

Original Article

Scaling Graph Partitioning Techniques for Information Retrieval: A Methodological Exploration

Ibrahim Atoum¹, Tarek Kanan², Maraw Fayiz Hamza³

¹Department of Artificial Intelligence, Faculty of Science and Information Technology, Al-Zaytoonah University of Jordan, Amman, Jordan.

²University of Jordan, Amman, Jordan.

³Department of Computer Science, Faculty of Science and Information Technology, Al-Zaytoonah University of Jordan, Amman, Jordan.

¹Corresponding Author : i.atoum@zuj.edu.jo

Received: 08 February 2025

Revised: 10 March 2025

Accepted: 11 April 2025

Published: 29 April 2025

Abstract - This article examines scaling graph partitioning techniques through traditional and Machine Learning (ML) approaches while discussing challenges and solutions. Effective partitioning of large Information Retrieval (IR) datasets requires examining critical factors, including k selection choices alongside eigenvector selection and initial strategies while managing data uncertainty. The research strongly emphasizes integrating ML to allow systems to dynamically adapt and improve indexing, query processing, and clustering within large document collections and knowledge graphs. The article examines multiple methods based on their performance with large graphs, their community detection capabilities, parallelization convenience, and the flexibility derived from ML. Through Graph Neural Networks (GNNs) and Reinforcement Learning (RL), ML optimizes partitions by learning from evolving relationships and retrieving performance feedback. The research addresses conventional issues, including computational complexity and workload prediction, while examining ML limitations, which depend on labeled data and face interpretability concerns. The analysis covers several mitigation strategies that boost scalability, adaptive learning systems, and online learning approaches. The examination highlights the essential function of Graph Partitioning (GP) for IR system improvement and the growing influence of ML in this area. The discussion examines applications like social network analysis and fraud detection while exploring the potential advancements in dynamic GP for IR and the expected expansion of ML-based solutions. Through the discussions, it becomes clear that hybrid techniques combining ML methods with traditional approaches lead to better partitioning performance, which creates more resilient and scalable IR systems.

Keywords - Scaling, Graph partitioning, Machine learning, Information retrieval, Dynamic graph partitioning.

1. Introduction

The explosion of digital data has revolutionized IR systems, demanding efficient management of massive repositories. Graph theory offers a powerful approach to represent entities (e.g., documents, users) and their interrelationships, modelling complex connections often missed by traditional linear IR models like Boolean Retrieval Model (BRM), Vector Space Model (VSM), and Latent Semantic Analysis (LSA) [1]. While these models provide a foundation for IR, their scalability limitations become apparent with modern dataset sizes and complexities [2]. GP Methods offer a promising solution [3]. By partitioning a graph $G(V,E)$. They enhance IR efficiency in balanced subsets while minimizing edge cuts, particularly in distributed environments. These methods fall into deterministic (fixed algorithms) and stochastic (incorporating randomness) categories.

However, the sheer scale of modern graphs poses significant challenges [4]. The computational cost of partitioning massive graphs can create performance

bottlenecks in large-scale IR systems [5]. Efficient strategies are crucial, as partitioning enables parallel processing, reduces memory consumption, and improves query performance. The historical use of graph representations in IR underscores their potential for capturing complex relationships for enhanced search and question answering.

As graph-based IR systems grow, efficient GP becomes essential. Techniques like multilevel partitioning [6], spectral partitioning [1], and recursive bisection [2] address minimizing edge cuts and balancing subgraph sizes. Scaling these algorithms presents challenges, with traditional methods often reaching computational limits. Research efforts have explored parallel algorithms [3], distributed frameworks [5], and approximation techniques [1, 2] to mitigate computational costs while maintaining partition quality.

Utilizing ML presents effective solutions, especially for dynamic graphs. GNNs [6] facilitate adaptive partitioning techniques, whereas RL [7] describes partitioning as a



sequential decision process that optimizes retrieval performance. Enhanced clustering algorithms also show potential.

The evolving features of IR data require partitioning approaches to demonstrate adaptability. ML-based methods deliver adaptable solutions but require frequent model retraining and face interpretability problems while needing labelled data to avoid overfitting and high computational demands. Researchers in academia today examine the potential for online learning systems to create effective measurement tools for dynamic GP techniques. Traditional and ML methods require careful evaluation of k selection, eigenvector choices, initial partitioning strategies, and data uncertainty management.

This research thoroughly evaluates traditional and ML-based GP approaches by assessing their performance abilities, scalability, and suitability for large-scale dynamic IR systems. The study demonstrates the effects of dataset features on partitioning methods by thoroughly examining important metrics like partition quality, computational expenses, and IR-specific results, including precision and query response time. The study identifies effective methods to overcome ML challenges through adaptive strategies and online learning needs.

The spectral partitioning techniques, alongside recursive bisection, achieve a dependable reduction of edge cuts while maintaining balanced partitions. Dynamic datasets present a significant challenge for these methods, showing limited effectiveness when processing extensive data volumes. We developed an innovative methodology integrating GNNs with RL techniques to facilitate dynamic partition adjustments when graph structures evolve. GNNs use node embeddings to detect the real-time formation of new community patterns. Through retrieval feedback, IR systems gain effective scalability with dynamic large-scale datasets such as knowledge graphs and social networks by continuously utilizing the RL component to refine partitions.

The framework successfully tackles scalability challenges by employing advanced strategies like parallelized multilevel coarsening and adaptive online learning methods. According to published research findings, empirical studies demonstrate that these combined approaches surpass conventional single-method techniques. Transfer learning addresses ML challenges by diminishing the need for labeled data and enhancing system robustness throughout multiple IR contexts. Graph-based IR systems achieve greater capability through the validated combination of deterministic efficiency with ML flexibility across large-scale datasets. According to the research, hybrid methodologies present transformative opportunities to expand GP scalability.

The paper is organized as follows: Section 2 examines how deterministic methods affect scalability in GP. Section 3 provides insights into Spectral Partitioning along with

community structures in networks. Section 4 further examines Spectral Partitioning's implications. Section 5 provides a detailed comparison of Recursive Bisection techniques. The sixth section examines the application of stochastic models in data analysis. The paper presents its findings in Section 7 and offers the final conclusion in Section 8.

2. Scalability in Graph Partitioning: Deterministic Approaches

Deterministic GP [8] offers an efficient solution to graph partition scalability problems because it is deterministic. Structured methods produce replicable outcomes by using predefined algorithms that remove randomness.

The partitioning process divides a graph into separate divisions while optimizing specific criteria, minimizing edge cuts between partitions, and ensuring balanced partition sizes and consistent node connectivity [9].

K-way Partitioning represents a deterministic GP method that segments the graph into k partitions to minimize edge cuts between partitions and achieve balanced sizes across partitions. The Balanced Partitioning method generates equal-sized partitions so different workloads can be distributed efficiently across various situations. Additionally, Multilevel Partitioning employs a hierarchical strategy: The approach involves reducing the graph size through coarsening, partitioning the smaller graph, and refining these partitions to match the original graph's scale.

