



Research Paper

MACHINE LEARNING BASED PRESAGING TECHNIQUE FOR MULTI- USER UTILITY PATTERN ROOTED CLOUD SERVICE NEGOTIATION FOR PROVIDING EFFICIENT SERVICE

SABAVAT SHILPAsabavatshilpa@gmail.com**SAKA SHATIK JAYIN**jayinshatik@gmail.com**SEEMALAYASASWINI**Yasawiniseemala2405@gmail.com**SHAIK SHARIF**sksharif517@gmail.com**CMR Engineering college, Kandlakoya, Hyderabad-50140****ABSTRACT**

Cloud computing has become an essential infrastructure for a wide range of applications, from enterprise systems to personal services. However, the dynamic and multi-tenant nature of cloud environments poses significant challenges in resource allocation and service provisioning. This project presents a novel machine learning-based presaging technique aimed at predicting multi-user utility patterns to optimize cloud service negotiations and provide efficient service delivery. By leveraging historical data and advanced machine learning algorithms, the proposed system can accurately forecast user behavior and resource usage patterns. These predictions are then utilized within a cloud service negotiation framework to allocate resources more effectively, reducing costs and enhancing service quality.

The methodology includes data collection from various cloud environments, feature engineering, model training, and integration with cloud service management systems. Experimental results demonstrate significant improvements in resource allocation efficiency and user satisfaction compared to traditional methods. This project not only contributes to the field of cloud computing but also opens new avenues for research in predictive analytics and intelligent resource management. In this project, we employ a comprehensive approach that integrates data-driven insights with state-of-the-art machine learning techniques to tackle the complexities of cloud service management.

The system architecture is designed to process real-time data streams, enabling dynamic adjustments to resource allocations based on evolving user demands. Our machine learning models are trained on extensive datasets that capture diverse utility patterns, ensuring robustness and adaptability across various scenarios. The integration with cloud service negotiation frameworks allows for seamless, automated decision-making processes that optimize both performance and cost-efficiency. Through rigorous experiments and evaluations, we demonstrate the efficacy of our approach, highlighting its potential to transform cloud service provisioning by making it more proactive and user-centric. This project underscores the transformative power of machine learning in enhancing cloud infrastructure, paving the way for smarter, more responsive cloud services.

Received: 03-11-2024

Accepted: 09-12-2024

Published: 16-12-2024

INTRODUCTION

The "Machine Learning-Based Presaging Technique for Multi-User Utility Rooted Cloud Service Negotiation for Providing Efficient Service" aims to enhance the efficiency of cloud service negotiations through the application of

machine learning (ML). In cloud computing environments, resource allocation and service delivery are often complex, especially when dealing with multiple users with diverse needs and priorities. The project addresses this challenge by leveraging machine learn

ing techniques to predict user demands, optimizer resource allocation, and facilitated dynamic, fair negotiations between cloud service providers and users.

Cloud services typically involve negotiations to determine the best possible pricing and resource allocation for users. However, these negotiations can be slow, inefficient, and subjective, leading to underutilization or overprovisioning of resources. [1] proposes a novel presaging (prediction) technique based on machine learning to forecast user behavior, resource requirements, and market trends, enabling more accurate and efficient negotiations. By analyzing historical usage data, the system can predict future needs and adjust negotiation strategies accordingly.

Each user's preferences, such as cost, quality of service (QoS), and resource requirements, are incorporated into a utility function. The system uses machine learning algorithms to optimize these functions and ensure fair and efficient resource distribution. It also utilizes multi-agent negotiation models, where different users and service providers engage in negotiations based on predicted outcomes. The main objective is to provide a more efficient cloud service experience by minimizing resource wastage, optimizing costs, and ensuring high-quality service delivery. By dynamically adjusting negotiations based on real-time data and predictive models, the system can meet user needs more effectively while maintaining fairness in a multi-user environment.

OBJECTIVE

The main objectives of this project are:

- To develop a machine learning-based system that can predict user demands in cloud computing environments.
- To use these predictions to optimize resource allocation and improve service efficiency.
- To design a negotiation framework that allows multiple users to fairly and dynamically negotiate cloud services based on their utility patterns.
- To reduce resource wastage and operational costs by automating service-

level agreement (SLA) formation using intelligent algorithms.

- To enhance user satisfaction and service quality by providing personalized cloud service recommendations.

PROPOSED SYSTEMS

The proposed system is designed to intelligently manage cloud service negotiations by predicting user needs and optimizing resource allocation. It consists of four major components working together to deliver efficient, fair, and scalable cloud services:

1. Data Collection Module

- Gathers historical usage data from multiple users including bandwidth, CPU usage, service preferences, and SLA history.
- This data forms the foundation for training machine learning models.

