



UNIVERSIDADE FEDERAL RURAL DO SEMI-ÁRIDO
PRÓ-REITORIA DE GRADUAÇÃO
DEPARTAMENTO DE ENGENHARIAS E TECNOLOGIAS
CURSO DE ENGENHARIA DE COMPUTAÇÃO

MATHEUS DUARTE DA SILVA

**UM ESTUDO SOBRE A MIGRAÇÃO DE UM SISTEMA DE APOIO À
APRENDIZAGEM DO AMBIENTE *ON-PREMISE* PARA UMA NUVEM
PÚBLICA DA AWS.**

PAU DOS FERROS-RN

2023

MATHEUS DUARTE DA SILVA

**UM ESTUDO SOBRE A MIGRAÇÃO DE UM SISTEMA DE APOIO À
APRENDIZAGEM DO AMBIENTE *ON-PREMISE* PARA UMA NUVEM
PÚBLICA DA AWS.**

Trabalho de Conclusão de Curso apresentado
ao curso de Engenharia de Computação da
Universidade Federal Rural do Semi-Árido
(UFERSA), como parte dos requisitos para
obtenção do Título de Engenheiro de
Computação.

Orientador: Prof. Dr. Walber José Adriano
Silva

PAU DOS FERROS-RN

2023

© Todos os direitos estão reservados a Universidade Federal Rural do Semi-Árido. O conteúdo desta obra é de inteira responsabilidade do (a) autor (a), sendo o mesmo, passível de sanções administrativas ou penais, caso sejam infringidas as leis que regulamentam a Propriedade Intelectual, respectivamente, Patentes: Lei nº 9.279/1996 e Direitos Autorais: Lei nº 9.610/1998. O conteúdo desta obra tomar-se-á de domínio público após a data de defesa e homologação da sua respectiva ata. A mesma poderá servir de base literária para novas pesquisas, desde que a obra e seu (a) respectivo (a) autor (a) sejam devidamente citados e mencionados os seus créditos bibliográficos.

S586e Silva, Matheus Duarte da.
Um estudo sobre a migração de um sistema de
apoio à aprendizagem do ambiente on-premise para
uma nuvem pública da AWS / Matheus Duarte da
Silva. - 2023.
47 f. : il.

Orientador: Walber José Adriano Silva.
Monografia (graduação) - Universidade Federal
Rural do Semi-árido, Curso de Engenharia de
Computação, 2023.

1. Arquitetura. 2. AWS. 3. Custos. 4. Moodle.
5. Nuvem. I. Silva, Walber José Adriano, orient.
II. Título.

Ficha catalográfica elaborada por sistema gerador automático em conformidade
com AACR2 e os dados fornecidos pelo(a) autor(a).

Biblioteca Campus Pau dos Ferros
Bibliotecária: Erlanda Maria Lopes da Silva
CRB: 15/966

O serviço de Geração Automática de Ficha Catalográfica para Trabalhos de Conclusão de Curso (TCC's) foi desenvolvido pelo Instituto de Ciências Matemáticas e de Computação da Universidade de São Paulo (USP) e gentilmente cedido para o Sistema de Bibliotecas da Universidade Federal Rural do Semi-Árido (SISBI-UFERSA), sendo customizado pela Superintendência de Tecnologia da Informação e Comunicação (SUTIC) sob orientação dos bibliotecários da instituição para ser adaptado às necessidades dos alunos dos Cursos de Graduação e Programas de Pós-Graduação da Universidade.

MATHEUS DUARTE DA SILVA

**UM ESTUDO SOBRE A MIGRAÇÃO DE UM SISTEMA DE APOIO À
APRENDIZAGEM DO AMBIENTE *ON-PREMISE* PARA UMA NUVEM
PÚBLICA DA AWS.**

Trabalho de Conclusão de Curso apresentado
ao curso de Engenharia de Computação da
Universidade Federal Rural do Semi-Árido
(UFERSA), como parte dos requisitos para
obtenção do Título de Engenheiro de
Computação.

Defendida em: 13 / 10 / 2023.

BANCA EXAMINADORA

Prof. Dr. Walber José Adriano Silva (UFERSA)
Presidente

Prof. Dr. João Batista Borges Neto (UFRN)
Membro Examinador

Prof. Dr. Reudismam Rolim de Sousa (UFERSA)
Membro Examinador

AGRADECIMENTOS

Agradeço a Deus por me guiar nos meus caminhos e a São Jorge por me proteger.

Aos meus pais Antônia Hiudézia e Evanaldo José por ter me dado o dom da vida, por ter me orientado em toda a minha trajetória e por todo o apoio necessário para chegar até aqui.

Agradeço ao Prof. Walber, por ter acreditado em mim e me orientado durante a realização do projeto, ter me apresentado o mundo da infraestrutura, esse foi um fator importante na minha carreira profissional.

Aos professores, que fizeram parte da minha trajetória, Segundo, Cecílio, Adller, Pedro, Marco, vocês foram imprescindíveis para me tornar o profissional que eu sou hoje.

Ao meu primo Bruno Emerson, pelos conselhos, pelo ponto de apoio e toda ajuda necessária. Agradecer a minha irmã Sara Duarte, por todos os ensinamentos que me foi passado, apesar de ser mais nova.

Aos meus colegas Kaio Alencar, Kauann Jacome, Vítor Gomes, Lucas Abrantes e Augusto Vinicius, vocês me mostraram, que era possível sempre me ajudando com as disciplinas.

Aos meus amigos de treino Yosell, Gersival e Heris. A disciplina que construí treinando em muitos momentos foram por causa de vocês.

E, por último, mas não menos importante, expresso minha profunda gratidão à minha namorada, Vitória Gonçalves, que sempre me incentivou, serviu de suporte e me ajudou quando eu mais precisei, dentro e fora da universidade, palavras são poucas para descrever o quanto sou grato a você.

“Se eu vi mais longe, foi por que estava
em ombros de gigantes.”

Isaac Newton

RESUMO

O crescimento tecnológico vem impactando diretamente na forma como as empresas gerenciam seus dados, à medida que elas crescem os custos com infraestrutura aumentam, fazendo com que a migração para nuvem se torne tendência. Esta, por sua vez, possui características, como o serviço sobre demanda, provisionamento de recursos automático, amplo acesso à rede, agrupamento de recursos e elasticidade nos serviços. Sobre os seus modelos de implantação, destacam-se as nuvens públicas, privadas, híbridas e comunitárias e tem como modelos de serviços as *Infrastructure as a Service* (IaaS), *Platform as a Service* (PaaS) e *Software as a Service* (SaaS). Atualmente, o sistema Moodle vem sendo bem aceito em instituições de ensino, contribuindo assim para educação a distância com um Ambiente Virtual de Aprendizado, isso ocorre devido a sua ampla usabilidade e customização. Diante disso, faz necessário a possível migração do sistema Moodle implantado *on-premises* para nuvem pública. Foi usada a AWS como base para uma proposta de arquitetura e seus custos detalhados. A partir da proposta e implementação da arquitetura, foi realizado alguns testes de carga com o objetivo de avaliar o desempenho da infraestrutura e comparação dos custos sugeridos na infraestrutura e o modelo utilizado na UFERSA. Com base nos resultados, é notável uma melhora em se fazer o uso da nuvem, pois além do controle sobre os gastos, serviços são provisionados de forma automática de acordo com a carga de trabalho, fazendo com que a migração para ambiente de nuvem se torne vantajosa para a UFERSA.

Palavras-chave: arquitetura; AWS; custos; moodle; nuvem.

ABSTRACT

Technological Growth has directly impacted the way companies manage their data. As they grow, infrastructure costs increase, making migration to the cloud a trend. This, in turn, has characteristics such as on-demand service, automatic resource provisioning, broad network access, resource pooling and service elasticity. Regarding its deployment models, public, private, hybrid and community clouds stand out and its service models include Infrastructure as a Service (IaaS), Platform as a Service (PaaS) and Software as a Service (SaaS). Currently, the Moodle system has been well accepted in educational institutions, thus contributing to distance education with a Virtual Learning Environment, this is due to its wide usability and customization. Therefore, it is necessary to possibly migrate the Moodle system implemented on-premises to the public cloud. AWS was used as the basis for an architectural proposal and its detailed costs. Based on the proposal and implementation of the architecture, some load tests were carried out with the aim of evaluating the performance of the infrastructure and comparing the costs suggested in the infrastructure and the model used at UFERSA. Based on the results, there is a notable improvement in using the cloud, as in addition to controlling expenses, services are automatically provisioned according to the workload, making migration to a cloud environment a advantageous for UFERSA.

Keywords: architecture; AWS; costs; moodle; cloud.

LISTA DE FIGURAS

Figura 1: Comparação entre os modelos de serviços e o modelo <i>on-premisses</i>	19
Figura 2: Estrutura moodle.....	21
Figura 3: Arquitetura na nuvem.	23
Figura 4: Região, azs, <i>data centers</i>	24
Figura 5 : Teste de desempenho.....	32
Figura 6: Consumo de cpu das instâncias durante a execução dos testes.	35
Figura 7: Utilização de cpu durante 500 requisições, cenário 2.....	37
Figura 8: Servidores durante os testes do terceiro cenário.....	38
Figura 9: Uso da cpu com 5000 requisições.	39
Figura 10: Teste de performance para uma amostra	40

LISTA DE TABELAS

Tabela 1: Recursos utilizados	24
Tabela 2: Custo de vpc	27
Tabela 3: Custos do <i>route 53</i>	28
Tabela 4: Custos do <i>cloudfront</i>	28
Tabela 5: Custos do <i>load balancer</i>	29
Tabela 6: Custos com o rds	30
Tabela 7: Custo das instâncias ec2	30
Tabela 8: Custos mensais da arquitetura	31
Tabela 9: Teste de carga do sistema	34
Tabela 10: Amostra 1º cenário	36
Tabela 11: Tempo de resposta para o 2º cenário.	37
Tabela 12: Resultados dos testes para o terceiro cenário	38
Tabela 13: Tempo do servidor para 5000 requisições.....	39
Tabela 14: Aumento do número de máquinas ec2	40

LISTA DE ABREVIATURAS E SIGLAS

AMI	<i>Amazon Machine Image</i>
API	<i>Application Programming Interface</i>
AVA	Ambientes Virtuais de Aprendizado
AWS	<i>Amazon Web Service</i>
AZ	<i>Availability Zone</i>
EBS	<i>Elastic Block Store</i>
CAPEX	<i>Capital Expenditure</i>
CDN	<i>Content Delivery Network</i>
EaD	Educação à Distância
EC2	<i>Elastic Compute Cloud</i>
GCP	<i>Google Cloud Platform</i>
HTTP	<i>Hypertext Transfer Protocol</i>
IaaS	<i>Infrastructure as a Service</i>
LMS	<i>Learning Management System</i>
NAT	<i>Network Address Translation</i>
PaaS	<i>Platform as a Service</i>
PDA	<i>Personal Digital Assistant</i>
RDS	<i>Relational Database Service</i>
SaaS	<i>Software as a Service</i>
TICs	Tecnologias da Informação e Computação
TI	Tecnologia da Informação
UFERSA	Universidade Federal Rural do Semi-Árido
VPC	<i>Virtual Private Cloud</i>
VU	<i>Virtual User</i>

