

Design and Development of Great-Recital Processing Created Methodology for Refining Semantic-Created Amalgamated Data Dispensation Technology

S. Ravichandran

Professor, Department of Computer Science, Annai Fathima College of Arts and Science, Madurai, Tamil Nadu, India

E-mail: dravichandran6@gmail.com

(Received 6 February 2021; Revised 18 February 2021; Accepted 15 March 2021; Available online 23 March 2021)

Abstract - Recovering RDF datasets from appropriated information sources has become a fundamental interaction to accomplish the vision of the semantic web. The semantic web vision advances an everyday growing of the semantic diagram that requires some preparing upgrades including SPARQL end-point and combined questions to secure the interminable extension. The current progressed elite figuring engineering has been growing quickly to conquer numerous issues including execution. Elite registering can possibly be utilized to improve the unified SPARQL inquiry and generally speaking activity including the performance and handling. Hence, this work surveys the procedures to advance the improvement on isolated and distant combined semantic-based information. That is, the superior figuring climate including equipment engineering and extraordinary design registering has been overviewed. Besides, working semantic diagram combination necessities is dissected. Besides, the current strategies to demonstrate SPARQL execution are considered. Furthermore, an examination about the usage of cutting edge equipment figuring engineering is investigated to upgrade the execution of the current SPARQL alliance administration activity.
Keywords: SPARQL, CUDA, Linked Open Data, Parallelization, HPC, Semantic Web and Linked Data

I. INTRODUCTION

The technique of collaborating and accessing information over the web has evolved resulting in the new era of the web called Semantic Web (SW) [1]. SW is a newer version of the web that supports the existing functionality of the existing one and aimed to make the web machine-understandable [2]. It will develop dedicated pieces of data together that allows the machine to understand the contents [3]; that is; computers can interoperate and reason on behalf of users. SW has emerged as a strategic technology for formally and semantically linking data together forming a mesh of knowledge and its main goal is to provide a unified cloud-based environment [4] for data that holds the semantics in such a way that a machine can understand, process, and utilize. According to the W3C, semantic web technologies “enable people to create data stores on the Web, build vocabularies, and write rules for handling data” [5]. Linked data is defined as “a method of publishing structured data so that it can be inter-linked and become more useful through semantic queries, founded on HTTP, RDF and URIs” [6]. It aggregates some Semantic Web Technologies and provides a simple interface where an application can interact with these open data [7]. As

reviewed by the inventor the Semantic Web, the recommendation of helpful Linked Data should be as follows: “(1) use URIs as names for things, (2) use HTTP URIs so that people can look up those names, (3) when someone looks up a URI, useful information is provided using the standards (RDF, SPARQL), and (4) include links to other URIs so that they can discover more things”. Linked Open Data (LOD) cloud project shows the advanced level approached by the community. Each node in this cloud diagram represents a distinct data set published as Linked Data. As of March 2019, LOD collected 1,239 datasets with 16,147 links with many domains. For each one of these datasets, a minimum of 1000 triples are required at 50 connections to other datasets on the LOD cloud. It is worth mentioning that one dataset already contain 3,000,000,000 triples as will discussed later.

Computing techniques with high processing powers with a massive amount of memory are one of the main gains of High-Performance Computing (HPC) [8]. They utilize parallel computing to achieve many objectives. However, single-processor computers could not provide an effective processing capability that is required for the HPC applications [10]. Parallel computing [9] is a useful technique to enhance the computations efficiency for advanced tasks such as big data; that is, the design of algorithms, models, and software adopts the multi-core architecture computers for better performance and scalability [11]. The Semantic Web promotes the expansion of a massive open semantic linked data. This vision must reside in the availability and accessibility to avoid outdated data processing. In order to overcome this issue and its consequences challenges, fast and scalable processing capabilities should be involved. In this research, the nature of the evolving LOD from the federated query point of view with respect to the HPC available environments is studied. The rest of the paper is organized as follows: Section 2 shows federated data management and an overview of the basis of Semantic Web concepts and technologies and it introduces the evolution of data over the semantic web era [12]. In section 3, a spotlight is on the High-Performance Computing (HPC) techniques. In section 4, the semantic architecture is discussed together with a presentation of a case using advanced hardware computing. In section 5, the related literature is discussed. Finally, the conclusion of the paper is in section 6.

