

Maximizing edge connectivity in graph partitioning using hotspots

Isam A. Alobaidi^{a*}, Hiba G. Fareed^b, Jennifer L. Leopold^c, and Andrea E. Smith^c

^a*School of Computer Science and Information Systems, Northwest Missouri State University, Maryville, MO, United States*

^b*Mathematics Department, Mustansiriyah University, Baghdad, Iraq*

^c*Department of Computer Science, Missouri University of Science and Technology, Rolla, MO, United States*

CHRONICLE

Article history:

Received October 31, 2024

Received in revised format March 10, 2025

Accepted April 24 2025

Available online

April 24 2025

Keywords:

Graph partitioning

Graph data mining

Structures

Hotspot

ABSTRACT

Graphs have long been used to model relationships between entities. For some applications, a single graph is sufficient; for other problems, a collection of graphs may be more appropriate to represent the underlying data. Many contemporary problem domains, for which graphs are an ideal data model, contain an enormous amount of data (e.g., social networks). Hence, researchers frequently employ parallelized or distributed processing. The graph data must first be partitioned and assigned to the multiple processors in a way that the workload is balanced and inter-processor communication is minimized. The latter problem may be complicated by the existence of edges between vertices in a graph that have been assigned to different processors. Herein we introduce a strategy that combines vocabulary-based summarization of graphs (*VoG*) and detection of hotspots (i.e., vertices of high degree) to determine how a single undirected graph should be partitioned to optimize multi-processor load balancing and minimize the number of edges that exist between the partitioned subgraphs. We benchmark our method against another well-known partitioning algorithm (*METIS*) to demonstrate the benefits of our approach.

© 2025 by the authors; licensee Growing Science, Canada.

1. Introduction

Graphs are often used to model relationships between entities. For some applications, a single graph is sufficient; for other problems, a collection of graphs may be more appropriate to represent the underlying data. Some of these graphs may contain an enormous amount of data (e.g., social networks) so parallelized or distributed processing often is employed. Before the analysis can commence, the graph dataset is partitioned and a subset of data is assigned to each processor. The partitioning should be done in such a way that the ensuing workload will be balanced, and inter-processor communication will be minimized. These tasks can be particularly challenging for a single graph and consideration must be given to which vertices are assigned to which partitions (i.e., processors) and what edges originally existed between those vertices. Ideally, partitions should be of approximately equal size, and the number of edges between vertices in different partitions should be minimized. The problem of finding good partitions in these respects has been studied in graph theory. Despite the numerous algorithms that have been proposed and implemented, the complexity of this problem is still considered NP-complete (Verma et al., 2017; Chen et al., 2019; Sakouhi et al., 2018). In general, most graph partitioning algorithms utilize either edge-cut partitioning or vertex-cut partitioning. Edge-cut partitioning splits the vertices of a graph into disjoint sets of approximately equal size considering the minimum number of cut-edges (e.g., PowerGraph (Gonzalez et al., 2012), Spark GraphX (Gonzalez et al., 2014), and Chaos (Roy et al., 2015). In contrast, vertex-cut partitioning splits the edges of a graph into equal-sized sets. In this approach, the partitioning of a single graph must satisfy two requirements: the quality graph partitioning criterion (which guarantees no lost data) and load balancing. Many studies have shown that edge-cut partitioning produces more accurate results on large real-world graphs (Gonzalez et al., 2012; Gonzalez et al., 2014; Chen et al., 2019).

* Corresponding author

E-mail address ialobaidi@nwmissouri.edu (I. A. Alobaidi)

Herein we introduce a novel vertex-cut partitioning strategy that determines how a single, undirected graph should be partitioned to optimize multi-processor load balancing and minimize the number of edges that exist between the partitioned subgraphs. Our approach, *GraPH*, first uses vocabulary-based summarization (Koutra et al., 2015) to identify the most highly connected structures that exist in the graph (e.g., cliques, stars, and chains). We define hotspots as the vertices in these structures with the highest degree. The hotspots become the starting points from which subgraph partitions are formed.

This paper is organized as follows. In Section 2 we briefly discuss some of the related work in graph partitioning. We present the *GraPH* algorithm in Section 3 and include a discussion of the *VoG* summarization algorithm. In Section 4 we benchmark our method against another well-known partitioning algorithm (*METIS*) to demonstrate the benefits of our approach. Concluding remarks and a discussion of future work are provided in Section 5.

