

A study of preconditioners for network interior point methods

Joaquim J. Júdice (joaquim.judice@co.it.pt)

Departamento de Matemática, Universidade de Coimbra, 3000 Coimbra, Portugal.

João Patricio (joao.patricio@co.it.pt)

Escola Superior de Tecnologia de Tomar, Instituto Politécnico de Tomar, 2300 Tomar, Portugal.

Luis F. Portugal *

Departamento de Matemática, Universidade de Coimbra, 3000 Coimbra, Portugal.

Mauricio G. C. Resende (mgcr@research.att.com)

Shannon Laboratory, AT&T Labs Research, Florham Park, NJ 07932 USA.

Geraldo Veiga (gveiga@att.com)

AT&T, Middletown, NJ 07748 USA.

Abstract. We study and compare preconditioners available for network interior point methods. We derive upper bounds for the condition number of the preconditioned matrices used in the solution of systems of linear equations defining the algorithm search directions. The preconditioners are tested using PDNET, a state-of-the-art interior point code for the minimum cost network flow problem. A computational comparison using a set of standard problems improves the understanding of the effectiveness of preconditioners in network interior point methods.

Keywords: Interior point method, linear programming, network flows, conjugate gradient, preconditioners, experimental testing of algorithms.

1. Introduction

A number of implementations of interior point methods for the minimum cost network flow (MCNF) problem have been proposed [16, 19, 20]. These procedures show that interior point methods may be competitive with the classical algorithms in several classes of the MCNF problem. The efficiency of these codes is heavily dependent on the method used to compute the search directions within the interior point algorithm framework.

Finding a search direction in interior point methods for linear programming involves the solution of a system of linear equations. In the most common variants, this results in a system of normal equations:

$$A\Theta A^T x = b, \quad (1)$$

where A is the $m \times n$ matrix of the coefficients and Θ is a $n \times n$ diagonal matrix of positive elements. Either a direct factorization or an iterative

* Luis F. Portugal passed away on July 21, 1998.



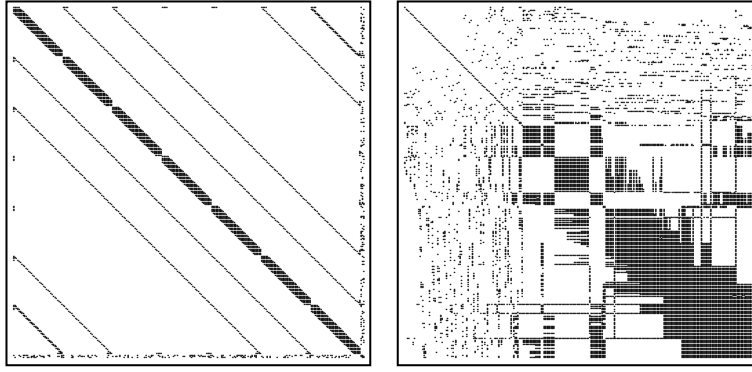


Figure 1. Sparsity patterns for $A\Theta A^T$ and $L + L^T$

method can be applied to solve (1). If we assume that A has full row rank, then $A\Theta A^T$ is a symmetric positive definite matrix. The most commonly used direct and iterative approaches to solve this system are, respectively, the Cholesky factorization and the preconditioned conjugate gradient (PCG) algorithm.

Even for some small scale MCNF problems, direct methods suffer from excessive memory requirements. This deficiency is illustrated in Figure 1 and Table I for a class of transshipment problems on a grid. Let L denote the triangular factor of $A\Theta A^T$. Figure 1 displays the nonzero structures of the matrices $A\Theta A^T$ and $L + L^T$ after minimum degree ordering, illustrating the large amount of fill-in that occurs during the factorization.

In the particular case of the MCNF problem, direct factorization methods are not competitive with the PCG algorithm. Efficient preconditioning strategies for $A\Theta A^T$ have been devised, making the PCG method a powerful approach in this context. Table I compares CPLEX and BAROPT (a state-of-the-art primal-dual interior point algorithm based on Cholesky factorization) and PDNET (a network flow conjugate gradient based primal-dual interior point algorithm). In this table, the label *nodes* indicates the number of nodes (the number of arcs is $8 \times \text{nodes}$), *nzlw* and *nz* denote the numbers of nonzero elements of the lower triangular part of $A\Theta A^T$ and factor L , respectively, *DSW* represents the size of the dense window [1], *ITRS* and *TIME* represent the number of interior-point iterations and the CPU time in seconds taken by CPLEX and PDNET, respectively, and *CGITRS* represents the total number of the conjugate gradient iterations required by PDNET. The experiment was done on a Silicon Graphics Challenge (196 MHz MIPS R10000 processor).

Table I shows the rapidly increasing difficulty encountered by the direct factorization method as the size of the network increases. Even for the small 4096 node instance, the amount of fill-in is overwhelming. Standard techniques such as the use of a dense window appear to be of little help. A

Table I. CPLEX & PDNET on transshipment on a grid

<i>nodes</i>	CPLEX					PDNET		
	<i>nzlw</i>	<i>DSW</i>	<i>nz</i>	<i>ITRS</i>	<i>TIME</i>	<i>ITRS</i>	<i>CGITRS</i>	<i>TIME</i>
256	1705	62	8.7×10^3	27	0.54	26	145	0.26
512	3505	209	4.8×10^4	20	2.84	33	164	0.71
1024	7004	82	4.8×10^4	44	5.38	33	137	1.59
2048	14048	113	1.3×10^6	54	19.57	35	211	4.51
4096	28652	1654	2.2×10^6	29	674.61	38	237	11.40
8192	57151	165	7.5×10^5	61	210.21	43	246	26.62
16384	112645	220	1.7×10^7	67	514.08	41	254	61.38
32768	229872	15306	1.6×10^9			67	1697	501.10
65536		did not run				47	370	593.26

conjugate gradient based method solves this same instance 60 times faster and is capable of solving much larger problems.

The key ingredients to construct an efficient conjugate gradient based interior point code are the choice of stopping criterion and preconditioners for the PCG algorithm. The stopping criterion should minimize the computational effort involved in computing search directions that guarantee global convergence for the interior point algorithms. This is the idea behind the truncated optimization methods [6, 7]. In [19], a truncated Newton interior point method for linear programming is described. PDNET is an implementation of this algorithm for the MCNF problem.

A preconditioning matrix M is used to indirectly solve (1) by transformed system $M^{-1}A\Theta A^\top x = M^{-1}b$. When no preconditioner is used, the efficiency and the convergence of the conjugate gradient method depends on the distribution of the eigenvalues of $A\Theta A^\top$. An appropriate starting solution usually results in well conditioned system matrices for the initial iterations of the interior point method. However, as the algorithm progresses, the values of the scaling diagonal Θ spread and the conditioning of $A\Theta A^\top$ degrades rapidly. An effective preconditioner must improve on the spectral properties of the original system matrix without adding significant computational effort to the PCG algorithm.

Several families of preconditioners have been proposed for the PCG applied to network interior point methods [16, 18, 21, 26]. In this paper, we review the diagonal [26], the maximum spanning tree (MST) [21], and the diagonally compensated maximum spanning tree (DCMST) [16] preconditioning procedures. We derive upper bounds for the condition numbers of the preconditioned matrices obtained with these schemes. PDNET is used

to compare the behavior of these preconditioning techniques in the initial, intermediate and final stages of the truncated Newton method.

Before concluding this introduction, we present some notation and outline the remainder of the paper. We denote a vector of ones by e . Let $x \in R^n$. We represent by X the $n \times n$ diagonal matrix having the elements of x in the diagonal. The number of elements of the set $\mathcal{H}$ is denoted by $|\mathcal{H}|$. Let B be a square matrix. $\underline{\lambda}(B)$, $\bar{\lambda}(B)$ and $\kappa(B)$, represent the smallest eigenvalue, the largest eigenvalue and the condition number of B , respectively. $\text{diag}(B)$ denotes the diagonal matrix having the diagonal elements of B in the diagonal.

In Section 2 we describe the truncated Newton algorithm for the linear optimization problem. In Section 3 the conjugate gradient algorithm is presented. In Section 4 we study the conditioning of $A\Theta A^\top$, describing preconditioners and derive upper bounds for the condition number of their corresponding matrices. The computational experience and the concluding remarks are reported in sections 6 and 7, respectively.

2. The truncated Newton algorithm

The minimum cost network flow problem is a linear optimization problem of the form

$$\min \{c^\top x \mid Ax = b; x + s = u; x, s \geq 0\}, \quad (2)$$

where A is an $m \times n$ matrix obtained from the node-arc incidence matrix after the redundant constraints are removed, and b , u , x and s are the vectors of supply/demand at the nodes, arc capacities, flow variables and primal slacks, respectively. The dual of (2) can be written as:

$$\max \{b^\top y - u^\top w \mid A^\top y - w + z = c; z, w \geq 0\}, \quad (3)$$

where y is the vector of dual variables and w and z are the vectors of dual slacks.

The Karush-Kuhn-Tucker conditions

$$\begin{aligned} Ax &= b, \\ x + s &= u, \\ A^\top y - w + z &= c, \\ Xz &= 0, \\ Sw &= 0, \\ x, s, w, z &\geq 0, \end{aligned} \quad (4)$$

are necessary and sufficient for (x, y, s, w, z) to be an optimal solution of the pair of problems (2)-(3). Primal-dual variants of interior point methods solve

(4) by following the central trajectory [15, 8] defined as

$$\begin{aligned}
 Ax &= b, \\
 x + s &= u, \\
 A^\top y - w + z &= c, \\
 Xz &= \mu e, \\
 Sw &= \mu e, \\
 x, s, w, z &\geq 0,
 \end{aligned} \tag{5}$$

where $\mu > 0$ is the trajectory parameter. At each iteration, the interior point algorithm sets the value of the trajectory parameter and computes the Newton search direction used to obtain the next iterate toward the central trajectory. Convergence to an optimal solution is attained as $\mu \rightarrow 0$.

The truncated dual feasible, primal infeasible variant described in [19] stems from the work of Dembo et al. [6], Dembo and Steihaug [7] and [13]. Customary with minimum cost network flow problems, we assume the existence of strictly interior dual solutions. The iterates generated by the algorithm are required to satisfy the dual constraints and strict positivity of primal and dual variables, but not the primal equality constraints.

Following the derivation by Kojima, Megiddo and Mizuno [13], let $\beta_0, \beta_1, \beta_2, \beta_3, \gamma$ and γ_p be such that $0 < \gamma < 1$, $\gamma_p > 0$ and $0 \leq \beta_0 < \beta_1 < \beta_2 < \beta_3 < 1$. Furthermore, let ε_c and ε_p represent tolerances required for the optimal complementarity gap and primal feasibility, respectively. The sequence of solutions produced by the interior point algorithm belong to a given neighborhood of the central trajectory,

$$\mathcal{K} = \{(x, s, y, w, z) : \tag{6}$$

$$(x, s, y, w, z) > 0, x + s = u, A^\top y - w + z = c, \tag{6}$$

$$x_i z_i \geq \gamma \mu(x, s, w, z) \quad (i = 1, 2, \dots, n), \tag{7}$$

$$w_i s_i \geq \gamma \mu(x, s, w, z) \quad (i = 1, 2, \dots, n), \tag{8}$$

$$\mu(x, s, w, z) \geq \gamma_p \|Ax - b\| \text{ or } \|Ax - b\| \leq \varepsilon_p \}, \tag{9}$$

where the central trajectory parameter is defined as

$$\mu(x, s, w, z) = \beta_1 \frac{x^\top z + w^\top s}{2n}.$$

In the definition of $\mathcal{K}$, relation (6) states the feasibility requirements for the algorithm iterates. The inequalities in (7) prevent the generated sequence from reaching the feasible boundary before convergence. Relations (8-9) exclude the possibility of convergence to an infeasible complementary solution.

At iteration k , starting with an interior, but possibly primal infeasible solution, $(x^0, y^0, s^0, w^0, z^0) \in \mathcal{K}$, the algorithm computes an approximate Newton

direction $(\Delta x^k, \Delta y^k, \Delta s^k, \Delta w^k, \Delta z^k)$ by solving the linear system

$$A\Delta x^k = b - Ax^k + r^k, \quad (10)$$

$$\Delta x^k + \Delta s^k = 0, \quad (11)$$

$$A^\top \Delta y^k - \Delta w^k + \Delta z^k = 0, \quad (12)$$

$$Z^k \Delta x^k + X^k \Delta z^k = \mu_k e - X^k Z^k e, \quad (13)$$

$$W^k \Delta s^k + S^k \Delta w^k = \mu_k e - W^k S^k e. \quad (14)$$

As indicated by (10), the primal equality constraints are not satisfied by the current iterate and the residual vector r^k reflects the inexact computation of the Newton direction, satisfying

$$\|r^k\| \leq \beta_0 \|Ax^k - b\|. \quad (15)$$

Observe that setting $\beta_0 = 0$ is equivalent to the replacement of (10) by

$$A\Delta x^k = b - Ax^k,$$

resulting in the exact Newton method of Kojima, Megiddo and Mizuno [13].

