4}I + \frac{1}{4}(\lambda I - C)^{-1}C = -I + \frac{1}{4}[I + (\lambda I - C)^{-1}C] \\ &= -I + \frac{\lambda}{4}(\lambda I - C)^{-1} = -I + \lambda[(4\lambda - 4)I + \Theta]^{-1}. \end{aligned} \quad (108)$$

Here we have used the definition $C = I - \frac{1}{4}W^T \hat{Q}W = I - \frac{1}{4}\Theta$. Because Θ is a symmetric matrix, Φ is also a symmetric matrix.

Define

$$\tilde{H} \triangleq \hat{Q}W\Phi W^T \in \mathbb{R}^{(N-d_n) \times (N-d_n)}, H \triangleq W^T \hat{Q}W\Phi = \Theta\Phi \in \mathbb{R}^{d_n \times d_n}. \quad (109)$$

As a well-known linear algebra result, $\tilde{H}$ and H have the same nonzero eigenvalues. Note that $\lambda_{\max}(H)$ may not be equal to $\lambda_{\max}(\tilde{H})$ owing to the possible zero eigenvalues. Nevertheless, we can define $\lambda_{\max}^+(X) \triangleq \max\{\lambda_{\max}(X), 0\}$, and then we have

$$\lambda_{\max}^+(\tilde{H}) = \lambda_{\max}^+(H).$$

According to (109) and (108), we know

$$\begin{aligned} H &= \Theta\Phi = \Theta(-I + \lambda[(4\lambda - 4)I + \Theta]^{-1}) \\ &= -\Theta + \lambda\Theta[(4\lambda - 4)I + \Theta]^{-1} \\ &= -\Theta + \lambda(I - (4\lambda - 4)[(4\lambda - 4)I + \Theta]^{-1}) \\ &= -\Theta + \lambda I - \lambda(4\lambda - 4)[(4\lambda - 4)I + \Theta]^{-1}. \end{aligned}$$

It is well known that if $\alpha I + \Theta$ is invertible, then Θ has an eigenvalue θ iff $(\alpha I + \Theta)^{-1}$ has an eigenvalue $(\alpha + \theta)^{-1}$, and the corresponding eigenvectors are the same. Similarly, because we already assumed $(4\lambda - 4)I + \Theta$ is invertible, θ is an eigenvalue of Θ iff $H = -\Theta + \lambda I - \lambda(4\lambda - 4)[(4\lambda - 4)I + \Theta]^{-1}$ has an eigenvalue $-\theta + \lambda - \lambda(4\lambda - 4)[(4\lambda - 4) + \theta]^{-1}$. Recall that $\Theta = W^T \hat{Q}W$ satisfies $0 \leq \Theta \leq \hat{\lambda}I$; thus any eigenvalue θ satisfies $0 \leq \theta \leq \hat{\lambda}$. Therefore

$$\lambda_{\max}(H) \leq \max_{\theta \in [0, \hat{\lambda}]} \left\{ -\theta + \lambda - \frac{\lambda(4\lambda - 4)}{(4\lambda - 4) + \theta} \right\} \triangleq g(\theta). \quad (110)$$

Because $v_1 \neq 0$, without loss of generality, we can assume $\|v_1\| = 1$. We have

$$\begin{aligned} \lambda &= v_1^T \hat{Q} \hat{B} v_1 + v_1^T \tilde{H} v_1 \leq \hat{\lambda} + v_1^T \tilde{H} v_1 \leq \hat{\lambda} + \lambda_{\max}^+(\tilde{H}) = \hat{\lambda} + \lambda_{\max}^+(H) \\ &\leq \hat{\lambda} + \max \left\{ 0, \max_{\theta \in [0, \hat{\lambda}]} \left\{ -\theta + \lambda - \frac{\lambda(4\lambda - 4)}{(4\lambda - 4) + \theta} \right\} \right\}, \end{aligned} \quad (111)$$

where the first equality is due to (107), the first inequality is due to the induction hypothesis, the second inequality uses the obvious relation $\lambda_{\max}(\tilde{H}) \leq \lambda_{\max}^+(\tilde{H})$, and the last inequality is due to (110).

To prove (104), we consider two cases.

Case 1. $\max_{\theta \in [0, \hat{\lambda}]} g(\theta) \leq 0$.

In this case, $\lambda \leq \hat{\lambda} < 4/3$, where the first inequality is due to (111), and the second inequality is due to the induction hypothesis. Thus, in Case 1 (104) holds.

Case 2. $\max_{\theta \in [0, \hat{\lambda}]} g(\theta) > 0$.

Then there exists some $\theta \geq 0$ such that $g(\theta) > 0$. Note that $g(\theta)$ can also be expressed as $g(\theta) = \theta(-1 + \frac{\lambda}{(4\lambda - 4) + \theta})$; thus

$$-1 + \frac{\lambda}{(4\lambda - 4) + \theta} > 0. \quad (112)$$

If $\lambda < 1$, then (104) already holds; so we can assume $\lambda > 1$. Thus (112) implies $1 < \frac{\lambda}{(4\lambda - 4) + \theta} \leq \frac{\lambda}{4\lambda - 4}$, which leads to $\lambda < \frac{4}{3}$. Thus in Case 2 (104) also holds. This finishes the proof of (104).

Remark 7. The proof of this subsection can lead to an alternative proof of Lemma 2. In particular, the induction step (Step 2) of Section 6.2 can be replaced by the proof here. The proof presented here is more complicated and less intuitive than the one in Section 6.2 (which is just a straightforward application of Strang's result Lemma 5), but the benefit is that it can help establish a stronger bound of λ , as done in the next subsection.

7.2.3. Step 3. More Precise Bound of λ . We will continue the proof in Section 7.2.2, to further prove

$$\lambda = \lambda_{\max}(AQ_n A^T) \leq 4/3 - \alpha_n. \quad (113)$$

We rewrite (111) as follows:

$$\lambda \leq \hat{\lambda} + \max \left\{ 0, \max_{\theta \in [0, \hat{\lambda}]} g(\theta) \right\}, \text{ where } g(\theta) = \lambda - \frac{\lambda(4\lambda - 4)}{4\lambda - 4 + \theta} - \theta. \quad (114)$$

If $\lambda < 1$, then we are done because $1 \leq 4/3 - \alpha_n$. Assume $1 \leq \lambda < 4/3$ from now on.

We first analyze the function $g(\theta)$. Taking the derivative of g , we get

$$g'(\theta) = \frac{\lambda(4\lambda - 4)}{(4\lambda - 4 + \theta)^2} - 1 = \frac{(\sqrt{\lambda(4\lambda - 4)} + 4\lambda - 4 + \theta)(\sqrt{\lambda(4\lambda - 4)} - 4\lambda + 4 - \theta)}{(4\lambda - 4 + \theta)^2}.$$

Because $\lambda > 1$ and $\theta \geq 0$, the term in the first bracket in the numerator is positive. Define

$$\theta^* = \sqrt{\lambda(4\lambda - 4)} - 4\lambda + 4 > 0,$$

where the inequality holds due to $\lambda < 4/3$. Then we have

$$g'(\theta) \begin{cases} \geq 0, & \theta \leq \theta^*; \\ \leq 0, & \theta > \theta^*. \end{cases}$$