SUMÁRIO

1. INTRODUÇÃO	13
1.1 Problematização	13
1.2 Justificativa	14
1.3 Metodologia	15
1.4 Objetivos	15
1.4.1 Objetivo Geral	15
1.4.2 Objetivos Específicos	15
2. FUNDAMENTAÇÃO TEÓRICA	16
2.1 Características da Computação em Nuvem	16
2.2 Tipos de Nuvem	17
2.2.1 Nuvem Pública	17
2.2.2 Nuvem Privada	18
2.2.3 Nuvem Híbrida	18
2.2.4 Nuvem Comunitária	18
2.3 Modelos de Serviço em Nuvem	18
2.3.1 IaaS	19
2.3.2 PaaS	20
2.3.3 SaaS	20
2.4 Sistema de apoio à aprendizagem na Educação à Distância	20
2.5 Arquitetura Moodle	21
3. PROPOSTA DE MIGRAÇÃO	22
3.1 Descrição da arquitetura	22
3.2 Precificação dos recursos	26
3.2.1 Estimativa para os custos	26
3.2.2 Cálculo dos custos	26
4. RESULTADOS E DISCUSSÕES	32

4.1 Configurações.....	32
4.2 Testes de carga	34
4.2.1 Configuração dos cenários	34
4.2.2 Cenário 1	35
4.2.3 Cenário 2	36
4.2.4 Cenário 3	38
4.2.5 Cenário 4	39
4.3 Comparações	41
5. CONCLUSÃO	42
6. REFERÊNCIAS	44

1. INTRODUÇÃO

O crescimento tecnológico vem impactando diretamente a maneira como empresas gerenciam suas aplicações e seus dados. Mediante a isso,

A constante evolução e busca por sistemas de informação que atendam melhor às necessidades de uma determinada organização em expansão é extremamente importante para o seu desenvolvimento (Vanderlei, Santana, Carosia, 2021 p.2).

Nesse contexto, à medida que a empresa cresce, é comum o aumento do *CAPital EXpenditure* (CAPEX), ou seja, as despesas relacionadas à aquisição de equipamentos, afim de garantir uma maior capacidade de processamento, memória e armazenamento. Além disso, a ampliação da infraestrutura da empresa pode ocasionar em uma maior complexidade nos seus servidores, pois terá que ser pensado em soluções que sejam redundantes a falhas e tenham boa eficiência nas cargas de trabalho, bem como a implantação de sistemas de *backup* e recuperação de desastres, para garantir uma segurança adequada dos dados (Costa; Santos, 2021; Queiroz et al., 2021).

Diante disso, a migração de um servidor local para a nuvem tem se tornado uma tendência crescente. Segundo Armbrust (2009, p. 4), “*Cloud Computing* refere-se tanto aos aplicativos fornecidos como serviços pela Internet quanto ao *hardware* e *software* de sistemas nos *data centers* que oferecem esses serviços”. Com isso, este trabalho terá como objetivo, investigar a viabilidade de uma possível migração do sistema da Universidade Federal Rural do Semi Árido (UFERSA), para a nuvem pública, utilizando os serviços da *Amazon Web Services* (AWS), provedor de computação em nuvem.

1.1 Problematização

O modelo de infraestrutura tecnológica implementado no Ambientes Virtuais de Aprendizado (AVA) da UFERSA segue o desenvolvimento de *hardware* e *software* em ambiente local (*on-premise*). Neste modelo de infraestrutura, quando o crescimento da demanda por recursos computacionais é exigido, a capacidade da infraestrutura limita o atendimento dessa demanda. Além disso, há uma dificuldade de provisionar recursos computacionais de maneira antecipada, visando prover a demanda futura.

Usualmente, a consequência natural é decidir por super provisionar a infraestrutura, considerando a necessidade máxima para uma determinada demanda, visando garantir o atendimento da demanda com a capacidade. Isso culmina com a

aquisição de equipamentos e a subutilização dos mesmos, gerando má alocação de recurso público.

Um modelo alternativo ao do *on-premise* é a computação em nuvem. Nela, os recursos computacionais podem ser provisionados conforme a demanda por poder computacional e serem liberados quando a demanda diminui. Além disso, o custo por esses recursos são proporcionais à alocação requerida por eles. Essa característica de computação como utilidade aos moldes do pagamento de energia elétrica, Internet e água residencial tem o potencial de reduzir custos e alavancar oportunidades não exploradas. Por exemplo, em um cenário onde há o aumento de discentes matriculados na UFERSA, como esse crescimento irá impactar na infraestrutura que fornece suporte aos sistemas de apoio ao aprendizado da instituição?

Diante desse contexto tem-se como problematização a investigação dos desafios de migração da infraestrutura computacional de um sistema de apoio ao aprendizado da UFERSA para um modelo de computação em nuvem pública.

1.2 Justificativa

A proposta de investigar a utilização da nuvem pública para ser adotada no sistema de apoio ao aprendizado ocorreu através do relato de problemas de disponibilidade no sistema da UFERSA. Há períodos do ano, quando o acesso aos sistemas institucionais é muito demandado, em que a capacidade dos recursos computacionais acaba demonstrando suas limitações. São exemplos desses relatos: comportamento inesperado dos sistemas, indisponibilidade de sites e serviços *on-line* ofertados pela universidade. Uma possível causa desses problemas é o compartilhamento da infraestrutura computacional para prover os serviços de Tecnologia da Informação (TI) da instituição (opção pelo modelo *on-premise*).

O modelo de computação em nuvem pública possui o potencial de sanar esse problema de limitação da capacidade do *on-premise*. Contudo, elementos como custo, segurança, confiabilidade, integração de serviços, etc, podem ser fatores impeditivos para essa mudança de modelo. Portanto, essa pesquisa buscará trazer luz a essas questões através de uma investigação sobre a migração para esse modelo de computação em nuvem ao invés do *on-premise*.

1.3 Metodologia

Este trabalho será realizado através de um estudo de caso, em que adotaremos uma abordagem de cunho quantitativa, conforme os seguintes procedimentos:

- Revisar a bibliografia dos principais conceitos de nuvem, elencando a viabilidade da sua utilização;
- Propor uma arquitetura de infraestrutura para implantação do sistema de apoio ao aprendizado da UFERSA;
- Propor uma estimativa de custo da arquitetura a ser implantada no ambiente de computação em nuvem;
- Realizar medições na proposta e realizar estimativas de desempenho para a implementação da arquitetura.

1.4 Objetivos

Essa seção descreve os objetivos que espera-se atender com a conclusão do trabalho.

1.4.1 Objetivo Geral

O objetivo da pesquisa é realizar um estudo sobre a migração de um sistema de apoio à aprendizagem do ambiente *on-premise* para nuvem pública utilizando os serviços da AWS.

1.4.2 Objetivos Específicos

- Realizar um estudo bibliográfico sobre os conceitos de computação em nuvem e trabalhos relacionados;
- Desenvolver uma arquitetura para um sistema de apoio ao aprendizado;
- Levantar os custos da infraestrutura proposta;
- Implantar a arquitetura e comparar custos e desempenho dessa arquitetura.

2. FUNDAMENTAÇÃO TEÓRICA

O termo *on-premise* é utilizado para designar uma infraestrutura hospedada e mantida localmente dentro de uma organização, sendo ela responsável pelo gerenciamento dos componentes de *hardware*, *software* e redes. Para esse tipo de modelo, dificuldades sobre ampliação, escalonamento de tarefas e otimização em recursos, ainda são consideradas tarefas complexas que envolvem o aumento do CAPEX da organização (Carissimi, 2015; Nascimento; Santos, 2021).

Diante desse cenário, o modelo de infraestrutura *on-premise* está sendo gradativamente substituído por outros modelos, como a computação em nuvem. Segundo Erl, Puttini e Mahmood (2013, p. 33), “Uma nuvem refere-se a um ambiente distinto projetado com o propósito de fornecer recursos de TI escaláveis e mensuráveis remotamente”. Esse modelo provê recursos computacionais, em que as organizações pagam pelos serviços consumidos, com a possibilidade de reduzir os custos com manutenções ou investimentos com servidores físicos (Bandeira, 2022; Carissimi, 2015).

2.1 Características da Computação em Nuvem

O modelo de computação em nuvem possui características que oferecem vantagens ao contratar recursos na forma de serviços. Dentre eles, destacam-se os serviços sob demanda, amplo acesso à rede, agrupamento de recursos, elasticidade e serviços mensuráveis. Nos serviços sob demanda, os recursos são provisionados automaticamente de acordo com a carga de trabalho. Segundo Zuffo et al. (2013, p.13), “Esses recursos computacionais podem ser capacidade computacional, armazenamento, impressão ou recursos virtualizados, como servidores, computadores *desktop* e redes virtuais” (Simmon et al., 2018; Souza; Moreira; Machado, 2021; Zuffo et al., 2013).

No que diz respeito ao amplo acesso à rede, segundo Machado, Moreira e Souza (2021, p. 6), “Recursos são disponibilizados por meio da rede e acessados através de mecanismos padronizados que possibilitam o uso por plataformas *thin* ou *thin client*, tais como celulares, *laptops* e PDAs”. Em outras palavras, a rede é essencial para a computação em nuvem pois permite que qualquer dispositivo que esteja interligado a ela faça o uso de seus recursos (Simmon et al., 2018; Souza; Moreira; Machado, 2021; Zuffo et al., 2013).

Outra característica é o agrupamento de recursos, nela a gestão e organização é feita em um *pool* dos recursos físicos e virtuais a fim de servir múltiplos inquilinos.

Segundo Zuffo et al. (2013 p. 13), “Esses recursos podem ser dinamicamente alocados de acordo com as necessidades de desempenho, escala e demanda”. Vale ressaltar, que a localização física dos servidores é omitida, dessa forma o usuário não sabe onde seus dados estão sendo armazenados (Simmon et al., 2018; Souza; Moreira; Machado, 2021; Zuffo et al., 2013).

A rápida elasticidade é a característica na qual os recursos são provisionados ou excluídos de forma rápida e automática à medida que a demanda ao serviço muda. Segundo Mell e Grace (2011 p. 2) “Para o consumidor, as capacidades disponíveis para provisionamento muitas vezes parecem ilimitadas e podem ser apropriadas em qualquer quantidade a qualquer momento”. Cabe destacar, que a elasticidade dos recursos influencia diretamente nos custos da organização (Simmon et al., 2018; Souza; Moreira; Machado, 2021; Zuffo et al., 2013).

Como última característica tem-se os serviços mensuráveis, ele resalta que os provedores em nuvem possuem um sistema de faturamento, com objetivo de monitorar e mensurar, apresentando de maneira clara e em tempo real todos os custos dos recursos alocados a respectiva infraestrutura para os usuários. É importante destacar, que a exposição desses dados é importante para a avaliação, gestão e alocação de forma eficiente e apropriada a organização que eventualmente irá consumir os recursos (Borges; Silva, 2019; Simmon et al., 2018; Souza; Moreira; Machado, 2021; Zuffo et al. 2013).

2.2 Tipos de Nuvem

Os modelos de implantação e operação em nuvem são definidos de acordo com o propósito e a necessidade considerada por cada organização. Sendo os principais a nuvem pública, privada, híbrida e comunitária (Mell; Grance, 2011).