Deterministic methods provide stable performance for applications through their reliable and predictable nature. Users can obtain predictions through the use of designated parameters with these methods. Researchers and practitioners achieve efficient graph analyses by implementing these methods to tackle the scalability issues found in GP. The following sections present each technique, explaining its methods and benefits while showing how it works in real-world scenarios.

All three partitioning methods share a common goal: Partitioning methods aim to break down graphs or datasets into smaller manageable units. Optimizing performance and resource utilization across partitioning methods requires balanced partition sizes. Balanced partition sizes help eliminate bottlenecks while improving partitioning strategy efficiency.

Partitioning techniques are widely used across multiple areas of application. Parallel computing involves distributing tasks to minimize idle periods while maximizing throughput performance. Database management performance improves through data sharding as it divides databases into smaller, more manageable segments. Cloud computing uses these partitioning techniques to achieve efficient load balancing and resource allocation [5, 10, 11].

The Kernighan-Lin algorithm and spectral partitioning represent two key partitioning algorithms. The two methods ensure balanced partitions and reduce edge cuts effectively.

All partitioning methods present specific challenges, especially regarding computational complexity. Balanced partitions become particularly challenging to sustain when dealing with irregular graph structures or vast amounts of data, which adds complexity to the partitioning task.

Multiple strategies can be implemented to overcome these challenges. Heuristic methods enable faster partitioning, and adaptive strategies provide flexibility in adapting to workload changes. The multilevel approach reduces the complexity of the problem while improving performance by dividing it into manageable steps.

The significance of empirical testing remains paramount throughout the process. Following established guidelines and performing empirical testing are essential steps to identify the best partition sizes and configurations and validate the effectiveness of selected methods within particular applications.

2.1. K-Way Partitioning: Techniques, Challenges, and Solutions

The primary goal of k -way partitioning is to minimize edges between partitions, reducing communication costs in distributed systems while maintaining balanced sizes [12, 13]. It is widely used in load balancing and resource allocation and is essential for enhancing IR systems. Although distinct from clustering, which groups similar data points, clustering results often serve as input for k -way partitioning to distribute data or workloads. These applications are vital for improving IR systems through parallel processing and efficient resource utilization [12].

While k -way partitioning offers significant advantages, the method's effectiveness hinges on selecting an appropriate value for k . The choice of k is influenced by several factors, such as application requirements, available resources, and graph size. Larger graphs may require a higher k to create manageable partitions.

While increasing k can enhance load balancing, lower communication costs, and introduce greater management overhead. Therefore, empirical testing and adherence to domain-specific guidelines are vital for determining the optimal value of k [14].

Finding the optimal k involves balancing various factors and ensuring that the resulting partitions meet formal criteria. These criteria guarantee both the completeness and correctness of the partitioning. Specifically, the conditions for this partitioning, where V represents the set of all vertices (nodes) in the graph and can be expressed in two ways. First, the union of all partitions must equal the original vertex set; $V_1 \cup V_2 \cup \dots \cup V_k = V$ (Completeness). Second, the partitions must be mutually exclusive; $V_i \cap V_j = \emptyset$ for $i \neq j$ (Correctness) [15, 16].

While the completeness and correctness conditions ensure valid partitions, the effectiveness of partitioning is measured by the number of edges cut. The objective function captures this, $cut(V_1, V_2, \dots, V_k) = \sum_{i=1}^k \sum_{j \neq i} |E(V_i, V_j)|$, where $|E(V_i, V_j)|$ denotes the number of edges connecting vertices in partitions V_i and V_j . This formula calculates the number of edges crossing between all pairs of distinct partitions. Additionally, to ensure balanced partition sizes, the sizes of the partitions should ideally be close to $\frac{|V|}{k}$, meaning that the number of vertices in each partition should be approximately equal [10, 12, 15].

The objective function creates a standard to gauge partition quality. Achieving good partitioning outcomes depends on the availability of efficient algorithms. The KL algorithm represents a vital approach through its local search process that employs greedy vertex swapping between partitions to reduce edge cuts iteratively. Spectral partitioning applies eigenvalues and eigenvectors of the Laplacian matrix to map the graph into a lower dimensional space, facilitating segmentation. Despite their superior performance with large graphs, spectral methods face greater computational complexity than the KL algorithm [15]. This section provides an extensive explanation of spectral partitioning methods.

Researchers identify that K -way partitioning offers many benefits but presents numerous challenges, according to sources [5, 12, 13]. Optimal solution discovery faces obstacles because the number of partitions in expanding graphs grows exponentially alongside computational complexity. Balanced partitioning becomes challenging with irregular graphs since their uneven vertex degrees lead to significant size differences. The dynamic nature of graphs presents challenges for K -way partitioning because structural modifications require costly re-partitioning tasks. The objective function minimizes edge cuts but overlooks important aspects like data locality and communication patterns. Appropriate partition count k needs experimental determination and a deep understanding of application-specific requirements.

K -way partitioning finds applications across multiple domains, including IR and cloud computing, besides data mining and parallel computing, which demonstrates its significance in terms of scalability and computational cost efficiency and precision [12]. Large graph management demands scalable algorithms that balance partition quality and resource use. Graph coarsening and distributed computing techniques provide benefits yet require accuracy tradeoffs. The level of accuracy shows how the partitions fulfill application requirements appropriately. While the edge-cut reduction is standard, it does not always indicate actual performance because specific application goals take precedence. Heuristic accuracy also varies. Determining optimal partitioning presents a significant computational challenge because it falls into the NP-hard category. Algorithms must balance speed and accuracy. Parallel and distributed computing lowers expenses yet creates

communication overhead challenges. Algorithm design relies heavily on balancing cost efficiency, accuracy levels, and scalable performance.

2.2. Balanced Partitioning: Strategies for Scalability and Efficiency in Graphs

The balanced partitioning method in GP handles scalability through dataset division into equally sized smaller workloads or tasks [17]. Effective performance optimization and resource utilization require this balance because it decreases execution time across distributed systems while improving overall efficiency. Balanced partitioning protects large-scale computations and data analysis from the drawback of overloaded partitions [18].

Multiple techniques exist to accomplish balanced partitioning. GP algorithms like KL and spectral partitioning work to reduce edge connections between subgraphs, achieving better balance in partitioning processes [19, 20]. The KL algorithm employs an iterative local search method to reduce edge cuts by swapping vertices between partitions. Load-balancing algorithms used in server management stand alongside graph-specific algorithms as essential components. The algorithms evenly distribute incoming requests to servers, which helps avoid bottlenecks and maintains balanced resource usage [21].

