

Intelligent Load Balancing in Web Servers: A Review of Emerging Algorithms and Architectures

Pushkar Bati¹, Mruthunjaya S², Prajwal Gowda H R³, Prem⁴, Mithun B N⁵, Krushitha Shetty⁶
^{1,2,3,4,5,6}APSCE, Bengaluru, Karnataka, India

Abstract—The swift expansion of web services and internet usage has posed considerable challenges in maintaining optimal server performance, particularly in data centers, necessitating the use of load balancing as a crucial technique for distributing workloads among multiple servers to ensure effective resource usage, scalability, and reliability. This examination investigates different load-balancing techniques, including Round-Robin, Least Connections, and Weighted Least Connections, as well as more sophisticated methods like Dynamic Feedback Mechanisms and Reordering-Robust Load Balancing, focusing on enhancing response times, optimizing resource distribution, and improving fault tolerance in both traditional and cloud-based settings. Furthermore, developments in distributed caching, container orchestration, and adaptable server configurations are analyzed to address the requirements of contemporary applications, while machine learning-based strategies such as Dynamic Weighted Polling and Proximity-Based Load Balancing are highlighted for their flexibility and capability to manage high traffic and intricate workloads, ultimately enhancing server efficiency, minimizing latency, and providing a smooth user experience even during peak usage times.

Index Terms—Load Balancing, web servers, dynamic feedback, resource utilization, cloud computing.

I. INTRODUCTION

A server is a robust computer system or software that provides services, resources, or information to clients across a network, managing multiple requests at once to facilitate web services, applications, and data centers. Its primary functions encompass hosting websites, operating applications, storing information, and facilitating communication, with performance evaluated through response time, resource usage, and throughput. Servers are classified based on their roles into categories such as web servers (for HTTP/HTTPS delivery), database servers (for data management),

file servers (for file storage and distribution), application servers (for handling business logic), and mail servers (for email services). Web servers, which can consist of hardware, software, or a combination of both, process and provide content to browsers, where hardware servers in data centers guarantee scalability and reliability, while software solutions like Apache, Nginx, or IIS manage requests, scripts, and backend processes, supporting both static and dynamic offerings such as e-commerce, social media, and cloud-based applications.

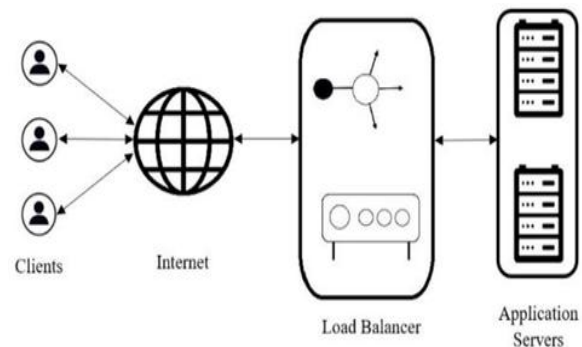


Figure 1. Basic architecture of the Load Balancer

The diagram represents a typical load balancing architecture where clients connect through the internet to a load balancer, which then distributes incoming requests to multiple application servers. This setup improves scalability, resource utilization, and fault tolerance by preventing overload on a single server. Load balancers employ strategies like Round Robin or Least Connection while also performing health checks to ensure traffic is routed only to active servers. In high-demand scenarios, such as e-commerce sales events, this architecture ensures faster response times, uninterrupted service, and an overall better user experience (see fig 1).

II. REVIEW OF LITERATURE

Load balancing has evolved from conventional techniques such as Round Robin and Least Connections to contemporary methods that incorporate machine learning, swarm intelligence, and hybrid metaheuristics. These flexible strategies utilize real-time decision-making, predictive analytics, and multi-objective optimization to improve resource utilization, reduce latency, and increase reliability in cloud, fog, and edge computing environments.

