

An Optimized IoT-enabled Big Data Analytics Architecture for Edge-Cloud Computing

Muhammad Babar, Mian Ahmad Jan*, Xiangjian He*, Muhammad Usman Tariq, Spyridon Mastorakis, Ryan Alturki

Abstract—The awareness of edge computing is attaining eminence and is largely acknowledged with the rise of Internet of Things (IoT). Edge-enabled solutions offer efficient computing and control at the network edge to resolve the scalability and latency-related concerns. Though, it comes to be challenging for edge computing to tackle diverse applications of IoT as they produce massive heterogeneous data. The IoT-enabled frameworks for Big Data analytics face numerous challenges in their existing structural design, for instance, the high volume of data storage and processing, data heterogeneity, and processing time among others. Moreover, the existing proposals lack effective parallel data loading and robust mechanisms for handling communication overhead. To address these challenges, we propose an optimized IoT-enabled big data analytics architecture for edge-cloud computing using machine learning. In the proposed scheme, an edge intelligence module is introduced to process and store the big data efficiently at the edges of the network with the integration of cloud technology. The proposed scheme is composed of two layers: IoT-edge and Cloud-processing. The data injection and storage is carried out with an optimized MapReduce parallel algorithm. Optimized Yet Another Resource Negotiator (YARN) is used for efficiently managing the cluster. The proposed data design is experimentally simulated with an authentic dataset using Apache Spark. The comparative analysis is decorated with existing proposals and traditional mechanisms. The results justify the efficiency of our proposed work.

Index Terms—Internet of Things (IoT), Big Data Analytics, Edge Computing, Backpropagation (BP) Neural Network, Machine Learning, Yet Another Resource Negotiator (YARN)

I. INTRODUCTION

THE Internet of Things (IoT) is consist of software utilities along with hardware. The hardware involves of sensor-embedded devices connected with each other. The software utilities, on the other hand, comprise data storage and analytics programs that provide assistance in presenting information to the users [1]. The idea of edge computing is gaining significant attention with the rise of IoT [2]–[4]. Edge-enabled solutions

offer efficient computing and provide the control near the edge of the network to resolve the scalability and latency issues [5], [6]. The IoT applications generate massive data, i.e., big data¹, which require effective aggregation, intelligent processing, and robust analysis to realize optimum outcomes for decision-making [7], [8]. In an IoT-enabled edge computing environment, the data is growing rapidly, and the computation relies heavily on this massive data. Moreover, the fast speed at which the data is generated make it become infeasible to stockpile the data into any explicit server. It is a challenging task to hold a gigantic amount of big data, which continues to raise exponentially at an extraordinary speed [9], [10]. In IoT-enabled edge computing environments, major sources of Big Data are the embedded sensors [11], [12].

With the rise of Big Data in IoT-enabled edge environments, the notion of Machine Learning (ML) model development is gaining prominence [13]–[15]. Moreover, the federated learning concept is also introduced at the edges of the network [16]. ML-based models deliver enriched performance in the market by utilizing the huge data [17], [18]. In this context, low-cost computing devices, sensors, and storing mechanisms integrated with convincingly speedy wireless equipment provide the foundation for managing the big data [19], [20]. However, an accessible substructure is still required to deal with the massive set of devices. Big Data analytics has the ability of providing massive opportunities to resolve problems related to various applications dealing with plethora of data [21], [22]. The future of this universe lies with ML-enabled IoT paradigms that aim to bring intelligence to the real-world physical objects [23]. Various IoT applications seem to be progressively integrated with ML techniques.

For edge computing, it becomes challenging to manage the diverse IoT applications as they produce massive amount of big data. Edge computing is the future of big data analytics. The domain of big data has undertaken a huge revolution over the decade. This revolution is creating a shift in computing infrastructure; where the computing is moving out of the data center to the edge. The architectures of IoT-enabled Big Data analytics face numerous challenges, such as the storage and processing of high volume of data, latency, data heterogeneity, inconsistent data patterns, and so forth, in their existing structural design. Moreover, the existing proposals lack an efficient parallel data ingestion and incur a high communication overhead [24]. The major challenges and problems that need to be undertaken are inadequate data ingestion into the traditional cluster supervision frameworks, e.g. Apache Hadoop/Spark.

Manuscript submitted on October 29, 2021 and asterisks indicate the corresponding authors.

Muhammad Babar is with the Department of Computer Science, Allama Iqbal Open University (AIQU), Islamabad, Pakistan.

Mian Ahmad Jan is with the Department of Computer Science, Abdul Wali Khan University Mardan, Pakistan (E-mail: mianjan@awkum.edu.pk).

Xiangjian He is with the Global Big Data Technologies Center (GBDTC), School of Electrical and Data Engineering, University of Technology Sydney, Australia. (E-mail: xiangjian.he@uts.edu.au).

Muhammad Usman is with the Abu Dhabi School of Management, Abu Dhabi, UAE.

Spyridon Mastorakis is with the College of Information Science & Technology, University of Nebraska Omaha, USA

Ryan Alturki is with the Department of Information Science, College of Computer and Information Systems, Umm Al-Qura University, Makkah, Saudi Arabia.

Correspondence: mianjan@awkum.edu.pk

¹Big Data and big data are used interchangeably

The customary data ingestion is time wasting, requires more storage, contains commands that are difficult to understand, has no appendage, and has no partial ingestion [25]. Similarly, the processing of traditional big data analytics platforms based on conventional cluster management face challenges such as difficult scheduling, inefficient load balancing, non-scalability, and responsibility unification.

This research aims to design a specific framework using edge computing to evaluate the huge data proficiently and overcome the data ingestion and computation issues. We adopt the optimization of YARN framework to customize for edge computation. The major contributions of this paper include:

- A framework for data analytics at the edge including parameter selection for modules affecting the distributed files for edge data availability, efficient pre-processing to prepare the data for quick processing and accurate prediction, and efficient technique to concurrently store the data at the edges.
- An algorithm for parallel data ingestion, an edge-map-reduce algorithm for parallel processing at every edge, an edge-yarn algorithm for efficient cluster resource management by optimizing the YARN, and an enhanced weighted BP algorithm for model parallelism.
- Simulation of the proposed scheme and its algorithms performed using the open-source Apache Spark 3.0 framework. We preferred the open source platform to achieve the cross platform applicability.

The rest of this paper is organized as follows. In Section II, we provide a comprehensive literature review. In Section III, we present our proposed IoT-enabled big data analytics framework. In Section IV, the experimental results and analysis are provided to validate our framework. Finally, we conclude this article in Section V.

