
Improved Large-Scale Graph Learning through Ridge Spectral Sparsification

Daniele Calandriello^{1,2} Ioannis Koutis³ Alessandro Lazaric⁴ Michal Valko¹

Abstract

The representation and learning benefits of methods based on graph Laplacians, such as *Laplacian smoothing* or *harmonic function solution for semi-supervised learning* (SSL), are empirically and theoretically well supported. Nonetheless, the exact versions of these methods scale poorly with the number of nodes n of the graph. In this paper, we combine a spectral sparsification routine with Laplacian learning. Given a graph $\mathcal{G}$ as input, our algorithm computes a sparsifier in a *distributed* way in $\mathcal{O}(n \log^3(n))$ time, $\mathcal{O}(m \log^3(n))$ work and $\mathcal{O}(n \log(n))$ memory, using only $\log(n)$ rounds of communication. Furthermore, motivated by the regularization often employed in learning algorithms, we show that constructing sparsifiers that preserve the spectrum of the Laplacian *only up to* the regularization level may drastically reduce the size of the final graph. By constructing a spectrally-similar graph, we are able to bound the error induced by the sparsification for a variety of downstream tasks (e.g., SSL). We empirically validate the theoretical guarantees on Amazon co-purchase graph and compare to the state-of-the-art heuristics.

1. Introduction

Graphs are a very effective data structure to represent relationships between entities (e.g., social and collaboration networks, influence graphs). Over the years, many machine learning problems have been defined and solved exploiting the graph representation, such as *graph-regularized least squares* (LAPRLS, Belkin et al. 2005), *Laplacian smoothing* (LAPSMO, Sadhanala et al. 2016) *graph semi-supervised learning* (SSL, Chapelle et al. 2010; Zhu et al. 2003), *laplacian embedding* (LE, Belkin & Niyogi 2001, and *spectral*

clustering (SC, Von Luxburg 2007). The intuition behind graph-based learning is that the information expressed by the graph helps to capture the underlying structure of the problem (e.g., a manifold), thus improving the learning. For instance, LAPSMO and SSL rely on the assumption that nodes that are *close* in the graph are more likely to have similar labels. Similarly, LE and SC try to find a low-dimensional representation of the nodes using the eigenvectors of the Laplacian of the graph. In general, given a graph $\mathcal{G}$ of n nodes and m edges, most of graph-based learning tasks require computing the minimum of a cost function based on the associated $n \times n$ Laplacian matrix $\mathbf{L}_{\mathcal{G}}$, which contains m non-zero entries. Solving *exactly* such optimization problems amounts to $\mathcal{O}(n^3)$ time and $\mathcal{O}(n^2)$ space complexity in the worst case and they become infeasible even for mildly large/dense graphs.

A complete review of the literature on large-scale graph learning is beyond the scope of this paper and we only consider methods that reduce learning space and time complexity starting from a given graph received as input.¹ We identify mainly three possible approaches. We can (1) reduce runtime replacing the pseudo-inverse operator $\mathbf{L}_{\mathcal{G}}^+$ with an *iterative solver*, (2) reduce time and space complexity replacing the large graph $\mathcal{G}$ with a sparser approximation $\mathcal{H}$, or (3) reduce runtime and increase memory capacity by *distributing* the computation across multiple machines.

Iterative solvers. Iterative methods can solve a number of learning problems without explicitly constructing $\mathbf{L}_{\mathcal{G}}^+$ (e.g., gradient descent, GD, for LAPSMO, iterative averaging for SSL, and the power method for SC). In this case we only need $\mathcal{O}(m)$ time per iteration. Unfortunately, all simple iterative methods (e.g., GD) converge in a number of iterations proportional to the condition number of the Laplacian, $\kappa = \lambda_{\max}(\mathbf{L}_{\mathcal{G}})/\lambda_{\min}(\mathbf{L}_{\mathcal{G}})$, which may grow linearly with the number of nodes n , thus removing the advantage of the iterative method, whose complexity tends to $\mathcal{O}(n^3)$ in the worst case. Advanced iterative methods, such as the *preconditioned conjugate gradient*, use preconditioning to find an accurate solution in a number of iterations independent of κ . Koutis et al. (2011) gives a nearly-linear solver for Laplacians or *strongly diagonally dominant* (SDD) matri-

¹SequeL team, INRIA Lille - Nord Europe, France ²LCSL, IIT, Italy, and MIT, USA. ³New Jersey Institute of Technology, USA ⁴Facebook AI Research, Paris, France. Correspondence to: Daniele Calandriello <daniele.calandriello@iit.it>.

¹Many algorithms reduce the complexity of graph learning *at construction* time but they cannot be applied to *natural* graphs (e.g., social graphs) and therefore we do not review them.

ces, that using a chain of preconditioners, converges in only $\mathcal{O}(m \log(n))$ time. As space and time costs scale with the number of edges, a natural desire is to reduce m by sparsifying and distributing the graph.

Graph sparsification. The objective of sparsification methods is to remove *redundant* edges, so that the resulting sparse sub-graph can be easily stored in memory and efficiently manipulated to compute final solutions. A simple *graph-sparsification* technique is to sample $n\bar{q}$ (with $\bar{q} > 1$) edges from $\mathcal{G}$ with probabilities proportional to the edge weights with replacement. While computationally very efficient, uniform sampling requires sampling a number of edges proportional to $\mathcal{O}(n\mu(\mathcal{G}))$ (i.e., $\bar{q} \propto \mu(\mathcal{G})$), where $\mu(\mathcal{G})$ is the *coherence* of the Laplacian matrix, and it can grow as large as n when the graph is highly structured (e.g., if there is a single edge e connecting two components of the graph we need to sample all of the edges of the graph—potentially $\mathcal{O}(n^2)$ —to guarantee that we do not exclude e and generate an inappropriate $\mathcal{H}$). A more refined approach is the k -neighbors (kN) sparsifier (Sadhanala et al., 2016), which performs local sparsifications node-by-node by keeping all edges at nodes with degree smaller than $\bar{q}$, and samples them proportionally to their weights whenever the degree is bigger than $\bar{q}$. While in certain structured graphs, this method may perform much better than uniform (Von Luxburg et al., 2014), in the general case $\bar{q}$, still needs to scale with the coherence $\mu(\mathcal{G})$. A more effective method is to sample edges proportionally to their *effective resistance*, which intuitively measures the importance of an edge in preserving the minimum distance between two nodes. As a result, only relevant edges are kept and the sparsified graph could be reduced to $\mathcal{O}(n \text{polylog}(n))$ edges. Nonetheless, computing effective resistances also requires the pseudo-inverse $\mathbf{L}_{\mathcal{G}}^+$, thus being as expensive as solving any graph-Laplacian learning problem.

Distributed computing. When the number of edges m is too large to fit the whole graph in a single machine, we are forced to distribute the edges across multiple machines. At the same time, if the sparsifier construction or the downstream inference can be parallelized, we can also reduce their runtime. Unfortunately, distributing data and computation across multiple machines can cause large communication costs. For example, simple GD or label propagation methods require $\mathcal{O}(\kappa)$ iterations (and communication rounds) to converge and access to non-local (e.g., neighbors in a graph) data. While preconditioned solvers reduce the number of iterations, almost none of their memory access is local, thus making difficult to have efficient distributed implementations.

Contribution. In this paper, we propose a new approach that aims at integrating the benefits of the three different methods above. Using the large memory and computational capacity of distributed computing and leveraging the sequen-

tial sparsification methods of Kelner & Levin (2013) and Candriello et al. (2017), we show how to compute an accurate sparsifier $\mathcal{H}$ of graph $\mathcal{G}$ in $\mathcal{O}(n \log^3(n))$ time, $\mathcal{O}(n \log^2(n))$ work and $\mathcal{O}(n \log(n))$ memory, using only $\log(n)$ rounds of communication. Afterwards, learning tasks can be solved directly on $\mathbf{L}_{\mathcal{H}}$ on a single machine using near-linear time solvers, resulting in an overall $\mathcal{O}(n \log^3(n))$ runtime. Moreover, we show that the regularization used in some graph-based learning algorithms allows using even sparser graphs. In particular, we introduce the notion of *ridge* effective resistance to obtain sparsifiers that are better adapted to solve Laplacian-regularized learning tasks (e.g., LAPSMO, SSL) and are smaller than standard spectral sparsifiers without compromising the performance of downstream tasks.

2. Background

We use lowercase letters a for scalars, bold lowercase letters $\mathbf{a}$ for vectors and uppercase bold letters $\mathbf{A}$ for matrices. We use $\mathbf{A} \preceq \mathbf{B}$ to denote that $\mathbf{B} - \mathbf{A}$ is positive semi-definite (PSD), $[\mathbf{A}]_{i,j}$ to indicate the (i, j) -th entry of $\mathbf{A}$, and ordered the eigenvalues as $\lambda_1(\mathbf{A}) \leq \dots \leq \lambda_n(\mathbf{A})$.

2.1. Graphs and graph Laplacian

We denote with $\mathcal{G} = (\mathcal{V}, \mathcal{E})$, an undirected weighted graph with n nodes $\mathcal{V}$ and m edges $\mathcal{E}$. Each edge $e_{i,j} \in \mathcal{E}$ has a weight $a_{e_{i,j}}$ measuring the “similarity” between nodes i and j . Given graphs $\mathcal{G}$ and $\mathcal{G}'$ over the same set of nodes $\mathcal{V}$, $\mathcal{G} + \mathcal{G}'$ denotes the graph obtained by summing the weights of their edges. For graph $\mathcal{G}$, we introduce the weighted adjacency matrix $\mathbf{A}_{\mathcal{G}}$ with entries $[\mathbf{A}_{\mathcal{G}}]_{i,j} = a_{e_{i,j}}$, the total weights $A = \sum_e a_e$, and the diagonal degree matrix $\mathbf{D}_{\mathcal{G}}$ with entries $[\mathbf{D}_{\mathcal{G}}]_{i,i} \triangleq \sum_j a_{e_{i,j}}$. The Laplacian of $\mathcal{G}$ is the PSD matrix $\mathbf{L}_{\mathcal{G}} \triangleq \mathbf{D}_{\mathcal{G}} - \mathbf{A}_{\mathcal{G}}$. Furthermore, we assume that $\mathcal{G}$ is connected and thus $\mathbf{L}_{\mathcal{G}}$ has only one eigenvalue equal to 0 and $\text{Ker}(\mathbf{L}_{\mathcal{G}}) = \mathbf{1}$. Let $\mathbf{L}_{\mathcal{G}}^+$ be the pseudoinverse of $\mathbf{L}_{\mathcal{G}}$ and $\mathbf{L}_{\mathcal{G}}^{-1/2} = (\mathbf{L}_{\mathcal{G}}^+)^{1/2}$. For any node $i = 1, \dots, n$, we denote with $\chi_i \in \mathbb{R}^n$, the indicator vector, so that $\mathbf{b}_e \triangleq \sqrt{a_e}(\chi_i - \chi_j)$ is the “edge” vector. If we denote with $\mathbf{B}_{\mathcal{G}}$ the $m \times n$ signed edge-vertex incidence matrix, then the Laplacian can be written as $\mathbf{L}_{\mathcal{G}} = \sum_e \mathbf{b}_e \mathbf{b}_e^T = \mathbf{B}_{\mathcal{G}}^T \mathbf{B}_{\mathcal{G}}$.

2.2. Learning on graphs

Given graph $\mathcal{G}$ and its Laplacian $\mathbf{L}_{\mathcal{G}}$, we denote with $\mathbf{f} \in \mathbb{R}^n$, a *labeling* of its nodes, where $[\mathbf{f}]_i$ is the value associated with the i -th node. Many graph learning algorithms assume that the optimal labeling $\mathbf{f}^*$ is *smooth* w.r.t. the graph, i.e., the quantity $\sum_e a_e ([\mathbf{f}^*]_{e_i} - [\mathbf{f}^*]_{e_j})^2 = \mathbf{f}^{*T} \mathbf{L}_{\mathcal{G}} \mathbf{f}^*$ is small. In the following, we review examples from the supervised, semi-supervised and unsupervised learning with graphs.

Laplacian smoothing (LAPSMO) with Gaussian noise.

Given a graph $\mathcal{G}$ on n nodes, let $\mathbf{y} \triangleq \mathbf{f}^* + \xi$ be a noisy

measurement of $\mathbf{f}^*$ with $[\xi]_i \sim \mathcal{N}(0, \sigma^2)$. The goal of LAPSMO is to find a vector $\hat{\mathbf{f}}$ that accurately reconstructs $\mathbf{f}^*$ under the graph smoothness assumption by solving

$$\begin{aligned}\hat{\mathbf{f}} &\triangleq \arg \min_{\mathbf{f} \in \mathbb{R}^n} (\mathbf{f} - \mathbf{y})^\top (\mathbf{f} - \mathbf{y}) + \lambda \mathbf{f}^\top \mathbf{L}_G \mathbf{f} \\ &= (\lambda \mathbf{L}_G + \mathbf{I})^{-1} \mathbf{y},\end{aligned}\quad (1)$$

where λ is a regularization parameter.

Graph semi-supervised learning (SSL). In SSL, the input $\mathbf{f}_\ell$ is a partial observation of the labels $\mathbf{f}^*$ for a subset $\mathcal{S} \subset [n]$ of nodes. The goal is to predict the labels $\mathbf{f}_u$ of the unrevealed nodes. The *harmonic function solution* (HFS) by Zhu et al. (2003) solves the optimization problem

$$\begin{aligned}\hat{\mathbf{f}}_{\text{HFS}} &\triangleq \arg \min_{\mathbf{f} \in \mathbb{R}^n} \frac{1}{\ell} (\mathbf{f} - \mathbf{y})^\top \ell_S (\mathbf{f} - \mathbf{y}) + \lambda \mathbf{f}^\top \mathbf{L}_G \mathbf{f} \\ &= (\lambda \ell \mathbf{L}_G + \mathbf{I}_S)^\dagger \mathbf{y}_S,\end{aligned}\quad (2)$$

where $\ell \triangleq |\mathcal{S}|$ is the number of labeled nodes received as input, $\mathbf{I}_S \in \mathbb{R}^{n \times n}$ is the identity matrix with zeros at nodes not in $\mathcal{S}$, and $\mathbf{y}_S \triangleq \mathbf{I}_S \mathbf{y} \in \mathbb{R}^n$. Similarly, in *local transductive regression* (LTR) (Cortes et al., 2008), the optimization problem is

$$\begin{aligned}\hat{\mathbf{f}}_{\text{LTR}} &\triangleq \arg \min_{\mathbf{f} \in \mathbb{R}^n} (\mathbf{f} - \mathbf{y})^\top \mathbf{C} (\mathbf{f} - \mathbf{y}) + \mathbf{f}^\top (\mathbf{L}_G + \lambda \mathbf{I}) \mathbf{f} \\ &= (\mathbf{C}^{-1} (\mathbf{L}_G + \lambda \mathbf{I}) + \mathbf{I})^{-1} \mathbf{y}_S,\end{aligned}\quad (3)$$

where $\mathbf{C}$ is a diagonal matrix with entries c_ℓ for nodes in $\mathcal{S}$, c_u for entries not in $\mathcal{S}$, and $c_\ell \geq c_u > 0$.

Spectral clustering (SC). Applying the Laplacian smoothness assumption, the goal of SC is to find k disjoint subset assignments such that the clusters are smooth w.r.t. the Laplacian. Let $\{\mathbf{f}_c\}_{c=1}^k$ be the cluster indicator vectors such that $[\mathbf{f}_c]_i \triangleq 1$ if node i is in the c -th cluster and $[\mathbf{f}_c]_i \triangleq 0$ otherwise. Denote with $\mathbf{F} \in \mathbb{R}^{n \times k}$, the matrix containing the assignments, and let $\mathcal{C}$ be the space of feasible clustering, such that all $\mathbf{f}_c$ are binary and each row of $\mathbf{F}$ contains only one non-zero entry. Since computing the minimum ratio-cut is NP-hard (Von Luxburg, 2007; Lee et al., 2014), even under constraints (Cucuringu et al., 2016), SC defines instead the relaxed problem

$$\hat{\mathbf{F}} \triangleq \arg \min_{\mathbf{F}: \mathbf{F}^\top \mathbf{F} = \mathbf{I}_k, \mathbf{f}_c \perp \mathbf{1}} \text{Tr}(\mathbf{F}^\top \mathbf{L} \mathbf{F}).$$

Once the relaxed solution is computed, we can use different heuristics to recover the clustering, such as thresholding or performing a k -means clustering on the $\hat{\mathbf{F}}$ matrix.

Computational complexity. The problems above require either to compute an eigendecomposition of the Laplacian $\mathbf{L}_G$ or to solve a linear system involving $\mathbf{L}_G$. Computing these exactly is not feasible when the number of nodes n and edges m grows. In particular, (a) storing $\mathbf{L}_G$ in memory

requires $\mathcal{O}(m)$ space, and it is not feasible when m is large, (b) even if $\mathbf{L}_G$ is sparse and m is small, the pseudo-inverse $\mathbf{L}_G^+$ might be dense, and thus computing and storing $\mathbf{L}_G^+$ exactly requires up to $\mathcal{O}(n^3)$ time and $\mathcal{O}(n^2)$ space.

3. Distributed Spectral Sparsification

In this section, we describe a new, sequential, distributed, and efficient algorithm for graph sparsification that can be used as a preprocessing step to solve a large variety of downstream learning tasks, without significantly affecting their performance. We point out that while distributing data-agnostic sparsifiers (e.g. uniform sampling) is straightforward, distributing the computation of sparsifiers based on effective resistances requires a careful merging procedure to guarantee satisfactory memory vs. accuracy tradeoff, which is what we provide in this section.

3.1. (ε, γ) -spectral sparsifiers

We start with the introduction of the notion of (ε, γ) -sparsifier that is adapted for the learning tasks that use sparsified graph Laplacian.

Definition 1. A (ε, γ) -spectral sparsifier of $\mathcal{G}$ is a re-weighted sub-graph $\mathcal{H} \subseteq \mathcal{G}$ whose Laplacian $\mathbf{L}_{\mathcal{H}}$ satisfies

$$(1 - \varepsilon) \mathbf{L}_G - \varepsilon \gamma \mathbf{I} \preceq \mathbf{L}_{\mathcal{H}} \preceq (1 + \varepsilon) \mathbf{L}_G + \varepsilon \gamma \mathbf{I}. \quad (4)$$

For $\gamma = 0$, this definition reduces to the standard notion of ε -spectral sparsifier (Spielman & Teng, 2011). The main difference is that an (ε, γ) -spectral sparsifier allows for an extra *additive* error of order $\varepsilon \gamma$. This change is directly motivated by the fact that the sparsifier $\mathcal{H}$ may be used in learning tasks whose solution may not be sensitive to small (additive) errors. As a result, (ε, γ) -spectral sparsifiers are able to further reduce the size of $\mathcal{H}$ w.r.t. $(\varepsilon, 0)$ -sparsifiers, without significantly affecting the final learning performance. Formally, an ε -sparsifier preserves all the quadratic forms up to a small multiplicative (constant) error, and thus can be used to provide an accurate approximation to many important quantities such as graph cuts or eigenvalues. In fact, for all $i \in [n]$, an ε -sparsifier guarantees that $(1 - \varepsilon) \lambda_i(\mathbf{L}_G) \leq \lambda_i(\mathbf{L}_{\mathcal{H}}) \leq (1 + \varepsilon) \lambda_i(\mathbf{L}_G)$. Nonetheless, in many learning tasks (e.g., LTR) the noise level in the signal $\mathbf{f}$ requires regularizing the solution so that the Laplacian $\mathbf{L}_G$ itself is eventually replaced by $\mathbf{L}_G + \lambda \mathbf{I}$ (e.g., Eq. 1). This corresponds to *soft-thresholding* the eigenvalues of the Laplacian, so that eigenvalues below λ are partially ignored. If λ is properly tuned w.r.t. the noise, the regularization increases stability and improves the learning performance. Therefore, constructing a sparsifier that accurately reconstructs *all* eigenvalues of $\mathbf{L}_G$ may be wasteful, as it may require keeping most of the edges. As a result, in tasks where $\mathbf{L}_G$ is regularized, it is better to use (ε, γ) -sparsifiers, as their additive error $\gamma \mathbf{I}$ is homogeneous with the regu-

larization $\lambda \mathbf{I}$ and their smaller size allows scaling to up.² We now extend the results of Spielman & Srivastava (2011) for the construction of (ε, γ) -sparsifiers. We redefine the edge effective resistance to account for the regularization.

Definition 2. The γ -effective resistance of an edge e in graph $\mathcal{G}$ is defined as

$$r_e(\gamma) \triangleq \mathbf{b}_e^\top (\mathbf{L}_{\mathcal{G}} + \gamma \mathbf{I})^{-1} \mathbf{b}_e. \quad (5)$$

The “effective dimension” of the graph is the total sum of the γ -effective resistances, $d_{\text{eff}}(\gamma) \triangleq \sum_e r_e(\gamma)$.

We can now construct a sparsifier $\mathcal{H}$ by sampling $\bar{q}$ times each edge with a probability proportional to its γ -effective resistance. More formally, the resulting (random) graph contains $q_e \sim \mathcal{B}(r_e(\lambda); \bar{q})$ copies of each edge, where $\mathcal{B}$ is the Binomial distribution, and its associated Laplacian is $\mathbf{L}_{\mathcal{H}} = \sum_{e \in \mathcal{H}} q_e / (\bar{q} r_e(\gamma)) \mathbf{b}_e \mathbf{b}_e^\top$, which is an unbiased estimator of $\mathbf{L}_{\mathcal{G}}$. We can then apply existing results from sketching of PSD matrices (Alaoui & Mahoney, 2015) to prove that $\mathcal{H}$ is a valid (ε, γ) -sparsifier.

Proposition 1 (Cohen et al. 2017). Let $\varepsilon > 0$ and $\gamma \geq 0$ be the accuracy parameters and $0 \leq \delta \leq 1$ the probability of error. Let $\mathcal{H}$ be the graph obtained by sampling edges in $\mathcal{G}$ with a probability proportional to their γ -effective resistances. If $\bar{q} \geq 4 \log(4n/\delta)/\varepsilon^2$, then w.p. $1 - \delta$, $\mathcal{H}$ is an (ε, γ) -sparsifier with $\mathcal{O}(d_{\text{eff}}(\gamma)\bar{q})$ edges.

We first notice that this result reduces to the one of Spielman & Srivastava (2011) for $\gamma = 0$. In fact, $d_{\text{eff}}(0) = n - 1$ for all graphs, thus matching the space requirement $\bar{q}$ for ε -sparsifiers. Nonetheless, as γ increases, the size of $\mathcal{H}$ reduces significantly. Using $\mathbf{L}_{\mathcal{G}} = \mathbf{B}_{\mathcal{G}}^\top \mathbf{B}_{\mathcal{G}}$, the effective dimension $d_{\text{eff}}(\gamma)$ can be conveniently rewritten as

$$d_{\text{eff}}(\gamma) = \text{Tr}(\mathbf{B}_{\mathcal{G}}^\top \mathbf{B}_{\mathcal{G}} (\mathbf{B}_{\mathcal{G}}^\top \mathbf{B}_{\mathcal{G}} + \gamma \mathbf{I})^{-1}) = \sum_{i=1}^n \frac{\lambda_i(\mathbf{L}_{\mathcal{G}})}{\lambda_i(\mathbf{L}_{\mathcal{G}}) + \gamma},$$

thus showing that $d_{\text{eff}}(\gamma)$ is the “soft” rank of the Laplacian, where γ significantly reduces the contribution of small eigenvalues to the total sum. While in the worst case $d_{\text{eff}}(\gamma)$ can be as large as $n - 1$, for a variety of graphs with rapidly decaying spectrum (Jamakovic & Miegheem, 2006; Samukhin et al., 2008; Zhan et al., 2010; Akoglu et al., 2015), $d_{\text{eff}}(\gamma)$ may be significantly smaller than $n - 1$, thus reducing the number of edges $\bar{q}$ required to obtain an (ε, γ) -sparsifier.

3.2. The algorithm

As pointed out in the introduction, the main limitation of effective-resistance-based sparsification is that the computation of r_e requires inverting the Laplacian matrix, thus resulting in a computational cost that already matches the cost

²Whenever no regularization is required in the learning task (i.e., HFS, SC), we set $\gamma = 0$ and consider “standard” ε -sparsifiers.

Algorithm 1 The DiSRe algorithm.