The [1] paper presents an in-depth analysis of dynamic Load Balancing techniques for web servers, concentrating on algorithms like Pending Job, IP Hash, Least Connection, and Weighted Least Connection, which are assessed using metrics such as response time, throughput, resource utilization, CPU usage, memory consumption, and disk I/O. The Pending Job algorithm, evaluated with request volumes ranging from 100 to 10,000, enhanced transactions per second by 23% and decreased overall processing time by 18.5% when compared to Round Robin. In a similar vein, IP Hash surpassed Least Connection by achieving lower latency (61 ms versus 73 ms), greater throughput (1.23 Mbps compared to 1.11 Mbps), and improved resource utilization (189.2 MB instead of 203.8 MB). The findings indicate that dynamic Load Balancing significantly enhances the performance of web servers, with Pending Job and IP Hash identified as the most effective methods, while additional optimization is suggested.

The [2] paper introduces an algorithm called Multiple Regression-Based Searching (MRBS) aimed at enhancing server selection and routing in Data Center Networks (DCN) through Software Defined Networking (SDN). It takes into account various factors such as load, response time, bandwidth, and server utilization, utilizing regression analysis to forecast response times and determine the most suitable server and route. By merging heuristic path searching with regression methods, MRBS facilitates dynamic load balancing and minimizes delays. It was implemented within a spine-leaf architecture using Mininet and the Floodlight controller, resulting in more than an 85% reduction in delays and notable enhancements in throughput, bandwidth utilization, and response times when compared to ECMP and MTSS. Additionally, MRBS efficiently copes with

abrupt message surges and unpredictable traffic, making it particularly effective for dynamic DCN scenarios.

The paper [3] examines Load Balancing based on Flow and Traffic Patterns within a distributed controller setup in Software Defined Networking (SDN), assessing factors such as response time, availability, and transaction rates. Flow-based balancing distinguishes between large and standard flows with adjustable thresholds, leading to a 94% reduction in response time and a 14% increase in transaction rates, making it particularly effective for demanding activities like media streaming. Conversely, traffic pattern-based balancing differentiates between TCP and UDP traffic, resulting in a 2.69% enhancement in availability and proving to be more suitable for web server applications. The research concludes that flow-based approaches outperform in terms of speed and efficiency, whereas traffic pattern-based methods offer improvements in availability and scalability.

The [4] paper evaluates AWS, Microsoft Azure, and Google Cloud Platform (GCP) by assessing their service offerings, pricing structures, and competitive approaches, taking into account elements such as hardware, maintenance, workforce, computing, storage, and distribution expenses. AWS stands out as the frontrunner thanks to its extensive range of services and worldwide reach, Azure is recognized for its hybrid cloud capabilities and strong integration with Microsoft products, while GCP is appreciated for its innovation and cost-effectiveness. The research underscores scalability, security, availability, and adaptable pricing as crucial factors. It concludes that AWS is ideal for large enterprises, Azure is advantageous for organizations dependent on Microsoft, and GCP is suitable for projects that prioritize cost and innovation.

The [5] paper introduces a Load Balancing algorithm designed for IaaS cloud environments aimed at enhancing resource distribution and task management while adhering to QoS requirements such as deadlines and completion times. It allocates workloads across virtual machines (VMs) and flexibly reassigns tasks to avoid SLA breaches. This approach, executed in CloudSim with configurations of 2 to 6 VMs and 10 to 40 tasks, achieved an average Resource Utilization

of 78%, surpassing the performance of the Dynamic LB Algorithm. Its ability to adapt to different workloads and larger tasks makes it well-suited for high QoS applications like media streaming and cloud services.

The [6] paper reviews load balancing algorithms used in Nginx, concentrating on IP_HASH, WRR, LEAST_CONN, and a new variant called NEW_HASH. By altering the hash function and employing quicksort with binary search for node selection, NEW_HASH minimizes collisions and enhances search efficiency. In tests conducted within a clustered Nginx-Keepalived configuration under different levels of concurrency, it exhibited up to a 94% reduction in response times, fewer access failures, and greater throughput compared to conventional methods. Nonetheless, NEW_HASH introduces memory overhead due to array operations, which impacts stability when subjected to extreme loads. The research recommends additional refinement to achieve a balance between performance and memory efficiency.

