

PAPER • OPEN ACCESS

Key Technologies of Media Big Data in-Depth Analysis System Based on 5G Platform

To cite this article: Qiang Lin and Xilin Zhang 2022 *J. Phys.: Conf. Ser.* **2294** 012007

View the [article online](#) for updates and enhancements.

You may also like

- [Lamb wave propagation modelling and simulation using parallel processing architecture and graphical cards](#)
P Pako, T Bielak, A B Spencer et al.
- [C³. A Command-line Catalog Cross-match Tool for Large Astrophysical Catalogs](#)
Giuseppe Riccio, Massimo Brescia, Stefano Cavuoti et al.
- [Accurate Estimation of Scattering Strength Distribution by Simultaneous Reception of Ultrasonic Echoes with Multichannel Transducer Array](#)
Yusaku Abe, Hideyuki Hasegawa and Hiroshi Kanai



ECS
The
Electrochemical
Society
Advancing solid state &
electrochemical science & technology

DISCOVER
how sustainability
intersects with
electrochemistry & solid
state science research

Key Technologies of Media Big Data in-Depth Analysis System Based on 5G Platform

*Qiang Lin^{1,2}, Xilin Zhang²

¹School of Maritime Economics and Management, Dalian Maritime University, Dalian 116026, China

²Dalian University of Science and Technology, Dalian 116052, China

*Email: lqchina@126.com

Abstract: To meet the needs of large-scale users for personalized streaming media services with high speed, low delay, and high quality in a 5G mobile network environment, this paper studies the resource allocation mechanism of streaming media based on a 5G network from the perspective of user demand prediction, which can alleviate the pressure of mobile network, improve the utilization rate of streaming media resources and the quality of user service experience. The augmented reality visualization of large-scale social media data must rely on the computing power of distributed clusters. This paper constructs a distributed parallel processing framework in a high-performance cluster environment, which adopts a loosely coupled organizational structure. Each module can be combined, called, and expanded arbitrarily under the condition of following a unified interface. In this paper, the algebraic method of parallel computing algorithm is innovatively proposed to describe parallel processing tasks and organize and call large-scale data-parallel processing operators, which effectively supports the business requirements of large-scale parallel processing of large-scale spatial social media data and solves the bottleneck of large-scale spatial social media data-parallel processing.

1. Introduction

5G network communication technology can be stably transmitted in multiple application scenarios, effectively improving the Quality of Experience (QoE) of user services. At the same time, the ultra-low delay of 5G network communication ensures the security of the 5G network communication remote control application. In 5G network, Device-to-Device (D2D) communication technology is a short-range direct data transmission technology based on a cellular system, which enables direct data transmission between mobile terminal devices, greatly alleviating the load pressure of base stations, improving the utilization of spectrum resources and the quality of user service experience [1].

The fundamental goal of the application and development of network communication is always to meet the requirements of users for the continuous improvement of network service quality. Since the 2G era, people's demand for network communication has become more and more diversified, and the demand for network service quality has become higher and higher. At the same time, the network communication technology is constantly improving, and the large-scale application of 5G network communication is bound to meet the different needs of users and further optimize user QoE [2]. Within the coverage of 5G communication network, users can enjoy high-quality network communication services regardless of whether they are located in remote areas or central cities.

Under the background of big data, the challenges brought by data-intensive computing are first



manifested in the source of data. Big data is different from well-designed structured data in relational databases, which is usually stored in the form of unstructured or semi-structured data [3]. First of all, this makes it difficult to use relational database to store this kind of data and also poses a serious challenge to the traditional computing model based on a relational database. On the other hand, traditional high-performance computing technology is usually good at dealing with small but compute-intensive scientific problems, while big data problems are typical data-intensive high-performance computing, which requires a new high-performance computing architecture for big data. Due to the lack of big data storage and processing means, the current use of big data presents a prominent contradiction that a large number of observation data but the ability to analyze and mine new knowledge is extremely scarce [4].

Reference [5] proposes a D2D service solution based on 5G backhaul cross-network architecture to solve the problems of complex resource management and difficult service quality guarantee in 5G mobile network scenarios, which ensures the stability of 5G transmission network and improves the quality of user service experience. To solve the problem of sharing cache resources among mobile network operators, a hierarchical cache architecture is proposed in reference[6], which maximizes the profit of operators through the competitive game method. Yang et al. [7] proposed an efficient hybrid multi-fault location algorithm based on Hopfield neural network to ensure real-time transmission services and improve the quality of user experience for the multi-fault problem faced in the operation of 5G mobile network.