Input: $\mathcal{G}$

Output: $\mathcal{H}_{\mathcal{G}}$

- 1: Partition $\mathcal{G}$ into k sub-graphs:
- 2: $\mathcal{H}_{1,\ell} \leftarrow \mathcal{G}_\ell \leftarrow \{(e_{i,j}, q_e = 1, \tilde{p}_{1,e} = 1)\}$
- 3: Initialize set $\mathcal{S}_1 = \{\mathcal{H}_{1,\ell}\}_{\ell=1}^k$
- 4: **for** $h = 1, \dots, k - 1$ **do**
- 5: Pick two sparsifiers $\mathcal{H}_{h,i}, \mathcal{H}_{h,i'}$ from $\mathcal{S}_h$
- 6: $\bar{\mathcal{H}} \leftarrow \text{Merge-Resparsify}(\mathcal{H}_{h,i}, \mathcal{H}_{h,i'})$
- 7: Place $\bar{\mathcal{H}}$ back into $\mathcal{S}_{h+1}$
- 8: **end for**
- 9: Return $\mathcal{H}_{\mathcal{G}}$, the last sparsifier in $\mathcal{S}_k$

Algorithm 2 Merge-Resparsify

Require: (ε, γ) -sparsifiers $\mathcal{H}_{h,i}, \mathcal{H}_{h,i'}$ of graphs $\mathcal{G}_{h,i}, \mathcal{G}_{h,i'}$

Ensure: $\bar{\mathcal{H}}$, an (ε, γ) sparsifier of $\mathcal{G}_{h,i} + \mathcal{G}_{h,i'}$

- 1: Initialize $\bar{\mathcal{H}} = \mathcal{H}_{h,i} + \mathcal{H}_{h,i'}$
- 2: For all $e \in \bar{\mathcal{H}}$, use a fast SDD solver to compute

$$\tilde{r}_{h+1,e}(\gamma) \leftarrow (1 - \varepsilon) \mathbf{b}_e^\top (\mathbf{L}_{\bar{\mathcal{H}}} + (1 + \varepsilon) \gamma \mathbf{I})^{-1} \mathbf{b}_e$$

- 3: Set probabilities $\tilde{p}_{h+1,e} \leftarrow \min\{\tilde{r}_{h+1,e}(\gamma), \tilde{p}_{h,e}\}$
- 4: Sample $q_{h+1,e}$ from $\mathcal{B}(\tilde{p}_{h+1,e}/\tilde{p}_{h,e}, q_{h,e})$
- 5: Return $\bar{\mathcal{H}} \leftarrow \{(e_{i,j}, q_{h+1,e}, \tilde{p}_{h+1,e})\}$ for all $q_{h+1,e} > 0$

of the learning tasks themselves. Moreover, large graphs cannot be stored in memory, and multiple passes over the graph would result in a disk access overhead larger than the computational cost. In order to avoid these problems, we adapt our previous work (Calandriello et al., 2017) in online sparsification and randomized linear algebra (see a thorough discussion and comparison at the end of the section) to obtain the distributed sequential resparsification (DiSRe) algorithm (Alg. 1).³

The structure. We represent a sparsifier $\mathcal{H}$ as a collection of weighted edges $\mathcal{H} \triangleq \{(e_{i,j}, q_e, \tilde{p}_e)\}$, and the Laplacian can be reconstructed as $\mathbf{L}_{\mathcal{H}} \triangleq \sum_{e \in \mathcal{H}} 1/\tilde{p}_e (q_e/\bar{q}) \mathbf{b}_e \mathbf{b}_e^\top$. Intuitively, each edge e has an associated weight based on its probability $\tilde{p}_e$, and a number of included copies q_e . Keeping multiple copies of each edge helps the random $\mathbf{L}_{\mathcal{H}}$ to concentrate towards $\mathbf{L}_{\mathcal{G}}$, where the maximum number of copies $\bar{q}$ for an edge trades-off success probability and the size of $\mathcal{H}$. We assume we have k machines. DiSRe begins by partitioning the graph $\mathcal{G}$ into k sub-graphs $\mathcal{G}_\ell$ on n vertices and $m_\ell \geq n$ edges, such that $\mathcal{G} = \{\mathcal{G}_\ell\}_{\ell=1}^k$. In other words, it splits the matrix $\mathbf{B}_{\mathcal{G}}$ into submatrices $\mathbf{B}_{\mathcal{G}_i}$ by arbitrarily selecting a subset of rows. The sub-graphs are small enough that they can be stored in memory,⁴ and they are

³Whenever the original graph contains $m \leq \tilde{\mathcal{O}}(d_{\text{eff}}(\gamma))$ edges, there is no need to run DiSRe as the (ε, γ) -sparsifiers would not reduce the size of the graph.

⁴Whenever this is not possible (i.e., m/k is too large to be

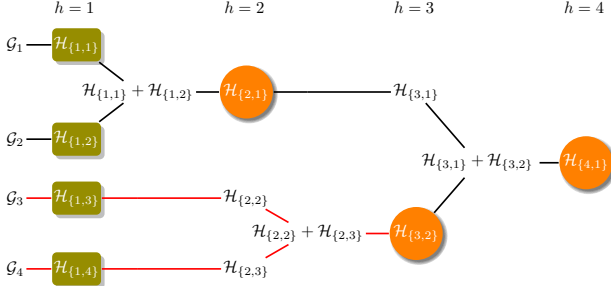


Figure 1. Merge tree for Alg. 1.

also obviously sparsifiers of themselves, therefore we can define an initial set of sparsifiers $\mathcal{S}_1 \triangleq \{\mathcal{H}_{1,\ell}\}_{\ell=1}^k$, with $\mathcal{H}_{1,\ell} \triangleq \{(e_{i,j}, q_{1,e} = \bar{q}, \tilde{p}_{1,e} = 1)\}_{e \in \mathcal{G}_i}$. With this definition, $\mathcal{H}_{1,\ell}$ contains edges $e_{i,j}$ with unit weight $\tilde{p}_{1,e} = 1$ and $\mathcal{H}_{1,\ell} = \mathcal{G}_\ell$. Starting from these initial sparsifiers, **DiSRe** proceeds through a sequence of *merge* and *sparsify* operations where two sparsifiers are first combined and then sparsified again to keep having manageable-size graphs at each step. While **DiSRe** can run on any arbitrary sequence of merges, we consider the most (computationally) effective scheme, where sparsifiers are merged two-by-two in parallel, thus inducing a *balanced* full binary merge tree (see Fig. 2). For notational convenience, we consider that at each iteration h , the inner loop of Alg. 1 only merges two arbitrary sparsifiers from the pool of available sub-graphs $\mathcal{S}_h$ and merges them into a new sparsifier. In practice, multiple merge-and-sparsify operations can be executed in a parallel and asynchronous way. The size of $\mathcal{S}_h$, number of sparsifiers present at layer h , is $|\mathcal{S}_h| = k - h + 1$. Therefore, a node in the tree corresponding to a sparsifier is uniquely identified by two indices $\{h, \ell\}$ where h is the height of the layer and $\ell \leq |\mathcal{S}_h|$ is the index of the node in the layer. We also define the graph $\mathcal{G}_{\{h,\ell\}}$ as the union of all sub-graphs $\mathcal{G}_{\ell'}$ that are reachable from node $\{h, \ell\}$ as leaves (descendants of $\{h, \ell\}$). For example, in Fig. 2, sparsifier $\mathcal{H}_{3,1}$ in node $\{3, 1\}$ approximates the graph $\mathcal{G}_{\{3,2\}} = \mathcal{G}_3 + \mathcal{G}_4$, where we highlight in red the descendant tree. **The resparsification.** In Alg. 2 we detail how two arbitrary sparsifiers are combined to obtain a temporary graph $\bar{\mathcal{H}}$. While the merge operation simply combines $\mathcal{H}_{h,i}$ and $\mathcal{H}_{h,i'}$ by summing their weights, the resparsification aims at generating a valid sparsifier from the “original” sub-graph $(\mathcal{G}_{h,i} + \mathcal{G}_{h,i'})$, as if it was directly sparsified at the beginning. We first compute estimates $\tilde{r}(\gamma)$ of the γ -effective resistance by using fast solvers to invert the strongly diagonal dominant $L_{\bar{\mathcal{H}}} + \gamma \mathbf{I}$ matrix. Instead of sampling edges in $\bar{\mathcal{H}}$ directly proportionally to $\tilde{r}(\gamma)$ (more precisely $\tilde{p}_{h+1,e}$), we perform a “resampling” scheme where an edge e is preserved with a “reweighted” probability $\tilde{p}_{h+1,e}/\tilde{p}_{h,e}$. Intuitively, the overall

stored on a single machine), we can simply apply the same merging scheme of **DiSRe** by loading *small enough* chunks of the graph and sparsifying them sequentially.

sequence of resampling guarantees that at each step $h + 1$, an edge $e \in (\mathcal{G}_{h,i} + \mathcal{G}_{h,i'})$ has the “correct” probability $\tilde{p}_{h+1,e}$ of being included in the sparsifier.

Performance. We now study the performance of **DiSRe** and its complexity. *Time* complexity refers to the amount of time necessary to compute the final solution and *work* complexity refers to the total amount of operations carried out by *all* machines to compute the final solution.

Theorem 1. *Let $\varepsilon > 0$ be the accuracy, $0 \leq \delta \leq 1$ the probability of error, and $\rho \triangleq (1 + 3\varepsilon)/(1 - \varepsilon)$. Given an arbitrary graph $\mathcal{G}$ and an arbitrary merge tree structure, if **DiSRe** is run with parameter $\bar{q} \triangleq 26\rho \log(3n/\delta)/\varepsilon^2$, then each sub-graphs $\mathcal{H}_{\{h,\ell\}}$ is an (ε, γ) -sparsifier of $\mathcal{G}_{\{h,\ell\}}$ with at most $3\bar{q}d_{\text{eff}}(\gamma)$ edges with probability $1 - \delta$. Whenever the merge tree is balanced and k is big enough such that $m/k \leq 3\bar{q}d_{\text{eff}}(\gamma)$,⁵ then merge operations can be run in parallel across the machines with an overall time complexity of $\mathcal{O}(d_{\text{eff}}(\gamma) \log^3(n))$, a total work $\mathcal{O}(m \log^3(n))$, and $\mathcal{O}(\log(n))$ rounds of communication.*

Discussion. [Kelner & Levin \(2013\)](#) proposed a sequential algorithm for graph sparsification that closely emulates the batch sampling of [Spielman & Srivastava \(2011\)](#) in a semi-streaming setting and incrementally constructs an ε -sparsifier. However, their proof had a flaw since they treated dependent variables as independent ([Calandriello et al., 2016](#)). [Kyng et al. \(2016\)](#) resolved the issues in the proof of [Kelner & Levin \(2013\)](#) and showed that a slightly modified algorithm can construct a sparsifier with $\mathcal{O}(n \log(n)/\varepsilon^2)$ edges in $\mathcal{O}(m \log^2(n)/\varepsilon^2)$ time, matching the space complexity of batch sampling. The method proposed by [Kyng et al. \(2016\)](#) can be further improved by parallelizing its computation over multiple machines. Using the parallel sparsification algorithm of [Koutis & Xu \(2016\)](#), the time complexity can be reduced up to $\tilde{\mathcal{O}}(\log^6(n))$. Nonetheless, since these methods require *random access to the edges*, they cannot be easily distributed (it would have $\mathcal{O}(m \text{polylog}(n))$ communication cost) and scaled to graphs that cannot be stored on a single machine. Furthermore, the algorithm of [Kyng et al. \(2016\)](#) accurately reconstructs the whole spectrum of the Laplacian, which leads to sparsifiers whose number of edges scales linearly with n . On the other hand, in regularized learning tasks, the presence of multiplicative and additive spectral error allows creating smaller sparsifiers whose size scales with $d_{\text{eff}}(\gamma)$. Notice that, for γ large enough, this possibly means sparsifiers with less than $n - 1$ edges, necessarily leading to disconnected graphs. Finally, note that merging two traditional ε -sparsifiers gives an ε -sparsifier, merging two (γ, ε) -sparsifiers produces a less accurate $(2\gamma, \varepsilon)$ -sparsifier. Therefore simple merge-and-reduce strategies ([Feldman et al.,](#)

⁵This implies that there are enough machines so that the leaves in the merge tree already have relatively sparse sub-graphs.

2013), which address every resparsification as independent, would either cumulate errors or require multiple passes over the data. Similarly to Kyng et al. (2016), DiSRé’s sequential Merge-Resparsify solves this problem (Appendix B).

Mixed additive-multiplicative reconstruction is studied more extensively in randomized matrix algebra (Drineas & Mahoney, 2017). Cohen et al. (2016) developed an efficient method to spectrally sparsify generic matrices up to $(1 \pm \varepsilon)$ multiplicative and γ -additive errors using an incremental sampling method based on *ridge leverage scores* (i.e., the analog of γ -effective resistances for matrices). If applied to graph Laplacians, their method adds edges incrementally and returns an (ε, γ) -sparsified graph with $\mathcal{O}(d_{\text{eff}}(\gamma) \log^2(n))$ edges in $\mathcal{O}(m \log(n))$ time. Nonetheless, Cohen et al. (2016) provided only ε -sparsifiers, suggesting to set γ as small as possible, and did not explore the advantages possible in machine learning. Moreover, no existing (ε, γ) -sparsifier construction method can leverage both distribution and fast solvers. Cohen et al. (2016) can only add edges (but not remove them as DiSRé), preventing repeated merge-and-resparsify. Other streaming RLS sampling methods, such as by Cohen et al. (2017), use dense intermediate sketches, such as *frequent directions* (Ghashami et al., 2016), that are not Laplacians of a subgraph and cannot be easily paired with near-linear solvers for Laplacians.

4. Downstream Guarantees

We now show how the spectral reconstruction guarantees provided by (ε, γ) -sparsifiers translate into guarantees on the quality of the approximate solutions computed using $\mathcal{H}$ instead of $\mathcal{G}$. We first introduce a result for ε -sparsifiers in SSL and then show how for regularized problems, (ε, γ) -sparsification can further improve computational performance without loss in accuracy in LAPSMO.

4.1. Generalization bounds for SSL

Given the closed form solutions of HFS (Eq. 2) and LTR (Eq. 3), we simply replace $\mathbf{L}_{\mathcal{G}}$ with $\mathbf{L}_{\mathcal{H}}$ and then run a nearly-linear time solver to obtain approximate solutions $\tilde{\mathbf{f}}_{\text{HFS}}$ and $\tilde{\mathbf{f}}_{\text{LTR}}$. We compare approximate solutions to their exact counterparts in the context of algorithmic stability.

Definition 3. Let $\mathcal{L}$ be a transductive learning algorithm. We denote by $\mathbf{f}$ and $\mathbf{f}'$ the solutions obtained by running $\mathcal{L}$ on datasets $\mathcal{V} \triangleq (\mathcal{S}, \mathcal{T})$ and $\mathcal{V}' \triangleq (\mathcal{S}', \mathcal{T}')$ respectively. $\mathcal{L}$ is uniformly β -stable w.r.t. the squared loss if there exists $\beta \geq 0$ such that for any two partitions $(\mathcal{S}, \mathcal{T})$ and $(\mathcal{S}', \mathcal{T}')$ that differ by exactly one training (and test) point and for all $i \in [n]$, we have $|([\mathbf{f}]_i - [\mathbf{y}]_i)^2 - ([\mathbf{f}']_i - [\mathbf{y}]_i)^2| \leq \beta$.

The stability of LTR was proven by Cortes et al. (2008). On the other hand, the singularity of the Laplacian may

lead to unstable behavior in HFS due to the $(\gamma \ell \mathbf{L}_{\mathcal{G}} + \mathbf{I}_{\mathcal{S}})^+$ pseudo-inverse, with drastically different results for small perturbations of the dataset. For this reason, we take the Stable-HFS algorithm by Belkin et al. (2004), where an additional regularization term is introduced to restrict the space of admissible solutions to the space $\mathcal{F} \triangleq \{\mathbf{f} : \langle \mathbf{f}, \mathbf{1} \rangle = 0\}$ of solutions orthogonal to the null space of $\mathbf{L}_{\mathcal{G}}$ (i.e., centered functions). As shown by Belkin et al. (2004), to satisfy the constraint, it is sufficient to set an additional regularization parameter μ to $\mu \triangleq ((\gamma \ell \mathbf{L}_{\mathcal{G}} + \mathbf{I}_{\mathcal{S}})^+ \mathbf{y}_{\mathcal{S}})^{\top} \mathbf{1} / ((\gamma \ell \mathbf{L}_{\mathcal{G}} + \mathbf{I}_{\mathcal{S}})^+ \mathbf{1})^{\top} \mathbf{1}$, and compute the solution $\hat{\mathbf{f}}_{\text{STA}}$ as $\hat{\mathbf{f}}_{\text{STA}} \triangleq (\gamma \ell \mathbf{L}_{\mathcal{G}} + \mathbf{I}_{\mathcal{S}})^+ (\mathbf{y}_{\mathcal{S}} - \mu \mathbf{1})$. While Stable-HFS is more stable and thus more suited for theoretical analysis, its space and time requirement remains $\mathcal{O}(m)$ and cannot be applied to graphs with a large number of edges. Therefore, we again replace $\hat{\mathbf{f}}_{\text{STA}}$ with an approximate solution $\tilde{\mathbf{f}}_{\text{STA}}$ computed using $\mathbf{L}_{\mathcal{H}}$. Define $\hat{R}(\mathbf{f}) \triangleq \frac{1}{\ell} \sum_{i=1}^{\ell} (\mathbf{f}(x_i) - \mathbf{y}(x_i))^2$ as the empirical error and $R(\mathbf{f}) \triangleq \frac{1}{u} \sum_{i=1}^u (\mathbf{f}(x_i) - \mathbf{y}(x_i))^2$ as the generalization.

Theorem 2. Let $\mathcal{G}$ be a fixed (connected) graph with eigenvalues $0 = \lambda_1(\mathcal{G}) < \lambda_2(\mathcal{G}) \leq \dots \leq \lambda_n(\mathcal{G})$, and $\mathcal{H}$ an ε -sparsifier of $\mathcal{G}$. Let $\mathbf{y} \in \mathbb{R}^n$ be the labels of the nodes in $\mathcal{G}$ with $|\mathbf{y}(x)| \leq c$ and $\mathcal{F}$ be the set of centered functions such that $|\mathbf{f}(x) - \mathbf{y}(x)| \leq 2c$. Let $\mathcal{S} \subset \mathcal{V}$ be a random subset of labeled nodes, if the labels $\mathbf{y}_{\mathcal{S}}$ are centered, then w.p. at least $1 - \delta$ (w.r.t. the random generation of the sparsifier $\mathcal{H}$ and the random subset of labeled points $\mathcal{S}$) the resulting Stable-HFS solution satisfies

$$R(\tilde{\mathbf{f}}) \leq \hat{R}(\hat{\mathbf{f}}) + \beta + \left(2\beta + \frac{4c^2(\ell + u)}{\ell u} \right) \sqrt{\frac{\pi(\ell, u) \ln \frac{1}{\delta}}{2}} + \frac{1}{1 - \varepsilon} \left(\frac{2(1 + \varepsilon)\varepsilon \ell \gamma \lambda_2(\mathcal{G}) c}{((1 - \varepsilon)\ell \gamma \lambda_2(\mathcal{G}) - 1)^2} \right)^2, \quad (6)$$

where $\tilde{\mathbf{f}}$ and $\hat{\mathbf{f}}$ are computed on $\mathcal{H}$ and $\mathcal{G}$,

$$\pi(\ell, u) \triangleq \frac{\ell u}{\ell + u - 0.5} \frac{2 \max\{\ell, u\}}{2 \max\{\ell, u\} - 1} \quad \text{and} \quad \beta \leq \frac{3c\sqrt{\ell}}{((1 - \varepsilon)\ell \gamma \lambda_2(\mathcal{G}) - 1)^2} + \frac{4c}{(1 - \varepsilon)\ell \gamma \lambda_2(\mathcal{G}) - 1}.$$

Thm. 2 (full proof in Appendix A) shows how approximating $\mathcal{G}$ with $\mathcal{H}$ impacts the generalization error as the number of labeled samples ℓ increases. If we set $\varepsilon = 0$, we recover the bound of Cortes et al. (2008), which depends only on $\hat{R}(\hat{\mathbf{f}})$ and β . When $\varepsilon > 0$, we see from Eq. 6 that the two terms already present in the exact case are either unchanged ($\hat{R}(\hat{\mathbf{f}})$) or increase only by a constant factor β . Because of the approximation, a new error term (the last one in Eq. 6) is added to the bound, but we can see that it is negligible compared to β . In fact, it converges to zero as $\mathcal{O}(\varepsilon^2/\ell^2(1 - \varepsilon)^4)$ as ℓ grows and it is dominated by β for any constant value

of ε . This means that increasing ε corresponds to a constant increase in the bound, regardless of the size of the problem. Consequently, ε can be freely chosen to trade off accuracy and space complexity (Thm. 1) depending on the problem constraints. Finally, because the eigenvalues present in the bound are the ones of the original graph, any additional knowledge on the spectral properties of the input graph can be easily included in the analysis. Therefore, it is straightforward to provide stronger guarantees for Sparse-HFS when combined with assumptions on the graph generating model. Finally, we remark the level of generality of this result that holds for the integration between HFS and any ε -accurate spectral sparsification method. We postpone computational considerations to the following subsection.

4.2. Generalization bounds for LAPSMO

Starting from the closed form solution of LAPSMO (Eq. 1) we can replace the $\mathbf{L}_G$ matrix with a sparsified Laplacian $\mathbf{L}_H$ and using a fast linear solver, compute an approximate solution $\tilde{\mathbf{f}} = (\lambda \mathbf{L}_H + \mathbf{I})^{-1} \mathbf{y}$ in $\mathcal{O}(n \log^2(n))$ time and $\mathcal{O}(n \log(n))$ space. Finally, we can decompose the error as $\|\mathbf{f}^* - \tilde{\mathbf{f}}\|_2^2 \leq \|\mathbf{f}^* - \hat{\mathbf{f}}\|_2^2 + \|\hat{\mathbf{f}} - \tilde{\mathbf{f}}\|_2^2$. The first term can be bounded using classical results from empirical process theory (Bühlmann & Van De Geer, 2011). We bound the second term in the following theorem.

Theorem 3. *For an arbitrary graph $\mathcal{G}$ and its (ε, γ) -sparsifier, let $\hat{\mathbf{f}}$ be the LAPSMO solution computed using $\mathbf{L}_G$ and $\tilde{\mathbf{f}}$ the solution computed using $\mathbf{L}_H$. Then,*

$$\|\tilde{\mathbf{f}} - \hat{\mathbf{f}}\|_2^2 \leq \frac{\varepsilon^2}{1 - \varepsilon} (0.25 + \lambda\gamma) \left(\lambda \hat{\mathbf{f}}^\top \mathbf{L}_G \hat{\mathbf{f}} + \lambda\gamma \|\hat{\mathbf{f}}\|_2^2 \right),$$

where λ is the regularization of LAPSMO.

For ε -sparsifiers, Sadhanala et al. (2016) derive a similar bound $\|\tilde{\mathbf{f}} - \hat{\mathbf{f}}\|_2^2 \leq \mathcal{O}(\lambda \hat{\mathbf{f}}^\top \mathbf{L}_G \hat{\mathbf{f}})$. Setting $\gamma = 0$, we recover their bound up to constants. When $\gamma > 0$ instead, additional error terms emerge due to the introduced bias. In particular, the term $\lambda\gamma \|\hat{\mathbf{f}}\|_2^2$ depends on the norm of the exact solution $\hat{\mathbf{f}}$, which in turn depends on the value of λ . Nonetheless, when $\|\mathbf{f}^*\|_2^2$ is small, as is the case in our experiments, setting $\gamma = 1/\lambda$ makes this term a constant, which is reflected by the good empirical performance. Computationally, for both Stable-HFS and LAPSMO, passing from computing a solution on the full graph to computing a solution on the sparsifier reduces the number of edges, which makes the memory and runtime plummet. Moreover, carefully distributing the sparsification process across multiple machines allows computing a final solution in a time independent from the number of edges, since the preprocessing sparsification step takes only $\mathcal{O}(n \log^3(n))$ time, and the solution step only $\mathcal{O}(n \log^2(n))$. Up to logarithmic terms, this results in an overall $\mathcal{O}(n)$ near-linear runtime,

without any assumptions on the input graph. For graphs with a particularly favorable spectrum and problems with enough regularization, this is only $\tilde{\mathcal{O}}(d_{\text{eff}}(\gamma))$, resulting in a potentially sub-linear runtime. This result, only possible due to a particular structure of learning problems, opens up unexplored possibilities that would not be possible for general graph problems.

4.3. Bounds for other problems

Many other problems can be well approximated using (ε, γ) -sparsifiers. For example, the cost of a SC solution evaluated on $\mathbf{L}_H$ is very close to the cost evaluated on $\mathbf{L}_G$.

Proposition 2. *For any rank k orthogonal projection $\mathbf{F}^\top \mathbf{F}$, if $\mathcal{H}$ is an (ε, γ) -sparsifier of $\mathcal{G}$, we have*

$$\text{Tr}(\mathbf{F}^\top \mathbf{L}_H \mathbf{F}) \leq (1 + \varepsilon) \text{Tr}(\mathbf{F}^\top \mathbf{L}_G \mathbf{F}) + \varepsilon\gamma k.$$

Therefore, a clustering that well separates the sparsifier will also separate well the true graph. Similarly, we can obtain strong approximation guarantees for a variety of other Laplacian-based algorithms. Regularized problems such as LTR (Cortes et al., 2008), Laplacian-regularized least squares, and Laplacian SVM (Belkin et al., 2005) are of particular interest since the additive γ error is absorbed by the regularization and it is possible to provide strong generalization guarantees.

5. Experiments

We empirically validate our theoretical findings by testing how (ε, γ) -sparsifiers improves computational complexity without sacrificing final accuracy.

Dataset. We run experiments on the Amazon co-purchase graph (Sadhanala et al., 2016). This graph fits our setting: It cannot be generated from vectorial data and is only artificially sparse, since the crawler that created it had no access to the true private co-purchase network held by Amazon. To compensate, Gleich & Mahoney (2015) use a densification procedure that given the graph adjacency matrix $\mathbf{A}_G$, computes all k -step neighbors $\mathbf{A}_{G,k} \triangleq \sum_{s=1}^k \mathbf{A}_G^s$. We make the graph unweighted for numerical stability. The final graph has $n = 334,863$ nodes and $m = 98,465,352$ edges, with an average degree of 294. We followed an approach similar to Sadhanala et al. (2016) and introduce a hand-designed smooth signal as a target. We then perform 2000 iterations of the power method to compute an approximation of the smallest eigenvector $\mathbf{v}_{\min}$, which is used as a smooth function over the graph.

Baselines. For all setups, we compute an “exact” solution (up to convergence error) using a fast linear solver. Computing this EXACT baseline requires $\mathcal{O}(m \log(n))$ time and $\mathcal{O}(m)$ space and achieves the best performance. Afterwards, we compare three different sparsification procedures to eval-

Improved Large-Scale Graph Learning through Ridge Spectral Sparsification

Alg.	Parameters	$ \mathcal{E} $ ($\times 10^6$)	Err. SSL ($\ell=346$)	Err. SSL ($\ell=672$)	Err. $D(\tilde{\mathbf{f}})(\sigma=10^{-3})$	Err. $D(\tilde{\mathbf{f}})(\sigma=10^{-2})$
EXACT		98.5	0.312 ± 0.022	0.286 ± 0.010	0.067 ± 0.0004	0.756 ± 0.006
kN	$k=60$	15.7	0.329 ± 0.0143	0.311 ± 0.027	0.172 ± 0.0004	0.822 ± 0.002
kN	$k=90$	21.2	0.334 ± 0.024	0.311 ± 0.024	0.125 ± 0.0002	0.811 ± 0.003
DiSRe	$\gamma=0, \bar{q}=100$	15	0.314 ± 0.0165	0.296 ± 0.015	0.068 ± 0.0003	0.758 ± 0.005
DiSRe	$\gamma=0, \bar{q}=150$	22.8	0.314 ± 0.0158	0.310 ± 0.024	0.068 ± 0.0004	0.756 ± 0.005
DiSRe	$\gamma=10^3, \bar{q}=100$	7.3	—	—	0.072 ± 0.0003	0.789 ± 0.005
DiSRe	$\gamma=10^2, \bar{q}=100$	11.8	—	—	0.068 ± 0.0002	0.772 ± 0.004
DiSRe	$\gamma=10, \bar{q}=100$	14.4	—	—	0.068 ± 0.0004	0.760 ± 0.004

Table 1. Results for the SSL and the smoothing problems.

uate if they can accelerate computation while preserving accuracy. We run **DiSRe** with different values of γ depending on the setting. For empirically strong heuristics, we attempted to uniformly subsample the edges, but at the sparsity level achieved by the other methods, the uniformly sampled sparsifier is disconnected and highly inaccurate. Instead, we compare to the state-of-the-art k -neighbors (kN) heuristic by [Sadhanala et al. \(2016\)](#), which is just as fast as uniform sampling and more accurate in practice.