The [7] article presents a hybrid Load Balancing method for multipath routing networks by combining Genetic Algorithms (GA) with Particle Swarm Optimization (PSO). This integration improves resource distribution by dispersing traffic across multiple paths, thereby reducing congestion, latency, and packet loss. GA evolves solutions through generations, while PSO applies swarm intelligence for continuous refinement. Simulations under diverse traffic and topology scenarios showed improved throughput, QoS, and resource utilization compared to conventional methods. The study highlights its suitability for data-intensive applications and suggests exploring ACO and ML in future research.

The [8] paper examines various Load Balancing techniques in fog computing, motivated by the rising number of IoT devices and the necessity for effective resource management. It classifies approaches into static methods such as Round-Robin, Min-Min, and Tabu Search, as well as dynamic strategies that include traditional, nature-inspired, agent-based, real-time, and hybrid solutions, assessed using criteria like energy consumption, response time, costs, and resource utilization. A suggested framework distributes workloads among virtual machines to

enhance performance and lower energy usage, incorporating factors such as user requests and task loads, with the goal of minimizing latency, improving utilization, and efficiently managing IoT and real-time applications.

The [9] paper discusses Load Balancing in cloud computing, highlighting its importance in distributing workloads across different nodes to avoid bottlenecks and improve resource utilization. It assesses various methods using criteria such as throughput, fault tolerance, migration duration, response time, and scalability, analyzing static algorithms (Round Robin, Min-Min, Max-Min), dynamic algorithms (Honeybee Foraging, Ant Colony Optimization), and hybrid approaches that merge both for enhanced efficiency. Different workloads, from CPU-heavy to memory-heavy, are allocated among cloud nodes to boost performance, shorten execution times, and ensure equity. The research concludes that efficient Load Balancing improves resource distribution, reduces delays, and enhances the reliability and performance of cloud systems.

The [10] paper examines strategies for Load Balancing (LB) in Software-Defined Networking (SDN), concentrating on techniques applicable to both the data and control planes. It assesses various methods based on QoS metrics such as throughput, response time, packet loss, and resource utilization. Data plane LB encompasses server-based, link-based, static, dynamic, meta-heuristic (GA, ACO), and machine learning strategies (ANN, RL) to ensure efficient distribution of traffic, while control plane LB focuses on equalizing workloads over multiple controllers through hierarchical or horizontal configurations. Factors like traffic patterns, flow tables, and bandwidth usage are utilized to yield optimized routing, reduced latency, and enhanced fault tolerance. The research underscores the promise of machine learning in SDN LB but highlights the necessity for practical testing to enhance scalability and performance.

The [11] paper presents a strategy for Load Balancing in cloud computing based on Deep Reinforcement Learning (DRL), framing dynamic task scheduling as a Markov Decision Process (MDP). A DRL agent allocates workloads among virtual machines while evaluating metrics such as response time, server

utilization, and task success ratio as rewards. The results from simulations indicate that this method surpasses conventional approaches like Round Robin, Min-Min, and Genetic Algorithms, demonstrating enhanced adaptability, scalability, and real-time decision-making capabilities in fluctuating cloud environments.

The [12] paper introduces the Artificial Bee Colony (ABC) algorithm, which aims to enhance load balancing in cloud environments by simulating the foraging behavior of bees. In this model, tasks are considered as food sources while virtual machines act as foraging bees that evaluate the fitness of workloads for their distribution. This method minimizes makespan, decreases execution duration, and maximizes resource efficiency. When compared to Round Robin and Randomized allocations, the ABC method demonstrates improved CPU efficiency and reduced waiting times for virtual machines in the queue.

The [13] paper introduces a Load Balancing method for edge computing based on Q-learning, allowing edge devices to determine the best offloading strategies in environments where latency is critical. By spreading tasks across local edge nodes and centralized cloud servers, this model effectively decreases both congestion and energy consumption. In comparison to fixed and heuristic approaches, it achieves a latency reduction of 20–30% in real-time applications such as healthcare monitoring and smart surveillance.