2.2.1 Nuvem Pública

No modelo de nuvem pública os serviços e recursos são de propriedades dos provedores de nuvem, nas quais estes se encarregam de gerir, manter e administrar todos os componentes físicos dos *datacenters*. Disponibilizando recursos computacionais de forma dinâmica e granular através da Internet, permitindo aos clientes pagar pela utilização dos seus serviços e a segurança na nuvem de suas aplicações (Mather; Kumaraswamy; Latif, 2019; Mell; Grance, 2011).

2.2.2 Nuvem Privada

Segundo Mather, Kumaraswamy e Latif (2009, p. 23), “Nuvens privadas e nuvens internas são termos usados para descrever ofertas que emulam a computação em nuvem em redes privadas”. Nesse modelo, através de técnicas de virtualização, os recursos são provisionados e geridos em uma implantação local pela própria organização ou disponibilizados por terceiros, tendo uso dedicados exclusivamente à empresa solicitante, sendo hospedando internamente (*on-premises*) ou por provedor de serviços gerenciados (*off-premises*) (Amazon Web Service, 2023; Harding, 2011; Mell; Grance, 2011).

2.2.3 Nuvem Híbrida

No modelo de nuvem híbrida, faz-se o uso de dois ou mais tipos de nuvem no intuito de obter características específicas de ambos os ambientes. Ou seja, por meio da nuvem híbrida, empresas têm a capacidade de operar aplicações não críticas em ambientes de nuvem pública e utilizar da privada para alocação de aplicações sensíveis e informações confidenciais. Vale ressaltar que esse modelo, apesar de fazer o uso de distintas entidades, utilizam de tecnologias que padronizam a troca de informações, na qual, segundo Harding (2011, p. 20), “uma nuvem híbrida pode ser coordenada por um intermediário que federaliza dados, identidade, segurança e outros detalhes” (Harding, 2011; Mather; Kumaraswamy; Latif, 2019; Mell; Grance, 2011).

2.2.4 Nuvem Comunitária

Diferente da nuvem privada, que fornece recursos a uma única organização, esse modelo é composto por várias organizações, que têm preocupações compartilhadas no que diz respeito à segurança, privacidade e políticas de conformidade, além do compartilhamento de recursos e da troca de informações por meio de conexões entre as organizações associadas. Cabe destacar que segundo Harding (2011, p. 20), “Embora o ônus de criar e gerenciar a nuvem seja retirado dos ombros de cada organização membro, isso deve ser feito pela comunidade como um todo” (Mather; Kumaraswamy; Latif, 2019; Mell; Grance, 2011).

2.3 Modelos de Serviço em Nuvem

Segundo Zuffo et al. (2013), “A prestação de serviços convenientes *on-line*, sob demanda, em qualquer lugar e com oferta de recursos ilimitados, passa a ser o grande foco de atenção da computação em nuvem”. Diante disso, os modelos de serviços de

nuvem são definidos pelo tipo de controle que o usuário tem sobre a infraestrutura e o software fornecido pela nuvem. Sendo os principais *IaaS*, *PaaS*, *SaaS*. A Figura 1 apresenta os tipos de serviços em comparação com a infraestrutura *on-premises*. “Uma infraestrutura de TI *on-premise* representa a maior responsabilidade que você terá como usuário e gestor” (Red Hat, 2023).

Figura 1: Comparação entre os modelos de serviços e o modelo *on-premises*.

On-Premises	IaaS	PaaS	SaaS
Aplicação	Aplicação	Aplicação	Aplicação
Segurança	Segurança	Segurança	Segurança
Banco de Dados	Banco de Dados	Banco de Dados	Banco de Dados
Sistema Operacional	Sistema Operacional	Sistema Operacional	Sistema Operacional
Virtualização	Virtualização	Virtualização	Virtualização
Armazenamento	Armazenamento	Armazenamento	Armazenamento
Rede	Rede	Rede	Rede
Data Centers	Data Centers	Data Centers	Data Centers
Responsabilidade/gerenciado pelo usuário			
Responsabilidade/gerenciado pelo provedor			

Fonte: Produzida pelo autor (2023).

2.3.1 IaaS

A *Infrastructure as a Service* (IaaS), tem como principal objetivo, fornecer ao usuário controle sobre parâmetros do ambiente virtualizado na nuvem, ou seja, a administração dos recursos provisionados como sistema operacional, implantação da aplicação, segurança dos recursos na nuvem, armazenamento de dados, interações do sistema e criação de redes virtuais é de total responsabilidade de quem irá utilizar os serviços. Segundo Machado, Moreira e Souza (2021, p.11) “O usuário não administra ou controla a infraestrutura da nuvem, mas tem controle sobre os sistemas operacionais, armazenamento e aplicativos implantados, e, eventualmente, seleciona componentes de rede, tais como *firewalls*”. Vale ressaltar que o acesso, administração e manutenção dos recursos físicos é abstraído dos usuários, estes por sua vez é de total responsabilidade dos provedores da nuvem (Carissimi, 2015; Harding, 2011; Red Hat, 2023; Souza; Moreira; Machado, 2021; Sosinsky, 2010; Zuffo et al. 2013).

2.3.2 PaaS

Nos modelos *Platform as a Service* (PaaS) é ofertado a plataforma completa para o desenvolvimento e implantação de aplicativos, na qual os provedores oferecem ambientes pré-configurados incluindo toda a parte de infraestrutura, tornando a implantação de serviço mais rápida e fácil. Nesse tipo de modelo, o nível de controle dos recursos por parte dos usuários é restringido e delegado ao provedor. Na qual, segundo Sosinsky (2010, p.70), “Em um modelo PaaS, os clientes podem interagir com o *software* para inserir e recuperar dados, realizar ações, obter resultados e, na medida em que o fornecedor permite, personalizar a plataforma envolvida” (Carissimi, 2015; Harding, 2011; Red Hat, 2023; Souza; Moreira; Machado, 2021; Sosinsky, 2010; Zuffo et al. 2013).

2.3.3 SaaS

Para os *Software as a Service* (SaaS), os provedores de nuvem proporcionam sistemas já configurados. Nesse tipo de modelo, os usuários não fazem o controle ou administração do aplicativo, deixando todas essas responsabilidades ao provedor de nuvem. Segundo Red Hat (2023), “O usuário precisa apenas se conectar à aplicação por meio de um painel de controle ou uma API. Não é necessário instalar nenhum *software* em máquinas individuais” (Carissimi, 2015; Harding, 2011; Red Hat, 2023; Souza; Moreira; Machado, 2021; Sosinsky, 2010; Zuffo et al. 2013).

2.4 Sistema de apoio à aprendizagem na Educação à Distância

A evolução das Tecnologias da Informação e Computação (TICs) vem contribuindo diretamente na Educação a Distância (EaD), na qual segundo Lapa e Belloni (2012), “[...] as TIC apresentam recursos com potencial transformador para a educação, seja ela presencial ou pela modalidade à distância.” Diante disso, entende-se como EaD um método de construção do conhecimento que não necessita da interação física direta dos participantes (Burnham et al., 2009; Nona; Alves, 2003).

Dessa forma, há uma busca por Ambientes Virtuais de Aprendizagem (AVA) que auxiliem na EaD. Esses ambientes também são conhecidos como LMS (*Learning Management System*), na qual os ambientes virtuais trata-se de um espaço onde há uma interação do ser humano e objetos técnicos, no intuito de construir conhecimento. (Santos; Okada, 2003). Diante disso, mediante o ponto de vista pedagógico, o AVA é uma sala de aula *online* com o propósito de ser um ambiente de ensino/aprendizagem, operando de forma abrangente e significativa com o auxílio de *softwares* que colaboram com a criação

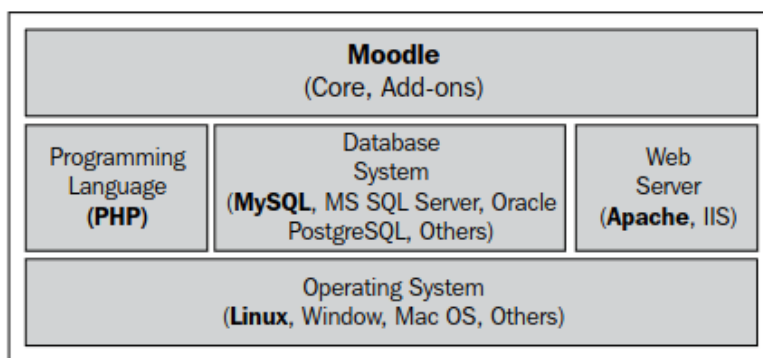
e execução de cursos para Internet como instrumento de gestão dos conteúdos educacionais (Alves, 2009; Vasconcelos; Jesus; Santos, 2020).

Dentre os AVAs um dos mais utilizados no espaço educacional é o Moodle¹. Este por sua vez é descrito como um ambiente com recursos que auxiliam na interação e no desenvolvimento das atividades (Vasconcelos; Jesus; Santos, 2020). Segundo Silva et al. (2018, p.1178), “O espaço virtual Moodle permite que o professor perceba com mais clareza e facilidade quem está participando ou não do espaço, além de oferecer muitos recursos para desenvolver a avaliação continuada”.

2.5 Arquitetura Moodle

Segundo Silva et al. (2018) “A plataforma Moodle é um recurso muito bem aceito nas instituições de ensino pela facilidade de personalizar e desenvolver customizações da sala de aula virtual”. Diante disso o Moodle foi desenvolvido com objetivo de escala e adaptar-se a elevados números de usuário, na qual em sua estrutura é composta por uma combinação de ferramentas chamada de LAMP (Linux, Apache, MySQL e PHP) (Buchner, 2011; Valencia, 2022). A Figura 2 é apresentada a arquitetura geral do sistema Moodle.

Figura 2: Estrutura Moodle.



Fonte: Buncher (2011)

¹ Disponível em:< <https://moodle.org/>>. Acesso em 19 de outubro de 2023.

3. PROPOSTA DE MIGRAÇÃO

A partir da compreensão teórica abordada anteriormente, foi feito um planejamento e a modelagem para investigação sobre a proposta de migração de um sistema de apoio à aprendizagem do ambiente *on-premise* para uma nuvem pública. Assim, seguindo a metodologia adotada (vide Capítulo 1), criou-se um modelo de uma arquitetura, com objetivo de acomodar as características da aplicação *on-premise* em uma possível migração para a nuvem pública da AWS. No intuito de facilitar o entendimento da arquitetura, serão detalhados todos os recursos utilizados na arquitetura.

3.1 Descrição da arquitetura

Como etapa inicial do planejamento da arquitetura, foi feito um levantamento sobre quais serviços e de que forma eles seriam acomodados na nuvem. Ahson e Ilyas (2010) sugerem um plano com algumas fases a fim de alcançar alguns objetivos para a utilização com sucesso da computação em nuvem. Dentre elas, destacam-se:

- O desenvolvimento da arquitetura de TI, fase que fornece um guia para o projeto e concretização do modelo em nuvem;
- Fase do desenvolvimento de requisitos de qualidade de serviço, que diz respeito aos requisitos de qualidade de serviço denominada de necessidades não funcionais, como o *failover* (tolerância a falhas), a confiabilidade e a segurança.