To formally define balanced partitioning, consider a graph G that is to be partitioned into k subsets $P_1, P_2, \dots, P_k$. The following conditions must be met: $\bigcup_{i=1}^k P_i = V$ (all vertices are included) and $P_i \cap P_j = \emptyset$ for all $i \neq j$ (partitions are mutually exclusive) [12, 15]. To maintain balance, the sizes of the partitions should satisfy $\left| |P_i| - \frac{|V|}{k} \right| \leq \epsilon$ for a small tolerance ϵ . Furthermore, to enhance efficiency, minimizing the number of edges that cut across the partitions, expressed as minimizing $\sum_{i \neq j} |E(P_i, P_j)|$, is often desirable [18].

Despite its importance, balanced partitioning faces several inherent challenges. One significant challenge is accurately predicting workloads, which can lead to imbalances [12]. Another challenge lies in the computational complexity of some partitioning algorithms, particularly with larger datasets, which can significantly limit efficiency [18, 22]. These challenges necessitate the use of sophisticated strategies.

Several effective solutions have been developed to address these challenges. For instance, heuristic methods can provide quick partitions, while adaptive strategies can respond to real-time workload changes [12]. ML techniques can be used for workload estimation [23].

Multilevel partitioning can simplify complex GP problems [15]. For example, multilevel partitioning has been explored for GNN processing [24]. Parallel processing can distribute tasks to reduce computation times [21]. For instance, parallel processing has been used in graph signal processing [25].

Balanced partitioning is essential for effectively operating multiple applications [12, 15]. Balanced workload distribution in cloud computing systems prevents bottlenecks from occurring. This approach divides graphs into equal-sized subgraphs while minimizing edge cuts during network analysis and community detection. Parallel computing systems distribute tasks across several processors, reducing idle time and enhancing throughput capacity. The implementation of sharding in database management systems leads to better query performance. Applications include distributing resources among virtual machines while using image segmentation to detect objects and optimizing code execution by sharing memory resources evenly. Balanced GP becomes computationally complex with uneven distributions because precise algorithms are required.

Scalability problems in balanced GP intensify as graphs grow in size [5, 13]. Splitting large graphs into equally sized subsets becomes more complex, leading to higher computational expenses. A GP algorithm must achieve both balanced partitions and high-quality cuts to be efficient [12]. GP scalability issues are generally resolved through the utilization of heuristic algorithms along with parallel computing techniques.

2.3. Multilevel Partitioning: Enhancing Graph Efficiency through Structured Approaches

Multilevel partitioning is an advanced method for GP that improves efficiency through a three-phase process: coarsening, partitioning, and refining [12, 16]. Graph coarsening reduces the original graph into a more straightforward version by merging vertices and edges and repeating them to create a more manageable representation. The fundamental structure of the original graph remains present in the coarsened graph, which allows for quicker calculations and efficient partitioning of large datasets.

Multilevel partitioning is an advanced method for GP that improves efficiency through a three-phase process: coarsening, partitioning, and refining [12, 16]. The graph coarsening process begins by merging vertices and edges in multiple iterations to transform the original graph into a compact and manageable version. The coarsened graph preserves the original structure while improving computational speed and partitioning efficiency for large-scale datasets. The partitioning algorithm starts on the reduced graph that has been minimized to an optimal size using deterministic methods such as spectral partitioning or recursive bisection to achieve efficient and balanced partitions. These algorithms aim to decrease the number of edge connections between distinct partitions, which are known as edge cuts [15]. During the refinement phase, partitions get applied to the original graph through local optimization techniques such as the Kernighan-Lin algorithm and simulated annealing, which help minimize edge cuts and enhance balance.

The multilevel partitioning process can be summarized as follows: When you coarsen the original graph $G = (V, E)$

to obtain $G' = (V', E')$, the result will always contain fewer vertices and edges. Computational efficiency and scalability benefits arise from this method [5, 12, 13] while it encounters challenges, including managing detail loss during coarsening to maintain partition quality [15]. Algorithms achieve different levels of effectiveness depending on graph characteristics, while computational complexity restricts efficiency, particularly with large graphs [5].

Multiple approaches exist which can be used to solve these challenges. Adaptive coarsening methods dynamically modify detail levels according to graph properties to safeguard key structural elements [26]. Multiple partitioning levels become possible through a hierarchical approach [26], and spectral methods provide insights for efficient coarsening [27]. Specialized algorithms help improve performance and lessen complexity, while quality metrics balance detail loss and simplification.

Multilevel partitioning finds applications across various domains. The method effectively reduces the complexity of large social network graphs to enable community analysis [28, 29]. Finite element analysis (FEA) uses more minor elements to divide complex structures, enabling efficient engineering simulation [30]. The approach functions as an essential component of parallel computing [31], network design [32], and data clustering [33] because it enhances task scheduling and load balancing in matrix operations while enabling effective segmentation of large data sets. Data management systems improve storage and retrieval capabilities [12], and image processing techniques gain accuracy when images are segmented into larger groups before processing [33].

Graph efficiency in large-scale applications requires scalable multilevel partitioning methods [5, 13]. The method begins with coarsening the graph for simplification. Then, it is divided into evenly distributed partitions that minimize edge cuts before optimizing these partitions as the graph returns to its full original scale. The method enhances graph organization while efficiently minimizing communication costs to control larger graphs [12, 13]. Selecting appropriate algorithms guarantees that the computational workload remains manageable while maximizing performance benefits to achieve effective results in large-scale applications.

3. Spectral Partitioning: Unveiling Community Structures in Graphs

Spectral Partitioning (SP) is a technique in graph theory that utilizes the spectral properties of a graph's Laplacian matrix to minimize edge cuts, often revealing natural community structures within the graph [27]. This method is widely applicable in fields such as social network analysis, image segmentation, and clustering, making it essential for understanding complex relationships in data.

Let $G = (V, E)$ be a graph consisting of n vertices. The Laplacian matrix L is defined as $L = D - A$ where D is the

degree matrix (a diagonal matrix where each entry D_{ii} reflects the degree of the vertex i), and A is the graph's adjacency matrix [34]. The eigenvalues $\lambda_1, \lambda_2, \dots, \lambda_n$, and corresponding eigenvectors $v_1, v_2, \dots, v_n$ of the Laplacian matrix L are computed. The goal of spectral partitioning is to minimize the edge cut [35], which can be represented as:

$$Cut(S, \bar{S}) = \sum_{i \in S, j \in \bar{S}} A_{ij}$$

Where S and S' (often denoted as $\bar{S}$) are two disjoint subsets of vertices. The partitioning is often achieved by selecting a subset of the eigenvectors corresponding to the smallest eigenvalues (excluding the first eigenvalue, zero). This approach effectively identifies natural community structures within the graph, leading to high-quality partitions. Thus, the spectral partitioning can be summarized as $Partition(G) \rightarrow Minimize Cut(S, \bar{S})$ Using L and its eigenvalues/eigenvectors [36].

While spectral partitioning is effective for dense graphs, yielding high-quality partitions [12], it presents challenges. Its computational intensity makes it less suitable for huge graphs [3, 4], and it may struggle with graphs lacking clear community structures, potentially resulting in less meaningful partitions [5, 6].

