



UNIVERSIDADE FEDERAL DE PERNAMBUCO
CENTRO DE INFORMÁTICA
PROGRAMA DE PÓS-GRADUAÇÃO EM CIÊNCIA DA COMPUTAÇÃO

PAULO HENRIQUE GOMES SILVA

**Análise de precificação de instâncias spot para minimização de custos em um
contexto multi-nuvem**

Recife

2025

PAULO HENRIQUE GOMES SILVA

Análise de precificação de instâncias spot para minimização de custos em um contexto multi-nuvem

Dissertação apresentada ao Programa de Pós-Graduação Acadêmico em Ciência da Computação do Centro de Informática da Universidade Federal de Pernambuco (UFPE), como requisito parcial para obtenção do título de mestre em Ciência da Computação.

Área de Concentração: Sistemas Distribuídos

Orientador (a): Carlos André Guimarães Ferraz

Recife

2025

.Catalogação de Publicação na Fonte. UFPE - Biblioteca Central

Silva, Paulo Henrique Gomes.

Análise de precificação de instâncias spot para minimização de custos em um contexto multi-nuvem / Paulo Henrique Gomes Silva. - Recife, 2025.

103 f.: il.

Dissertação (Mestrado) - Universidade Federal de Pernambuco, Centro de Informática, Programa de Pós-Graduação em Ciência da Computação, 2025.

Orientação: Carlos André Guimarães Ferraz.

Inclui referências e apêndices.

1. Computação em nuvem; 2. Instâncias Spot; 3. Precificação dinâmica; 4. Multi-nuvem; 5. Otimização de custos; 6. FinOps. I. Ferraz, Carlos André Guimarães. II. Título.

UFPE-Biblioteca Central

Paulo Henrique Gomes Silva

**“Análise de precificação de instâncias SPOT para minimização
de custos em um contexto multi-nuvem”**

Dissertação de mestrado apresentada ao Programa de Pós-Graduação em Ciência da Computação da Universidade Federal de Pernambuco, como requisito parcial para a obtenção do título de Mestre em Ciência da Computação. Área de Concentração: Redes de Computadores e Sistemas Distribuídos.

Aprovado em: 27/01/2025.

BANCA EXAMINADORA

Prof. Dr. Vinícius Cardoso Garcia
Centro de Informática / UFPE

Prof. Dr. Ioram Schechtman Sette
CESAR School

Prof. Dr. Carlos André Guimarães Ferraz
Centro de Informática / UFPE
(orientador)

Dedico esta dissertação, com todo o meu amor e gratidão,

A Deus, que me sustenta, me dá sabedoria e a quem devo tudo o que tenho e sou.

À minha avó, cuja força, sacrifício e sabedoria foram fundamentais para que eu pudesse alcançar este momento.

À minha mãe, que sempre acreditou no meu potencial e me incentiva a alcançar os objetivos com garra e determinação.

À minha esposa e aos meus filhos, pela paciência, apoio e compreensão durante todas as horas que dediquei a este trabalho.

E, com destaque especial, à minha irmã Heloísa, que, mesmo enfrentando um diagnóstico de saúde desafiador nas últimas semanas de finalização deste trabalho, tem nos ensinado a encarar as adversidades com alegria, perseverança e resiliência. Sua força é uma inspiração que transcende palavras e que carregarei comigo para sempre. Este título é também para você, e aguardo com alegria o dia em que o seu será igualmente concretizado.

A todos vocês, esta conquista é tanto minha quanto de vocês.

AGRADECIMENTOS

Ao meu orientador, professor Carlos André Guimarães Ferraz, por ter acreditado na minha proposta desde o processo de seleção e que, ao longo do curso, me constrangeu com tanto zelo, humanidade e ensinamentos. Sua orientação foi fundamental para o desenvolvimento deste trabalho, e sou eternamente grato por toda a dedicação.

A todos os professores com quem tive a oportunidade de aprender ao longo das disciplinas. Seus ensinamentos foram pilares fundamentais para a construção do conhecimento que deu base a esta dissertação.

Ao meu companheiro de turma, Pedro Caminha, por todo apoio e parceria ao longo das disciplinas e nas reuniões semanais.

A Kyungyong Lee, Jaeil Hwang e todos os autores do SpotLake, que, desde o primeiro contato, foram extremamente receptivos e disponibilizaram um conjunto expressivo de dados. Sem essa generosidade e suporte, não seria possível alcançar os resultados apresentados neste trabalho.

A todos os amigos que sempre me incentivaram e desejaram o bem.

À minha família, que, com paciência e compreensão, me deu suporte durante todo o mestrado. A vocês, que são meu alicerce, devo grande parte desta conquista.

Por fim, a todas as pessoas que, de maneira direta ou indireta, contribuíram para que eu chegasse até aqui, seja com palavras de incentivo, conselhos ou simplesmente acreditando no meu potencial. Cada gesto fez diferença nesta jornada.

RESUMO

A computação em nuvem revolucionou a forma como organizações de todos os tamanhos gerenciam seus recursos computacionais, oferecendo escalabilidade, flexibilidade e eficiência de custos. Entre as opções disponíveis, as instâncias spot se destacam como uma alternativa econômica ao reutilizar capacidade computacional ociosa. Contudo, sua adoção apresenta desafios devido à alta variabilidade de preços, à imprevisibilidade na disponibilidade e às diferenças regionais entre os provedores. Essas características tornam complexa a implementação de estratégias otimizadas em ambientes multi-nuvem, onde decisões baseadas em dados são fundamentais. Este trabalho analisa a precificação de instâncias spot e sob demanda nos três maiores provedores globais de computação em nuvem — AWS, Azure e GCP —, utilizando o *dataset* SpotLake, que disponibilizou 73.431 arquivos e 14.804.169 registros de preços coletados ao longo de 12 meses. Os resultados desta análise revelaram que, em determinadas regiões e tipos de instâncias, os preços spot podem ser até 90% inferiores aos preços sob demanda, evidenciando o potencial de economias significativas para organizações que adotam estratégias bem informadas. Além disso, estratégias multi-nuvem demonstraram economias adicionais de até 25%, ao permitir a seleção de regiões e provedores com preços mais vantajosos. A análise revelou também padrões importantes de correlação entre regiões e provedores, fornecendo *insights* sobre a interação entre suas estratégias de precificação. Foi possível identificar que regiões específicas, como *us_east_1*, apresentam maior estabilidade e competitividade de preços, enquanto outras exibem maior volatilidade, reforçando a relevância de análises regionais na tomada de decisão. Além de explorar as dinâmicas de preços spot, este trabalho contribuiu para uma melhor compreensão de fatores como localização geográfica, sazonalidade e características das instâncias que influenciam as decisões de alocação. Esses resultados não apenas destacam as oportunidades de otimização de custos para empresas, mas também fornecem subsídios para futuras pesquisas, incluindo o desenvolvimento de ferramentas automatizadas para monitoramento e previsão de preços em tempo real.