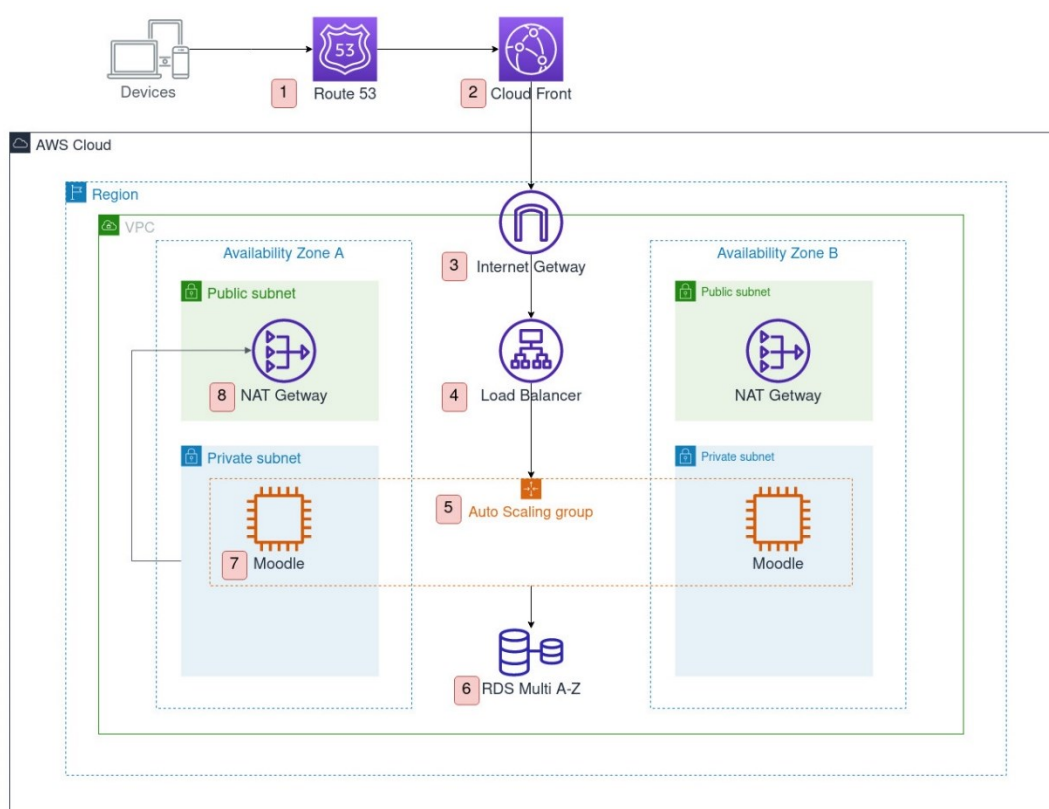
A *Amazon Web Service* (2023), por sua vez, compilou e documentou um conjunto de melhores práticas e princípios, na qual torna possível avaliar e projetar uma infraestrutura. Constituída de 6 pilares, são eles:

- Excelência operacional, apoia o desenvolvimento, execução e melhoria contínua dos processos;
- Segurança, diz respeito à proteção das informações utilizando de boas práticas, a fim de melhorar a segurança na nuvem;
- Confiabilidade, descreve como as cargas de trabalho devem ser projetadas e executadas sempre se preocupando com o *failover* e a disponibilidade da aplicação;

- Eficiência de desempenho, remete-se a como um sistema mantém seu desempenho de acordo com a mudança de demanda e evolução das tecnologias;
- Otimização dos custos, descreve como ter uma redução nos custos dos serviços;
- Sustentabilidade, trata sobre o impacto ambiental, econômico e social da empresa.

A modelagem da arquitetura foi planejada fazendo o uso das boas práticas que a AWS disponibiliza no documento da *Well-Architected Framework*². Na Figura 3 é apresentado um diagrama da arquitetura sugerida.

Figura 3: Arquitetura na Nuvem.



Fonte: Produzido pelo autor (2023).

A acomodação da arquitetura foi feita dentro de uma região ou locais geográficos onde são hospedados os servidores físicos da AWS, na qual definiu-se uma rede virtual privada também chamada de VPC, composta por zonas de disponibilidade ou AZs, que consistem em *data centers* físicos redundantes conectados por *links* de baixa latência e

² Disponível em < https://wa.aws.amazon.com/index.pt_BR.html>. Acesso em 19 de outubro de 2023

alta disponibilidade. É proposto que em cada zona de disponibilidade sejam criadas uma sub-rede pública que possui um *gateway* de Internet e uma rede privada que tem acesso aos recursos internos. A Figura 4 mostra a relação dos servidores e as zonas de disponibilidades.

Figura 4: Região, AZs, *Data Centers*.



Fonte: Amazon Web Service (2023).

No que diz respeito ao modelo da arquitetura do sistema, ao fazer a análise da Figura 3, temos enumerados os serviços e recursos que são utilizados da AWS. Na Tabela 1 apresenta a listagem dos componentes que foram usados na arquitetura.

Tabela 1: Recursos Utilizados

Componentes	Descrição
1. Route 53	Serviço de DNS, que converte endereços de domínios para identificadores de recursos computacionais.
2. CloudFront	Serviço de Rede de entrega de conteúdo (CDN).
3. Internet Gateway	Componente da VPC que permite comunicação dos elementos dentro da VPC com a Internet
4. Load Balancer	Distribui o tráfego de entrada entre vários destinos
5. Auto Scaling Group	Monitora as aplicações e escala recursos computacionais de maneira automática
6. RDS	Serviço de Banco de Dados Relacional
7. EC2	Serviço de máquinas virtuais (instâncias) baseado em sistemas operacionais modernos (e.g., Linux, Windows)
8. NAT Gateway	Serviço de <i>Network Address Translation</i> (NAT) permite redes privadas, dentro da VPC, terem acesso à Internet.

Fonte: Produzido pelo autor (2023).

Na arquitetura proposta, as requisições às instâncias (servidores virtuais) são enviadas ao *Route 53* que terá como função fazer a resolução do DNS e direcionar o tráfego para o *Amazon CloudFront* um serviço de CDN. Este, por sua vez, tem a finalidade de fazer a distribuição do conteúdo do serviço *web* do Moodle de forma rápida e com baixa latência. Segundo *Amazon Web Service* (2023), essa entrega de conteúdo é feita através dos pontos de presença, que são serviços de *datacenters* conectados mundialmente; esses pontos tem por objetivo encurtar a distância entre o conteúdo e usuários, reduzindo a latência e o tempo de carregamento.

Para a comunicação entre o tráfego da Internet e a aplicação dentro da VPC na nuvem, utiliza-se o *gateway* de Internet. Este, por sua vez, tem em sua saída um *Elastic Load Balancer*, que será responsável por fazer a distribuição das requisições entre as instâncias alocadas em AZs distintas redirecionando as cargas de trabalho entre elas, garantindo uma tolerância a falhas na aplicação *web* e uma alta disponibilidade do serviço. Vale ressaltar que a política de escalonamento das *Workloads* é feita através do *Round Robin*, segundo Tamanaha et al. (2021, p. 3), “Sua decisão consiste em atribuir as requisições em sequência com base no grupo de servidores disponíveis”.

Como forma de garantir a segurança ao acesso das instâncias do Moodle em cada AZ, as EC2 serão criadas em sub-rede privada, na qual o acesso à saída dos dados dos recursos privados para a Internet foi feito através de um NAT *gateway* entre a sub-rede privada e a pública.

As máquinas virtuais EC2 se encontrarão dentro de um *Auto Scaling Group* a fim de garantir um escalonamento e gerenciamento automático no número de instâncias de acordo com a demanda nas suas cargas de trabalho entre as AZs. Vale ressaltar que ao utilizar o *Load Balancer* em conjunto com *Auto Scaling* o dimensionamento vertical e horizontal é feito automaticamente, bem como a distribuição dos *workloads* entre as máquinas.

Segundo Paula e Dian (2021), “Escalabilidade é a capacidade do sistema se expandir sem perder desempenho.” Dentre elas destaca-se a escalabilidade horizontal e vertical. Na escalabilidade horizontal se aumenta o número de recursos do mesmo tipo, em que é distribuída a carga de trabalho entre eles. Na escalabilidade vertical o recurso é aumentado para um melhor, ou seja, em termos de máquinas virtuais é aumentado a RAM, vCPU ou armazenamento. (Erl; Putinni; Mahmood, 2013)

Vale ressaltar que, a persistência dos dados será feita através de um banco relacional RDS MySQL multi A-Z. Este, por sua vez, cria uma cópia redundante dos

dados em *standby* em outra AZ, com o intuito de garantir alta disponibilidade, redundância dos dados e segurança contra falhas, para a aplicação.

3.2 Precificação dos recursos

Uma das vantagens da computação em nuvem é o controle sobre custos da sua infraestrutura, em que você paga pelo que usar (Mell; Grance, 2011; Sullivan, 2010). Desse modo, os custos podem ser inferidos de forma granular e separadamente, podendo ser mensurado o valor gasto na infraestrutura no final de mês. Adiante será apresentada a precificação sobre cada serviço que será utilizado na arquitetura proposta.

3.2.1 Estimativa para os custos

Para estimativas dos cálculos da arquitetura proposta, levou-se em consideração uma base de dados dos orçamentos, fornecidos pela Universidade Federal Rural do Semi-Árido, tendo sido pela última vez atualizada em 09/01/2023. Segundo UFERSA (2023) há cerca de 9,26 mil discentes ativos, para uma melhor estimativa, considerou-se 10 mil matrículas ativas e foi adicionado mais 50% a esse valor. Assim, neste trabalho irá assumir o quantitativo de 15 mil usuários para fins de estimativa de custos e desempenho, mesmo a instituição, alvo desse trabalho, ter menos usuários ativos do que esse valor escolhido.

Para gerar as estimativas de custos, outras informações são necessárias. Uma delas é o perfil de consumo de dados e interação de cada usuário com a arquitetura que está sendo proposta. Dessa forma, assumiu, aqui, que consultas ao site que cada indivíduo realizava foi inferida em 10 solicitações HTTP por dia e gerava uma quantidade de 25 kB/s de dados por consulta³. Diante disso, era gerado por dia 150 mil consultas ao site, ou seja, 15mil usuários x 10 requisições/dia, resultando em um total aproximado de dados de 26 kB/s x 150 mil consultas/dia, gerando um total de 3,75 GB dados. No mês o site chegaria a cerca de 4,5 milhões de requisições, resultado da 150mil consultas em 30 dias com um volume aproximado de dados das requisições de 112,5 GB, resultado do produto das 4,5milhões de solicitação pelos 26 kB/s.

3.2.2 Cálculo dos custos

Para os cálculos dos custos gerados por cada serviço, utilizou-se os dados calculados na seção 3.2.1 e a arquitetura proposta. Vale ressaltar que esses valores podem

³ É válido destacar que esses dados são estimativas médias da quantidade de uso do sistema. Como essa pesquisa não teve acesso aos dados reais do AVA da UFERSA, esses valores são considerados razoáveis pelos pesquisadores desse trabalho.

alterar, à medida que a demanda das máquinas virtuais aumenta, os cálculos da conversão foram realizados baseado no valor do dólar no dia 11 de outubro de 2023, na qual um dólar estava valendo R\$ 5,049.