II. RELATED WORK

A. Federated Data Management

RDF represents Resource Description Framework [13] that is a coordinated labelled diagram making SPARQL is a semantic chart coordinating with RDF inquiry language [14]. It is a World Wide Web Consortium (W3C) standard for de-scribing assets over the web. In this manner, utilize metadata portrays information to convey more data about components. This metadata empowers the automatic preparing of web assets. RDF language can be viewed as an information model for developing the information on the web and relations existing on the planet. RDF addresses information in the game plan of Subject-Predicate-Object as demonstrated in figure 1. It is perceived as RDF significantly increases $\langle S, P, O \rangle$ that number of information together characterizes a huge assortment of triples as a RDF chart. Hence, connecting more than one RDF chart structure countless semantically connected diagrams.

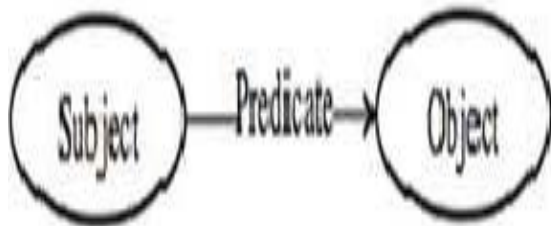


Fig. 1 Data Representation of RDF

The Unique Resource Identifier (URI) gives an assigned area of the information depiction or literals (i.e., strings, date, numbers, and so on) or on the other hand basically a clear hub. RDF can effectively associate the information to fabricate a chart of associated information as the semantic web local area as of now utilizes it and practices [15][16]. It additionally works with information sharing and reusing across applications, endeavours, and local area limits [17]. As of late, RDF and Linked information acquired higher consideration; i.e., Dbpedia. This uncovered new progressed prerequisites to get to, recover, and handle circulated and dynamic semantically and soundly associated datasets with advancement procedures for suitable question preparing.

These triples require an extraordinary inquiry language to be gotten to and prepared. This is going on by means of a standard semantic language called SPARQL. The question language is utilized as a RDF inquiry language. It works with isolated and far off RDF datasets. Practically speaking, SPARQL gives a chart design idea that al-lows clients to coordinate against a given RDF dataset. SPARQL which represents SPARQL Protocol and RDF Query Language [17], is a normalized RDF charts inquiry language created by the World Wide Web Consortium. SPARQL convention methodically gets to different information sources [17] through their far off endpoints to create lexical inquiries [16]. Since RDF depends on diagrams, SPARQL utilizes

chart design coordinating [17] by using significantly increases. The SPARQL inquiry can be organized as the command PREFIX assertion to permit contract the portrayal URIs, the announcing re-mission and design, and the question chart pat-tern for certain conditions to adjust the outcomes as added in figure 2 snapped from [14].