Experimental procedure. We repeat each experiment 10 times with different sparsifiers and report the average performance of $\tilde{\mathbf{f}}$ on the specific task and its standard deviation. More details on experiments are given in the Appendix C.

5.1. Laplacian smoothing with Gaussian noise

We set $\mathbf{f}^* = \mathbf{v}_{\min}$ and test different levels of noise, $\log_{10}(\sigma) \in \{-3, -2, -1, 0\}$. After constructing the sparsifier $\mathcal{H}$, we compute an approximate solution $\tilde{\mathbf{f}}$ using LAPSMO (Eq. 2) with $\lambda \in \{10^{-3}, 10^{-2}, 10^{-1}, 1, 10\}$. We measure the performance by the squared error $D(\tilde{\mathbf{f}}) = \|\mathbf{f}^* - \tilde{\mathbf{f}}\|_2^2$. As $\|\mathbf{f}^*\|_2^2 = \|\mathbf{v}_{\min}\|_2^2 = 1$, good values of $D(\tilde{\mathbf{f}})$ should be below 1.

Accuracy. In the interest of space, in Tab. 1, we report results for $\sigma = \{0.001, 0.01\}$ and the best regularization λ for each method. We first notice that all sparsifiers are considerably smaller than the original graph, keeping only a small fraction of its edges. The smallest sparsifiers are obtained by **DiSRe** when γ is large. The comparison with **DiSRe** with $\gamma = 0$ (i.e., ε -sparsifier) confirms that the additive error translates into an extra compression of the resulting sparsifier. This also impacts the accuracy which degrades as γ increases. Nonetheless, we notice that while ε -sparsifiers perfectly match the accuracy of the exact method, even for large γ (and thus much smaller graph), **DiSRe** still outperforms kN, which has a significantly worse accuracy. Finally, we note that for $\gamma = 0$, the impact of $\bar{q}$ is as expected: Increasing $\bar{q}$ increases the size of the sparsifier and slightly improves the performance.

Computational complexity. All algorithms require 90s to load the graph from disk. The preprocessing phase of kN takes slightly less than 1min, while **DiSRe**’s takes 12min

on 4 machines. For the solving step, EXACT is unsurprisingly the slowest, requiring 12min to compute an $\hat{\mathbf{f}}$ solution. Both kN and (ε, γ) -sparsifiers require 1–2min, depending on the number of edges preserved. Overall, preprocessing the graph with **DiSRe** before computing a solution does not introduce any overhead compared to EXACT (both take roughly 12min). We notice that while kN is overall faster, the time for **DiSRe** could be easily reduced by increasing the number of parallel processes when computing effective resistances or with a better network topology allowing point-to-point communication. Moreover, once we have access to an accurate ε -sparsifier, it is easier to solve problem repeatedly, e.g., to cross-validate regularization. For example, computing a solution for 4 different values of λ (see the appendix) is crucial for good performance and requires 48min for EXACT and only 20min for **DiSRe**. Finally, memory usage is reduced by a factor of 3 as EXACT requires over 30GB of memory to execute while **DiSRe** never exceeds 10GB. We expect these advantages to only grow larger as we scale to larger graphs.

5.2. SSL with harmonic function solution

We also test **DiSRe** on a SSL problem. The labels are generated taking the sign of $\mathbf{f}^* = \mathbf{v}_{\min}$ and $\ell \in \{20, 346, 672, 1000\}$ labels are revealed. The labeled nodes are chosen at random so that 0 and 1 labels are balanced in the dataset. We run **Stable-HFS** with $\lambda \in \{10^{-6}, 10^{-4}, 10^{-2}, 1\}$. In Tab. 1, we report results for $\ell = \{346, 672\}$ and the best λ for each method. We run **DiSRe** with $\gamma = 0$ as **Stable-HFS** does not have any regularization and ε -sparsifiers are preferable. The average size of the sparsifiers is the same as before as they are agnostic to the learning task. Similar to the smoothing case, **DiSRe** achieves a performance that closely approximates the exact solution, despite the significant compression of the original graph. Furthermore, the effectiveness of the ε -sparsifier returned by **DiSRe** is confirmed by its comparison with kN, whose error is significantly worse. Finally, we notice that the computational analysis in the previous section holds for SSL as well. In fact, although the learning task is different, we use the same SSD solver to compute the HFS and thus the running time are comparable in the two tasks.

Acknowledgements

Experiments presented in this paper were carried out using the Grid'5000 testbed, supported by a scientific interest group hosted by Inria and including CNRS, RENATER, and several universities as well as other organizations (see <https://www.grid5000.fr>). The research presented was also supported by European CHIST-ERA project DELTA, French Ministry of Higher Education and Research, Nord-Pas-de-Calais Regional Council, Inria and Otto-von-Guericke-Universität Magdeburg associated-team north-european project Allocate, and French National Research Agency projects ExTra-Learn (n.ANR-14-CE24-0010-01) and BoB (n.ANR-16-CE23-0003). I. Koutis is supported by NSF CAREER award CCF-1149048.

References

- Akoglu, L., Tong, H., and Koutra, D. [Graph based anomaly detection and description: a survey](#). *Data Mining and Knowledge Discovery*, 29(3):626–688, 2015.
- Alaoui, A. E. and Mahoney, M. W. [Fast randomized kernel methods with statistical guarantees](#). In *Neural Information Processing Systems*, 2015.
- Belkin, M. and Niyogi, P. [Laplacian eigenmaps and spectral techniques for embedding and clustering](#). In *Neural Information Processing Systems*, 2001.
- Belkin, M., Matveeva, I., and Niyogi, P. [Regularization and Semi-Supervised Learning on Large Graphs](#). In *Conference on Learning Theory*, 2004.
- Belkin, M., Niyogi, P., and Sindhvani, V. [On manifold regularization](#). In *International conference on Artificial Intelligence and Statistics*, 2005.
- Bühlmann, P. and Van De Geer, S. [Statistics for high-dimensional data: methods, theory and applications](#). Springer Science & Business Media, 2011.
- Calandriello, D., Lazaric, A., and Valko, M. [Analysis of keller and levin graph sparsification algorithm for a streaming setting](#). *arXiv:1609.03769*, 2016.
- Calandriello, D., Lazaric, A., and Valko, M. [Distributed sequential sampling for kernel matrix approximation](#). In *International conference on Artificial Intelligence and Statistics*, 2017.
- Chapelle, O., Schölkopf, B., and Zien, A. [Semi-supervised learning](#). The MIT Press, 2010.
- Cohen, M. B., Musco, C., and Pachocki, J. W. [Online row sampling](#). In *Approximation, Randomization, and Combinatorial Optimization. Algorithms and Techniques*, 2016.
- Cohen, M. B., Musco, C., and Musco, C. [Input sparsity time low-rank approximation via ridge leverage score sampling](#). In *Symposium on Discrete Algorithms*, 2017.
- Cortes, C., Mohri, M., Pechyony, D., and Rastogi, A. [Stability of transductive regression algorithms](#). In *International Conference on Machine Learning*, 2008.
- Cucuringu, M., Koutis, I., Chawla, S., Miller, G., and Peng, R. [Simple and scalable constrained clustering: a generalized spectral method](#). In *International Conference on Artificial Intelligence and Statistics*, 2016.
- Drineas, P. and Mahoney, M. W. [Lectures on randomized numerical linear algebra](#). *arXiv:1712.08880*, abs/1712.08880, 2017.
- Feldman, D., Schmidt, M., and Sohler, C. [Turning big data into tiny data: Constant-size coresets for k-means, pca and projective clustering](#). In *Symposium on Discrete Algorithms*, 2013.
- Ghashami, M., Liberty, E., Phillips, J. M., and Woodruff, D. P. [Frequent directions: Simple and deterministic matrix sketching](#). *SIAM Journal on Computing*, 45(5):1762–1792, 2016.
- Gleich, D. F. and Mahoney, M. W. [Using local spectral methods to robustify graph-based learning algorithms](#). In *Knowledge Discovery and Data Mining*, pp. 359–368, 2015.
- Jamakov, A. and Mieghem, P. V. [The Laplacian spectrum of complex networks](#). In *European Conference on Complex Systems*, 2006.
- Kelner, J. A. and Levin, A. [Spectral sparsification in the semi-streaming setting](#). *Theory of Computing Systems*, 53(2):243–262, 2013.
- Koutis, I. and Xu, S. C. [Simple parallel and distributed algorithms for spectral graph sparsification](#). *Transactions on Parallel Computing*, 3(2):14, 2016.
- Koutis, I., Miller, G. L., and Peng, R. [A nearly-m log n time solver for SDD linear systems](#). In *Symposium on Foundations of Computer Science*, pp. 590–598, 2011.
- Kyng, R. and Sachdeva, S. [Approximate gaussian elimination for laplacians-fast, sparse, and simple](#). In *Foundations of Computer Science*, 2016.
- Kyng, R., Pachocki, J., Peng, R., and Sachdeva, S. [A framework for analyzing resparsification algorithms](#). In *Symposium on Theory of Computing*, 2016.
- Lee, J. R., Gharan, S. O., and Trevisan, L. [Multiway spectral partitioning and higher-order cheeger inequalities](#). *Journal of the ACM*, 61(6):37:1–37:30, 2014.

- Levy, H. *Stochastic dominance: Investment decision making under uncertainty*. Springer, 2015.
- Sadhanala, V., Wang, Y.-X., and Tibshirani, R. [Graph sparsification approaches for laplacian smoothing](#). In *International conference on Artificial Intelligence and Statistics*, 2016.
- Samukhin, A. N., Dorogovtsev, S. N., and Mendes, J. F. F. [Laplacian spectra of, and random walks on, complex networks: Are scale-free architectures really important?](#) *Physical Review E*, 77(3):036115, 2008.
- Spielman, D. A. and Srivastava, N. [Graph sparsification by effective resistances](#). *SIAM Journal on Computing*, 40(6), 2011.
- Spielman, D. A. and Teng, S.-H. [Spectral sparsification of graphs](#). *SIAM Journal on Computing*, 40(4):981–1025, 2011.
- Tropp, J. A. [Freedman’s inequality for matrix martingales](#). *Electronic Communications in Probability*, 16:262–270, 2011.
- Tropp, J. A. [An introduction to matrix concentration inequalities](#). *Foundations and Trends® in Machine Learning*, 8 (1-2):1–230, 2015.
- Von Luxburg, U. [A tutorial on spectral clustering](#). *Statistics and computing*, 17(4):395–416, 2007.
- Von Luxburg, U., Radl, A., and Hein, M. [Hitting and commute times in large random neighborhood graphs](#). *Journal of Machine Learning Research*, 15(1):1751–1798, 2014.
- Zhan, C., Chen, G., and Yeung, L. F. [On the distributions of Laplacian eigenvalues versus node degrees in complex networks](#). *Physica A: Statistical Mechanics and its Applications*, 389(8):1779–1788, 2010.
- Zhu, X., Ghahramani, Z., and Lafferty, J. [Semi-Supervised Learning Using Gaussian Fields and Harmonic Functions](#). In *International Conference on Machine Learning*, 2003.

A. Proofs for the statements in Sec. 4

In the proofs, we make several uses the following reformulation of Def. 1.

Proposition 3. A sub-graph $\mathcal{H}$ is a (ε, γ) -sparsifier of $\mathcal{G}$ if and only if

$$(1 - \varepsilon)\mathbf{L}_{\mathcal{G}} - \varepsilon\gamma\mathbf{I} \preceq \mathbf{L}_{\mathcal{H}} \preceq (1 + \varepsilon)\mathbf{L}_{\mathcal{G}} + \varepsilon\gamma\mathbf{I} \iff \|(\mathbf{L}_{\mathcal{G}} + \gamma\mathbf{I})^{-1/2}(\mathbf{L}_{\mathcal{H}} - \mathbf{L}_{\mathcal{G}})(\mathbf{L}_{\mathcal{G}} + \gamma\mathbf{I})^{-1/2}\|_2^2 \leq \varepsilon.$$

Theorem 2. Let $\mathcal{G}$ be a fixed (connected) graph with eigenvalues $0 = \lambda_1(\mathcal{G}) < \lambda_2(\mathcal{G}) \leq \dots \leq \lambda_n(\mathcal{G})$, and $\mathcal{H}$ an ε -sparsifier of $\mathcal{G}$. Let $\mathbf{y} \in \mathbb{R}^n$ be the labels of the nodes in $\mathcal{G}$ with $|\mathbf{y}(x)| \leq c$ and $\mathcal{F}$ be the set of centered functions such that $|\mathbf{f}(x) - \mathbf{y}(x)| \leq 2c$. Let $\mathcal{S} \subset \mathcal{V}$ be a random subset of labeled nodes, if the labels $\mathbf{y}_{\mathcal{S}}$ are centered, then w.p. at least $1 - \delta$ (w.r.t. the random generation of the sparsifier $\mathcal{H}$ and the random subset of labeled points $\mathcal{S}$) the resulting *Stable-HFS* solution satisfies

$$\begin{aligned} R(\tilde{\mathbf{f}}) &\leq \hat{R}(\hat{\mathbf{f}}) + \beta + \left(2\beta + \frac{4c^2(\ell + u)}{\ell u}\right) \sqrt{\frac{\pi(\ell, u) \ln \frac{1}{\delta}}{2}} \\ &\quad + \frac{1}{1 - \varepsilon} \left(\frac{2(1 + \varepsilon)\varepsilon\ell\gamma\lambda_2(\mathcal{G})c}{((1 - \varepsilon)\ell\gamma\lambda_2(\mathcal{G}) - 1)^2} \right)^2, \end{aligned} \quad (6)$$

where $\tilde{\mathbf{f}}$ and $\hat{\mathbf{f}}$ are computed on $\mathcal{H}$ and $\mathcal{G}$,

$$\begin{aligned} \pi(\ell, u) &\triangleq \frac{\ell u}{\ell + u - 0.5} \frac{2 \max\{\ell, u\}}{2 \max\{\ell, u\} - 1} \quad \text{and} \\ \beta &\leq \frac{3c\sqrt{\ell}}{((1 - \varepsilon)\ell\gamma\lambda_2(\mathcal{G}) - 1)^2} + \frac{4c}{(1 - \varepsilon)\ell\gamma\lambda_2(\mathcal{G}) - 1}. \end{aligned}$$

Proof of Thm. 2. Step 1 (generalization of stable algorithms). Let β be the stability of *Stable-HFS* when using the sparsified Laplacian $L_{\mathcal{H}}$ in place of $L_{\mathcal{G}}$. Then using the result of Cortes et al. (2008), we have that with probability at least $1 - \delta$ (w.r.t. the randomness of the labeled set $\mathcal{S}$) the solution $\tilde{\mathbf{f}}$ satisfies

$$R(\tilde{\mathbf{f}}) \leq \hat{R}(\tilde{\mathbf{f}}) + \beta + \left(2\beta + \frac{c^2(\ell + u)}{\ell u}\right) \sqrt{\frac{\pi(\ell, u) \log(1/\delta)}{2}}.$$

In order to obtain the final result, we first derive an upper bound on the stability β and relate the empirical error of $\tilde{\mathbf{f}}_{\text{STA}}$ to the one of $\hat{\mathbf{f}}_{\text{STA}}$. Furthermore, it can be shown that if we center the vector of labels $\tilde{\mathbf{y}}_{\mathcal{S}} \triangleq \mathbf{y}_{\mathcal{S}} - \bar{\mathbf{y}}_{\mathcal{S}}$, with $\bar{\mathbf{y}} \triangleq \frac{1}{\ell} \mathbf{y}_{\mathcal{S}}^T \mathbf{1}$, then the solution of *Stable-HFS* can be rewritten in closed form as $\hat{\mathbf{f}}_{\text{STA}} = (\gamma \triangleq \mathbf{L}_{\mathcal{G}} + \mathbf{I}_{\mathcal{S}})^+(\tilde{\mathbf{y}}_{\mathcal{S}} - \mu \mathbf{1}) = (\mathbf{P}(\gamma \mathbf{L}_{\mathcal{G}} + \mathbf{I}_{\mathcal{S}}))^+ \tilde{\mathbf{y}}_{\mathcal{S}}$.

Step 2 (stability). The bound on the stability follows similar steps as in the analysis of *Stable-HFS* by Belkin et al. (2004) integrated with the properties of spectral sparsifiers reported in Def. 1. Let $\mathcal{S}$ and $\mathcal{S}'$ be two labeled sets only that only differ by one element and $\tilde{\mathbf{f}}$ and $\tilde{\mathbf{f}}'$ be the solutions obtained by running *Stable-HFS* using $L_{\mathcal{H}}$ and $\mathcal{S}$ and $\mathcal{S}'$ respectively.

Without loss of generality, we assume that $\mathbf{I}_{\mathcal{S}}(\ell, \ell) = 1$ and $\mathbf{I}_{\mathcal{S}}(\ell + 1, \ell + 1) = 0$, and the opposite for $\mathbf{I}_{\mathcal{S}'}$. The original proof of Cortes et al. (2008) showed that the stability β can be bounded as $\beta \leq \|\tilde{\mathbf{f}} - \tilde{\mathbf{f}}'\|$. In the following, we show that the difference between the solutions $\tilde{\mathbf{f}}$ and $\tilde{\mathbf{f}}'$ and thus the stability of the algorithm, is strictly related to eigenvalues of the sparse graph $\mathcal{H}$. Let $\mathbf{A} \triangleq \mathbf{P}(\ell\gamma\mathbf{L}_{\mathcal{H}} + \mathbf{I}_{\mathcal{S}})$ and $\mathbf{B} \triangleq \mathbf{P}(\ell\gamma\mathbf{L}_{\mathcal{H}} + \mathbf{I}_{\mathcal{S}'})$. We remind that if the labels are centered, the solutions of *Stable-HFS* can be conveniently written as $\tilde{\mathbf{f}} = \mathbf{A}^{-1}\tilde{\mathbf{y}}_{\mathcal{S}}$ and $\tilde{\mathbf{f}}' = \mathbf{B}^{-1}\tilde{\mathbf{y}}_{\mathcal{S}'}$. As a result, the difference between the solutions can be written as

$$\|\tilde{\mathbf{f}} - \tilde{\mathbf{f}}'\| = \|\mathbf{A}^{-1}\tilde{\mathbf{y}}_{\mathcal{S}} - \mathbf{B}^{-1}\tilde{\mathbf{y}}_{\mathcal{S}'}\| \leq \|\mathbf{A}^{-1}(\tilde{\mathbf{y}}_{\mathcal{S}} - \tilde{\mathbf{y}}_{\mathcal{S}'})\| + \|\mathbf{A}^{-1}\tilde{\mathbf{y}}_{\mathcal{S}'} - \mathbf{B}^{-1}\tilde{\mathbf{y}}_{\mathcal{S}'}\|. \quad (7)$$

Let us consider any vector $\mathbf{f} \in \mathcal{F}$, since the null space of a Laplacian $\mathbf{L}_{\mathcal{H}}$ is the one vector $\mathbf{1}$ and $\mathbf{P} = \mathbf{L}_{\mathcal{H}}\mathbf{L}_{\mathcal{H}}^+$, then $\mathbf{P}\mathbf{f} = \mathbf{f}$. Thus, we have

$$\|\mathbf{P}(\ell\gamma\mathbf{L}_{\mathcal{H}} + \mathbf{I}_{\mathcal{S}})\mathbf{f}\| \stackrel{(1)}{\geq} \|\mathbf{P}\ell\gamma\mathbf{L}_{\mathcal{H}}\mathbf{f}\| - \|\mathbf{P}\mathbf{I}_{\mathcal{S}}\mathbf{f}\| \stackrel{(2)}{\geq} \|\mathbf{P}\ell\gamma\mathbf{L}_{\mathcal{H}}\mathbf{f}\| - \|\mathbf{f}\| \stackrel{(3)}{\geq} (\ell\gamma\lambda_2(\mathcal{H}) - 1)\|\mathbf{f}\|, \quad (8)$$

where (1) follows from the triangle inequality and (2) follows from the fact that $\|\mathbf{P}\mathbf{I}_S\mathbf{f}\| \leq \|\mathbf{f}\|$, since the largest eigenvalue of the projection matrix $\mathbf{P}$ is one and the norm of $\mathbf{f}$ restricted on $\mathcal{S}$ is smaller than the norm of $\mathbf{f}$. Finally, (3) follows from the fact that $\|\mathbf{P}\mathbf{L}_\mathcal{H}\mathbf{f}\| = \|\mathbf{L}_\mathcal{H}\mathbf{L}_\mathcal{H}^+\mathbf{L}_\mathcal{H}\mathbf{f}\| = \|\mathbf{L}_\mathcal{H}\mathbf{f}\|$ and since $\mathbf{f}$ is orthogonal to the null space of $\mathbf{L}_\mathcal{H}$ then $\|\mathbf{L}_\mathcal{H}\mathbf{f}\| \geq \lambda_2(\mathcal{H})\|\mathbf{f}\|$, where $\lambda_2(\mathcal{H})$ is the smallest non-zero eigenvalue of $\mathbf{L}_\mathcal{H}$. At this point we can exploit the spectral guarantees of the sparsified Laplacian $L_\mathcal{H}$ and we have that $\lambda_2(\mathcal{H}) \geq (1 - \varepsilon)\lambda_2(\mathcal{G})$. As a result, we have an upper bound on the spectral radius of the inverse operator $(\mathbf{P}(\ell\gamma\mathbf{L}_\mathcal{H} + \mathbf{I}_S))^{-1}$ and thus

$$\|\mathbf{A}^{-1}(\mathbf{y}_S - \mathbf{y}_{S'})\| \leq \frac{4M}{\ell\gamma(1 - \varepsilon)\lambda_2(\mathcal{G}) - 1},$$

where the first step follows from Eq. 8 since both $\tilde{\mathbf{y}}_S$ and $\tilde{\mathbf{y}}_{S'}$ are centered and thus $(\mathbf{y}_S - \mathbf{y}_{S'}) \in \mathcal{F}$ and the second step is obtained by bounding $\|\tilde{\mathbf{y}}_S - \tilde{\mathbf{y}}_{S'}\| \leq \|\mathbf{y}_S - \mathbf{y}_{S'}\| + \|\bar{\mathbf{y}}_S - \bar{\mathbf{y}}_{S'}\| \leq 4M$. The second term in Eq. 7 can be bounded as

$$\|\mathbf{A}^{-1}\tilde{\mathbf{y}}_{S'} - \mathbf{B}^{-1}\tilde{\mathbf{y}}_{S'}\| = \|\mathbf{B}^{-1}(\mathbf{B} - \mathbf{A})\mathbf{A}^{-1}\tilde{\mathbf{y}}_{S'}\| = \|\mathbf{B}^{-1}\mathbf{P}(\mathbf{I}_S - \mathbf{I}_{S'})\mathbf{A}^{-1}\tilde{\mathbf{y}}_{S'}\| \leq \frac{1.5M\sqrt{\ell}}{(\ell\gamma(1 - \varepsilon)\lambda_2(\mathcal{G}) - 1)^2},$$

where we used $\|\tilde{\mathbf{y}}_{S'}\| \leq \|\mathbf{y}_{S'}\| + \|\bar{\mathbf{y}}_{S'}\| \leq 2M\sqrt{\ell}$, $\|\mathbf{P}(\mathbf{I}_S - \mathbf{I}_{S'})\| \leq \sqrt{2} < 1.5$ and we applied Eq. 8 twice. Putting it all together we obtain the stated bound.