3.2.2.1 VPC

Segundo *Amazon Web Service* (2023), apesar de não haver cobrança a VPC, os recursos na qual se adiciona são taxados baseado no tempo de uso, na qual para a utilização de dois *NAT Gateways* na região do Norte da Virgínia o preço por sua utilização é de US\$ 0.045 por hora para cada 1GB de dados que passam pelo serviço. Temos que o total dos custos da VPC estão detalhados na Tabela 2.

Tabela 2: Custo de VPC

Item	Descrição dos custos	Valor total [US\$]
1	Produto da hora de uso do <i>gateway</i> de 730 horas em um mês pelo valor da utilização de US\$ 0.045.	32.85
2	Produto do processamento de dados do <i>NAT gateway</i> 112,50 GB por mês pelo valor da utilização de US\$ 0.045.	5.06
3	Hora de processamento e mensais do <i>NAT gateway</i> , soma dos Itens 1 e 2.	37.91
4	Uso total e custo mensal do processamento de dois <i>NAT gateways</i> .	75.82

Fonte: Produzida pelo autor (2023).

3.2.2.2 Route 53

No *Route 53*, a precificação se dá através de cada zona hospedada da aplicação, em que será pago uma quantia de US\$ 0.50 e por consulta padrão ao DNS será sobrado um valor de US\$ 0.40 pelo primeiro milhão de consultas e US\$ 0.20 por mais de 1 milhão de consultas. Os custos de uso do serviço, estão detalhados na Tabela 3.

Tabela 3: Custos do *Route 53*

Item	Descrição dos custos	Valor total [US\$]
1	Duas zonas hospedadas.	1.00
2	Produto das 4,50 milhões de consultas padrão por valor de utilização do recurso de US\$ 0.0000004.	1.80
3	Custo das zonas hospedada no <i>Route 53</i> , soma dos Itens 1 e 2.	2.80

Fonte: Produzida pelo autor (2023).

3.2.2.3 *CloudFront*

Para o *CloudFront* a cobrança é feita da transferência externa para a internet, em que nos primeiros 10 TB será pago US\$ 0.085 por GB e por número de solicitações HTTP ou HTTPS, na qual por cada 10000 solicitações HTTP serão cobradas US\$ 0.0075 e HTTPS serão cobradas US\$ 0.010. Vale ressaltar que para esse serviço existem níveis sempre gratuitos, nos primeiros 1 TB de saída de dados, 10000000 solicitações HTTP e HTTPS serão gratuitas. A Tabela 4 apresenta os custos para o serviço de *CloudFront*.

Tabela 4: Custos do *CloudFront*

Item	Descrição dos custos	Valor total [US\$]
1	Produto da transferência de dados de saída para a internet no total de 112,5 GB por um valor de utilização de US\$ 0.085.	9.56
2	Produto dos custos das 4,5 milhões de solicitações HTTP pelo valor da utilização do recurso de US\$ 0.000001.	4.50
3	Custos mensal com o <i>CloudFront</i> no Estados Unidos (mensal), soma dos Itens 1 e 2.	14.06

Fonte: Produzida pelo autor (2023).

3.2.2.4 *Load Balancer*

A cobrança do *Load Balancer* se dá por hora parcial ou completa de sua execução e transferência em GB por meio do recurso. Em que a cada hora parcial é cobrado US\$ 0.025 e a cada GB de dados processados US\$ 0.008. Os custos com o balanceador e carga e apresentado na Tabela 5.

Tabela 5: Custos do *Load Balancer*

Item	Descrição dos custos	Valor total [US\$]
1	Produto do Balanceador de carga, por US\$ 0.025 por hora parcial e 730 horas de uso em um mês.	18.25
2	Produto de um <i>Classic Load Balancer</i> , por um total de 112,5 GB de dados e US\$ 0.008 por GB de dados processados.	0.90
3	Custos com o balanceador de carga mensalmente, soma dos Itens 1 e 2.	19.15

Fonte: Produzida pelo autor (2023).

3.2.2.5 RDS

O banco utilizado na arquitetura foi o *Amazon Relational Database Service* (*Amazon RDS*), para esse tipo de recurso você paga pelo que usar e pelo tipo de instância sob demanda, na qual será pago por horas de capacidade computacional. Diante disso para um tipo de instância db.t3.micro será cobrado US\$ 0.034 por hora, com os custo de armazenamento SSD e uma capacidade de armazenar de 20 GB a 64TiB, as instâncias de banco de dados gp2 de uso geral serão cobrado US\$ 0.23 por GB/mês. vale ressaltar que aos serviços associados com o RDS destaca-se o *proxy* RDS este por suas vezes fica entre a aplicação e o banco de dados gerenciando as conexões, para instâncias MySQL o *proxy* é cobrado por hora consumida a uma taxa de US\$ 0.015 por vCPU/hora. Tem-se que para o RDS os custos descritos na Tabela 6.

Tabela 6: Custos com o RDS

Item	Descrição dos custos	Valor total [US\$]
1	Produto de uma instância do RDS com <i>engine</i> do MySQL, por 730 horas em um mês a um valor de utilização de recurso de US\$ 0.034 por hora.	24.82
2	Produto do armazenamento de 30 GB do EBS no mês por US\$ 0.23 do custo de GB armazenado por instância.	6.90
3	Custo mensal do RDS <i>proxy</i> é o produto de 1 instância com 2 vCPUs, por 730 horas de utilização em um mês e US\$ 0.015 de utilização de vCPU/hora.	21.90
4	Custo total com o serviço de RDS é a soma dos itens 1, 2 e 3.	53.62

Fonte: Produzida pelo autor (2023).

3.2.2.6 Instâncias EC2

As instâncias EC2 são taxadas de acordo com o seu tipo, na qual para a arquitetura modelo, utilizou-se o tipo sob demanda, permitindo apenas pagar pelo que usar, utilizando instâncias do tipo t2.micro com 1 vCPU e 1GB de memória a taxa paga será de US\$ 0.0116 por horas, a um custo médio de US\$ 8.73 por mês de cada instância. Vale ressaltar que, para o serviço do AWS de *auto scaling* não há custos, pagando pelos recursos AWS. Os custos com duas instâncias EC2 rodando sobre demanda é apresentado na Tabela 7.

Tabela 7: Custo das instâncias EC2

Descrição dos custos	Valor total [US\$]
Custos de duas instâncias por US\$ 0.0116 por hora demandada e um total de 730 horas por mês.	16.936

Fonte: Produzida pelo autor (2023).

3.2.2.7 Custos totais

A estimativa realizada pelo AWS *calculator*⁴, apresentou a soma de cada serviço individualmente na qual pode-se obter os gastos mensais e anuais que a infraestrutura iria

⁴ Disponível em <<https://calculator.aws/#/estimate?id=21873367185b35c3f464baae8dd6496f5a9e7c83>>. Acesso em 20 de outubro de 2023.

gerar, sendo que US\$ 182.39, aproximadamente R\$ 920,89, realizando a sua conversão de custo mensal e com valor anual teria US\$ 2188.68 ou R\$11050,64. Cabe salientar que foi utilizado a conversão das moedas de dólar para real no dia 11 de outubro de 2023, na qual um dólar estava valendo R\$ 5,049. A Tabela 8 apresenta um compilado dos custos mensais calculados.

Tabela 8: Custos mensais da arquitetura

Serviços	Valor [US\$]	Valor [R\$]
VPC	75.82	382,81
<i>Route 53</i>	2.80	14,14
<i>CloudFront</i>	14.06	70,99
<i>Elastic Load Balancing</i>	19.15	96,69
<i>Amazon RDS</i>	53.62	270,73
<i>EC2</i>	16.94	85,53

Fonte: Produzida pelo autor (2023).

4. RESULTADOS E DISCUSSÕES

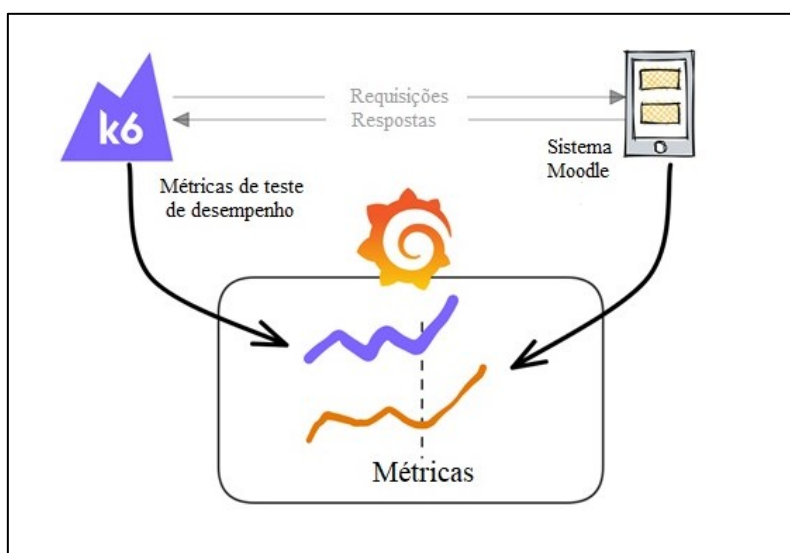
Neste capítulo é abordado todas as configurações e análises dos testes de cargas feitos, além das comparações entre os modelos de arquitetura proposta e o modelo utilizado pela UFERSA.

4.1 Configurações

Os resultados foram obtidos por meio de teste de carga utilizando a ferramenta chamada Grafana⁵. Ela tem como finalidade, fazer a análise de métricas através de gráficos, além de criar painéis que permitem uma visualização desejada, como base nos recursos monitorados, e disponibiliza um período gratuito para testar alguns recursos como Grafana Cloud K6 (Grafana, 2023). Vale ressaltar que, com a ferramenta K6, é possível a realização dos testes de desempenho em alvos, como sites disponíveis na Internet. Os testes de desempenho têm como objetivo descrever aspectos dinâmicos do sistema de tal maneira que seja possível a identificação de gargalos no sistema.

Diante disso, para os testes de desempenho utiliza-se uma instância do Grafana Cloud alimentado com o motor K6 OSS, uma ferramenta de teste de carga de código aberto desenvolvida pela Grafana (K6, 2023; Grafana, 2023). Para a realização dos testes, o *script* ou manifesto é carregado em uma instância do Grafana Cloud, dando a possibilidade de execução de testes em diferentes localizações geográficas (Grafana, 2023). Como é destacado na Figura 5.

Figura 5 : Teste de desempenho.



Fonte: Elaborado pelo autor (2023).

⁵ Disponível em: < <https://grafana.com/> >. Acesso em 22 de outubro de 2023.

As métricas do sistema são apresentadas na Figura 5, em que, ao executar os testes de performance pelo K6 é enviado requisições HTTP e espera-se receber respostas, gerando um conjunto de métricas por parte da ferramenta de teste e do sistema Moodle, sendo possível a visualização dos resultados em um único painel (Grafana, 2023). Segundo a K6 (2023), para a realização dos testes são executados os seguintes passos:

1. Executa os testes de Nuvem;
2. A *Cloud K6* valida o teste e gera metadados, além de criar um plano de execução;
3. Solicita recursos de nuvem no intuito de fazer a ativação dos geradores de cargas necessários e enviam os testes para eles;
4. Executa-se os testes, orquestrando o seu comportamento, a fim de coordená-lo em diferentes executores e zonas de carga;
5. À medida que os dados são recebidos é feito o seu processamento;
6. São exibidos os resultados dos testes.