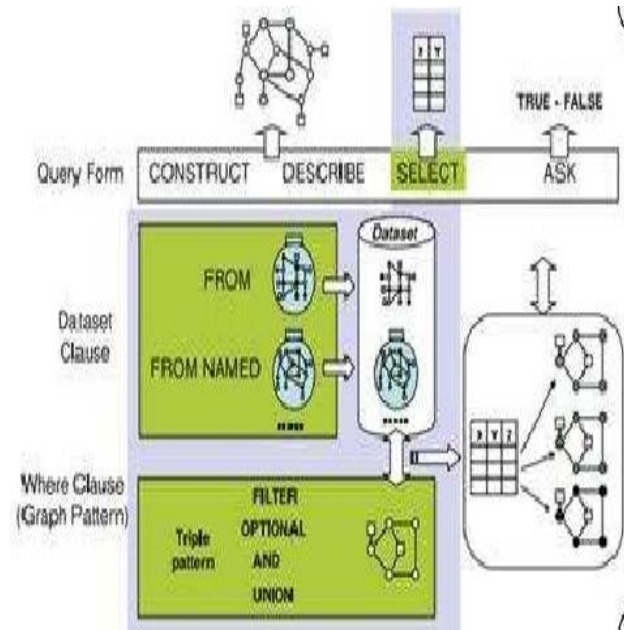


Fig. 2 Sample Diagram of SPARQL Query

Presently, SPARQL 1.1 is the more current form that permits progressed highlights including "united question" that permits inquiries from circulated SPARQL endpoints; that is, important for the actual question can be conjured against a given SPARQL endpoint. That is, in a united question, each The SPARQL inquiry is deco presented into sub-inquiries, and every one of them is shipped off given SPARQL endpoints to be executed. When each solicitation is finished, the outcomes are consolidated and introduced. SPARQL permits RDFS (RDF Schema) and Ontology questioning and moving toward triple stores and unified end-focuses making diverse area is prepared immediately. The federator is answerable for the understandable sub-questioning preparing and gathering.

III. PROBLEM STATEMENT

A. High-Performance Computing

The solid methodology is normally continued in questioning circulated semen-spasm chart that is creating holding up time in an extraordinary mode. Notwithstanding, the idea of the circulated semantic diagram climate upholds numerous re-missions simultaneously. Subsequently, parallelism in such manner uses the force of figuring as far as preparing a high number of solicitations in a brief time frame. The way toward planning a High Processing Computing (HPC) framework requires a fitting equipment design with the important equal algorithm and for the most part utilized for

logical registering, industry, wellbeing, protection, designing, security, or public research centres. The association, micro connectivity, size of the assets, and accessibility of the vigorous programming to deal with activity at that given scale are what make HPC did not quite the same as commonplace PCs [12]. The cycles like circulation, calculation, and last assortment of use information are considered in the planning stage to expand the productivity of a framework. Around there, monstrous equal preparing is needed for raising registering handling execution. The idea of multicore PCs is getting well known among programming planers [12]. Three regular kinds of parallelization models: shared memory parallelization [29], an idea that has a common memory for all favourable to processor and appropriated memory parallelization [17], an idea that has committed memory for every processor as depicted beneath, and Special Architecture as talked about later.

B. Shared Memory Parallelization

Shared memory parallelism is subject to divided memory to give straightforward correspondence among the distinctive preparing hubs. Every processor will play out a devoted errand. A common variable can be gotten to by processors as indicated by the prerequisites. The various varieties of shared memory are Uniform Memory Access (UMA), Non-uniform Memory Access (NUMA), and Cash-Only Memory Architecture (COMA). Strings are a lightweight cycle that requires a modest quantity of advancement time and less endeavours to be booked by the working frameworks. A cycle can make more than one string that runs at a similar location space as the parent favourable to cess. Multi-centre design is a solitary registering unit that contains some preparing unit (called centre) at the very chip that may share the memory and run the directions in equal [16].

C. Distributed Memory Parallelization

Dispersed memory parallelization is a one of a kind methodology that uses a dedicated memory for every processor [13]. It doesn't have a typical memory for processors. It utilizes a message interface for processors to build up correspondence between them. Accordingly, every processor will have an exceptional ID in the correspondence bunch. At the point when a processor x needs to speak with a processor y, it sends an immediate message to y utilizing its location. The Message Passing Interface (MPI) is a notable determination of an equal library for setting up correspondences among processors. It will give an assortment of APIs that can be called by the application code to speak with different processors. The determination of MPI will give two kinds of calls: highlight point and aggregate calls.