2. Related works

In this section, we briefly review some of the research that has been done in graph partitioning. Graph partitioning is considered NP-Complete despite the numerous sequential, distributed, and parallel algorithms that have been developed. One of the most significant challenges of the problem continues to be minimizing the loss of information from the original graph dataset when the partitions are formed by minimizing the number of edges that exist between vertices in different partitions. This situation is more likely to occur as the number of partitions increases.

Some heuristic methods for sequential graph partitioning of a single graph are discussed in (Echbarthi & Kheddouci, 2016; Karypis & Kumar, 1998a). One offline method (wherein the entire graph resides in memory), *METIS*, is proposed in (Karypis & Kumar, 1998). This method produces high-quality partitions in terms of uniformity of partition size and minimization of “lost” edges. However, it cannot handle large graphs because of the offline setting. The *METIS* algorithm consists of three phases: coarsening, partitioning, and refinement. During each phase, a sequence of specialized algorithms is applied. These algorithms select the maximal matchings in the coarsening phase, partition the coarse graph in the partitioning phase, and project the graph back to the original graph in the refinement phase. Others propose the Streaming *METIS* partitioning (SMP) method by extending *METIS* with an online setting (Echbarthi & Kheddouci, 2016; Bourse et al., 2014). *SMP* provides the ability to adjust the memory capacity and subsequently decrease computational requirements by applying the partitioning method to small subgraphs.

Some graph partitioning techniques are designed for specific application problems. A technique for local (i.e., memory-resident, sequential processing) graph partitioning (Bonnet et al., 2015) specifically targets fixed cardinality problems such as k -densest subgraph and max k -vertex cover by developing a fixed parameter algorithm using a greediness-for-parameterization technique. Zhang et al. (2017) propose a heuristic graph edge partitioning strategy called Neighbor Expansion (NE) with polynomial running time. Their goal was to reduce the running time and communication cost for some specific applications such as triangle counting and PageRank.

Graph partitioning in a distributed environment is addressed by different researchers (Karypis & Kumar, 1998b; Kiveris et al., 2014; Park et al., 2016; Rahimian et al., 2015; Wang et al., 2014). Rahimian et al. (2015) propose a fully distributed algorithm called JA-BE-JA. This algorithm is built on two types of partitioning: vertex-cut and edge-cut partitioning; the absence of central coordination and the processing of each vertex independently make this algorithm well-designed for distributed processing. Another distributed algorithm, PACC (Partition-Aware Connected Components), based on graph partitioning for edge-filtering and load-balancing, is proposed in (Park et al., 2016). Wang et al. (2014) propose a multi-level label propagation (MLP) method that uses distributed memory across several machines for partitioning the graphs. PARallel Submodular Approximation (PARSA) was developed by Li et al. (2015) to partition a graph to fit the storage and computation ability of each machine. One important characteristic of graph partitioning algorithms is the strategy employed for selecting the vertex around which the subgraph will be built for each partition. Many algorithms select such vertices randomly. Our approach is motivated by *MELT* (Ward et al., 2017), MapReduce-based Efficient Large-scale Trajectory anonymization. The main objective of that work was to examine paths traveled by people in a geographical space and then partition the space into regions around popular locations (e.g., a coffee house, an exercise center, etc.); those locations are referred to as hotspots. The utilization of hotspots as a basis for forming partitions is a novel feature of our partitioning strategy.

3. Methodology

In this section, we present the *GraPH* strategy for partitioning a single undirected graph. We begin with some preliminary definitions that will facilitate this discussion. An explanation of the vocabulary-based summarization of graphs (*VoG*) technique developed in (Koutra et al., 2015) then follows; this is a key component for our approach as it is used to determine subgraphs of high connectivity (e.g., cliques, stars, and chains). Finally, we present our complete set of algorithms, detailing how the vocabulary-based summarization and identification of hotspots lead to the creation of optimal partitioning.

3.1 Preliminaries

Definition 1. Graph: A graph G is a tuple (V, E, L) where V is a finite set of nodes called the vertex set of G , and E is a set of 2-element subsets of V ($E \subseteq V \times V$) called the edge set of G . The nodes and edges are labeled by the function L .

Definition 2. Graph partitioning: A graph $G = (V, E)$ will be partitioned into k subgraphs $G'_{sub} = (V', E')$, $sub = 1, \dots, k$. Each $V'_{subset} \subset V_{set}$ where $V_i \cap V_j = \emptyset$ for $i \neq j$, and each $E'_{subset} \subset E_{set}$.