The [14] study utilizes Ant Colony Optimization (ACO) for dynamic Load Balancing in distributed systems, with virtual ants leveraging pheromone trails to discover the best routes for task allocation. Simulations indicate that ACO surpasses deterministic techniques such as Min-Min and Max-Min by enhancing throughput, optimizing resource distribution, and minimizing migration time, demonstrating its effectiveness in dynamic settings.

A comparative study of Load Balancing algorithms reveals that each one performs better in particular situations: Enhanced Weighted Join enhances response times, PBLB reduces latency by directing traffic to the closest server, and Pending Jobs combined with IP Hash guarantees high throughput during significant loads. At the same time, DWLC

improves the utilization of processors and memory but has a complex structure, whereas Dynamic Feedback and DWRS offer scalable, adaptive optimization in real-time.

III. INFERENCES CAPTIVATED

The analyzed research presents a variety of methods aimed at enhancing web server performance during periods of high traffic. The EWMC algorithm minimizes response times and server delays by adaptively adjusting weights according to real-time usage. PBLB directs traffic towards the nearest geographical servers, outperforming the traditional Round Robin method in scenarios with distributed clients. To achieve scalability, DCH improves throughput by altering virtual nodes, although this adds a layer of configuration complexity. IP Hash offers a consistent distribution by associating client IP addresses with specific servers, thereby exceeding the Least Connections method in terms of latency and effectiveness. Dynamic feedback and polling techniques contribute to throughput optimization through real-time modifications. Algorithms such as Pending Job and DWLC emphasize increasing transaction volumes, while DWRS fosters greater parallelism by directing more requests to underused servers. Approximate Web Server Queuing lowers drop rates by strategically prioritizing incoming requests. Innovative techniques are incorporating intelligence, with DRL providing flexibility in cloud workloads, and bio-inspired methodologies like ABC and ACO enhancing resource efficiency. Hybrid frameworks such as WOA-BA and ML-driven Load Balancing appear promising for reducing latency and facilitating predictive automation.

IV. PROPOSED SYSTEM

The proposed system introduces a Web Server Load Balancer to address congestion, slow responses, scalability issues, and single points of failure. It uses Round Robin and dynamic Load Balancing based on server metrics, with health checks and failover for reliability. Scalability is ensured by adding servers as needed, while DRL supports adaptive decisions. Predictive models and hybrid methods like WOA-BA improve task allocation, reducing latency and energy consumption in fog-cloud environments.

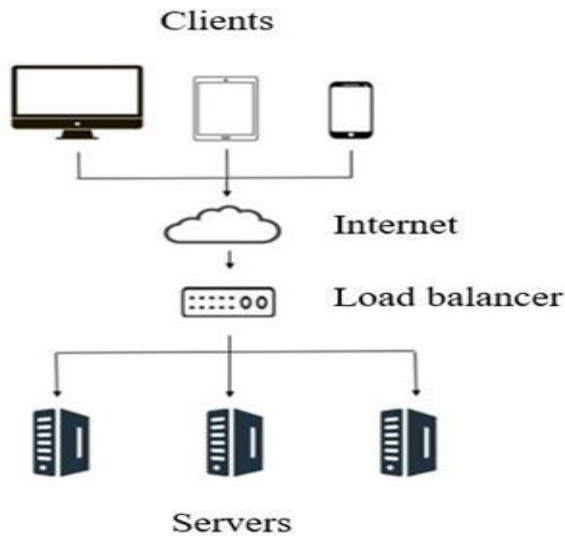


Figure 2 Proposed Load Balancer

A client-server architecture with a Load Balancer enables users on desktops, tablets, or smartphones to access services over the internet. The Load Balancer distributes requests evenly, preventing server overload and ensuring efficient resource usage. Scalability is achieved by adding or removing servers without downtime to handle traffic fluctuations. Tools like NGINX enhance load balancing, Docker supports containerized deployments, and PostgreSQL secures data and sessions. Wireshark aids network monitoring, while automated failover ensures uninterrupted service, enhancing performance, reliability, and business continuity (see fig 2).