The system of linear equations (10-14) can be rewritten in a more compact form, suitable . Let

$$\Theta^k = \left(Z^k (X^k)^{-1} + W^k (S^k)^{-1} \right)^{-1} \quad (16)$$

and

$$\bar{b} = -A\Theta^k \left(\mu_k (X^k)^{-1} e - \mu_k (S^k)^{-1} e - c + A^\top y^k \right) + (b - Ax^k), \quad (17)$$

allowing the original system to be solved in two stages. First, Δy^k is computed by solving the linear system

$$A\Theta^k A^\top \Delta y^k = \bar{b} + r^k. \quad (18)$$

Recovering Δx^k , Δs^k , Δw^k and Δz^k can be achieved by simple affine expressions,

$$\Delta x^k = \Theta^k A^\top \Delta y^k + \Theta^k (\mu_k (X^k)^{-1} e - \mu_k (S^k)^{-1} e - c + A^\top y^k),$$

$$\Delta s^k = -\Delta x^k,$$

$$\Delta z^k = -z^k + \mu_k (X^k)^{-1} e - Z^k (X^k)^{-1} \Delta x^k,$$

$$\Delta w^k = -w^k + \mu_k (S^k)^{-1} e - W^k (S^k)^{-1} \Delta s^k.$$

When applying an iterative method for solving (18), the computation of this search direction amounts to solving inexactly the system of normal equations

$$A\Theta^k A^\top \Delta y^k = \bar{b}. \quad (19)$$

The relation in (15) provides a stopping criterion for the iterative method, while assuring global convergence of the interior point method,

$$\|A\Theta^k A^\top \Delta y^k - \bar{b}\| \leq \beta_0 \|Ax^k - b\|. \quad (20)$$

After obtaining the search directions, the new iterate is computed as

$$\begin{aligned} x^{k+1} &= x^k + \alpha_p \Delta x^k, \\ s^{k+1} &= s^k + \alpha_p \Delta s^k, \\ y^{k+1} &= y^k + \alpha_d \Delta y^k, \\ w^{k+1} &= w^k + \alpha_d \Delta w^k, \\ z^{k+1} &= z^k + \alpha_d \Delta z^k, \end{aligned}$$

where distinct step sizes $\alpha_p, \alpha_d \in (0, 1]$ are chosen in the primal and dual spaces. Still in accordance with the derivation of Kojima, Megiddo and Mizuno [13], the step sizes are computed in two stages. Initially, the algorithm computes a common step size which guarantees a reduction of the trajectory parameter and complementary slackness for the next iterate, satisfying the global convergence requirements,

$$\begin{aligned} \bar{\alpha} &= \max\{0 < \alpha \leq 1 : (x^k(\alpha), s^k(\alpha), y^k(\alpha), w^k(\alpha), z^k(\alpha)) \in \mathcal{K}, \\ &\quad \mu(x^k(\alpha), s^k(\alpha), y^k(\alpha), w^k(\alpha), z^k(\alpha)) \leq (1 - \alpha(1 - \beta_2))\mu_k\}, \end{aligned}$$

where

$$\begin{aligned} x^k(\alpha) &= x^k + \alpha \Delta x^k, \\ s^k(\alpha) &= s^k + \alpha \Delta s^k, \\ y^k(\alpha) &= y^k + \alpha \Delta y^k, \\ w^k(\alpha) &= w^k + \alpha \Delta w^k, \\ z^k(\alpha) &= z^k + \alpha \Delta z^k. \end{aligned}$$

Subsequently, more effective primal and dual step sizes can be selected, as long as the new iterate satisfies

$$\begin{aligned} (x^{k+1}, s^{k+1}, y^{k+1}, w^{k+1}, z^{k+1}) &\in \mathcal{K}, \\ \mu_{k+1} &\leq (1 - \bar{\alpha}(1 - \beta_3))\mu_k. \end{aligned}$$

Most practical implementations of this algorithm [14, 19] select large step sizes,

$$\begin{aligned} \alpha_p &= \rho_p \max\{\alpha > 0 : x^k + \alpha \Delta x^k \geq 0, s^k + \alpha \Delta s^k \geq 0\}, \\ \alpha_d &= \rho_d \max\{\alpha > 0 : w^k + \alpha \Delta w^k \geq 0, z^k + \alpha \Delta z^k \geq 0\}, \end{aligned}$$

with $\rho_p = \rho_d = 0.9995$. However, the algorithm described in this section cannot ensure global convergence for this choice of step sizes.

Infeasible exact Newton methods possess global, polynomial and even super linear convergence. Kojima, Megiddo and Mizuno [13] have shown that when $\beta_0 = 0$, the algorithm described in this section produces, in a finite number of iterations, a solution that either satisfies the given optimality and feasibility tolerances or exceeds a prescribed upper bound for its norm. In [27], Zhang presented an exact Newton interior point algorithm which guarantees to return an approximate solution by assuming that the set of optimal solutions is nonempty. This algorithm was also proved to have polynomial complexity. The convergence theory of infeasible exact Newton methods is based on the fact that the generated sequence of iterates satisfies the equality

$$Ax^k - b = v_k(Ax^0 - b), \quad (21)$$

where v_k is positive and converges to zero as $k \rightarrow \infty$.

The inexact Newton method presented in this section was introduced by Portugal et al. [19]. In this algorithm, the norm $\|A(x^k + \Delta x^k) - b\|$ measures the closeness between the exact and the approximated Newton directions. The sequence generated by the algorithm does not satisfy equality (21). Instead, the weaker condition

$$\|Ax^k - b\| = v_k \|Ax^0 - b\| \quad (22)$$

holds. A slight modification of the analysis presented in [13] establishes the following convergence result for the truncated Newton method.

Theorem 1. *Either there is an iteration k where $(x^k, s^k, y^k, w^k, z^k)$ is an approximate solution satisfying*

$$\mu_k < \varepsilon_c \quad \text{and} \quad \|Ax^k - b\| < \varepsilon_p$$

or

$$\|(w^k, z^k)\| \rightarrow \infty \quad \text{as} \quad k \rightarrow \infty.$$

Proof. The proof differs from that presented in [13] for the exact Newton algorithm in only one detail. Instead of satisfying equality (11a) of [13], the sequence generated by this inexact Newton method satisfies

$$A(x^k + \alpha \Delta x^k) - b = (1 - \alpha)(Ax^k - b) + \alpha r^k,$$

where r^k is such that

$$\|r^k\| \leq \beta_0 \|Ax^k - b\|.$$

By taking this into consideration, we obtain for every $\alpha \in [0, 1]$ that (see [13], page 271)

$$\begin{aligned} g_p(\alpha) &\geq (\beta_1 - \beta_0)\varepsilon^*\alpha - \eta\alpha^2 \text{ if } g_p(0) \geq 0, \\ (1 - \alpha(1 - \beta_0))\|Ax^k - b\| &\leq \varepsilon_p \text{ if } g_p(0) < 0, \end{aligned}$$

where

$$g_p(\alpha) = (x^k + \alpha \Delta x^k)^\top (z^k + \alpha \Delta z^k) + (w^k + \alpha \Delta w^k)^\top (s^k + \alpha \Delta s^k) + \gamma_p(1 - \alpha) \|Ax^k - b\|.$$

Consequently, the primal and dual step sizes are bounded away from zero, that is, there exists a $\alpha^* > 0$ such that $\bar{\alpha} \geq \alpha^*$ holds for every k . $\square$

It follows from this theorem that the algorithm converges to an approximate optimal solution if the sequence of dual iterates is bounded. This assumption, commonly used in the analysis of nonlinear programming algorithms, is reasonable when solving feasible linear programs, especially in the case of network flows, where detection of infeasibility is straightforward. Computational experience described in [18, 19] support this conjecture. We have not been able to establish this result theoretically (see also [4]). Since dual feasibility is assumed in our algorithm, the linear program has a finite optimal solution.

After the design of our algorithm, several attempts have been made to develop globally convergent truncated interior-point methods for linear programs and other related problems [9, 17, 5, 4]. We have no intention in this paper to comment on the benefits and drawbacks of these alternative techniques in practice. Instead, our main objective is to provide bounds for the condition number of the matrices $M^{-1}(A\Theta^k A^\top)$ associated with the system of normal equations (19), where Θ^k is given by (16), M is a preconditioning matrix to be discussed later and A is a full-row rank node-arc incidence matrix. We believe that the results included in this paper may also have an important impact on the solution of linear network flow problems by these alternative interior-point techniques.

Another important property of the interior point algorithms is concerned with the limit points of the iteration sequence. Let $\mathcal{S}$ denote the set of pairs of primal and dual optimal solutions. Classical linear programming theory guarantees the existence of a strictly complementary pair [22], which induces the optimal partition of the primal variables index set $\{1, 2, \dots, n\}$ into

$$\begin{aligned} \mathcal{B} &= \{i : z_i = 0 \text{ and } w_i = 0 \ \forall (x, y, s, w, z) \in \mathcal{S}\}, \\ \mathcal{L} &= \{i : x_i = 0 \text{ and } w_i = 0 \ \forall (x, y, s, w, z) \in \mathcal{S}\}, \\ \mathcal{U} &= \{i : s_i = 0 \text{ and } z_i = 0 \ \forall (x, y, s, w, z) \in \mathcal{S}\}, \end{aligned}$$

where $\mathcal{B}$ indicates primal variables and corresponding slacks with strictly positive values, $\mathcal{L}$ primal variables in the lower bound and $\mathcal{U}$ primal variables in the upper bound. constitute the optimal partition of . Equality (21) can be used to prove that there exists a constant τ such that the sequences of iterates generated by infeasible exact Newton algorithms progressing inside

neighborhoods of the central path similar to set $\mathcal{K}$ satisfy, for all k sufficiently large [23], [25],

$$\begin{aligned} \tau \leq s_i^k, x_i^k \leq \frac{1}{\tau} \text{ and } \gamma\tau\mu_k \leq w_i^k, z_i^k \leq \frac{\mu_k}{\tau} \text{ if } i \in \mathcal{B}, \\ \tau \leq s_i^k, z_i^k \leq \frac{1}{\tau} \text{ and } \gamma\tau\mu_k \leq x_i^k, w_i^k \leq \frac{\mu_k}{\tau} \text{ if } i \in \mathcal{L}, \\ \tau \leq x_i^k, w_i^k \leq \frac{1}{\tau} \text{ and } \gamma\tau\mu_k \leq s_i^k, z_i^k \leq \frac{\mu_k}{\tau} \text{ if } i \in \mathcal{U}. \end{aligned} \quad (23)$$

Therefore, in degenerate cases, the limiting behavior of these trajectories guarantees the identification of optimal primal and dual facets

We have not been able to establish a similar result for the truncated version of the algorithm. Nevertheless, we observed that the generated iteration sequence satisfies conditions (23) in practice. Our implementation of the truncated Newton method uses these conditions for early detection of the optimal solution.

3. The preconditioned conjugate gradient algorithm

A fundamental feature of an interior point implementation is the method selected for solving the linear system which determines the search direction at each iteration. As described in Section 2, primal-dual interior point variants use the system of normal equations

$$A\Theta^k A^\top \Delta y^k = \bar{b}, \quad (24)$$

where Θ^k and $\bar{b}$ are given by (16) and (17).

In PDNET, the preconditioned conjugate gradient algorithm is used to solve this linear system, using a stopping criterion based on the convergence properties of the truncated algorithm as stated in (20). Our implementation of the PCG algorithm follows the pseudo-code presented in Figure 2. This iterative solution method for linear systems attempts to increase the rate of convergence of the standard conjugate gradient method [3, 10] by transforming the original system in Equation 24 as follows:

$$M^{-1}A\Theta^k A^\top \Delta y^k = M^{-1}\bar{b}.$$

The preconditioning matrix M is symmetric and positive definite, sufficient conditions for the PCG algorithm to inherit all convergence properties of the standard conjugate gradient method. The aim of preconditioning is to produce a system matrix $M^{-1}A\Theta^k A^\top$ with more favorable spectral properties than those of $A\Theta^k A^\top$. The result is improved efficiency of the conjugate gradient algorithm by reducing the total number of iterations.

```

procedure pcg( $A, \Theta^k, \bar{b}$ )
1  Compute  $\Delta y_0$  as a guess of  $\Delta y^k$ ;
2   $r_0 := \bar{b} - A\Theta^k A^\top \Delta y_0$ ;
3  solve  $Mz_0 = r_0$ ;
4   $p_0 := z_0$ ;
5   $i := 0$ ;
6  do stopping criterion not satisfied  $\rightarrow$ 
7       $q_i := A\Theta^k A^\top p_i$ ;
8       $\alpha_i := z_i^\top r_i / p_i^\top q_i$ ;
9       $\Delta y_{i+1} := \Delta y_i + \alpha_i p_i$ ;
10      $r_{i+1} := r_i - \alpha_i q_i$ ;
11     solve  $Mz_{i+1} = r_{i+1}$ ;
12      $\beta_i := z_{i+1}^\top r_{i+1} / z_i^\top r_i$ ;
13      $p_{i+1} := z_{i+1} + \beta_i p_i$ ;
14      $i := i + 1$ 
15 od;
16  $\Delta y^k := \Delta y_i$ 
end pcg;

```

Figure 2. The preconditioned conjugate gradient algorithm

The usual strategies for building preconditioners consist of computing approximations of either the system matrix or its inverse. In the second case, the linear systems in steps 3 and 11 of the PCG algorithm are replaced by the matrix-vector products $z_0 := M^{-1}r_0$ and $z_{i+1} := M^{-1}r_{i+1}$, respectively. Effective preconditioners reduce the iteration count of the conjugate gradient algorithm but incur in the additional computational cost of constructing either M or M^{-1} . In addition, at each iteration of the PCG algorithm either solves the linear system $Mz_{i+1} = r_{i+1}$ or computes $z_{i+1} := M^{-1}r_{i+1}$. The gain in the convergence speed achieved with a preconditioning scheme must be superior to the extra computational costs.