The primary limitation of spectral partitioning is its computational demand [4, 5]. To address this, approximation algorithms or multilevel techniques can be employed to reduce computational burdens while preserving the advantages of spectral analysis [6, 30]. Preprocessing the graph to create coarser representations allows faster computations without significantly compromising partition quality [3, 4]. Additionally, integrating parallel computing can enhance efficiency by distributing the computational load across multiple processors [12, 26].

Spectral partitioning finds application across various domains. Social networks benefit from spectral methods, which identify user groups with similar interests in targeted marketing through clustering analysis of user preferences and interactions [2]. Medical image segmentation benefits from spectral methods, which distinguish tissues or tumors, thereby improving diagnostic accuracy by separating tissue types based on their distinct characteristics [1]. ML applies clustering techniques to high-dimensional datasets, improving marketing strategies by enabling more precise targeting of consumer segments. The field of computer vision supports object recognition through image segmentation into relevant regions, which facilitates object identification and classification of digital images [2]. The approach enhances parallel computing performance by optimizing GP to distribute processing tasks across multiple processors. Network structure robustness and efficient data routing benefit from node organization through spectral partitioning methods. Bioinformatics uses gene expression profiles to group genes together, providing critical information for diagnosing and treating diseases by

highlighting gene activity patterns corresponding to specific conditions.

Effective management of large datasets alongside complex graphs demands scalable spectral partitioning techniques. Preserving algorithm efficiency through approximations [5, 7] along with enabling parallelization for concurrent processing [12, 26] and implementing dimensionality reduction methods to simplify problem complexity [4, 8] stand out as the key aspects.

The system demands efficient memory management to prevent overload [5, 13] and adaptability to dynamic graph changes to sustain real-time application effectiveness [3, 31]. Spectral partitioning delivers scalable performance within social networks, bioinformatics, and computer vision applications.

4. Optimizing Graph Structures: The Role of Recursive Bisection

In both computer science and engineering fields, researchers commonly use recursive bisection as a basic GP technique [37]. The recursive bisection method creates equal-sized subsets called “bisections” from a graph at every stage while constructing a hierarchical partition structure. Parallel computing circuit design and network optimization applications depend on this method, necessitating efficient data organization [38].

The recursive bisection process, beginning with the original graph G , can be described as follows: During each recursion level, the original graph splits into two distinct subgraphs. $G_1 = (V_1, E_1)$ and $G_2 = (V_2, E_2)$, where node sets V_1 and V_2 together make up the entire node set V without any overlap [39]. The goal is to find partitions that ensure the size difference between V_1 and V_2 remains below a small tolerance value ϵ , as given by the condition $||V_1| - |V_2|| \leq \epsilon$. The essential task is to reduce the number of edges crossing between partitions, which can be formulated as minimizing $|E(G_1, G_2)|$.

This technique enables multiple processors in parallel computing environments to share tasks efficiently. Recursive bisection achieves reduced communication overhead and better processing efficiency through equal partition sizes and minimized inter-partition edges. This approach strengthens network design through performance enhancement and reliability improvements when component connections are reduced.

Recursive bisection has multiple advantages but also has substantial drawbacks [40]. The quality of results for irregular and complex graphs depends heavily on the initial partitioning strategy because it is crucial in determining the outcome. Their recursive algorithm structure increases the computational overhead in processing massive graphs. Smart initialization methods using heuristics produce better initial partitions that address these issues. Refinement algorithms improve partition quality by redistributing workload, while parallel processing implementation reduces computational overhead [41].

The recursive bisection method has applications in numerous scientific and engineering fields. Optimal task division by parallel computing load balancing enhances the performance of processors during matrix multiplication operations. FEA helps engineers transform complex geometries into essential elements, which allows for more efficient structural calculations. Network analysis enables large graph structures to split into more minor elements by making minimal edge cuts, which helps identify community structures in social networks [42].

Medical imaging requires precise structure delineation, achieved through recursive bisection techniques that segment images based on color or texture criteria by continuous division. The method partitions large databases into smaller segments using database attributes to improve query speed and overall system performance [18]. Using recursive bisection methods in geographical information systems enables spatial data to be segmented, improving query performance and analytical capabilities and assisting in resource management and planning tasks [43, 44]. Signal processing utilizes this algorithm to process audio signals by reducing noise and extracting key features during segmentation tasks.

Recursive bisection achieves balance, minimizes edge cuts, and produces unbalanced partitions when vertex degrees vary. As graph sizes expand, computational resources become more necessary, which requires exact optimization throughout each bisection step. Heuristics and parallel processing algorithms achieve balanced partition quality and efficiency by scaling to match unique graph properties.

The effectiveness of recursive bisection scalability in GP depends on various determining factors. The quest for optimal bisection cuts continues to be crucial, while spectral methods have demonstrated performance enhancements for large-scale graphs. Parallel processing enables simultaneous processing of multiple partitions, maintaining balance and preventing bottleneck formation. Managing sparse representation is vital for large graphs, and real-time application performance depends on system flexibility through changes to nodes and edges. When these factors are addressed, recursive bisection scales across multiple applications, including parallel computing and social network analysis.

5. Stochastic Models in Data Analysis: Navigating Uncertainty and Complexity

Analyzing complex networks that exhibit uncertainty demands implementing probabilistic and statistical methods throughout GP tasks [45]. Stochastic models integrate randomness into their framework for scenarios where outcomes rely on probabilistic events [46]. By employing these models, analysts create partitions that reflect the inherent variability in graph structures, leading to improved partitioning techniques [47]. Stochastic models enhance decision-making processes across multiple domains, such as

ML and operations research, using statistical characteristics like expected values and variances.

Stochastic models serve as mathematical frameworks incorporating random elements, allowing them to accurately represent complex systems that depend on probabilistic events [48]. The models create partitions based on data points' probabilistic distributions, which enables them to handle both variability and uncertainty in the data. Stochastic partitioning methods rely on random sampling techniques or probabilistic algorithms that enhance partitions through improvements in iterative data observation. These models demonstrate effective performance in dynamic environments through their capability to adjust to changing data patterns.

A stochastic model can be represented using probability distributions and random variables. Let X be a random variable representing the data points to be partitioned, drawn from a probability distribution $P(X)$. The goal is to partition the data into k subsets $P_1, P_2, \dots, P_k$ such that specific criteria are met. The union of all partitions must cover the entire data space, expressed as $\bigcup_{i=1}^k P_i = X$ and $P_i \cap P_j = \emptyset$ for all $i \neq j$. The objective function can be defined based on a cost or similarity measure; for instance, if $C(P_i)$ represents the cost associated with the partition P_i , the overall cost may be minimized as $\sum_{i=1}^k C(P_i)$. This cost could depend on factors such as intra-partition variance or the number of inter-partition edges in a graph [46].