II. LITERATURE REVIEW

Superfast network connectivity and novel frameworks of IoT are just some of the things transforming applications around the globe. This sort of transformation does not happen in isolation and overnight. It is easy to lose the benefits of IoT with the presence of Big Data, but connectivity at the network edges offer alternative transmission options to the end users. Nevertheless, there are challenges faced by Big Data and IoT in the edge computing paradigm related to interoperability issues, heterogeneity problems, data preparation and cleaning, data format challenges, normalization of data, data filtration, and loading data into processing frameworks. Researchers have been working to tackle the challenges of big data in complex environments for over a decade now. A summary of existing proposals for big data processing of IoT is provided in Table I. The related works are tabulated along with the key challenges. A specific solution is suggested to process the Big Data in a healthcare domain that contains the separation of dynamic data into subgroups, which are based on the hypothetical data fusion working in Hadoop to improve computational effectiveness [26]. The key challenges identified for the proposed scheme are the use of customary Map Reduce (MR) mechanism, insufficient data loading, and

TABLE I
STATE-OF-THE-ART MODELS FOR IOT DATA PROCESSING

Proposals	Challenges
[26]	• Used Traditional Cluster management • Data loading efficiency is ignored in the architecture • Only health dataset is evaluated
[27]	• Causes delay in processing • Data accumulation prior to data ingestion is highlighted while data ingestion proficiency is overlook • Classical map-reduce framework used
[28]	• Computation at edge is overlooked during processing
[29]	• The data ingestion proficiency is overlooked and the utilization of MRV1 framework • Limited architecture
[31]	• Conceptual and logical framework • Implementation is not performed
[32]	• It causes additional delay This scheme is also insufficient to load data efficiently to the Hadoop • Conventional cluster resource handling • Only transportation dataset are evaluated
[33]	• Customary and conventional cluster resource handling • It fails to load data efficiently
[34]	• Data loading mechanism exist but delay in loading • Another issue is the utilization of classical map-reduce framework
[35]	• This scheme is also insufficient to load data to the Hadoop • Traditional cluster resource handling is utilized in this scheme
[36]	• Data collection is preferred and data ingestion prior to analysis is overlooked • Only health dataset is evaluated to test the proposed system

impalpable structure. The data loading efficiency is ignored in the architecture and only the health dataset is evaluated. A scheme is proposed for data management, which is composed of different tiers [27]. Tiers are liable for different steps of data processing. Nevertheless, it is a comprehensive design that consists of four tiers from big data gathering to analysis of big data, it also causes computation delays, and the use of a legacy MR framework slows down the performance. Moreover, the focus is mainly on data aggregation before data ingestion, while data loading competence is overlooked [28].

In [29], the authors proposed an IoT framework using Hadoop-based Big Data analytics with different layers from data acquisition to application. The major issue in the design of this architecture is that data loading effectiveness was overlooked. Some researchers proposed a model based on data analytics that encourages the concept of smart and connected societies (SCC), which is developing from the notion of smart cities [30]. The SCC model is a theoretical structure that is not yet implemented for the physical world. Likewise, the authors in [31] presented a framework, however, this is also a theoretical prototype. These researchers also overlooked the data loading process in the context of a parallel and distributed domain. There are few designs recommended to tackle the challenges of Big Data analytics within a smart infrastructure [32]–[34], but these solutions utilized the customary cluster management mechanisms and unsatisfactory data ingestion to Apache Hadoop. Furthermore, a design was proposed to explore the data in an accessible and efficient transport environment [35]. However, the proposed work causes additional processing delay. Moreover, the proposed work was only

verified using transport data, and data loading competence was ignored during Big Data loading to Hadoop. A proposal was presented that promotes the data aggregation of results but the data loading before analysis was overlooked [36]. These studies conclude with the following challenges that need careful consideration.

- Effective data loading into Hadoop Distributed File System (HDFS) for edge data availability is missing.
- Lack of efficient and improved method to load and store the datasets at the edge.
- Lack of efficient cluster resource management.
- Lack of efficient processing for edge devices while taking into consideration the parallel and distributed paradigm.

Architectures of Big Data analytics for complex environments e.g. edge-cloud, is a novel initiative to switch the traditional form of services and facilities to a well-designed mode [37]–[41]. This changeover seems to be challenging as there are no predefined instructions available to construct a robust architecture for Big Data computation. IoT offers services using sensors and other objects that supply a huge quantity of data for analysis. Big Data analytics using Hadoop involves several varied phases including data loading and data processing that bring issues and challenges. Although, numerous research works exist to analyze the Big Data efficiently and measure its performance, studies dealing with data loading efficiency are exceedingly rare. The literature suggests that several proposals have been presented to analyze the Big Data in smart city environment using the Hadoop framework, but significant challenges exist that need to be dealt with, e.g. efficient data loading and processing using an effective cluster management technique. Therefore, we discover the need for a resourceful model of Big Data analytics and proposes an improved architecture for this purpose.

III. AN IOT-ENABLED BIG DATA ANALYTICS FOR EDGE-CLOUD

In this section, we design a big data processing scheme using machine learning models. The machine learning models are trained using the huge big data produced within the IoT-Edge environment.

A. System Model

The system model of the proposed IoT-enabled Big Data analytics framework is discussed in this section. The model is composed of two different layers: IoT-Edge layer, and Cloud-processing layer as depicted in Fig. 1.

The first layer is comprised of IoT sensors and embedded devices integrated with edge servers. The second layer is composed of cloud and processing units that perform various activities, e.g. big data aggregation, distributed storage of big data, and big data processing with model training. The cloud storage is responsible for holding various sources of data. The data is produced by various IoT sensors such as environmental observing, security monitoring [42], facility monitoring, traffic, power observing, and transportation observing sensors. The generated data is received by edge devices

and servers. The edge data is properly gathered by various sophisticated processes known as edge caching. Finally, the edge data is collected by cloud that performs distributed storage. The gathered data is feed into a distributed architecture for parallel processing. It involves the overall data management that includes collection, aggregation, and storage. The data aggregation offers the grouped data with a precise shaping for further analysis, which is useful for enormous data that lacks valued information. The preferred technique for aggregation is divide-and-conquer that divides the huge dataset into smaller chunks and groups the similar data in the form of blocks.

The preprocessing mechanisms, e.g. normalization, filtration, and queuing, are utilized to prepare the data for efficient processing and training. The normalization of data is performed using a min-max normalization method as depicted in Algorithm 1. The filtration is performed using an optimized Kalman Filter (KF) and the queuing is carried out using the hybrid M/M/1 queuing model. Algorithm 2 depicts the working of filtration process to speed the actual processing. KF provides only the quality information to be processed. It filters the inferiority information that affects the quality. Moreover, KF is ideal for real-time processing that can easily be integrated with the Apache Spark platform and avoids heterogeneity. The message queue is also utilized to speed up the data processing. The message queue functions in a particular operational mode to obtain the message M at time t and then forwards it to the resultant component. It is controlled by H (a specific handler). This enhances the efficient exploitation of big data for better results. The data is loaded to the processing server with parallel loading using a map-only parallel algorithm. It speeds up the overall computation time. The number of parallel channels is kept balanced by optimizing the block size of processing unit. The big data processing is performed by dividing the big datasets into small chunks of fixed block sizes. The blocks are processed by each node in parallel. The default block size of Hadoop (the rationale of Apache Spark) takes more time to be processed and provides less number of parallelism. The default size is optimized and modified to enhance the data ingestion efficacy. The default block size in the Hadoop latest version is 128MB. The default size creates too many replicas and communication overhead. Therefore, the default block size is tuned using the block-size parameter.