Palavras-chaves: Computação em nuvem, Instâncias Spot, Precificação dinâmica, Multi-nuvem, Otimização de custos, FinOps.

ABSTRACT

Cloud computing has revolutionized how organizations of all sizes manage their computational resources, offering scalability, flexibility, and cost efficiency. Among the available options, spot instances stand out as an economical alternative by reutilizing idle computational capacity. However, their adoption presents challenges due to high price variability, unpredictable availability, and regional differences among providers. These characteristics complicate the implementation of optimized strategies in multi-cloud environments, where data-driven decisions are essential. This study examines the pricing of spot and on-demand instances across the three largest global cloud providers — AWS, Azure, and GCP — using the SpotLake dataset, which comprises 73,431 files and 14,804,169 pricing records collected over 12 months. The results revealed that, in specific regions and instance types, spot prices can be up to 90% lower than on-demand prices, highlighting the potential for significant savings for organizations that adopt well-informed strategies. Furthermore, multi-cloud strategies demonstrated additional savings of up to 25% by enabling the selection of regions and providers with more advantageous pricing. The analysis also uncovered important correlation patterns between regions and providers, offering insights into the interaction between their pricing strategies. Specific regions, such as *us_east_1*, were found to exhibit greater stability and competitive pricing, while others showed higher volatility, emphasizing the importance of regional analyses in decision-making processes. Beyond exploring spot price dynamics, this work contributes to a better understanding of factors such as geographic location, seasonality, and instance characteristics that influence allocation decisions. These findings not only highlight cost optimization opportunities for companies but also provide foundations for future research, including the development of automated tools for real-time price monitoring and prediction.

Keywords: Cloud computing, Spot instances, Dynamic pricing, Multi-cloud, Cost optimization, FinOps.

LISTA DE FIGURAS

Figura 1 – Principais provedores de nuvem.	15
Figura 2 – Árvore da estrutura dos arquivos disponibilizados.	34
Figura 3 – Estrutura da base MongoDB do experimento.	36
Figura 4 – Distribuição de preços sob demanda e spot para os provedores AWS, Azure e GCP na região <i>us_east_1</i> , considerando máquinas de propósito geral no período de maio de 2023 a maio de 2024 – Observe as variações ao interpretar os gráficos, pois elas reforçam a análise das dinâmicas de preços descritas no texto.	52
Figura 5 – Histogramas e KDEs de preços sob demanda e spot para os provedores AWS, Azure e GCP na região <i>us_east_1</i> , considerando máquinas de Propósito Geral no período de maio de 2023 a maio de 2024.	56
Figura 6 – Variação de preços sob demanda e spot na AWS para a região <i>us_east_1</i> , instâncias Otimizadas para Memória – Período: 05/2023 a 05/2024.	70
Figura 7 – Variação de preços sob demanda e spot na AWS para a região <i>sa_east</i> , instâncias Otimizadas para Memória – Período: 05/2023 a 05/2024.	70
Figura 8 – Variação de preços sob demanda e spot na Azure para a região <i>us_west_1</i> , instâncias de Propósito Geral – Período: 05/2023 a 05/2024.	71
Figura 9 – Variação de preços sob demanda e spot na Azure para a região <i>ca_central</i> , instâncias de Propósito Geral – Período: 05/2023 a 05/2024.	71
Figura 10 – Variação de preços sob demanda e spot na GCP para a região <i>sa_east</i> , instâncias Otimizadas para Computação – Período: 05/2023 a 05/2024.	72
Figura 11 – Variação de preços sob demanda e spot na GCP para a região <i>us_west_2</i> , instâncias Otimizadas para Computação – Período: 05/2023 a 05/2024.	72

LISTA DE CÓDIGOS

Código Fonte 1	– Exemplo de documento da coleção <i>instances_type</i>	36
Código Fonte 2	– Exemplo de documento da coleção <i>regions</i>	37
Código Fonte 3	– Exemplo de documento da coleção <i>pricing</i>	37

LISTA DE TABELAS

Tabela 1 – Comparação entre os Trabalhos Relacionados e a Presente Dissertação . . .	32
Tabela 2 – Agrupamento das regiões	39
Tabela 3 – Agrupamento de instâncias de propósito geral	42
Tabela 4 – Correlação de Pearson entre provedores nas regiões destacadas	69
Tabela 5 – Resumo Quantitativo e Metodológico da Pesquisa	80
Tabela 6 – Agrupamento de instâncias de propósito geral	89
Tabela 7 – Agrupamento de instâncias otimizadas para memória	90
Tabela 8 – Agrupamento de instâncias otimizadas para computação	91
Tabela 9 – Agrupamento de instâncias de Computação Acelerada	92
Tabela 10 – Agrupamento de instâncias otimizadas para armazenamento	99
Tabela 11 – Correlação de Pearson entre provedores por região - Computação Acelerada	100
Tabela 12 – Correlação de Pearson entre provedores por região - Propósito Geral	101
Tabela 13 – Correlação de Pearson entre provedores por região - Otimizada para Com- putação	102
Tabela 14 – Correlação de Pearson entre provedores por região - Otimizada para Memória	102

SUMÁRIO

1	INTRODUÇÃO	14
1.1	CONTEXTUALIZAÇÃO	14
1.2	MOTIVAÇÃO	15
1.3	OBJETIVOS	16
1.3.1	Objetivo Geral	16
1.3.2	Objetivos Específicos	16
1.4	ESTRUTURA DA DISSERTAÇÃO	17
2	FUNDAMENTAÇÃO TEÓRICA	19
2.1	COMPUTAÇÃO EM NUVEM	19
2.1.1	Características Principais e Benefícios	19
2.1.2	Modelos de Serviço	20
2.2	MODELOS DE PRECIFICAÇÃO	20
2.2.1	Modelos de Precificação Fixa	21
2.2.2	Modelos de Precificação Dinâmica	21
2.3	MÁQUINAS VIRTUAIS	22
2.3.1	Modelos de Precificação de Máquinas Virtuais	22
2.3.1.1	<i>Precificação Sob-Demanda</i>	22
2.3.1.2	<i>Instâncias Reservadas</i>	23
2.3.1.3	<i>Modelo de Precificação Spot</i>	23
2.3.1.3.1	<i>Instâncias Spot na AWS</i>	24
2.3.1.3.2	<i>Instâncias Spot na Azure</i>	24
2.3.1.3.3	<i>Instâncias Spot no Google Cloud Platform (GCP)</i>	24
2.4	TRABALHOS RELACIONADOS	25
2.4.1	Previsão de Preços Spot com Redes Neurais LSTM Empilhadas	25
2.4.2	Estimativa da Densidade de Probabilidade dos Preços Spot na AWS	27
2.4.3	Estratégias Combinadas de Análise de Custo e Disponibilidade em Mercados Spot	28
2.4.4	Previsão de Preços de Curto Prazo Usando Transições Probabilísticas em Instâncias Spot	29