Probabilistic methods enable partition refinement through adjustments derived from random samples taken from the distribution. Markov Chain Monte Carlo techniques investigate partition spaces by iterating partition adjustments using acceptance mechanisms that evaluate configuration cost and likelihood probabilities. Calculating the expected value of costs for each partition offers valuable support for making informed decisions.

Stochastic models handle data variability and uncertainty, making them ideal for applications in dynamic environments experiencing changing data distribution patterns over time [48]. Robust partitioning strategies benefit from random elements that help optimization processes escape local optima. These models offer statistical insights into data, which allow people to make better decisions.

The main drawback of stochastic models lies in their computational demands, which become especially noticeable with Markov Chain Monte Carlo (MCMC) methods because they need numerous iterations to reach convergence. Using variational inference techniques with approximate Bayesian approaches leads to quicker convergence alongside lower computational requirements [48]. Ensemble methods combine results from several runs to maintain consistency, whereas initial analyses boost accuracy by elucidating underlying probability distributions [22]. Stochastic models become more accessible through intuitive tools and visualization techniques [5].

Stochastic models play an essential role in handling uncertainty in numerous disciplines. The Black-Scholes model demonstrates how stochastic models aid finance by enabling option pricing and risk evaluation. Queueing theory allows the prediction of customer wait times in call centers, which helps optimize staff allocation and improve customer satisfaction [44]. Ecological models use stochastic approaches to examine species population changes and the influence of random events. Healthcare uses the SIR model to simulate how diseases spread, which supports the evaluation of public health responses [49]. The manufacturing sector uses stochastic modeling to align production and inventory levels to fluctuating demand patterns. ML depends on stochastic models for algorithm training, while sports analytics applies these models to assess player performance and team tactics.

Stochastic models use random probabilistic approaches to separate graphs into balanced subsets while minimizing edge cuts during partitioning. The algorithm rapidly finds effective partitions but produces variable results because of built-in randomness [18]. Stochastic models are typically less demanding regarding computational resources than deterministic methods. Still, they need numerous iterations to produce dependable results, which impacts their performance stability for extensive applications.

The results generated by stochastic partitioning methods vary according to the probabilistic model utilized. Well-designed systems achieve high accuracy with complex graph structures where traditional methods encounter difficulties. The model's inherent randomness introduces variability across different runs, which requires assessing its robustness and performing multiple iterations to guarantee reliable results [22].

The scalability of stochastic partitioning leverages randomization to manage large graphs efficiently. Probabilistic models use random sampling techniques to explore different configurations, which accelerates finding the best cuts.

The significance of computational complexity management grows with graph size expansion, which typically involves adopting sampling methods and parallel processing to improve performance. Larger graphs require precise tuning of model parameters to maintain accuracy and quality because they present structural challenges.

6. ML Techniques for Dynamic Graph Partitioning in IR

ML improves dynamic GP within IR systems through adaptive responses to graph structure changes, which optimize partition performance for retrieval tasks [50]. GNNs excel at this task because their learning process creates node and edge representations, allowing dynamic partitioning adjustments in response to changing graph relationships [51]. These learned embeddings enable GNNs to discover related documents or user clusters and enhance IR effectiveness.

RL enables dynamic partition adjustment by treating partitioning as a sequence of decisions [52]. Through this method, agents use feedback from metrics like query response time or document relevance to learn how to optimize partitions. This adaptive strategy enables ongoing enhancements to partitioning methods.

ML techniques enhance traditional clustering algorithms like K-means and hierarchical clustering [53]. Supervised learning can generate partitions that better align with retrieval objectives, while semi-supervised learning proves helpful when access to fully labeled data is limited. Community detection algorithms such as spectral clustering and modularity optimization combined with ML techniques work to identify community changes over time so that partitions stay applicable for IR tasks [6]. ML systems can merge with dynamic programming methodologies to enhance partitioning choices throughout graph evolution [54]. The system gains the ability to develop optimal partition update strategies in real-time by preserving historical data of past partitions and retrieval performance while balancing efficiency and accuracy.

Adaptive GP techniques apply ML to analyze the unique features of a graph during real-time evaluation. Adjusting clustering techniques according to user interactions, such as frequently accessed documents and user queries, enhances IR relevance [50]. Ensemble learning methods play a role in dynamic GP as well. These methods merge multiple ML models, which results in more resilient partitioning strategies. Ensemble methods combine outputs from different partitioning algorithms to help the system dynamically select the optimal approach based on the graph's current state and retrieval needs.

Dynamic GP powered by ML algorithms keeps partitions updated alongside the evolving data landscape. The ability to adapt to changing conditions plays a key role in preserving IR quality, tailoring partitions to fit specialized tasks, and accelerating query response times. GNNs assist in the discovery of document and user clusters that enhance search result relevance [53].

The scalability of these methods enables them to process large datasets while automatically learning from past data and user behavior to advance partitioning techniques without human intervention. GNNs stand out for their ability to capture node relationships within graphs, which proves helpful in detecting clusters among related documents or users.

The structural information of graphs enables GNNs to develop representations for nodes and their connections, which suitably match their relevance to particular queries and tasks. The dynamic adjustment capability of GNNs to modify partitions according to data changes results in improved search relevance. GNNs exhibit superior scalability that allows them to manage huge datasets efficiently. The system automates learning while persistently refining partitioning strategies through analysis of historical data and user interactions.

While these benefits exist, we must acknowledge some critical limitations. Some ML models, intense learning approaches, demand substantial computational resources because of their high complexity. These methods' effectiveness heavily depends on obtaining high-quality labelled data for training. Third, overfitting is another concern. According to literature references, understanding complex ML models presents interpretability challenges [6, 12]. The partitioning process becomes complex due to fast changes in the underlying graph, which require continuous model retraining.

Despite their advantages, there remain limitations that should be evaluated. Because of their high complexity, some ML models and intense learning methods need substantial computational resources. The success of these techniques relies heavily on access to high-quality labeled data for model training purposes [52]. Third, overfitting is another concern. The ability to interpret complex ML models remains a significant challenge. Frequent alterations to the base graph structure increase partitioning complexity, requiring continuous model retraining [58].

Despite these limitations, the applications of ML techniques in dynamic GP for IR remain varied. The applications of ML techniques for dynamic GP in IR extend to social network analysis, recommender systems, fraud detection in financial institutions, biomedical research, and IR itself [52]. Dynamic partitioning on e-learning platforms enables better resource retrieval and delivers personalized content to users [58]. IoT systems that monitor environmental conditions use adaptive partitioning to handle real-time data from sensors in forest fire detection networks [59]. Partitioning refers to data distribution in ML processes into separate training, validation, and testing groups. Good partitioning practices improve model results, yet inadequate partitioning can cause models to overfit or underfit [53]. Adaptive learning of features and relationships enables ML to achieve high accuracy when applied to dynamic graphs. The effectiveness of this method depends significantly on both the training data quality and the algorithms' suitability. ML scalability for GP seeks to improve partitioning quality through effective algorithm utilization. While parallel processing and scalable architectures enhance effectiveness, high-quality partitioning requires models to adapt specifically to graph data features while scaling.