4. The condition number of the normal equations matrix

In this section we exploit the special structure of the node-arc incidence matrix to derive an upper-bound for the condition number of the normal equations matrix $A\Theta^k A^\top$ that depends on the size of the diagonal elements of the diagonal matrix Θ^k . By assuming property (23), we further show that the condition numbers of the matrices $A\Theta^k A^\top$ used by the truncated interior-point methods are uniformly bounded in the last iterations provided this algorithm converges to a primal nondegenerate solution. However, the boundedness of

the condition number is not guaranteed under a primal degenerate optimal solution.

We start by recalling some properties that will be useful to establish the main results of this and next sections. We also omit the superscripts k for easier explanation.

Proposition 1. *Let B be a $m \times m$ symmetric matrix. Then*

$$\min_{1 \leq l \leq m} \left\{ b_{ll} - \sum_{\substack{h=1 \\ h \neq l}}^m |b_{hl}| \right\} \leq \underline{\lambda}(B) \leq \bar{\lambda}(B) \leq \max_{1 \leq l \leq m} \left\{ b_{ll} + \sum_{\substack{h=1 \\ h \neq l}}^m |b_{hl}| \right\}.$$

Proof. This is the well known Gerschgorin Theorem (see [3]). $\square$

Proposition 2. *Let B and C be symmetric matrices, and let $\lambda_i(B)$ denote the i^{th} eigenvalue, where we order the eigenvalues in an increasing order. Also, let $\underline{\lambda}(C)$ and $\bar{\lambda}(C)$ denote the extreme eigenvalues of C . Then,*

$$\lambda_i(B) + \underline{\lambda}(C) \leq \lambda_i(B+C) \leq \lambda_i(B) + \bar{\lambda}(C).$$

Proof. This is a Corollary of the Courant-Fischer lemma (see [3]). $\square$

Proposition 3. *Let B be a symmetric positive semi-definite matrix and C be a symmetric positive definite matrix.*

- (i) *If ρ_1 is a number such that $\bar{\lambda}(\rho_1 C - B) \leq 0$, then $\underline{\lambda}(C^{-1}B) \geq \rho_1$.*
- (ii) *If ρ_2 is a number such that $\underline{\lambda}(\rho_2 C - B) \geq 0$, then $\bar{\lambda}(C^{-1}B) \leq \rho_2$.*

Proof. See [3], page 404. $\square$

Let $\mathcal{G} = (\mathcal{N}, \mathcal{A})$ denote the network associated to the MCNF problem, where $\mathcal{N}$ and $\mathcal{A}$ represent the set of nodes and the set of directed arcs, respectively. Each element (i, j) of $\mathcal{A}$ corresponds to one column of A and vice-versa. The same correspondence between elements of $\mathcal{A}$ and diagonal elements of Θ exists. Suppose that $\mathcal{H}$ is a subset of $\mathcal{A}$ and let $A_{\mathcal{H}}$ and $\Theta_{\mathcal{H}}$ be the corresponding submatrices of A and Θ , respectively. The diagonal element $[A_{\mathcal{H}} \Theta_{\mathcal{H}} A_{\mathcal{H}}^{\top}]_{ll}$ and the off-diagonal element $[A_{\mathcal{H}} \Theta_{\mathcal{H}} A_{\mathcal{H}}^{\top}]_{lh}$ of $A_{\mathcal{H}} \Theta_{\mathcal{H}} A_{\mathcal{H}}^{\top}$ are given by

$$[A_{\mathcal{H}} \Theta_{\mathcal{H}} A_{\mathcal{H}}^{\top}]_{ll} = \sum_{(i,j):(i,j) \in \mathcal{H}_l} \theta_{ij}$$

and

$$[A_{\mathcal{H}} \Theta_{\mathcal{H}} A_{\mathcal{H}}^{\top}]_{lh} = \begin{cases} \sum_{(i,j):(i,j) \in \mathcal{H}_l \cap \mathcal{H}_h} -\theta_{ij} & \text{if } \mathcal{H}_l \cap \mathcal{H}_h \neq \emptyset \\ 0 & \text{if } \mathcal{H}_l \cap \mathcal{H}_h = \emptyset \end{cases}$$

respectively, where $\mathcal{H}_l$ represents the set of arcs in $\mathcal{H}$ incident to node l and θ_{ij} denotes the diagonal element of $\Theta_{\mathcal{H}}$ associated to arc (i, j) .

Theorem 2. $\bar{\lambda}(A_{\mathcal{H}}\Theta_{\mathcal{H}}A_{\mathcal{H}}^{\top}) \leq 2\bar{\lambda}(\Theta_{\mathcal{H}})|\mathcal{H}|.$

Proof. We note that

$$[A_{\mathcal{H}}\Theta_{\mathcal{H}}A_{\mathcal{H}}^{\top}]_{ll} \geq \sum_{\substack{h=1 \\ h \neq l}}^m |[A_{\mathcal{H}}\Theta_{\mathcal{H}}A_{\mathcal{H}}^{\top}]_{hl}|, \quad l = 1, \dots, m. \quad (25)$$

So, from Proposition 1 we obtain

$$\begin{aligned} \bar{\lambda}(A_{\mathcal{H}}\Theta_{\mathcal{H}}A_{\mathcal{H}}^{\top}) &\leq 2 \max_{1 \leq l \leq m} \{[A_{\mathcal{H}}\Theta_{\mathcal{H}}A_{\mathcal{H}}^{\top}]_{ll}\} = 2 \max_{1 \leq l \leq m} \left\{ \sum_{(i,j):(i,j) \in \mathcal{H}_l} \theta_{ij} \right\} \\ &\leq 2\bar{\lambda}(\Theta_{\mathcal{H}}) \max_{1 \leq l \leq m} \{|\mathcal{H}_l|\} \leq 2\bar{\lambda}(\Theta_{\mathcal{H}})|\mathcal{H}| \end{aligned}$$

as desired. $\square$

Since we assume that A is a full row rank matrix, $A\Theta A^{\top}$ is positive definite and there exists at least one basis of A . However, $A_{\mathcal{H}}\Theta_{\mathcal{H}}A_{\mathcal{H}}^{\top}$ is not positive definite for every subset $\mathcal{H}$ of $\mathcal{A}$. A necessary and sufficient condition for the positive semi-definite matrix $A_{\mathcal{H}}\Theta_{\mathcal{H}}A_{\mathcal{H}}^{\top}$ to be nonsingular is the existence of a set of arcs $\mathcal{T} \subseteq \mathcal{H}$ such that $\bar{\mathcal{G}} = (\mathcal{N}, \mathcal{T})$ is a spanning tree.

Theorem 3. $A_{\mathcal{H}}\Theta_{\mathcal{H}}A_{\mathcal{H}}^{\top}$ is positive definite and

$$\underline{\lambda}(A_{\mathcal{H}}\Theta_{\mathcal{H}}A_{\mathcal{H}}^{\top}) \geq \frac{\underline{\lambda}(\Theta_{\mathcal{T}})}{m^2},$$

if and only if there exists $\mathcal{T} \subseteq \mathcal{H}$ such that $\bar{\mathcal{G}} = (\mathcal{N}, \mathcal{T})$ is a spanning tree.

Proof. The rank of $A_{\mathcal{H}}$ equals the cardinality of $\mathcal{T}$, where $\mathcal{T} \subseteq \mathcal{H}$ is the maximum cardinality set such that $\bar{\mathcal{G}} = (\mathcal{N}, \mathcal{T})$ does not contain any cycle [2]. If $\bar{\mathcal{G}} = (\mathcal{N}, \mathcal{T})$ is not a spanning tree, then $|\mathcal{T}| < m$, $A_{\mathcal{H}}\Theta_{\mathcal{H}}A_{\mathcal{H}}^{\top}$ is singular and $\underline{\lambda}(A_{\mathcal{H}}\Theta_{\mathcal{H}}A_{\mathcal{H}}^{\top}) = 0$. If $\bar{\mathcal{G}} = (\mathcal{N}, \mathcal{T})$ is a spanning tree, then $|\mathcal{T}| = m$ and $A_{\mathcal{H}}\Theta_{\mathcal{H}}A_{\mathcal{H}}^{\top}$ is positive definite. In this case, let $\bar{\mathcal{T}} = \mathcal{H} - \mathcal{T}$. Since $A_{\bar{\mathcal{T}}}\Theta_{\bar{\mathcal{T}}}A_{\bar{\mathcal{T}}}^{\top}$ and $A_{\mathcal{T}}(\Theta_{\mathcal{T}} - \underline{\lambda}(\Theta_{\mathcal{T}})I)A_{\mathcal{T}}^{\top}$ are positive semi-definite matrices and the latter is singular, then from Proposition 2 we have

$$\begin{aligned} \underline{\lambda}(A_{\mathcal{H}}\Theta_{\mathcal{H}}A_{\mathcal{H}}^{\top}) &\geq \underline{\lambda}(A_{\mathcal{T}}\Theta_{\mathcal{T}}A_{\mathcal{T}}^{\top}) + \underline{\lambda}(A_{\bar{\mathcal{T}}}\Theta_{\bar{\mathcal{T}}}A_{\bar{\mathcal{T}}}^{\top}) \geq \underline{\lambda}(A_{\mathcal{T}}\Theta_{\mathcal{T}}A_{\mathcal{T}}^{\top}) \\ &\geq \underline{\lambda}(A_{\mathcal{T}}(\Theta_{\mathcal{T}} - \underline{\lambda}(\Theta_{\mathcal{T}})I)A_{\mathcal{T}}^{\top}) + \underline{\lambda}(\Theta_{\mathcal{T}})\underline{\lambda}(A_{\mathcal{T}}A_{\mathcal{T}}^{\top}) \\ &= \frac{\underline{\lambda}(\Theta_{\mathcal{T}})}{\underline{\lambda}(A_{\mathcal{T}}^{-\top}A_{\mathcal{T}}^{-1})}. \end{aligned}$$

Now, $A_{\mathcal{T}}^{-1}$ is a matrix having nonzero elements equal to 1 or -1 [2]. So, according to Proposition 1, $\underline{\lambda}(A_{\mathcal{T}}^{-\top}A_{\mathcal{T}}^{-1}) \leq m^2$. This gives the above lower bound for $\underline{\lambda}(A_{\mathcal{H}}\Theta_{\mathcal{H}}A_{\mathcal{H}}^{\top})$ when $A_{\mathcal{H}}\Theta_{\mathcal{H}}A_{\mathcal{H}}^{\top}$ is nonsingular. $\square$

The next result establishes an upper bound for the condition number of $A\Theta A^\top$.

Corollary 1. *Suppose that the elements of Θ represent weights associated to the arcs of $\mathcal{A}$ and let $\bar{G} = (\mathcal{N}, \bar{T})$ be a maximum spanning tree of $G = (\mathcal{N}, \mathcal{A})$ and $\bar{T} = \mathcal{A} - T$. Then*

$$\kappa(A\Theta A^\top) \leq 2m^3 \frac{\bar{\lambda}(\Theta_T)}{\underline{\lambda}(\Theta_T)} + 2m^2(n-m) \frac{\bar{\lambda}(\Theta_{\bar{T}})}{\underline{\lambda}(\Theta_{\bar{T}})}. \quad (26)$$

Proof. From Theorem 3 we have that

$$\underline{\lambda}(A\Theta A^\top) \geq \frac{\underline{\lambda}(\Theta_T)}{m^2},$$

since $A\Theta A^\top$ is positive definite. On the other hand, Proposition 2 and Theorem 2 give

$$\begin{aligned} \bar{\lambda}(A\Theta A^\top) &\leq \bar{\lambda}(A_T \Theta_T A_T^\top) + \bar{\lambda}(A_{\bar{T}} \Theta_{\bar{T}} A_{\bar{T}}^\top) \\ &\leq 2m\bar{\lambda}(\Theta_T) + 2(n-m)\bar{\lambda}(\Theta_{\bar{T}}). \end{aligned}$$

Condition (26) is obtained by dividing the upper bound by the lower bound estimated for the eigenvalues of $A\Theta A^\top$. $\square$

At the first iterations of the truncated Newton algorithm, we can manage to obtain matrices $A\Theta A^\top$ well conditioned by choosing an appropriate starting point for this method. As the algorithm progresses, the values of the diagonal elements of Θ spread and the conditioning of $A\Theta A^\top$ degrades. At the final stages, it may happen that $A\Theta A^\top$ is very ill conditioned.