Definition 3. Full-clique: Let $G = (V, E)$ be an undirected graph. A set FC of vertices in G is called a Full-clique if any two distinct vertices in FC are adjacent in G , when $k \geq 1$. The Full-clique term may refer to the subgraph in some cases. If several edges are missing, this will be defined as a Near-clique.

Definition 4. Full bipartite core: Let $G = (V, E)$ be an undirected graph. A set Fb of vertices in G is called full-bipartite if two sets of vertices S_1 and S_2 , $S_1 \cap S_2 = \emptyset$, have edges between them, where each vertex in S_1 will be connected to every edge in S_2 but not within the same set. When the core is not fully connected this will be defined as a Near-bipartite core.

Definition 5. Star: A Star consists of one internal vertex in set S_1 connected to k edges of other sets S_{i+1} (spokes). A Star is considered as a special case of a Full bipartite core.

Definition 6. Chain: A Chain is a sequence of vertices such that all vertices have degree 2, except two which have degree 1.

Fig. 1 shows examples of these structure types.

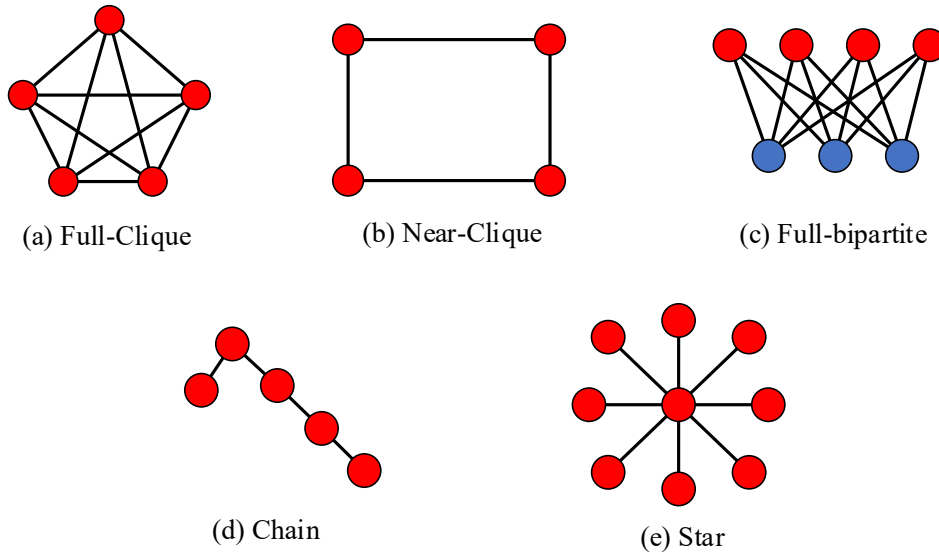


Fig. 1. Types of Structures

3.2 VOG Graph Summarization

The ability to summarize information about highly connected groups of nodes within a graph can significantly improve our overall comprehension of the graph's structure and relationships. Vocabulary-based summarization of Graphs (VoG) (Koutra et al., 2015) is a formal methodology developed for this purpose. Using a set of terms (i.e., a vocabulary) like those defined in section 3.1, VoG provides a summary of the most highly connected and frequently occurring structures in a graph. For problem domains like social networks and communication networks, these are typically the most interesting structures.

Algorithm 1 VoG

Input Graph G

Output Graph summary M , encoding cost.

- 1: **Subgraph Generation.** Using graph decomposition methods, produce a set of candidate subgraphs, which may overlap with each other
 - 2: **Subgraph Labeling.** Characterize each subgraph as one of the vocabulary structure types.
 - 3: **Summary Assembly.** From the candidate structures, select a non-redundant subset to instantiate the graph model M . Utilizing a heuristic model (e.g., PLAIN, TOP10, TOP100, GREEDY'nFORGET), the set of structures with the lowest description cost will be selected.
-

Algorithm 1 outlines the main steps that are performed in *VoG*; see (Koutra et al., 2015) for a more detailed discussion. Using graph decomposition methods, candidate subgraphs are generated then classified as various connected structures such as cliques, stars, and chains; if a subgraph qualifies as more than one of these structure types, a scoring method based on minimum description length (MDL) is used to determine which structure type that subgraph best fits. *VoG* then uses another scoring system to determine which collection of those structures best characterizes the graph. This is called the summary model, and could include all the structures (PLAIN), just the k structures with the best scores (TOP10, TOP100), or a combination of structures whose total score is best (GREEDY'nFORGET).