D. Special Architecture

Exceptional Architecture either uses equal processing or utilizes equal figuring by an imaginative methodology. In

bunches [12], a bunch of computers are associated together to play out a specific undertaking. In some design, all CPUs are on a similar board and offer a similar framework transport. In an-other design, a CPU can have its board and fringe gadgets. A fast organization association is needed to join these in an unexpected way. Scalability is one of the primary benefits here; that is, new hubs can be added to the bunch easily to build the group framework computational force [13]. Groups may use a conveyed memory design where MPI vehicles out the correspondences among bunch hubs. Be that as it may, assuming the bunch utilizes circulated memory engineering, the association medium among group hubs is presented to high correspondence overhead. On the other hand, bunches May use shared memory structures or a mixture model architecture where a few hubs share memory (multi-code processors) while the individual hubs utilize circulated memory engineering. In networks processing [15], geologically separate computational assets are used to take care of a computational issue in a heterogeneous climate way that may not really have a similar equipment/programming structures; that is, each errand is handled in very surprising machines permitting by and large application parallelism. Through message passing and a unique program, these remotely relegated machines can be imparting and dealing with the running undertakings on a similar matrix.

Framework processing requires high undertaking planning super-vision to such an extent that assets on the lattice be altogether used. Distributed computing [13] is an arising figuring model that embraces the disseminated and virtual utilization of assets advancing expense decrease and convenience. In the distributed computing model, the Cloud Service Provider (CSP) offers certain indecencies to singular clients and business organizations. There are three layers in the cloud structure that addresses a layer of calculation. These layers are Software-as-a-Service (SaaS) layer, Platform-as-a-Service (PaaS) layer, and Infrastructure-as-a-Service (IaaS) layer. One significant contrast here among cloud, matrix, and group registering is that distributed computing doesn't just offer a model of calculation. Nonetheless, it can offer different types of assistance; i.e., capacity, exceptional asset, or computational pool leasing. Haze registering is an all-inclusive rendition of distributed computing that extends and improves on the idea of distributed computing to beat some distributed computing constraints; i.e., inertness and supports smaller than normal exercises including capacity, communicating among gadgets and server farms [10].

Wilderness registering is additionally unique engineering that exceptional the idea of dispersed figuring by means of synchronous union of various designs to lessen the programming intricacy and smooth out computational assignments among bunches to accomplish top execution[11]. The race of supercomputer treatment is a functioning field because of reasons referenced previously. As per [5] [9], table I shows the main ten quickest superior PCs on the planet concerning November 2019. Rmax and

Rpeak are scores used to rank supercomputers reliant upon their presentation using the LINPACK Benchmark. In the table underneath, Rmax represents Maximal LINPACK execution accomplished, Rpeak represents Theoretical pinnacle execution, and GFlops represents Giga (billion) Floating Point Operations per Second.

IV. IMPLEMENTATION

A. The Semantic Graph Requirements and Leveraging

The vision of the semantic web is to permit information distributed in a coordinated way so that they interlaced and serve the semantic questions. In this manner, it requires datasets to be gotten too distantly and handled the integration of these information; that is, semantic incorporation permits arrangement of the changed assertions from triple stores and measures in a unified stage. Handling datasets disconnected is fine. In any case, it is against the progression of understanding the vision of the semantic web. As expressed in [15] "Joining requires spending assets on planning heterogeneous information things, settling clashes, cleaning the information, etc. Such expenses can likewise be tremendous.

The expense of incorporating a few sources may not be advantageous if the increase is limited, particularly within the sight of excess information and bad quality information." The absence of legitimate handling of connected datasets on the cloud, for example, looking or summing up was prompting issues in getting to and mentioning these information. For these datasets to be prepared productively, combination is critical; that is, the test is in joining the required datasets that are appropriated in better places utilize assorted URIs and models while addressing their entities and pattern components.

The incorporation of records from free and disseminated datasets is needed to carries out a bunch of SPARQL re-journeys. SPARQL question streamlining handles the assortments of joins strategies to coordinate RDF datasets including tie join, settled circle join, hash join, symmetric join, various hash join, sub query building, join strategy choice, also, join requesting. Combination substance (or incorporation design) can either be appeared or virtual [7].