Therefore, $g(\theta)$ is increasing in $[0, \theta^*]$ and decreasing in $[\theta^*, \infty)$. This implies

$$g(\theta) \leq g(\theta^*), \quad \forall \theta \geq 0. \quad (115)$$

According to $0 < \lambda < 4/3$, we have $\lambda > \sqrt{\lambda(4\lambda - 4)} = 4\lambda - 4 + \theta^* \Rightarrow -1 + \frac{\lambda}{4\lambda - 4 + \theta^*} > 0 \Rightarrow g(\theta^*) > 0$. Together with (115) we obtain $\max\{0, \max_{\theta \in [0, \hat{\lambda}]} g(\theta)\} \leq g(\theta^*)$. Substituting into (114), we obtain

$$\lambda \leq \hat{\lambda} + g(\theta^*).$$

We will derive an inequality on λ and $\hat{\lambda}$ from the above relation as below. Substituting the expression of $g(\cdot)$ into the relation, we obtain

$$\lambda \leq \hat{\lambda} + \lambda - \frac{\lambda(4\lambda - 4)}{4\lambda - 4 + \theta^*} - \theta^* \Rightarrow \hat{\lambda} \geq \frac{\lambda(4\lambda - 4)}{4\lambda - 4 + \theta^*} + \theta^* = \sqrt{\lambda(4\lambda - 4)} + \theta^* = 2\sqrt{\lambda(4\lambda - 4)} - 4\lambda + 4.$$

This implies

$$\begin{aligned} & \hat{\lambda}^2 + (4\lambda - 4)^2 + 2\hat{\lambda}(4\lambda - 4) \geq 4\lambda(4\lambda - 4) \\ \iff & \hat{\lambda}^2 - \lambda^2 + 2(\hat{\lambda} - \lambda)(4\lambda - 4) + (\lambda - (4\lambda - 4))^2 \geq 0 \\ \iff & (\hat{\lambda} - \lambda)(\hat{\lambda} + \lambda) + 2(\hat{\lambda} - \lambda)(4\lambda - 4) + (4 - 3\lambda)^2 \geq 0. \end{aligned} \quad (116a)$$

Define

$$\delta = 4/3 - \lambda \in (0, 1/3), \quad \hat{\delta} = 4/3 - \hat{\lambda} \in (0, 4/3). \quad (117)$$

Substituting into (116a), we obtain

$$\begin{aligned} & (\delta - \hat{\delta})(8/3 - \delta - \hat{\delta}) + (\delta - \hat{\delta})(8/3 - 8\delta) + 9\delta^2 \geq 0 \\ \iff & (\delta - \hat{\delta})(16/3 - 9\delta - \hat{\delta}) + 9\delta^2 \geq 0 \\ \iff & \frac{16}{3}\delta - \frac{16}{3}\hat{\delta} + 8\hat{\delta}\delta + \hat{\delta}^2 \geq 0 \\ \iff & \delta \geq \frac{\hat{\delta}(16 - 3\hat{\delta})}{8(2 + 3\hat{\delta})} = h(\hat{\delta}). \end{aligned}$$

It is easy to verify that $h(t)$ is increasing in $t \in [0, 4/3]$; in fact, $h'(t) = \frac{36}{(2+3t)^2} - 1 = \frac{(8+3t)(4-3t)}{(2+3t)^2} \geq 0$ for $t \in [0, 4/3]$.

According to (93), we have $\hat{\delta} = 4/3 - \hat{\lambda} \geq \alpha_{n-1}$. Applying the monotonicity of h , we have

$$\delta \geq h(\hat{\delta}) \geq h(\alpha_{n-1}) = \alpha_n,$$

which combined with (117) leads to (113). This finishes the proof of Proposition 4.

7.2.4. Proof of Claim 3. Define another sequence as $\omega_k = \frac{16}{3\alpha_k} - 9k$. Then $\alpha_k = \frac{16}{3} \frac{1}{9k + \omega_k}$ and $\omega_1 = 7$, $\omega_2 = 38/5$. We then derive the recurrence equation of ω_k . According to (80), we have

$$\begin{aligned} \frac{16}{3} \frac{1}{9k + 9 + \omega_{k+1}} &= \frac{2}{3} \frac{1}{9k + \omega_k} \frac{16 - 16/(9k + \omega_k)}{2 + 16/(9k + \omega_k)} = \frac{2}{3} \frac{1}{9k + \omega_k} \frac{16(9k + \omega_k - 1)}{2(9k + \omega_k + 8)} \\ \Rightarrow 9k + 9 + \omega_{k+1} &= \frac{(9k + \omega_k)(9k + \omega_k + 8)}{9k + \omega_k - 1} \\ \Rightarrow \omega_{k+1} &= \omega_k + \frac{1}{9k + \omega_k - 1} [(9k + \omega_k)(9k + \omega_k + 8) - (9k + \omega_k - 1)(9k + 9 + \omega_k)] \\ \Rightarrow \omega_{k+1} &= \omega_k + \frac{9}{9k + \omega_k - 1}. \end{aligned}$$

It is easy to see that $\omega_k > 0 \Rightarrow \omega_{k+1} > \omega_k > 0$, thus

$$\omega_k > \omega_1 = 7, \forall k.$$

Furthermore, $\omega_{k+1} = \omega_k + \frac{9}{9k + \omega_k - 1} \leq \omega_k + \frac{1}{k}$, thus

$$\omega_k \leq \omega_1 + \sum_{j=1}^{k-1} \frac{1}{j} \leq 8 + \log(k-1).$$

The lower bound and upper bound on ω_k imply upper and lower bounds on α_k :

$$\frac{16}{3} \frac{1}{9k + 7} \geq \alpha_k \geq \frac{16}{3} \frac{1}{9k + 8 + \log(k-1)}. \quad (118)$$

As a side comment, this implies that $\lim_{k \rightarrow \infty} \alpha_k = \frac{16}{27k} \approx \frac{0.59}{k}$. For our purpose, we need a universal lower bound on α_k . When $k \geq 3$, we have $3k \geq 8 + \log(k-1)$, thus $12k \geq 9k + 8 + \log(k-1)$, which further implies

$$\frac{16}{3} \frac{1}{9k + 8 + \log(k-1)} \geq \frac{4}{9k}, \forall k \geq 3.$$

Combining with the bound (118), we obtain

$$\alpha_k \geq \frac{4}{9k} > \frac{4}{9(k+1)}, \forall k \geq 3.$$

Notice that $\alpha_1 = \frac{1}{3} > \frac{4}{9} \cdot \frac{1}{2}$, and $\alpha_2 = \frac{5}{24} > \frac{4}{9} \cdot \frac{1}{3}$; we have $\alpha_k > \frac{4}{9(k+1)}$ for any $k \geq 1$. This finishes the proof of the claim.

7.3. Proof of Lemma 3

We rewrite the lemma statement below. Suppose $P_i \in \mathbb{R}^{N \times N}$, $i = 1, \dots, n$ are projection matrices; then the lemma claims that

$$\frac{1}{n!} \sum_{\sigma=(\sigma_1, \dots, \sigma_n) \in \Gamma} P_{\sigma_n} P_{\sigma_{n-1}} \dots P_{\sigma_1} \leq \frac{1}{n} \sum_i P_i. \quad (119)$$

We first prove the case $n = 2$, $n = 3$, and $n = 4$, then prove the general case $n = 2k$ and $n = 2k + 1$ separately.

When $n = 2$, (119) reduces to $P_1 P_2 + P_2 P_1 \leq P_1 + P_2$. Notice that $P_i = P_i^2$ because P_i is a projection matrix, we have $P_1 + P_2 - P_1 P_2 + P_2 P_1 = P_1^2 + P_2^2 - P_1 P_2 + P_2 P_1 = (P_1 - P_2)^2 = (P_1 - P_2)^T (P_1 - P_2) \geq 0$.