3.3 Proposed Algorithm

In *GRAPH*, we first use *VoG* to identify the most highly connected, and frequently occurring, subgraphs. That produces a set of structures (i.e., the model summary), S . Algorithm 2 is then used to select a subset of S which we call the majority structures, $MajS$. The number of majority structures depends on the desired number of partitions, n . The n structures in S that have the largest number of vertices become the majority structures.

Algorithm 2 Select the Majority Structures

Input S is set of structures produced by *VoG*,

n is number of desired partitions

Output $MajS$ contains n structures in S that have the largest number of vertices

```

1:   SortedS = Sort structures in  $S$  in descending order by number of vertices
2:   for  $i = 1$  to  $n$  do
3:       MajS[i] = SortedS[i]
4:   end-for
5:   return MajS

```

For each majority structure, Algorithm 3 is applied to identify the vertex that has the highest degree; in the case of a tie, an arbitrary choice between those qualifying vertices is made. These vertices of highest degree are called hotspots.

Algorithm 3 Assign the HotSpot

Input $S = (V_S, E_S)$ is a structure

Output $HotSpots$ is a vertex in V_S that is the hotspot vertex for structure $S = (V_S, E_S)$

```

1:   for  $i = 1$  to  $|V_S|$  do
2:       degree[i] = 0
3:   end-for
4:   for  $i = 1$  to  $|V_S|$  do
5:       for  $j = 1$  to  $|V_S|$  do
6:           if there is an  $edge(i, j)$  in  $E_S$  then
7:               degree[i] = degree[i] + 1
8:           end-if
9:       end-for
10:  end-for
11:  HotSpot = 1
12:  for  $i = 2$  to  $|V_S|$  do
13:      if degree[HotSpot] ≤ degree[i] then
14:          HotSpot = i
15:      end-if
16:  end-for
17:  return HotSpot

```

After assigning the hotspots, the actual partitioning commences. The subgraph that will be assigned to a partition consists of all the vertices in a hotspot's structure unless that number of vertices exceeds the total number of vertices in the graph divided by the number of desired partitions; that is considered the ideal partition size.

In Algorithm 4, we start a depth-first search from a hotspot vertex (denoted as *Hotspot*). The *MajS* denoted in the algorithm is the set of structures from which the hotspot was selected.

There are two discontinuation criteria for building a subgraph partition; the expansion will stop when either of those conditions is satisfied:

1. The current size of a partition subgraph has reached the ideal partition size.
2. The path length from the current vertex to the hotspot has reached a maximum threshold (i.e., the total number of desired partitions).

Some vertices from the original graph may not be included in any partition using these conditions. To handle those cases, we perform a breadth-first search starting from each hotspot until all nodes are included in some partition.

Algorithm 4 Graph Partitioning

Input Graph $G = (V, E)$, *HotSpot*, and *MajS*

HotSpot is a vertex in the structure connected to the largest number of edges

MajS is a set containing structures that have the largest number of vertices

n is the number of partitions

Output All *SubGraph* of G , where $|V|$ of each subgraph $\geq PartitionSize$

```

1:  PartitionSize =  $|V|/n$ 
2:  if  $|MajS_i| \leq PartitionSize$  then
3:      Include all nodes of  $MajS_i$  in  $Partition_i$ 
4:  end-if
5:  Perform DFS starting from each HotSpot
6:       $SubGraph_{DFS} \leftarrow DFS$ 
7:  Perform BFS starting from each HotSpot
8:       $SubGraph_{BFS} \leftarrow BFS$ 
9:   $SubGraph \leftarrow SubGraph_{DFS} \cup SubGraph_{BFS}$ 
10: return SubGraph

```

3.4 Computational Complexity

The complexity of one well-known partitioning method that is considered to produce high-quality partitions, *METIS* (implemented as *kmetis*), is approximately $O(V + E + k \log k)$ where V is the number of nodes, E the number of edges, and k is the number of partitions (G. Karypis, Accessed: 2019-22-01). In contrast, the complexity of *GrAPH* is approximately $O(V + E + n \log n)$ where V is the number of nodes, E is the number of edges, and n is the number of structures. Contributing to the overall complexity of *GrAPH* is the complexity of BFS and DFS, which are $O(V + E)$, and the complexity of sorting n structures, which is $O(n \log n)$. We are not including the complexity of the *VoG* processing, which has not been published by its authors.