RDF chart datasets should coordinate with LOD standards for utilizing the Semantic Web vision; that is, downloading the necessary arrangement of RDF diagrams to be privately handled would take into consideration restricted information legitimacy and negate the vision of the semantic web. In "sub-mists" of LOD, spaces are marked to help the documentation of on tologies. Figure 3 shows the topography space from the LOD cloud.

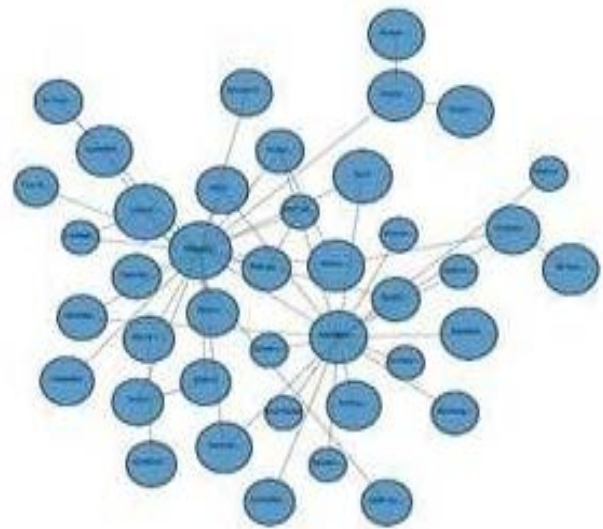


Fig. 3 Basic Diagram of Geography sub-cloud

One of the RDF datasets from geology sub-cloud is Yahoo! Geo Planet that gives intelligent foundation to geo-referring to information. The absolute size of this informational index just is 49,734,022 triples. Another is Linked Geo Data that uses Open-Street Map projects information with all out size 3,000,000,000 triples. Concerning the capacity to request far off endpoints through the SERVICE proviso, these datasets can be reached at the same time to pass on demands. The unified SPARQL question can be shipped off various dataset stockpiles as Linked Open Data as demonstrated in figure 4 by [16]. The solid condition of registering was moderate and had execution issues; i.e., decrease of speed [9]. The explanation was the execution of an undertaking has finished totally before another errand. Every one of the undertakings ought to be masterminded in a line and executed all together. There-front, if there are various undertakings to be done, say thousands, a more noteworthy time would be needed for their finishing [16]. HPC pointed toward expanding the presentation and abilities. Various advancements have occurred to accelerate the equipment handling. Reviewing the unique engineering computing and the extraordinary figuring power conditions accessible at supercomputer focuses, legitimate appropriation of joint techniques would bring about improving generally semantic web execution activities including united SPARQL question preparing. Some registering designs are accessible as Message Passing Interface (MPI) that backings group hubs with their people RAM [8]. OpenMP is another alternative that upkeep numerous stages shared memory multiprocessing [11]. Likewise, there is OpenCL (represents Open Computing Language) that gives incredible abilities in task-planning on GPU [6]. For this situation, legitimate question request is essential here to accelerate the cycling and abbreviate the holding up time. Additionally, using the force of equal registering propels the pattern of the semantic advancement.