In this paper, a Parallel Computing Algorithmic Algebra Theory (PCAAT) for large-scale spatial social media data parallel processing is proposed, which is used to describe the call relationship of different operators in the parallel environment. In the PCAAT description language system, each parallel processing task is described by an equation, and the relationship between algorithms is described by algebraic operation symbols. In the distributed cluster environment, the ideal processing architecture requires reducing data migration as much as possible, and maintaining flexible and extensible parallel processing capability in certain business processing by deploying and invoking processing operators.

2. A parallel computing framework for spatial social media data

HPGCF (High Performance GeoARMedia Computing Framework) is an algorithm integration framework based on the high-performance cluster environment, which can greatly reduce the complexity of constructing a parallel algorithm. The framework adopts a loosely coupled structure. Each module can evolve independently without affecting the rest under the condition of following a unified interface. The processing operators can be expanded arbitrarily. The data division mode is also a separate library to meet the needs of different algorithms [8].

A multimedia parallel processing algorithm a is abstractly expressed as:

$$a = f(e, p, s, c) \quad (1)$$

Where, e is the computation intensity estimation model. p is the data partitioning method. s is the computation task scheduling method. c is the processing operator, and f is the parallel processing framework.

In this paper, the `hpgc _ Algorithm` base class is constructed to implement a high-performance parallel processing framework. The specific parallel processing operators can inherit the `hpgc _ Algorithm` interface and implement their own methods to expand the processing capacity of the algorithm.

The base class `geo _ algorithm` is the implementation interface class of the serial operator. The serial operator can inherit this class to implement its own processing and give it to the parallel operator for unified scheduling. Task partition scheduling is called by the interface defined by the base class `hpgc _ scheduler`, such as the common master-slave scheduling mode, which is responsible for processing task scheduling under the unified framework of `hpgc _ Algorithm` [9]. Partition of large-scale data is handled by the interface defined by the base class `hpgc _ partition`, and the relevant data partitioning methods inherit the implementation data partitioning of this class. The data

partitioning method mainly refers to the computing intensity of each partition, and its estimation model is managed by the Computing _ Evaluation class and implemented by its base class. The parallel architecture framework is shown in Figure 1.

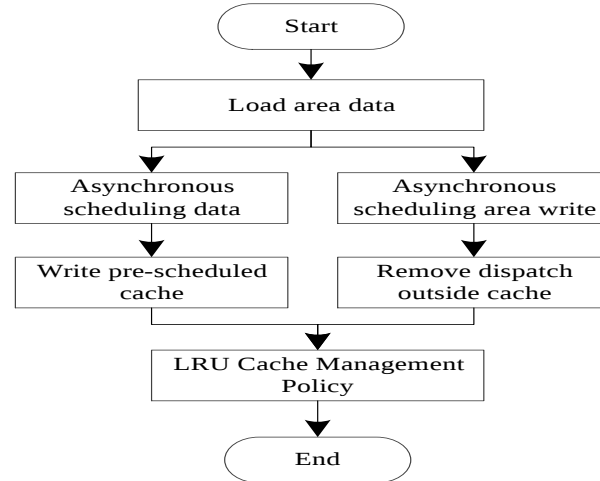


Fig. 1 Parallel computing architecture for spatial social media data

Multi-scale multimedia processing services are mostly data-intensive computing, and data partitioning methods, which plays an important role in the whole parallel processing architecture, even directly determine the processing efficiency and load balancing. If the data is divided according to the number of data, it is easy to cause data skew and uneven load due to the uneven size of each piece of data, so some more accurate data division methods are usually needed [10]. This paper proposes a parallel computing framework that uses a unified Feature Computing Index (FCI) to assess the actual computational intensity of each data partition. One of the simpler ways to define FCI is to measure the calculation intensity by the amount of data. The formula is as follows:

$$F_i = Atr + G_i \quad (2)$$

Where, Atr represents the attribute data of the data entity, and the unit is Bytes; G_i represents the multimedia data entity, such as picture, audio, video, 3D model, etc., and the unit is Bytes; F represents the computational intensity of the entire data entity [11].