4. Results and Analysis

In this section we compare the results of partitioning three datasets using *GrAPH* and another well-known partitioning method, *METIS*, which was discussed in Section 2. The *GrAPH* algorithms presented in Section 3.2 and 3.3 were (collectively) implemented in Matlab and C++. A C++ implementation of *METIS* was downloaded from the Karypis Lab website [G. Karypis, Accessed: 2019-22-01]. Our experiments were executed on an Intel(R) Core(TM) i7-6700 CPU@3.40GHz computer with 32 GB memory.

4.1 Data Description

Three single undirected graphs were used to evaluate our approach. Table 1 lists descriptive information about the graphs. One graph was synthetically generated; the second graph represented a two-dimensional finite element mesh; the third graph represented a three-dimensional finite element mesh.

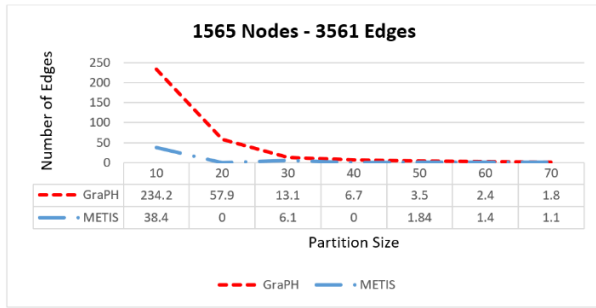
Table 1

Description of the Graphs Tested

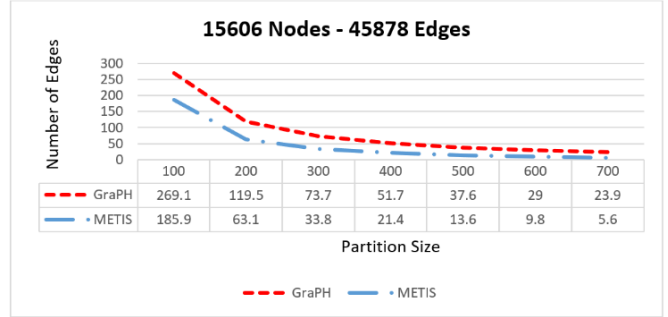
Graph Name	Number of Nodes	Number of Edges	Description
Synthetic	1565	3561	Synthetically generated
4ELT	15606	45878	2D Finite element mesh
COPTER2	55476	352238	3D Finite element mesh

4.2 Experiment and Results

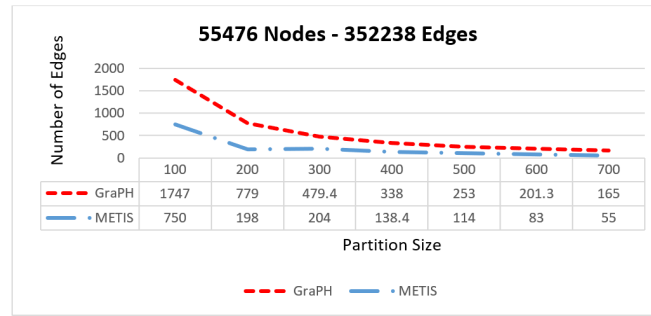
We executed *GraPH* and *METIS* on each of the graphs listed in Table 1, testing seven different numbers of partitions for each graph. The results from each test were analyzed in terms of three different metrics: the number of interior edges per partition (i.e., edges in a partition's graph), the number of exterior edges per partition (i.e., edges between vertices in a partition and vertices assigned to other partitions), and the total number of edges lost (i.e., edges from the original graph that were not represented in any of the partition graphs). Seven tests were conducted to create 10, 20, 30, 40, 50, 60, and 70 partitions, respectively, of the Synthetic graph. *METIS* failed to partition this graph into either 20 or 40 partitions; the program simply failed to return any results. *GraPH* produced results for all the tested numbers of partitions for this graph. The representation of edges amongst partitions was not well distributed when 10 partitions were requested. Specifically, the number of interior edges in one of those partitions was much higher than in the other partitions, which was not an optimal partitioning. This was likely since when a hotspot is selected from a structure, if the structure can fit entirely into a partition, all nodes from that structure automatically will be added to the partition before the depth-first search algorithm is run. This can then prevent other partitions from growing during depth-first search (as would be the case in unconnected components), encouraging disproportionate partition sizes. Because the 4ELT and COPTER2 graphs were much larger than the Synthetic graph, we tested larger numbers of partitions for those graphs, namely: 100, 200, 300, 400, 500, 600, and 700.