2.4.5	Análise de Influência Geográfica na Precificação de Instâncias Spot da AWS	30
3	METODOLOGIA	33
3.1	CONJUNTO DE DADOS	33
3.2	SPOTLAKE	34
3.3	TRATAMENTO E TRANSFORMAÇÃO DOS DADOS	36
3.3.1	Agrupamento das regiões	38
3.3.2	Agrupamento dos tipos de instâncias	39
3.3.2.1	<i>Propósito Geral</i>	40
3.3.2.2	<i>Otimizado para Memória</i>	40
3.3.2.3	<i>Otimizado para Computação</i>	40
3.3.2.4	<i>Computação Acelerada</i>	40
3.3.2.5	<i>Otimizado para Armazenamento</i>	41
3.4	AMBIENTE E FERRAMENTAS ADOTADAS PARA EXECUÇÃO DO EXPERIMENTO	43
3.4.1	Ferramentas de desenvolvimento e bibliotecas utilizadas	43
3.4.2	Estratégias adotadas para otimização de desempenho	44
3.5	FLUXO GERAL DO EXPERIMENTO	45
3.6	MÉTODOS ESTATÍSTICOS APLICADOS	46
3.6.1	Análise Descritiva	46
3.6.2	Distribuição de Preços	46
3.6.3	Boxplots e Visualização de Outliers	47
3.6.4	Correlação e Análise de Dependência	47
3.6.5	Séries Temporais	47
4	ANÁLISES E RESULTADOS	49
4.1	ACESSO AOS GRÁFICOS GERADOS	49
4.2	ANÁLISE DESCRITIVA DOS DADOS	50
4.2.1	Histogramas e Distribuições de Preços	54
4.3	COMPARAÇÃO DOS PREÇOS SOB DEMANDA E SPOT ENTRE OS PROVEDORES	57
4.3.1	Análise dos preços sob demanda e spot: AWS	57
4.3.2	Análise dos preços sob demanda e spot: Azure	61
4.3.3	Análise dos preços sob demanda e spot: GCP	63

4.3.4	Análise comparativa dos mapas de calor entre os provedores, famílias de máquinas e regiões	66
4.4	VALIDAÇÃO ESTATÍSTICA E ANÁLISE DE CORRELAÇÃO	67
4.4.1	Análise das Correlações entre Provedores	68
4.5	AVALIAÇÃO DE PADRÕES TEMPORAIS E SAZONAIS	69
4.6	DISCUSSÃO E APLICABILIDADE PRÁTICA	74
4.6.1	Regiões mais vantajosas para cada provedor	74
4.6.2	Aplicabilidade prática: Estratégias para alocação multi-nuvem	75
4.6.3	Casos de Uso: Aplicação Prática dos Resultados	75
4.6.3.1	<i>Cenário 1: Processamento de Dados em Batch</i>	<i>76</i>
4.6.3.2	<i>Cenário 2: Aplicação Web Escalável</i>	<i>77</i>
4.6.4	Relevância para a pesquisa acadêmica e o mercado	78
4.6.5	Diferenciais da Pesquisa	78
4.7	INTEGRAÇÃO DOS RESULTADOS COM OS OBJETIVOS DA PESQUISA	79
4.7.1	Sumarização da Análise	80
5	CONCLUSÃO E TRABALHOS FUTUROS	81
5.1	PRINCIPAIS LIMITAÇÕES	81
5.2	TRABALHOS FUTUROS	82
5.3	CONTRIBUIÇÕES ACADÊMICAS E IMPLICAÇÕES TEÓRICAS	83
5.4	CONSIDERAÇÕES FINAIS	84
	REFERÊNCIAS	85
	APÊNDICE A – TABELAS DE AGRUPAMENTO DE INSTÂNCIAS	89
	APÊNDICE B – TABELAS DE CORRELAÇÃO DE PEARSON ENTRE PROVEDORES, REGIÕES E TIPOS DE INSTÂNCIA	100

1 INTRODUÇÃO

Nos últimos anos, o crescimento exponencial da computação em nuvem transformou a maneira como empresas e pesquisadores acessam recursos computacionais (MUSHTAQ et al., 2017). Entre os diversos serviços oferecidos, considerando o serviço de máquinas virtuais, o uso de instâncias spot e sob demanda se destaca como uma solução flexível e econômica, especialmente em ambientes multi-nuvem. Este trabalho busca analisar estratégias eficientes de precificação dessas instâncias, considerando ambientes compostos por múltiplos provedores de nuvem.

Para alcançar esse objetivo, realizou-se uma análise abrangente dos preços praticados pelos principais provedores, identificando padrões gerais e relações entre provedores e regiões, capazes de fornecer informações importantes para a tomada de decisão em alocação de recursos.

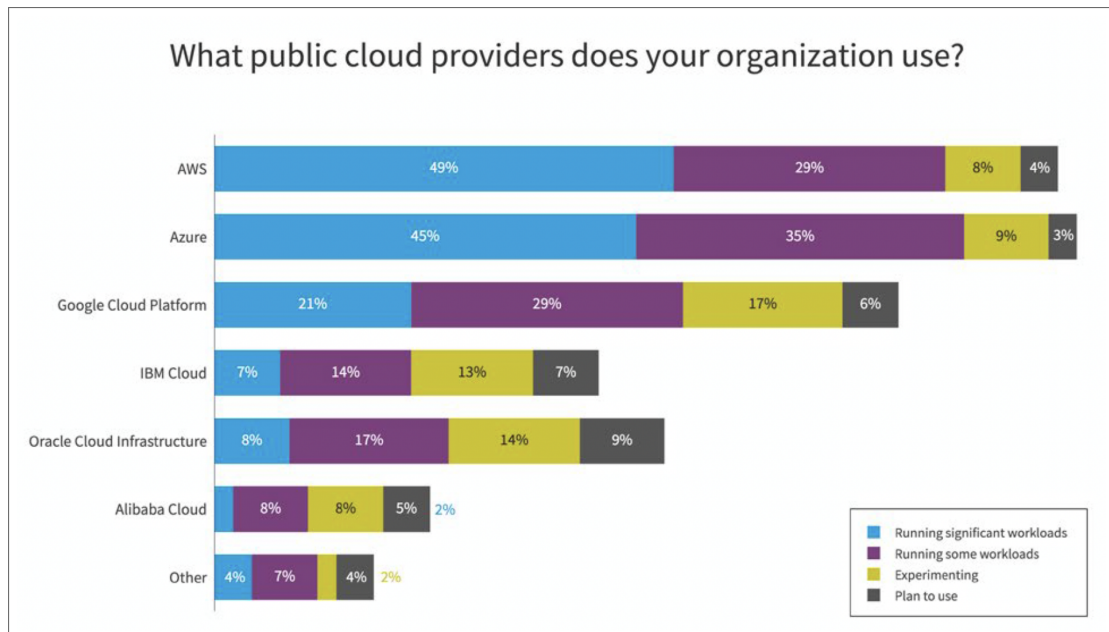
Os resultados obtidos oferecem subsídios para uma possível otimização de custos e aumento da eficiência operacional em cenários variados, contribuindo de forma prática tanto para o mercado quanto para futuras pesquisas acadêmicas na área.