When $n = 3$, (119) reduces to $\frac{1}{6} \sum_{i,j,k \text{ are distinct}} P_i P_j P_k \leq \frac{1}{3} (P_1 + P_2 + P_3)$. Note that $(P_i - P_k) P_j (P_i - P_k) \geq 0$, thus

$$P_i P_j P_i + P_k P_j P_k \geq P_i P_j P_k + P_k P_j P_i.$$

Summing up the above inequality for all possible triples (i, j, k) , we get

$$\sum_{i \neq j} P_i P_j P_i \geq \sum_{i,j,k \text{ are distinct}} P_i P_j P_k. \quad (120)$$

We then need to bound the left-hand-side of the above inequality. Because $I - P_j \geq 0$, we have $P_i(I - P_j)P_i \geq 0$, which implies $P_i \geq P_iP_jP_i$. Summing up this inequality for all pairs $i \neq j$, we obtain $\frac{1}{6} \sum_{i \neq j} P_iP_jP_i \leq \frac{1}{3}(P_1 + P_2 + P_3)$. Combining with (120), we obtain the desired inequality $\frac{1}{6} \sum_{i,j,k \text{ are distinct}} P_iP_jP_k \leq \frac{1}{3}(P_1 + P_2 + P_3)$.

The proof for $n = 4$ illustrates partially the gist of a general proof, so we present this proof. When $n = 4$, (119) reduces to $\frac{1}{24} \sum_{i,j,k,l \text{ are distinct}} P_iP_jP_kP_l \leq \frac{1}{4}(P_1 + P_2 + P_3 + P_4)$. Similar to (120) in the $n = 3$ case, we first prove

$$\frac{1}{24} \sum_{i,j,k,l \text{ are distinct}} P_iP_jP_kP_l \leq \frac{1}{12} \sum_{i \neq j} P_iP_jP_i. \quad (121)$$

To prove this inequality, we need the following two basic inequalities:

$$\begin{aligned} (P_i - P_l)(P_j + P_k)^2(P_i - P_l) &\geq 0, \\ (P_i + P_l)(P_j - P_k)^2(P_i + P_l) &\geq 0. \end{aligned}$$

Summing up these two inequalities, we can eliminate terms like $P_iP_jP_kP_l$ (with three distinct subscripts) and keep the terms like $P_iP_jP_i$ (with two distinct subscripts) and $P_iP_jP_kP_l$ (with four distinct subscripts), to obtain

$$P_iP_jP_i + P_iP_kP_i + P_lP_jP_l + P_lP_kP_l \geq P_iP_jP_kP_l + P_iP_kP_jP_l + P_lP_jP_kP_i + P_lP_kP_jP_i.$$

Summing up this inequality for all possible (i, j, k, l) that are distinct, we obtain (121). Similar to the proof of $n = 3$ case, we have $\frac{1}{12} \sum_{i \neq j} P_iP_jP_i \leq \frac{1}{4}(P_1 + P_2 + P_3 + P_4)$; thus combining with (121) we obtain the desired result.

We next prove the case $n = 2k$, where $k \geq 2$ is a positive integer. We will prove that

$$E_{\sigma \in \Gamma} (P_{\sigma_n} P_{\sigma_{n-1}} \dots P_{\sigma_1}) \leq E_{\pi \in \Gamma_k} (P_{\pi_1} \dots P_{\pi_{k-1}} P_{\pi_k} P_{\pi_{k-1}} \dots P_{\pi_1}), \quad (122)$$

where Γ_k is the set of k -permutations of $1, 2, \dots, n$ (here, a k -permutation is a permutation of k distinct numbers chosen from $1, 2, \dots, n$), and $E_{\sigma \in \Gamma}$ and $E_{\pi \in \Gamma_k}$ denote the expectation over a uniform distribution on Γ and Γ_k , respectively.

To prove (122), we need the following fact: for any $\epsilon = (\epsilon_1, \dots, \epsilon_k) \in \{1, -1\}^k$, we have

$$G_{\sigma, \epsilon} \triangleq (P_{\sigma_n} + \epsilon_1 P_{\sigma_1})(P_{\sigma_{n-1}} + \epsilon_2 P_{\sigma_2}) \dots (P_{\sigma_{k+1}} + \epsilon_k P_{\sigma_k})(P_{\sigma_{k+1}} + \epsilon_k P_{\sigma_k}) \dots (P_{\sigma_n} + \epsilon_1 P_{\sigma_1}) \geq 0. \quad (123)$$

This relation holds because for any positive-semidefinite matrix X and any symmetric matrix Y , we have $YXY = Y^TXY \geq 0$. Applying this fact k times leads to (123).

The expression of $G_{\sigma, \epsilon}$ in (123) involves 2^k terms in the form of $P_{i_1}P_{i_2} \dots P_{i_n}$. To prove (122), only two terms are of interest to us. The strategy is to pick ϵ_i 's properly so that summing up a bunch of relations of the form (123) will eliminate all but the two desired terms. We elaborate this strategy below.

Define

$$\begin{aligned} \Lambda_k &\triangleq \{(\epsilon_1, \dots, \epsilon_k) \in \{1, -1\}^k \mid \text{the number of } -1 \text{ in } \epsilon_1, \dots, \epsilon_k \text{ is odd}\}, \\ \Lambda_k^c &\triangleq \{(\epsilon_1, \dots, \epsilon_k) \in \{1, -1\}^k \mid \text{the number of } -1 \text{ in } \epsilon_1, \dots, \epsilon_k \text{ is even}\}. \end{aligned}$$

For example, when $k = 3$, $\Lambda_3 = \{(-1, 1, 1), (1, -1, 1), (1, 1, -1), (-1, -1, -1)\}$, and the complement $\Lambda_3^c = \{(1, 1, 1), (-1, -1, 1), (-1, 1, -1), (1, -1, -1)\}$. As a well-known fact,

$$|\Lambda_k^c| - |\Lambda_k| = \sum_{i \text{ is even}, 0 \leq i \leq n} \binom{n}{i} - \sum_{i \text{ is odd}, 0 \leq i \leq n} \binom{n}{i} = (1 - 1)^k = 0. \quad (124)$$

This matrix $G_{\sigma, \epsilon}$ can be expressed as the sum of 2^k terms, and each term is of the form $\pm P_{\pi_1} \dots P_{\pi_n}$, where $\pi_i \in \{\sigma_i, \sigma_{n+1-i}\}$. For the fixed permutation σ , define a set

$$\Omega(\sigma) = \{(\pi_1, \dots, \pi_n) \mid \pi_i \in \{\sigma_i, \sigma_{n+1-i}\}, \forall i\}.$$

We partition the set into three subsets:

$$\begin{aligned} \Omega_0(\sigma) &= \{(\pi_1, \dots, \pi_n) \in \Omega \mid \pi_i = \pi_{n+1-i}, \forall i\}, \\ \Omega_1(\sigma) &= \{(\pi_1, \dots, \pi_n) \in \Omega \mid \pi_i \neq \pi_{n+1-i}, \forall i\}, \\ \Omega_2(\sigma) &= \Omega \setminus (\Omega_0 \cup \Omega_1). \end{aligned}$$