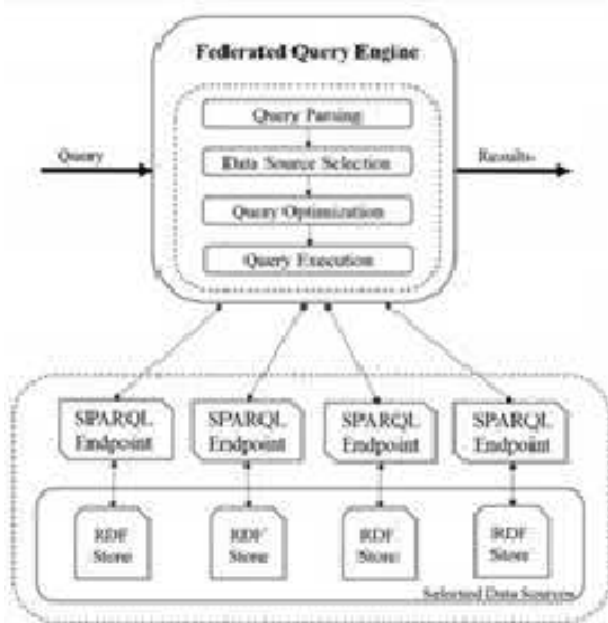


Fig. 4 SPARQL Federated Query Processing

HPC writing presented a few different ways of settling errands simultaneously to extraordinarily diminish the measure of execution time henceforth expanding the presentation in organization administrations or systems fitted with the alliance administrations. The implementation of shared memory engineering permits composing consecutive projects that divide similar memory between all figuring hubs. Unique consideration ought to be taken during the entrance of shared memory by various hubs. Also, shared memory design won't force any computational overhead as correspondence among hubs happens through the common memory. Actually, dispersed memory design doesn't expect that labourer hubs divide any sort of memory among them. Hence, processors can exploit passing messages to speak with one another, which may acquaint correspondence overhead due with the continuous correspondence between the labourer hubs. One of the renowned gigantic equal design is called Compute Unified Device Architecture (CUDA) that offers coverings for cutting edge programming dialects like Java and Python. Accordingly, this work is considered to act as an illustration of carrying out extraordinary coordinated SPARQL questioning execution preparing.

B. Compute Unified Device Architecture (CUDA)

In this part, a specialized conversation is introduced to examine one of the notable equal processing design (Compute Unified Device Architecture CUDA). This conversation customizes the accessibility of this architecture. It offers coverings for cutting edge programming dialects, for example, Java and Python that backings SPARQL question preparing. CUDA is implemented in the Graphics Processing Units (GPU). As a rule, GPU enjoys numerous benefits including minimal expense, straightforward, crude computational force, and

incredible memory transmission capacity [15]. It uses the utilization of equal computing innovation or ability. This innovation upgrades the presentation of the PC game industry by improving the nature of designs of a game and physical science estimation. The premier benefits of CUDA are the quick execution time and precise numerical estimations that outcome in the creation of sensible gaming and age of great pictures. The NVidias CUDA innovation upholds adaptable equal projects with enabling a wide assortment of utilizations like scanty framework, web search tools, models in material science, and calculations in science among others. The architecture is appeared in figure 5 from [14]. CUDA utilizes similar portions, which are equal, to late General Purpose Graphical Processing Unit (GPGPU) models. The contrast between the two models is that CUDA gives flexible strings creation, shared and worldwide memory, string blocks, and better synchronization [10]. In CUDA engineering, there are two sections: the host and the gadget. The program is running into a host program in CPU. That is, at least one equal portions for the gadget execution on GPU. A reasonable game plan among strings is required. String is a fundamental unit for the information control inside CUDA. The string listed as a multidimensional pinnacle (one-dimensional, two-dimensional, or three-dimensional string file). Every one of the strings access simultaneously hinders shared memory. Each square can be distinguished utilizing a one-dimensional or two-dimensional list and all networks will share a worldwide memory [5].

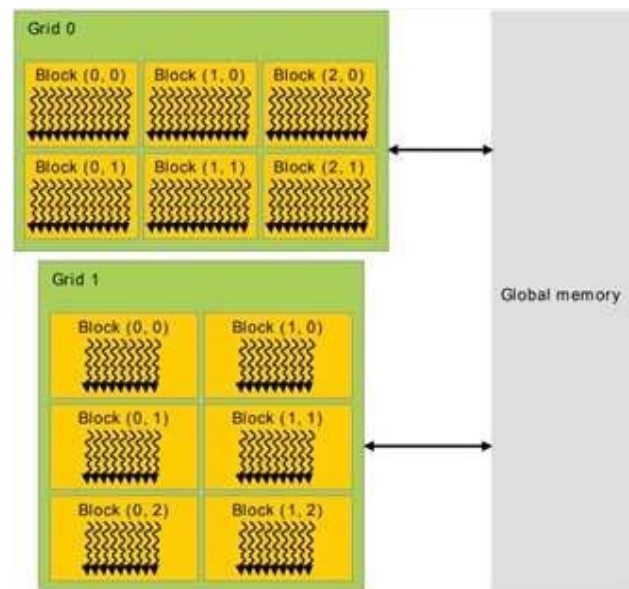


Fig. 5 Architecture of CUDA

In CUDA, shared memory is on a chip that expands effectiveness when contrasted with nearby worldwide memory (dormancy is multiple times lesser than the worldwide common memory) [15] Here, strings block approach shared memory. The strings should be synchronized so that right outcomes are given during equal strings collaboration. To support the introduction of this