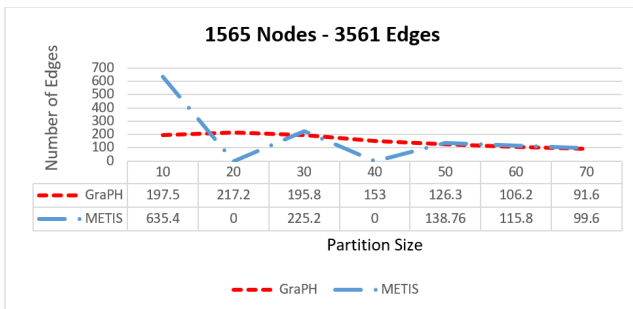
(a)



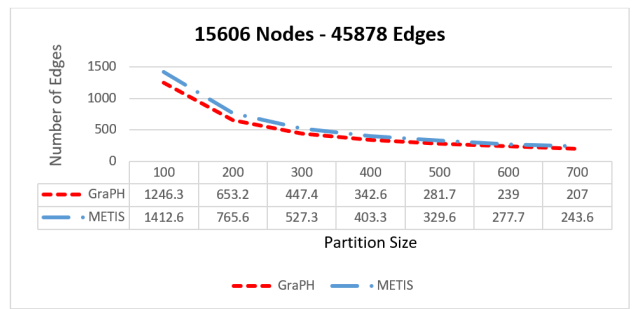
(b)



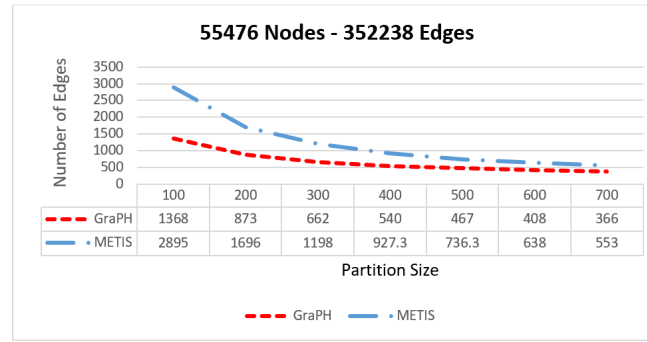
(c)

Fig. 2. (a, b, and c) Interior Edges per Partition

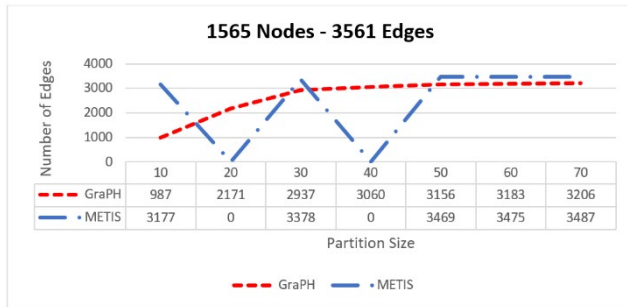
(a)



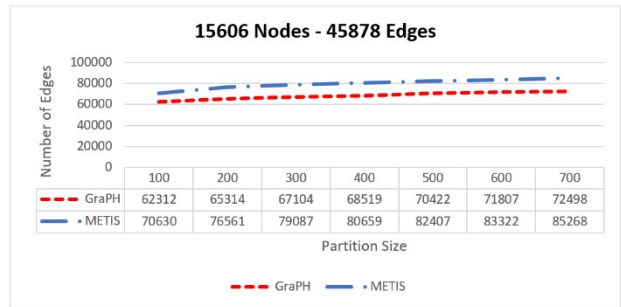
(b)



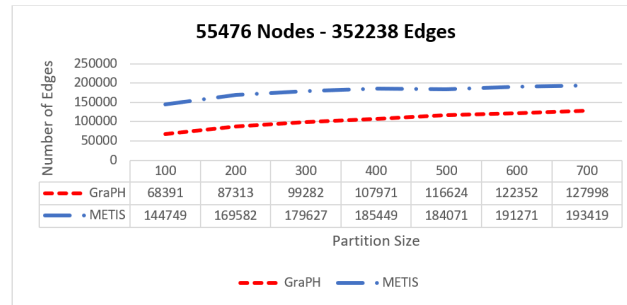
(c)