Este capítulo apresenta o contexto geral em que o estudo está inserido, a motivação da pesquisa e os objetivos que orientaram sua execução. Além disso, destaca-se brevemente a relevância acadêmica e prática da pesquisa, ressaltando suas principais contribuições ao avanço do estado da arte em computação em nuvem.

1.1 CONTEXTUALIZAÇÃO

Com a crescente adoção de estratégias digitais, a computação em nuvem tornou-se um pilar fundamental para empresas e organizações que buscam agilidade, escalabilidade e redução de custos. Os três maiores provedores de nuvem pública — AWS, Azure e GCP, de acordo com o relatório Flexera Software (2024), vide figura 1 — oferecem uma ampla gama de serviços, dentre os quais se destacam as instâncias spot, conhecidas por sua disponibilidade variável e preços reduzidos em relação às instâncias sob demanda (WU et al., 2021).

Figura 1 – Principais provedores de nuvem.



Fonte: Flexera Software (2024)

Essas instâncias aproveitam recursos ociosos dos provedores, oferecendo custos atrativos para aplicações tolerantes a falhas ou de baixa criticidade. No entanto, a variabilidade nos preços spot, associada a diferenças regionais e políticas específicas de cada provedor, impõe desafios significativos para organizações que buscam otimizar seus gastos (LI et al., 2016). O uso de estratégias multi-nuvem surge como uma alternativa para explorar vantagens de diferentes provedores e regiões, tornando essencial uma análise detalhada das dinâmicas de precificação.

Além disso, a ausência de estudos que integrem os três principais provedores simultaneamente deixa lacunas importantes na literatura, dificultando a comparação direta entre suas estratégias de precificação. Este trabalho busca preencher essa lacuna, oferecendo uma visão abrangente e fundamentada que apoie tanto a academia quanto o mercado na tomada de decisões informadas.

1.2 MOTIVAÇÃO

A motivação para este trabalho surge da necessidade crescente de soluções eficientes para o gerenciamento de custos em ambientes multi-nuvem. Organizações enfrentam desafios para balancear a previsibilidade e a economia de custos ao escolherem entre instâncias reservadas, sob demanda e spot.

Embora instâncias spot ofereçam uma economia significativa, sua variabilidade de preços e a possibilidade de interrupções podem limitar sua aplicabilidade em determinados cenários. Por outro lado, o uso de estratégias multi-nuvem apresenta uma oportunidade de mitigar essas limitações, permitindo a exploração de preços mais baixos em diferentes provedores e regiões.

Outro ponto de motivação é a relevância acadêmica de estudar a precificação de instâncias em um contexto integrado entre os três maiores provedores. A análise de correlações entre regiões e provedores, somada à identificação de padrões de variabilidade de preços, contribui diretamente para o entendimento do funcionamento dos mercados de computação em nuvem e para o desenvolvimento de estratégias de alocação otimizadas.

Por fim, a possibilidade de construir ferramentas dinâmicas para análise e previsão de preços, com base nos resultados apresentados, representa uma oportunidade de impacto prático significativo para empresas que buscam maximizar a eficiência de seus recursos computacionais.

1.3 OBJETIVOS

Esta seção apresenta o objetivo geral e os objetivos específicos deste trabalho. A definição desses elementos busca orientar a condução da pesquisa, delimitar seu escopo e estabelecer os resultados que se pretendem alcançar com a análise proposta.

1.3.1 Objetivo Geral

O objetivo geral deste trabalho é analisar as flutuações de preços de instâncias spot e sob demanda e suas correlações entre regiões e provedores em um contexto multi-nuvem.

1.3.2 Objetivos Específicos

- Identificar regiões geográficas e tipos de instâncias mais vantajosas para cada provedor, considerando os preços spot e sob demanda.
- Analisar as características específicas que influenciam as dinâmicas de precificação entre provedores e regiões.
- Propor cenários práticos de economia, simulando cargas de trabalho em diferentes es-

estratégias de alocação.

- Explorar o potencial de estratégias multi-nuvem para reduzir custos e aumentar a eficiência operacional em ambientes de nuvem.

Ao longo dos capítulos subsequentes, serão abordados os métodos aplicados na pesquisa, bem como os resultados obtidos, oferecendo contribuições significativas tanto para o mercado quanto para a academia.

1.4 ESTRUTURA DA DISSERTAÇÃO

Esta dissertação está estruturada em cinco capítulos, que são organizados da seguinte forma:

- **Capítulo 1: Introdução** — Apresenta o contexto da pesquisa, a motivação que impulsionou este estudo, os objetivos gerais e específicos, bem como a relevância acadêmica e prática do tema abordado.
- **Capítulo 2: Fundamentação Teórica** — Explora os conceitos centrais da computação em nuvem, os diferentes modelos de precificação de máquinas virtuais, e uma revisão dos trabalhos relacionados, destacando as contribuições e limitações de estudos prévios.
- **Capítulo 3: Metodologia** — Detalha o conjunto de dados utilizado, o ambiente experimental e as ferramentas adotadas para a execução dos experimentos. Este capítulo também descreve os métodos estatísticos e de análise aplicados ao longo da pesquisa.
- **Capítulo 4: Análises e Resultados** — Apresenta os resultados obtidos a partir das análises realizadas, incluindo a descrição das flutuações de preços, correlações entre provedores e regiões, e os padrões temporais e sazonais identificados. Também são discutidas a aplicabilidade prática dos resultados e suas implicações no mercado e na academia.
- **Capítulo 5: Conclusão e Trabalhos Futuros** — Resume as principais contribuições desta pesquisa, identifica suas limitações e propõe direções para estudos futuros. Este capítulo também discute o impacto dos resultados na literatura acadêmica e no mercado.

Adicionalmente, são incluídos apêndices com tabelas complementares que detalham agrupamentos de instâncias e correlações entre provedores, oferecendo suporte adicional às análises realizadas nesta dissertação.

2 FUNDAMENTAÇÃO TEÓRICA

A transformação digital e a adoção crescente de tecnologias de nuvem têm remodelado as estratégias de Tecnologia da Informação (TI) em organizações de todos os tamanhos e setores. No centro dessa transformação está a computação em nuvem, que oferece escalabilidade, eficiência e capacidade de serviço sem precedentes. Um componente crucial e inovador dentro desse espectro é o uso de instâncias spot, que permite uma abordagem de precificação dinâmica baseada em oferta e demanda real do mercado. Esta seção explora os fundamentos teóricos que sustentam o estudo das instâncias spot e sua precificação em ambientes multi-nuvem, fornecendo o embasamento necessário para entender como diferentes modelos de precificação podem coexistir e influenciar a economia da computação em nuvem.