As discussed in Section 2, we assume that there exists a constant τ such that the conditions in (23) hold when μ becomes small.

Corollary 2. *Suppose that conditions (23) hold with μ small. Then, at the later iterations of the truncated Newton method*

$$\kappa(A\Theta A^\top) \leq \begin{cases} \tau_1 + \tau_2 \frac{1}{\frac{1}{2\gamma\mu^2} + \frac{1}{2}} & \text{if } |\mathcal{B}| = m \\ \tau_1 + \tau_2 & \text{if } |\mathcal{B}| > m \\ \tau_1(\frac{1}{2\gamma\mu^2} + \frac{1}{2}) + \tau_2 & \text{if } |\mathcal{B}| < m, \end{cases} \quad (27)$$

where the constants τ_1 and τ_2 are given by

$$\tau_1 = 2m^3 \frac{1}{\gamma\tau^4} \quad \text{and} \quad \tau_2 = 2m^2(n-m) \frac{1}{\gamma\tau^4}.$$

Proof. Since

$$\Theta = (ZX^{-1} + WS^{-1})^{-1},$$

where Z , X , W and S are diagonal matrices with diagonal elements z_{ij} , x_{ij} , w_{ij} and s_{ij} respectively, then it immediately follows from the conditions (23) that

$$\begin{aligned} \frac{\tau^2}{2\mu} &\leq \underline{\lambda}(\Theta_{\mathcal{B}}) \leq \bar{\lambda}(\Theta_{\mathcal{B}}) \leq \frac{1}{\gamma\tau^2 2\mu} \quad \text{and} \\ \frac{1}{\frac{1}{\gamma\tau^2\mu} + \frac{\mu}{\tau^2}} &\leq \underline{\lambda}(\Theta_{\mathcal{L} \cup \mathcal{U}}) \leq \bar{\lambda}(\Theta_{\mathcal{L} \cup \mathcal{U}}) \leq \frac{1}{\frac{\tau^2}{\mu} + \gamma\tau^2\mu}. \end{aligned}$$

Furthermore for small and positive μ we have

$$\begin{aligned} \frac{1}{\frac{1}{\gamma\tau^2\mu} + \frac{\mu}{\tau^2}} &= \frac{\tau^2}{2\mu} \left[\frac{2}{1 + \frac{1}{\mu^2\gamma}} \right] < \frac{\tau^2}{2\mu} \quad \text{and} \\ \frac{1}{\frac{\tau^2}{\mu} + \gamma\tau^2\mu} &= \frac{1}{2\gamma\tau^2\mu} \left[\frac{2}{1 + \frac{1}{\gamma^2}} \right] < \frac{1}{2\gamma\tau^2\mu}. \end{aligned}$$

Now the conditions (27) follow from (26) by setting $\mathcal{T} = \mathcal{B}$, $\mathcal{T} \subset \mathcal{B}$ and $\mathcal{T} \supset \mathcal{B}$, respectively. $\square$

We note that $|\mathcal{B}| = m$, $|\mathcal{B}| < m$ and $|\mathcal{B}| > m$ correspond to the cases of primal and dual nondegeneracy, dual degeneracy and primal degeneracy respectively. It follows from this theorem that under the assumption (23) the condition number of $A\Theta A^\top$ is uniformly bounded at the last iterations of the algorithm provided the optimal solution is primal nondegenerate, confirming in a more formal approach the analysis presented in [11]. The upper bound found for $\kappa(A\Theta A^\top)$ when the primal and dual optimal solution are unique, is smaller than that obtained when dual degeneracy is present, since in the first case the second term in (27) vanishes as $\mu \rightarrow 0$. If the problem is primal degenerate, then the right-hand side of (27) converges to infinity as $\mu \rightarrow 0$. In this case, $A\Theta A^\top$ is usually poorly conditioned when μ becomes small. In certain implementations of interior point methods for linear programming where Cholesky factorization is used to compute the search directions several schemes are used to overcome this problem, namely, procedures to handle negative and zero pivots encountered during the factorization process [24]. In the conjugate gradient based interior point algorithm we propose for the MCNF problem, preconditioning schemes for the matrix $A\Theta A^\top$ are used. As we show in the next section, we can manage to reduce to 1 the first term in (27) by using spanning tree based preconditioning techniques.

5. Condition Number of the Preconditioned Matrix

As is discussed in [19] and in this paper, the system of normal equations (19) is solved by using the Preconditioned Conjugate Gradient Method, with a preconditioner M . The Diagonal, the Maximum Spanning Tree (MST) and the Diagonal Compensated Maximum Spanning Tree (DCMST) Preconditioners have been the most recommended in practice to be used in interior-point network flow problems [16, 19, 26]. The success of these techniques depends on the boundedness of the condition number of the preconditioned matrix $M^{-1}(A\Theta_k A^\top)$. In this section we derive some condition number upper bounds for each one of the preconditioners stated before. These theoretical results seem to support our strategy for choosing the preconditioner discussed in [19]. The diagonal preconditioner

$$M = \text{diag}(A\Theta A^\top)$$

was first used by Yeh [26] in the context of interior point algorithms for network flows. The cost of constructing M^{-1} is $O(n)$ additions and $O(m)$ divisions. The cost of computing $z_{i+1} = M^{-1}r_{i+1}$ is $O(m)$ multiplications. The next result gives an upper bound for $\kappa(M^{-1}A\Theta A^\top)$.

Theorem 4. *Let $\bar{G} = (\mathcal{N}, \mathcal{T})$ be a maximum spanning tree of $G = (\mathcal{N}, \mathcal{A})$, $\mathcal{T} = \mathcal{A} - \mathcal{T}$ and $M = \text{diag}(A\Theta A^\top)$. Then,*

$$\kappa(M^{-1}A\Theta A^\top) \leq 2m^3 \frac{\bar{\lambda}(\Theta_{\mathcal{T}})}{\underline{\lambda}(\Theta_{\mathcal{T}})} + 2m^2(n-m) \frac{\bar{\lambda}(\Theta_{\bar{\mathcal{T}}})}{\underline{\lambda}(\Theta_{\mathcal{T}})}. \quad (28)$$

Proof. Since $2M - A\Theta A^\top$ is positive semi-definite, then

$$\underline{\lambda}(2M - A\Theta A^\top) \geq 0$$

and, by Proposition 3,

$$\bar{\lambda}(M^{-1}A\Theta A^\top) \leq 2.$$

Now, let

$$\rho_1 = \frac{1}{m^3 \frac{\bar{\lambda}(\Theta_{\mathcal{T}})}{\underline{\lambda}(\Theta_{\mathcal{T}})} + m^2(n-m) \frac{\bar{\lambda}(\Theta_{\bar{\mathcal{T}}})}{\underline{\lambda}(\Theta_{\mathcal{T}})}}.$$

From Propositions 1, 2 and Theorem 3 we have

$$\begin{aligned} \bar{\lambda}(\rho_1 M - A\Theta A^\top) &\leq \rho_1 \bar{\lambda}(M) - \underline{\lambda}(A\Theta A^\top) \\ &\leq \rho_1 m \bar{\lambda}(\Theta_{\mathcal{T}}) + \rho_1(n-m) \bar{\lambda}(\Theta_{\bar{\mathcal{T}}}) - \frac{\underline{\lambda}(\Theta_{\mathcal{T}})}{m^2} = 0. \end{aligned}$$

Thus, $\frac{2}{\rho_1}$ is an upper bound for the condition number of $M^{-1}A\Theta A^\top$. $\square$

It is interesting to note that the upper-bound achieved for the preconditioned matrix $M^{-1}(A\Theta A^\top)$ coincides with that of the matrix $A\Theta A^\top$. This result is not surprising. Computational experience described in [20] has shown that the diagonal preconditioner is effective at the first stages of the interior point algorithms when $A\Theta A^\top$ is itself well conditioned. However, as the algorithms progress, it performs similarly to the identity matrix preconditioner.

The lack of effectiveness of the diagonal preconditioner at the intermediate and final iterations of the interior point methods motivated the search for other preconditioning schemes. The maximum spanning tree (MST) preconditioner was introduced by Resende and Veiga [21]. It has the form

$$M = A_{\mathcal{T}} \Theta_{\mathcal{T}} A_{\mathcal{T}}^\top,$$

where $\bar{\mathcal{G}} = (\mathcal{N}, \mathcal{T})$ is a maximum spanning tree of $\mathcal{G} = (\mathcal{N}, \mathcal{A})$ with edge weights Θ . Contrary to the diagonal preconditioner, this preconditioning matrix takes advantage of the structure of the optimal solution of the MCNF problem. Constructing M involves computing $\bar{\mathcal{G}} = (\mathcal{N}, \mathcal{T})$ and inverting $\Theta_{\mathcal{T}}$ ($\Theta_{\mathcal{T}}^{-1}$ is explicitly computed in order to perform multiplications instead of divisions when the linear systems $Mz_{i+1} = r_{i+1}$ are solved in the PCG algorithm). In [21], the Kruskal algorithm is used to compute an approximate maximum spanning tree, while in [19] an exact maximum spanning tree is obtained by employing an implementation of Prim's algorithm with a Fibonacci heap data structure. The running time for this last method is $O(n + m \log(m))$. The cost of solving $Mz_{i+1} = r_{i+1}$ is $O(m)$ additions and multiplications.

In [16], Mehrotra and Wang proposed the Diagonally Compensated Maximum Spanning Tree (DCMST) Preconditioner $M = A_{\mathcal{T}} \Theta_{\mathcal{T}} A_{\mathcal{T}}^\top + D$ where $\bar{\mathcal{G}} = (\mathcal{N}, \mathcal{T})$ is a maximum spanning tree of $\bar{\mathcal{G}} = (\mathcal{N}, \mathcal{A})$, $\mathcal{T} = \mathcal{A} - \mathcal{T}$, and $D = \phi \text{diag}(A_{\bar{\mathcal{T}}} \Theta_{\bar{\mathcal{T}}} A_{\bar{\mathcal{T}}}^\top)$ with ϕ a nonnegative parameter. We further notice that the MST preconditioner is a special case of this last one when $\phi = 0$. Next we establish an upper-bound for the condition number of the preconditioned matrix when M is the DCMST preconditioner.

Theorem 5. *Let $\bar{\mathcal{G}} = (\mathcal{N}, \mathcal{T})$ be a maximum spanning tree of $\mathcal{G} = (\mathcal{N}, \mathcal{A})$ and $\mathcal{T} = \mathcal{A} - \mathcal{T}$. Define*

$$M = A_{\mathcal{T}} \Theta_{\mathcal{T}} A_{\mathcal{T}}^\top + D,$$

where

$$D = \phi \text{diag}(A_{\bar{\mathcal{T}}} \Theta_{\bar{\mathcal{T}}} A_{\bar{\mathcal{T}}}^\top) \quad \text{with} \quad 0 \leq \phi \leq 2.$$

Then,

$$\kappa(M^{-1}A\Theta A^\top) \leq \min \left\{ \frac{1 + m^2 \frac{\bar{\lambda}(D)}{\underline{\lambda}(\Theta_T)}}{1 + m^2 \frac{\underline{\lambda}(D)}{\underline{\lambda}(\Theta_T)}} \left(1 + 2m^2(n-m) \frac{\bar{\lambda}(\Theta_{\bar{T}})}{\underline{\lambda}(\Theta_T)} \right), \right. \\ \left. \frac{1 + m^2 \frac{\bar{\lambda}(D)}{\underline{\lambda}(\Theta_T)}}{1 + \frac{1}{2m} \frac{\underline{\lambda}(D)}{\underline{\lambda}(\Theta_T)}} \left(1 + m(n-m) \frac{\bar{\lambda}(\Theta_{\bar{T}})}{\underline{\lambda}(\Theta_T)} \right) \right\}. \quad (29)$$

Proof. Let

$$\rho_1 = \frac{1}{1 + m^2 \frac{\bar{\lambda}(D)}{\underline{\lambda}(\Theta_T)}}$$

and

$$\hat{\rho}_2 = \frac{1 + 2m^2(n-m) \frac{\bar{\lambda}(\Theta_{\bar{T}})}{\underline{\lambda}(\Theta_T)}}{1 + m^2 \frac{\underline{\lambda}(D)}{\underline{\lambda}(\Theta_T)}}, \quad \bar{\rho}_2 = \frac{1 + m(n-m) \frac{\bar{\lambda}(\Theta_{\bar{T}})}{\underline{\lambda}(\Theta_T)}}{1 + \frac{1}{2m} \frac{\underline{\lambda}(D)}{\underline{\lambda}(\Theta_T)}}.$$