Fig. 3. (a, b, and c) Exterior Edges per Partition

(a)



(b)



(c)

Fig. 4. (a, b, and c) Total Edges Lost

For all three of the graphs listed in Table 1, in the majority of the tests, the partitions produced by *GraPH* had a higher number of interior edges in each partition than the partitions produced by *METIS*. It can be seen in Figure 2 that more edges from the original graph were retained within the partitions produced by *GraPH*. As shown in Figure 3, the *GraPH* partitioning resulted in fewer exterior edges (between partitions) than what occurred in the *METIS* partitioning. Additionally, as shown in Figure 4, *GraPH* outperformed *METIS* in terms of reducing the total number of edges lost from the original graph. It should be noted that as the desired number of partitions grew, the difference in partition quality (in terms of the three metrics) obtained from both methods became less distinct.

Because of the use of two methods (depth-first/breadth-first search) in *GraPH* for the extension process that include vertices in/out of partition boundaries, we also evaluated different variations of our method. We ran *GraPH* on the three test graphs using four different orders of processing:

1. Depth-first search extension for vertices inside the partition boundaries followed by breadth-first search extension for vertices outside the partition boundaries.
2. Breadth-first search extension for vertices inside the partition boundaries followed by depth-first search extension for vertices outside the partition boundaries.
3. Depth-first search extension for vertices inside the partition boundaries followed by depth-first search extension for vertices outside the partition boundaries.
4. Breadth-first search extension for vertices inside the partition boundaries followed by breadth-first search extension for vertices outside the partition boundaries.

We found that more consistent partitions were obtained (in terms of more interior edges and fewer external edges per partition) when we utilized the depth-first search extension process for vertices inside the boundaries followed by breadth-first search extension processing for vertices outside the boundaries. We also tested random assignment of hotspots. This was found to be unreliable in generating high-quality partitions. Interestingly, although the number of internal edges was not balanced across partitions utilizing randomization, *GrAPH* still outperformed *METIS* in terms of producing partitions with more internal edges and fewer external edges.

5. Conclusion and Future Work

With the proliferation of data in our technological world and the usefulness of modeling some problems using graphs, it is becoming increasingly difficult to process an entire graph dataset in memory. It is more efficient to partition a single large graph, and process multiple smaller subgraphs. However, in doing so, the partitioning of what may be highly interconnected data must be done in such a way as to balance the workload amongst the individual processes, minimize inter-process communication, and minimize loss of information from the original dataset. The latter problems can occur if, in the original graph, there is an edge that exists between vertices assigned to different partitions.

Herein we have presented an algorithm, *GrAPH*, for partitioning a single, undirected graph. Our algorithm strives to produce quality partitions in terms of uniformity of the size of each partition, maximization of the number of edges from the original graph that are included in each partition, and minimization of the number of edges from the original graph that effectively exist between partitions. Our approach is novel; we first utilize vocabulary-based summarization (*VoG*) to find the most highly connected structures and then find the vertices of highest degree (known as hotspots) within those structures. A benchmark comparison of *GrAPH* with another well-known, high-quality partitioning algorithm (*METIS*) demonstrated the benefits of our strategy.

In the future, we plan to explore ways to distribute or parallelize the *GrAPH* algorithms so that we can process even larger graphs than those tested for this study. To that end, we also may explore the use of some approximation (e.g., sampling) methods that may increase the efficiency of the assignment of vertices to partitions after identification of structures and hotspots.

Acknowledgement

This work was inspired by discussions with Dr. Danai Koutra, whose insights were instrumental in shaping the research direction.