Step 3 (empirical error). The other element effected by the sparsification is the empirical error $\hat{R}(\tilde{\mathbf{f}})$. We first recall that $\mathbf{P} = \mathbf{L}_\mathcal{G}^+\mathbf{L}_\mathcal{G} = \mathbf{L}_\mathcal{G}^{-1/2}\mathbf{L}_\mathcal{G}\mathbf{L}_\mathcal{G}^{-1/2}$ (and equivalently with $\mathcal{G}$ replaced by $\mathcal{H}$) and we introduce $\tilde{\mathbf{P}} = \mathbf{L}_\mathcal{G}^{-1/2}\mathbf{L}_\mathcal{H}\mathbf{L}_\mathcal{G}^{-1/2}$. Let $\tilde{\mathbf{A}} \triangleq \mathbf{P}(\ell\gamma\mathbf{L}_\mathcal{H} + \mathbf{I}_S)$, $\hat{\mathbf{A}} \triangleq \mathbf{P}(\ell\gamma\mathbf{L}_\mathcal{G} + \mathbf{I}_S)$, then we can rewrite the empirical error as

$$\begin{aligned} \hat{R}(\tilde{\mathbf{f}}) &= \frac{1}{\ell} \|\mathbf{I}_S\tilde{\mathbf{f}} - \mathbf{I}_S\hat{\mathbf{f}} + \mathbf{I}_S\hat{\mathbf{f}} - \tilde{\mathbf{y}}_S\|^2 \\ &\leq \frac{1}{\ell} \|\mathbf{I}_S\hat{\mathbf{f}} - \tilde{\mathbf{y}}_S\|^2 + \frac{1}{\ell} \|\mathbf{I}_S\tilde{\mathbf{f}} - \mathbf{I}_S\hat{\mathbf{f}}\|^2 \\ &\leq \hat{R}(\hat{\mathbf{f}}) + \frac{1}{\ell} \|\mathbf{I}_S(\tilde{\mathbf{A}}^{-1} - \hat{\mathbf{A}}^{-1})\tilde{\mathbf{y}}_S\|^2 \\ &\leq \hat{R}(\hat{\mathbf{f}}) + \frac{1}{\ell} \|\hat{\mathbf{A}}^{-1}(\hat{\mathbf{A}} - \tilde{\mathbf{A}})\tilde{\mathbf{A}}^{-1}\tilde{\mathbf{y}}_S\|^2 \\ &= \hat{R}(\hat{\mathbf{f}}) + \frac{\ell^2\gamma^2}{\ell} \|\hat{\mathbf{A}}^{-1}(\mathbf{P}(\mathbf{L}_\mathcal{G} - \mathbf{L}_\mathcal{H}))\tilde{\mathbf{A}}^{-1}\tilde{\mathbf{y}}_S\|^2 \\ &= \hat{R}(\hat{\mathbf{f}}) + \ell\gamma^2 \|\hat{\mathbf{A}}^{-1}(\mathbf{P}(\mathbf{L}_\mathcal{G} - \mathbf{L}_\mathcal{H})\mathbf{P})\tilde{\mathbf{A}}^{-1}\tilde{\mathbf{y}}_S\|^2, \end{aligned}$$

where in the last passagen we use $\mathbf{L}_\mathcal{G}\mathbf{P} = \mathbf{L}_\mathcal{G}$ and $\mathbf{L}_\mathcal{H}\mathbf{P} = \mathbf{L}_\mathcal{H}$. To bound the second term, we derive

$$\begin{aligned} \|\hat{\mathbf{A}}^{-1}\mathbf{P}(\mathbf{L}_\mathcal{G} - \mathbf{L}_\mathcal{H})\mathbf{P}\tilde{\mathbf{A}}^{-1}\tilde{\mathbf{y}}_S\|^2 &\stackrel{(1)}{=} \|\hat{\mathbf{A}}^{-1}\mathbf{L}_\mathcal{G}^{1/2}\mathbf{L}_\mathcal{G}^{-1/2}(\mathbf{L}_\mathcal{G} - \mathbf{L}_\mathcal{H})\mathbf{L}_\mathcal{G}^{-1/2}\mathbf{L}_\mathcal{G}^{1/2}\tilde{\mathbf{A}}^{-1}\tilde{\mathbf{y}}_S\|^2 \\ &\stackrel{(2)}{=} \|\hat{\mathbf{A}}^{-1}\mathbf{L}_\mathcal{G}^{1/2}\mathbf{P}\mathbf{L}_\mathcal{G}^{-1/2}(\mathbf{L}_\mathcal{G} - \mathbf{L}_\mathcal{H})\mathbf{L}_\mathcal{G}^{-1/2}\mathbf{P}\mathbf{L}_\mathcal{G}^{1/2}\tilde{\mathbf{A}}^{-1}\tilde{\mathbf{y}}_S\|^2 \\ &\stackrel{(3)}{=} \|\hat{\mathbf{A}}^{-1}\mathbf{L}_\mathcal{G}^{1/2}\mathbf{P}(\mathbf{P} - \tilde{\mathbf{P}})\mathbf{P}\mathbf{L}_\mathcal{G}^{1/2}\tilde{\mathbf{A}}^{-1}\tilde{\mathbf{y}}_S\|^2 \\ &\leq \|\hat{\mathbf{A}}^{-1}\mathbf{L}_\mathcal{G}^{1/2}\mathbf{P}\|^2 \|\mathbf{P} - \tilde{\mathbf{P}}\|^2 \|\mathbf{P}\mathbf{L}_\mathcal{G}^{1/2}\tilde{\mathbf{A}}^{-1}\|^2 \|\tilde{\mathbf{y}}_S\|^2 \\ &\stackrel{(4)}{\leq} \ell M^2 \varepsilon^2 \|\hat{\mathbf{A}}^{-1}\mathbf{L}_\mathcal{G}^{1/2}\mathbf{P}\|^2 \|\tilde{\mathbf{A}}^{-1}\mathbf{L}_\mathcal{G}^{1/2}\mathbf{P}\|^2, \end{aligned}$$

where in (1) and (2), we use the definition of $\mathbf{P}$, in (3) we use the definition of $\tilde{\mathbf{P}}$, while in (4) we use the fact that Def. 1 implies that $(1 - \varepsilon)\mathbf{P} \preceq \tilde{\mathbf{P}} \preceq (1 + \varepsilon)\mathbf{P}$ and thus the largest eigenvalue of $\mathbf{P} - \tilde{\mathbf{P}}$ is ε and $\|\mathbf{P} - \tilde{\mathbf{P}}\|^2 \leq \varepsilon^2$. We need now to bound $\|\tilde{\mathbf{A}}^{-1}\mathbf{L}_\mathcal{G}^{1/2}\mathbf{P}\|^2 = \|\tilde{\mathbf{A}}^{-1}\mathbf{L}_\mathcal{G}^{1/2}\|^2$. From the definition of spectral norm,

$$\begin{aligned} \|\tilde{\mathbf{A}}^{-1}\mathbf{L}_\mathcal{G}^{1/2}\|^2 &= \max_{\|\mathbf{x}\|=1} \mathbf{x}^\top \mathbf{L}_\mathcal{G}^{1/2} \tilde{\mathbf{A}}^{-1} \tilde{\mathbf{A}}^{-1} \mathbf{L}_\mathcal{G}^{1/2} \mathbf{x} = \max_{\|\mathbf{x}\|=1} \mathbf{x}^\top \tilde{\mathbf{A}}^{-1} \mathbf{L}_\mathcal{G} \tilde{\mathbf{A}}^{-1} \mathbf{x} \\ &\leq \frac{1}{1 - \varepsilon} \max_{\|\mathbf{x}\|=1} \mathbf{x}^\top \tilde{\mathbf{A}}^{-1} \mathbf{L}_\mathcal{H} \tilde{\mathbf{A}}^{-1} \mathbf{x} = \frac{1}{1 - \varepsilon} \|\tilde{\mathbf{A}}^{-1}\mathbf{L}_\mathcal{H}^{1/2}\|^2 = \frac{1}{1 - \varepsilon} \|\tilde{\mathbf{A}}^{-1}\mathbf{L}_\mathcal{H}^{1/2}\mathbf{P}\|^2. \end{aligned}$$

Similarly to Eq. 8, finding a lower bound on $\|\tilde{\mathbf{A}}\mathbf{L}_{\mathcal{H}}^{-1/2}\mathbf{P}\mathbf{x}\|$ for all $\mathbf{x}$ is equivalent to find a lower bound for all $\mathbf{f} \in \mathcal{F}$ to

$$\begin{aligned}
 \|\mathbf{P}(\ell\gamma\mathbf{L}_{\mathcal{H}} + \mathbf{I}_{\mathcal{S}})\mathbf{L}_{\mathcal{H}}^{-1/2}\mathbf{f}\| &\geq \|\mathbf{P}\ell\gamma\mathbf{L}_{\mathcal{H}}\mathbf{L}_{\mathcal{H}}^{-1/2}\mathbf{f}\| - \|\mathbf{P}\mathbf{I}_{\mathcal{S}}\mathbf{L}_{\mathcal{H}}^{-1/2}\mathbf{f}\| \\
 &\geq \|\mathbf{P}\ell\gamma\mathbf{L}_{\mathcal{H}}\mathbf{L}_{\mathcal{H}}^{-1/2}\mathbf{f}\| - \|\mathbf{L}_{\mathcal{H}}^{-1/2}\mathbf{f}\| \\
 &\geq \left(\ell\gamma\sqrt{\lambda_2(\mathcal{H})} - \frac{1}{\sqrt{\lambda_2(\mathcal{H})}} \right) \|\mathbf{f}\| \\
 &\geq \frac{1}{\sqrt{\lambda_2(\mathcal{H})}} (\ell\gamma\lambda_2(\mathcal{H}) - 1) \|\mathbf{f}\| \\
 &\geq \frac{1}{\sqrt{(1+\varepsilon)\lambda_2(\mathcal{G})}} (\ell\gamma(1-\varepsilon)\lambda_2(\mathcal{G}) - 1) \|\mathbf{f}\|.
 \end{aligned}$$

Similarly, we can show that

$$\begin{aligned}
 \|\mathbf{P}(\ell\gamma\mathbf{L}_{\mathcal{G}} + \mathbf{I}_{\mathcal{S}})\mathbf{L}_{\mathcal{G}}^{-1/2}\mathbf{f}\| &\geq \|\mathbf{P}\ell\gamma\mathbf{L}_{\mathcal{G}}\mathbf{L}_{\mathcal{G}}^{-1/2}\mathbf{f}\| - \|\mathbf{P}\mathbf{I}_{\mathcal{S}}\mathbf{L}_{\mathcal{G}}^{-1/2}\mathbf{f}\| \\
 &\geq \|\mathbf{P}\ell\gamma\mathbf{L}_{\mathcal{G}}\mathbf{L}_{\mathcal{G}}^{-1/2}\mathbf{f}\| - \|\mathbf{L}_{\mathcal{G}}^{-1/2}\mathbf{f}\| \\
 &\geq \left(\ell\gamma\frac{\lambda_2(\mathcal{G})}{\sqrt{\lambda_2(\mathcal{G})}} - \frac{1}{\sqrt{\lambda_2(\mathcal{G})}} \right) \|\mathbf{f}\| \\
 &\geq \frac{1}{\sqrt{\lambda_2(\mathcal{G})}} (\ell\gamma\lambda_2(\mathcal{G}) - 1) \|\mathbf{f}\| \\
 &\geq \frac{1}{\sqrt{(1+\varepsilon)\lambda_2(\mathcal{G})}} (\ell\gamma(1-\varepsilon)\lambda_2(\mathcal{G}) - 1) \|\mathbf{f}\|.
 \end{aligned}$$

Taking this and putting all together gives

$$\hat{R}(\tilde{\mathbf{f}}) \leq \hat{R}(\hat{\mathbf{f}}) + \frac{1}{1-\varepsilon} \frac{(1+\varepsilon)^2 \varepsilon^2 \ell^2 \gamma^2 \lambda_2(\mathcal{G})^2 M^2}{(\ell\gamma(1-\varepsilon)\lambda_2(\mathcal{G}) - 1)^4}.$$

Combining the three steps above concludes the proof. $\square$

Theorem 3. For an arbitrary graph $\mathcal{G}$ and its (ε, γ) -sparsifier, let $\hat{\mathbf{f}}$ be the LAPSMO solution computed using $\mathbf{L}_{\mathcal{G}}$ and $\tilde{\mathbf{f}}$ the solution computed using $\mathbf{L}_{\mathcal{H}}$. Then,

$$\|\tilde{\mathbf{f}} - \hat{\mathbf{f}}\|_2^2 \leq \frac{\varepsilon^2}{1-\varepsilon} (0.25 + \lambda\gamma) \left(\lambda \hat{\mathbf{f}}^T \mathbf{L}_{\mathcal{G}} \hat{\mathbf{f}} + \lambda\gamma \|\hat{\mathbf{f}}\|_2^2 \right),$$

where λ is the regularization of LAPSMO.

Proof. We need to bound the distance between $\tilde{\mathbf{f}}$ and $\hat{\mathbf{f}}$. Using the definition, the fact that $\mathbf{A}^{-1} - \mathbf{B}^{-1} = \mathbf{B}^{-1}(\mathbf{B} - \mathbf{A})\mathbf{A}^{-1}$ and collecting $(\mathbf{L}_{\mathcal{G}} + \gamma\mathbf{I})^{1/2}$ we have

$$\begin{aligned}
 \|\tilde{\mathbf{f}} - \hat{\mathbf{f}}\|_2^2 &= \|(\lambda\mathbf{L}_{\mathcal{H}} + \mathbf{I})^{-1} - \lambda\mathbf{L}_{\mathcal{G}} + \mathbf{I})^{-1}\mathbf{y}\|_2^2 = \|(\lambda\mathbf{L}_{\mathcal{H}} + \mathbf{I})^{-1}(\lambda\mathbf{L}_{\mathcal{H}} - \lambda\mathbf{L}_{\mathcal{G}})(\lambda\mathbf{L}_{\mathcal{G}} + \mathbf{I})^{-1}\mathbf{y}\|_2^2 \\
 &= \lambda^2 \|(\lambda\mathbf{L}_{\mathcal{H}} + \mathbf{I})^{-1}(\lambda\mathbf{L}_{\mathcal{H}} - \mathbf{L}_{\mathcal{G}})(\lambda\mathbf{L}_{\mathcal{G}} + \mathbf{I})^{-1}\mathbf{y}\|_2^2 \\
 &= \lambda^2 \|(\lambda\mathbf{L}_{\mathcal{H}} + \mathbf{I})^{-1}(\mathbf{L}_{\mathcal{G}} + \gamma\mathbf{I})^{1/2}(\mathbf{L}_{\mathcal{G}} + \gamma\mathbf{I})^{-1/2}(\mathbf{L}_{\mathcal{H}} - \mathbf{L}_{\mathcal{G}})(\mathbf{L}_{\mathcal{G}} + \gamma\mathbf{I})^{-1/2}(\mathbf{L}_{\mathcal{G}} + \gamma\mathbf{I})^{1/2}\hat{\mathbf{f}}\|_2^2.
 \end{aligned}$$

Then, using Prop. 3,

$$\begin{aligned}
 & \lambda^2 \|(\lambda \mathbf{L}_{\mathcal{H}} + \mathbf{I})^{-1} (\mathbf{L}_{\mathcal{G}} + \gamma \mathbf{I})^{1/2} (\mathbf{L}_{\mathcal{G}} + \gamma \mathbf{I})^{-1/2} (\mathbf{L}_{\mathcal{H}} - \mathbf{L}_{\mathcal{G}}) (\mathbf{L}_{\mathcal{G}} + \gamma \mathbf{I})^{-1/2} (\mathbf{L}_{\mathcal{G}} + \gamma \mathbf{I})^{1/2} \hat{\mathbf{f}}\|_2^2 \\
 & \leq \varepsilon^2 \lambda^2 \|(\lambda \mathbf{L}_{\mathcal{H}} + \mathbf{I})^{-1} (\mathbf{L}_{\mathcal{G}} + \gamma \mathbf{I})^{1/2}\|_2^2 \hat{\mathbf{f}}^\top (\mathbf{L}_{\mathcal{G}} + \gamma \mathbf{I}) \hat{\mathbf{f}} \\
 & = \frac{\varepsilon^2}{1 - \varepsilon} \lambda^2 \|(\lambda \mathbf{L}_{\mathcal{H}} + \mathbf{I})^{-1} ((1 - \varepsilon) \mathbf{L}_{\mathcal{G}} - \varepsilon \gamma \mathbf{I} + \gamma \mathbf{I}) (\lambda \mathbf{L}_{\mathcal{H}} + \mathbf{I})^{-1}\|_2 \hat{\mathbf{f}}^\top (\mathbf{L}_{\mathcal{G}} + \gamma \mathbf{I}) \hat{\mathbf{f}} \\
 & \leq \frac{\varepsilon^2}{1 - \varepsilon} \lambda^2 \|(\lambda \mathbf{L}_{\mathcal{H}} + \mathbf{I})^{-1} (\mathbf{L}_{\mathcal{H}} + \gamma \mathbf{I}) (\lambda \mathbf{L}_{\mathcal{H}} + \mathbf{I})^{-1}\|_2 \hat{\mathbf{f}}^\top (\mathbf{L}_{\mathcal{G}} + \gamma \mathbf{I}) \hat{\mathbf{f}} \\
 & = \frac{\varepsilon^2}{1 - \varepsilon} \lambda^2 \max_i \left\{ \frac{\lambda_i(\mathbf{L}_{\mathcal{H}}) + \gamma}{(\lambda \lambda_i(\mathbf{L}_{\mathcal{H}}) + 1)^2} \right\} \hat{\mathbf{f}}^\top (\mathbf{L}_{\mathcal{G}} + \gamma \mathbf{I}) \hat{\mathbf{f}} \\
 & = \frac{\varepsilon^2}{1 - \varepsilon} \lambda \max_i \left\{ \frac{\lambda \lambda_i(\mathbf{L}_{\mathcal{H}})}{(\lambda \lambda_i(\mathbf{L}_{\mathcal{H}}) + 1)^2} + \frac{\lambda \gamma}{(\lambda \lambda_i(\mathbf{L}_{\mathcal{H}}) + 1)^2} \right\} \hat{\mathbf{f}}^\top (\mathbf{L}_{\mathcal{G}} + \gamma \mathbf{I}) \hat{\mathbf{f}} \\
 & \leq \frac{\varepsilon^2}{1 - \varepsilon} \lambda \left(\max_i \left\{ \frac{\lambda \lambda_i(\mathbf{L}_{\mathcal{H}})}{(\lambda \lambda_i(\mathbf{L}_{\mathcal{H}}) + 1)^2} \right\} + \max_i \left\{ \frac{\lambda \gamma}{(\lambda \lambda_i(\mathbf{L}_{\mathcal{H}}) + 1)^2} \right\} \right) \hat{\mathbf{f}}^\top (\mathbf{L}_{\mathcal{G}} + \gamma \mathbf{I}) \hat{\mathbf{f}} \\
 & \leq \frac{\varepsilon^2}{1 - \varepsilon} \lambda (0.25 + \lambda \gamma) \hat{\mathbf{f}}^\top (\mathbf{L}_{\mathcal{G}} + \gamma \mathbf{I}) \hat{\mathbf{f}} \leq \frac{\varepsilon^2}{1 - \varepsilon} ((0.25 + \lambda \gamma) \lambda \hat{\mathbf{f}}^\top \mathbf{L}_{\mathcal{G}} \hat{\mathbf{f}} + (0.25 + \lambda \gamma) \lambda \gamma \mathbf{I}),
 \end{aligned}$$

which concludes the proof. $\square$

B. Proof of Thm. 1

Theorem 1. Let $\varepsilon > 0$ be the accuracy, $0 \leq \delta \leq 1$ the probability of error, and $\rho \triangleq (1 + 3\varepsilon)/(1 - \varepsilon)$. Given an arbitrary graph $\mathcal{G}$ and an arbitrary merge tree structure, if **DiSRe** is run with parameter $\bar{q} \triangleq 26\rho \log(3n/\delta)/\varepsilon^2$, then each sub-graphs $\mathcal{H}_{\{h,\ell\}}$ is an (ε, γ) -sparsifier of $\mathcal{G}_{\{h,\ell\}}$ with at most $3\bar{q}d_{\text{eff}}(\gamma)$ edges with probability $1 - \delta$. Whenever the merge tree is balanced and k is big enough such that $m/k \leq 3\bar{q}d_{\text{eff}}(\gamma)$,⁶ then merge operations can be run in parallel across the machines with an overall time complexity of $\mathcal{O}(d_{\text{eff}}(\gamma) \log^3(n))$, a total work $\mathcal{O}(m \log^3(n))$, and $\mathcal{O}(\log(n))$ rounds of communication.

The proof of Thm. 1 assembles the techniques from two sources. First, it is based on the analysis of the algorithm of (Kelner & Levin, 2013), that shows it produces a spectral sparsifier in high probability. We have previously published this analysis as a technical report (Calandriello et al., 2016) for the case of $\gamma = 0$. In this alternative proof, we rigorously take into account the dependencies across subsequent resparsifications using martingale inequalities, fixing a flaw in the original analysis. Second, it also copies the steps for the analysis of SQUEAK (Calandriello et al., 2017) which is an algorithm for general kernel sparsification. In particular, Alg. 1 is an instantiation of SQUEAK to the special case of graph sparsification. While SQUEAK’s analysis by Calandriello et al. (2017) holds in general, in **DiSRe** we can exploit the specific structure of graph Laplacians to perform a few optimizations. For completeness and ease of verification, we restate here SQUEAK’s original proof with the necessary modifications to the notation, constants, and relevant quantities to target the graph learning setting.

Merge trees We first formalize the random process induced by Alg. 1.

We partition $\mathcal{G}$ into k disjoint sub-graphs $\mathcal{G}_i$ of size n_i , such that $\mathcal{G} = \cup_{i=1}^k \mathcal{G}_i$. For each sub-graph $\mathcal{G}_i$, we construct an initial sparsifier $\mathcal{H}_{\{1,i\}} \triangleq \{(j, \tilde{p}_{0,i} = 1, q_{0,i} = \bar{q}) : j \in \mathcal{G}_i\}$ by inserting all edges from $\mathcal{G}_i$ into $\mathcal{H}_{1,i}$ with weight $\tilde{p}_{0,i} \triangleq 1$ and number of copies $q_{0,i} \triangleq \bar{q}$. It is easy to see that $\mathcal{H}_{\{1,i\}}$ is an $(0, 0)$ -accurate sparsifier, and we can split the graph into small enough sub-graphs to make sure that it can be easily stored and manipulated in memory. Afterwards, the initial sparsifiers $\mathcal{H}_{\{1,i\}}$ are included into the sparsifier pool $\mathcal{S}_1$.

At iteration h , the inner loop of Alg. 1 arbitrarily chooses two sparsifiers from $\mathcal{S}_h$ and merges them into a new sparsifier. Any arbitrary sequence of merges can be described by a full binary tree, i.e., a binary tree where each node is either a leaf or has exactly two children. Fig. 2 shows several different merge trees corresponding to different choices for the order of the merges. Note that starting from k leaves, a full binary tree will always have exactly $k - 1$ internal nodes. Therefore, regardless of the structure of the merge tree, we can always transform it into a tree of depth k , with all the initial sparsifiers $\mathcal{H}_{1,i}$ as leaves on its deepest layer. After this transformation, we index the tree nodes using their height (longest path from the node to a leaf, also defined as the depth of the tree minus the depth of the node), where leaves have height 1 and the root has height k . We can also see that at each layer, there is a single sparsifier merge, and the size of $\mathcal{S}_h$ (number of sparsifiers

⁶This implies that there are enough machines so that the leaves in the merge tree already have relatively sparse sub-graphs.

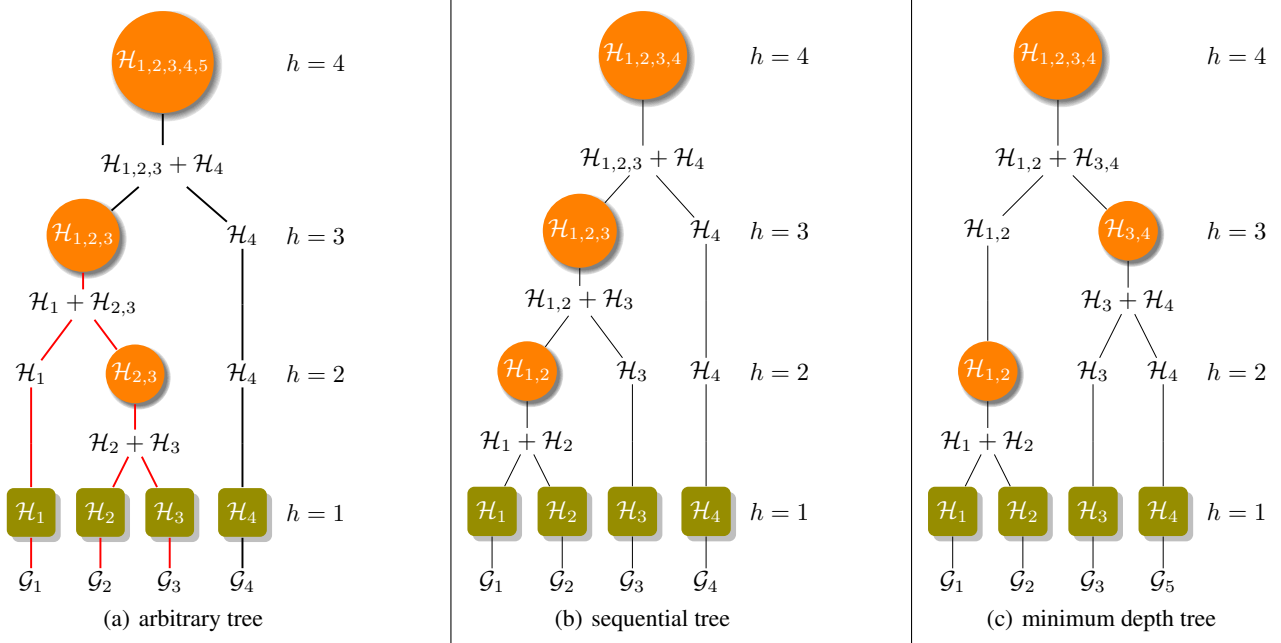


Figure 2. Merge trees for Alg. 1.

present at layer h) is $|\mathcal{S}_h| = k - h + 1$. Therefore, a node corresponding to a sparsifier is uniquely identified with two indices $\{h, \ell\}$, where h is the height of the layer and $\ell \leq |\mathcal{S}_h|$ is the index of the node in the layer. For example, in Fig. 2(a), the node containing $\mathcal{H}_{1,2,3}$ is indexed as $\{3, 1\}$, and the highest node containing $\mathcal{H}_4$ is indexed as $\{3, 2\}$.