For most of the proof, we will use the abbreviation $\Omega_t = \Omega_t(\sigma)$, $t = 0, 1, 2$. For any $\pi = (\pi_1, \dots, \pi_n) \in \Omega$, define an indicator vector of π as $\delta(\pi) = (\delta_1, \dots, \delta_k)$, where each δ_i is determined by

$$\delta_i = \mathbb{I}(\pi_i - \pi_{n+1-i}) = \begin{cases} 0, & \pi_i = \pi_{n+1-i}, \\ 1, & \pi_i \neq \pi_{n+1-i}, \end{cases} \quad (125)$$

where $\mathbb{I}(z)$ equals 0 if $z = 0$ and equals 1 if $z \neq 0$. For example, when $n = 6$ and $\pi = (\sigma_1, \sigma_2, \sigma_3, \sigma_4, \sigma_2, \sigma_6)$, the corresponding indicator vector is $(0, 0, 1)$; when $\pi = (\sigma_1, \sigma_2, \sigma_3, \sigma_3, \sigma_2, \sigma_1)$, the indicator vector is $(0, 0, 0)$. Clearly, we have

$$\delta(\pi) = (0, 0, \dots, 0), \quad \forall \pi \in \Omega_0; \quad \delta(\pi) = (1, 1, \dots, 1), \quad \forall \pi \in \Omega_1; \quad \delta(\pi) \notin \{0_k, 1_k\}, \quad \forall \pi \in \Omega_2. \quad (126)$$

In the expression of $G_{\sigma, \epsilon}$, half of the terms have coefficient 1 and the other half have coefficient -1 . To understand which terms have coefficient 1 and which have coefficient -1 , consider a special $\epsilon = (-1, 1, \dots, 1)$ (i.e., $\epsilon_1 = -1$ and all other $\epsilon_i = 1$). A term with coefficient -1 has the form $P_{\sigma_1} P_{\pi_2} \dots P_{\pi_{n-1}} P_{\sigma_n}$ or $P_{\sigma_n} P_{\pi_2} \dots P_{\pi_{n-1}} P_{\sigma_1}$ (i.e., with an indicator vector whose first element $\delta_1 = 1$), and a term with coefficient 1 has the form $P_{\sigma_1} P_{\pi_{n-1}} \dots P_{\pi_2} P_{\sigma_1}$ or $P_{\sigma_n} P_{\pi_{n-1}} \dots P_{\pi_2} P_{\sigma_n}$ (i.e., with an indicator vector whose first element $\delta_1 = 0$). We can see that the coefficient is in fact $\epsilon_1^{\delta_1}$. For general $\epsilon \in \Lambda$ and $\pi \in \Omega$, the coefficient of $P_{\pi_n} P_{\pi_{n-1}} \dots P_{\pi_2} P_{\pi_1}$ in $G_{\sigma, \epsilon}$ is $(\epsilon_1)^{\delta_1} \dots (\epsilon_k)^{\delta_k}$, where $\delta = \delta(\pi)$ is defined as in (125). We can then write the expression of $G_{\sigma, \epsilon}$ as

$$G_{\sigma, \epsilon} = \sum_{\pi \in \Omega} \epsilon_1^{\delta_1} \dots \epsilon_k^{\delta_k} P_{\pi_n} P_{\pi_{n-1}} \dots P_{\pi_1}.$$

Summing up this relation for all ϵ in Λ_k , we have

$$\sum_{\epsilon \in \Lambda_k} G_{\sigma, \epsilon} = \sum_{\epsilon \in \Lambda_k} \sum_{\pi \in \Omega} \epsilon_1^{\delta_1} \dots \epsilon_k^{\delta_k} P_{\pi_n} P_{\pi_{n-1}} \dots P_{\pi_1} = \sum_{\pi \in \Omega} P_{\pi_n} P_{\pi_{n-1}} \dots P_{\pi_1} \left(\sum_{\epsilon \in \Lambda_k} \epsilon_1^{\delta_1} \dots \epsilon_k^{\delta_k} \right). \quad (127)$$

Note that in this expression, $\delta_1, \dots, \delta_k$ depend on π .

Denote $0_k = (0, 0, \dots, 0) \in \mathbb{R}^k$, $1_k = (1, 1, \dots, 1) \in \mathbb{R}^k$. Define

$$g_k(\delta) \triangleq \sum_{\epsilon \in \Lambda_k} \epsilon_1^{\delta_1} \dots \epsilon_k^{\delta_k}, \quad h_k(\delta) \triangleq \sum_{\epsilon \in \Lambda_k^c} \epsilon_1^{\delta_1} \dots \epsilon_k^{\delta_k}.$$

For any $\delta \neq 0_k$, we have $g_k(\delta) + h_k(\delta) = \sum_{\epsilon \in \{1, -1\}^k} \epsilon_1^{\delta_1} \dots \epsilon_k^{\delta_k} = (1^{\delta_1} + (-1)^{\delta_1}) \dots (1^{\delta_k} + (-1)^{\delta_k}) = 0$, thus

$$h_k(\delta) = -g_k(\delta), \quad \forall \delta \neq 0_k. \quad (128)$$

It is easy to see that

$$\frac{1}{|\Lambda_k|} g_k(\delta) = \begin{cases} 1, & \delta = (0, 0, \dots, 0), \\ -1, & \delta = (1, 1, \dots, 1). \end{cases} \quad (129)$$

We will prove that for any $\delta \notin \{0_k, 1_k\}$,

$$g_k(\delta) = \sum_{\epsilon \in \Lambda_k} \epsilon_1^{\delta_1} \dots \epsilon_k^{\delta_k} = 0. \quad (130)$$

We prove (130) by induction on k . When $k = 2$, $\Lambda_2 = \{(-1, 1), (1, -1)\}$, we have:

$$\begin{aligned} \text{when } \delta = (0, 1), \quad g_2(\delta) &= (-1)^0 1^1 + 1^0 (-1)^1 = 1 - 1 = 0, \\ \text{when } \delta = (1, 0), \quad g_2(\delta) &= (-1)^1 1^0 + 1^1 (-1)^0 = -1 + 1 = 0. \end{aligned}$$

Assume (130) holds for $k - 1$, that is,

$$g_{k-1}(\hat{\delta}) = 0, \quad \forall \hat{\delta} \in \{0, 1\}^{k-1} \setminus \{0_{k-1}, 1_{k-1}\}. \quad (131)$$

According to (128), we have

$$h_{k-1}(\hat{\delta}) = 0, \quad \forall \hat{\delta} \in \{0, 1\}^{k-1} \setminus \{0_{k-1}, 1_{k-1}\}. \quad (132)$$

Now consider k . Because $\delta \neq 0_k$, there must exist some j such that $\delta_j = 1$; without loss of generality, we assume

$$\delta_k = 1. \quad (133)$$

Partition Γ_k into two sets:

$$\Lambda_{k,1} \triangleq \{\epsilon \in \Lambda_k \mid \epsilon_k = 1\}, \quad \Lambda_{k,2} \triangleq \{\epsilon \in \Lambda_k \mid \epsilon_k = -1\}. \quad (134)$$

If ϵ contains an odd number of -1 and the last element $\epsilon_k = 1$ (or $\epsilon_k = -1$), then the first $k-1$ elements contain an odd (or even) number of -1 . Thus

$$\Lambda_{k,1} = \{(\hat{\epsilon}, 1) \mid \hat{\epsilon} \in \Lambda_{k-1}\}, \quad \Lambda_{k,2} = \{(\hat{\epsilon}, -1) \mid \hat{\epsilon} \in \Lambda_{k-1}^c\}.$$

Split $g_k(\delta)$ into two parts $g_k(\delta) = g_{k,1}(\delta) + g_{k,2}(\delta)$, where

$$g_{k,1}(\delta) = \sum_{\epsilon \in \Lambda_{k,1}} \epsilon_1^{\delta_1} \dots \epsilon_k^{\delta_k}, \quad g_{k,2}(\delta) = \sum_{\epsilon \in \Lambda_{k,2}} \epsilon_1^{\delta_1} \dots \epsilon_k^{\delta_k}.$$

Denote $\hat{\delta} = (\delta_1, \dots, \delta_{k-1})$. We already assume $\delta \neq 1_k$ and $\delta_k = 1$, so we know

$$\hat{\delta} \neq 1_{k-1}. \quad (135)$$

But it is possible that $\hat{\delta} = 0_{k-1}$. Consider two cases.