2.1 COMPUTAÇÃO EM NUVEM

Computação em nuvem é um modelo tecnológico que permite o acesso sob demanda a recursos de computação, como servidores, armazenamento, bancos de dados, redes e *software*, através da internet. Este modelo oferece uma alternativa eficaz ao armazenamento e processamento de dados em *hardware* local, proporcionando maior flexibilidade, escalabilidade e eficiência (KAUR; KAMBOJ, 2023). A ideia de computação como serviço começou a ganhar tração com a introdução de *mainframes* na década de 1950 e foi posteriormente impulsionada por tecnologias como virtualização e servidores dedicados (CHARD et al., 2017).

2.1.1 Características Principais e Benefícios

A principal característica da computação em nuvem é sua capacidade de escalabilidade e elasticidade, permitindo aos usuários expandir ou reduzir recursos conforme necessário, pagando apenas pelo que usam. Isso elimina a necessidade de grandes investimentos iniciais em infraestrutura de TI e reduz os custos operacionais, pois o gerenciamento e a manutenção da infraestrutura são responsabilidade dos provedores de serviços em nuvem. Além disso, a computação em nuvem oferece alta disponibilidade e confiabilidade, pois os dados podem ser replicados em múltiplos servidores em localizações geográficas distintas, garantindo continuidade de serviço mesmo em caso de falha de um servidor ou desastre natural (ABUALKISHIK;

ALWAN; GULZAR, 2020).

John McCarthy, um pioneiro da inteligência artificial, teve um papel crucial na concepção da computação em nuvem durante uma palestra no Massachusetts Institute of Technology (MIT) em 1961. Na ocasião, ele imaginou um futuro em que a computação poderia ser organizada como um serviço público — um conceito que ressoa fortemente com o modelo moderno de computação em nuvem. McCarthy propôs que, assim como a eletricidade e o telefone, os recursos computacionais poderiam ser usados e pagos conforme a necessidade, sugerindo a ideia de que o acesso à capacidade computacional poderia ser tão fácil e onipresente quanto o acesso à água corrente. Esta visão revolucionária antecipou a infraestrutura de serviços em nuvem de hoje, onde recursos de TI são oferecidos como serviços escaláveis e elásticos sob demanda, transformando fundamentalmente a maneira como as tecnologias da informação são consumidas e gerenciadas (KAUR; KAMBOJ, 2023).

2.1.2 Modelos de Serviço

Existem três modelos principais de serviço na computação em nuvem: Infraestrutura como Serviço (*IaaS - Infrastructure as a Service*), Plataforma como Serviço (*PaaS - Platform as a Service*) e Software como Serviço (*SaaS - Software as a Service*). IaaS fornece recursos de computação básicos, como capacidade de processamento e armazenamento, permitindo aos usuários executar qualquer *software* ou sistema operacional. PaaS oferece um ambiente de desenvolvimento e hospedagem na nuvem, facilitando para os desenvolvedores criarem aplicações web sem precisar gerenciar a infraestrutura subjacente. SaaS entrega aplicações de *software* como um serviço on-line, eliminando a necessidade de instalar e executar aplicações nos computadores individuais dos usuários (BERENBERG; CALDER, 2022).

2.2 MODELOS DE PRECIFICAÇÃO

A precificação de serviços em nuvem é um aspecto crítico do modelo de negócios de provedores de serviços em nuvem, que influencia tanto a adoção de serviços por consumidores quanto a rentabilidade dos provedores. Os modelos de precificação são desenvolvidos para refletir o custo de operação e manutenção da infraestrutura em nuvem, ao mesmo tempo em que maximizam a utilidade tanto para o provedor quanto para o usuário. Tradicionalmente, esses modelos podem ser classificados em duas categorias principais: precificação fixa e precificação

dinâmica (ALSHAREEF, 2023).

Os modelos de precificação adotados pelos provedores de nuvem desempenham um papel fundamental na utilização eficiente de máquinas virtuais. Cada modelo apresenta diferentes implicações sobre o custo, a previsibilidade e a acessibilidade dos recursos computacionais. Por exemplo, enquanto o modelo sob demanda oferece maior flexibilidade para atender a picos imprevisíveis de carga, ele também impõe custos significativamente mais altos em comparação com os modelos baseados em instâncias spot. Por outro lado, as instâncias spot representam uma oportunidade única de redução de custos, mas com o *trade-off* de possíveis interrupções (ZHANG et al., 2021).

Essas variações na precificação influenciam diretamente as estratégias de uso de máquinas virtuais, exigindo que os usuários escolham cuidadosamente entre custo e disponibilidade de recursos. Assim, compreender a relação entre os modelos de precificação e os diferentes tipos de máquinas virtuais torna-se essencial para uma análise fundamentada e para a seleção de instâncias mais adequadas a diferentes cenários de aplicação (ZHANG et al., 2020). Na seção 2.3.1, será discutido como os tipos de precificação de máquinas virtuais são categorizados e como suas características impactam essas decisões estratégicas.

2.2.1 Modelos de Precificação Fixa

A precificação fixa é o modelo mais simples e mais comumente adotado pelos provedores de serviços em nuvem. Neste modelo, os preços são estabelecidos antecipadamente e não variam em resposta às mudanças na demanda ou oferta. Os clientes pagam uma taxa fixa por um conjunto específico de recursos, como CPU, memória e espaço em disco, independentemente do uso real. Essa abordagem simplifica o processo de vendas e o planejamento financeiro para os clientes, mas pode não refletir eficientemente a utilização real dos recursos, levando a potenciais ineficiências econômicas e desperdício de recursos (ZHANG et al., 2021).

2.2.2 Modelos de Precificação Dinâmica

Por outro lado, a precificação dinâmica ajusta os preços com base em variáveis de mercado, como a demanda atual e a disponibilidade de recursos. Este modelo é mais flexível e pode levar a uma gestão de recursos mais eficiente. Um dos exemplos mais notáveis de precificação dinâmica é o modelo de instâncias spot, no qual os preços variam em tempo real de acordo

com as flutuações na oferta e na demanda. Este modelo permite que os provedores de serviços em nuvem otimizem a utilização de seus recursos ociosos, oferecendo-os a preços mais baixos quando a demanda é menor (LI et al., 2022). Trataremos desse modelo de precificação de máquinas virtuais em detalhes neste trabalho.

Do ponto de vista econômico, a precificação dinâmica pode ser mais vantajosa tanto para provedores quanto para consumidores. Para os provedores, maximiza a receita ajustando os preços de acordo com as condições de mercado e gerenciando melhor a capacidade. Para os consumidores, oferece a possibilidade de aproveitar preços mais baixos durante períodos de menor demanda. No entanto, essa abordagem requer que os consumidores sejam mais proativos e estratégicos em seu planejamento de uso, o que pode aumentar a complexidade da gestão de TI.

2.3 MÁQUINAS VIRTUAIS

Máquinas virtuais, do inglês (*Virtual Machines - VMs*) são emulações de computadores físicos que executam sistemas operacionais e aplicativos, comportando-se como uma entidade independente com recursos atribuídos de uma máquina física subjacente. Elas são fundamentais na computação em nuvem porque permitem a virtualização de recursos, o que significa que o *hardware* do servidor é abstraído e dividido em várias máquinas virtuais independentes que podem ser alocadas dinamicamente para atender a diferentes usuários e aplicações. Esta tecnologia é crucial para a elasticidade da nuvem, permitindo que os provedores de serviços em nuvem ofereçam recursos computacionais de maneira flexível e escalável.