7. Discussions

GP presents key challenges: computational complexity, dependence on initial conditions, and the tradeoff between simplification and accuracy. Table 1 summarizes various GP techniques' characteristics, strengths, weaknesses, and suitable applications.

Computational complexity measures the time and memory resources required to resolve a problem according to the size of its input. The extensive search spaces present significant challenges when employing Balanced, Multilevel, and K-Way methods. The computational cost of

solving large datasets for optimal solutions demands using heuristics or approximations and parallel computing techniques. Spectral partitioning requires substantial resources due to its dependence on eigenvalue computations. Stochastic techniques, including MCMC, incur high computational costs through their iterative methods, which restricts their utility on large datasets unless approximations or parallel processing are applied.

In optimization problems, the search space can grow significantly; for n items, the number of possible combinations is 2^n , complicating the search for optimal solutions. Balanced methods aim for equal-sized partitions, increasing complexity through numerous combinations. Multilevel methods simplify problems and refine solutions but escalate complexity due to multiple transformation levels.

K-Way methods partition input into k groups, where increasing k can lead to factorial growth in partitions, thus expanding the search space. The number of possible partitions is approximated by $P(n) \approx n! / k! (n - k)!$. Time complexity is expressed as $T(n) = O(f(n) \cdot g(n))$, with $f(n)$ for generating partitions and $g(n)$ for evaluation costs. Space complexity is often represented as $S(n) = O(n \cdot k)$.

Dependence on initial conditions significantly impacts partition quality, particularly for Recursive Bisection and Balanced partitioning. Poor starting strategies or inaccurate workload predictions can produce suboptimal or unbalanced results. Iterative algorithms are particularly sensitive, with poor starts causing suboptimal partitions, imbalances, and convergence issues. Mitigation strategies include smart initialization (e.g., spectral methods), multiple runs, and adaptive algorithms.

While various techniques address computational complexity, smart initialization also helps mitigate the impact of initial conditions if the starting point is represented as a vector. x_0 and the final solution as $x^* = f(x_0)$, a poorly chosen x_0 Can yield suboptimal results.

For Recursive Bisection, if the initial partition P_0 is unbalanced, subsequent partitions $P_1, P_2, \dots$, may also be imbalanced. The imbalance of a partition can be expressed as $I(P) = \left| \frac{W_1 - W_2}{W_1 + W_2} \right|$, where W_1 and W_2 are the weights of the partitions. A high $I(P)$ propagates through recursive splits, leading to poor outcomes. The convergence rate R can be defined as $R(P_0) = \frac{1}{1 + I(P_0)}$. A high initial imbalance results in a low convergence rate, making the algorithm slower or unstable.

Several strategies can be employed to mitigate issues related to initial conditions in partitioning algorithms. Smart Initialization uses methods like spectral techniques to provide better starting points, represented as $x_0 = \argmin \sum_{i,j} A_{ij} \cdot (x_i - x_j)^2$. Multiple Runs involve executing the algorithm multiple times and selecting the best

result, expressed as $R^* = \min(R_1, R_2, \dots, R_n)$. Additionally, Adaptive Algorithms adjust parameters based on performance feedback, enabling dynamic improvements. Together, these strategies enhance partition quality by addressing the effects of initial conditions.

The tradeoff between simplification and accuracy in GP can be expressed mathematically. Let S represent simplification (or efficiency), and A denote accuracy. The equation can define this tradeoff $T(S, A) = \alpha S - \beta A$, where α and β are weights indicating the importance of each aspect.

In multilevel partitioning, detail D can be affected by coarsening. As efficiency S increases, detail D decreases, represented as $D = D_0 - kS$, with D_0 as the initial detail and k as a constant. For stochastic methods, variability V impacts accuracy, expressed as $A = A_0 - cV$, where A_0 is the baseline accuracy, and c indicates the effect of variability.

To achieve consistent outcomes, ensemble methods can be employed, defined as $A_{ensemble} = \frac{1}{n} \sum_{i=1}^n A_i$, where A_i is the accuracy of each model, and n is the number of models.

As detailed in Table 1, GP techniques vary. Deterministic methods are efficient for large graphs but can be complex. Spectral methods are effective for community detection but computationally expensive.

The recursive bisection is simple but may lose details. Stochastic methods handle uncertainty but require significant computational resources. ML methods adapt to dynamic graphs and offer high performance but often require labeled data. The table further compares computational efficiency, scalability, and accuracy techniques.

Combining conventional GP with ML techniques markedly boosts performance efficiency and accuracy. The hybrid technique utilizes the advantages of both methodologies to tackle important issues like scalability and computational complexity.

Efficiency improves when traditional methods merge with adaptive ML techniques because algorithms optimize starting parameters, decreasing the time required to determine optimal partitions. Hybrid methods achieve better outcomes by refining traditional results through learned patterns, ensuring both simplification and precision.

Hybrid techniques demonstrate strong robustness and scalability while effectively processing multiple data types. The partitioning strategies improve because ML components learn from new data continuously during graph evolution. ML insights generate smart initialization methods that enhance starting points during partitioning and reduce the likelihood of suboptimal outcomes. Adapting to multiple domains allows practitioners to create effective and customized solutions for different challenges.

Table 1. characteristics of different graph partitioning techniques

Feature	Deterministic	Spectral	Recursive Bisection	Stochastic	ML
Definition	k -subset partitioning.	Spectral Laplacian partitioning.	Recursive two-subset division.	Probabilistic partitioning.	Adaptive, learning-based partitioning.
Objective	Minimize cuts and balance.	Minimize cuts and find communities.	Minimize cuts, balanced partitions.	Robust, flexible partitions.	Optimize IR effectiveness on dynamic graphs.
Applications	Load balancing.	Social networks, images.	Parallel computing.	Data analysis.	Dynamic IR recommendations.
Considerations	k selection, balance.	Eigenvectors.	Initial strategy.	Model, uncertainty.	Features, model, data, drift.
Strengths	Efficient large graphs.	Good communities.	Simple.	Handles uncertainty.	Adapts, high IR performance.
Challenges	Complexity.	Workload prediction.	Detail loss.	Intensity.	Labelled data, overfitting.
Solutions	Heuristics, parallel.	Approximations.	Adaptive coarsening.	Smart initialization	Online learning, transfer learning.
Computational Efficiency	Generally high, depending on a specific algorithm.	Computationally intensive (eigenvalue decomposition).	Relatively efficient.	It varies greatly; it can be very intensive.	It can be computationally expensive, especially training.
Scalability	Good, especially with parallelization.	It can be challenging for huge graphs.	Moderate; recursion can become expensive.	Highly variable; some scale well, others don't.	Can scale well with appropriate techniques (e.g., distributed training).
Accuracy	Varies can be suboptimal if the structure is ignored.	High for community detection but computationally costly.	It can be good, but it may be unbalanced.	It depends on the model; it can be high with good models.	Potentially very high, especially for dynamic graphs.
Application Suitability	Resource allocation, parallel computing.	Community detection and image segmentation.	Parallel computing, hierarchical data.	Data mining, clustering.	Dynamic IR, personalized recommendations, evolving networks.