Case 1. $\hat{\delta} = 0_{k-1}$, That is, $\delta = (0_{k-1}, 1)$.

In this case

$$\begin{aligned} g_{k,1}(\delta) &= \sum_{\epsilon \in \Lambda_{k,1}} \epsilon_1^{\delta_1} \dots \epsilon_k^{\delta_k} = \sum_{\epsilon \in \Lambda_{k,1}} \epsilon_1^0 \dots \epsilon_{k-1}^0 \epsilon_k^1 = \sum_{\epsilon \in \Lambda_{k,1}} \epsilon_k^1 = \sum_{\epsilon \in \Lambda_{k,1}} 1^1 = |\Lambda_{k,1}| = |\Lambda_{k-1}|, \\ g_{k,2}(\delta) &= \sum_{\epsilon \in \Lambda_{k,2}} \epsilon_1^{\delta_1} \dots \epsilon_k^{\delta_k} = \sum_{\epsilon \in \Lambda_{k,2}} \epsilon_1^0 \dots \epsilon_{k-1}^0 \epsilon_k^1 = \sum_{\epsilon \in \Lambda_{k,2}} \epsilon_k^1 = \sum_{\epsilon \in \Lambda_{k,2}} (-1)^1 = -|\Lambda_{k,2}| = -|\Lambda_{k-1}^c|, \end{aligned}$$

thus

$$g_k(\delta) = g_{k,1}(\delta) + g_{k,2}(\delta) = |\Lambda_{k-1}| - |\Lambda_{k-1}^c| = 0,$$

where the last step is due to (124).

Case 2. $\hat{\delta} \neq 0_{k-1}$.

Together with (135), we have

$$\hat{\delta} \notin \{0_{k-1}, 1_{k-1}\}.$$

This enables us to apply the induction hypothesis (131) and its corollary (132). In fact,

$$\begin{aligned} g_{k,1}(\delta) &= \sum_{\epsilon \in \Lambda_{k,1}} \epsilon_1^{\delta_1} \dots \epsilon_k^{\delta_k} \stackrel{(133),(134)}{=} \sum_{\epsilon \in \Lambda_{k,1}} \epsilon_1^{\delta_1} \dots \epsilon_{k-1}^{\delta_{k-1}} 1^1 = \sum_{\hat{\epsilon} \in \Lambda_{k-1}} \hat{\epsilon}_1^{\delta_1} \dots \hat{\epsilon}_{k-1}^{\delta_{k-1}} = g_{k-1}(\hat{\delta}) \stackrel{(131)}{=} 0, \\ g_{k,2}(\delta) &= \sum_{\epsilon \in \Lambda_{k,2}} \epsilon_1^{\delta_1} \dots \epsilon_k^{\delta_k} \stackrel{(133),(134)}{=} \sum_{\epsilon \in \Lambda_{k,2}} \epsilon_1^{\delta_1} \dots \epsilon_{k-1}^{\delta_{k-1}} (-1)^1 = - \sum_{\hat{\epsilon} \in \Lambda_{k-1}^c} \hat{\epsilon}_1^{\delta_1} \dots \hat{\epsilon}_{k-1}^{\delta_{k-1}} = h_{k-1}(\hat{\delta}) \stackrel{(132)}{=} 0. \end{aligned}$$

Thus $g_k(\delta) = g_{k,1}(\delta) + g_{k,2}(\delta) = 0$.

In both cases, we have proved $g_k(\delta) = 0$, which finishes the induction step. Therefore (130) holds for any k .

Next, we analyze the sum $\sum_{\epsilon \in \Lambda_k} G_{\sigma, \epsilon}$. According to (127), we have

$$\begin{aligned} \sum_{\epsilon \in \Lambda_k} G_{\sigma, \epsilon} &= \sum_{\pi \in \Omega} P_{\pi_n} P_{\pi_{n-1}} \dots P_{\pi_1} \left(\sum_{\epsilon \in \Lambda_k} \epsilon_1^{\delta_1} \dots \epsilon_k^{\delta_k} \right) \\ &= \sum_{\pi \in \Omega} P_{\pi_n} P_{\pi_{n-1}} \dots P_{\pi_1} g_k(\delta(\pi)) \\ &\stackrel{(i)}{=} \sum_{\pi \in \Omega_0} P_{\pi_n} P_{\pi_{n-1}} \dots P_{\pi_1} \cdot g_k(0_k) + \sum_{\pi \in \Omega_1} P_{\pi_n} P_{\pi_{n-1}} \dots P_{\pi_1} \cdot g_k(1_k) + \sum_{\pi \in \Omega_2} P_{\pi_n} P_{\pi_{n-1}} \dots P_{\pi_1} \cdot g_k(\delta(\pi)) \\ &\stackrel{(ii)}{=} \sum_{\pi \in \Omega_0} P_{\pi_n} P_{\pi_{n-1}} \dots P_{\pi_1} \cdot |\Gamma_k| + \sum_{\pi \in \Omega_1} P_{\pi_n} P_{\pi_{n-1}} \dots P_{\pi_1} \cdot (-1)|\Gamma_k| + \sum_{\pi \in \Omega_2} P_{\pi_n} P_{\pi_{n-1}} \dots P_{\pi_1} \cdot 0 \\ &= |\Gamma_k| \left(\sum_{\pi \in \Omega_0} P_{\pi_n} P_{\pi_{n-1}} \dots P_{\pi_1} - \sum_{\pi \in \Omega_1} P_{\pi_n} P_{\pi_{n-1}} \dots P_{\pi_1} \right), \end{aligned}$$

where (i) is due to (126) and (ii) is due to (129), (130). According to (123), any $G_{\sigma,\epsilon} \geq 0$; thus the above relation implies the following important relation

$$\sum_{\pi \in \Omega_0} P_{\pi_n} P_{\pi_{n-1}} \dots P_{\pi_1} \geq \sum_{\pi \in \Omega_1} P_{\pi_n} P_{\pi_{n-1}} \dots P_{\pi_1}. \quad (136)$$

Note that this relation holds for a fixed permutation σ and the corresponding set $\Omega_0 = \Omega(\sigma)$ and $\Omega_1(\sigma)$. Each $\pi \in \Omega_0$ corresponds to a k -permutation χ of $(12 \dots n)$ determined by $\pi = (\chi_1 \dots \chi_{k-1} \chi_k \chi_k \chi_{k-1} \dots \chi_1)$, and each $\pi \in \Omega_1$ corresponds to a permutation of $(12 \dots n)$. We rewrite (136) as

$$\sum_{\pi \in \Omega_0(\sigma)} P_{\pi_n} P_{\pi_{n-1}} \dots P_{\pi_1} \geq \sum_{\pi \in \Omega_1(\sigma)} P_{\pi_n} P_{\pi_{n-1}} \dots P_{\pi_1},$$

and summing up this relation for all possible permutations $\sigma \in \Gamma$ leads to

$$E_{\chi \in \Gamma_k} (P_{\chi_1} \dots P_{\chi_{k-1}} P_{\chi_k} P_{\chi_k} P_{\chi_{k-1}} \dots P_{\chi_1}) \geq E_{\sigma \in \Gamma} (P_{\sigma_n} P_{\sigma_{n-1}} \dots P_{\sigma_1}),$$

which is exactly (122).