Vale ressaltar que as Zonas de cargas é um local que executa os teste do Grafana Cloud, em que para os teste de carga do sistema utilizamos a zona de US East (Ohio) da Amazon com uma instância do tipo m5.large, que por sua vez possui equilíbrio em seus recursos para uma variedade de fluxos de trabalho, contam com processadores Intel Xeon Platinum série 8000, são constituídas com 2 vCPUs e 8 GiB de memória RAM (Amazon Web Service, 2023; Grafana, 2023).

Para os testes foi implementado a arquitetura proposta e utilizou-se duas máquinas virtuais EC2 do tipo t2.micro que possui 1 vCPU e 1GB de memória RAM e uma família de processadores Intel Xeon. Em seu uso são consideradas instâncias de uso geral e de baixo custo, com uma *Amazon Machine Image-AMI* do tipo Ubuntu server de arquitetura X86_64. Baseado na proposta de infraestrutura conforme descrita na Seção 3.1, utilizou-se duas instâncias e zonas de disponibilidades diferentes, no intuito de reduzir o *failover*.

O sistema foi o Moodle na versão 4.2.2 (Moodle, 2023). Fazendo o uso do banco de dados foi o recurso da RDS com a *engine* MySQL da versão 8, com uma classe de instâncias do tipo db.t3.micro, com 2 CPUs, 1GB de memória RAM e 20GB de armazenamento. Para o teste não considerou o tipo Multi-AZ, sendo criado em uma única zona de disponibilidade.

Afim de simular uma escalabilidade horizontal, utilizou um *auto scaling group*, capaz de aumentar o número de instâncias de acordo com regras pré-estabelecidas com política de utilização de 80% da vCPU e um *Load Balancer*, para balancear a carga igualmente entre todas as instâncias, no intuito de melhorar a disponibilidade do sistema,

distribuindo a carga de trabalho aos recursos computacionais. Vale ressaltar que com objetivo de investigar a arquitetura, fixou-se o número de instâncias em dois, para análise de respostas das requisições.

Salienta-se que todos esses serviços da AWS foram criados em laboratório virtual fornecido pela nuvem para fins de estudo. As criações dos recursos de *Route 53* e *CloudFront* não foram consideradas nos testes, visto que estas se tratavam de CDN e entrega de conteúdo respectivamente. Cabe destacar que a implantação da arquitetura e escolhas dos serviços utilizados, seguem os modelos permitido pelo nível *free tier*, o nível gratuito disponibilizado pela AWS.

4.2 Testes de carga

No intuito de alcançar os objetivos propostos e de garantir uma maior confiabilidade dos testes, a escolha dos valores se deu baseado nos experimentos feitos por Wiechork, Silva e Antoni (2021), em que sua pesquisa tinha como objetivo analisar desempenho em serviços na nuvem, utilizando o GCP e a AWS.

4.2.1 Configuração dos cenários

Para cada cenário, foram colhidas 6 amostras para cada um dos 4 cenários e, através do cálculo da média de cada teste, encontrou-se o tempo de resposta ao servidor, na Tabela abaixo são descritos os cenários sobre cada teste.

Tabela 9: Teste de carga do Sistema

Cenários	Virtual Users – Vus	Iterações	Total de requisições
1º	1	50	50
2º	10	50	500
3º	50	50	2500
4º	100	50	5000

Fonte: Produzida pelo autor (2023).

A Tabela 9 descreve quatro cenários, nos quais foram realizados os experimentos. O número de *Virtual Users* (VUs), indica uma simulação da quantidade de usuários testando o sistema. Essa medida tem como objetivo enviar requisições HTTP ao sistema alvo. O número de iterações, por sua vez, refere-se ao total de requisições enviadas por cada usuário simultaneamente.

Diante disso, variou-se o número de VUs aumentando o total de requisições e feita as análises em cima desses valores. Além dos cenários descritos na Tabela 9, foi

realizada a adição de mais um caso, no intuito de investigar as respostas do servidor. Neste, por sua vez, variou-se o número de instâncias verificando a resposta ao servidor e o total de requisições falhas, na qual o número de requisições foi semelhante ao cenário 4 da Tabela 9.

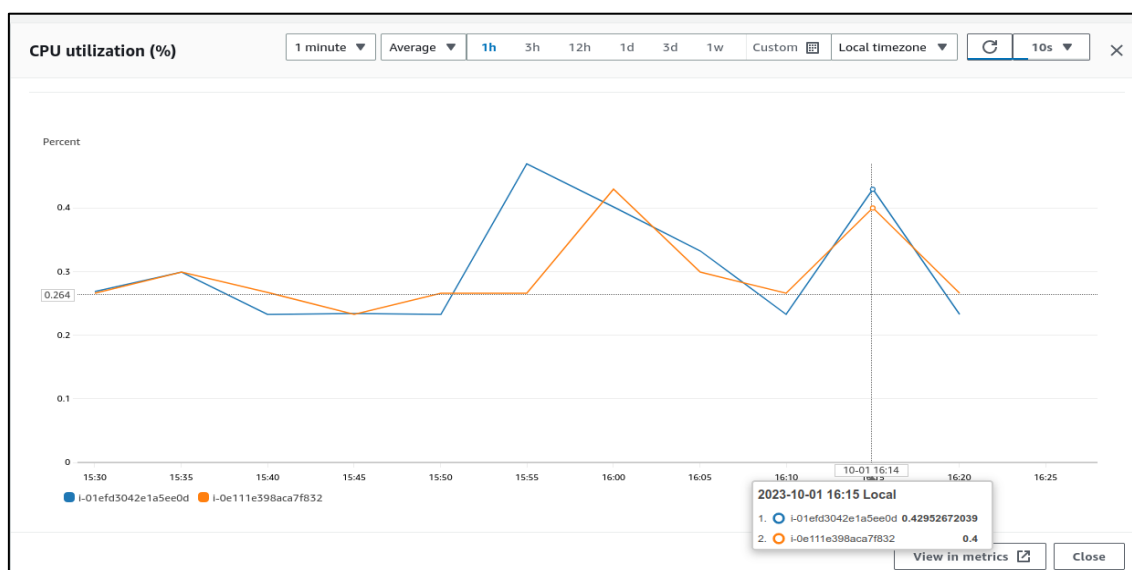
No intuito de obter valores condizentes em cada execução dos testes na coleta das amostras, verificou-se o estado inicial do consumo de CPU em que se constatou uma utilização entre 0,23% a 0,35%. Vale ressaltar que esses valores obtidos foram colhidos através do recurso de monitoramento do *CloudWatch*⁶.

Durante a realização dos testes de desempenho no servidor, foi obtido dados do consumo de CPU de cada amostra dos cenários descritos na Tabela 9. Em que, para cada amostra estressava-se a arquitetura implementada (com duas instâncias) realizando o envio de requisições HTTP, por consequência o consumo de CPU das instâncias aumentava e era obtido outra amostra quando os níveis de utilização das instâncias baixassem.

4.2.2 Cenário 1

Para o primeiro cenário os testes foram realizados com 1 VU enviado 50 requisições HTTP ao site (para a página inicial do sistema). Durante a execução do teste de desempenho, a utilização de CPU do sistema testado encontrou-se entre 0,39% a 0,46% de uso. Na Figura 6 são apresentadas informações do seu uso durante a execução de uma amostra.

Figura 6: Consumo de CPU das instâncias durante a execução dos testes.



Fonte: Elaborado pelo autor (2023).

⁶ Disponível em: < <https://aws.amazon.com/pt/cloudwatch/> >. Acesso em 22 de outubro de 2023

Ao fazer a análise da Figura 6 percebe-se que as 50 requisições estão sendo distribuídas entre as duas instâncias pelo *Load Balancer*, em que ambas se encontram com um valor aproximado de 0.4% durante a execução da carga. Vale ressaltar que o balanceador de carga utiliza do algoritmo de escalonamento com política de *Round Robin*, para fazer a distribuição das requisições. Após a experimentação de 6 amostras, no teste de desempenho, a Tabela 10 informa os resultados encontrados.

Tabela 10: Amostra 1º cenário

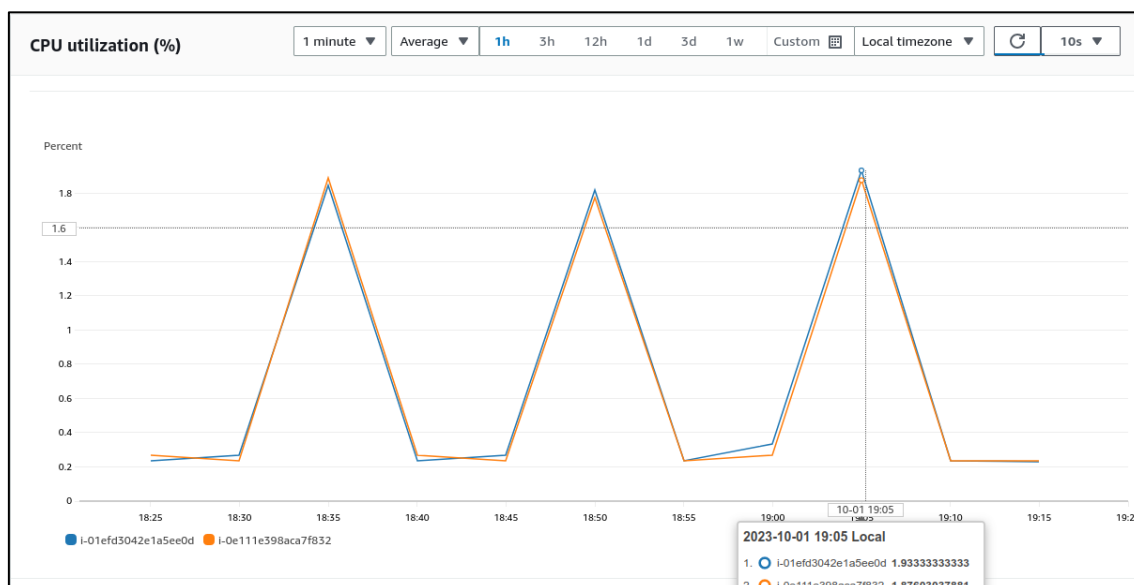
Amostra	Cenário	Duração	Resposta HTTP [ms]
1 ^a	1VU – 50 requisições	9s	77,4
2 ^a	1VU – 50 requisições	9s	75,0
3 ^a	1VU – 50 requisições	9s	75,3
4 ^a	1VU – 50 requisições	9s	89,5
5 ^a	1VU – 50 requisições	9s	79,6
6 ^a	1VU – 50 requisições	9s	72,7

Fonte: Elaborado pelo autor (2023).

Ao fazer a análise dos valores, percebe-se que a duração dos testes foi de 9s, em que se calculou a média das requisições HTTP ao site, obtendo um valor médio de 78,25ms na resposta do servidor, ao analisar o Grafana cloud durante a execução dos testes de carga, percebe-se que a taxa de requisições ao servidor foi de 100%, ou seja, das 50 solicitações todas foram atendidas com sucesso.