We also define the graph $\mathcal{G}_{\{h, \ell\}}$ as the union of all sub-graph $\mathcal{G}_{\ell'}$ that are reachable from node $\{h, \ell\}$ as leaves. For example, in Fig. 2(a), sparsifier $\mathcal{H}_{1,2,3}$ in node $\{3, 1\}$ is constructed starting from all edges in $\mathcal{G}_{\{3, 1\}} = \mathcal{G}_1 \cup \mathcal{G}_2 \cup \mathcal{G}_3$, where we highlight in red the descendant tree. We now define $\mathbf{L}^h$ as the block diagonal matrix where each diagonal block $\mathbf{L}_{\mathcal{G}_{\{h, \ell\}}}$ is the Laplacian constructed on $\mathcal{G}_{\{h, \ell\}}$. Without loss of generality, we will assume that each of the sub-graphs $\mathcal{G}_i$ is connected and spans all the n nodes in the graph. This simplifies the notation for the $\mathbf{L}_{\mathcal{G}_{\{h, \ell\}}}$ matrices, making them all $n \times n$ matrices. If this is not the case, the whole proof still follows through with different number of nodes $n_{\{h, \ell\}}$ for each $\mathcal{G}_{\{h, \ell\}}$. Again, from Fig. 2, $\mathbf{L}^3$ is a $2n \times 2n$ matrix with two blocks on the diagonal, a first $n \times n$ block $\mathbf{L}_{\mathcal{G}_{\{3, 1\}}}$ constructed on $\mathcal{G}_{\{3, 1\}} = \mathcal{G}_1 \cup \mathcal{G}_2 \cup \mathcal{G}_3$, and a second $n \times n$ block $\mathbf{L}_{\mathcal{G}_{\{3, 2\}}}$ constructed on $\mathcal{G}_{\{3, 2\}} = \mathcal{G}_4$. Similarly, we combine Def. 1 and Prop. 3 to define

$$\begin{aligned} \mathbf{P}_{\{h, \ell\}} &\triangleq (\mathbf{L}_{\mathcal{G}_{\{h, \ell\}}} + \gamma \mathbf{I})^{-1/2} \mathbf{L}_{\mathcal{G}_{\{h, \ell\}}} (\mathbf{L}_{\mathcal{G}_{\{h, \ell\}}} + \gamma \mathbf{I})^{-1/2} \quad \text{and} \\ \tilde{\mathbf{P}}_{\{h, \ell\}} &\triangleq (\mathbf{L}_{\mathcal{G}_{\{h, \ell\}}} + \gamma \mathbf{I})^{-1/2} \mathbf{L}_{\mathcal{H}_{\{h, \ell\}}} (\mathbf{L}_{\mathcal{G}_{\{h, \ell\}}} + \gamma \mathbf{I})^{-1/2}, \end{aligned}$$

and have $\mathbf{P}^h$ as a block diagonal projection matrix, where each block $\mathbf{P}_{\{h, \ell\}}$ is defined using $\mathbf{L}_{\mathcal{G}_{\{h, \ell\}}}$, and block diagonal $\tilde{\mathbf{P}}^h$, where each block $\tilde{\mathbf{P}}_{\{h, \ell\}}$ is defined using $\mathbf{L}_{\mathcal{H}_{\{h, \ell\}}}$ and $\mathcal{H}_{\{h, \ell\}}$.

The statement. Since $\mathbf{P}^h - \tilde{\mathbf{P}}^h$ is block diagonal, we have that a bound on its largest eigenvalue implies an equal bound on each matrix on the diagonal, i.e.,

$$\|\mathbf{P}^h - \tilde{\mathbf{P}}^h\| = \max_{\ell} \|\mathbf{P}_{\{h, \ell\}} - \tilde{\mathbf{P}}_{\{h, \ell\}}\| \leq \varepsilon \implies \|\mathbf{P}_{\{h, \ell\}} - \tilde{\mathbf{P}}_{\{h, \ell\}}\| \leq \varepsilon$$

for all blocks ℓ on the diagonal, and since each block corresponds to a sparsifier $\mathcal{H}_{\{h, \ell\}}$, this means that if $\|\mathbf{P}^h - \tilde{\mathbf{P}}^h\| \leq \varepsilon$, all sparsifiers at layer ℓ are (ε, γ) -sparsifiers of their respective graphs. Let $d_{\text{eff}}^{\{h, \ell\}}(\gamma)$ be the effective dimension of $\mathbf{L}_{\mathcal{G}_{\{h, \ell\}}}$.

Our goal is to show

$$\begin{aligned} & \mathbb{P}\left(\exists h \in \{1, \dots, k\} : \|\mathbf{P}^h - \tilde{\mathbf{P}}^h\|_2 \geq \varepsilon \cup \max_{\ell=1, \dots, |\mathcal{S}_h|} |\mathcal{H}_{\{h, \ell\}}| \geq 3\bar{q}d_{\text{eff}}^{\{h, \ell\}}(\gamma)\right) \\ &= \mathbb{P}\left(\exists h \in \{1, \dots, k\} : \underbrace{\left(\max_{\ell=1, \dots, |\mathcal{S}_h|} \|\mathbf{P}_{\{h, \ell\}} - \tilde{\mathbf{P}}_{\{h, \ell\}}\|_2\right)}_{A_h} \geq \varepsilon \cup \underbrace{\left(\max_{\ell=1, \dots, |\mathcal{S}_h|} |\mathcal{H}_{\{h, \ell\}}| \geq 3\bar{q}d_{\text{eff}}^{\{h, \ell\}}(\gamma)\right)}_{B_h}\right) \leq \delta, \quad (9) \end{aligned}$$

where event A_h refers to the case when some sparsifier $\mathcal{H}_{\{h, \ell\}}$ at an intermediate layer h fails to accurately approximate $\mathbf{L}_{\{h, \ell\}}$ and event B_h considers the case when the memory requirement is not met (i.e., too many edges are kept in one of the sparsifiers $\mathcal{H}_{\{h, \ell\}}$ at a certain layer h). After reformulating and a union bound we obtain

$$\begin{aligned} & \mathbb{P}\left(\exists h \in \{1, \dots, k\} : \|\mathbf{P}^h - \tilde{\mathbf{P}}^h\|_2 \geq \varepsilon \cup \max_{\ell=1, \dots, |\mathcal{S}_h|} |\mathcal{H}_{\{h, \ell\}}| \geq 3\bar{q}d_{\text{eff}}^{\{h, \ell\}}(\gamma)\right) \\ & \leq \sum_{h=1}^k \sum_{\ell=1}^{|\mathcal{S}_h|} \mathbb{P}\left(\|\mathbf{P}_{\{h, \ell\}} - \tilde{\mathbf{P}}_{\{h, \ell\}}\|_2 \geq \varepsilon\right) \\ & \quad + \sum_{h=1}^k \sum_{\ell=1}^{|\mathcal{S}_h|} \mathbb{P}\left(|\mathcal{H}_{\{h, \ell\}}| \geq 3\bar{q}d_{\text{eff}}^{\{h, \ell\}}(\gamma) \cap \left\{\forall h' \in \{1, \dots, h\} : \|\mathbf{P}^{h'} - \tilde{\mathbf{P}}^{h'}\|_2 \leq \varepsilon\right\}\right) \leq \delta. \quad (10) \end{aligned}$$

The accuracy of the sparsifier (first term in the previous bound) is guaranteed by the fact that given an (ε, γ) -accurate sparsifier we obtain γ -effective resistance estimates (i.e. RLS estimates) which are at least a fraction of the true ones, thus forcing the algorithm to sample each column *enough*. On the other hand, the space complexity bound is achieved by exploiting the fact that estimates are always upper-bounded by the true γ -effective resistance, thus ensuring that Alg. 1 does not oversample columns w.r.t. the sampling process following the exact γ -effective resistance.

In the reminder of the proof, we will show that both events happen with probability smaller than $\delta/(2k^2)$. Since $|\mathcal{S}_h| = k - h + 1$, we have

$$\sum_{h=1}^k \sum_{\ell=1}^{|\mathcal{S}_h|} \frac{\delta}{2k^2} = \sum_{h=1}^k (k - h + 1) \frac{\delta}{2k^2} = k(k+1) \frac{\delta}{4k^2} \leq k^2 \frac{\delta}{2k^2} = \frac{\delta}{2},$$

and the union bound over all events is smaller than δ . The main advantage of splitting the failure probability as we did in Eq. 10 is that we can now analyze the processes that generated each $\mathbf{P}_{\{h, \ell\}} - \tilde{\mathbf{P}}_{\{h, \ell\}}$ (and each sparsifier $\mathcal{H}_{\{h, \ell\}}$) separately. Focusing on a single node $\{h, \ell\}$ restricts our problem on a well defined graph $\mathcal{G}_{\{h, \ell\}}$, where we can analyze the evolution of $\mathcal{H}_{\{h, \ell\}}$ sequentially.

B.1. Bounding the projection error $\|\mathbf{P}_{\{h, \ell\}} - \tilde{\mathbf{P}}_{\{h, \ell\}}\|$

The sequential process. Thanks to the union bound in Eq. 10, instead of having to consider the whole merge tree followed by Alg. 1, we can focus on each individual node $\{h, \ell\}$ and study the sequential process that generated its sparsifier $\mathcal{H}_{\{h, \ell\}}$. We will now map more clearly the actions taken by Alg. 1 to the process that generated $\mathbf{P}_{\{h, \ell\}} - \tilde{\mathbf{P}}_{\{h, \ell\}}$. We begin by focusing on $\tilde{\mathbf{P}}_{\{h, \ell\}}$, which is a random matrix defined starting from the fixed graph Laplacian $\mathbf{L}_{\mathcal{G}_{\{h, \ell\}}}$ and the random sparsifier Laplacian $\mathbf{L}_{\mathcal{H}_{\{h, \ell\}}}$, where the randomness influences both which edges are included in $\mathcal{H}_{\{h, \ell\}}$, and the weight with which they are added.

Note that since the merge tree is decided in advance, the graph $\mathcal{G}_{\{h, \ell\}}$ is not a random object and is fixed for the whole process. Consider now an edge $e \in \mathcal{G}_{\{h, \ell\}}$. Again for simplicity and without loss of generality⁷, we will assume that the starting graphs in the leaves are edge-disjoint. Therefore, there is a single path in the tree, with length h , from the leaves to $\{h, \ell\}$. This means that for all $s < h$, we can properly define a unique $\tilde{p}_{s, e}$ and $q_{s, e}$ associated with that point. More in detail, if at layer s point i is present in $\mathcal{G}_{\{s, \ell'\}}$, it means that either (1) Alg. 1 used $\mathcal{H}_{\{s, \ell'\}}$ to compute $\tilde{p}_{s, e}$, and $\tilde{p}_{s, e}$ to compute $q_{s, e}$, or (2) at layer h , Alg. 1 did not have any merge scheduled for point i , and we simply propagate $\tilde{p}_{s, e} = \tilde{p}_{s-1, i}$ and $q_{s, e} = q_{s-1, i}$. Consistently with the algorithm, we initialize $\tilde{p}_{0, i} = 1$ and $q_{0, i} = \bar{q}$.

⁷Alternatively, we can assign an index to each of the edges in the leaf graphs, requiring at most $km \leq kn^2$ indices.

Denote $m_{\{h,\ell\}} = |\mathcal{G}_{\{h,\ell\}}|$ so that we can use index $i \in [m_{\{h,\ell\}}]$ to index all edges in $\mathcal{G}_{\{h,\ell\}}$. Given the $n \times m_{\{h,\ell\}}$ matrix $\mathbf{Q} = (\mathbf{L}_{\mathcal{G}_{\{h,\ell\}}} + \gamma \mathbf{I})^{-1/2} \mathbf{B}_{\mathcal{G}_{\{h,\ell\}}}$ with its e -th column $\mathbf{q}_e = (\mathbf{L}_{\mathcal{G}_{\{h,\ell\}}} + \gamma \mathbf{I})^{-1/2} \mathbf{B}_{\mathcal{G}_{\{h,\ell\}}} \mathbf{e}_{m_{\{h,\ell\}},e}$, we can rewrite the projection matrix as that $\mathbf{P}_{\{h,\ell\}} = \mathbf{Q} \mathbf{Q}^\top = \sum_{e=1}^{m_{\{h,\ell\}}} \mathbf{q}_e \mathbf{q}_e^\top$. Note that

$$\|\mathbf{q}_e \mathbf{q}_e^\top\| = \mathbf{q}_e^\top \mathbf{q}_e = \mathbf{e}_{m_{\{h,\ell\}},e}^\top \mathbf{Q}^\top \mathbf{Q} \mathbf{e}_{m_{\{h,\ell\}},e} = \mathbf{e}_{m_{\{h,\ell\}},e}^\top \mathbf{Q} \mathbf{Q}^\top \mathbf{e}_{m_{\{h,\ell\}},e} = \mathbf{e}_{m_{\{h,\ell\}},e}^\top \mathbf{P}_{\{h,\ell\}} \mathbf{e}_{m_{\{h,\ell\}},e} = r_e(\gamma),$$

or, in other words, the norm $\|\mathbf{q}_e \mathbf{q}_e^\top\|$ is equal to the γ -effective resistance of the e -th edge w.r.t. to graph $\mathcal{G}_{\{h,\ell\}}$. Note that since e is present only in node l on layer h , its γ -effective resistance is uniquely defined w.r.t. $\mathcal{G}_{\{h,\ell\}}$ and can be shortened as $r_{h,e}$. Using $\mathbf{q}_e$, we can also introduce the random matrix $\tilde{\mathbf{P}}_s^{\{h,\ell\}}$ as

$$\tilde{\mathbf{P}}_s^{\{h,\ell\}} = \sum_{e=1}^{m_{\{h,\ell\}}} \frac{q_{s,e}}{\tilde{q} \tilde{p}_{s,e}} \mathbf{q}_e \mathbf{q}_e^\top = \sum_{e=1}^{m_{\{h,\ell\}}} \sum_{j=1}^{\tilde{q}} \frac{z_{s,e,j}}{\tilde{q} \tilde{p}_{s,e}} \mathbf{q}_e \mathbf{q}_e^\top.$$

where $z_{s,e,j}$ are $\{0, 1\}$ r.v. such that $q_{s,e} = \sum_{j=1}^{\tilde{q}} z_{s,e,j}$, or in other words $z_{s,e,j}$ are the Bernoulli random variables that compose the Binomial $q_{s,e}$ associated with edge e , with j indexing each individual copy of the edge. Note that when $s = h$, we have that $\tilde{\mathbf{P}}_h^{\{h,\ell\}} = \tilde{\mathbf{P}}_{\{h,\ell\}}$ and we recover the definition of the approximate projection matrix from Alg. 1. But, for a general $s \neq h$, $\tilde{\mathbf{P}}_s^{\{h,\ell\}}$ does not have a direct interpretation in the context of Alg. 1. It combines the vectors $\mathbf{q}_e$, which are defined using $\mathbf{L}_{\mathcal{G}_{\{h,\ell\}}}$ at layer h , with the weights $\tilde{p}_{s,e}$ computed by Alg. 1 across multiple nodes at layer s , which are potentially stored in different machines that cannot communicate. Nonetheless, $\tilde{\mathbf{P}}_s^{\{h,\ell\}}$ is a useful tool to analyze Alg. 1.

Taking into account that we are now considering a specific node $\{h, \ell\}$, we can drop the index from the graphs $\mathcal{G}_{\{h,\ell\}} = \mathcal{G}$, γ -effective resistances $\tau_{h,e}$, and size $m_{\{h,\ell\}} = m$. Using this shorter notation, we can reformulate our objective as bounding $\|\mathbf{P}_{\{h,\ell\}} - \tilde{\mathbf{P}}_{\{h,\ell\}}\|_2 = \|\mathbf{P}_{\{h,\ell\}} - \tilde{\mathbf{P}}_h^{\{h,\ell\}}\|_2$, and reformulate the process as a sequence of matrices $\{\mathbf{Y}_s\}_{s=1}^h$ defined as

$$\mathbf{Y}_s = \mathbf{P}_{\{h,\ell\}} - \tilde{\mathbf{P}}_s^{\{h,\ell\}} = \frac{1}{\tilde{q}} \sum_{e=1}^m \sum_{j=1}^{\tilde{q}} \left(1 - \frac{z_{s,e,j}}{\tilde{p}_{s,e}}\right) \mathbf{q}_e \mathbf{q}_e^\top,$$

where $\mathbf{Y}_h = \mathbf{P}_{\{h,\ell\}} - \tilde{\mathbf{P}}_h^{\{h,\ell\}} = \mathbf{P}_{\{h,\ell\}} - \tilde{\mathbf{P}}_{\{h,\ell\}}$, and $\mathbf{Y}_1 = \mathbf{P}_{\{h,\ell\}} - \tilde{\mathbf{P}}_0^{\{h,\ell\}} = \mathbf{0}$ since $\tilde{p}_{0,i} = 1$ and $q_{0,i} = \tilde{q}$.

B.2. Bounding $\mathbf{Y}_h$

We transformed the problem of bounding $\|\mathbf{P}_{\{h,\ell\}} - \tilde{\mathbf{P}}_{\{h,\ell\}}\|$ into the problem of bounding $\mathbf{Y}_h$, which we modeled as a random matrix process, connected to Alg. 1 by the fact that both algorithm and random process $\mathbf{Y}_h$ make use of the same weight $\tilde{p}_{s,e}$ and multiplicities $q_{s,e}$.

The frozen process. Inspired by Cohen et al. (2016), we will now replace the process $\mathbf{Y}_s$ with an alternative process $\bar{\mathbf{Y}}_s$ defined as

$$\bar{\mathbf{Y}}_s = \mathbf{Y}_{s-1} \mathbb{I} \{ \|\bar{\mathbf{Y}}_{s-1}\| \leq \varepsilon \} + \bar{\mathbf{Y}}_{s-1} \mathbb{I} \{ \|\bar{\mathbf{Y}}_{s-1}\| \geq \varepsilon \}.$$

This process starts from $\bar{\mathbf{Y}}_0 = \mathbf{Y}_0 = \mathbf{0}$, and is identical to $\mathbf{Y}_s$ until a step $\bar{s}$ where for the first time $\|\bar{\mathbf{Y}}_{\bar{s}}\| \leq \varepsilon$ and $\|\bar{\mathbf{Y}}_{\bar{s}+1}\| \geq \varepsilon$. After this failure happen the process $\bar{\mathbf{Y}}_s$ is “frozen” at $\bar{s}$ and $\bar{\mathbf{Y}}_s = \bar{\mathbf{Y}}_{\bar{s}+1}$ for all $\bar{s} + 1 \leq s \leq h$. Consequently, if any of the intermediate elements of the sequence violates the condition $\|\mathbf{Y}_s\| \leq \varepsilon$, the last element will violate it too. For the rest, $\bar{\mathbf{Y}}_s$ behaves exactly like $\mathbf{Y}_s$. Therefore,

$$\mathbb{P}(\|\mathbf{Y}_h\| \geq \varepsilon) \leq \mathbb{P}(\|\bar{\mathbf{Y}}_h\| \geq \varepsilon),$$

and if we can bound $\mathbb{P}(\|\bar{\mathbf{Y}}_h\| \geq \varepsilon)$ we will have a bound for the failure probability of Alg. 1, even though after “freezing” the process $\bar{\mathbf{Y}}_h$ does not make the same choices as the algorithm.

We will see now how to construct the process $\bar{\mathbf{Y}}_s$ starting from $z_{s,e,j}$ and $\tilde{p}_{s,e,j}$. We recursively define the indicator ($\{0, 1\}$) random variable $\bar{z}_{s,e,j}$ as

$$\bar{z}_{s,e,j} = \mathbb{I} \left\{ u_{s,e,j} \leq \frac{\bar{p}_{s,e,j}}{\bar{p}_{s-1,e,j}} \right\} \bar{z}_{s-1,e,j},$$

where $u_{s,e,j} \sim \mathcal{U}(0, 1)$ is a $[0, 1]$ uniform random variable and $\bar{p}_{s,e,j}$ is defined as

$$\bar{p}_{s,e,j} = \tilde{p}_{s,e} \mathbb{I} \{ \|\bar{\mathbf{Y}}_{s-1}\| \leq \varepsilon \cap z_{s-1,e,j} = 1 \} + \bar{p}_{s-1,e,j} \mathbb{I} \{ \|\bar{\mathbf{Y}}_{s-1}\| \geq \varepsilon \cup z_{s-1,e,j} = 0 \}.$$

This definition of the process satisfies the freezing condition, since if $\|\mathbf{Y}_{\bar{s}+1}\| \geq \varepsilon$ (we have a failure at step $\bar{s}$), for all $s' \geq \bar{s} + 1$ we have $\bar{z}_{s',i,j} = \bar{z}_{\bar{s}+1,i,j}$ with probability 1 ($\bar{p}_{s+1,i,j}/\bar{p}_{\bar{s},i,j} = \bar{p}_{\bar{s},i,j}/\bar{p}_{\bar{s},i,j} = 1$), and the weights $1/(\bar{q}\bar{p}_{s+1,i,j}) = 1/(\bar{q}\bar{p}_{\bar{s},i,j})$ never change.

Introducing a per-copy weight $\bar{p}_{s,e,j}$ and enforcing that $\bar{p}_{s+1,i,j} = \bar{p}_{s,e,j}$ when $z_{s,e,j} = 0$ avoids subtle inconsistencies in the formulation. In particular, not doing so would semantically correspond to reweighting dropped copies. Although this does not directly affect $\mathbf{Y}_s$ (since the ratio $z_{s,e,j}/\tilde{p}_{s,e}$ is zero for dropped copies), and therefore the relationship $\mathbb{P}(\|\mathbf{Y}_h\| \geq \varepsilon) \leq \mathbb{P}(\|\bar{\mathbf{Y}}_h\| \geq \varepsilon)$ still holds, we will see later how maintaining consistency helps us bound the second moment of our process.

We can now arrange the indices s , e , and j into a linear index $t = s$ in the range $[1, \dots, m^2\bar{q}]$, obtained as $t = \{s, e, j\} = (s-1)m\bar{q} + (e-1)\bar{q} + j$. We also define the difference matrix as

$$\bar{\mathbf{X}}_{\{s,e,j\}} = \frac{1}{\bar{q}} \left(\frac{z_{s-1,e,j}}{\bar{p}_{s-1,e,j}} - \frac{z_{s,e,j}}{\bar{p}_{s,e,j}} \right) \mathbf{q}_e \mathbf{q}_e^\top,$$

which allows writing the cumulative matrix as $\bar{\mathbf{Y}}_{\{s,e,j\}} = \sum_{r=1}^{\{s,e,j\}} \bar{\mathbf{X}}_{\{s,e,j\}}$ where the checkedges $\{s, m, \bar{q}\}$ correspond to $\bar{\mathbf{Y}}_s$,

$$\bar{\mathbf{Y}}_{\{s,m,\bar{q}\}} = \bar{\mathbf{Y}}_s = \frac{1}{\bar{q}} \sum_{e=1}^m \sum_{j=1}^{\bar{q}} \left(1 - \frac{z_{s,e,j}}{\bar{p}_{s,e,j}} \right) \mathbf{q}_e \mathbf{q}_e^\top.$$

Let $\mathcal{F}_s$ be the filtration containing all the realizations of the uniform random variables $u_{s,e,j}$ up to the step s , that is $\mathcal{F}_s = \{u_{s',e',j'}, \forall \{s', e', j'\} \leq s\}$. Again, we notice that $\mathcal{F}_s$ defines the state of the algorithm after completing iteration s because, unless a “freezing” happened, Alg. 1 and $\bar{\mathbf{Y}}_s$ flip coins with the same probability, and generate the same sparsifiers. Since $\tilde{r}_{s,e}$ and $\bar{p}_{s,e,j}$ are computed at the beginning of iteration s using the sparsifier $\mathcal{H}_{\{s,\ell'\}}$ (for some ℓ' unique at layer s), they are fully determined by $\mathcal{F}_{s-1}$. Furthermore, since $\mathcal{F}_{s-1}$ also defines the values of all indicator variables $\bar{z}_{s',e,j}$ up to $\bar{z}_{s-1,e,j}$ for any i and j , we have that all the Bernoulli variables $\bar{z}_{s,e,j}$ at iteration s are conditionally independent given $\mathcal{F}_{s-1}$. In other words, we have that for any e' , and j' such that $\{s, 1, 1\} \leq \{s, e', j'\} < s$ the following random variables are equal in distribution,

$$\bar{z}_{s,e,j} | \mathcal{F}_{\{s,e',j'\}} = \bar{z}_{s,e,j} | \mathcal{F}_{\{s-1,m,\bar{q}\}} \sim \mathcal{B}\left(\frac{\bar{p}_{s,e,j}}{\bar{p}_{s-1,e,j}}\right), \quad (11)$$

and for any e' , and j' such that $\{s, 1, 1\} \leq \{s, e', j'\} \leq \{s, m, \bar{q}\}$ and $s \neq \{s, e', j'\}$ we have the independence

$$\bar{z}_{s,e,j} | \mathcal{F}_{\{s-1,m,\bar{q}\}} \perp \bar{z}_{s,e',j'} | \mathcal{F}_{\{s-1,m,\bar{q}\}}. \quad (12)$$

While knowing that $\|\mathbf{Y}_s\| \leq \varepsilon$ is not sufficient to provide guarantees for the approximate probabilities $\tilde{p}_{s,e}$, we can show that it is enough to prove that the frozen probabilities $\bar{p}_{s,e,j}$ are never too small.

Lemma 1. *Let $\alpha \triangleq (1 + 3\varepsilon)/(1 - \varepsilon)$ and $\bar{p}_{s,e,j}$ be the sequence of probabilities generated by the freezing process. Then for any s, e , and j , we have $\bar{p}_{s,e,j} \geq p_{h,e}/\alpha = r_{h,e}/\alpha$.*

Proof. Let $\bar{s}$ be the step where the process freezes ($\bar{s} = h$ if it does not freeze), or, in other words, $\|\mathbf{Y}_{\bar{s}}\| < \varepsilon$ and $\|\mathbf{Y}_{\bar{s}+1}\| \geq \varepsilon$. From the definition of $\bar{p}_{s,e,j}$, we have that

$$\begin{aligned} \bar{p}_{s,e,j} &\geq \bar{p}_{\bar{s},e} = \tilde{p}_{\bar{s},e} = \max \{ \min \{ \tilde{r}_{\bar{s},e}, \tilde{p}_{\bar{s}-1,e} \}, \tilde{p}_{\bar{s}-1,e}/2 \} \\ &\geq \min \{ \tilde{r}_{\bar{s},e}, \tilde{p}_{\bar{s}-1,e} \} = \min \{ \tilde{r}_{\bar{s},e}, \tilde{p}_{\bar{s}-2,e} \} = \min \{ \tilde{r}_{\bar{s},e}, \tilde{p}_{\bar{s}-3,e} \} \dots = \min \{ \tilde{r}_{\bar{s},e}, \tilde{p}_{0,e} \} = \tilde{r}_{\bar{s},e}, \end{aligned}$$

and therefore $\bar{p}_{s,e,j} \geq \tilde{r}_{\bar{s},e}$. Now let $\{\bar{s}, \ell'\}$ be the node where $\tilde{r}_{\bar{s},e}$ was computed. From Alg. 1 we know it is computed using the sparsifier $\bar{\mathcal{H}}$ generated by the union of the two children of node $\{\bar{s}, \ell'\}$, which is an $(\varepsilon, 2\gamma)$ -sparsifier of $\mathcal{G}_{\{\bar{s}, \ell'\}}$.