Therefore, on the basis of evaluating the calculation intensity of each calculation entity, the global index can be used to quickly perform data partitioning on large-scale data. If the calculation intensity of each data partitioning result is limited to F_0 , there is a data partitioning method:

$$\sum_{i=1}^n f_i = n \times Atr + \sum_{i=1}^n G_i \leq f_0 \quad (3)$$

Where, Atr represents the attribute data of the data entity, and the unit is Bytes; G_i represents the multimedia data entity, such as picture, audio, video, 3D model, and the unit is Bytes; F_0 represents the calculation intensity of the whole data entity; F_o represents the upper limit of calculation intensity of each data division.

Therefore, when the FC/value of each data partition reaches a given upper limit, the HPGCF performs the next data partition, and the data partition process can be faster and more efficient by using the metadata and the database index. After large-scale data partitioning is completed, data tasks are queued for distribution one by one. They have considerable data processing intensity, and each data block is usually processed independently without communication.

3. Experimental analysis

3.1. Experimental data

In this paper, 1 million pieces of social media data are randomly generated as a test data set, including

text, pictures, audio, video and some 3D model data. The size of a single piece of media data is about 10 MB. The data is distributed in 1MB ~ 200MB, totaling about 10 TB. In order to verify the effectiveness of HPGCF in processing large-scale spatial social media data, according to the spatial social media data, it is stored on the distributed NoSQL cluster data server to test the parallel processing of images that meet the query conditions, and to observe the response efficiency of large-scale data processing.

The operation of compressing to the half size of the original image is operator A. The operator of adding watermark is operator B, and the test data set is D. Then, according to the PCAAT rule, a parallel processing process can be described as follows:

$$C = B^n + 2A^n = B^n + A^n + A^n = (B + 2A)^n \quad (4)$$

Where n represents the number of parallel computing processes. The number 2 represents that the compression operation A is performed twice. The data set D may be divided into m portions, which are assigned by the management process to n processes for processing. Each process first performs the watermarking operation B, and then performs the compression operation A twice.

The test program runs on the Lustre parallel computer cluster (1 MDT, 9 OSTs, each node mainly consists of 8-core Intel (R) Xeon (R) CPU E5-2603, 16GB memory, and 8TB Lustre external memory can be shared). This paper tests the processing efficiency of the data under serial algorithm, single-node parallel (multi-proc) processing and multi-node parallel (multi-nodes) processing. This operation involves about 20 GB of images. The results are shown in Figure 2.

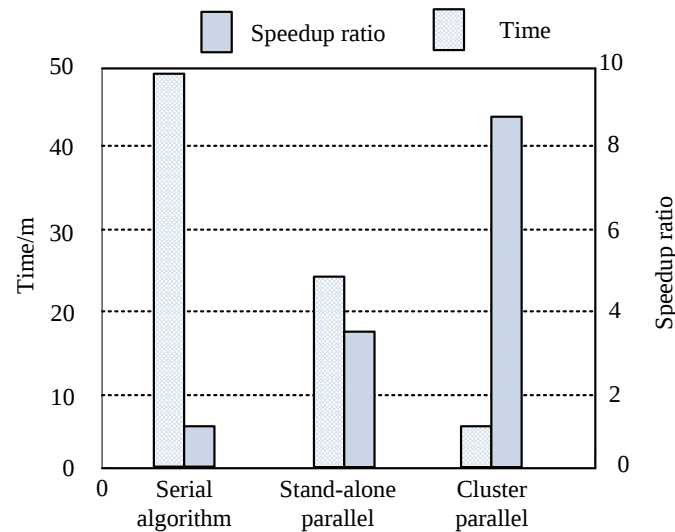


Fig. 2 Histogram of parallel processing efficiency test

3.2. Test analysis

In order to prove the good performance of the proposed algorithm in the paper, it is compared with Random Queue Algorithm (QQA) and Longest Queue Algorithm (LQA) from the aspects of system stability and system energy consumption. The main parameters are shown in Table 1.

Table 1 Main parameters of simulation

Parameter	Definition	Numerical value
num	Queue	15
θ	Trade-off factor	1
B	Edge calculation data package maximum	50
E	Unit energy consumption	15

The unit energy consumption represents the average energy consumption for transmitting one packet. It is assumed that the time average energy consumption of the edge cloud is $\overline{W_j}$, and $\overline{W_j} \leq W_j^{av}$ is satisfied. W_j^{av} is a continuous non-negative number, which represents the average energy consumption for transmitting a data packet; W is the core network; j is the edge cloud; av is the number of requests.