We first note that $0 \leq \rho_1 \leq 1$. Furthermore, it follows from the definition of the matrix D and $0 \leq \phi \leq 2$ that $\bar{\lambda}(D) \leq 2(n-m)\bar{\lambda}(\Theta_{\bar{T}})$ and this guarantees that $\hat{\rho}_2, \bar{\rho}_2 \geq 1$. From Proposition 2, Theorem 3 and the positive semi-definiteness of $A_{\bar{T}}^\top \Theta_{\bar{T}} A_{\bar{T}}^\top$, we have

$$\begin{aligned} \bar{\lambda}(\rho_1 M - A\Theta A^\top) &= \bar{\lambda} \left((\rho_1 - 1) A_T \Theta_T A_T^\top + \rho_1 D - A_{\bar{T}}^\top \Theta_{\bar{T}} A_{\bar{T}}^\top \right) \\ &\leq (\rho_1 - 1) \underline{\lambda}(A_T \Theta_T A_T^\top) + \rho_1 \bar{\lambda}(D) \\ &\leq (\rho_1 - 1) \frac{\underline{\lambda}(\Theta_T)}{m^2} + \rho_1 \bar{\lambda}(D) \\ &= \left(\frac{\underline{\lambda}(\Theta_T)}{m^2} + \bar{\lambda}(D) \right) \rho_1 - \frac{\underline{\lambda}(\Theta_T)}{m^2} = 0. \end{aligned}$$

On the other hand, from Proposition 2 and Theorems 2 and 3 we obtain

$$\begin{aligned} \underline{\lambda}(\hat{\rho}_2 M - A\Theta A^\top) &= \underline{\lambda} \left((\hat{\rho}_2 - 1) A_T \Theta_T A_T^\top + \hat{\rho}_2 D - A_{\bar{T}}^\top \Theta_{\bar{T}} A_{\bar{T}}^\top \right) \\ &\geq (\hat{\rho}_2 - 1) \underline{\lambda}(A_T \Theta_T A_T^\top) + \hat{\rho}_2 \underline{\lambda}(D) - \bar{\lambda}(A_{\bar{T}}^\top \Theta_{\bar{T}} A_{\bar{T}}^\top) \\ &\geq (\hat{\rho}_2 - 1) \frac{\underline{\lambda}(\Theta_T)}{m^2} + \hat{\rho}_2 \underline{\lambda}(D) - 2(n-m) \bar{\lambda}(\Theta_{\bar{T}}) \\ &= \left(\frac{\underline{\lambda}(\Theta_T)}{m^2} + \underline{\lambda}(D) \right) \hat{\rho}_2 - \\ &\quad \left(\frac{\underline{\lambda}(\Theta_T)}{m^2} + 2(n-m) \bar{\lambda}(\Theta_{\bar{T}}) \right) = 0. \end{aligned}$$

Now, it is well known that each element of $A_{\mathcal{T}}^{-1}A_{\bar{\mathcal{T}}}^{-\top}$ equals either 0, 1 or -1 [2]. This implies that

- (i) $0 \leq [A_{\mathcal{T}}^{-1}A_{\bar{\mathcal{T}}}^{-\top}A_{\mathcal{T}}^{-\top}]_{ll} \leq n-m, \quad l = 1, \dots, m,$
- (ii) $-(n-m) \leq [A_{\mathcal{T}}^{-1}A_{\bar{\mathcal{T}}}^{-\top}A_{\mathcal{T}}^{-\top}]_{lh} \leq n-m \quad l, h = 1, \dots, m \quad l \neq h.$

Consequently, $\bar{\lambda}(A_{\mathcal{T}}^{-1}A_{\bar{\mathcal{T}}}^{-\top}A_{\mathcal{T}}^{-\top}) \leq m(n-m)$ according to Proposition 1. Furthermore,

$$\begin{aligned}
& \underline{\lambda} \left((\bar{\rho}_2 - 1)\Theta_{\mathcal{T}} + \bar{\rho}_2 A_{\mathcal{T}}^{-1} D A_{\mathcal{T}}^{-\top} - A_{\mathcal{T}}^{-1} A_{\bar{\mathcal{T}}}^{-\top} \Theta_{\bar{\mathcal{T}}} A_{\bar{\mathcal{T}}}^{-\top} A_{\mathcal{T}}^{-\top} \right) \\
& \geq (\bar{\rho}_2 - 1) \underline{\lambda}(\Theta_{\mathcal{T}}) + \bar{\rho}_2 \underline{\lambda}(A_{\mathcal{T}}^{-1} D A_{\mathcal{T}}^{-\top}) - \bar{\lambda} \left(A_{\mathcal{T}}^{-1} A_{\bar{\mathcal{T}}}^{-\top} \Theta_{\bar{\mathcal{T}}} A_{\bar{\mathcal{T}}}^{-\top} A_{\mathcal{T}}^{-\top} \right) \\
& \geq (\bar{\rho}_2 - 1) \underline{\lambda}(\Theta_{\mathcal{T}}) + \bar{\rho}_2 \frac{1}{\underline{\lambda}(A_{\mathcal{T}} D^{-1} A_{\mathcal{T}}^{\top})} - \bar{\lambda}(\Theta_{\bar{\mathcal{T}}}) \bar{\lambda} \left(A_{\mathcal{T}}^{-1} A_{\bar{\mathcal{T}}}^{-\top} A_{\mathcal{T}}^{-\top} \right) \\
& \geq (\bar{\rho}_2 - 1) \underline{\lambda}(\Theta_{\mathcal{T}}) + \bar{\rho}_2 \frac{\underline{\lambda}(D)}{2m} - \bar{\lambda}(\Theta_{\bar{\mathcal{T}}}) m(n-m) \\
& \geq \left(\underline{\lambda}(\Theta_{\mathcal{T}}) + \frac{\underline{\lambda}(D)}{2m} \right) \bar{\rho}_2 - (\underline{\lambda}(\Theta_{\mathcal{T}}) + \bar{\lambda}(\Theta_{\bar{\mathcal{T}}}) m(n-m)) = 0.
\end{aligned}$$

Thus, $\underline{\lambda}(\bar{\rho}_2 M - A\Theta A^{\top}) \geq 0$. According to Proposition 3, $\frac{\hat{\rho}_2}{\rho_1}$ and $\frac{\bar{\rho}_2}{\rho_1}$ are upper bounds for the condition number of $M^{-1}A\Theta A^{\top}$. $\square$

Consider now the MST preconditioner. Since $\phi = 0$ in this case, then $\underline{\lambda}(D) = \bar{\lambda}(D) = 0$. Furthermore $2m^2 > m$ for all $m > 1$ and this leads to

$$\kappa(M^{-1}A\Theta A^{\top}) \leq 1 + m(n-m) \frac{\bar{\lambda}(\Theta_{\bar{\mathcal{T}}})}{\underline{\lambda}(\Theta_{\mathcal{T}})} \quad (30)$$

It is interesting to further note that this upper bound is smaller than the one for the DCMST independently of the choice of $\phi > 0$. In fact, as $\underline{\lambda}(D) < \bar{\lambda}(D)$ for any $\phi > 0$ and $\frac{1}{2m} < m^2$ for any $m > 1$, then

$$1 + m(n-m) \frac{\bar{\lambda}(\Theta_{\bar{\mathcal{T}}})}{\underline{\lambda}(\Theta_{\mathcal{T}})} < up$$

where up is the upper-bound provided by (29).

As before, if we assume that the conditions (23) hold, it is possible to show that the condition number of $M^{-1}(A\Theta A^{\top})$ is uniformly bounded in the last stages of the algorithm for both the MST and DCMST preconditioners. This is shown in the next theorem.

Corollary 3. *Suppose that conditions (23) hold with μ small. Let $\bar{\mathcal{G}} = (\mathcal{N}, \mathcal{T})$ be a maximum spanning tree of $\mathcal{G} = (\mathcal{N}, \mathcal{A})$ and $\mathcal{T} = \mathcal{A} - \mathcal{T}$. Define*

$$M = A_{\mathcal{T}} \Theta_{\mathcal{T}} A_{\mathcal{T}}^{\top} + D,$$

where

$$D = \phi \operatorname{diag}(A_{\bar{\mathcal{T}}} \Theta_{\bar{\mathcal{T}}} A_{\bar{\mathcal{T}}}^\top) \quad \text{with} \quad 0 \leq \phi \leq 2.$$

Then, at the later iterations of the truncated Newton method

$$\kappa(M^{-1}A\Theta A^\top) \leq \begin{cases} 1 + \tau_1 \frac{1}{\frac{1}{2\mu^2} + \frac{1}{2}} + \tau_2 \frac{1}{\left(\frac{1}{2\mu^2} + \frac{1}{2}\right)^2} & \text{if } |\mathcal{B}| = m \\ 1 + \tau_1 + \tau_2 & \text{if } |\mathcal{B}| > m \\ 1 + \tau_1 + \tau_2 & \text{if } |\mathcal{B}| < m, \end{cases}$$

where

$$\tau_1 = (\phi m^2 + m) \frac{1}{\gamma \tau^4} \quad \text{and} \quad \tau_2 = \phi m^3.$$

Proof. From Theorem 5 we have

$$\kappa(M^{-1}A\Theta A^\top) \leq \left(1 + m(m-1) \frac{\bar{\lambda}(D)}{\underline{\lambda}(\Theta_{\mathcal{T}})}\right) \left(1 + m(n-m) \frac{\bar{\lambda}(\Theta_{\bar{\mathcal{T}}})}{\underline{\lambda}(\Theta_{\mathcal{T}})}\right).$$

By the definition of D and (25) it follows that $\bar{\lambda}(D) \leq \phi(n-m)\bar{\lambda}(\Theta_{\bar{\mathcal{T}}})$. Hence

$$\begin{aligned} \kappa(M^{-1}A\Theta A^\top) &\leq \left(1 + \phi m^2(n-m) \frac{\bar{\lambda}(\Theta_{\bar{\mathcal{T}}})}{\underline{\lambda}(\Theta_{\mathcal{T}})}\right) \left(1 + m(n-m) \frac{\bar{\lambda}(\Theta_{\bar{\mathcal{T}}})}{\underline{\lambda}(\Theta_{\mathcal{T}})}\right) \\ &\leq 1 + (\phi m^2 + m)(n-m) \frac{\bar{\lambda}(\Theta_{\bar{\mathcal{T}}})}{\underline{\lambda}(\Theta_{\mathcal{T}})} + \\ &\quad \phi m^3(n-m)^2 \left(\frac{\bar{\lambda}(\Theta_{\bar{\mathcal{T}}})}{\underline{\lambda}(\Theta_{\mathcal{T}})}\right)^2. \end{aligned}$$

The upper bounds are obtained by using this inequality an the same arguments of Corollary 2. $\square$

This last result shows that, even in the presence of primal degeneracy, the condition number of $M^{-1}A\Theta A^\top$ is uniformly bounded after μ becomes small. Furthermore, when the problem is nondegenerate, $\kappa(M^{-1}A\Theta A^\top) \rightarrow 1$ as $\mu \rightarrow 0$. This nice property of the MST and DCMST preconditioners at the later stages of the interior point methods has been confirmed in practice.

6. Computational experience

In this section, we illustrate the theoretical results presented in this paper with examples solving linear systems with five different preconditioners: diagonal, maximum spanning tree (MST), and diagonally compensated maximum spanning tree (DCMST) using three values for the parameter ϕ (.1, 1, and 10).

Each preconditioner was tested on six instances taken from the First DIMACS Implementation Challenge [12] that were solved with PDNET in [19]: a 512-node instance from the Grid-Density-8 class, a 512-node instance from the Grid-Density-16 class, a 514-node instance from the Grid-Long class, a 514-node instance from the Grid-Wide class, a 1024-node instance from the Mesh-1 class, and a 512-node instance from the Netgen-Lo class.

For each instance, the preconditioners were used to solve three linear systems. The first system corresponds to an early interior point iteration (iteration 1). The second system is that of an intermediate iteration (iteration 15). The third system is of a late iteration (iteration 30). For each problem, a set of linear systems was generated by running PDNET with the default settings in [19]. In this way, identical systems were solved with each of the preconditioners. The conjugate gradient method stops when the norm of the residual satisfies a predetermined tolerance (we use the tolerance 10^{-10}). A limit of 1000 conjugate gradient iterations was imposed, with the purpose of detecting numerical difficulties.

Figures 9, 10, and 11 show the residual as a function of the conjugate gradient iteration for the 512-node instance from the Grid-Density-8 class and Figures 12, 13, and 14 show the residual as a function of the conjugate gradient iteration for the 512-node instance from the Grid-Density-16 class. On these six linear systems, the maximum spanning tree preconditioner is the only preconditioner to lower the residual monotonically. Diagonal compensation of the maximum spanning tree preconditioner degrades performance. Pure diagonal preconditioning can only solve the instances from iteration 1.