From its definition, we know that $\tilde{r}_{\bar{s},e}$ is computed by Alg. 1 as

$$\tilde{r}_{s,e} = (1 - \varepsilon)\phi_e^\top (\mathbf{L}_{\overline{\mathcal{H}}} + (1 + \varepsilon)\gamma\mathbf{I})^{-1} \phi_i,$$

using only the edges in $\overline{\mathcal{H}}$ that are available at node $\{\bar{s}, \ell'\}$. Let once again $\mathcal{G} = \mathcal{G}_{\{h,\ell\}}$ by dropping the subscript, and define $\overline{\mathcal{H}}^c$ as the complement set of all edges in other sparsifiers at level $\bar{s}$ not in $\overline{\mathcal{H}}$. From the definition of $\tilde{\mathbf{P}}_{\bar{s}}^{\{h,\ell\}}$ and Prop. 3 we know that

$$\|\mathbf{Y}_{\bar{s}}\| = \left\| \mathbf{P}_{\{h,\ell\}} - \tilde{\mathbf{P}}_{\bar{s}}^{\{h,\ell\}} \right\|_2 = \left\| (\mathbf{L}_{\mathcal{G}} + 2\gamma\mathbf{I})^{-1/2} (\mathbf{L}_{\mathcal{G}} - (\mathbf{L}_{\overline{\mathcal{H}}} + \mathbf{L}_{\overline{\mathcal{H}}^c})) (\mathbf{L}_{\mathcal{G}} + 2\gamma\mathbf{I})^{-1/2} \right\|_2 \leq \varepsilon$$

and we know that this implies

$$\mathbf{L}_{\overline{\mathcal{H}}} \preceq \mathbf{L}_{\mathcal{G}} + \varepsilon(\mathbf{L}_{\mathcal{G}} + 2\gamma\mathbf{I}) - \mathbf{L}_{\overline{\mathcal{H}}^c} \preceq \mathbf{L}_{\mathcal{G}} + \varepsilon(\mathbf{L}_{\mathcal{G}} + 2\gamma\mathbf{I}).$$

Plugging it in the initial definition,

$$\begin{aligned} \tilde{r}_{s,e} &= (1 - \varepsilon)\mathbf{b}_e^\top (\mathbf{L}_{\overline{\mathcal{H}}} + (1 + \varepsilon)\gamma\mathbf{I})^{-1} \mathbf{b}_i \\ &\geq (1 - \varepsilon)\mathbf{b}_e^\top (\mathbf{L}_{\mathcal{G}} + \varepsilon(\mathbf{L}_{\mathcal{G}} + 2\gamma\mathbf{I}) + (1 + \varepsilon)\gamma\mathbf{I})^{-1} \mathbf{b}_e \\ &= (1 - \varepsilon) \frac{1}{1 + 3\varepsilon} \mathbf{b}_e^\top (\mathbf{B}_{\mathcal{G}}\mathbf{B}_{\mathcal{G}}^\top + \gamma\mathbf{I})^{-1} \mathbf{b}_e \geq \frac{1 - \varepsilon}{1 + 3\varepsilon} r_{h,e} \geq \frac{r_{h,e}}{\alpha}. \end{aligned}$$

□

We now proceed by studying the process $\{\overline{\mathbf{Y}}_s\}_{s=1}^h$ and showing that it is a bounded martingale. In order to show that $\overline{\mathbf{Y}}_s$ is a martingale, it is sufficient to verify the following (equivalent) conditions,

$$\mathbb{E} [\overline{\mathbf{Y}}_s \mid \mathcal{F}_{s-1}] = \overline{\mathbf{Y}}_{s-1} \Leftrightarrow \mathbb{E} [\overline{\mathbf{X}}_{\{s,e,j\}} \mid \mathcal{F}_{s-1}] = \mathbf{0}.$$

We begin by inspecting the conditional random variable $\overline{\mathbf{X}}_{\{s,e,j\}} \mid \mathcal{F}_{s-1}$. Given the definition of $\overline{\mathbf{X}}_{\{s,e,j\}}$, the conditioning on $\mathcal{F}_{s-1}$ determines the values of $\bar{z}_{s-1,e,j}$ and the approximate probabilities $\bar{p}_{s-1,e,j}$ and $\bar{p}_{s,e,j}$. In fact, remember that these quantities are fully determined by the realizations in $\mathcal{F}_{s-1}$ which are contained in $\mathcal{F}_{s-1}$. As a result, the only stochastic quantity in $\overline{\mathbf{X}}_{\{s,e,j\}}$ is the variable $\bar{z}_{s,e,j}$. Specifically, if $\|\overline{\mathbf{Y}}_{s-1}\| \geq \varepsilon$, then we have $\bar{p}_{s,e,j} = \bar{p}_{s-1,e,j}$ and $\bar{z}_{s,e,j} = \bar{z}_{s-1,e,j}$ (the process is stopped), and the martingale requirement $\mathbb{E} [\overline{\mathbf{X}}_{\{s,e,j\}} \mid \mathcal{F}_{s-1}] = \mathbf{0}$ is trivially satisfied. On the other hand, if $\|\overline{\mathbf{Y}}_{s-1}\| \leq \varepsilon$ we have

$$\begin{aligned} \mathbb{E}_{u_{s,e,j}} \left[\frac{1}{\bar{q}} \left(\frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s-1,e,j}} - \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}} \right) \mathbf{q}_e \mathbf{q}_e^\top \mid \mathcal{F}_{s-1} \right] &= \frac{1}{\bar{q}} \left(\frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s-1,e,j}} - \frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s,e,j}} \mathbb{E} \left[\mathbb{I} \left\{ u_{s,e,j} \leq \frac{\bar{p}_{s,e,j}}{\bar{p}_{s-1,e,j}} \right\} \mid \mathcal{F}_{s-1} \right] \right) \mathbf{q}_e \mathbf{q}_e^\top \\ &= \frac{1}{\bar{q}} \left(\frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s-1,e,j}} - \frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s,e,j}} \frac{\bar{p}_{s,e,j}}{\bar{p}_{s-1,e,j}} \right) \mathbf{q}_e \mathbf{q}_e^\top = \mathbf{0}, \end{aligned}$$

where we use the recursive definition of $\bar{z}_{s,e,j}$ and the fact that $u_{s,e,j}$ is a uniform random variable in $[0, 1]$. This proves that $\overline{\mathbf{Y}}_s$ is indeed a martingale. We now compute an upper bound R on the norm of the values of the difference process as

$$\|\overline{\mathbf{X}}_{\{s,e,j\}}\| = \frac{1}{\bar{q}} \left\| \left(\frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s-1,e,j}} - \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}} \right) \mathbf{q}_e \mathbf{q}_e^\top \right\| \leq \frac{1}{\bar{q}} \frac{1}{\bar{p}_{s,e,j}} \|\mathbf{q}_e \mathbf{q}_e^\top\| = \frac{1}{\bar{q}} \frac{1}{\bar{p}_{s,e,j}} \tau_{h,i} \leq \frac{1}{\bar{q}} \frac{\alpha}{\tau_{h,i}} \tau_{h,i} = \frac{\alpha}{\bar{q}} \triangleq R,$$

where we used Lem. 1 to bound $\bar{p}_{s,e,j} \leq r_{h,e}/\alpha$. If instead, $\|\overline{\mathbf{Y}}_{s-1}\| \geq \varepsilon$, the process is stopped and $\|\overline{\mathbf{X}}_s\| = \|\mathbf{0}\| = 0 \leq R$.

We are now ready to use a Freedman matrix inequality from Tropp (2011) to bound the norm of $\overline{\mathbf{Y}}$.

Proposition 4 (Tropp, 2011, Theorem 1.2). *Consider a matrix martingale $\{\mathbf{Y}_k : k = 0, 1, 2, \dots\}$ whose values are self-adjoint matrices with dimension d , and let $\{\mathbf{X}_k : k = 1, 2, 3, \dots\}$ be the difference sequence. Assume that the difference sequence is uniformly bounded in the sense that*

$$\|\mathbf{X}_k\|_2 \leq R \quad \text{almost surely for } k = 1, 2, 3, \dots$$

Define the predictable quadratic variation process of the martingale as

$$\mathbf{W}_k \triangleq \sum_{j=1}^k \mathbb{E} \left[\mathbf{X}_j^2 \mid \{\mathbf{X}_s\}_{s=0}^{j-1} \right], \quad \text{for } k = 1, 2, 3, \dots$$

Then, for all $\varepsilon \geq 0$ and $\sigma^2 > 0$,

$$\mathbb{P}(\exists k \geq 0 : \|\mathbf{Y}_k\|_2 \geq \varepsilon \cap \|\mathbf{W}_k\| \leq \sigma^2) \leq 2d \cdot \exp \left\{ -\frac{\varepsilon^2/2}{\sigma^2 + R\varepsilon/3} \right\}.$$

In order to use the previous inequality, we develop the probability of error for any fixed h as

$$\begin{aligned} \mathbb{P}(\|\mathbf{Y}_h\| \geq \varepsilon) &\leq \mathbb{P}(\|\bar{\mathbf{Y}}_h\| \geq \varepsilon) = \mathbb{P}(\|\bar{\mathbf{Y}}_h\| \geq \varepsilon \cap \|\mathbf{W}_h\| \leq \sigma^2) + \mathbb{P}(\|\bar{\mathbf{Y}}_h\| \geq \varepsilon \cap \|\mathbf{W}_h\| \geq \sigma^2) \\ &\leq \underbrace{\mathbb{P}(\|\bar{\mathbf{Y}}_h\| \geq \varepsilon \cap \|\mathbf{W}_h\| \leq \sigma^2)}_{(a)} + \underbrace{\mathbb{P}(\|\mathbf{W}_h\| \geq \sigma^2)}_{(b)}. \end{aligned}$$

Using the bound on $\|\bar{\mathbf{X}}_{\{s,e,j\}}\|_2$, we can directly apply Prop. 4 to bound (a) for any fixed σ^2 . To bound the part (b), we use the following lemma, proved later in Sec. B.3.

Lemma 2 (Low probability of the large norm of the predictable quadratic variation process). *We have that*

$$\mathbb{P} \left(\|\mathbf{W}_h\| \geq \frac{6\alpha}{\bar{q}} \right) \leq n \cdot \exp \left\{ -\frac{2\bar{q}}{\alpha} \right\}.$$

Since $\mathbf{P}_{\{h,\ell\}}$ is defined at most on n nodes, combining Prop. 4 with $\sigma^2 = 6\alpha/\bar{q}$, Lem 2, the fact that $2\varepsilon/3 \leq 1$ and the value used by Alg. 1, $\bar{q} = 39\alpha \log(2n/\delta)/\varepsilon^2$ we obtain

$$\begin{aligned} \mathbb{P} \left(\|\mathbf{P}_{\{h,\ell\}} - \tilde{\mathbf{P}}_{\{h,\ell\}}\|_2 \geq \varepsilon \right) &= \mathbb{P}(\|\mathbf{Y}_h\| \geq \varepsilon) \leq \mathbb{P}(\|\bar{\mathbf{Y}}_h\| \geq \varepsilon \cap \|\mathbf{W}_h\| \leq \sigma^2) + \mathbb{P}(\|\mathbf{W}_h\| \geq \sigma^2) \\ &\leq 2n \cdot \exp \left\{ -\frac{\varepsilon^2 \bar{q}}{\alpha} \left(\frac{1}{12 + 2\varepsilon/3} \right) \right\} + n \cdot \exp \left\{ -\frac{2\bar{q}}{\alpha} \right\} \\ &\leq 3n \cdot \exp \left\{ -\frac{\varepsilon^2}{13\alpha} \bar{q} \right\} = 3n \cdot \exp \left\{ -3 \log \left(\frac{2n}{\delta} \right) \right\} \\ &= 3n \cdot \exp \left\{ -\log \left(\left(\frac{2n}{\delta} \right)^3 \right) \right\} = \frac{3n\delta^3}{8n^3} \leq \frac{\delta}{2n^2}. \end{aligned}$$

This, combined with the fact that $k \leq n^2/n \leq n$ since at most we can split our graph into n parts each containing n edges, concludes this part of the proof.

B.3. Proof of Lem. 2 (bound on predictable quadratic variation)

Step 1 (a preliminary bound). We start by writing out $\mathbf{W}_t$ for the process $\bar{\mathbf{Y}}_s$,

$$\mathbf{W}_t = \frac{1}{\bar{q}^2} \sum_{\{s,e,j\} \leq t} \mathbb{E} \left[\left(\frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s-1,e,j}} - \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}} \right)^2 \mid \mathcal{F}_{\{s,e,j\}-1} \right] \mathbf{q}_e \mathbf{q}_e^\top \mathbf{q}_e \mathbf{q}_e^\top.$$

We rewrite the expectation terms in the equation above as

$$\begin{aligned}
 & \mathbb{E} \left[\left(\frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s-1,e,j}} - \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}} \right)^2 \middle| \mathcal{F}_{\{s,e,j\}-1} \right] \\
 &= \mathbb{E} \left[\frac{\bar{z}_{s-1,e,j}^2}{\bar{p}_{s-1,e,j}^2} - 2 \frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s-1,e,j}} \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}} + \frac{\bar{z}_{s,e,j}^2}{\bar{p}_{s,e,j}^2} \middle| \mathcal{F}_{\{s,e,j\}-1} \right] \\
 &\stackrel{(a)}{=} \mathbb{E} \left[\frac{\bar{z}_{s-1,e,j}^2}{\bar{p}_{s-1,e,j}^2} - 2 \frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s-1,e,j}} \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}} + \frac{\bar{z}_{s,e,j}^2}{\bar{p}_{s,e,j}^2} \middle| \mathcal{F}_{s-1} \right] \\
 &= \frac{\bar{z}_{s-1,e,j}^2}{\bar{p}_{s-1,e,j}^2} - 2 \frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s-1,e,j}} \frac{1}{\bar{p}_{s,e,j}} \mathbb{E}[\bar{z}_{s,e,j} \mid \mathcal{F}_{s-1}] + \frac{1}{\bar{p}_{s,e,j}^2} \mathbb{E}[\bar{z}_{s,e,j}^2 \mid \mathcal{F}_{s-1}] \\
 &\stackrel{(b)}{=} \frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s-1,e,j}} - 2 \frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s-1,e,j}} \frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s-1,e,j}} + \frac{1}{\bar{p}_{s,e,j}^2} \mathbb{E}[\bar{z}_{s,e,j} \mid \mathcal{F}_{s-1}] \\
 &= \frac{1}{\bar{p}_{s,e,j}^2} \mathbb{E}[\bar{z}_{s,e,j} \mid \mathcal{F}_{s-1}] - \frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s-1,e,j}} \\
 &\stackrel{(c)}{=} \frac{1}{\bar{p}_{s,e,j}} \frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s-1,e,j}} - \frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s-1,e,j}^2} = \frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s-1,e,j}} \left(\frac{1}{\bar{p}_{s,e,j}} - \frac{1}{\bar{p}_{s-1,e,j}} \right),
 \end{aligned}$$

where in (a) we use the fact that the approximate probabilities $\bar{p}_{s-1,e,j}$ and $\bar{p}_{s,e,j}$ and $\bar{z}_{s-1,e,j}$ are fixed at the end of the previous iteration, while in (b) and (c) we use the fact that $\bar{z}_{s,e,j}$ is a Bernoulli of parameter $\bar{p}_{s,e,j}/\bar{p}_{s-1,e,j}$ (whenever $\bar{z}_{s-1,e,j}$ is equal to 1). Therefore, we can write $\mathbf{W}_t$ at the end of the process as

$$\mathbf{W}_h = \mathbf{W}_{\{h,m,\bar{q}\}} = \frac{1}{\bar{q}^2} \sum_{j=1}^{\bar{q}} \sum_{e=1}^m \sum_{s=1}^h \frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s-1,e,j}} \left(\frac{1}{\bar{p}_{s,e,j}} - \frac{1}{\bar{p}_{s-1,e,j}} \right) \mathbf{q}_e \mathbf{q}_e^\top \mathbf{q}_e \mathbf{q}_e^\top.$$

We can now upper-bound $\mathbf{W}_h$ as

$$\begin{aligned}
 \mathbf{W}_h &\preceq \frac{1}{\bar{q}^2} \sum_{j=1}^{\bar{q}} \sum_{e=1}^m \sum_{s=1}^h \frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s-1,e,j}} \left(\frac{1}{\bar{p}_{s,e,j}} - \frac{1}{\bar{p}_{s-1,e,j}} \right) \mathbf{q}_e \mathbf{q}_e^\top \mathbf{q}_e \mathbf{q}_e^\top \\
 &= \frac{1}{\bar{q}^2} \sum_{j=1}^{\bar{q}} \sum_{e=1}^m \left(\frac{\bar{z}_{h,e,j}}{\bar{p}_{h,e,j}^2} - \frac{\bar{z}_{h,e,j}}{\bar{p}_{h,e,j}^2} + \sum_{s=1}^h \frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s-1,e,j}} \left(\frac{1}{\bar{p}_{s,e,j}} - \frac{1}{\bar{p}_{s-1,e,j}} \right) \right) \mathbf{q}_e \mathbf{q}_e^\top \mathbf{q}_e \mathbf{q}_e^\top \\
 &= \frac{1}{\bar{q}^2} \sum_{j=1}^{\bar{q}} \sum_{e=1}^m \left(\frac{\bar{z}_{h,e,j}}{\bar{p}_{h,e,j}^2} + \left(\sum_{s=1}^h -\frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}^2} + \frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s,e,j} \bar{p}_{s-1,e,j}} \right) - \frac{\bar{z}_{0,e,j}}{\bar{p}_{0,e,j}^2} \right) \mathbf{q}_e \mathbf{q}_e^\top \mathbf{q}_e \mathbf{q}_e^\top \\
 &\preceq \frac{1}{\bar{q}^2} \sum_{j=1}^{\bar{q}} \sum_{e=1}^m \left(\frac{\bar{z}_{h,e,j}}{\bar{p}_{h,e,j}^2} + \left(\sum_{s=1}^h \frac{\bar{z}_{s-1,e,j}}{\bar{p}_{s,e,j} \bar{p}_{s-1,e,j}} - \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j} \bar{p}_{s-1,e,j}} \right) \right) \mathbf{q}_e \mathbf{q}_e^\top \mathbf{q}_e \mathbf{q}_e^\top \\
 &= \frac{1}{\bar{q}^2} \sum_{j=1}^{\bar{q}} \sum_{e=1}^m \left(\frac{\bar{z}_{h,e,j}}{\bar{p}_{h,e,j}^2} + \sum_{s=1}^h \frac{\bar{z}_{s-1,e,j}(1 - \bar{z}_{s,e,j})}{\bar{p}_{s,e,j} \bar{p}_{s-1,e,j}} \right) \mathbf{q}_e \mathbf{q}_e^\top \mathbf{q}_e \mathbf{q}_e^\top,
 \end{aligned}$$

where in the inequality we use the fact $\bar{p}_{s,e,j} \leq \bar{p}_{s-1,e,j}$. From the definition of $\bar{p}_{s,e,j}$, we know that when $\bar{z}_{s,e,j} = 0$, $\bar{p}_{s,e,j} = \bar{p}_{s-1,e,j}$. Therefore $\frac{\bar{z}_{s-1,e,j}(1 - \bar{z}_{s,e,j})}{\bar{p}_{s,e,j} \bar{p}_{s-1,e,j}} = \frac{\bar{z}_{s-1,e,j}(1 - \bar{z}_{s,e,j})}{\bar{p}_{s-1,e,j}^2}$, since the term is non-zero only when $\bar{z}_{s,e,j} = 0$. Finally,

we see that only one of the $\bar{z}_{s-1,e,j}(1 - \bar{z}_{s,e,j})$ terms can be active for $s \in [h]$ and thus

$$\begin{aligned} \mathbf{W}_h &\preceq \frac{1}{\bar{q}^2} \sum_{j=1}^{\bar{q}} \sum_{e=1}^m \left(\frac{\bar{z}_{h,e,j}}{\bar{p}_{h,e,j}^2} + \sum_{s=1}^h \frac{\bar{z}_{s-1,e,j}(1 - \bar{z}_{s,e,j})}{\bar{p}_{s-1,e,j}^2} \right) \mathbf{q}_e \mathbf{q}_e^\top \mathbf{q}_e \mathbf{q}_e^\top \\ &= \frac{1}{\bar{q}^2} \sum_{j=1}^{\bar{q}} \sum_{e=1}^m \left(\max \left\{ \max_{s=1,\dots,h} \left\{ \frac{\bar{z}_{s-1,e,j}(1 - \bar{z}_{s,e,j})}{\bar{p}_{s-1,e,j}^2} \right\}; \frac{\bar{z}_{h,e,j}}{\bar{p}_{h,e,j}^2} \right\} \right) \mathbf{q}_e \mathbf{q}_e^\top \mathbf{q}_e \mathbf{q}_e^\top \\ &= \frac{1}{\bar{q}^2} \sum_{j=1}^{\bar{q}} \sum_{e=1}^m \mathbf{q}_e \mathbf{q}_e^\top \mathbf{q}_e \mathbf{q}_e^\top \left(\max_{s=0,\dots,h} \left\{ \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}^2} \right\} \right). \end{aligned} \quad (13)$$

Step 2 (introduction of a stochastically dominant process). We want to study $\max_{s=0,\dots,h} \left\{ \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}^2} \right\}$. To simplify notation, we will consider $\max_{s=0,\dots,h} \left\{ \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}} \right\}$, where we removed the square, which will be re-added in the end. We know trivially that this quantity is larger or equal than, 1 because $\bar{z}_{0,e,j}/\bar{p}_{0,e,j} = 1$, but upper-bounding this quantity is not trivial as the evolution of the various $\bar{p}_{s,e,j}$ depends in a complex way on the interaction between the random variables $\bar{z}_{s,e,j}$. Nonetheless, whenever $\bar{p}_{s,e,j}$ is significantly smaller than $\bar{p}_{s-1,e,j}$, the probability of keeping a copy of edge e at iteration s (i.e., $\bar{z}_{s,e,j} = 1$) is also very small. As a result, we expect the ratio $\frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}}$ to be still small with high probability.

Unfortunately, due to the dependence between different copies of the edge at different iterations, it seems difficult to exploit this intuition directly to provide an overall high-probability bound on $\mathbf{W}_h$. For this reason, we simplify the analysis by replacing each of the (potentially dependent) chains $\{\bar{z}_{s,e,j}/\bar{p}_{s,e,j}\}_{s=0}^h$ with a set of (independent) random variables $w_{0,e,j}$ that will stochastically dominate them.

We define the random variable $w_{s,e,j}$ using the following conditional distribution,⁸

$$\mathbb{P} \left(\frac{1}{w_{s,e,j}} \leq a \mid \mathcal{F}_s \right) = \begin{cases} 0 & \text{for } a < 1/\bar{p}_{s,e,j} \\ 1 - \frac{1}{\bar{p}_{s,e,j}a} & \text{for } 1/\bar{p}_{s,e,j} \leq a < \alpha/p_{h,e} \\ 1 & \text{for } \alpha/p_{h,e} \leq a \end{cases}.$$

To show that this distribution is well defined, we use Lem. 1 to guarantee that $1/\bar{p}_{s,e,j} \leq a < \alpha/p_{h,e}$. Note that the distribution of $\frac{1}{w_{s,e,j}}$ conditioned on $\mathcal{F}_s$ is determined by only $\bar{p}_{s,e,j}$, $p_{h,e}$, and α , where $p_{h,e}$ and α are fixed. Remembering that $\bar{p}_{s,e,j}$ is a function of $\mathcal{F}_{s-1}$ (computed using the previous iteration), we have that

$$\mathbb{P} \left(\frac{1}{w_{s,e,j}} \leq a \mid \mathcal{F}_s \right) = \mathbb{P} \left(\frac{1}{w_{s,e,j}} \leq a \mid \mathcal{F}_{s-1} \right).$$

Notice that in the definition of $w_{s,e,j}$, none of the other $w_{s',e',j'}$ (for any different s' , e' , or j') appears and $\bar{p}_{s,e,j}$ is a function of $\mathcal{F}_{s-1}$. It follows that given $\mathcal{F}_{s-1}$, $w_{s,e,j}$ is independent from all other $w_{s',e',j'}$ (for any different s' , e' , or j'). This is easier to see in the probabilistic graphical model reported in Fig. 3, which illustrates the dependence between the various variables.

Finally for the special case $w_{0,e,j}$ the definition above reduces to

$$\mathbb{P} \left(\frac{1}{w_{0,e,j}} \leq a \right) = \begin{cases} 0 & \text{for } a < 1 \\ 1 - \frac{1}{a} & \text{for } 1 \leq a < \alpha/p_{h,e} \\ 1 & \text{for } \alpha/p_{h,e} \leq a \end{cases}, \quad (14)$$

since $\bar{p}_{0,e,j} = 1$ by definition. From this definition, $w_{0,e,j}$ and $w_{0,e',j'}$ are all independent, and this will allow us to use stronger concentration inequalities for independent random variables.