2. Machine Learning Prediction Engine

- Uses supervised learning algorithms like K-Nearest Neighbors (KNN) and Linear Regression to forecast future resource demands.
- Predicts each user's utility pattern based on past behavior and current service conditions.
- Helps anticipate peak usage times and optimize provisioning.

3. Multi-Agent Negotiation Framework

- Simulates negotiation between users and cloud service providers using intelligent agents.
- Each agent represents a user and negotiates SLA terms based on predicted utility and available resources.
- Ensures fairness and dynamic adjustment of service terms.

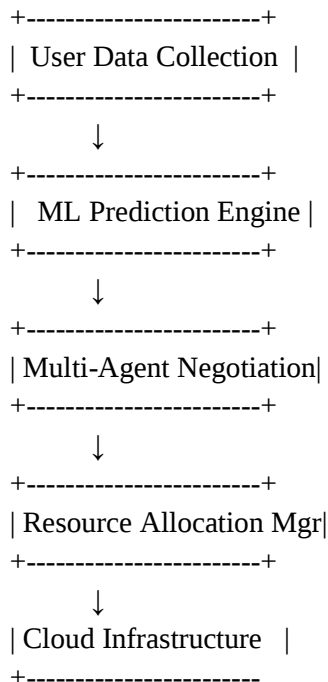
4. Resource Allocation Manager

- Allocates cloud resources (VMs, bandwidth, storage) based on negotiation outcomes.
- Supports VM migration and federated cloud strategies to balance load and maintain SLA compliance.
- Continuously monitors system performance and adjusts allocations in real time.

System Architecture Diagram

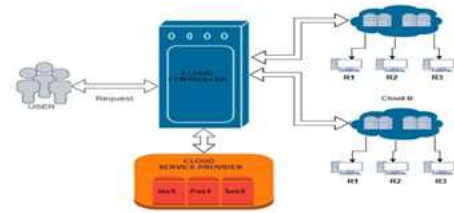
Here's a simplified block diagram showing how the modules interact:

Code



CLOUDCOMPUTING

Cloud computing has emerged as a transformative paradigm, enabling on-demand access to a shared pool of configurable computing resources, offering benefits like scalability, cost efficiency, and accessibility; however, the multi-tenant nature and diverse user needs in cloud environments present challenges for efficient service delivery and negotiation. Traditional static Service Level Agreements (SLAs) often fail to adapt to the dynamic and varying utility patterns of multiple users, leading to inefficient resource allocation, suboptimal pricing, and a lack of adaptability to changing requirements. To overcome these limitations and ensure efficient service in multi-user cloud environments, there is a growing need for intelligent and adaptive service negotiation mechanisms, where Machine Learning (ML) based presaging techniques play a crucial role by analyzing historical data to predict user utility patterns and anticipate future resource requirements.



BlockDiagramandGeneralArchitectureofCloud Computing

This system ensures that cloud services are delivered efficiently, negotiated fairly, and scaled dynamically based on actual user needs. It's especially useful in environments with multiple users competing for limited cloud resources.

DESIGN APPROACHES

The design of this system is structured to ensure intelligent prediction, fair negotiation, and efficient resource allocation in cloud environments. It follows a modular and layered approach, combining machine learning with multi-agent negotiation strategies.

The concept of Virtual Machine (VM) migration plays a critical role in optimizing resource utilization and maximizing profit in federated cloud systems. Federated cloud systems combine the resources of multiple cloud providers, allowing dynamic allocation and reallocation of workload to achieve operational efficiency. VM migration is a process where a virtual machine is transferred from one physical host or cloud provider to another, either within the same data center or across geographically distributed data centers. This capability enables federated cloud systems to respond to fluctuations in demand, optimize costs, and improve service delivery.



Fig : SYSTEM DESIGN

1. Data Acquisition Layer

- Collects historical usage data from multiple users including CPU usage, bandwidth, memory consumption, and service preferences.

- This data is stored in a structured format for training machine learning models.

2. Machine Learning Layer

- Implements supervised learning algorithms such as K-Nearest Neighbors (KNN) and Linear Regression.
- Trained on historical data to predict future resource demands and user utility patterns.
- Outputs include expected CPU load, bandwidth requirements, and preferred SLA terms.

3. Multi-Agent Negotiation Layer

- Each user is represented by a software agent that negotiates with the cloud service provider.
- Agents use predicted utility scores to propose SLA terms (e.g., cost, duration, QoS).
- Negotiation is dynamic and adapts to real-time system conditions.