2.3.1 Modelos de Precificação de Máquinas Virtuais

Os modelos de precificação de VMs em ambientes de nuvem variam amplamente, mas geralmente se agrupam em três categorias principais: precificação sob demanda, instâncias reservadas e modelos de precificação spot.

2.3.1.1 Precificação Sob-Demanda

A precificação sob demanda é o modelo mais simples e flexível, onde os usuários pagam pelas máquinas virtuais pelo tempo que as utilizam, sem compromissos de longo prazo. Este

modelo é ideal para empresas e desenvolvedores que precisam de recursos imediatos e estão lidando com cargas de trabalho variáveis ou experimentais. A vantagem do modelo sob demanda é que ele oferece máxima flexibilidade e elimina a necessidade de investimentos antecipados ou suposições sobre o uso futuro. No entanto, essa conveniência vem com um preço mais alto comparado a outras opções de precificação mais comprometidas (LEE; LIAN, 2017).

2.3.1.2 *Instâncias Reservadas*

Instâncias reservadas permitem que os usuários comprem capacidade de computação por um período estendido (tipicamente um ou três anos) em troca de um desconto significativo sobre as taxas sob demanda. Essas reservas são ideais para cargas de trabalho estáveis e previsíveis onde o uso contínuo é esperado. Os usuários se comprometem a usar uma certa quantidade de capacidade durante o período da reserva, o que permite que os provedores de nuvem planejem sua capacidade e otimizem a utilização de seus *data centers*. Esse modelo também beneficia os usuários com custos reduzidos, tornando a computação em nuvem economicamente viável para operações de longo prazo.

2.3.1.3 *Modelo de Precificação Spot*

As instâncias spot emergiram como uma solução inovadora para o aproveitamento da capacidade computacional ociosa em *data centers* de nuvem. Introduzidas inicialmente pela AWS (*Amazon Web Services*), de acordo com Ben-Yehuda et al. (2013a), elas representam um modelo de precificação dinâmico onde o preço é variável, baseado na oferta e demanda de recursos computacionais não utilizados. Esse modelo permite que os provedores de nuvem otimizem a utilização de seus recursos, enquanto oferecem aos clientes uma opção mais econômica para executar cargas de trabalho que são flexíveis em termos de tempo de execução.

O modelo de precificação spot é uma abordagem mais complexa e dinâmica. Ele permite que os usuários aproveitem a capacidade não utilizada dos provedores de serviços em nuvem a preços significativamente mais baixos que os modelos sob demanda. Este modelo é ideal para processos que podem tolerar interrupções, como cargas de trabalho *batch*, tarefas de processamento de dados em grande escala e cenários de teste e desenvolvimento que não são sensíveis ao tempo.

2.3.1.3.1 Instâncias Spot na AWS

A AWS foi pioneira no conceito de instâncias spot, lançando-as como parte de seu ecossistema EC2 (*Elastic Compute Cloud*). Essas instâncias permitem que os usuários façam lances de preço para a capacidade de máquina virtual que desejam usar. Se o preço do lance exceder o preço spot atual, que varia conforme as mudanças na oferta e na demanda, a instância é lançada automaticamente. No entanto, se o preço spot exceder o lance em qualquer momento durante a execução, a instância é terminada ou interrompida com uma notificação prévia de dois minutos (AWS, 2025).

Em 2017, a AWS implementou uma mudança específica em sua política operacional, conforme reportado em publicações de 2018 em blogs da AWS (A. W. S. Blogs, 2018). Essa alteração resultou em flutuações de preço menos frequentes e maior estabilidade de preços (BAUGHMAN et al., 2019). Esse ajuste nas políticas de precificação veio em resposta à necessidade de oferecer aos clientes um ambiente de cobrança mais previsível e estável, facilitando assim o planejamento financeiro e operacional das empresas que utilizam a infraestrutura da AWS. Essa mudança foi significativa, pois ajudou a consolidar a confiança dos usuários nos serviços da AWS, ao reduzir a incerteza associada aos custos de operação na nuvem.

2.3.1.3.2 Instâncias Spot na Azure

O Microsoft Azure oferece uma funcionalidade similar através de suas “Instâncias Spot”, que também aproveitam a capacidade computacional ociosa do Azure para oferecer custos reduzidos. Semelhante à AWS, essas instâncias estão disponíveis a preços significativamente menores em comparação com as opções de preço fixo. Entretanto, assim como as instâncias spot na AWS, elas podem ser interrompidas com pouca antecedência se a demanda por capacidade aumentar. O tempo prévio para emissão do sinal de interrupção no Azure é de trinta segundos (Microsoft, 2025).

2.3.1.3.3 Instâncias Spot no Google Cloud Platform (GCP)

No *Google Cloud Platform* (GCP), as instâncias spot são conhecidas como “*Preemptible VMs*”. Essas VMs oferecem uma opção ainda mais econômica para rodar cargas de trabalho de computação de alta escala que podem tolerar interrupções. As *Preemptible VMs* do GCP são instâncias que podem ser desligadas e reiniciadas pelo Google dependendo de suas necessidades

operacionais, mas, em troca, oferecem preços até 80% mais baratos do que as VMs padrão (Google Cloud Platform, 2025). O GCP não garante um aviso prévio para a interrupção dessas VMs e pode ou não emitir esse sinal em até um minuto de antecedência (Google Cloud, 2025).

O uso de instâncias spot requer uma estratégia bem pensada, pois a possibilidade de interrupção pode afetar significativamente o processamento se não for gerenciada corretamente. As empresas devem garantir que suas aplicações e dados possam ser rapidamente salvos e restaurados. Ferramentas de gestão de falhas e estratégias de *checkpointing* são críticas para minimizar o impacto de possíveis interrupções. Além disso, uma abordagem híbrida, usando tanto instâncias spot quanto instâncias reservadas ou sob demanda, pode ajudar a equilibrar custo e disponibilidade.

Em resumo, as instâncias spot revolucionaram o modo como as capacidades computacionais são vendidas e consumidas em plataformas de nuvem. Elas permitem uma utilização mais eficiente dos recursos de *data center*, oferecendo simultaneamente aos clientes uma maneira significativamente mais barata de processar cargas de trabalho que não são sensíveis ao tempo. Com cada provedor de nuvem oferecendo suas próprias variações e inovações nesse modelo, as instâncias spot continuam a ser uma área de crescente interesse e desenvolvimento no mundo da computação em nuvem.