4.2.3 Cenário 2

Para o segundo cenário, houve o aumento do número de usuários virtuais para 10 VUs. Foi mantido a quantidade de 50 consultas HTTP em simultâneo, totalizando 500 solicitações ao site. A Figura 7 apresenta o uso da CPU durante algumas amostras dos servidores da infraestrutura do Moodle.

Figura 7: Utilização de CPU durante 500 requisições, cenário 2.

Fonte: Elaborado pelo autor (2023).

Ao fazer a análise da Figura 7 percebe-se um aumento no consumo de CPU das duas instâncias bem próximas entre 1.75% e 1.95% o que era esperado pois o número de requisições foi maior. Essa faixa de valores foi encontrada a partir da execução das 6 amostras. Como citado anteriormente, após os servidores voltarem para um baixo consumo era obtida outra amostra, o que explica o comportamento periódico do gráfico. A Tabela 11 apresenta o tempo de resposta em cada execução dos testes.

Tabela 11: Tempo de resposta para o 2º cenário.

Amostra	Cenário	Duração	Resposta HTTP [ms]
1ª	10VU – 50 requisições	12s	162
2ª	10VU – 50 requisições	15s	153
3ª	10VU – 50 requisições	12s	154
4ª	10VU – 50 requisições	12s	144
5ª	10VU – 50 requisições	12s	148
6ª	10VU – 50 requisições	12s	146

Fonte: Elaborado pelo autor (2023).

Ao fazer a análise da Tabela 11, percebeu-se que ao aumentar o número de requisições enviadas ao servidor, há um aumento no tempo de resposta. Calculando a média das respostas HTTP do servidor, obteve-se um total de 151,17 ms, com um tempo de duração médio de cada teste de 13,5 s, com base nas análises do grafana cloud, pode-se verificar que durante os testes do segundo cenário, a taxa de resposta as requisições HTTP foram de 100%, dessa forma todas solicitação ao servidor foi respondida.

4.2.4 Cenário 3

Para o terceiro cenário aumentou-se o número de VUs para 50 com 50 iterações simultaneamente totalizando 2500 requisições. A Figura 8 apresenta a utilização da CPU durante os testes de carga do terceiro cenário.

Figura 8: Servidores durante os testes do terceiro cenário



Fonte: Elaborado pelo autor (2023).

Ao analisar o uso de CPU pode-se perceber um aumento comparado aos outros dois cenários anteriores entre 9.20% e 8.20%, ou seja, mesmo tendo aumentado o consumo de CPU das instâncias, os valores aproximados foram em torno de 10% da capacidade de processamento. A Tabela 12 apresenta os resultados dos testes.

Tabela 12: Resultados dos testes para o terceiro cenário

Amostra	Cenário	Duração	Resposta HTTP [ms]
1 ^a	50VU – 50 requisições	33s	945
2 ^a	50VU – 50 requisições	33s	926
3 ^a	50VU – 50 requisições	36s	770
4 ^a	50VU – 50 requisições	33s	909
5 ^a	50VU – 50 requisições	36s	950
6 ^a	50VU – 50 requisições	33s	860

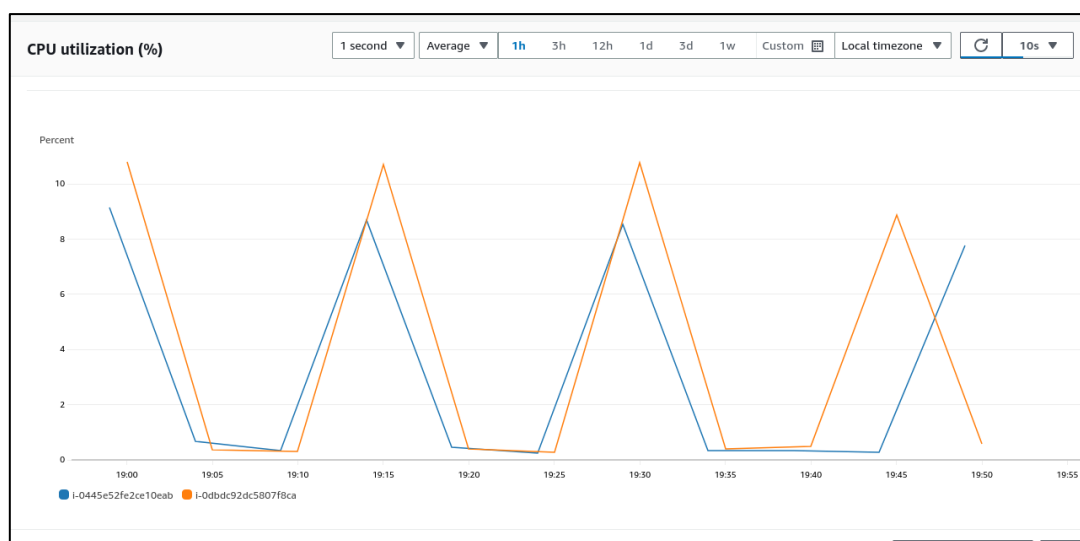
Fonte: Elaborado pelo autor (2023).

Ao analisar a Tabela 12, houve um aumento no tempo de resposta HTTP de 894,83ms, sem requisições falhadas, na qual a duração dos testes das amostras foi em torno de 34s, ocasionado pelo alto número de requisições ao site.

4.2.5 Cenário 4

No intuito de estressar a arquitetura, utilizamos um cenário de testes com 100 VUs fazendo cerca de 50 iterações, totalizando 5000 requisições ao site. Percebeu-se que a cada teste realizado o consumo de CPU para uma das instâncias ultrapassava de 10%, enquanto a outra ficava entre 8% e 9% do consumo. Como o cenário 3 e 4 são bastante semelhantes em termos de utilização CPU, uma forma de justificar essa similaridade é que há outro gargalo na arquitetura. Possivelmente I/O, mas esse comportamento não foi alvo desse trabalho, a Figura 9 mostra esse comportamento.

Figura 9: Uso da CPU com 5000 requisições.



Fonte: Elaborado pelo autor (2023).

Ao fazer as análises da Figura 9, nota-se que apesar das 5000 requisições o estresse da CPU não foi tão elevado. Na Tabela 13 é apresentado o tempo de resposta de cada amostra.

Tabela 13: Tempo do servidor para 5000 requisições

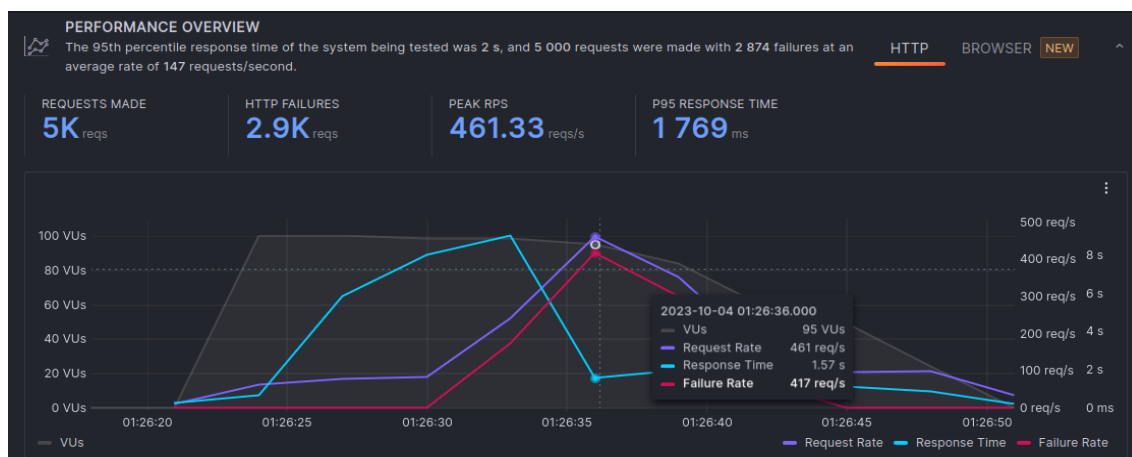
Amostra	Cenário	Duração	Resposta HTTP [s]
1 ^a	100VU – 50 requisições	39s	2,36
2 ^a	100VU – 50 requisições	36s	2,26
3 ^a	100VU – 50 requisições	33s	2,15
4 ^a	100VU – 50 requisições	39s	2,38
5 ^a	100VU – 50 requisições	39s	2,33
6 ^a	100VU – 50 requisições	39s	1,97

Fonte: Elaborado pelo autor (2023).

Ao fazer a análise do tempo de resposta do site na Tabela 13, tivemos um tempo médio de 2,24 s, e uma duração média na execução dos testes de aproximadamente 37,5

s. A Figura 10 mostra os resultados do teste de carga em uma das amostras do 4º cenário, nesse teste em comparação com os outros cenários há um aumento no total de requisições HTTP, para 5000 solicitações simultâneas. Devido ao número elevado de requisições houve uma taxa de aproximadamente 58% de requisições falhadas, como é mostrado na Figura 10.

Figura 10: Teste de performance para uma amostra



Fonte: Elaborado pelo autor (2023).

Ao analisar a Figura 10 percebeu-se que à medida que as requisições aumentavam representada pela linha azul, a taxa de solicitação representada pela linha roxa e a taxa de falhas destacada como a linha vermelha ainda estava baixa, porém à medida que a taxa de requisições aumenta, havia uma diminuição no tempo de resposta do site e por consequência a taxa de falhas aumentava. Em termos de valores quando a taxa de solicitação estava mandando 461 req/s a resposta ao servidor era de 1,57s e uma taxa de falha de 417 req/s.

No intuito de aferir quantas máquinas virtuais seriam necessárias para atender às 5 mil requisições, e garantir que houvesse uma diminuição no tempo de resposta, e por consequência uma diminuição da taxa de HTTP falhos. Fez-se o aumento no número de instâncias na arquitetura, o que resultou em uma melhora na resposta ao servidor. A Tabela 14, apresenta o aumento do número de máquinas.

Tabela 14: Aumento do número de máquinas EC2

Nº de instâncias	Cenário	Duração	Resposta HTTP [s]
2	100VU – 50 requisições	33s	1,77
4	100VU – 50 requisições	30s	1,09
6	100VU – 50 requisições	24s	0,738

Fonte: Elaborado pelo autor (2023).

Ao fazer a análise da Tabela 14, percebe-se que com 2 instâncias a resposta ao servidor era de 1,77s, com 4 instâncias o tempo de resposta diminuía para 1,09s e como 6 máquinas era consumido tinha uma resposta de 738ms, o que nos mostra que ao fazer o uso de mais instâncias a resposta ao servidor melhorava, porém, a taxa de requisições falhadas continuava alta.

Uma das explicações que pode ser tomada é que ao fazer 5000 requisições simultâneas com instâncias EC2 do tipo t2.micro consideradas máquinas de uso geral, na qual, possui uma baixa capacidade de processamento, para altas taxas de requisições HTTP, o que poderia explicar o elevado número de falhas. Uma solução para esse problema seria escalonamento vertical e horizontal da arquitetura, a fim de aumentar o número de instâncias e a capacidade de processamentos.

Analisando a Tabela 14 percebe-se para as amostras executada nas 6 instâncias, ao realizar a 5000 requisições o tempo de resposta é aproximado, quando comparado ao cenário 3. Vale ressaltar que o escalonamento de máquinas aumenta o custo do serviço, na qual para duas instâncias seria de US\$ 16,936, para 4 instâncias a precificação do EC2 seria de US\$ 33,872 e com 6 instâncias o custo de EC2 aumentaria por US\$ 50,808.