Algorithm 1 Data Normalization.

BEGIN

Input: = Each value of data block

Max value of corresponding data block

Min value of corresponding data block

Output: = scaled value

Set upper and lower limits (m , n)

// specific range

Identify Max and Min values (X_{max} , X_{min})

For each (data item) do

$Y = \frac{(X - X_{min})}{(X_{max} - X_{min})} * (n - m) + m$

END

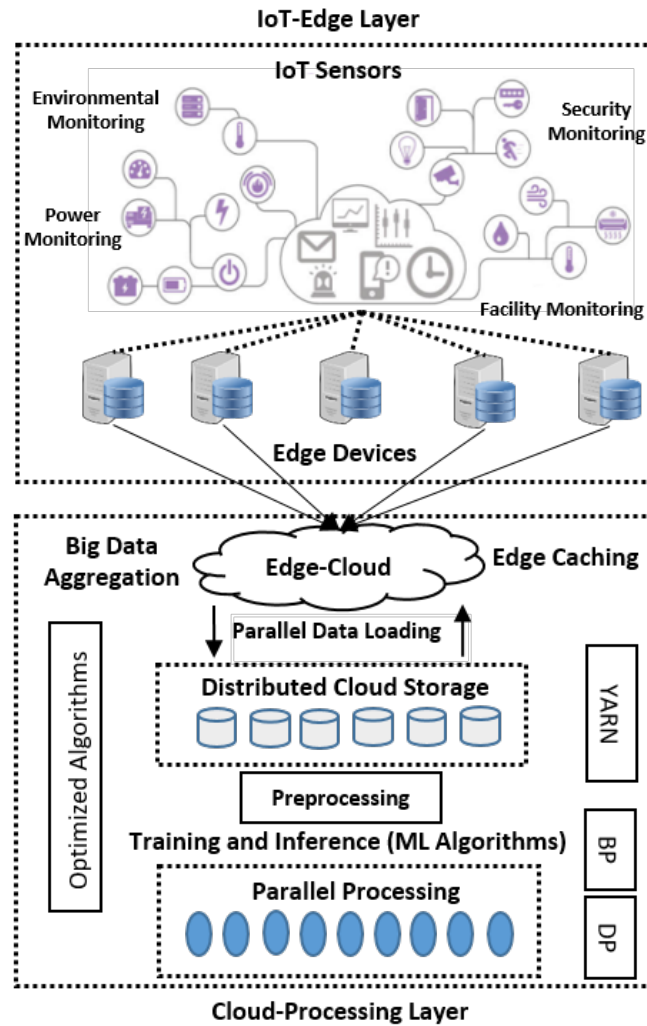


Fig. 1. System Model of our IoT-enabled Data Analytics Framework

Algorithm 2 Filtration using KF.

```

BEGIN
  Instantiation
    instantiate transition model ( $T$ )
    instantiate observation model ( $O$ )
    instantiate noise co-variance ( $CN$ )
    instantiate OBS co-variance ( $CO$ )
    instantiate control input ( $CI$ )
  Calculating first-hand DataSplit ( $D$ )
    Search previous DataSplit ( $DI$ )
    Search new DataSplit
  Approximating standing DataSplit
    DataSplit prediction ( $DP$ )
    co-variance prediction ( $CP$ )
  Merging present observation with prediction
    Discover present observation ( $X$ )
    Discover co-variance of observation ( $S$ )
    Discover optimal gain ( $K$ )
    Update DataSplit ( $PD$ )
  Update co-variance
    Prediction
    Observation
END
    
```

The processing and computation of big data are performed using the parallel and distributed framework of big data analytics. The proposed architecture is realized using the Apache Spark 3.0 framework. The traditional Spark is customized using an optimized YARN (Yet Another Resource Negotiator) cluster manager. The execution time of the overall big data computation using the proposed design is compared with the state-of-the-art existing related works. The parallel processing is equipped with a machine-learning algorithm of backpropagation (BP) neural network. This module includes the training and inference of BP neural network's machine learning model. The training of the model is performed with 70:30 ratio of the dataset. This module provides the customization of BP algorithm integrated with the proposed architecture. The novelty of BP network is the model parallelism. The proposed BP network is trained and validated in parallel with several processing nodes. The partial results produced by the nodes are merged using the ensemble learning for better results and improvement.

B. Hybrid M/M/1 Queue

To speed up the data processing for efficient exploitation of big data, we optimized the M/M/1 queuing model. This model performs various operations when it receives the Data Split D at time t . At this time, a system is thought to be in the steady state. Fig. 2 and Fig. 3 demonstrate the data processing and transmission of an M/M/I queue.

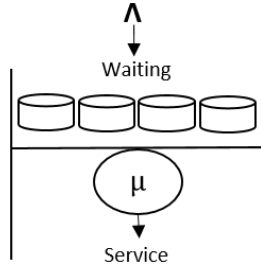


Fig. 2. M/M/1 Queuing Process.

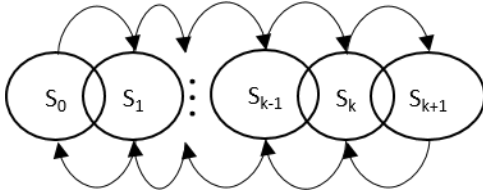


Fig. 3. M/M/1 Queue Computation.