2.4 TRABALHOS RELACIONADOS

A análise de precificação de instâncias em ambientes de nuvem é um campo amplamente estudado, com foco em diferentes abordagens e técnicas que buscam otimizar o uso de recursos computacionais e reduzir custos. Os trabalhos relacionados abrangem desde modelos preditivos de preços até análises estratégicas para alocação de cargas de trabalho em múltiplos provedores e regiões. Nesta seção, discutimos estudos relevantes que se alinham ao escopo desta dissertação, destacando suas contribuições e limitações em relação ao presente trabalho. A seguir, apresentamos uma discussão detalhada desses estudos.

2.4.1 Previsão de Preços Spot com Redes Neurais LSTM Empilhadas

O estudo apresentado por Chittora e Gupta (2020) investiga métodos avançados para prever preços de instâncias spot em serviços de nuvem, utilizando uma abordagem de redes neurais LSTM (*Long Short-Term Memory*) empilhadas. O objetivo principal é aprimorar a

precisão das previsões de preços spot, permitindo aos usuários realizar lances mais eficazes e econômicos em ambientes de nuvem. Para isso, os autores utilizaram um conjunto de dados de preços spot da *Amazon Web Services* (AWS) coletados ao longo de três meses, implementando um modelo LSTM de duas camadas. Esse modelo demonstrou resultados superiores quando comparado a modelos LSTM tradicionais e de três camadas, reduzindo o erro quadrático médio (RMSE) e o erro percentual absoluto médio (MAPE) de forma significativa.

Os resultados do trabalho de Chittora e Gupta (2020) indicam que o modelo proposto aumentou a precisão da previsão dos preços spot do dia seguinte em mais de 11% em relação a um LSTM padrão e mais de 52% em relação a um modelo de três camadas. Além disso, o RMSE e o MAPE foram consistentemente inferiores a 20% e 10%, respectivamente, em todas as cinco instâncias analisadas. No entanto, o estudo também identificou uma limitação significativa: o modelo empilhado de LSTM mostrou uma tendência a sobreajustar-se quando exposto a conjuntos de dados maiores, destacando a necessidade de melhorias para generalizar melhor os padrões de preços em diferentes condições.

Apesar dos avanços apresentados por Chittora e Gupta (2020), há limitações quando se compara o escopo e os objetivos do trabalho relacionado com esta dissertação. Enquanto o estudo de Chittora e Gupta (2020) restringiu-se a prever preços em curto prazo (próximo dia) com base em dados de três meses de uma única região da AWS, esta dissertação realizou uma análise mais ampla e estratégica. Utilizando um ano de dados históricos de preços, a pesquisa abrangeu três provedores globais — AWS, Azure e GCP — permitindo a exploração de dinâmicas de precificação em um contexto multi-nuvem.

Além disso, o trabalho relacionado concentrou-se exclusivamente na AWS, enquanto esta dissertação adotou uma abordagem multi-nuvem, investigando correlações entre regiões e provedores. Essa abordagem multi-nuvem é crucial para organizações que buscam otimizar suas estratégias de alocação de recursos em ambientes híbridos, oferecendo *insights* que vão além da previsão de preços em curto prazo.

Outro ponto de divergência significativa é o foco analítico. Enquanto o estudo de Chittora e Gupta (2020) priorizou a precisão preditiva para preços spot do dia seguinte, esta dissertação analisou de forma estratégica as dinâmicas de preços em diferentes tipos de carga de trabalho, considerando tanto custos quanto fatores de previsibilidade e confiabilidade das instâncias.

Por fim, o trabalho de Chittora e Gupta (2020) não aborda correlações entre regiões ou provedores, uma lacuna que esta dissertação preenche ao identificar padrões de correlação que podem auxiliar na alocação eficiente de cargas de trabalho em um contexto multi-nuvem.

Dessa forma, a presente dissertação não apenas complementa o estudo de Chittora e Gupta (2020), mas também amplia seu escopo ao fornecer uma análise abrangente e prática das dinâmicas de precificação em ambientes de computação em nuvem.

2.4.2 Estimativa da Densidade de Probabilidade dos Preços Spot na AWS

O estudo apresentado por Kabir et al. (2019) desenvolve uma técnica para calcular a densidade de probabilidade dos preços de instâncias spot da Amazon EC2. O objetivo principal do trabalho é fornecer uma ferramenta que permita aos usuários definirem lances mais eficientes, levando em consideração a urgência da tarefa e as condições do mercado. A abordagem utiliza o ajuste de curvas e a análise de similaridades históricas para estimar a densidade de probabilidade dos preços, permitindo uma melhor compreensão das incertezas associadas às variações de preços.

Os autores demonstraram que os padrões diários e semanais dos preços, bem como considerações relacionadas a feriados, podem prever de forma eficaz o comportamento dos valores das instâncias em determinados intervalos. Por exemplo, foi identificado que os preços tendem a ser mais altos durante a tarde nos dias de semana, com picos noturnos diários. Embora esses padrões ajudem a ajustar estratégias de lance, os autores reconhecem que mudanças nas condições do mercado ou nos padrões de uso global podem limitar a eficácia dessas previsões.

Ao comparar este trabalho com a presente dissertação, algumas diferenças cruciais podem ser destacadas. O estudo de Kabir et al. (2019) concentra-se exclusivamente no ambiente da AWS, limitando-se a padrões locais de preço e estratégias de previsão baseadas em dados históricos de 15 a 90 dias. Em contrapartida, esta dissertação adota uma perspectiva multi-nuvem, abrangendo três provedores globais (AWS, Azure e GCP) e um intervalo de análise de 12 meses. Isso permite uma avaliação mais abrangente e robusta das dinâmicas de precificação, capturando sazonalidades e correlações entre provedores e regiões.

Além disso, enquanto o trabalho relacionado foca na densidade de probabilidade como ferramenta de suporte à decisão em lances, o presente estudo amplia o escopo ao explorar correlações entre provedores e estratégias de alocação multi-nuvem. Essas análises oferecem subsídios para otimizar a alocação de cargas de trabalho, considerando não apenas o custo, mas também a previsibilidade e a confiabilidade das instâncias spot.

Uma limitação do estudo de Kabir et al. (2019) é a ausência de uma análise mais profunda sobre a influência de políticas de mercado e estratégias regionais na variabilidade dos preços.

Por outro lado, a pesquisa desta dissertação preenche essa lacuna ao explorar a influência do contexto regional e das características das instâncias na variabilidade dos preços. Isso fornece um panorama mais completo e alinhado às necessidades de empresas que buscam estratégias otimizadas em ambientes multi-nuvem.

Por fim, enquanto Kabir et al. (2019) abordam estratégias específicas para lances baseados em probabilidades, a presente dissertação propõe uma abordagem estratégica mais ampla, voltada para a análise de longo prazo e aplicações práticas. Isso destaca a relevância de ambos os estudos, com esta dissertação complementando e expandindo os resultados apresentados no trabalho relacionado.