4.3 Comparações

Após a realização dos testes na infraestrutura, foi comparado os custos da arquitetura proposta detalhados na seção 3.2.2.7, com os custos que a UFERSA gastou com os serviços de IaaS, apresentado através de um painel orçamentário⁷, tendo sido atualizado do dia 09/10/2023. Para a infraestrutura proposta os valores escolhidos foram superestimados, como é descrito na seção 3.2.1, tendo como custos anual de R\$ 11.050,64 e analisando o painel orçamentário, utilizamos os gastos com IaaS de R\$ 15.448,79. Tem-se uma redução de R\$ 4.398,15, aproximadamente 28% na diminuição dos custos da arquitetura proposta, comparado com os valores disponibilizados pela UFERSA.

Vale ressaltar que a seção 3.2.2.7 apresenta os custos com duas instâncias EC2 em execução, porém à medida que os recursos são provisionados e é feito uma escalabilidade horizontal das máquinas como foi realizado no cenário 4, cabe destacar, que são incrementados, mas apenas para o serviço de computação. Todos os demais serviços são praticamente mantidos os mesmos custos.

⁷ Disponível em < <https://transparencia.ufersa.edu.br/execucao-orcamentaria-e-financeira/>>. Acessado em 21 de outubro de 2023.

5. CONCLUSÃO

Mediante os aspectos pesquisados e aqui relatados, podemos compreender o crescimento tecnológico na atualidade, e conseqüentemente a alta migração de servidores locais para a nuvem. Nesse ínterim, tivemos como objetivo investigar a viabilidade de uma possível migração do sistema da Universidade Federal Rural do Semi-Árido-UFERSA, com servidores locais (*on-premises*) para nuvem pela AWS, realizando um estudo bibliográfico sobre os conceitos de computação em nuvem, desenvolvendo e apresentando os custos de uma arquitetura para um sistema de apoio ao aprendizado e implantando a arquitetura a fim de comparar os custos e desempenho da mesma. Além disso, conhecemos os tipos de nuvem dispostos pela AWS, com a percepção de que a computação em nuvem possui distintas características e possibilidades ao contratar esses recursos como serviços.

Dessa forma, entendemos que o modelo de infraestrutura tecnológica usado nos sistemas de apoio ao aprendizado da UFERSA segue o desenvolvimento de hardware e software em um ambiente local, e isso acaba por dificultar o provisionamento dos recursos computacionais, por isso tem-se a necessidade de se ter um modelo alternativo e singular do *on-premises*, como a computação em nuvem.

Sendo assim, como descrito no decorrer do trabalho, são inúmeras as vantagens da computação em nuvem, como: sua infraestrutura, controle de gastos, escalabilidade, etc. E para melhor compreensão e análise durante o trabalho, desenvolvemos investigações, comparações e estimativas de custos da arquitetura a ser implementada na nuvem.

Com isso, houve o levantamento e escolha de serviços e suas formas de implantação na nuvem, mostrando um modelo de arquitetura que fizesse uso das boas práticas que a AWS disponibiliza, acomodada dentro de uma região onde se hospeda servidores físicos da AWS, e assim percebemos a existência de grandes benefícios da migração do sistema UFERSA (*on-premises*) para nuvem pela AWS.

Em síntese, os benefícios são o controle de gastos da infraestrutura, que como mencionado no trabalho, é pago somente pelo que usar, e com os dados obtidos pela instituição foi feito à base de custos de cada serviço, assim também, a capacidade de suportar um alto número de requisições HTTP e maior resistência a falhas.

Os testes foram realizados utilizando o grafana fazendo análise e expondo por meio de gráficos, com testes de cargas com o objetivo de garantir maior amplitude e

veracidade dos mesmos, e com isso foi realizado comparações de custos e de arquiteturas dos serviços utilizados pela UFERSA os serviços da nuvem.

Dessa forma, com a construção desse trabalho, é perceptível o número de bonificações com a migração do sistema da instituição para o modelo de computação em nuvem pública, assim também como esse trabalho é relevante e de estudo pertinente para a comunidade acadêmica e social que possua interesse em compreender mais sobre os serviços da nuvem e suas finalidades.

Como trabalhos futuros, pretende-se melhorar a resposta ao servidor a um elevado número de requisições, alterando o tipo de instâncias, fazer uma investigação sobre aumentar a segurança através dos recursos que a nuvem oferece como a *Web Application Firewall* (WAF) e o *AWS Shield*⁸, bem como realização teste de penetração dentro da infraestrutura proposta, garantindo a segurança dos dados dentro da AWS.

⁸ Disponível em: < <https://aws.amazon.com/pt/shield/> >. Acesso em 22 de outubro de 2023.

6. REFERÊNCIAS

AHSON, Syed A.; ILYAS, Mohammad (Ed.). **Cloud computing and software services: theory and techniques**. CRC Press, 2010.

ALVES, Lynn. **Um olhar pedagógico das interfaces do Moodle**. 2009.
ARMBRUST, Michael et al. **Above the clouds: A berkeley view of cloud computing**. Technical Report UCB/EECS-2009-28, EECS Department, University of California, Berkeley, 2009.

Amazon Web Service. **Amazon Web Services (AWS) - Cloud Computing Services**. Disponível em: <<https://docs.aws.amazon.com/>>. Acesso em 11 de outubro de 2023.

BANDEIRA, Waliff Cordeiro. **Análise do Custo-benefício de Funções como Serviço e Infraestrutura como Serviço**. 2022.

BORGES, Filipe Afonso Nogueira; SILVA, Gabriel Santos Tavares da. **Proposta de modelo de migração de sistemas on-premises para nuvem pública no Brasil**. 2019.

BUCHNER, Alex. **Moodle 2 Administration**. Packt Publishing Ltd, 2011.

COSTA, Gisele Maria Freire; SANTOS, Clayton Eduardo dos. **Utilização de containers: Um contexto histórico**. Revista Científica e-Locução, v. 1, n. 20, p. 23-23, 2021.

CARISSIMI, Alexandre. **Desmistificando a computação em nuvem**. Instituto de Informática–Universidade Federal do Rio Grande do Sul (UFRGS)–Porto Alegre–RS, 2015.

ERL, Thomas; PUTTINI, Ricardo; MAHMOOD, Zaigham. **Cloud computing: concepts, technology & architecture**. Pearson Education, 2013.

GRAFANA. **Grafana Cloud documentation**. Disponível em: <<https://grafana.com/docs/grafana-cloud/>>. Acesso em 1 de outubro de 2023

HARDING, Chris. **Cloud computing for business-The open group guide**. Van Haren, 2011.

K6. **Load testing for engineering teams**. Disponível em: <<https://k6.io/>>. Acesso em 1 de outubro de 2023.

MATHER, Tim; KUMARASWAMY, Subra; LATIF, Shahed. **Cloud security and privacy: an enterprise perspective on risks and compliance**. " O'Reilly Media, Inc.", 2009.

MELL, Peter; GRANCE, Timothy. The NIST definition of cloud computing NIST Special Publication. **National Institute of Standards and Technology, Gaithersburg, Maryland, USA**, 2011.

MOODLE. **Moodle - Open-source learning platform**. Disponível em: <<https://moodle.org/>>. Acesso em 14 de setembro de 2023.

NASCIMENTO, Antonio Carlos Aparecido; SANTOS, Clayton Eduardo dos. **Implantação do E-Sus atenção primária em ambiente de nuvem: Estudo de caso**. Revista Científica e-Locução, v. 1, n. 19, p. 25-25, 2021.

PAULA, Laís de; DIAN, Mauricio de Oliveira. **Computação em nuvem: os desafios das empresas ao migrar para a nuvem**. Revista Interface Tecnológica, v. 18, n. 2, p. 304-315, 2021.

QUEIROZ, Carlos Felipe; CONCEIÇÃO, Lucas Maués Calandrine; BARRETO, Marcos Vinicius Sadala. **Estudo dos aspectos para a disponibilização de dados locais na nuvem: Estudo de caso Netdocto**. RECIMA21-Revista Científica Multidisciplinar-ISSN 2675-6218, v. 2, n. 6, p. e26474-e26474, 2021.

RED HAT. **O que são serviços IaaS, PaaS e SaaS?** Disponível em: <<https://www.redhat.com/pt-br/topics/cloud-computing/iaas-vs-paas-vs-saas>>. Acesso em 11 de outubro de 2023.

SANTOS, Edméa Oliveira dos; OKADA, Alexandra Lilavati Pereira. **A construção de ambientes virtuais de aprendizagem: por autorias plurais e gratuitas no ciberespaço**. ANPED, GT: Educação e Comunicação, n. 16, 2003.

SILVA, Débora de Sales Fontoura da et al. **Ensino híbrido com a utilização da plataforma Moodle**. Revista Thema, v. 15, n. 3, p. 1175-1186, 2018.

SIMMON, Eric et al. **Evaluation of cloud computing services based on NIST SP 800-145**. NIST Special Publication, v. 500, n. 322, p. 500-322, 2018.
SOSINSKY, Barrie. **Cloud computing bible**. John Wiley & Sons, 2010.

SOUSA, Flávio RC; MOREIRA, Leonardo O.; MACHADO, Javam C. **Computação em nuvem: Conceitos, tecnologias, aplicações e desafios**. II Escola Regional de Computação Ceará, Maranhão e Piauí (ERCEMAPI), p. 150-175, 2009.

SULLIVAN, Dan. **The definitive guide to cloud computing**. Real Time Nexus, n. 1, p. 4-11, 2010.

TAMANAH, L. K. G. et al. **Algoritmos para Balanceamento de Carga no NGINX**. 2021.

UFERSA. **Números Ufersa**. Disponível em: <<https://numeros.ufersa.edu.br/graduacao/>>. Acesso em: 11 out. 2023.

VALENCIA, Anthony Alexander Aveiga. **Propuesta de implementación de MOODLE de alta demanda balanceada y escalable**. 2022. Tese de Doutorado. Ecuador-Pucese-Escuela de Tecnologías de la Información.

VANDERLEI, Diego Borgheti; SANTANA, Jhonatam Wilson; CAROSIA, Jaciara Silva. **Estudo de caso sobre implantação de sistema de infraestrutura de redes e gestão de assinantes em empresas do setor de telecomunicações**. In: Congresso de Tecnologia-Fatec Mococa. 2021.

VASCONCELOS, Cristiane Regina Dourado; JESUS, Ana Lúcia Paranhos de; SANTOS, Carine de Miranda. **Ambiente virtual de aprendizagem (AVA) na educação a distância (EAD): um estudo sobre o Moodle**. Brazilian Journal of Development, v. 6, n. 3, p. 15545-15557, 2020.

WIECHORK, Karina; SILVA, Darlan Eziquiel Felisberto da; ANTONI, Marco. **Análise de desempenho em serviços gratuitos na nuvem utilizando os provedores AWS e GCP**. In: **Anais do XX Workshop em Desempenho de Sistemas Computacionais e de Comunicação**. SBC, 2021. p. 131-136.

ZUFFO, Marcelo Knörich et al. **A computação em nuvem na Universidade de São Paulo**. Revista USP, n. 97, p. 9-18, 2013.