The two splits, e.g. S_1 and S_2 , are balanced during the steady state. Here, $S_2, S_3, \dots, S_{k-1}, S_k, S_{k+1}$, can be calculated as:

$$\Lambda S_0 = \mu S_1. \quad (1)$$

$$\Lambda S_1 = \mu S_2. \quad (2)$$

$$\Lambda S_{k-1} = \mu S_k. \quad (3)$$

hence,

$$S_n = (\Lambda/\mu)^k S_{k-1} = (\Lambda/\mu)^k S_0 \quad (4)$$

As the probability is equal to 1,

$$\sum_{j=0}^{\infty} S_j = 1 \quad (5)$$

$$S_0 \left(1 + (\Lambda/\mu)^1 + (\Lambda/\mu)^2 + \dots \right) = 1 \quad (6)$$

Summing the series,

$$S = 1 - \Lambda/\mu \quad (7)$$

As, $S=1-S_0$, thus

$$S = 1 - (1 - \Lambda/\mu) \quad (8)$$

$$S = \Lambda/\mu \quad (9)$$

$$S_k = (\Lambda/\mu)^k S_0 \quad (10)$$

$$(S)_n = (1 - S) \quad (11)$$

$$S_n = S R^k (1 - S) \quad (12)$$

$$MS = \sum_{m=0}^{\infty} m \cdot S_j \quad (13)$$

$$MS = \sum_{m=0}^{\infty} m \cdot S^n (1 - S) \quad (14)$$

$$MS = (1 - S) S \frac{d}{ds} \sum_{m=0}^{\infty} S^m \quad (15)$$

$$MS = (1 - S) S \frac{d}{ds} \left(\frac{1}{1-S} \right) = \frac{1}{1-S} \quad (16)$$

The average is,

$$T = \frac{N}{\lambda} = \frac{S}{1-S} \frac{1}{\Lambda} = \frac{1}{\mu - \Lambda} \quad (17)$$

As we know $MQ = M \cdot \text{Tasks}$,

$$MQ = \left(\frac{S}{1-S} \right) \times S = \frac{S^2}{1-S} = \frac{\Lambda^2}{(\mu - \Lambda)^2} \quad (18)$$

Hence, the waiting time is

$$W_Q = T - \frac{1}{\mu} = \frac{1}{\mu - \Lambda} - \frac{1}{\mu} = \frac{\Lambda}{(\mu - \Lambda)^2} \quad (19)$$

Finally, the average number of tasks in a system becomes

$$M = \frac{S}{1-S} = S + \frac{S^2}{1-S} \quad (20)$$

C. Map-Only Algorithm for Parallel Data Ingestion

The Apache Hadoop and Apache Spark process the data available inside the Hadoop in HDFS format. The data ingestion is the key stage for both batch and stream processing. The data can either be loaded in batch or stream form. The data loading is dependent on the kind of processing available inside the parallel and distributed platform. The data must be loaded to the parallel processing platform before processing. The traditional mechanism of data loading is sequential and time-consuming. Therefore, the Map-Reduce programming paradigm is optimized for parallel data loading. The data loading also improves the overall data processing. The data analysis within the big data platforms, e.g. Apache Hadoop and Apache Spark, include the loading of data. The big data analytics platforms carry the processing of data by loading it in the required format. As a result, the data loading efficiency improves the overall data processing. We integrate the Sqoop utility with the map job of Map-Reduce paradigm

and introduced a map-only algorithm. In addition, the split size and replication factor of the traditional approach are also tuned with the optimized map-only algorithm for parallel data ingestion. The split size is defined using the unit byte. The replication based on the file can also be modified accordingly. First, we need to create a new directory at the root. Next, we need to verify and add a file to the directory. Afterward, a command is executed to change the replication. Similarly, the replication of all files can be changed within a directory. For this purpose, specific and generic arguments can be used to control the operations of Sqoop tool.

The general command-line parameters of Hadoop must precede any tool-specific parameters. This specific tool of Sqoop utility is utilized in the proposed architecture to import a particular table (source) to HDFS from an RDBMS. The HDFS format for processing is mandatory as the processing is carried out in the HDFS format by Hadoop and Apache Spark in the proposed framework. Every table line is symbolized as a separate (detach) record that can be saved as text files in HDFS or it can also be saved in binary form as Avro or Sequence Files. The `-table` argument is used to choose the table that is about to be loaded. A specific subset of the columns can also be selected and can be controlled using `-columns` argument if required. Similarly, specific rows can also be imported using the SQL WHERE clause. SQL WHERE clause is a very useful tool that can enhance the utilization of Sqoop commands. One of the most beneficial quality of Sqoop tool is the incremental import. Sqoop offers `-append` argument that is used for loading in an incremental import manner, which can be utilized to get only newer records or data.

D. Map-Reduce for Parallel Data Processing

The parallel processing in the edge environment is achieved by the Map-Reduce mechanism. It provides a parallel mechanism for data processing. Big data is huge and need to be divided into blocks or splits for distributed storage and parallel processing. Therefore, the map-reduce paradigm is preferred. The map-reduce programming paradigm is preferred as it is more scalable, flexible, cost-effective, fast, simple, and resilient. The process of mapping and reducing is depicted in Fig. 4. The traditional MapReduce paradigm is optimized for edge data processing by introducing the edge-map and edge-reduce phases. This edge-based MapReduce is introduced to realize the edge-based big data processing. The traditional cluster management is also replaced with optimized edge-YARN. The edge-YARN performs only the resource management inside the cluster along with other operations in the cluster. It is responsible for the overall cluster management. YARN is optimized for the edges of the system. The previous forms of distributed podiums, e.g. customary Hadoop by Apache, utilized the MR for computation and the management of clusters. Job dissemination in the YARN distributed cluster management framework is performed with optimized and tuned parameters. Earlier platforms, e.g. the previous versions of traditional Hadoop and map-reduce were used for cluster management along with parallel processing. However, they generated higher communication overhead and reduces the

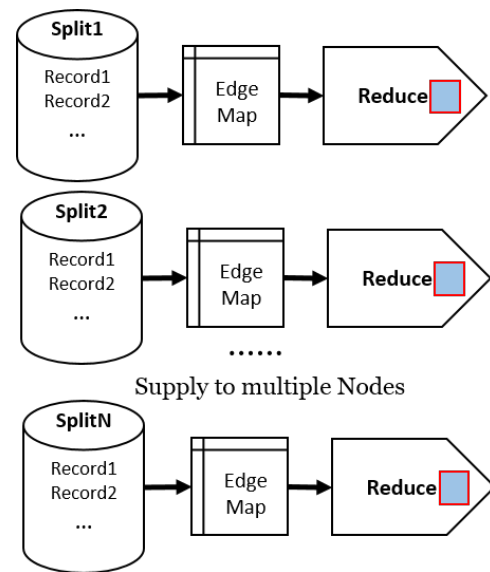


Fig. 4. MapReduce operation for parallel processing.

overall performance. To overcome these challenges, YARN is favored and optimized as it achieves the separate functionalities for cluster management. We considered three major parameters: number of replicas, block size, and the level of parallelism. The big data is processed using the parallel and distributed platform. Here, the data is distributed into several blocks (distributions) and each block is inserted in parallel fashion. It is mandatory to choose the amount of parallelism, i.e., how many nodes (blocks) are required to be inserted for processing and computation. This concept is known as level of parallelism. The YARN framework is furnished with dynamic programming (DP) for job distribution and managing the cluster. DP is used for the recursion function. The enormous dataset is repeatedly distributed into various blocks. In this context, DP is preferred as it is an ideal algorithm for recursive problems by breaking the problem into simpler sub-problems to produce an optimal solution. YARN operation in an edge computing paradigm is reflected in Fig. 5. The MRAppMaster delivers the implementation of Application Master (AM) in the YARN-enabled framework. Besides, interweaving is probable among map and reduce stages; hence, the reduce stage may start prior to the mapping stage ends.