methodology, the equal calculation on the graphical preparing units (GPUs) with engineering like NVidia's Compute Unified Device Architecture (CUDA) would help with moving a wide assortment of information in less measure of time. It utilizes the equal execution of information to execute numerous assignments simultaneously to speed up. That is, with additional tuning it would oblige the idea of the distributed semantic diagram. In the question motor, a preoccupied layer of manual setup would be required. For this situation here, a few coverings for Java, Python, and other programming dialects are accessible. It will save time and backing various applications on a solitary PC. It is an essential strategy for a steadily developing world [3]. Thusly, each SPARQL inquiry organization motor can arrange a nearby specially appointed alliance activity instead of stacking information that keep down the documentation of the semantic web. The earth shattering development of current semantic-based innovation and advancement will be expected later on. The new regions will involve the capacity of the information handled from the LOD and the showing of the knowledge with human interface, just as semantic and ontological information.

V. CONCLUSION

This investigation has given an outline of a way to deal with advance the favourable to censing of unified information the board from different exploration papers. The arising advances of uncommon HPC engineering would use the highlights of cutting edge preparing to offer support to end-clients or to consolidate favourable to censing hubs from various areas and give a uniform and lucidness computational forces to end-clients. In this work, the notable HPC architectures have been summed up to impart to the local area the conceivable crossing point between fields to progress semantic web errands like perusing, summing up, questioning, looking, sorting out, and so forth. Additionally, a spot-light to the field is presented for additional profound design to use HPC with the advancement of semantic diagrams. To end this, question optimization of inquiry league motors have been concentrated to improve the performance. Notwithstanding, arranging and utilizing elite registering engineering to improve in general inquiry organization handling is as yet an open test. Subsequently, one objective of this paper is to give motivation to in-corporate that which would permit expanding the semantic web applications.

It merits referencing that the blend of two models of equal supportive of censing procedures can deliver a half breed model where a few parts of the equal framework utilize shared memory while different segments will utilize a disseminated approach. The mixture model will consolidate both the Central Processing Unit (CPUs) and GPUs to take care of an issue. For this situation, CPUs may use circulated memory design and GPUs will utilize a common memory model. In certain situations, one can only with significant effort choose the kind of HPC design to execute a similar

application. At the end of the day, the application needs to handle an enormous number of information simultaneously. There is a need to move a measure of information among the handling hubs. HPC architecture can be a critical factor for the age of better outcomes. Thusly, the choice to convey certain applications utilizing an equal methodology ought to be application-subordinate. The equal preparing strategy works in a characterized climate. That is, controlling the preparing of execution will be quicker and produce precise outcomes. The question league motors de-pend on: no pre-processing per inquiry and unbound predicate inquiries. That is, the choice of information sources utilizing a metadata list may foster some exhibition issues if no legitimate setup on the HPC climate. In such manner, inquiry improvement with suitable HPC plan design would allow more prominent abilities in overseeing semantic combined information including improving the presentation. The overall standards of such models are including uniform semantic questioning of the dissemination and heterogeneity. Also, it has explored the strategies for equal dissemination of information to move a few types of information simultaneously. The difficulties of inquiry enhancement are expected to increment with moving semantic chart development. Semantics web innovations are growing that to assess people to share information in applications or sites. Accordingly, the exceptional engineering of elite registering with legitimate question improvement methods would lessen the expense of handling and increment the presentation.