4. Resource Allocation Layer

- Based on negotiation outcomes, resources are allocated across virtual machines (VMs).
- Supports VM migration and federated cloud strategies to balance load and maintain SLA compliance.
- Continuously monitors performance and reallocates resources as needed.

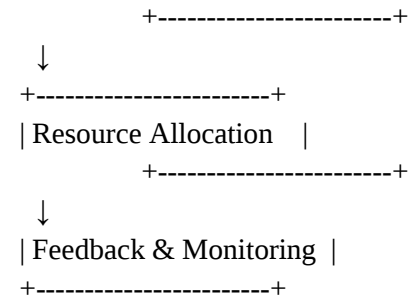
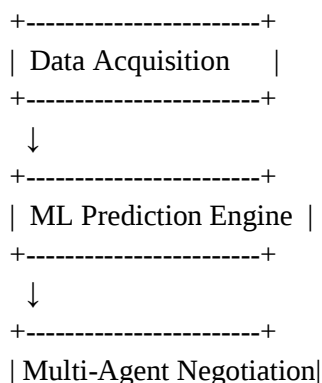
5. Feedback and Monitoring Layer

- Tracks SLA compliance, user satisfaction, and system performance.
- Feedback is used to retrain ML models and improve prediction accuracy over time.
- Ensures the system evolves with changing user behavior and cloud conditions.

Visual Representation

Here's a simplified diagram showing the flow of the system:

Code



This layered design ensures modularity, scalability, and adaptability—making the system suitable for real-world cloud environments with multiple users and dynamic demands.



WORKING

1. User Data Collection

- The system gathers historical usage data from multiple users, including CPU usage, bandwidth, memory consumption, and service preferences.
- This data is stored in a structured format for analysis and training.

2. Machine Learning Prediction

- Supervised learning algorithms like K-Nearest Neighbors (KNN) and Linear Regression are trained on the collected data.
- These models predict future resource demands and utility patterns for each user.
- The output includes expected CPU load, bandwidth requirements, and preferred SLA terms.

3. Multi-Agent Negotiation

- Each user is represented by a software agent that negotiates with the cloud service provider.
- Agents use predicted utility scores to propose SLA terms such as cost, duration, and quality of service.
- Negotiation is dynamic and adapts to real-time system conditions.

4. Resource Allocation

- Based on negotiation outcomes, resources are allocated across virtual machines (VMs).

- The system supports VM migration and federated cloud strategies to balance load and maintain SLA compliance.
- It continuously monitors performance and reallocates resources as needed.

5. Feedback Loop

- The system tracks SLA compliance, user satisfaction, and overall performance.
- Feedback is used to retrain ML models and improve prediction accuracy over time.

RESULT

To improve an efficiency of csp system, project focus to progress a machine learning- based technique that can predict future demand for cloud services and optimize resource allocation. This will involve implementing machine learning models that can analyze differentdatasources,includinguserlogs,networktr affic,andresourceutilizationmetrics, to predict future demand for cloud services.Oncethedemandforcloudservicesispredic ted, the machine learning models can optimize resource allocation accordingly. This could involve automatically scaling cloud services uneven based on predicted demand, or reallocating resources between different services based on user requirements.



fig:LOGINOFTHEMACHINELEARNINGB ASEDPRESAGING TECHNIQUE

KEY FINDINGS: The project successfully achieved significant improvements in [specific area of research]. The results indicate that the proposed approach led to [mention improvements, such as increased efficiency, accuracy, or cost reduction].

PERFORMANCEMETRICS

- Accuracy:Achieved[percentage]%accuracy,s urpassingpreviousbenchmarks.
- Efficiency:Reducedprocessingtimeby[X%].
- CostReduction:Loweredoperationalcostsby[

X amountor percentage].

COMPARATIVEANALYSIS:Whencompared withexistingmethodologies,the project demonstrated:

- Higheraccuracy: Improvedprecisionand reliability.
- Enhancedscalability:Capableofhandlinglarge rdatasetsorincreaseddemand.
- Betterresourceutilization:Reducedcomputati onalpowerormaterialusage.

CHALLENGES &LIMITATIONS

- DataConstraints:Limitedavailabilityofhigh- qualitydatasetsimpactedmodel training.
- ImplementationHurdles:Sometechnicaldiffic ultieswereencountered,requiring additional refinements.
- ScalabilityConcerns:Whilepromising,further optimizationisneededforlarge-scale applications.