The processing time of a request is set as $T1 = T2 + T3 + T4$. If the queue is long and the request is added to the queue, the consumer cannot process it immediately. From the perspective of production, the request processing time is $T1 = T2 + T3 + T4 + WaitTime$.

3.2.1. Stability study

Figure 3 shows the system stability of the three algorithms.

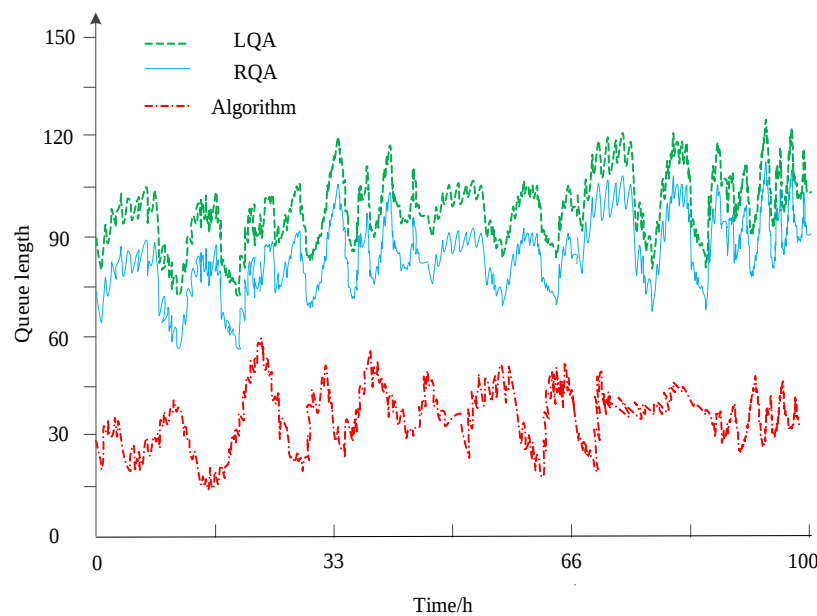


Fig. 3 Stability comparison results

Compared with the other two algorithms, the proposed algorithm in the paper can achieve the shortest queue length. With the increase of time, the queue length obtained by the algorithm in this paper is stable around 40, and the queue length is the shortest, that is, the algorithm in this paper has better system stability. This is because the algorithm in this paper considers the access control and transmission control in the process of data packet transmission. In addition, from the previous analysis, it can be seen that if the value of the queue length is small, the network congestion generated during data transmission will be reduced; on the contrary, the LQA selects the longest queue for data transmission, so it is most likely to cause network congestion.

3.2.2. System energy consumption

Figure 4 reflects the system consumption corresponding to the algorithm in this paper, QQA and LQA.

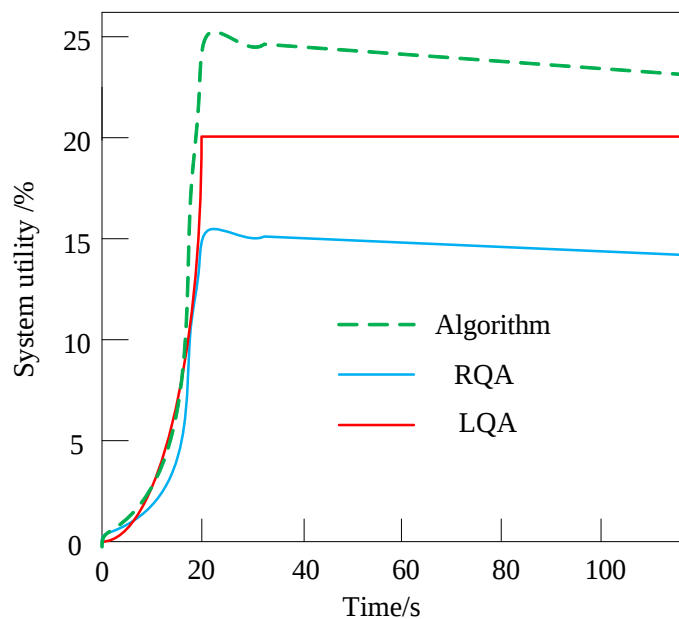


Figure 4 Comparison results of system energy consumption

The system energy consumption obtained by using the resource allocation strategy of the algorithm in this paper is always the minimum. On the contrary, LQA selects the longest queue for data transmission, so it generates the largest system energy consumption. The system energy consumption obtained by RQA is between that of the algorithm in this paper and that of LQA.