8. Conclusion

The article examines scaling techniques for graph partitioning through classic and machine-learning approaches. The increasing need for effective extensive dataset partitioning within IR emphasizes critical factors, including k -selection decisions and eigenvector options, alongside data uncertainty control. Machine learning techniques that utilize GNNs and RL provide enhanced adaptability and optimized strategies for essential IR tasks such as indexing and clustering. Traditional methods encounter computational complexity issues, while ML techniques present difficulties with labelled data dependence and interpretability problems. Our analysis shows significant variations in computational efficiency: Spectral partitioning demands high computational

resources, but recursive bisection executes efficiently, while stochastic methods show broad variations in their computational requirements. Scalability is also a concern, with general methods scaling well, while spectral partitioning struggles with large graphs. Accuracy varies, with spectral methods excelling in community detection at a high cost and recursive bisection sometimes leading to unbalanced partitions. Future research should focus on developing hybrid approaches that combine traditional and ML techniques, refining algorithms for better adaptability, and exploring ML models needing less labeled data. Understanding the tradeoffs between computational efficiency and accuracy in real-world contexts is crucial for advancing graph partitioning and improving data management in complex information environments.

References

- [1] Venkat N. Gudivada, Dhana L. Rao, and Amogh R. Gudivada, *Information Retrieval: Concepts, Models, and Systems*, Handbook of Statistics, vol. 38, pp. 331-401, 2018. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]

- [2] V. Gupta, D.K. Sharma, and A. Dixit, "Review of Information Retrieval: Models, Performance Evaluation Techniques and Applications," *International Journal of Sensors Wireless Communications and Control*, vol. 11, no. 9, pp. 896-909, 2021. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [3] José Devezas, and Sérgio Nunes, "A Review of Graph-Based Models for Entity-Oriented Search," *SN Computer Science*, vol. 2, 2021. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [4] Jiakang Li et al., "A Comprehensive Review of Community Detection in Graphs," *Neurocomputing*, vol. 600, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [5] Siddhartha Sahu et al., "The Ubiquity of Large Graphs and Surprising Challenges of Graph Processing: Extended Survey," *The VLDB Journal*, vol. 29, pp. 595-618, 2020. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [6] Bo-young Lim et al., "Multi-Level Graph Representation Learning Through Predictive Community-Based Partitioning," *Proceedings of the ACM on Management of Data*, vol. 3, no. 1, pp. 1-27, 2025. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [7] Qi Wang, Kenneth H. Lai, and Chunlei Tang, "Solving Combinatorial Optimization Problems Over Graphs with BERT-Based Deep Reinforcement Learning," *Information Sciences*, vol. 619, pp. 930-946, 2023. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [8] Rong Chen et al., "Powerlyra: Differentiated Graph Computation and Partitioning on Skewed Graphs," *ACM Transactions on Parallel Computing*, vol. 5, no. 3, pp. 1-39, 2019. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [9] Zhanzhe Li et al., "MDL: Maximum Density Label-Cut Graph Partitioning," *10th International Conference on Big Data and Information Analytics (BigDIA)*, Chiang Mai, Thailand, pp. 125-132, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [10] Rafael M.S. Siqueira et al., "Graph Partitioning Algorithms: A Comparative Study," *International Conference on Information Technology-New Generations*, pp. 513-520, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [11] Sivakumar Ponnusamy, and Pankaj Gupta, "Scalable Data Partitioning Techniques for Distributed Data Processing in Cloud Environments: A Review," *IEEE Access*, vol. 12, pp. 26735-26746, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [12] Tewodros Alemu Ayall et al., "Graph Computing Systems and Partitioning Techniques: A Survey," *IEEE Access*, vol. 10, pp. 118523-118550, 2022. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [13] Christian Mayer, "Scalable Graph Partitioning for Distributed Graph Processing," Institute for Parallel and Distributed Systems (IPVS), University of Stuttgart, pp. 1-161, 2019. [[Google Scholar](#)] [[Publisher Link](#)]
- [14] Anil Pacaci, and M. Tamer Özsu, "Experimental Analysis of Streaming Algorithms for Graph Partitioning," *Proceedings of the 2019 International Conference on Management of Data*, pp. 1375-1392, 2019. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [15] Charles-Edmond Bichot, and Patrick Siarry, *Graph Partitioning*, John Wiley & Sons, 2013. [[Google Scholar](#)] [[Publisher Link](#)]
- [16] Reinhard Diestel, *Graph Theory*, Springer Berlin, Heidelberg, pp. 1-455, 2025. [[CrossRef](#)] [[Publisher Link](#)]
- [17] Ravikant Diwakar et al., "Optimizing Load Distribution in Big Data Ecosystems: A Comprehensive Survey," *AI and the Revival of Big Data*, pp. 177-200, 2025. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [18] Ümit Çatalyürek et al., "More Recent Advances in (hyper) Graph Partitioning," *ACM Computing Surveys*, vol. 55, no. 12, pp. 1-38, 2023. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [19] J.T. Yan, "Fuzzy-Based Balanced Partitioning Under Capacity and Size-Tolerance Constraints in Distributed Quantum Circuits," *IEEE Transactions on Quantum Engineering*, vol. 4, pp. 1-15, 2023. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [20] Diego Recalde et al., "An Exact Approach for the Balanced K-Way Partitioning Problem with Weight Constraints and its Application to Sports Team Realignment," *Journal of Combinatorial Optimization*, vol. 36, pp. 916-936, 2018. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [21] Lingkai Meng et al., "A Survey of Distributed Graph Algorithms on Massive Graphs," *ACM Computing Surveys*, vol. 57, no. 2, pp. 1-39, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [22] Hamid Hadian, and Mohsen Sharifi, "GT-Scheduler: A Hybrid Graph-Partitioning and Tabu-Search Based Task Scheduler for Distributed Data Stream Processing Systems," *Cluster Computing*, vol. 27, pp. 5815-5832, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [23] Kezhao Huang et al., "WiseGraph: Optimizing GNN with Joint Workload Partition of Graph and Operations," *Proceedings of the Nineteenth European Conference on Computer Systems*, Athens, Greece pp. 1-17, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [24] Christy Alex Panicker, and M. Geetha, "Exploring Graph Partitioning Techniques for GNN Processing on Large Graphs: A Survey," *4th International Conference on Communication, Computing and Industry 6.0 (C2I6)*, Bangalore, India, pp. 1-7, 2023. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [25] Daniela Dapena, Daniel L. Lau, and Gonzalo R. Arce, "Parallel Graph Signal Processing: Sampling and Reconstruction," *IEEE Transactions on Signal and Information Processing Over Networks*, vol. 9, pp. 190-206, 2023. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [26] T. Heuer, "Scalable High-Quality Graph and Hypergraph Partitioning," Doctoral Thesis, Department of Informatics of the Karlsruhe Institute of Technology, pp. 1-254, 2022. [[Google Scholar](#)] [[Publisher Link](#)]
- [27] George Karypis et al., "Recent Trends in Graph Decomposition," Dagstuhl Reports, vol. 13, no. 8, pp. 1-45, 2024. [[Google Scholar](#)] [[Publisher Link](#)]
- [28] Lars Gottesbüren et al., "Deep Multilevel Graph Partitioning," *Arxiv*, pp. 1-19, 2021. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [29] Alan Valejo et al., "A Critical Survey of the Multilevel Method in Complex Networks," *ACM Computing Surveys (CSUR)*, vol. 53, no. 2, pp. 1-35, 2020. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]