REFERENCES

- [1] T. Berners-Lee, J. Hendler, and O. Lassila, "The Semantic Web," *Scientific American*, 2001.
- [2] G. Antoniou, P. Groth, F. Van Harmelen, R. Hoekstra and A. Semantic Web Primer, Third ed., MIT Press, 2012.
- [3] C. Weiss, P. Karras, and A. Bernstein, "Hexastore: Sextup leindexing for semantic web data management," *Proc. VLDB Endow.*, Vol. 2, No. 5, pp. 45-55, 2008.
- [4] S. Schlobach and C. A. Knoblock, "Dealing with the messiness of the web of data," *Journal of Web Semantics*, Vol. 4, No. 3, pp. 109-119, 2012.
- [5] X. Li, Z. Niu, and C. Zhang, "Towards efficient distributed SPARQL queries on linked data," in *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bio informatics)*, 2012.
- [6] A. I. Torre-Bastida, E. Villar Rodriguez, M. N. Bilbao and J. DelSer, "Intelligent SPARQL end points: Optimizing execution performance by auto maticquery relaxation and queue scheduling," in *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, 2016.
- [7] M. Acosta, M. E. Vidal, T. Lampo, J. Castillo and E. Ruckhaus, "ANAP- SID: An adaptive query processing engine for SPARQL end points," in *Lecture Notes in Computer Science (including subseries Lecture Notes In Artificial Intelligence and Lecture Notes in Bioinformatics)*, 2011.
- [8] M. Saleem and A. C. Ngonga Ngomo, "HiBISCuS: Hyper graph-based source selection for SPARQL end point federation," in *Lecture Notes in Computer Science (including sub series Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, Vol. 3, No. 6, pp. 90-102, 2014.
- [9] B. Quilitz and U. Leser, "Querying distributed RDF data sources with SPARQL," in *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, Vol. 2, No. 8, pp. 214-225, 2008.

- [10] M. Saleem and A. C. Ngonga Ngomo, "HiBISCuS: Hyper graph based Source selection for SPARQL end point federation," in *Lecture Notes In Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, 2014.
- [11] A. Bernstein, C. Kiefer and M. Stocker, "Opt ARQ: ASPARQL Optimization Approach based on Triple Pattern Selectivity Estimation," *Univ. Zurich*, Vol. 3, No. 1, pp. 56-66, Dep., 2007.
- [12] O. Hartig, "SQUIN: A Traversal Based Query Execution System for the Web of Linked Data," in *Proceedings of the 2013 International conference on Management of data – SIGMOD*, Vol. 13, No. 7, pp. 108-120, 2013.
- [13] J. Sathiamoorthy, Dr. B. Ramakrishnan and M. Usha "A Trusted Waterfall Frame work Based Peer to Peer Protocol for Reliable and Energy Efficient Data Transmission in MANETs," *Springer Journal of Wireless Personal Communication*, Article First Online: Vol. 3, No. 4, pp. 78-89, 17 May 2018.
- [14] D. Bednarek, J. Dokulil, J. Yahoo and F. Zavoral, "Using Methods of Parallel Semi-structured Data Processing for Semantic Web," in *Third International Conference on Advances in Semantic Processing*, Vol. 6, No. 7, pp. 44-49, 2009.
- [15] J. Sathiamoorthy, and Dr. B. Ramakrishna "A competent three tier fuzzy cluster algorithm for enhanced data transmission in cluster EAACK MANETs," *Springer Soft Computing*, Vol. 8, No. 6, pp. 123-135, July 2017.
- [16] X. Wang, T. Tiropanis, and H.C. Davis, "LHD: Optimising linked data query processing using parallelisation," in *CEUR Work shop Proceed*, Vol. 5, No. 3, pp. 71-82, 2013.
- [17] M. Hajibaba and S. Gorgin "A Review on Modern Distributed Computing Paradigms: Cloud Computing, Jungle Computing and Fog Computing," *J. Compute. Inf. Tech.*, Vol. 22, No. 2, pp. 69-84, 2014.