4. Conclusion

This paper constructs a new distributed data partition method based on a stable curve, which can effectively use the advantages of a NoSQL distributed database. It also makes the coding structure more compact and convenient for operation and supports smooth cross-scale reading and writing of database. The steady-state curve data partitioning method takes into account the special effect of spatial aggregation and the load balance of distributed data storage, effectively avoids the tilt of node data in the distributed cluster environment, optimizes the physical distribution of large-scale social media data in the distributed environment, and effectively supports the business requirements of rapid access to large-scale data.

In the future, the dynamic data tracking method focusing on deep learning is expected to break through the boundary of SFIT algorithm. The method is also expected to enable the augmented reality target recognition to play a role in the spatial social media data. Thus, the application scope of augmented reality technology and the stability of on-the-spot experience will be broadened.

Acknowledgement

The study was supported by the Liaoning Provincial Department of Education. The project number is L2020004.

References

- [1] Gartner Identifies Five Emerging Technology Trends That Will Blur the Lines Between Human and Machine [EB/OL]. 2018.
<https://electroi.com/gartner-identifies-five-emerging-technology-trends-that-will-blur-the-lines-between-human-and-machine/>.
- [2] Bradski, G.R., S.A. Miller, R. Abovitz, Methods and systems for creating virtual and augmented reality. 2019, Google Patents.

- [3] S. Chen, B. Yang, J. Yang and L. Hanzo, Dynamic Resource Allocation for Scalable Video Multirate Multicast Over Wireless Networks[J], IEEE Transactions on Vehicular Technology, Sept.2020, vol. 69, no. 9, pp. 10227-10241.
- [4] J. Yang, B. Yang, S. Chen, Y. Zhang, Y. Zhang and L. Hanzo, Dynamic Resource Allocation for Streaming Scalable Videos in SDN-Aided Dense Small-Cell Networks[J], IEEE Transactions on Communications, March 2019, vol. 67, no. 3, pp. 2114-2129.
- [5] A. Yousafzai, I. Yaqoob, M. Imran, A. Gani and R. Md Noor, Process Migration-Based Computational Offloading Framework for IoT-Supported Mobile Edge/Cloud Computing[J], IEEE Internet of Things Journal, May 2020, vol. 7, no. 5, pp. 4171-4182.
- [6] J. Zeng, J. Sun, B. Wu and X. Su, Mobile edge communications, computing, and caching(MEC3) technology in the maritime communication network[J], China Communications, May 2020, vol. 17, no. 5, pp. 223-234.
- [7] J. Ren, G. Yu, Y. He and G. Y. Li, Collaborative Cloud and Edge Computing for Latency Minimization[J], IEEE Transactions on Vehicular Technology, May 2019, vol. 68, no. 5, pp.5031-5044.
- [8] L. Wang, L. Jiao, J. Li, J. Gedeon and M. Mhlh-user, MOERA: Mobility-Agnostic Online Resource Allocation for Edge Computing[J], IEEE Transactions on Mobile Computing, Aug. 2019, vol. 18, no. 8, pp. 1843-1856.
- [9] Zhi Li, Qi Zhu. Genetic Algorithm-Based Optimization of Offloading and Resource Allocation in Mobile-Edge Computing. 2020, 11(2).
- [10] Z. Yan, M. Zhao, C. Westphal and C. W. Chen, Toward Guaranteed Video Experience: Service-Aware Downlink Resource Allocation in Mobile Edge Networks[J], IEEE Transactions on Circuits and Systems for Video Technology, June 2019, vol. 29, no. 6, pp. 1819-1831.
- [11] M. Yan, G. Feng, J. Zhou, Y. Sun and Y. -C. Liang, Intelligent Resource Scheduling for 5G Radio Access Network Slicing[J], IEEE Transactions on Vehicular Technology, Aug. 2019, vol. 68, no. 8, pp. 7691-7703.