V. CONCLUSION AND FUTURE ENHANCEMENTS

The review highlights the essential function of Load Balancing in boosting the performance of web server clusters by effectively allocating workloads and minimizing response times. Techniques such as Pending Job and IP Hash enhance resource utilization, while dynamic feedback and weighted polling enhance scalability and throughput. Future developments may include the use of AI and machine learning for forecasting load management, as well as energy-efficient strategies for sustainability. Integrating practical testing and stronger security measures like DDoS defense will guarantee more dependable, scalable, and resilient web server operations.

REFERENCES

- [1] M. Ibrahim, S. Y. Ameen, H. M. Yasin, N. Omar, and S. F. Kak, "Web server performance improvement using dynamic load balancing techniques: A review," *Asian Journal of Research in Computer Science*, vol. 10, no. 1, pp. 47–62, 2021. DOI: 10.9734/AJRCOS/2021/V10I130234.
- [2] G. S. Begam, M. Sangeetha, and N. R. Shanker, "Load balancing in DCN servers through SDN machine learning algorithm," *Arabian Journal for Science and Engineering*, vol. 47, no. 2, pp. 1423–1434, 2021. DOI: 10.1007/s13369-021-05911-1.
- [3] S. Prabakaran and R. Ramar, "Software defined network: Load balancing algorithm design and analysis," 2021.
- [4] T. Mufti, P. Mittal, and B. Gupta, "A review on Amazon Web Service (AWS), Microsoft Azure & Google Cloud Platform (GCP) services," in *Proc. Int. Conf. ICT for Digital, Smart, and Sustainable Development (ICIDSSD 2020)*, New Delhi, India, Feb. 2020, published Jan. 2021. DOI: 10.4108/eai.27-2-2020.2303255.
- [5] D. A. Shafiq, N. Z. Jhanjhi, A. Abdullah, and M. A. AlZain, "A load balancing algorithm for the data centres to optimize cloud computing applications," *IEEE Access*, Mar. 2021. DOI: 10.1109/ACCESS.2021.3065308.
- [6] C. Ma and Y. Chi, "Evaluation test and improvement of load balancing algorithms of Nginx," *IEEE Access*, vol. 10, 2022. DOI: 10.1109/ACCESS.2022.3146422.
- [7] S. Yoheswari, "Optimization techniques for load balancing in data-intensive applications using multipath routing networks," *Journal of Science and Technology Research*, vol. 5, no. 1, pp. 377–382, 2024.
- [8] M. Kaur and R. Aron, "A systematic study of load balancing approaches in the fog computing environment," 2021.
- [9] R. Kaur and P. Luthra, "Load balancing in cloud computing," Institutional Technical Report, SBS State Technical Campus, 2020.
- [10] M. Hamdana, E. Hassana, A. Abdelaziz, A. Elhigazi, B. Mohammed, and S. Khan, "A comprehensive survey of load balancing techniques in software-defined network,"

- Institutional Technical Report, University Technology Malaysia, 2021.
- [11] J. Guo, Y. Liang, *et al.*, “An efficient load balancing method based on deep reinforcement learning in cloud computing,” 2020.
 - [12] P. Saini and A. Kaur, “Load balancing in cloud computing using artificial bee colony algorithm,” *International Journal of Advanced Research in Computer Science and Software Engineering*, vol. 6, no. 4, 2016.
 - [13] K. Li, Y. Chen, *et al.*, “Load balancing in edge computing: A reinforcement learning approach,” *IEEE Internet of Things Journal*, 2021.
 - [14] N. Padhy and S. C. Satapathy, “Dynamic load balancing in distributed systems with ant colony optimization,” *Procedia Computer Science*, vol. 57, 2015.