The Resource Manager (RM) allocates a container (CON) when required by the Application Master (AM). Any specific request is entertained by the CON to exploit the required resources on a particular host. AM requests the Node Manager (NM) accountable for controlling the host (holding the CON) to take off the dedicated task of an application. This task could be any Map Reduce (MR) task. NM only assesses the resource usage and destroys any resource that utilizes extra memory than originally assigned. AM pays out its entire existence by negotiating various CONs to initiate request(s) to complete its application. AM is also accountable for restarting the job in new containers (CONs) if the previous one fails. It provides the progress back to the user. AM ends the job and relieves its CON on completion of the job. However, RM does not

$$\Lambda(n+1) = \begin{cases} m_{++} \Lambda(n), & X(n+1) < X(n) \\ m_{--} \Lambda(n), & X(n+1) > X(n) \\ \Lambda(n) & \end{cases} \quad (22)$$

where, the incremental factor m_{++} is greater than 1 and the decremental factor m_{--} is between 0 and 1. Here, $X(n+1)$ and $X(n)$ denotes the total error sum of squares after the $(n+1)^{th}$ and n^{th} iteration. Finally, the Δ symbolizes the learning rate. The philosophy managing the learning directions of BP is the adjustment to the weight and threshold of network should be performed with the direction of gradient.

$$S_{i+1} = S_i - \Delta_i g_i \quad (23)$$

where, S_i denotes the matrix of existing weight TLV, g_i denotes the gradient of current operation, and Δ_i denotes the learning rate. Suppose the three-layer BP model with input node is y_j , the node of hidden layer is x_j , and the node of output layer is z_i , then we can have,

$$x_j = f \sum_{m=0}^{\infty} w_{jm} y_j - \theta_j = f(NE T_j) \quad (24)$$

where,

$$NE T_j = \sum_{m=0}^{\infty} w_{jm} y_j - \theta_j \quad (25)$$

The computational output of the output node is

$$z_j = f \sum_{m=0}^{\infty} v_{jm} y_j - \theta_j = f(NE T_j) \quad (26)$$

The output node's error is

$$ERROR = \frac{1}{2} \sum_{m=0}^{\infty} (t_m - z_m)^2 \quad (27)$$

$$ERROR = \frac{1}{2} \sum_{m=0}^{\infty} (t_m - F)^2 \quad (28)$$

where,

$$F = \left(\sum_{m=0}^{\infty} v_{jm} y_j - \theta_j \right)^2 \quad (29)$$

IV. EXPERIMENTAL RESULTS AND DISCUSSION

In this section, the experimental results are discussed in detail. The proposed scheme is implemented using the Apache Spark 3.0 parallel and distributed framework. Spark is customized using an optimized YARN cluster management approach. The Spark MLlib library is utilized for the execution of machine learning neural network model. The MLlib library is provided by Spark. We used the MLlib for ML model execution with a 70-30 ratio of training and testing. We can also split the dataset for training and testing using 80-20 ratio, but most of the research studies follow 70-30 ratio of split. Training a model with 80 percent of the dataset can affect the model

performance. The efficiency of the ML model is assessed using the key parameters including accuracy, specificity, precision, recall, F-score, specificity, and sensitivity, respectively. The values of the parameters are calculated based on the confusion matrix. The data ingestion results are carried out using the traditional as well as the proposed approaches. The data loading results are compared with the traditional approach of data ingestion. The execution time and throughput results are compared with state-of-the-art existing proposals. The BP model evaluation is performed using the confusion matrix values to find the accuracy of the customized model. Authentic IoT-enabled datasets with big data characteristics are utilized to evaluate the applicability and performance of the system. The datasets include the IoT-enabled health, traffic, pollution, and water datasets [43]–[46]. The selected parameters, e.g. selection criteria for the results and validation section, is the selection criteria utilized by the base paper. We used similar setup of the existing studies. Moreover, we used identical datasets for evaluation and comparative analysis. Furthermore, the identical configuration (including CPU, storage, nodes, and cluster) is favored for evaluation. The proposed framework is interoperable and supports the cross-platform applicability as it is implemented in the Apache Spark, which is an open source framework.

A. Data Ingestion

The data ingestion is performed using an optimized mapping algorithm of MapReduce paradigm to load the data in parallel and speed up the overall execution time of big data processing and training module. We optimized the map-reduce programming paradigm and proposed an optimized map-only algorithm for data ingestion. The map-only algorithm is integrated with proposed architecture for data loading. In map-only algorithm, all the tasks in the system are performed in parallel that improves the data ingestion competence of our proposed system. It is observed that it gets approximately no time to load the data split into the distributed mechanism of Spark when the size of the data is small, e.g. up to 2GB. The data loading efficiency is highlighted in Fig. 8 that reveals drastic changes when the data size increases. It is evident from the figure that the proposed data loading is better than the traditional, i.e., classical approach. The proposed method loads the data with less time at each processing node of the cluster. The overall efficiency and improvement of the proposed big data ingestion module with a comparison to the traditional approach is also depicted in Fig. 9. In addition, the threshold value is highlighted in Fig.10. The Threshold Limit Value (TLV) is a specific range where the effect of the proposed approach is noticed. The TLV of the proposed approach is approximately 2 gigabytes. The data loading efficiency is revealed when the data exceeds the 2 gigabytes size.

B. ML Model Performance Evaluation

The ML model is implemented using the Apache Spark MLlib library. The evaluation of the ML model is performed using the confusion matrix. The parameters include accuracy, precision, recall, and F-measure. In addition, sensitivity, and

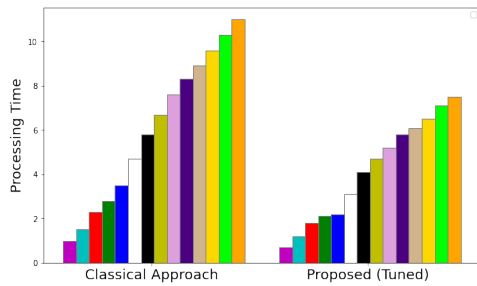


Fig. 8. Execution Time of Proposed Data Loading Approach.

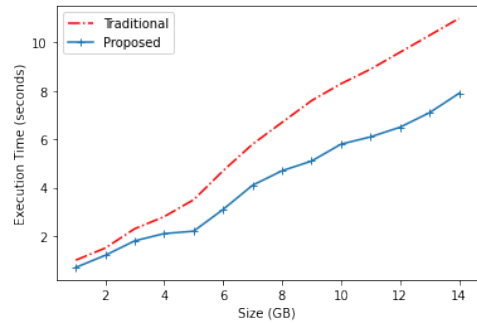


Fig. 9. Execution Time Comparison.