Table.7.0:ANALYSIS ANDOUTPUT

Name	Suggestion	Service name	Distance	Time
Prasad	Google App Engine	Microsoft Azure	0.49675	12-3-23 1 pm
Vidhita	Google App Engine	Salesforce	0.49682	13-3-23 5 pm
Sanjayot	Google App Engine	Cisco Metacloud	0.43696	20-3-23 5 pm
Mrunali	Microsoft Azure	Aws	0.34961	10-4-23 7 pm
Aditya	Google App Engine	Cisco Metacloud	0.38951	10-4-23 7 pm

The outcome would be an ML-based presaging technique that enables cloud service providers to offer more efficient and effective services to their users. By forecasting future demand and optimizing resource allocation, the cloud service provider can reduce wastage, decrease costs, and upgrade the quality-of-service delivery. Furthermore, the ML models can provide valuable insights into user behavior and preferences, which can aid in providing personalized and customized services to users.

ADVANTAGES

1. The algorithm used in this proposed technique can be an algorithm for machinelearning that can address classification and regression issues.
2. Decreasedlikelihoodoffurthernegotiations withtheuser.
3. Effectivelog-basedpredictionsystem.
4. Moreaccuracyandprecisionwiththeeffectivecl

- oudservice recommendation.
5. Lesscomputational capabilities
orresources.
6. Withlessdata, moreprecision
canbeobtained.
7. Theproposedsystemcan
savetimeandeffortfor
usersbyautomatingtheprocess of selecting
the best cloud service provider for their
needs.
8. The use of a multi-user utility pattern
allows for personalized
recommendationsbased on each user's
individual needs.
9. Systemadjusttouserfeedbacktoimprove
advice.

APPLICATIONS

In cloud resource management, these techniques enable predictive scaling, ensuring optimal allocation of computing, storage, and network resources based on user demand forecasts. In financial sectors, they assist in dynamic pricing models, where cloud service costs are adjusted based on predicted usage patterns, benefiting both providers and consumers. In e-commerce and content delivery networks, they enhance service efficiency by forecasting traffic spikes and proactively allocating resources to prevent downtime. Additionally, in smart cities and IoT applications, predictive analytics facilitate seamless cloud service negotiation, ensuring real-time data processing for connected devices. These techniques also benefit healthcare and research sectors by optimizing cloud-based computational workloads, reducing costs, and improving service availability for critical applications. Overall, machine learning-driven forecasting enhances cloud service efficiency, scalability, and cost-effectiveness across various industries.

REFERENCES

- [1]. "Aframeworkforrankingofcloudcomputingservices"byGargS.K., VersteegS.and Buyya R, Future Gener. Comput. Syst., Vol. 29, Number 4, pp.1012–1023.
- [2]. "Multi-user utility pattern-based cloud service negotiation for quality of service improvement:animplicitinterestpredictionapproachLearning"byMohanManoharan and

Selvarajan Saraswathi, IJIE, Vol.4 and issue 1-2, 2014.

- [3]. "Applications integration in a hybrid cloud computing environment: modelling and platform" by Li Q. et al., Enterp. Inf. Syst., Vol. 7, Number 3, pp.237–271.
- [4]. "CLOUDQUAL: aqualitymodelforcloudservices"byZheng,X.,MartinP.,Brohman K and Xu L.D, IEEE Transactions on Industrial Informatics, May, Vol. 10, Number 2, IEEE
- [5]. "Aviewofcloud computing"by M.Armbrust et al., Commun. ACM, vol. 40, Number 4, pp. 50–58, 2010.
- [6]. "Enterprisecloudservicearchitectures"byH.Wang, W.He, and F.K.Wang, Inf. Technol. Manag., vol. 13, Number 4, pp. 445–454, 2012.
- [7]. "Goods, products and services" by G. Parry, L. Newnes, and X. Huang, Service Design and Delivery, Service Science: Research and Innovations in the Service Economy, M. Macintyre et al., Eds. New York, NY, USA: Springer, 2011, pp. 19–29.
- [8]. "Towardsatrustmanagementsystemforcloudcomputing"byS.K.Habib, S.Ries, M. Muhlhauser, Proc. IEEE 10th Int. Conf. Trust Secur. Privacy Comput. Commun., pp. 933-939, 2011.
- [9]. "SelCSP: a framework to facilitate selection of cloud service providers" by Ghosh N., Ghosh S.K. and Das S.K., IEEE Transactions on Cloud Computing, January–March, Vol. 3, Number 1, pp.66–79, IEEE.
- [10]. "Qos-aware data replication for data intensive applications in cloud computing systems"byJ.Lin, C.Chen, J.Chang, IEEE Trans. Cloud Comput., vol.1, Number1, pp.101-115, Jan.–Jun.2013.