It remains to prove

$$E_{\chi \in \Gamma_k} (P_{\chi_1} \dots P_{\chi_{k-1}} P_{\chi_k} P_{\chi_{k-1}} \dots P_{\chi_1}) \leq \frac{1}{n} \sum_i P_i. \quad (137)$$

In fact, for any positive-semidefinite matrix X and any symmetric matrix Y , we have $YXY = Y^T XY \geq 0$. Applying this fact $k-1$ times leads to (137).

Combining (122) and (137), we immediately obtain the desired result (119) for the case $n = 2k$.

The case that $n = 2k - 1$ is an odd number is almost the same, except that the key quantity $G_{\sigma,\epsilon}$ is now defined as

$$G_{\sigma,\epsilon} \triangleq (P_{\sigma_n} + \epsilon_1 P_{\sigma_1}) \dots (P_{\sigma_{k+1}} + \epsilon_{k-1} P_{\sigma_{k-1}}) P_{\sigma_k} (P_{\sigma_{k+1}} + \epsilon_{k-1} P_{\sigma_{k-1}}) \dots (P_{\sigma_n} + \epsilon_1 P_{\sigma_1}). \quad (138)$$

In words, we pair P_{σ_i} with $P_{\sigma_{n+1-i}}$ for $i = 1, \dots, k-1$ and leave P_{σ_k} alone (following the same rule it would have been paired with itself). The rest of the proof is almost the same as the even case, so we skip it. Q.E.D.

8. Numerical Experiments

In this section, we test the performance of cyclic ADMM and RP-ADMM for solving various kinds of linear systems. As a benchmark, we also test the gradient descent method (GD) with a constant stepsize $\alpha = 1/\lambda_{\max}(A'A)$ for solving the least square problem $\min_{x \in \mathbb{R}^N} \|Ax - b\|^2/2$. Of course there are many other advanced algorithms for solving the least square problem, such as the conjugate gradient method, but we do not consider them because our focus is on testing the two ADMM algorithms. These two ADMM algorithms can be used to solve far more general problems than just linear systems, and we believe that the performance comparison for solving linear systems can shed light on more general scenarios.

In the numerical experiments, we set $b = 0$, thus the unique optimal solution is $x^* = 0$. The coefficient matrix A will be generated according to one of the random distributions below:

- Gauss: independent Gaussian entries $A_{ij} \sim \mathcal{N}(0, 1)$.
- Log-normal: independent log-normal entries $A_{ij} \sim \exp(\mathcal{N}(0, 1))$.
- Uniform: each entry is drawn independently from a uniform distribution on $[0, 1]$.
- Circulant Hankel: circulant Hankel matrix with independent standard Gaussian entries. More specifically, generate $\delta_1, \delta_2, \dots, \delta_N \sim \mathcal{N}(0, 1)$ and let $A_{ij} = \delta_{i+j-1}$ (define $\delta_k = \delta_{k-N}$ if $k > N$). Note that the entries of the circulant Hankel matrix are not independent because one δ_i can appear in multiple positions.

For the two ADMM algorithms, we only consider the n -coordinate versions (i.e., each block consists of only one coordinate). We let the three tested algorithms start from the same random initial point $y^0 = [x^0; \lambda^0]$ (GD will start from x^0). To measure the performance, we define the epoch complexity k to be the minimum k so that the relative error

$$\|Ax^k - b\|/\|Ax^0 - b\| < \epsilon,$$

where ϵ is a desired accuracy (we consider 10^{-2} and 10^{-3} ¹¹). For the two ADMM algorithms, one epoch refers to one round of primal and dual steps; for GD, one epoch refers to one gradient step. The total computation

Table 2. Results of Solving Linear Systems by Cyclic ADMM, RP-ADMM, and GD

N	Diverg. ratio	Epochs for $\epsilon = 0.01$			Epochs for $\epsilon = 0.001$		
		CycADMM ¹³	RPADMM	GD	CycADMM	RPADMM	GD
Gaussian							
3	0.7%	1.4e01	3.4e01	5.0e01	3.2e01	8.8e01	1.4e02
10	1.1%	4.1e01	1.8e02	2.0e02	1.2e02	1.1e03	1.5e03
100	3%	1.7e02	4.3e02	3.6e02	1.0e03	7.4e03	6.5e03
Log-normal							
3	0.8%	1.5e01	3.7e01	5.7e01	3.3e01	9.6e01	1.7e02
10	39.2%	1.2e02	3.4e02	6.4e02	3.2e02	2.4e03	6.3e03
100	100%	N/A	5.5e02	5.4e03	N/A	8.8e03	1.0e05
Uniform							
3	3.2%	2.8e01	7.4e01	1.5e02	7.0e01	2.6e02	6.0e02
10	83.0%	2.1e02	4.1e02	1.2e03	5.2e02	3.0e03	9.1e03
100	100%	N/A	9.1e02	1.4e04	N/A	1.4e04	9.7e04
Circulant Hankel							
3	5.6%	1.2e01	1.7e01	1.5e01	1.7e01	2.8e01	2.6e01
10	54.3%	4.2e01	6.0e01	6.5e01	7.5e01	1.3e02	1.7e02
100	100%	N/A	1.3e02	1.7e02	N/A	2.9e02	6.5e02

Note. For the two ADMM variants, one epoch refers to one round of primal and dual steps; for GD, one epoch refers to one gradient step.

time should be proportional to the epoch complexity because GD and the two ADMM variants have similar per-epoch cost¹²: a gradient descent step $x^{k+1} = x^k - \alpha A^T(Ax - b)$ contains two matrix-vector multiplications and thus takes time $2N^2 + O(N)$, and an ADMM round also takes time $2N^2 + O(N)$ (the primal update step of ADMM takes time $2N^2 + O(N)$, and the dual update step of ADMM takes time $O(N)$). We test 1,000 random instances for $N \in \{3, 10\}$ and 300 random instances for $N = 100$ and record the geometric mean of the number of epochs. In Table 2, “Diverg. ratio” represents the percentage of tested instances for which cyclic ADMM diverges, and “CycADMM” represents “cyclic ADMM” (note that RP-ADMM converges in all instances we tested, so its divergence ratio is 0). Note that for cyclic ADMM we only report the epoch complexity when it converges, while for RPADMM and GD we report the epoch complexity in all tested instances. If restricting to the successful instances of cyclic ADMM, we find that the epoch complexity of RPADMM does not change too much, whereas the epoch complexity of GD will be reduced (significantly in some settings).

The simulation results are summarized in Table 2. The main observations from the simulation are:

- For all random distributions of A we tested, cyclic ADMM does not always converge even when N is fixed to be 3. For $N = 100$ and many random distributions, cyclic ADMM diverges with probability 1. This means that the divergence of cyclic ADMM is not merely a “worst-case” phenomenon but actually quite common. When the dimension increases, the divergence ratio will increase.
- For standard Gaussian entries, cyclic ADMM converges with high probability. When cyclic ADMM converges, it converges faster than RP-ADMM and sometimes much faster.
- RPADMM typically converges faster than the basic gradient descent method and sometimes more than 10 times faster.