Figures 15, 16, and 17 show the residual as a function of the conjugate gradient iteration for the 514-node instance from the Grid-Long class. Figures 18, 19, and 20 show the residual as a function of the conjugate gradient iteration for the 514-node instance from the Grid-Wide class. Figures 3, 4, and 5 show the residual as a function of the conjugate gradient iteration for the 1024-node instance from the Mesh-1 class. On these instances, diagonal compensation appears to help more on the cases for which the diagonal preconditioner is able to solve the system, i.e. the iteration 1 systems. When the diagonal preconditioner fails to solve the system, diagonal compensation degrades performance. Again, on eight of these nine instances, the only preconditioner able to decrease the residual monotonically is the maximum spanning tree. The only exception is on the iteration 30 Mesh-1 instance. The nonmonotone convergence patterns of the diagonally compensated preconditioners mimic those of the diagonal preconditioner.

Figures 6, 7, and 8 show the residual as a function of the conjugate gradient iteration for the 512-node instance from the Netflo-Lo class. The behavior of the diagonally compensated preconditioners when the diagonal preconditioner is able to solve the linear system is similar to the previous cases. Again, on these three instances, the only preconditioner able to monotonically

cally decrease the residual is the maximum spanning tree. Oddly, even when the diagonal preconditioner fails to solve the system (iteration 30 instance), diagonal compensation appears to help.

7. Concluding remarks

As suggested by the theoretical bounds developed for the system matrix condition number, the MST preconditioner dominates the diagonal preconditioner in practice. However, countering the theoretical relationships of the condition number upper bounds, the DCMST preconditioner is sometimes slightly more effective than MST.

The computational experiments suggest that the effectiveness of the DCMST preconditioner depends strongly on the behavior of the diagonal preconditioner applied to the same linear system. In those cases where a PCG with a diagonal preconditioner is able to solve the linear system, a diagonal compensation appears to improve the performance of the maximum spanning tree preconditioner. For the instances considered, the diagonal preconditioner was only effective during the early interior point iterations and only for some instances.

The MST preconditioner appears to be the most robust. It was the only preconditioner to monotonically reduced the residual on all linear systems considered. This fact, along with the strong dependence of the DCMST preconditioners on the suitability of the diagonal preconditioner, leads us to conclude that the loss of robustness of the DCMST preconditioners outweighs any gains that may be achieved with these preconditioners.

References

1. Adler, I., N. Karmarkar, M. Resende, and G. Veiga: 1989, 'Data structures and programming techniques for the implementation of Karmarkar's algorithm'. *ORSA Journal on Computing* **1**, 84–106.
2. Ahuja, N. K., T. L. Magnanti, and J. B. Orlin: 1993, *Network Flows*. Englewood Cliffs, NJ: Prentice Hall.
3. Axelsson, O.: 1994, *Iterative Solution Methods*. Cambridge University Press.
4. Baryamureeba, V. and T. Steihaug: 2000, 'On the convergence of an inexact primal-dual interior point method for linear programming'. Technical report, Department of Informatics - University of Bergen, Bergen - Norway.
5. Bellavia, S. and M. Maconni: 1999, 'An inexact interior-point method for monotone NCP'. Technical report, University of Firenze, Italy.
6. Dembo, R. S., S. C. Eisenstat, and T. Steihaug: 1982, 'Inexact Newton methods'. *SIAM Journal on Numerical Analysis* **19**, 400–408.
7. Dembo, R. S. and T. Steihaug: 1983, 'Truncated-Newton algorithms for large scale unconstrained optimization'. *Mathematical Programming* **26**, 190–212.

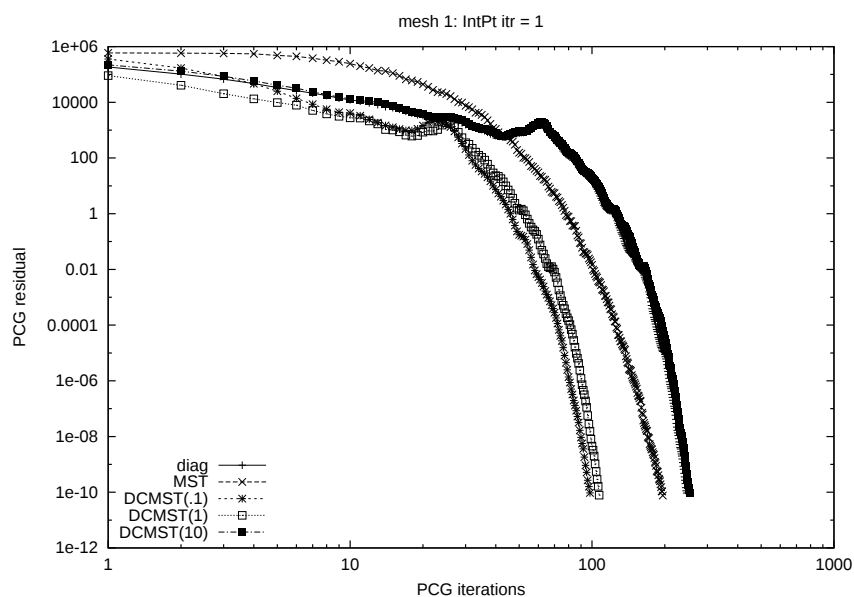


Figure 3. Convergence of PCG on 1024-node mesh-1 instance on interior point iteration 1.

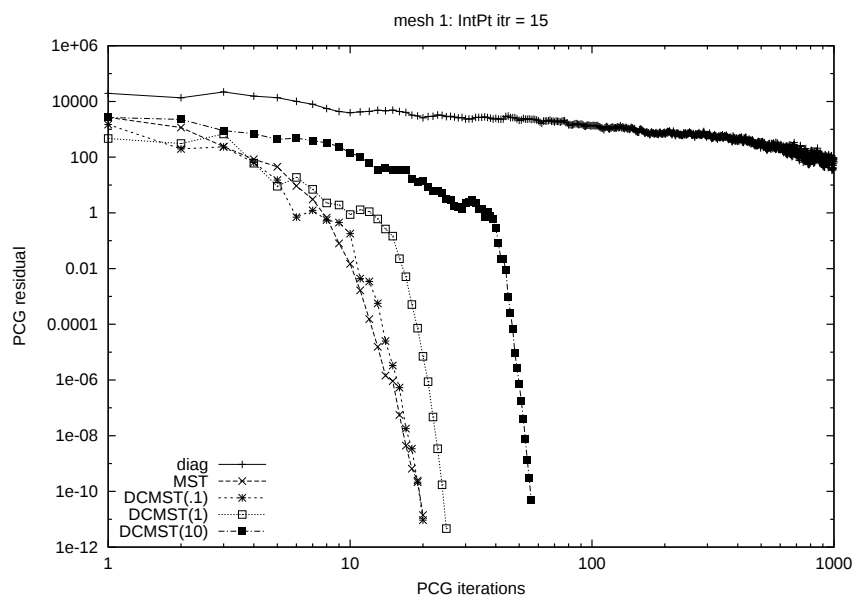


Figure 4. Convergence of PCG on 1024-node mesh-1 instance on interior point iteration 15.

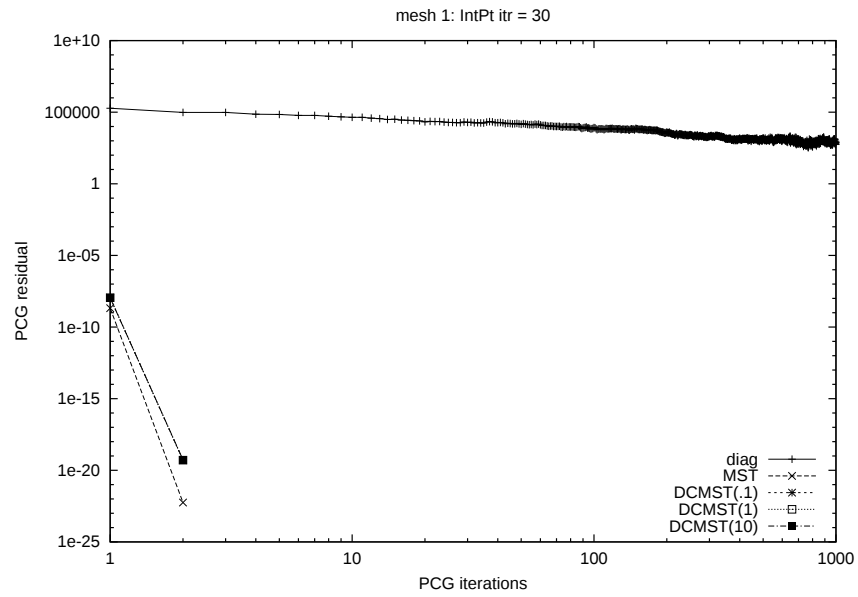


Figure 5. Convergence of PCG on 1024-node mesh-1 instance on interior point iteration 30.

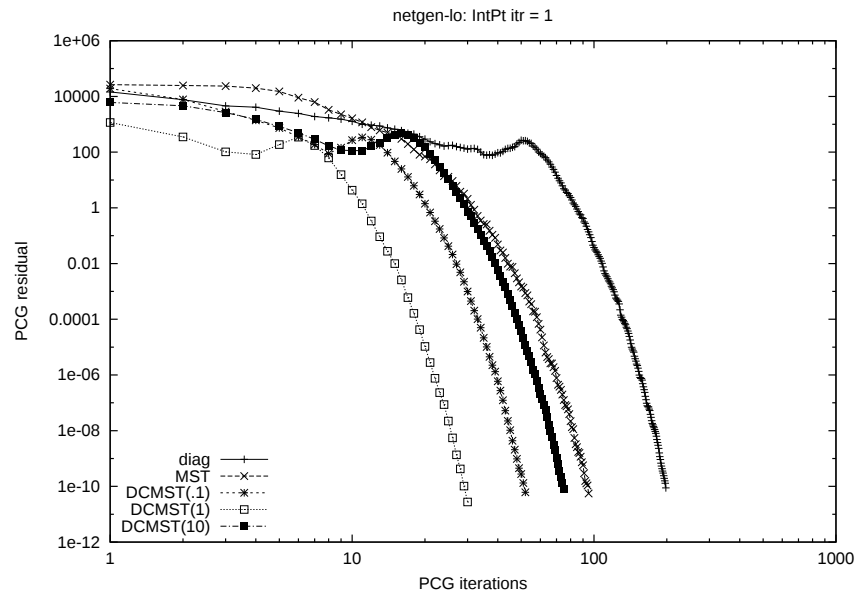


Figure 6. Convergence of PCG on 512-node netgen-lo instance on interior point iteration 1.

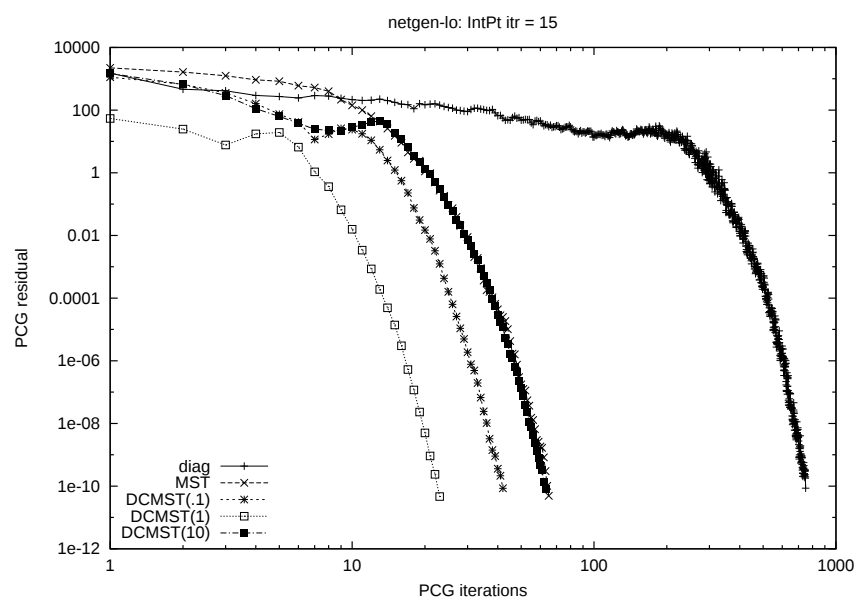


Figure 7. Convergence of PCG on 512-node netgen-lo instance on interior point iteration 15.

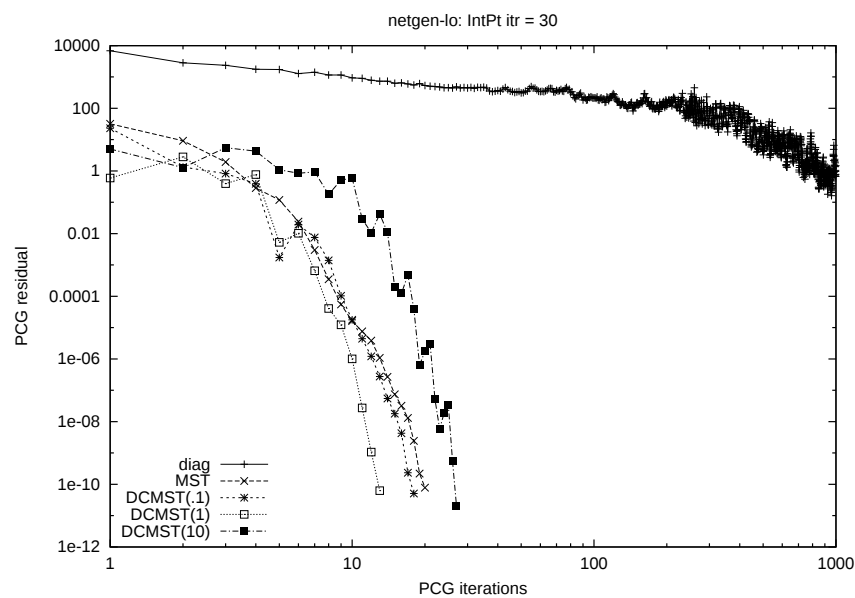


Figure 8. Convergence of PCG on 512-node netgen-lo instance on interior point iteration 30.