⁸ Notice that unlike $\bar{z}_{s,e,j}$, $w_{s,e,j}$ is no longer $\mathcal{F}_s$ -measurable but it is $\mathcal{F}'_s$ -measurable, where

$$\mathcal{F}'_{\{s,e,j\}} = \{u_{s',e',j'}, \forall \{s',e',j'\} \leq \{s,e,j\}\} \cup \{w_{s,e,j}\} = \mathcal{F}_{\{s,e,j\}} \cup \{w_{s,e,j}\}.$$

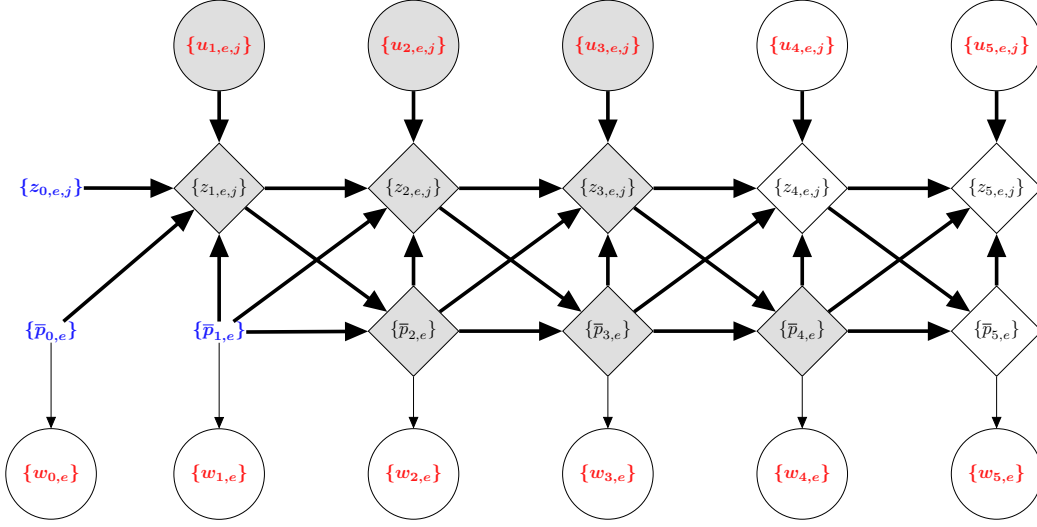


Figure 3. The dependence graph of the considered variables. **Red** variables are **random**. Black variables are deterministically computed using their input (a function of their input), with bold lines indicating the deterministic (functional) relation. **Blue** variables are **constants**. A grey filling indicates that a random variable is **observed** or a function of observed variables.

Step 3 Proving the dominance. We remind the reader that a random variable A stochastically dominates random variable B , if for all values a the two equivalent conditions are verified,

$$\mathbb{P}(A \geq a) \geq \mathbb{P}(B \geq a) \Leftrightarrow \mathbb{P}(A \leq a) \leq \mathbb{P}(B \leq a).$$

As a consequence, if A dominates B , the following implication holds,

$$\mathbb{P}(A \geq a) \geq \mathbb{P}(B \geq a) \implies \mathbb{E}[A] \geq \mathbb{E}[B],$$

while the reverse (A dominates B , if $\mathbb{E}[A] \geq \mathbb{E}[B]$) is not true in general. Following this definition of stochastic dominance, our goal is to prove

$$\mathbb{P}\left(\max_{s=0}^h \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}} \leq a\right) \geq \mathbb{P}\left(\frac{1}{w_{0,e,j}} \leq a\right).$$

We prove this inequality by proceeding backwards with a sequence of conditional probabilities. We first study the distribution of the maximum conditional to the state of the algorithm at the end of iteration h , i.e., $\mathcal{F}_h$. From the definition of $w_{h,e,j}$, we know that, w.p. 1, $1/\bar{p}_{h,e} \leq 1/w_{h,e,j}$. Therefore,

$$\mathbb{P}\left(\max_{s=0,\dots,h} \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}} \leq a\right) \geq \mathbb{P}\left(\max\left\{\max_{s=0,\dots,h-1} \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}}; \frac{\bar{z}_{h,e,j}}{w_{h,e,j}}\right\} \leq a\right).$$

Now focus on an arbitrary intermediate step $1 \leq k \leq h$, where we fix $\mathcal{F}_{k-1}$. Since $u_{k,e,j}$ and $w_{k,e,j}$ are independent given

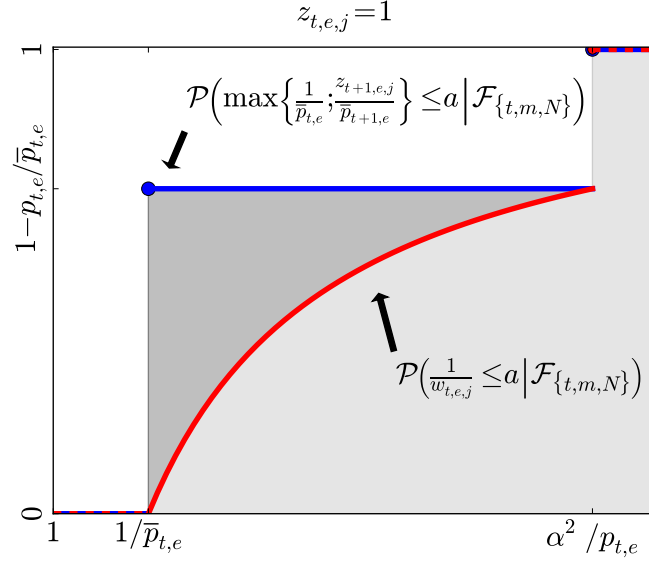


Figure 4. C.d.f. of $\max\{\bar{z}_{k-1,e,j}/\bar{p}_{t,e,j}; \bar{z}_{k,e,j}/\bar{p}_{k,e,j}\}$ and $\bar{z}_{k-1,e,j}/w_{k-1,e,j}$ conditioned on $\mathcal{F}_{\{k-1\}}$. For conciseness, we omit the e, j indices.

$\mathcal{F}_{k-1}$, we have

$$\begin{aligned}
 \mathbb{P}\left(\frac{\bar{z}_{k,e,j}}{w_{k,e,j}} \leq a \mid \mathcal{F}_{k-1}\right) &= \mathbb{P}\left(\mathbb{I}\left\{u_{k,e,j} \leq \frac{\bar{p}_{k,e,j}}{\bar{p}_{k-1,e,j}}\right\} \frac{1}{w_{k,e,j}} \leq a \mid \mathcal{F}_{k-1}\right) \\
 &= \begin{cases} 0 & \text{for } a \leq 0 \\ 1 - \frac{\bar{p}_{k,e,j}}{\bar{p}_{k-1,e,j}} & \text{for } 0 \leq a < 1/\bar{p}_{k,e,j} \\ 1 - \frac{\bar{p}_{k,e,j}}{\bar{p}_{k-1,e,j}} + \frac{\bar{p}_{k,e,j}}{\bar{p}_{k-1,e,j}} \left(1 - \frac{1}{\bar{p}_{k,e,j}a}\right) = 1 - \frac{1}{\bar{p}_{k-1,e,j}a} & \text{for } 1/\bar{p}_{k,e,j} \leq a < \alpha/p_{h,e} \\ 1 & \text{for } \alpha/p_{h,e} \leq a \end{cases} \\
 &\geq \begin{cases} 0 & \text{for } a < 1/\bar{p}_{k-1,e,j} \\ 1 - \frac{1}{\bar{p}_{k-1,e,j}a} & \text{for } 1/\bar{p}_{k-1,e,j} \leq a < 1/\bar{p}_{k,e,j} \\ 1 - \frac{1}{\bar{p}_{k-1,e,j}a} & \text{for } 1/\bar{p}_{k,e,j} \leq a < \alpha/p_{h,e} \\ 1 & \text{for } \alpha/p_{h,e} \leq a \end{cases} \quad (15) \\
 &= \mathbb{P}\left(\frac{1}{w_{k-1,e,j}} \leq a \mid \mathcal{F}_{k-2}\right) = \mathbb{P}\left(\frac{1}{w_{k-1,e,j}} \leq a \mid \mathcal{F}_{k-1}\right),
 \end{aligned}$$

where the inequality is also represented in Fig. 4. We now proceed by peeling off layers from the end of the chain one by one, taking advantage of the dominance we just proved. Fig. 4 visualizes one step of the peeling when $\bar{z}_{k-1,e,j} = 1$ (note that the peeling is trivially true when $\bar{z}_{k-1,e,j} = 0$ since the whole chain terminated at step $\bar{z}_{k-1,e,j}$). We show how to move

from an iteration $k \leq h$ to $k - 1$.

$$\begin{aligned}
 \mathbb{P} \left(\max \left\{ \max_{s=0 \dots k-1} \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}}; \frac{\bar{z}_{k,e,j}}{w_{k,e,j}} \right\} \leq a \right) &= \mathbb{E}_{\mathcal{F}_{k-1}} \left[\mathbb{P} \left(\max \left\{ \max_{s=0 \dots k-1} \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}}; \frac{\bar{z}_{k,e,j}}{w_{k,e,j}} \right\} \leq a \mid \mathcal{F}_{k-1} \right) \right] \\
 &\stackrel{(a)}{\geq} \mathbb{E}_{\mathcal{F}_{k-1}} \left[\mathbb{P} \left(\max \left\{ \max_{s=0 \dots k-1} \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}}; \frac{\bar{z}_{k-1,e,j}}{w_{k-1,e,j}} \right\} \leq a \mid \mathcal{F}_{k-1} \right) \right] \\
 &= \mathbb{E}_{\mathcal{F}_{k-1}} \left[\mathbb{P} \left(\max \left\{ \max_{s=0 \dots k-2} \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}}; \frac{\bar{z}_{k-1,e,j}}{\bar{p}_{k-1,e,j}}; \frac{\bar{z}_{k-1,e,j}}{w_{k-1,e,j}} \right\} \leq a \mid \mathcal{F}_{k-1} \right) \right] \\
 &= \mathbb{E}_{\mathcal{F}_{k-1}} \left[\mathbb{P} \left(\max \left\{ \max_{s=0 \dots k-2} \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}}; \bar{z}_{k-1,e,j} \max \left\{ \frac{1}{\bar{p}_{k-1,e,j}}; \frac{1}{w_{k-1,e,j}} \right\} \right\} \leq a \mid \mathcal{F}_{k-1} \right) \right] \\
 &\stackrel{(b)}{=} \mathbb{E}_{\mathcal{F}_{k-1}} \left[\mathbb{P} \left(\max \left\{ \max_{s=0 \dots k-2} \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}}; \frac{\bar{z}_{k-1,e,j}}{w_{k-1,e,j}} \right\} \leq a \mid \mathcal{F}_{k-1} \right) \right] = \mathbb{P} \left(\max \left\{ \max_{s=0 \dots k-2} \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}}; \frac{\bar{z}_{k-1,e,j}}{w_{k-1,e,j}} \right\} \leq a \right),
 \end{aligned}$$

where in (a), given $\mathcal{F}_{k-1}$, everything is fixed except $u_{k,e,j}$ and $w_{k,e,j}$ and we can use the stochastic dominance in (15), and in (b) we use the fact that the inner maximum is always attained by $1/w_{k,e,j}$ since by definition $1/w_{k-1,e,j}$ is lower-bounded by $1/\bar{p}_{k-1,e,j}$. Applying the inequality recursively from $k = h$ to $k = 1$ removes all $\bar{z}_{s,e,j}$ from the maximum and we are finally left with only $w_{0,e,j}$ as we wanted,

$$\mathbb{P} \left(\max_{s=0, \dots, h} \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}} \leq a \right) \geq \mathbb{P} \left(\max \left\{ \frac{\bar{z}_{0,e,j}}{\bar{p}_{0,e,j}}; \frac{\bar{z}_{0,e,j}}{w_{0,e,j}} \right\} \leq a \right) \geq \mathbb{P} \left(\frac{1}{w_{0,e,j}} \leq a \right),$$

where in the last inequality we used that $\bar{z}_{0,e,j} = 1$ from the definition of the algorithm and $\bar{p}_{0,e,j} = 1$ while $w_{0,e,j} \leq 1$ by (14).

Step 4 (stochastic dominance on $\mathbf{W}_h$). Now that we proved the stochastic dominance of $1/w_{0,e,j}$, we plug this result in the definition of $\mathbf{W}_h$. For the sake of notation, we introduce the term $\bar{p}_{h',e,j}^{\max}$ to indicate the maximum over the first h' step of copy e, j such that

$$\max_{s=0, \dots, h'} \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}} = \frac{1}{\bar{p}_{h',e,j}^{\max}}.$$

We first notice that while $\bar{\mathbf{Y}}_h$ is not necessarily PSD, $\mathbf{W}_h$ is a sum of PSD matrices. Introducing the function $\Lambda(\{1/\bar{p}_{h,e,j}^{\max}\}_{e,j})$ we can restate Eq. 13 as

$$\|\mathbf{W}_h\| = \lambda_{\max}(\mathbf{W}_h) \leq \Lambda(\{1/\bar{p}_{h,e,j}^{\max}\}_{e,j}) \triangleq \lambda_{\max} \left(\frac{1}{\bar{q}^2} \sum_{j=1}^{\bar{q}} \sum_{e=1}^m \left(\frac{1}{\bar{p}_{h,e,j}^{\max}} \right)^2 \mathbf{q}_e \mathbf{q}_e^T \mathbf{q}_e \mathbf{q}_e^T \right).$$

In Step 4, we showed that $1/\bar{p}_{h,e,j}^{\max}$ is stochastically dominated by $1/w_{0,e,j}$ for every e and j . In order to bound $\Lambda(\{1/\bar{p}_{h,e,j}^{\max}\}_{e,j})$, we need to show that this dominance also applies to the summation over all columns inside the matrix norm. We can reformulate $\Lambda(\{1/\bar{p}_{h,e,j}^{\max}\}_{e,j})$ as

$$\begin{aligned}
 \lambda_{\max} \left(\frac{1}{\bar{q}^2} \sum_{j=1}^{\bar{q}} \sum_{e=1}^m \left(\frac{1}{\bar{p}_{h,e,j}^{\max}} \right)^2 \mathbf{q}_e \mathbf{q}_e^T \mathbf{q}_e \mathbf{q}_e^T \right) &= \max_{\mathbf{x}: \|\mathbf{x}\|=1} \mathbf{x}^T \left(\frac{1}{\bar{q}^2} \sum_{j=1}^{\bar{q}} \sum_{e=1}^m \left(\frac{1}{\bar{p}_{h,e,j}^{\max}} \right)^2 \mathbf{q}_e \mathbf{q}_e^T \mathbf{q}_e \mathbf{q}_e^T \right) \mathbf{x} \\
 &= \max_{\mathbf{x}: \|\mathbf{x}\|=1} \frac{1}{\bar{q}^2} \sum_{j=1}^{\bar{q}} \sum_{e=1}^m \left(\frac{1}{\bar{p}_{h,e,j}^{\max}} \right)^2 \|\mathbf{q}_e\|_2^2 \mathbf{x}^T \mathbf{q}_e \mathbf{q}_e^T \mathbf{x} = \max_{\mathbf{x}: \|\mathbf{x}\|=1} \frac{1}{\bar{q}^2} \sum_{j=1}^{\bar{q}} \sum_{e=1}^m \left(\frac{1}{\bar{p}_{h,e,j}^{\max}} \right)^2 (\|\mathbf{q}_e\|_2 \mathbf{q}_e^T \mathbf{x})^2.
 \end{aligned}$$

From this reformulation, it is easy to see that, because $1/\bar{p}_{h,e,j}^{\max}$ is strictly positive, the function $\Lambda(\{1/\bar{p}_{h,e,j}^{\max}\}_{e,j})$ is monotonically increasing w.r.t. the individual $1/\bar{p}_{h,e,j}^{\max}$, or in other words that increasing an $1/\bar{p}_{h,e,j}^{\max}$ without decreasing the others can only increase the maximum. Introducing $\Lambda(\{1/w_{0,e,j}\}_{e,j})$ as

$$\Lambda(\{1/w_{0,e,j}\}_{e,j}) \triangleq \max_{\mathbf{x}: \|\mathbf{x}\|=1} \frac{1}{\bar{q}^2} \sum_{j=1}^{\bar{q}} \sum_{e=1}^m \left(\frac{1}{w_{0,e,j}} \right)^2 (\|\mathbf{q}_e\|_2 \mathbf{q}_e^T \mathbf{x})^2,$$

we now need to prove the stochastic dominance of $\Lambda(\{1/w_{0,e,j}\}_{e,j})$ over $\Lambda(\{1/\bar{p}_{h,e,j}^{\max}\}_{e,j})$. Using the definition of $1/\bar{p}_{h,e,j}^{\max}$, $w_{h,e,j}$, and the monotonicity of Λ we have

$$\begin{aligned} \mathbb{P}\left(\Lambda\left(\left\{\frac{1}{\bar{p}_{h,e,j}^{\max}}\right\}_{e,j}\right) \leq a\right) &= \mathbb{P}\left(\Lambda\left(\left\{\max\left\{\max_{s=0,\dots,h-1} \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}}; \frac{\bar{z}_{h,e,j}}{\bar{p}_{h,e,j}}\right\}\right\}_{e,j}\right) \leq a\right) \\ &\geq \mathbb{P}\left(\Lambda\left(\left\{\max\left\{\max_{s=0,\dots,h-1} \frac{\bar{z}_{s,e,j}}{\bar{p}_{s,e,j}}; \frac{\bar{z}_{h,e,j}}{w_{h,e,j}}\right\}\right\}_{e,j}\right) \leq a\right). \end{aligned}$$

Now pick $1 \leq k \leq h$, for a fixed $\mathcal{F}_{k-1}$, $\frac{1}{\bar{p}_{k-1,e,j}^{\max}}$ is a constant and $\max\left\{\frac{1}{\bar{p}_{k-1,e,j}^{\max}}; x\right\}$ is a monotonically increasing function in x , making $\Lambda\left(\max\left\{\frac{1}{\bar{p}_{k-1,e,j}^{\max}}; x\right\}\right)$ also an increasing function. Therefore, we have

$$\begin{aligned} \mathbb{P}\left(\Lambda\left(\left\{\max\left\{\frac{1}{\bar{p}_{k-1,e,j}^{\max}}; \frac{\bar{z}_{k,e,j}}{w_{k,e,j}}\right\}\right\}_{e,j}\right) \leq a\right) &= \mathbb{E}_{\mathcal{F}_{k-1}}\left[\mathbb{P}\left(\Lambda\left(\left\{\max\left\{\frac{1}{\bar{p}_{k-1,e,j}^{\max}}; \frac{\bar{z}_{k,e,j}}{w_{k,e,j}}\right\}\right\}_{e,j}\right) \leq a \mid \mathcal{F}_{k-1}\right)\right] \\ &\stackrel{(a)}{\geq} \mathbb{E}_{\mathcal{F}_{k-1}}\left[\mathbb{P}\left(\Lambda\left(\left\{\max\left\{\frac{1}{\bar{p}_{k-1,e,j}^{\max}}; \frac{\bar{z}_{k-1,e,j}}{w_{k-1,e,j}}\right\}\right\}_{e,j}\right) \leq a \mid \mathcal{F}_{k-1}\right)\right] \\ &\stackrel{(b)}{=} \mathbb{E}_{\mathcal{F}_{k-1}}\left[\mathbb{P}\left(\Lambda\left(\left\{\max\left\{\frac{1}{\bar{p}_{k-2,e,j}^{\max}}; \frac{\bar{z}_{k-1,e,j}}{w_{k-1,e,j}}\right\}\right\}_{e,j}\right) \leq a \mid \mathcal{F}_{k-1}\right)\right], \end{aligned}$$

where inequality (a) follows from the fact that stochastic dominance is preserved by monotonically increasing functions (Levy, 2015), such as Λ , combined with the fact that for a fixed $\mathcal{F}_{k-1}$ the variables $\bar{z}_{k,e,j}$ and $w_{k,e,j}$ are all independent and (b) from the definition of $1/\bar{p}_{k-1,e,j}^{\max}$ and the fact that by definition $1/w_{k-1,e,j}$ is lower-bounded by $1/\bar{p}_{k-1,e,j}$. We can iterate this inequality to obtain the desired result

$$\mathbb{P}(\|\mathbf{W}_h\| \geq \sigma^2) \leq \mathbb{P}\left(\Lambda\left(\left\{\frac{1}{\bar{p}_{h,e,j}^{\max}}\right\}_{e,j}\right) \geq \sigma^2\right) \leq \mathbb{P}\left(\lambda_{\max}\left(\frac{1}{\bar{q}^2} \sum_{j=1}^{\bar{q}} \sum_{e=1}^m \left(\frac{1}{w_{0,e,j}}\right)^2 \mathbf{q}_e \mathbf{q}_e^T \mathbf{q}_e \mathbf{q}_e^T\right) \geq \sigma^2\right).$$

Step 5 (concentration inequality). Since all $w_{0,e,j}$ are (unconditionally) independent from each other, we can apply the following theorem.

Proposition 5 (Tropp, 2015, Theorem 5.1.1). Consider a finite sequence $\{\mathbf{X}_k : k = 1, 2, 3, \dots\}$ whose values are independent, random, PSD Hermitian matrices with dimension d . Assume that each term in the sequence is uniformly bounded in the sense that

$$\lambda_{\max}(\mathbf{X}_k) \leq L \quad \text{almost surely for } k = 1, 2, 3, \dots$$

Introduce the random matrix $\mathbf{V} \triangleq \sum_k \mathbf{X}_k$, and the maximum eigenvalue of its expectation

$$\mu_{\max} \triangleq \lambda_{\max}(\mathbb{E}[\mathbf{V}]) = \lambda_{\max}\left(\sum_k \mathbb{E}[\mathbf{X}_k]\right).$$

Then, for all $h \geq 0$,

$$\begin{aligned} \mathbb{P}(\lambda_{\max}(\mathbf{V}) \geq (1+h)\mu_{\max}) &\leq d \cdot \left[\frac{e^h}{(1+h)^{1+h}}\right]^{\frac{\mu_{\max}}{L}} \\ &\leq d \cdot \exp\left\{-\frac{\mu_{\max}}{L}((h+1)\log(h+1) - h)\right\}. \end{aligned}$$

In our case, we have

$$\mathbf{X}_{\{e,j\}} = \frac{1}{\bar{q}^2} \frac{1}{w_{0,e,j}} \mathbf{q}_e \mathbf{q}_e^T \mathbf{q}_e \mathbf{q}_e^T \preceq \frac{1}{\bar{q}^2} \frac{\alpha^2}{p_{h,e}^2} \mathbf{q}_e \mathbf{q}_e^T \mathbf{q}_e \mathbf{q}_e^T \preceq \frac{1}{\bar{q}^2} \frac{\alpha^2}{p_{h,e}^2} \|\mathbf{q}_e \mathbf{q}_e^T\|^2 \mathbf{I} \preceq \frac{\alpha^2}{\bar{q}^2} \mathbf{I},$$

where the first inequality follows from the definition of $w_{0,e,j}$ in Eq. 14, the second from the PSD ordering, and the third from the definition of $\|\mathbf{q}_e \mathbf{q}_e^\top\|$.

Therefore, we can use $L \triangleq \alpha^2/\bar{q}^2$ for the purpose of Prop. 5. We need now to compute $\mathbb{E}[\mathbf{X}_k]$, that we can use in turn to compute $\mu_{\max}$. We begin by computing the expected value of $1/w_{0,e,j}$. Let us denote the c.d.f. of $1/w_{0,e,j}^2$ as

$$F_{1/w_{0,e,j}^2}(a) = \mathbb{P}\left(\frac{1}{w_{0,e,j}^2} \leq a\right) = \mathbb{P}\left(\frac{1}{w_{0,e,j}} \leq \sqrt{a}\right) = \begin{cases} 0 & \text{for } a < 1 \\ 1 - \frac{1}{\sqrt{a}} & \text{for } 1 \leq a < \alpha^2/p_{h,e}^2 \\ 1 & \text{for } \alpha^2/p_{h,e}^2 \leq a \end{cases}.$$

Since $\mathbb{P}(1/w_{0,e,j}^2 \geq 0) = 1$, we have that

$$\begin{aligned} \mathbb{E}\left[\frac{1}{w_{0,e,j}}\right] &= \int_{a=0}^{\infty} [1 - F_{1/w_{0,e,j}}(a)] da \\ &= \int_{a=0}^1 (1 - F_{1/w_{0,e,j}}(a)) da + \int_{a=1}^{\alpha^2/p_{h,e}^2} (1 - F_{1/w_{0,e,j}}(a)) da + \int_{a=\alpha^2/p_{h,e}^2}^{\infty} (1 - F_{1/w_{0,e,j}}(a)) da \\ &= \int_{a=0}^1 (1 - 0) da + \int_{a=1}^{\alpha^2/p_{h,e}^2} \left(1 - \left(1 - \frac{1}{\sqrt{a}}\right)\right) da + \int_{a=\alpha^2/p_{h,e}^2}^{\infty} (1 - 1) da \\ &= \int_{a=0}^1 da + \int_{a=1}^{\alpha^2/p_{h,e}^2} \frac{1}{\sqrt{a}} da = 1 + [2\sqrt{a}]_1^{\alpha^2/p_{h,e}^2} = 2\alpha/p_{h,e} - 1. \end{aligned}$$

Therefore,

$$\begin{aligned} \mu_{\max} &= \lambda_{\max}(\mathbb{E}[\mathbf{V}]) = \lambda_{\max}\left(\sum_{\{e,j\}} \mathbb{E}[\mathbf{X}_{\{e,j\}}]\right) = \lambda_{\max}\left(\frac{1}{\bar{q}^2} \sum_{j=1}^{\bar{q}} \sum_{e=1}^m \mathbb{E}\left[\frac{1}{w_{0,e,j}^2}\right] \mathbf{q}_e \mathbf{q}_e^\top \mathbf{q}_e \mathbf{q}_e^\top\right) \\ &= \lambda_{\max}\left(\frac{1}{\bar{q}} \sum_{e=1}^m \left(\frac{2\alpha}{p_{h,e}} - 1\right) p_{h,e} \mathbf{q}_e \mathbf{q}_e^\top\right) \leq \lambda_{\max}\left(\frac{2\alpha}{\bar{q}} \sum_{e=1}^m \mathbf{q}_e \mathbf{q}_e^\top\right) = \frac{2\alpha}{\bar{q}} \lambda_{\max}(\mathbf{P}) \leq \frac{2\alpha}{\bar{q}} \triangleq L. \end{aligned}$$

Therefore, selecting $h = 2$, $\sigma^2 = 6\alpha/\bar{q}$ and applying Prop. 5 we have

$$\begin{aligned} \mathbb{P}(\|\mathbf{W}_h\| \geq \sigma^2) &\leq \mathbb{P}\left(\lambda_{\max}\left(\frac{1}{\bar{q}^2} \sum_{j=1}^{\bar{q}} \sum_{e=1}^m \frac{1}{w_{0,e,j}^2} \mathbf{q}_e \mathbf{q}_e^\top \mathbf{q}_e \mathbf{q}_e^\top\right) \geq (1+2) \frac{2\alpha}{\bar{q}}\right) \\ &\leq m \cdot \exp\left\{-\frac{2\alpha}{\bar{q}} \frac{\bar{q}^2}{\alpha^2} (3\log(3) - 2)\right\} \\ &\leq n \cdot \exp\left\{-\frac{2\bar{q}}{\alpha}\right\}. \end{aligned}$$