We have also tested BR-ADMM for solving the same problems, though the simulation results are not listed in Table 2. As expected, BR-ADMM also always converges for solving these linear systems. The convergence speed is usually slower than RP-ADMM. Nevertheless, BR-ADMM can save some sampling time compared with RP-ADMM and may be more favorable if random permutation is not available owing to system architecture constraint. The detailed comparison of BR-ADMM and RP-ADMM and the design of other randomized schemes or even deterministic schemes that outperform RP schemes are left as future work.

9. Concluding Remarks

In this paper, we prove the expected convergence of randomly permuted ADMM (RP-ADMM) for solving a nonsingular square system of equations (extension to nonsquare systems is straightforward). We also prove a bound on the expected convergence rate of RP-ADMM for solving linear systems and the expected convergence rate of RP-BCD for solving quadratic problems. The motivation is to resolve the divergence issue of cyclic multiblock ADMM. Our result shows that RP-ADMM may serve as a simple remedy, and we expect RP-ADMM to be one of the important solvers in large-scale optimization. One interesting finding along the path is that the update matrix of RP-BCD has spectrum lying in $(-1/3, 1)$ instead of the commonly seen $(-1, 1)$.

Random permutation is widely known to be empirically better than independently randomized versions, but little was known about its theoretical properties in general. Note that most existing analyses of BCD (e.g., Beck and Tetruashvili [1], Sun and Hong [39], Tseng [42]) are applicable to both the cyclic update rule and the random permutation update rule. However, in light of a recent study that established an up to $O(n^2)$ gap between cyclic CD and R-CD (Sun and Ye [41]), it is unlikely that RP-CD will have the same rate as cyclic CD. Our result in this paper established, for the first time, an $O(n)$ gap between RP-CD and cyclic-CD for general quadratic problems, making some progress toward the conjecture that RP-CD is faster than R-CD.

We emphasize that the convergence speed analysis of large-scale optimization has mostly been limited to independently randomized update order in the past decade. Going beyond independent randomized order is an important topic for enlarging the scope of large-scale optimization. Not only is the analysis of random permutation quite challenging, even the analysis of the most classic cyclic order is highly nontrivial (Sun and Ye [41]). There are quite a few open questions regarding the convergence rate of non-independent-randomized order. Regarding the random permutation order, a very interesting open question is the worst-case convergence rate of RP-BCD for quadratic problems. Owing to the close relation with matrix AM-GM inequality, this problem seems to be a quite fundamental problem. Moving to ADMM, the similar questions about the convergence rate of various variants of ADMM, including RP-ADMM and BR-ADMM, are also open.

Acknowledgments

The authors thank an anonymous reviewer for many helpful comments on the manuscript, which enabled improvement of the presentation of the paper. This paper is a strengthened version of a previous technical report, “On the expected convergence of randomly permuted ADMM,” with several new results, mainly the ones on the expected convergence rate. Part of this work was done when R. Sun was visiting Stanford University as a student of University of Minnesota and when R. Sun was a postdoctoral student at Stanford University.

Endnotes

¹ Rigorously speaking, these two bounds are not directly comparable because the result for the randomized version only holds with high probability, whereas the result for the cyclic version always holds; in any event, this $O(n^2)$ gap is still meaningful if ignoring this difference between deterministic and randomized algorithm.

² This paper appeared after the first version of the present paper.

³ There has been some effort in using random matrix theory to tackle this problem, but no progress has been reported to our knowledge. This is partially because the desired result seems to be rather tight, such that even a small relaxation can lead to failure.

⁴ Suppose A is an $m \times n$ full column-rank matrix, where $m \geq n$, and the system $Ax = b$ is feasible. The update formula is $y^{k+1} = (I - L_g^{-1} A^T A) y^k$, which is same as the update formula for solving a square system of equations $\tilde{A}x = b$, where $\tilde{A} \in \mathbb{R}^{n \times n}$ is the square root matrix of the matrix $A^T A \in \mathbb{R}^{n \times n}$. Now the matrix $\tilde{A}$ is a square invertible matrix; thus by applying the result for square system of equations, we can obtain the convergence of the sequence $\phi^k = E(y^k)$.

⁵ Rigorously speaking this is not a fair comparison because the complexity of C-CD is deterministic complexity.

⁶ The computation complexity equals the iteration complexity times the per-iteration cost. We do not present iteration complexity because there may be confusion about whether “one iteration” means n coordinate updates or 1 coordinate update. Presenting iteration complexity is better if one considers a general convex problem, but then one needs to discuss the per-iteration cost. We are considering quadratic problems throughout the paper, so we believe it is more clear to stick to computation complexity.

⁷ Here, the relative accuracy ϵ means $\|E(x^k) - x^*\|/\|x^0 - x^*\|$ or $\|E(f(x^k)) - f^*\|/\|f(x^0) - f^*\|$.

⁸ Note that this “improvement” is valid when the convergence rate is characterized by only κ_{CD} and n . It is common to use other parameters, such as the maximum eigenvalue, to characterize the convergence rate (see Sun and Ye [41] for a detailed discussion), and our result here does not provide improvement for other kinds of convergence rate.

⁹ For the purpose of proving Theorem 2, we do not need to prove this direction. Here we present the proof because it is quite straightforward and makes the result more comprehensive.

¹⁰ A more detailed argument is as follows. Because $-1/3 \leq \hat{Z}_k$, we can let $\beta_1 = -1/3 + \epsilon$ for a sufficiently small positive number ϵ , while keeping $\alpha_1 = 0, \alpha_n = 1, \beta_n = 1$. The set (76) now becomes $\{0, 0, -2/3 + 2\epsilon, 2, -\frac{(4/3-\epsilon)^2}{4(2/3+\epsilon)}\}$. Both $-2/3 + 2\epsilon$ and $-\frac{(4/3-\epsilon)^2}{4(2/3+\epsilon)}$ are strictly larger than $2/3$, thus the extreme value $-2/3$ cannot be achieved. By a similar argument the other extreme value 2 also cannot be achieved.

¹¹ For high accuracy such as $\epsilon = 10^{-6}$, it takes too many epochs for the algorithms to converge when $n = 100$ because most matrices we generated are highly ill-conditioned, so we do not report the results. Based on the limited experiments for high accuracy, similar gaps between RP-ADMM and GD are observed.

¹² In MATLAB simulations each epoch of GD takes much less time than a round of ADMM because MATLAB implements matrix operations much faster than a “for” loop. For a more fair central processing unit time comparison, one should use other programming languages such as C.

¹³ For cyclic ADMM, only record the iteration complexity in convergent instances.