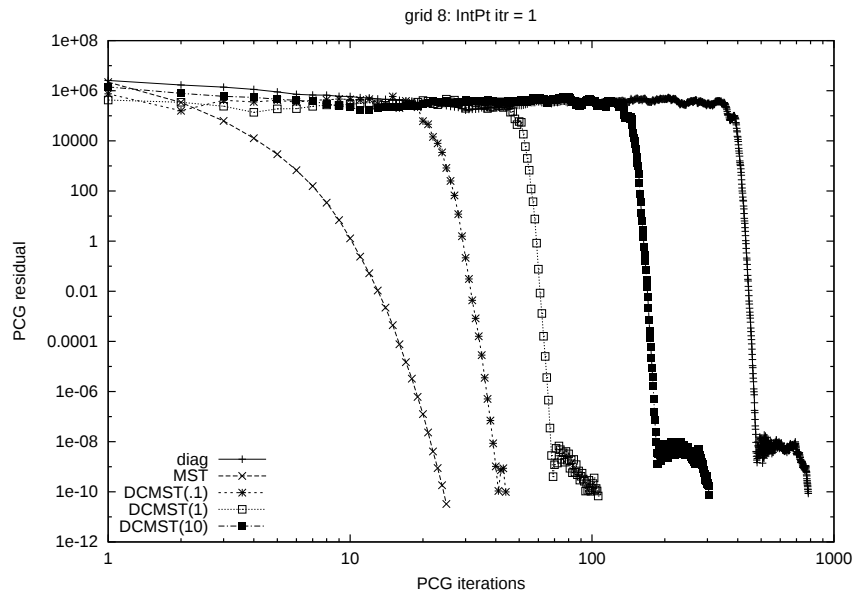


Figure 9. Convergence of PCG on 512-node grid-8 instance on interior point iteration 1.

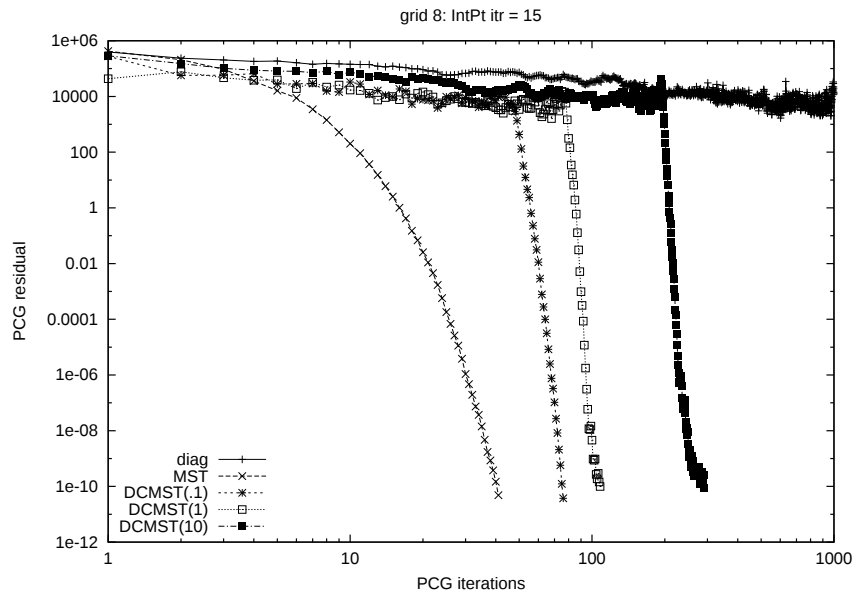


Figure 10. Convergence of PCG on 512-node grid-8 instance on interior point iteration 15.

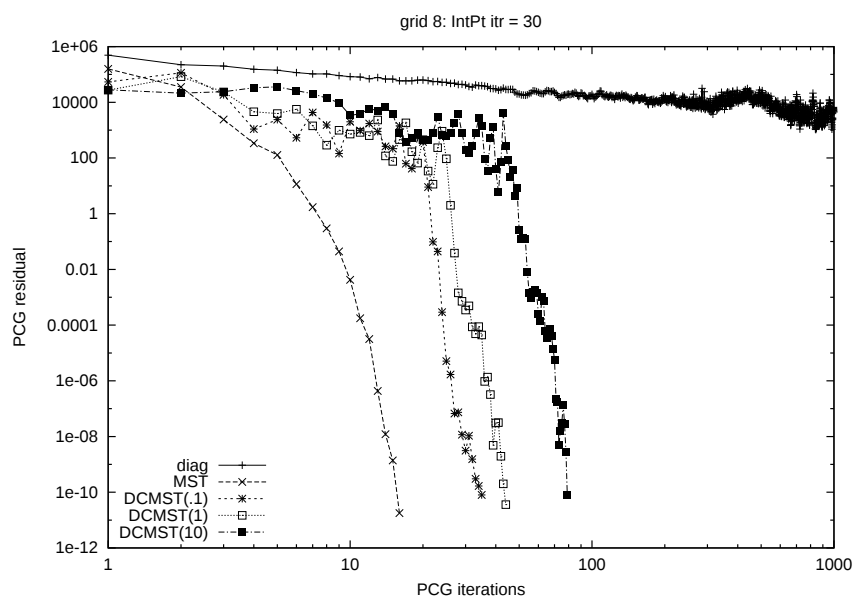


Figure 11. Convergence of PCG on 512-node grid-8 instance on interior point iteration 30.

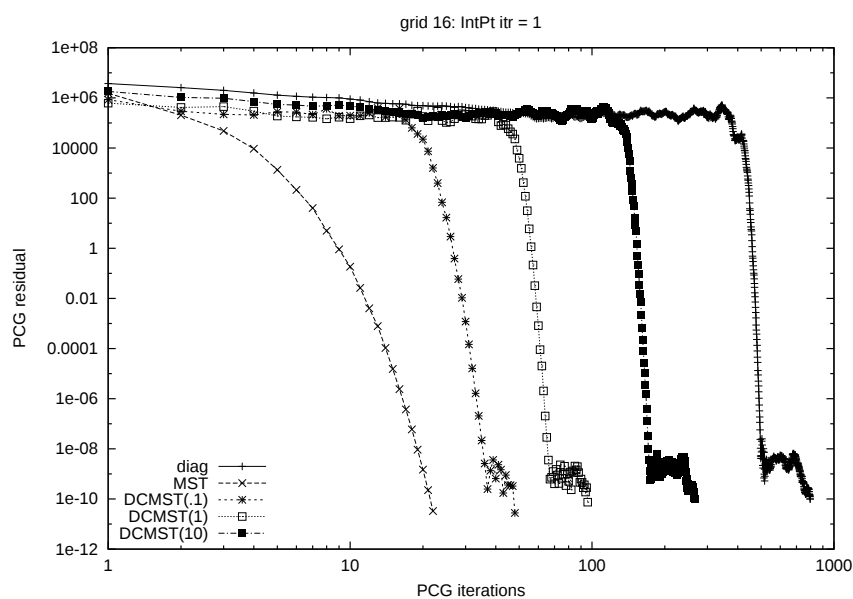


Figure 12. Convergence of PCG on 512-node grid-16 instance on interior point iteration 1.

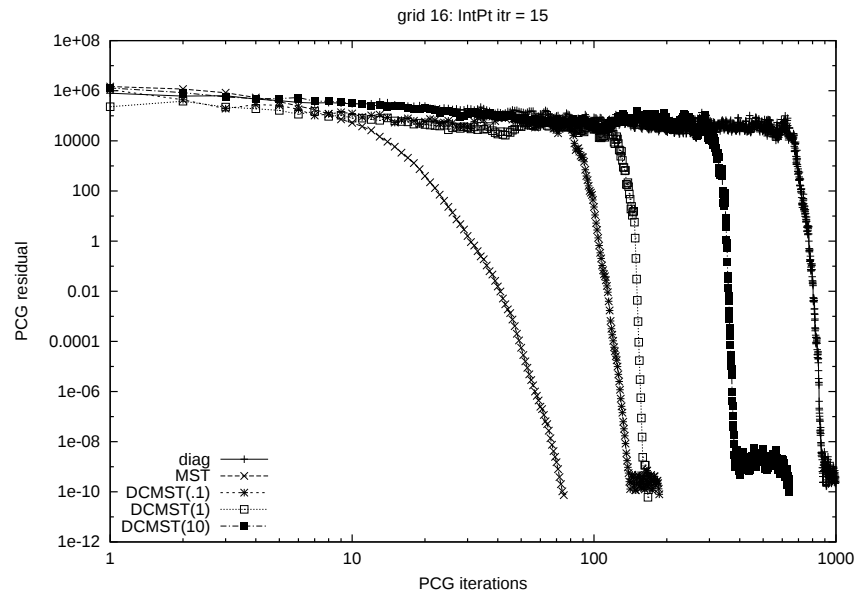


Figure 13. Convergence of PCG on 512-node grid-16 instance on interior point iteration 15.

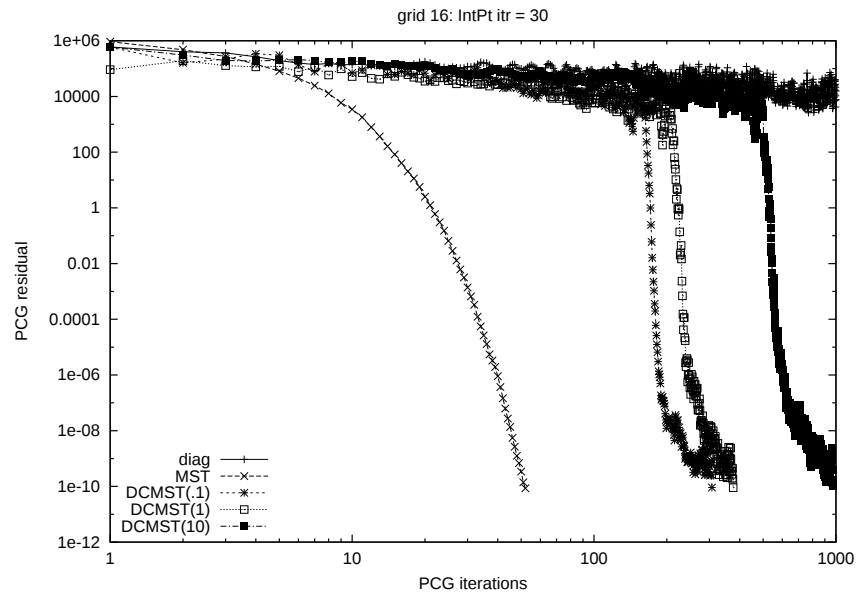


Figure 14. Convergence of PCG on 512-node grid-16 instance on interior point iteration 30.

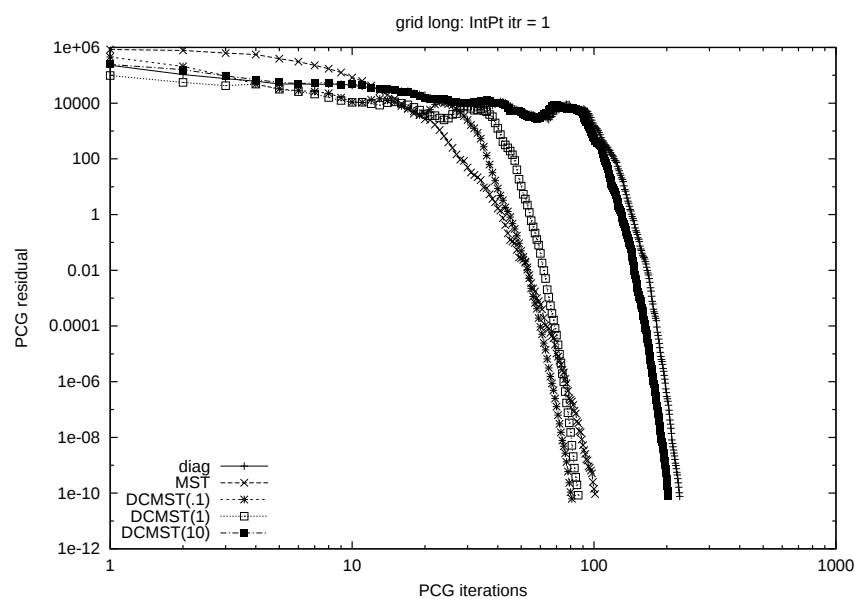


Figure 15. Convergence of PCG on 514-node grid-long instance on interior point iteration 1.