B.4. Space complexity bound

Denote with A the event $A = \{\forall h' \in \{1, \dots, h\} : \|\mathbf{P}^{h'} - \tilde{\mathbf{P}}^{h'}\|_2 \leq \varepsilon\}$, and again $m = |\mathcal{G}_{\{h,\ell\}}|$. To begin with, we can show that under event A , the approximate γ -effective resistances (i.e. RLS) are accurate (Lem. 2, [Calandriello et al. 2017](#)). Let $\mathcal{H} = \mathcal{H}_{\{h,\ell\}}$ be a $(\varepsilon, 2\gamma)$ -sparsifier of $\mathcal{G} = \mathcal{G}_{\{h,\ell\}}$. Using Prop. 3 it is straightforward to see that under A we have

$$\begin{aligned} \tilde{p}_{h+1,e} &\leq \tilde{r}_{h+1,e}(\gamma) = (1 - \varepsilon) \mathbf{b}_e^\top (\mathbf{L}_{\mathcal{H}} + (1 + \varepsilon) \gamma \mathbf{I})^{-1} \mathbf{b}_e \\ &\leq (1 - \varepsilon) \mathbf{b}_e^\top ((1 - \varepsilon) \mathbf{L}_{\mathcal{G}} - 2\varepsilon \gamma \mathbf{I} + (1 + \varepsilon) \gamma \mathbf{I})^{-1} \mathbf{b}_e = \frac{(1 - \varepsilon)}{(1 - \varepsilon)} \mathbf{b}_e^\top (\mathbf{L}_{\mathcal{G}} + \gamma \mathbf{I})^{-1} \mathbf{b}_e = r_{h+1,e} = p_{h+1,e}. \end{aligned}$$

Letting $q = |\mathcal{H}_{\{h,\ell\}}| = \sum_{e=1}^m q_{h,e} = \sum_{j=1}^{\bar{q}} \sum_{e=1}^m z_{h,e,j}$ be the random number of edges in $\mathcal{H}_{\{h,\ell\}}$, we reformulate

$$\begin{aligned} \mathbb{P} \left(|\mathcal{H}_{\{h,\ell\}}| \geq 3\bar{q}d_{\text{eff}}^{\{h,\ell\}}(\gamma) \cap \left\{ \forall h' \in \{1, \dots, h\} : \left(\|\mathbf{P}^{h'} - \tilde{\mathbf{P}}^{h'}\|_2 \leq \varepsilon \right) \leq \varepsilon \right\} \right) \\ = \mathbb{P} \left(|\mathcal{H}_{\{h,\ell\}}| \geq 3\bar{q}d_{\text{eff}}^{\{h,\ell\}}(\gamma) \cap A \right) \\ = \mathbb{P} \left(\sum_{j=1}^{\bar{q}} \sum_{e=1}^m z_{h,e,j} \geq 3\bar{q}d_{\text{eff}}^{\{h,\ell\}}(\gamma) \cap A \right) \\ = \mathbb{P} \left(\sum_{j=1}^{\bar{q}} \sum_{e=1}^m z_{h,e,j} \geq 3\bar{q}d_{\text{eff}}^{\{h,\ell\}}(\gamma) \mid A \right) \mathbb{P}(A). \end{aligned}$$

While we do know that the $z_{h,e,j}$ are Bernoulli random variables (since they are either 0 or 1), it is not easy to compute the success probability of each $z_{h,e,j}$, and in addition there could be dependencies between $z_{h,e,j}$ and $z_{h,e',j'}$. Similarly to Lem. 2, we are going to find a stochastic variable to dominate $z_{h,e,j}$. Denoting with $u'_{s,e,j} \sim \mathcal{U}(0, 1)$ a uniform random variable, we will define $w'_{s,e,j}$ as

$$w'_{s,e,j} | \mathcal{F}_{\{s,e',j'\}} \triangleq w'_{s,e,j} | \mathcal{F}_{s-2} \triangleq \mathbb{I} \left\{ u'_{s,e,j} \leq \frac{p_{h,e}}{\tilde{p}_{s-1,e}} \right\} \sim \mathcal{B} \left(\frac{p_{h,e}}{\tilde{p}_{s-1,e}} \right)$$

for any e' and j' such that $\{s, 1, 1\} \leq \{s, e', j'\} < \{s, e, j\}$. Note that $w'_{s,e,j}$, unlike $z_{s,e,j}$, does not have a recursive definition, and its only dependence on any other variable comes from $\tilde{p}_{s-1,e}$. First, we peel off the last step

$$\begin{aligned} \mathbb{P} \left(\sum_{j=1}^{\bar{q}} \sum_{e=1}^m z_{h,e,j} \geq g \mid A \right) &= \mathbb{E}_{\mathcal{F}_{t-1}|A} \left[\mathbb{P} \left(\sum_{j=1}^{\bar{q}} \sum_{e=1}^m \mathbb{I} \left\{ u_{h,e,j} \leq \frac{\tilde{p}_{h,e}}{\tilde{p}_{t-1,e}} \right\} z_{h-1,e,j} \geq g \mid \mathcal{F}_{t-1} \cap A \right) \right] \\ &\leq \mathbb{E}_{\mathcal{F}_{t-1}|A} \left[\mathbb{P} \left(\sum_{j=1}^{\bar{q}} \sum_{e=1}^m \mathbb{I} \left\{ u'_{h,e,j} \leq \frac{p_{h,e}}{\tilde{p}_{t-1,e}} \right\} z_{h-1,e,j} \geq g \mid \mathcal{F}_{t-1} \cap A \right) \right] = \mathbb{P} \left(\sum_{j=1}^{\bar{q}} \sum_{e=1}^m w'_{h,e,j} z_{h-1,e,j} \geq g \mid A \right), \end{aligned}$$

where we used the fact that conditioned on A , $\mathcal{H}_{\{h,\ell\}}$ is an (ε, γ) -sparsifier of $\mathcal{G}_{\{h,\ell\}}$, which guarantees that $\tilde{p}_{h,e} \leq p_{h,e}$. Plugging this in the previous bound,

$$\mathbb{P} \left(\sum_{j=1}^{\bar{q}} \sum_{e=1}^m z_{h,e,j} \geq g \mid A \right) \mathbb{P}(A) \leq \mathbb{P} \left(\sum_{j=1}^{\bar{q}} \sum_{e=1}^m w'_{h,e,j} z_{h-1,e,j} \geq g \cap A \right) \leq \mathbb{P} \left(\sum_{j=1}^{\bar{q}} \sum_{e=1}^m w'_{h,e,j} z_{h-1,e,j} \geq g \right).$$

We now proceed by peeling off layers from the end of the chain one by one. We show how to move from an iteration $s \leq h$ to $s-1$,

$$\begin{aligned} \mathbb{P} \left(\sum_{j=1}^{\bar{q}} \sum_{e=1}^m w'_{s,e,j} z_{s-1,e,j} \geq g \right) &= \mathbb{E}_{\mathcal{F}_{s-2}} \left[\mathbb{P} \left(\sum_{j=1}^{\bar{q}} \sum_{e=1}^m \mathbb{I} \left\{ u'_{s,e,j} \leq \frac{p_{h,e}}{\tilde{p}_{s-1,e}} \right\} z_{s-1,e,j} \geq g \mid \mathcal{F}_{s-2} \right) \right] \\ &= \mathbb{E}_{\mathcal{F}_{s-2}} \left[\mathbb{P} \left(\sum_{j=1}^{\bar{q}} \sum_{e=1}^m \mathbb{I} \left\{ u'_{s,e,j} \leq \frac{p_{h,e}}{\tilde{p}_{s-1,e}} \right\} \mathbb{I} \left\{ u_{s-1,e,j} \leq \frac{\tilde{p}_{s-1,e}}{\tilde{p}_{s-2,e}} \right\} z_{s-2,e,j} \geq g \mid \mathcal{F}_{s-2} \right) \right] \\ &= \mathbb{E}_{\mathcal{F}_{s-2}} \left[\mathbb{P} \left(\sum_{j=1}^{\bar{q}} \sum_{e=1}^m \mathbb{I} \left\{ u'_{s-1,e,j} \leq \frac{p_{h,e}}{\tilde{p}_{s-2,e}} \right\} z_{s-2,e,j} \geq g \mid \mathcal{F}_{s-2} \right) \right] = \mathbb{P} \left(\sum_{j=1}^{\bar{q}} \sum_{e=1}^m w'_{s-1,e,j} z_{s-2,e,j} \geq g \right). \end{aligned}$$

Applying this repeatedly from $s = h$ to $s = 2$, we have

$$\mathbb{P} \left(\sum_{j=1}^{\bar{q}} \sum_{e=1}^m w'_{h,e,j} z_{h-1,e,j} \geq g \right) = \mathbb{P} \left(\sum_{j=1}^{\bar{q}} \sum_{e=1}^m w'_{1,e,j} z_{0,e,j} \geq g \right) = \mathbb{P} \left(\sum_{j=1}^{\bar{q}} \sum_{e=1}^m w'_{1,e,j} \geq g \right).$$

Now, all the $w'_{1,e,j}$ are independent Bernoulli random variables, and we can bound their sum with a Hoeffding-like bound using Markov inequality,

$$\begin{aligned}
 \mathbb{P}\left(\sum_{j=1}^{\bar{q}} \sum_{e=1}^m w'_{1,e,j} \geq g\right) &= \inf_{\theta > 0} \mathbb{P}\left(e^{\sum_{j=1}^{\bar{q}} \sum_{e=1}^m \theta w'_{1,e,j}} \geq e^{\theta g}\right) \\
 &\leq \inf_{\theta > 0} \frac{\mathbb{E}\left[e^{\sum_{j=1}^{\bar{q}} \sum_{e=1}^m \theta w'_{1,e,j}}\right]}{e^{\theta g}} = \inf_{\theta > 0} \frac{\mathbb{E}\left[\prod_{j=1}^{\bar{q}} \prod_{e=1}^m e^{\theta w'_{1,e,j}}\right]}{e^{\theta g}} = \inf_{\theta > 0} \frac{\prod_{j=1}^{\bar{q}} \prod_{e=1}^m \mathbb{E}\left[e^{\theta w'_{1,e,j}}\right]}{e^{\theta g}} \\
 &= \inf_{\theta > 0} \frac{\prod_{j=1}^{\bar{q}} \prod_{e=1}^m (p_{h,e} e^{\theta} + (1 - p_{h,e}))}{e^{\theta g}} = \inf_{\theta > 0} \frac{\prod_{j=1}^{\bar{q}} \prod_{e=1}^m (1 + p_{h,e}(e^{\theta} - 1))}{e^{\theta g}} \\
 &\leq \inf_{\theta > 0} \frac{\prod_{j=1}^{\bar{q}} \prod_{e=1}^m e^{p_{h,e}(e^{\theta} - 1)}}{e^{\theta g}} \leq \inf_{\theta > 0} \frac{e^{\bar{q}(e^{\theta} - 1) \sum_{e=1}^m p_{h,e}}}{e^{\theta g}} = \inf_{\theta > 0} e^{(d_{\text{eff}}^{\{h,\ell\}}(\gamma) \bar{q}(e^{\theta} - 1) - \theta g)} \leq \inf_{\theta > 0} e^{(d_{\text{eff}}^{\{h,\ell\}}(\gamma) \bar{q}(e^{\theta} - 1) - \theta g)},
 \end{aligned}$$

where we use the fact that $1 + x \leq e^x$, $w'_{1,e,j} \sim \mathcal{B}(p_{h,e})$ and by Def. 2, $\sum_{e=1}^m p_{h,e} = \sum_{e=1}^m r_{h,e} = d_{\text{eff}}^{\{h,\ell\}}(\gamma)$. The choice of θ minimizing the previous expression is obtained as

$$\frac{d}{d\theta} e^{(\bar{q} d_{\text{eff}}^{\{h,\ell\}}(\gamma)(e^{\theta} - 1) - \theta g)} = e^{(\bar{q} d_{\text{eff}}^{\{h,\ell\}}(\gamma)(e^{\theta} - 1) - \theta g)} (\bar{q} d_{\text{eff}}^{\{h,\ell\}}(\gamma) e^{\theta} - g) = 0,$$

and thus $\theta = \log(g / (\bar{q} d_{\text{eff}}^{\{h,\ell\}}(\gamma)))$. Plugging this in the previous bound,

$$\begin{aligned}
 \inf_{\theta} \exp\left\{\bar{q} d_{\text{eff}}^{\{h,\ell\}}(\gamma)(e^{\theta} - 1) - \theta g\right\} &= \exp\left\{g - \bar{q} d_{\text{eff}}^{\{h,\ell\}}(\gamma) - g \log\left(\frac{g}{\bar{q} d_{\text{eff}}^{\{h,\ell\}}(\gamma)}\right)\right\} \\
 &= \exp\left\{-g \left(\log\left(\frac{g}{\bar{q} d_{\text{eff}}^{\{h,\ell\}}(\gamma)}\right) - 1\right)\right\} e^{-\bar{q} d_{\text{eff}}^{\{h,\ell\}}(\gamma)},
 \end{aligned}$$

and choosing $g = 3\bar{q} d_{\text{eff}}^{\{h,\ell\}}(\gamma)$, we conclude our proof.

C. Further details on experiments

Software and hardware All our code is implemented in Julia and runs on a cluster of 4 machines with 128 GB of RAM and a 10-core Xeon E5-2630. The distributed computation is achieved using a suboptimal but simple producer-consumer queue.

Linear solver We use approximate Gaussian elimination scheme by [Kyng & Sachdeva \(2016\)](#) from `Laplacians.jl`⁹ package.

Details on DiSRe. To compute all (ε, γ) sparsifiers, we use **DiSRe** after splitting the input graph into 8 sub-graphs,¹⁰ resulting in 3 rounds of resparsifications using 4 machines. For each resparsification, we compute the effective resistance in parallel on each machine using 10 processes.¹¹

Additional results. For completeness, in the following we provide the tables containing all the combinations of hyper-parameters $(\varepsilon, \gamma, \bar{q}, k, \sigma, l)$ used in our experiments. As the supplementary files, we also provide two OpenDocument spreadsheets containing these results, one for the Laplacian smoothing experiment (`extra_results_smoothing.ods`) and one for the SSL experiment (`extra_results_ssl.ods`).

⁹<http://github.com/danspielman/Laplacians.jl>

¹⁰To guarantee that each of the sub-graphs is defined on the same set of nodes, we pre-construct a spanning tree of $\mathcal{G}$ and include it in each of the sub-graphs. Note that the analysis holds even if each sub-graph is disconnected from the others. We choose this approach to avoid complicating the code with additional searches of connected components in the sparsifiers.

¹¹`Laplacian.jl` is strictly single-threaded. Therefore, we use multiple processes on a single machine. Faster runtime and lower memory usage could be achieved by sharing the memory with threads.

Improved Large-Scale Graph Learning through Ridge Spectral Sparsification

# labels	λ	mean	std	# labels	λ	mean	std	# labels	λ	mean	std
20	1.00E-006	0.4117	0.0539	20	1.00E-006	0.4308	0.0527	20	1.00E-006	0.4354	0.0447
20	1.00E-004	0.3989	0.0548	20	1.00E-004	0.4121	0.0537	20	1.00E-004	0.4316	0.0402
20	1.00E-002	0.4489	0.0662	20	1.00E-002	0.4274	0.0656	20	1.00E-002	0.4251	0.0426
20	1.00E+000	0.4109	0.0464	20	1.00E+000	0.4343	0.0688	20	1.00E+000	0.4186	0.0343
346	1.00E-006	0.3145	0.0165	346	1.00E-006	0.3144	0.0158	346	1.00E-006	0.3293	0.0143
346	1.00E-004	0.3327	0.0477	346	1.00E-004	0.3161	0.0245	346	1.00E-004	0.3480	0.0333
346	1.00E-002	0.3526	0.0813	346	1.00E-002	0.3783	0.0981	346	1.00E-002	0.3425	0.0218
346	1.00E+000	0.4357	0.1084	346	1.00E+000	0.3867	0.1009	346	1.00E+000	0.3493	0.0349
672	1.00E-006	0.2967	0.0154	672	1.00E-006	0.2957	0.0089	672	1.00E-006	0.3112	0.0277
672	1.00E-004	0.3334	0.0534	672	1.00E-004	0.3063	0.0424	672	1.00E-004	0.3219	0.0219
672	1.00E-002	0.3718	0.0971	672	1.00E-002	0.3710	0.1022	672	1.00E-002	0.3112	0.0141
672	1.00E+000	0.3197	0.0481	672	1.00E+000	0.3107	0.0244	672	1.00E+000	0.3235	0.0121
1000	1.00E-006	0.2795	0.0053	1000	1.00E-006	0.2830	0.0078	1000	1.00E-006	0.3127	0.0225
1000	1.00E-004	0.2963	0.0542	1000	1.00E-004	0.3200	0.0719	1000	1.00E-004	0.3028	0.0142
1000	1.00E-002	0.3418	0.0706	1000	1.00E-002	0.3284	0.0769	1000	1.00E-002	0.2947	0.0104
1000	1.00E+000	0.3578	0.0937	1000	1.00E+000	0.3257	0.0786	1000	1.00E+000	0.3025	0.0128

SSL, **Left: DiSRe**, $\bar{q} = 100$, **Middle: DiSRe**, $\bar{q} = 150$, **Right: kN**, $k = 60$

# labels	λ	mean	std	# labels	λ	mean	std
20	1.00E-006	0.4477	0.0299	20	1.00E-006	0.4047	0.0544
20	1.00E-004	0.4389	0.0352	20	1.87E-003	0.4381	0.0841
20	1.00E-002	0.4269	0.0364	20	1.23E-002	0.4499	0.0762
20	1.00E+000	0.4208	0.0441	20	5.34E-001	0.4255	0.0641
346	1.00E-006	0.3533	0.0378	346	1.00E-006	0.3120	0.0218
346	1.00E-004	0.3341	0.0248	346	1.87E-003	0.3553	0.0770
346	1.00E-002	0.3431	0.0388	346	1.23E-002	0.3654	0.0915
346	1.00E+000	0.3628	0.0261	346	5.34E-001	0.3318	0.0544
672	1.00E-006	0.3248	0.0263	672	1.00E-006	0.2863	0.0101
672	1.00E-004	0.3119	0.0237	672	1.87E-003	0.3293	0.0388
672	1.00E-002	0.3127	0.0138	672	1.23E-002	0.4105	0.1056
672	1.00E+000	0.3253	0.0333	672	5.34E-001	0.3470	0.0695
1000	1.00E-006	0.3107	0.0147	1000	1.00E-006	0.2772	0.0049
1000	1.00E-004	0.3076	0.0115	1000	1.87E-003	0.3181	0.0414
1000	1.00E-002	0.3073	0.0238	1000	1.23E-002	0.3070	0.0533
1000	1.00E+000	0.2952	0.0146	1000	5.34E-001	0.3063	0.0860

SSL, **Left: kN**, $k = 90$, **Right: EXACT**

σ	λ	mean	std	σ	λ	mean	std	σ	λ	mean	std
0.001	0.001	0.1908	0.0006	0.001	0.001	0.1902	0.0004	0.001	0.001	0.1910	0.0007
0.001	0.01	0.0681	0.0003	0.001	0.01	0.0681	0.0004	0.001	0.01	0.0723	0.0004
0.001	0.1	0.2845	0.0006	0.001	0.1	0.2846	0.0005	0.001	0.1	0.2622	0.0020
0.001	1	0.7991	0.0002	0.001	1	0.7993	0.0003	0.001	1	0.7198	0.0034
0.001	10	0.9728	0.0001	0.001	10	0.9729	0.0001	0.001	10	0.9096	0.0047
0.01	0.001	19.0389	0.0464	0.01	0.001	18.9854	0.0510	0.01	0.001	19.0615	0.0551
0.01	0.01	5.3282	0.0273	0.01	0.01	5.2975	0.0240	0.01	0.01	5.5206	0.0143
0.01	0.1	0.7587	0.0059	0.01	0.1	0.7565	0.0053	0.01	0.1	0.9020	0.0088
0.01	1	0.8129	0.0025	0.01	1	0.8124	0.0014	0.01	1	0.7890	0.0050
0.01	10	0.9731	0.0003	0.01	10	0.9731	0.0004	0.01	10	0.9523	0.0056
0.1	0.001	1904.4359	5.9497	0.1	0.001	1899.5487	8.1249	0.1	0.001	1906.4671	2.9946
0.1	0.01	530.8727	2.7780	0.1	0.01	528.6970	1.6710	0.1	0.01	550.5067	1.7484
0.1	0.1	48.0577	0.4268	0.1	0.1	48.0165	0.3385	0.1	0.1	64.5503	0.6398
0.1	1	2.1887	0.0547	0.1	1	2.1773	0.0563	0.1	1	7.4914	0.3399
0.1	10	1.0044	0.0136	0.1	10	1.0041	0.0086	0.1	10	5.3122	0.3776
1	0.001	190532.1581	437.1137	1	0.001	189582.9265	413.1919	1	0.001	190769.1154	348.4353
1	0.01	53197.4003	187.2996	1	0.01	52905.9901	169.3004	1	0.01	54985.9953	241.7914
1	0.1	4770.8654	39.7099	1	0.1	4775.0883	18.4023	1	0.1	6484.4296	73.4393
1	1	140.4328	3.7919	1	1	137.7268	2.7805	1	1	688.6995	39.1172
1	10	5.2586	2.0544	1	10	4.7664	1.9795	1	10	416.7218	40.3694

LAPSMO, **Left: DiSRe**, $\bar{q} = 100$, **Middle: DiSRe**, $\bar{q} = 150$, **Right: DiSRe**, $\bar{q} = 100$, $\gamma = 1000$

Improved Large-Scale Graph Learning through Ridge Spectral Sparsification

σ	λ	mean	std	σ	λ	mean	std	σ	λ	mean	std
0.001	0.001	0.1905	0.0004	0.001	0.001	0.1907	0.0007	0.001	0.001	0.3024	0.0010
0.001	0.01	0.0689	0.0003	0.001	0.01	0.0681	0.0004	0.001	0.01	0.1724	0.0004
0.001	0.1	0.2814	0.0010	0.001	0.1	0.2841	0.0005	0.001	0.1	0.2838	0.0008
0.001	1	0.7925	0.0010	0.001	1	0.7983	0.0002	0.001	1	0.7856	0.0006
0.001	10	0.9716	0.0002	0.001	10	0.9727	0.0001	0.001	10	0.9706	0.0001
0.01	0.001	19.0622	0.0340	0.01	0.001	19.0317	0.0489	0.01	0.001	30.1645	0.0703
0.01	0.01	5.3400	0.0122	0.01	0.01	5.3175	0.0253	0.01	0.01	15.7825	0.0399
0.01	0.1	0.7727	0.0047	0.01	0.1	0.7606	0.0046	0.01	0.1	1.9238	0.0107
0.01	1	0.8086	0.0029	0.01	1	0.8127	0.0027	0.01	1	0.8227	0.0027
0.01	10	0.9719	0.0004	0.01	10	0.9730	0.0006	0.01	10	0.9712	0.0005
0.1	0.001	1906.6166	6.2333	0.1	0.001	1905.0821	4.0734	0.1	0.001	3022.1232	7.4520
0.1	0.01	533.4745	2.2259	0.1	0.01	531.3563	1.5535	0.1	0.01	1575.6447	3.3189
0.1	0.1	49.6360	0.3576	0.1	0.1	48.2788	0.3639	0.1	0.1	166.1367	0.6022
0.1	1	2.3184	0.0440	0.1	1	2.1985	0.0424	0.1	1	4.5508	0.0500
0.1	10	1.0092	0.0084	0.1	10	1.0191	0.0240	0.1	10	1.0276	0.0111
1	0.001	190451.4911	600.4823	1	0.001	190205.0504	370.3744	1	0.001	302265.6761	375.8875
1	0.01	53479.7617	251.9642	1	0.01	53118.5249	176.1894	1	0.01	157503.6242	369.7657
1	0.1	4961.0662	44.4829	1	0.1	4809.2633	34.7541	1	0.1	16585.2588	43.9628
1	1	152.0926	4.5554	1	1	140.9587	3.7546	1	1	380.2291	5.5260
1	10	4.4540	2.1627	1	10	4.8436	1.5676	1	10	6.8515	1.1969

LAPSMO, **Left:** DiSRe, $\bar{q} = 100, \gamma = 100$, **Middle:** DiSRe, $\bar{q} = 100, \gamma = 10$, **Right:** kN, $k = 60$

σ	λ	mean	std	σ	λ	mean	std
0.001	0.001	0.2823	0.0007	0.001	0.001	0.1896	0.0002
0.001	0.01	0.1259	0.0003	0.001	0.01	0.0676	0.0005
0.001	0.1	0.2834	0.0008	0.001	0.1	0.2856	0.0002
0.001	1	0.7929	0.0005	0.001	1	0.8000	0.0002
0.001	10	0.9719	0.0001	0.001	10	0.9730	0.0000
0.01	0.001	28.2517	0.0560	0.01	0.001	18.9408	0.0554
0.01	0.01	11.0908	0.0309	0.01	0.01	5.2789	0.0147
0.01	0.1	1.1154	0.0087	0.01	0.1	0.7566	0.0067
0.01	1	0.8119	0.0029	0.01	1	0.8127	0.0028
0.01	10	0.9722	0.0006	0.01	10	0.9737	0.0003
0.1	0.001	2823.5932	6.3830	0.1	0.001	1894.0375	5.7404
0.1	0.01	1108.9582	5.3674	0.1	0.01	527.5392	1.0052
0.1	0.1	84.0189	0.4347	0.1	0.1	47.7877	0.4702
0.1	1	2.8117	0.0394	0.1	1	2.1503	0.0282
0.1	10	1.0124	0.0178	0.1	10	1.0083	0.0184
1	0.001	282388.0211	754.7190	1	0.001	189515.7627	293.8406
1	0.01	110729.9273	250.8183	1	0.01	52639.4113	346.7699
1	0.1	8373.5778	55.8408	1	0.1	4724.3469	44.4066
1	1	200.9318	3.7598	1	1	138.3458	2.9040
1	10	4.8285	1.0723	1	10	4.1240	0.9370

LAPSMO, **Left:** kN, $k = 90$, **Right:** EXACT