- [30] Aman Garg et al., "A Review on Artificial Intelligence-Enabled Mechanical Analysis of 3D Printed and FEM-Modelled Auxetic Metamaterials," *Virtual and Physical Prototyping*, vol. 20, no. 1, pp. 1-50, 2025. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [31] Wan Luan Lee et al., "HyperG: Multilevel GPU-Accelerated K-Way Hypergraph Partitioner," *Proceedings of the 30th Asia and South Pacific Design Automation Conference*, Tokyo, Japan, pp. 1031-1040, 2025. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [32] Sheng Xiang et al., "Scalable Learning-Based Community-Preserving Graph Generation," *IEEE Transactions on Big Data*, pp. 1-14, 2025. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [33] Mohammad Amiriebrahimabadi, Zhina Rouhi, and Najme Mansouri, "A Comprehensive Survey of Multi-Level Thresholding Segmentation Methods for Image Processing," *Archives of Computational Methods in Engineering*, vol. 31, pp. 3647-3697, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [34] Ling Dingcet al., "Survey of Spectral Clustering Based on Graph Theory," *Pattern Recognition*, vol. 51, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [35] Ji Wang et al., "Scalable Spectral Clustering with Group Fairness Constraints," *Proceedings of The 26th International Conference on Artificial Intelligence and Statistics*, pp. 6613-6629, 2023. [[Google Scholar](#)] [[Publisher Link](#)]
- [36] Brahim Laassem et al., "A Spectral Method to Detect Community Structure Based on Coulomb's Matrix," *Social Network Analysis and Mining*, vol. 13, 2023. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [37] Swati A. Bhavsar, Varsha H. Patil, and Aboli H. Patil, "Graph Partitioning and Visualization in Graph Mining: A Survey," *Multimedia Tools and Applications*, vol. 81, pp. 43315-43356, 2022. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [38] Joel Mackenzie, Matthias Petri, and Alistair Moffat, "Tradeoff Options for Bipartite Graph Partitioning," *IEEE Transactions on Knowledge and Data Engineering*, vol. 35, no. 8, pp. 8644-8657, 2022. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [39] Chao-Wei Ou, and S. Ranka, "Parallel Incremental Graph Partitioning," *IEEE Transactions on Parallel and Distributed Systems*, vol. 8, no. 8, pp. 884-896, 1997. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [40] Douglas O. Cardoso et al., "Greedy Recursive Spectral Bisection for Modularity-Bound Hierarchical Divisive Community Detection," *Statistics and Computing*, vol. 34, pp. 1-18, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [41] Teague Tomesh et al., "Quantum Divide and Conquer for Combinatorial Optimization and Distributed Computing," *arXiv*, pp. 1-12, 2021. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [42] Sung-eui Yoon et al., *Real-Time Massive Model Rendering*, Springer International Publishing, pp. 1-112, 2022. [[Google Scholar](#)] [[Publisher Link](#)]
- [43] D. Slavchev, S. Margenov, and I.G. Georgiev "On the Application of Recursive Bisection and Nested Dissection Reorderings for Solving Fractional Diffusion Problems Using HSS Compression," *AIP Conference Proceedings*, vol. 2302, no. 1, 2020. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [44] Rajit Nair, and Amit Bhagat, *An Introduction to Clustering Algorithms in Big Data*, Encyclopedia of Information Science and Technology, 5th ed., pp. 559-576, 2021. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [45] Haixia Wu et al., "Link Prediction on Complex Networks: An Experimental Survey," *Data Science and Engineering*, vol. 7, pp. 253-278, 2022. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [46] Pierre Brémaud, *Probability Theory and Stochastic Processes*, Springer Cham, pp. 1-713, 2020. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [47] Abdul Samad, "Enhancing Community Detection and Data Clustering in Weighted Graphs using Gumbel Softmax-Based Approaches," 2024. [[Google Scholar](#)]
- [48] Sifeng Bi et al., "Stochastic Model Updating with Uncertainty Quantification: An Overview and Tutorial," *Mechanical Systems and Signal Processing*, vol. 204, 2023. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [49] James Alexander Scott, "MCMC Methods: Graph Samplers, Invariance Tests and Epidemic Models," Doctoral Thesis, Imperial College London, pp. 1-186, 2023. [[Google Scholar](#)] [[Publisher Link](#)]
- [50] Zehuan Hu et al., "Self-Learning Dynamic Graph Neural Network with Self-Attention Based on Historical Data and Future Data for Multi-Task Multivariate Residential Air Conditioning Forecasting," *Applied Energy*, vol. 364, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [51] Dilong Li et al., "Graph Neural Networks in Point Clouds: A Survey," *Remote Sensing*, vol. 16, no. 14, pp. 1-44, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [52] Mohamed Massaoudi et al., "Advancing Coherent Power Grid Partitioning: A Review Embracing Machine and Deep Learning," *IEEE Open Access Journal of Power and Energy*, vol. 12, pp. 59-75, 2025. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [53] Jianwu Long, and Luping Liu, "K*-Means: An Efficient Clustering Algorithm with Adaptive Decision Boundaries," *International Journal of Parallel Programming*, vol. 53, 2025. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [54] Harley Wiltzer et al., "Foundations of Multivariate Distributional Reinforcement Learning," *Advances in Neural Information Processing Systems*, vol. 37, pp. 101297-101336, 2025. [[Google Scholar](#)] [[Publisher Link](#)]
- [55] Mahmoud E. Farfoura et al., "A Novel Lightweight Machine Learning Framework for IoT Malware Classification Based on Matrix Block Mean Downsampling," *Ain Shams Engineering Journal*, vol. 16, no. 1, pp. 1-11, 2025. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]

- [56] Mohammad Abdallah et al., "An E-Learning Portal Quality Model: From Al-Zaytoonah University Students' Perspective," *International Conference on Information Technology*, Amman, Jordan, pp. 553-557, 2021. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [57] Ahmad A.A. Alkhatib, and Khalid Mohammad Jaber, "FDPA Internet of Things System for Forest Fire Detection, Prediction, and Behavior Analysis," *IET Wireless Sensor Systems*, vol. 14, no. 3, pp. 47-83, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]