References

- Bonnet, E., Escoffier, B., Paschos, V. T., & Tourniaire, E. (2015). Multi-parameter analysis for local graph partitioning problems: Using greediness for parameterization. *Algorithmica*, 71(3), 566-580.
- Bourse, F., Lelarge, M., & Vojnovic, M. (2014, August). Balanced graph edge partition. In *Proceedings of the 20th ACM SIGKDD international conference on Knowledge discovery and data mining* (pp. 1456-1465).
- Chen, R., Shi, J., Chen, Y., Zang, B., Guan, H., & Chen, H. (2019). Powerlyra: Differentiated graph computation and partitioning on skewed graphs. *ACM Transactions on Parallel Computing (TOPC)*, 5(3), 1-39.
- Echbarthi, G., & Kheddouci, H. (2016, August). Streaming METIS partitioning. In *2016 IEEE/ACM International Conference on Advances in Social Networks Analysis and Mining (ASONAM)* (pp. 17-24). IEEE.
- G. Karypis (Accessed: 2019-22-01). *Complexity of pmetis and kmetis Algorithms*. <http://glaros.dtc.umn.edu/gkhome/node/419>
- Gonzalez, J. E., Low, Y., Gu, H., Bickson, D., & Guestrin, C. (2012). {PowerGraph}: Distributed {Graph-Parallel} computation on natural graphs. In *10th USENIX symposium on operating systems design and implementation (OSDI 12)* (pp. 17-30).
- Gonzalez, J. E., Xin, R. S., Dave, A., Crankshaw, D., Franklin, M. J., & Stoica, I. (2014). {GraphX}: Graph processing in a distributed dataflow framework. In *11th USENIX symposium on operating systems design and implementation (OSDI 14)* (pp. 599-613).
- Karypis, G., & Kumar, V. (1998a). A fast and high quality multilevel scheme for partitioning irregular graphs. *SIAM Journal on scientific Computing*, 20(1), 359-392.
- Karypis, G., & Kumar, V. (1998b). A parallel algorithm for multilevel graph partitioning and sparse matrix ordering. *Journal of parallel and distributed computing*, 48(1), 71-95.
- Kiveris, R., Lattanzi, S., Mirrokni, V., Rastogi, V., & Vassilvitskii, S. (2014, November). Connected components in mapreduce and beyond. In *Proceedings of the ACM Symposium on Cloud Computing* (pp. 1-13).
- Koutra, D., Kang, U., Vreeken, J., & Faloutsos, C. (2015). Summarizing and understanding large graphs. *Statistical Analysis and Data Mining: The ASA Data Science Journal*, 8(3), 183-202.
- Li, M., Andersen, D. G., & Smola, A. J. (2015). Graph partitioning via parallel submodular approximation to accelerate distributed machine learning. *arXiv preprint arXiv:1505.04636*.

- Park, H. M., Park, N., Myaeng, S. H., & Kang, U. (2016, December). Partition aware connected component computation in distributed systems. In *2016 IEEE 16th International Conference on Data Mining (ICDM)* (pp. 420-429). IEEE.
- Rahimian, F., Payberah, A. H., Girdzijauskas, S., Jelasity, M., & Haridi, S. (2015). A distributed algorithm for large-scale graph partitioning. *ACM Transactions on Autonomous and Adaptive Systems (TAAS)*, 10(2), 1-24.
- Roy, A., Bindschaedler, L., Malicevic, J., & Zwaenepoel, W. (2015, October). Chaos: Scale-out graph processing from secondary storage. In *Proceedings of the 25th Symposium on Operating Systems Principles* (pp. 410-424).
- Sakouhi, C., Khaldi, A., & Ghezal, H. B. (2018). An overview of recent graph partitioning algorithms. In *Proceedings of the International Conference on Parallel and Distributed Processing Techniques and Applications (PDPTA)* (pp. 408-414). The Steering Committee of The World Congress in Computer Science, Computer Engineering and Applied Computing (WorldComp).
- Verma, S., Leslie L., Shin Y., & Gupta I. (2017). An Experimental Comparison of Partitioning Strategies in Distributed Graph Processing. In *Proceedings of the 43rd International Conference on Very Large Data Bases (VLDB) Endowment*, vol. 10, no. 5, pp. 493–504.
- Wang, L., Xiao, Y., Shao, B., & Wang, H. (2014, March). How to partition a billion-node graph. In *2014 IEEE 30th International Conference on Data Engineering* (pp. 568-579). IEEE.
- Ward, K., Lin, D., & Madria, S. (2017, June). Melt: Mapreduce-based efficient large-scale trajectory anonymization. In *Proceedings of the 29th International Conference on Scientific and Statistical Database Management* (pp. 1-6).
- Zhang, C., Wei, F., Liu, Q., Tang, Z. G., & Li, Z. (2017, August). Graph edge partitioning via neighborhood heuristic. In *Proceedings of the 23rd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp. 605-614).



© 2025 by the authors; licensee Growing Science, Canada. This is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC-BY) license (<http://creativecommons.org/licenses/by/4.0/>).