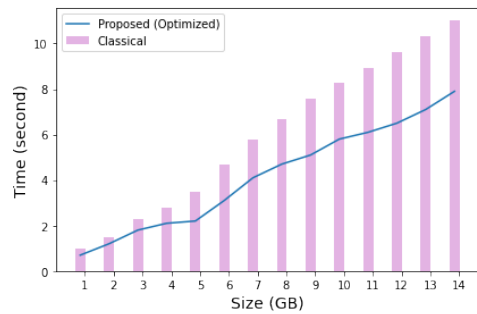


Fig. 10. Execution Time TLV.

specificity, are also considered. Accuracy is a ratio of the correctly classified records to the total number of records. Precision is a ratio between the correctly predicted positive records over the total positive predicted records. Recall is a ratio between correctly predicted positive records over the total related records. A combination of weighted average precision and recall indicators tell about the overall measured accuracy. F1 score of 1 is considered the perfect and the model is considered failed if the score is equal to 0.

The accuracy and precision of the proposed system are presented in Fig. 11. The figure highlights the comparative analysis of accuracy and precision between different processing nodes. The recall and F-measure of the system for multiple nodes are also depicted in Fig. 12. Similarly, the comparison of all the KPIs is provided in Fig. 13. Moreover, we provided the sensitivity and specificity of the proposed system model in Fig. 14. The proposed optimized learning model is also realized without Spark. The model accuracy can be further enhanced by expending dissimilar batch-size and progressively

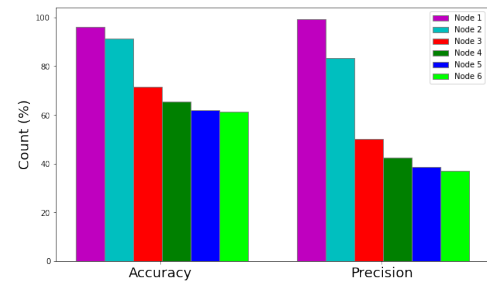


Fig. 11. BP Model Accuracy and Precision.

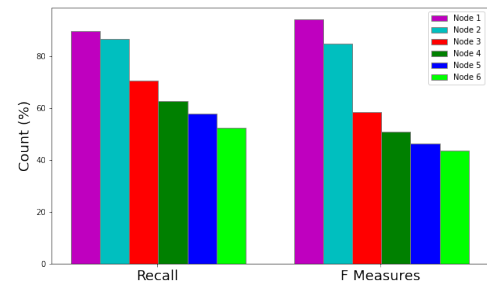


Fig. 12. BP Model Recall and F-Measure.

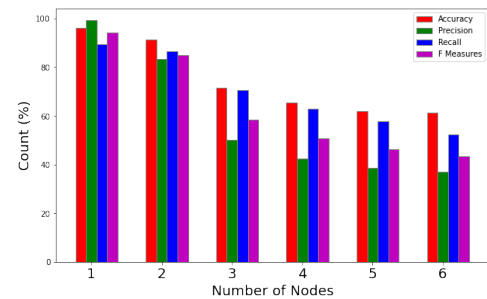


Fig. 13. KPIs Comparison.

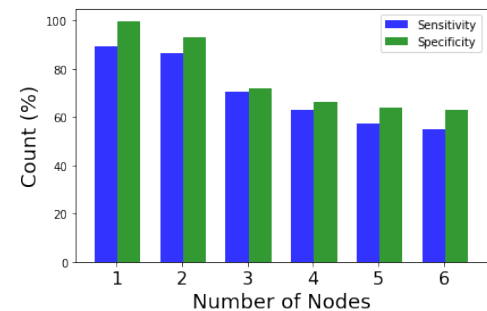


Fig. 14. Sensitivity and Specificity of Weighted Model.

increasing the epochs number. It is revealed that we achieve 81.07% accuracy on batch size 200 and the number of epochs 170 as depicted in Fig. 15. The accuracy refers to how often the model's predicted results are correct. Accuracy is the ratio between correctly classified images to the total number of images.

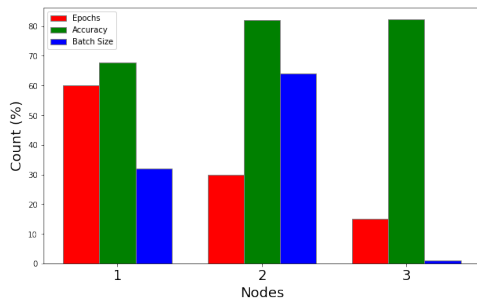


Fig. 15. Epoch-Batch Size-Accuracy Comparison.

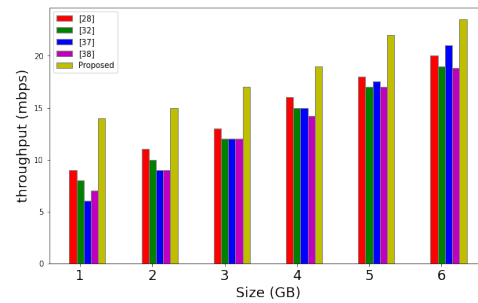


Fig. 18. Throughput Comparison.

C. Optimized YARN

The experimental results of our proposed architecture for execution time and throughput are presented in Fig. 16 and Fig. 17, respectively. The customization in the YARN cluster management improves the processing time as shown in Fig. 16. The processing time in the context of parallel nodes are also elaborated in Fig. 17. The throughput increases with a reduction in the computation time. The computation time reduces as the parallel and distributed framework is customized with optimized replica number, block sizes, and the usage of a aligned utility, as shown in Fig. 18.

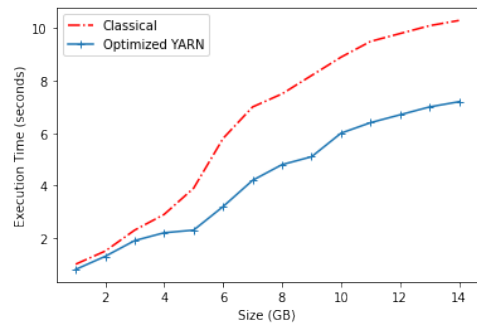


Fig. 16. Tailored YARN execution time.

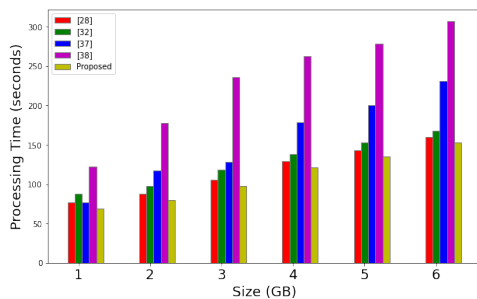


Fig. 17. Processing Time Comparison.