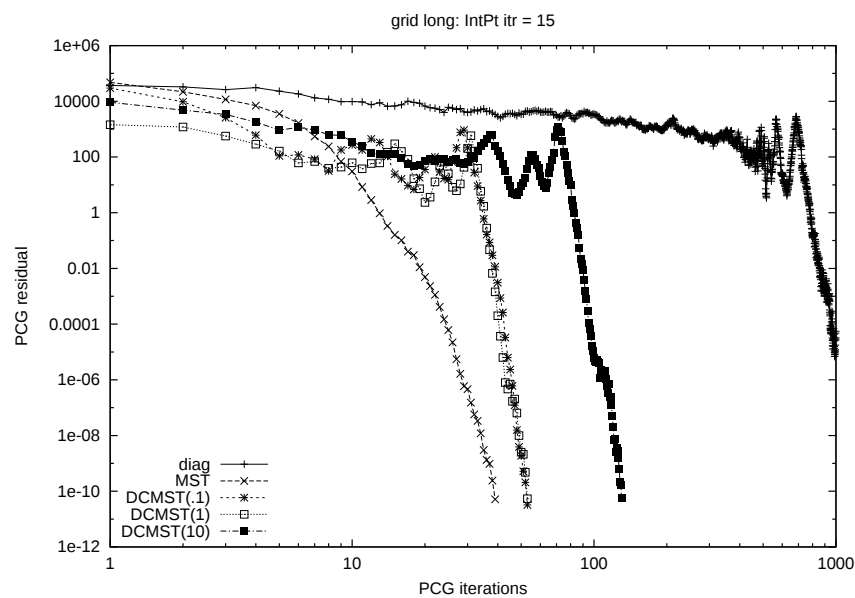


Figure 16. Convergence of PCG on 514-node grid-long instance on interior point iteration 15.

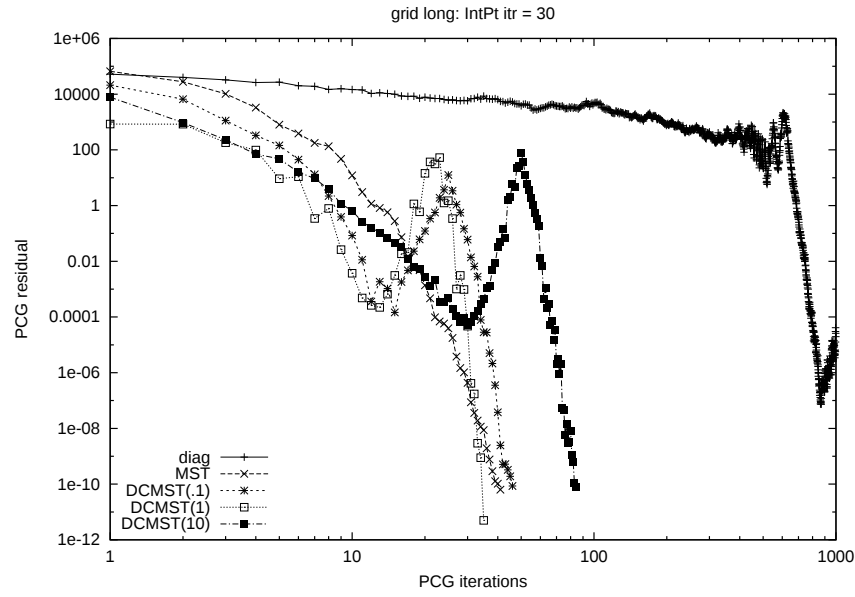


Figure 17. Convergence of PCG on 514-node grid-long instance on interior point iteration 30.

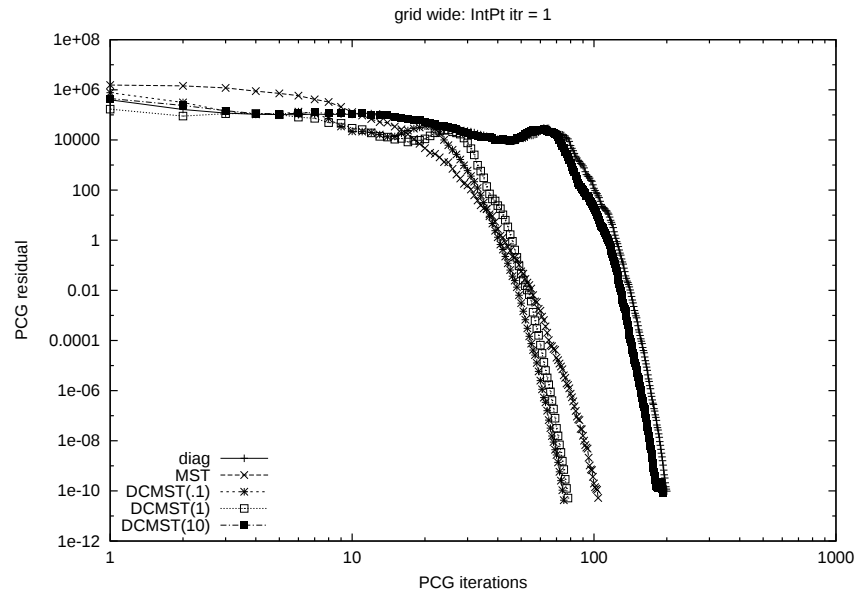


Figure 18. Convergence of PCG on 514-node grid-wide instance on interior point iteration 1.

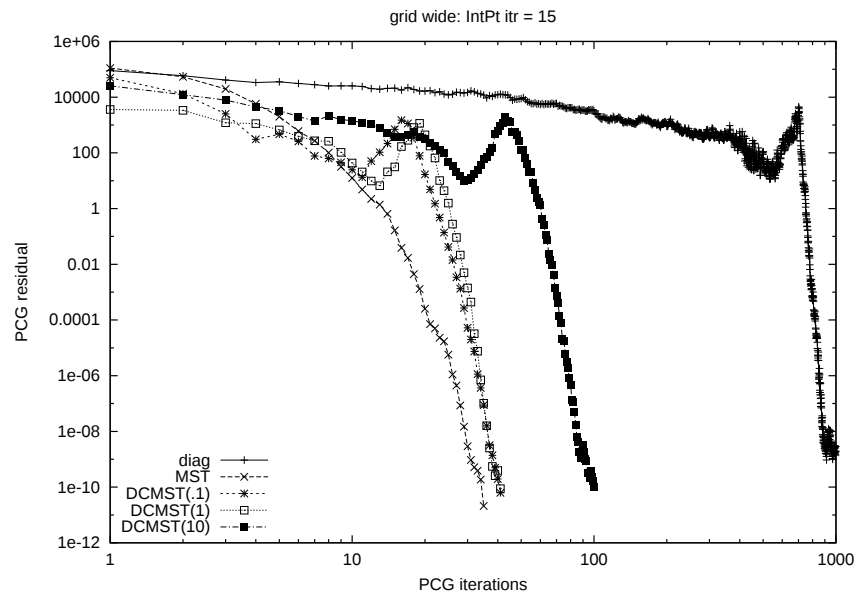


Figure 19. Convergence of PCG on 514-node grid-wide instance on interior point iteration 15.

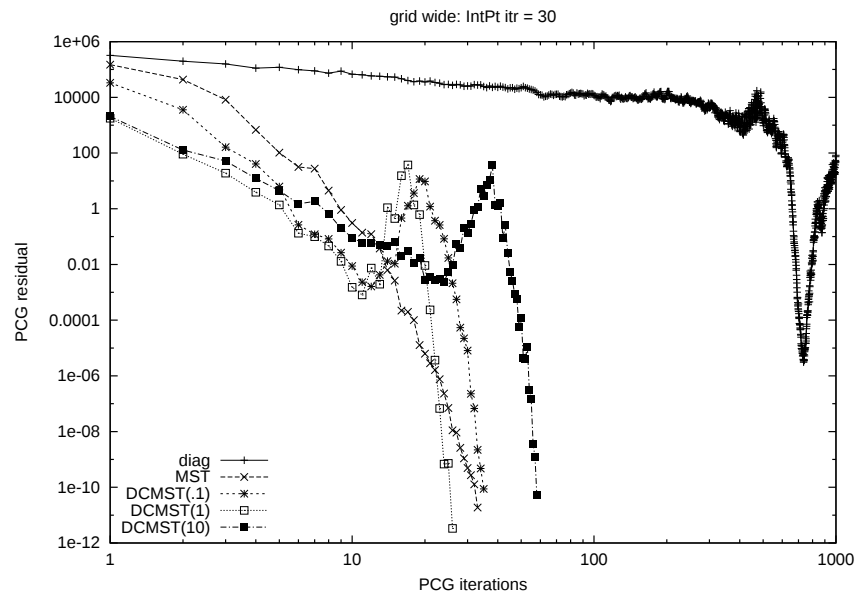


Figure 20. Convergence of PCG on 514-node grid-wide instance on interior point iteration 30.

8. Fiacco, A. and G. McCormick: 1968, *Nonlinear programming: Sequential unconstrained minimization technique*. New York: Wiley.
9. Freund, R., F. Jarre, and S. Mizuno: 1996, 'Convergence of a class of inexact interior-point algorithms for linear programs'. Technical report, Lucent Technologies, 600 Mountain Avenue, Murray Hill, New Jersey 07974, USA.
10. Golub, G. and C. van Loan: 1989, *Matrix Computations*. Baltimore, MD: The Johns Hopkins University Press.
11. Güler, O., D. den Hertog, C. Roos, T. Terlaky, and T. Tsuchiya: 1993, 'Degeneracy in interior point methods for linear programming: A survey'. *Annals of Operations Research* **46**, 107–138.
12. Johnson, D. and C. McGeoch (eds.): 1993, *Network flows and matching: First DIMACS implementation challenge*, Vol. 12 of *DIMACS Series in Discrete Mathematics and Theoretical Computer Science*. American Mathematical Society.
13. Kojima, M., N. Megiddo, and S. Mizuno: 1993, 'A primal-dual infeasible-interior-point algorithm for linear programming'. *Mathematical Programming* **61**, 263–280.
14. McShane, K., C. Monma, and D. Shanno: 1989, 'An implementation of a primal-dual interior point method for linear programming'. *ORSA Journal on Computing* **1**, 70–83.
15. Megiddo, N.: 1989, 'Pathways to the optimal set in linear programming'. In: N. Megiddo (ed.): *Progress in Mathematical Programming, Interior Point and Related Methods*. Springer, pp. 131–158.
16. Mehrotra, S. and J. Wang: 1995, 'Conjugate gradient based implementation of interior point methods for network flow problems'. In: L. Adams and J. Nazareth (eds.): *Linear and Nonlinear Conjugate Gradient Related Methods*.
17. Mizuno, S. and F. Jarre: 1996, 'Global and polynomial time convergence of an infeasible interior point algorithm using inexact computation'. Research Memorandum 605, Institute of Statistical Mathematics, Tokyo, USA.
18. Portugal, L., F. Bastos, J. Júdice, J. Paixão, and T. Terlaky: 1996, 'An investigation of interior point algorithms for the linear transportation problem'. *SIAM Journal on Scientific Computing* **17**, 1202–1223.
19. Portugal, L., M. Resende, G. Veiga, and J. Júdice: 2000, 'A truncated primal-infeasible dual-feasible interior point network flow method'. *Networks* **35**, 91–108.
20. Resende, M. and G. Veiga: 1993a, 'An efficient implementation of a network interior point method'. In: D. S. Johnson and C. C. McGeoch (eds.): *Network Flows and Matching: First DIMACS Implementation Challenge*, Vol. 12 of *DIMACS Series in Discrete Mathematics and Theoretical Computer Science*. American Mathematical Society, pp. 299–348.
21. Resende, M. and G. Veiga: 1993b, 'An implementation of the dual affine scaling algorithm for minimum cost flow on bipartite uncapacitated networks'. *SIAM Journal on Optimization* **3**, 516–537.
22. Schrijver, A.: 1986, *Theory of linear and integer programming*. John Wiley & Sons.
23. Wright, S.: 1994, 'Stability of linear algebra computations in interior point methods for linear programming'. Preprint ANL/MCS- P446-0694, Mathematics and Computer Science Division, Argonne National Laboratory, Argonne, IL 60439, USA.
24. Wright, S.: 1996, 'Modified Cholesky factorizations in interior point algorithms for linear programming'. Preprint ANL/MCS- P600-0596, Mathematics and Computer Science Division, Argonne National Laboratory, Argonne, IL 60439, USA.
25. Wright, S.: 1997, *Primal dual interior point methods*. SIAM.
26. Yeh, Q.-J.: 1989, 'A reduced dual affine scaling algorithm for solving assignment and transportation problems'. Ph.D. thesis, Columbia University, New York, NY.

27. Zhang, Y.: 1994, 'On the convergence of a class of infeasible interior point methods for the horizontal linear complementarity problem'. *SIAM Journal on Optimization* **4**, 208–227.