References

- [1] Beck A, Tetrushvili L (2013) On the convergence of block coordinate descent type methods. *SIAM J. Optim.* 23(4):2037–2060.
- [2] Blatt D, Hero AO, Gauchman H (2007) A convergent incremental gradient method with a constant step size. *SIAM J. Optim.* 18(1):29–51.
- [3] Boyd S, Parikh N, Chu E, Peleato B, Eckstein J (2011) Distributed optimization and statistical learning via the alternating direction method of multipliers. *Foundations Trends Machine Learning* 3(1):1–122.
- [4] Cai X, Han D, Yuan X (2016) The direct extension of ADMM for three-block separable convex minimization models is convergent when one function is strongly convex. *Math. Programming* 155(1):57–59.
- [5] Chan TF, Glowinski R (1978) *Finite Element Approximation and Iterative Solution of a Class of Mildly Non-linear Elliptic Equations* (Computer Science Department, Stanford University, Stanford, CA).
- [6] Chen C, He B, Ye Y, Yuan X (2014) The direct extension of ADMM for multi-block convex minimization problems is not necessarily convergent. *Math. Programming* 155(1/2):1–23.
- [7] Chen C, Li M, Liu X, Ye Y (2017) Extended ADMM and BCD for nonseparable convex minimization models with quadratic coupling terms: convergence analysis and insights. Working paper, Nanjing University, Nanjing, China.
- [8] Chen C, Shen Y, You Y (2013) On the convergence analysis of the alternating direction method of multipliers with three blocks. *Abstract and Applied Analysis*, vol. 2013 (Hindawi Publishing, London).
- [9] Davis D, Yin W (2017) A three-operator splitting scheme and its optimization applications. *Set Valued Var. Anal.* 25(4):829–858.
- [10] Deng W, Lai MJ, Peng Z, Yin W (2017) Parallel multi-block admm with $o(1/k)$ convergence. *J. Scientific Comput.* 71(2):712–736.
- [11] Gabay D, Mercier B (1976) A dual algorithm for the solution of nonlinear variational problems via finite element approximation. *Comput. Math. Appl.* 2(1):17–40.
- [12] Glowinski R, Marroco A (1975) Approximation par éléments finis d'ordre un, et la résolution, par pénalisation-dualité d'une classe de problèmes de dirichlet non linéaires. *ESAIM: Math. Model. Numerical Anal.* 9(R2):41–76.
- [13] Gürbüzbalaban M, Ozdaglar A, Parrilo P (2015) Why random reshuffling beats stochastic gradient descent. Preprint, submitted October 29, <https://arXiv:1510.08560>.
- [14] Han D, Yuan X (2012) A note on the alternating direction method of multipliers. *J. Optim. Theory Appl.* 155(1):227–238.
- [15] Han D, Yuan X, Zhang W (2014) An augmented Lagrangian based parallel splitting method for separable convex minimization with applications to image processing. *Math. Comput.* 83(289):2263–2291.
- [16] He B, Hou L, Yuan X (2015) On full Jacobian decomposition of the augmented Lagrangian method for separable convex programming. *SIAM J. Optim.* 25(4):2274–2312.
- [17] He B, Tao M, Yuan X (2012) Alternating direction method with Gaussian back substitution for separable convex programming. *SIAM J. Optim.* 22(2):313–340.
- [18] He B, Tao M, Yuan X (2017) Convergence rate and iteration complexity on the alternating direction method of multipliers with a substitution procedure for separable convex programming. *Math. Oper. Res.* 42(3):662–691.
- [19] He B, Xu HK, Yuan X (2016) On the proximal Jacobian decomposition of ALM for multiple-block separable convex minimization problems and its relationship to ADMM. *J. Sci. Comput.* 66(3):1204–1217.
- [20] Hong M, Chang T-H, Wang X, Razaviyayn M, Ma S, Luo Z-Q (2014) A block successive upper bound minimization method of multipliers for linearly constrained convex optimization. Preprint, submitted January 28, <https://arXiv:1401.7079>.
- [21] Hong M, Luo Z-Q (2012) On the linear convergence of the alternating direction method of multipliers. Preprint, submitted August 20, <https://arXiv:1208.3922>.
- [22] Kittaneh F (2006) Spectral radius inequalities for Hilbert space operators. *Proc. Amer. Math. Soc.* 134(02):385–390.
- [23] Lee C-P, Wright SJ (2016) Random permutations fix a worst case for cyclic coordinate descent. Preprint, submitted July 28, <https://arXiv:1607.08320>.
- [24] Leventhal D, Lewis AS (2010) Randomized methods for linear constraints: Convergence rates and conditioning. *Math. Oper. Res.* 35(3):641–654.
- [25] Li X, Sun D, Toh K-C (2014) A Schur complement based semi-proximal ADMM for convex quadratic conic programming and extensions. *Math. Programming* 155(1/2):333–373.
- [26] Li M, Sun D, Toh K-C (2015) A convergent 3-block semi-proximal ADMM for convex minimization problems with one strongly convex block. *Asia Pacific J. Oper. Res.* 32(04):1550024.
- [27] Lin T, Ma S, Zhang S (2014) On the convergence rate of multi-block ADMM. Preprint, submitted August 19, <https://arXiv:1408.4265>.
- [28] Lin T, Ma S, Zhang S (2014) On the global linear convergence of the ADMM with multi-block variables. Preprint, submitted August 19, <https://arXiv:1408.4266>.
- [29] Lin T, Ma S, Zhang S (2015) Iteration complexity analysis of multi-block ADMM for a family of convex minimization without strong convexity. Preprint, submitted April 13, <https://arXiv:1504.03087>.
- [30] Nesterov Y (2012) Efficiency of coordinate descent methods on huge-scale optimization problems. *SIAM J. Optim.* 22(2):341–362.
- [31] Recht B, Ré C (2012) Beneath the valley of the noncommutative arithmetic-geometric mean inequality: Conjectures, case-studies, and consequences. Preprint, submitted February 19, <https://arXiv:1202.4184>.
- [32] Recht B, Ré C (2013) Parallel stochastic gradient algorithms for large-scale matrix completion. *Math. Programming Comput.* 5(2):201–226.
- [33] Schmidt M, Le Roux N, Bach F (2013) Minimizing finite sums with the stochastic average gradient. Preprint, submitted September 10, <https://arXiv:1309.2388>.
- [34] Shalev-Shwartz S, Zhang T (2013) Stochastic dual coordinate ascent methods for regularized loss. *J. Machine Learning Res.* 14(1):567–599.
- [35] Shi HJM, Tu S, Xu Y, Yin W (2016) A primer on coordinate descent algorithms. Preprint, submitted September 30, <https://arXiv:1610.00040>.
- [36] Strang WG (1962) Eigenvalues of Jordan products. *Amer. Math. Monthly* 69(1):37–40.
- [37] Sun D, Toh KC, Yang L (2014) A convergent proximal alternating direction method of multipliers for conic programming with 4-block constraints. Preprint, submitted April 22, <https://arXiv:1404.5378>.
- [38] Sun R (2015) Matrix completion via nonconvex factorization: Algorithms and theory. Unpublished doctoral dissertation, University of Minnesota, Minneapolis.

- [39] Sun R, Hong M (2015) Improved iteration complexity bounds of cyclic block coordinate descent for convex problems. Cortes C, Lawrence ND, Lee DD, Sugiyama M, Garnett R, eds. *Advances in Neural Information Processing Systems 28* (MIT Press, Cambridge, MA), 1306–1314.
- [40] Sun R, Luo ZQ, Ye Y (2015) On the expected convergence of randomly permuted ADMM. Preprint, submitted March 22, <https://arXiv:1503.06387>.
- [41] Sun R, Ye Y (2016) Worst-case complexity of cyclic coordinate descent: $O(n^2)$ gap with randomized version. Preprint, submitted April 25, <https://arXiv:1604.07130>.
- [42] Tseng P (2001) Convergence of a block coordinate descent method for nondifferentiable minimization. *J. Optim. Theory Appl.* 109(3):475–494.
- [43] Wang H, Banerjee A, Luo Z-Q (2014) Parallel direction method of multipliers. *Proc. 27th Internat. Conf. Neural Inform. Processing Systems*, vol. 1 (MIT Press, Cambridge, MA), 181–189.
- [44] Wright SJ (2015) Coordinate descent algorithms. *Math. Programming* 151(1):3–34.
- [45] Wright SJ, Lee C-P (2017) Analyzing random permutations for cyclic coordinate descent. Preprint, submitted June 3, <https://arXiv:1706.00908>.