V. CONCLUSION

The massive data generated by IoT devices grow at an exponential rate at the network edge. Edge-enabled solutions offer efficient computing and control nearer to the network edge to resolve the scalability and latency challenges for IoT devices.

The existing solutions face numerous challenges in their structural design, for instance, the high volume of data storage and processing, data heterogeneity, and processing time among others. Moreover, parallel data ingestion and robust mechanism for handling communication overhead is also missing in the existing approaches. To address these challenges, this article proposes an IoT-enabled big data analytics framework for edge-centric computing. In the proposed scheme, an edge intelligence module is introduced and equipped with an optimized weighted-tuned ML model. To process and store the big data efficiently at the network edges, the optimized processing module is introduced with the integration of cloud platform. The proposed scheme is composed of two layers, i.e., IoT-edge and cloud-processing. The data injection and storage are carried out with an optimized MapReduce parallel algorithm. Optimized YARN is used for efficient management of the cluster. The proposed data design is simulated based on a real-world dataset using the Apache Spark and has much better results as compared to the traditional approaches.

ACKNOWLEDGMENTS

This work is partially supported by the NIH (NIGMS/P20GM109090), NSF under awards CNS-2016714, CNS-2104700, and CBET-2124918, the Nebraska University Collaboration Initiative, and the Nebraska Tobacco Settlement Biomedical Research Development Funds.

REFERENCES

- [1] Liu, Yongxin, Jian Wang, Jianqiang Li, Houbing Song, Thomas Yang, Shuteng Niu, and Zhong Ming. "Zero-bias deep learning for accurate identification of Internet-of-Things (IoT) devices." *IEEE Internet of Things Journal* 8, no. 4 (2020): 2627-2634.
- [2] Chen, Baotong, Jiafu Wan, Antonio Celesti, Di Li, Haider Abbas, and Qin Zhang. "Edge computing in IoT-based manufacturing." *IEEE Communications Magazine* 56, no. 9 (2018): 103-109.
- [3] Yu, Wei, Fan Liang, Xiaofei He, William Grant Hatcher, Chao Lu, Jie Lin, and Xinyu Yang. "A survey on the edge computing for the Internet of Things." *IEEE access* 6 (2017): 6900-6919.
- [4] Li, He, Kaoru Ota, and Mianxiong Dong. "Learning IoT in edge: Deep learning for the Internet of Things with edge computing." *IEEE network* 32, no. 1 (2018): 96-101.
- [5] Zhao, Zhiwei, Geyong Min, Weifeng Gao, Yulei Wu, Hancong Duan, and Qiang Ni. "Deploying edge computing nodes for large-scale IoT: A diversity aware approach." *IEEE Internet of Things Journal* 5, no. 5 (2018): 3606-3614.
- [6] Hassan, Najmul, Saira Gillani, Ejaz Ahmed, Ibrar Yaqoob, and Muhammad Imran. "The role of edge computing in internet of things." *IEEE communications magazine* 56, no. 11 (2018): 110-115.

- [7] J. Hajjaji, Yosra, Wadii Boulila, Imed Riadh Farah, Imed Romdhani, and Amir Hussain. "Big data and IoT-based applications in smart environments: A systematic review." *Computer Science Review* 39 (2021): 100318.
- [8] Zheng, Xu, Ling Tian, Bei Hui, and Xin Liu. "Distributed and Privacy Preserving Graph Data Collection in Internet-of-Thing Systems." *IEEE Internet of Things Journal* (2021).
- [9] Ranjan, Jayanthi, and Cyril Foropon. "Big data analytics in building the competitive intelligence of organizations." *International Journal of Information Management* 56 (2021): 102231.
- [10] Singh, Sanjay Kumar, and Abdul-Nasser El-Kassar. "Role of big data analytics in developing sustainable capabilities." *Journal of cleaner production* 213 (2019): 1264-1273.
- [11] Blazquez, Desamparados, and Josep Domenech. "Big Data sources and methods for social and economic analyses." *Technological Forecasting and Social Change* 130 (2018): 99-113.
- [12] Lv, Zhihan, Ranran Lou, Jinhua Li, Amit Kumar Singh, and Houbing Song. "Big data analytics for 6G-enabled massive internet of things." *IEEE Internet of Things Journal* 8, no. 7 (2021): 5350-5359.
- [13] Zhou, Lina, Shimei Pan, Jianwu Wang, and Athanasios V. Vasilakos. "Machine learning on big data: Opportunities and challenges." *Neuro-computing* 237 (2017): 350-361.
- [14] Li, Wei, Yuanbo Chai, Fazlullah Khan, Syed Rooh Ullah Jan, Sahil Verma, Varun G. Menon, and Xingwang Li. "A comprehensive survey on machine learning-based big data analytics for IoT-enabled smart healthcare system." *Mobile Networks and Applications* (2021): 1-19.
- [15] Marculescu, Radu, Diana Marculescu, and Umit Ogras. "Edge AI: Systems Design and ML for IoT Data Analytics." In *Proceedings of the 26th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 3565-3566. 2020.
- [16] Nour, Boubakr, Soumaya Cherkaoui, and Zoubair Mlika. "Federated Learning and Proactive Computation Reuse at the Edge of Smart Homes." *IEEE Transactions on Network Science and Engineering* (2021).
- [17] Roh, Yuji, Geon Heo, and Steven Euijong Whang. "A survey on data collection for machine learning: a big data-ai integration perspective." *IEEE Transactions on Knowledge and Data Engineering* (2019).
- [18] Zhou, Lina, Shimei Pan, Jianwu Wang, and Athanasios V. Vasilakos. "Machine learning on big data: Opportunities and challenges." *Neuro-computing* 237 (2017): 350-361.
- [19] Malek, Y. Nait, Abdelhak Kharbouch, H. El Khoukhi, Mohamed Bakhrouy, Vincenzo De Florio, Driss El Ouadghiri, Steven Latré, and Chris Blondia. "On the use of IoT and big data technologies for real-time monitoring and data processing." *Procedia computer science* 113 (2017): 429-434.
- [20] Iqbal, Rahat, Faiyaz Doctor, Brian More, Shahid Mahmud, and Usman Yousuf. "Big data analytics: Computational intelligence techniques and application areas." *Technological Forecasting and Social Change* 153 (2020): 119253.
- [21] Taherkordi, Amir, Frank Eliassen, and Geir Horn. "From IoT big data to IoT big services." In *Proceedings of the Symposium on Applied Computing*, pp. 485-491. 2017.
- [22] Hajjaji, Yosra, Wadii Boulila, Imed Riadh Farah, Imed Romdhani, and Amir Hussain. "Big data and IoT-based applications in smart environments: A systematic review." *Computer Science Review* 39 (2021): 100318.
- [23] Chaudhry, Abdul Aziz, Rafia Mumtaz, Syed Mohammad Hassan Zaidi, Muhammad Ali Tahir, and Syed Hassan Muzammil School. "Internet of Things (IoT) and machine learning (ML) enabled livestock monitoring." In *2020 IEEE 17th International Conference on Smart Communities: Improving Quality of Life Using ICT, IoT and AI (HONET)*, pp. 151-155. IEEE, 2020.
- [24] Porkodi, V., D. Yuvaraj, Juveriya Khan, S. Anbu Karuppusamy, Pallavi Murghai Goel, and M. Sivaram. "A Survey on Various Machine Learning Models in IOT Applications." In *2020 International Conference on Computing and Information Technology (ICIT-1441)*, pp. 1-4. IEEE, 2020.
- [25] Porkodi, V., D. Yuvaraj, Juveriya Khan, S. Anbu Karuppusamy, Pallavi Murghai Goel, and M. Sivaram. "A Survey on Various Machine Learning Models in IOT Applications." In *2020 International Conference on Computing and Information Technology (ICIT-1441)*, pp. 1-4. IEEE, 2020.
- [26] S. Din, H. Ghayvat, A. Paul, A. Ahmad, M. M. Rathore, and I. Shafi. "An architecture to analyze big data in the internet of things." In *ICST*, pages 677-682, Dec 2015.
- [27] M.M.Rathore, A. Ahmad A.Paul, S.Rho, Urban planning and building smart cities based on the internet of things using big data analytics, *Comput. Networks* 101(2016) 63-80
- [28] M. Mazhar Rathore, Anand Paul, Awais Ahmad, Marco Anisetti, and Gwanggil Jeon. Hadoop-Based Intelligent Care System (HICS): Analytical approach for big data in IoT. *ACM Transactions on Internet Technology (TOIT)*, December 2017. ISSN 1533-5399 (print), 1557-6051 (electronic).
- [29] Hussain, W., Merigó, J.M., Raza, M.R. and Gao, H., 2022. A new QoS prediction model using hybrid IOWA-ANFIS with fuzzy C-means, subtractive clustering and grid partitioning. *Information Sciences*, 584, pp.280-300.
- [30] Sun, H. Song, A. J. Jara, and R. Bie, "Internet of Things and big data analytics for smart and connected communities," *IEEE Access*, vol. 4, pp. 766-773, Mar. 2016, doi: 10.1109/ACCESS.2016.2529723
- [31] Tönjes, R., Barnaghi, P., Ali, M., Mileo, A., Hauswirth, M., Ganz, F., Pui, D. Real time iot stream processing and large-scale data analytics for smart city applications. *Proceedings of the European Conference on Networks and Communications (poster session)*, Bologna, Italy. (2014)
- [32] Rathore, M.M., Paul, A., Ahmad, A., Jeon, G.: IoT-based big data: from smart city towards next generation super city planning. *Int. J. Semant. Web Inf. Syst.* 13(1), 28-47 (2017)
- [33] B. N Silva, M. Khan, and K. Han, "Big Data Analytics Embedded Smart City Architecture for Performance Enhancement through Real-Time Data Processing and Decision-Making", "Wireless Communications and Mobile Computing, Vol 2017, Article ID 9429676,
- [34] B.Cheng,S.Longo,F.Cirillo,M.Bauer,andE.Kovacs,"Building a big data platform for smart cities: experience and lessons from santander," in *Proceedings of the 4th IEEE International Congress on Big Data (BigData Congress '15)*, pp. 592-599, New York, NY, US A, July 2015.
- [35] M. Rathore, A. Ahmad, A. Paul and G. Jeon, "Efficient Graph-Oriented Smart Transportation Using Internet of Things Generated Big Data," 2015 11th International Conference on Signal-Image Technology and Internet-Based Systems (SITIS), Bangkok, 2015, pp. 512-519
- [36] J. M. MazharRathore, Anand Paul, Awais Ahmad, Marco Anisetti, and Gwanggil Jeon. Hadoop-Based Intelligent Care System (HICS): Analytical approach for big data in IoT. *ACM Transactions on Internet Technology (TOIT)*, December 2017. ISSN 1533-5399 (print), 1557-6051 (electronic).
- [37] Hussain, W., Merigo, J.M., Gao, H., Alkalbani, A.M. and Rabhi, F.A., 2021. Integrated AHP-IOWA, POWA framework for ideal cloud provider selection and optimum resource management. *IEEE Transactions on Services Computing*.
- [38] Wang, Tian, Haoxiong Ke, Xi Zheng, Kun Wang, Arun Kumar Sangaiah, and Anfeng Liu. "Big data cleaning based on mobile edge computing in industrial sensor-cloud." *IEEE Transactions on Industrial Informatics* 16, no. 2 (2019): 1321-1329.
- [39] Leung, Carson K., Deyu Deng, Calvin SH Hoi, and Wookey Lee. "Constrained big data mining in an edge computing environment." In *International Conference on Big Data Applications and Services*, pp. 61-68. Springer, Singapore, 2017.
- [40] A. C. Nicolaescu, S. Mastorakis and I. Psaras, "Store Edge Networked Data (SEND): A Data and Performance Driven Edge Storage Framework," *IEEE INFOCOM 2021 - IEEE Conference on Computer Communications*, 2021, pp. 1-10, doi: 10.1109/INFOCOM42981.2021.9488804.
- [41] Al Azad, Md Washik, Susmit Shannigrahi, Nicholas Stergiou, Francisco R. Ortega, and Spyridon Mastorakis. "CLEEDGE: A Hybrid Cloud-Edge Computing Framework over Information Centric Networking." In *2021 IEEE 46th Conference on Local Computer Networks (LCN)*, pp. 589-596. IEEE, 2021.
- [42] Jan, Mian Ahmad, Fazlullah Khan, Rahim Khan, Spyridon Mastorakis, Varun G. Menon, Mamoun Alazab, and Paul Watters. "Lightweight Mutual Authentication and Privacy-Preservation Scheme for Intelligent Wearable Devices in Industrial-CPS." *IEEE Transactions on Industrial Informatics* 17, no. 8 (2020): 5829-5839.
- [43] UCI Machine Learning Repository." UCI. 01 01. Accessed 04 01, 2021. <https://archive.ics.uci.edu/ml/machine-learning-databases/heart-disease/>.
- [44] T. Dataset, Dataset Collection, <http://iot.ee.surrey.ac.uk:8080/datasets.html/> traffic.
- [45] P. Dataset, Dataset Collection, <http://iot.ee.surrey.ac.uk:8080/datasets.html/> pollution
- [46] Dataset Water meters, <http://data.surrey.ca/dataset/water-meters>