2.4.3 Estratégias Combinadas de Análise de Custo e Disponibilidade em Mercados Spot

O estudo apresentado por Portella et al. (2023) propõe uma abordagem combinada utilizando métodos estatísticos e redes neurais LSTM para analisar mercados de instâncias spot na Amazon EC2. O objetivo central é oferecer uma estratégia baseada em utilidade que equilibre custo e disponibilidade para previsões de curto prazo, enquanto utiliza um mecanismo baseado em LSTM para analisar tendências de longo prazo nos preços spot. Com dados históricos de preços do primeiro semestre de 2020, os autores avaliam a eficácia do modelo em fornecer sugestões de preços competitivos e garantir durabilidade para as instâncias spot.

Entre os principais resultados, destaca-se que o modelo sugeriu preços máximos equivalentes a 37% do preço sob demanda para instâncias do tipo *r5.2xlarge*, garantindo, em média, disponibilidade superior a 5,73 horas com erro médio quadrático (MSE) inferior a 10^{-6} . No entanto, limitações como a dependência de dados históricos e a suscetibilidade a mudanças nas políticas de mercado são apontadas como desafios para a generalização do modelo em contextos mais amplos.

Ao comparar este trabalho com a pesquisa realizada nesta dissertação, observa-se uma abordagem complementar, mas com diferenças notáveis. Primeiramente, enquanto o estudo de Portella et al. (2023) foca exclusivamente na Amazon EC2, esta dissertação adota uma perspectiva multi-nuvem, abrangendo também os provedores Azure e *Google Cloud Platform* (GCP). Essa abordagem mais ampla é essencial para organizações que operam em ambientes híbridos e buscam estratégias de otimização em provedores variados.

Além disso, o trabalho relacionado concentra-se principalmente na predição de preços para

curto e médio prazo, utilizando análises temporais específicas. Diferentemente, esta dissertação explora não apenas as flutuações de preços, mas também as correlações entre regiões e provedores, um aspecto fundamental para identificar padrões de custo que podem ser explorados em estratégias multi-regionais.

Outra diferença relevante está na aplicação prática. Enquanto o estudo relacionado aplica sua abordagem combinada em um caso de estudo específico de bioinformática, esta dissertação se posiciona de maneira mais ampla ao oferecer análises aplicáveis a diversos cenários organizacionais, sem estar vinculada a um único domínio de aplicação.

Por fim, embora ambos os trabalhos contribuam para o entendimento das dinâmicas de precificação no mercado spot, esta dissertação amplia significativamente o escopo da análise, propondo *insights* e metodologias que atendem a um público mais diverso de usuários e organizações.

2.4.4 Previsão de Preços de Curto Prazo Usando Transições Probabilísticas em Instâncias Spot

Em Mishra, Kesarwani e Yadav (2019), os autores apresentam uma abordagem para prever preços de curto prazo de instâncias preemptíveis, com foco nas instâncias spot da Amazon EC2. Utilizando probabilidades de transição baseadas em históricos de preços, o algoritmo desenvolvido pelos autores busca prever o preço da próxima época (período de tempo). Essa técnica é baseada em uma matriz de probabilidade que modela as transições entre diferentes níveis de preços observados historicamente.

O objetivo principal do estudo é fornecer uma ferramenta que auxilie usuários a realizar lances mais eficazes e econômicos para instâncias spot, otimizando os custos de execução. Os resultados indicam que o modelo obteve uma taxa de erro percentual médio (MAPE) bastante baixa, com valores como 0,187% para a instância c3.2xlarge, demonstrando a eficácia do método proposto na previsão de preços com alta precisão.

Apesar dos resultados promissores, o estudo apresenta algumas limitações. A dependência de dados históricos pode ser problemática em cenários onde há mudanças abruptas nos padrões de preço, como alterações nas políticas de mercado ou picos de demanda inesperados. Além disso, a abordagem foca exclusivamente em previsões de curto prazo, deixando de explorar padrões sazonais ou tendências de longo prazo.

A presente dissertação amplia essa discussão ao abordar as dinâmicas de preços spot de

forma mais abrangente. Diferentemente do trabalho de Mishra, Kesarwani e Yadav (2019), que se concentra exclusivamente na Amazon EC2, a pesquisa realizada aqui adota uma perspectiva multi-nuvem, abrangendo três dos maiores provedores globais — AWS, Azure e *Google Cloud Platform* (GCP). Isso permite não apenas prever preços, mas também identificar correlações entre provedores e explorar estratégias de alocação de recursos em contextos híbridos e multi-nuvem.

Além disso, enquanto o estudo relacionado foca exclusivamente em previsões para a próxima época, esta dissertação oferece uma análise mais estratégica, investigando como padrões de precificação podem ser utilizados para otimizar custos em diferentes tipos de carga de trabalho. Outro diferencial é a extensão temporal: enquanto o trabalho relacionado analisa apenas dois meses de dados históricos, esta pesquisa abrange um período de doze meses, possibilitando a identificação de padrões sazonais e tendências de longo prazo.

Por fim, a abordagem probabilística de Mishra, Kesarwani e Yadav (2019) fornece previsões com alta precisão em curtos intervalos, mas limita-se à previsão de preços em uma única região e instância. A análise multi-nuvem apresentada nesta dissertação amplia significativamente o escopo, considerando variabilidade entre regiões e provedores, e oferece uma contribuição prática para organizações que buscam maximizar a eficiência de seus recursos computacionais em ambientes de nuvem dinâmica.

2.4.5 Análise de Influência Geográfica na Precificação de Instâncias Spot da AWS

Em Ekwe-Ekwe e Barker (2018), os autores investigam como a localização geográfica impacta os preços de instâncias spot na Amazon EC2. O estudo representa uma das primeiras análises detalhadas sobre o papel da localização nos custos gerais de implantação de instâncias spot. Utilizando dados de preços de instâncias coletados de diversas regiões e zonas de disponibilidade da AWS ao longo de 60 dias, os autores avaliaram padrões de volatilidade de preços e confiabilidade média. O objetivo principal era identificar regiões mais econômicas e com menor risco de interrupção para diferentes tipos de instâncias.

Os resultados indicaram que as regiões do Canadá apresentaram os preços mais baixos e estáveis para instâncias mais potentes, como a *i3.large*, enquanto regiões como Ásia-Pacífico exibiram maior confiabilidade para instâncias de menor potência. Essa variabilidade sugere que a escolha estratégica da região pode gerar economias significativas, dependendo dos requisitos das cargas de trabalho.

Embora o estudo forneça valiosos *insights* sobre as dinâmicas regionais da AWS, algumas limitações são evidentes. O período de análise de apenas 60 dias restringe a observação de padrões sazonais ou mudanças estruturais nos preços. Além disso, o escopo foi limitado à AWS, excluindo provedores concorrentes e uma análise comparativa mais ampla.

Em contrapartida, esta dissertação aborda a questão da precificação de instâncias spot em um contexto multi-nuvem, considerando não apenas a AWS, mas também Azure e GCP. Ao expandir o escopo para três provedores globais, o estudo proporciona uma visão mais abrangente e prática para organizações que operam em ambientes multi-nuvem. Essa abordagem permite identificar estratégias de alocação mais eficientes, aproveitando as melhores oportunidades econômicas entre provedores e regiões.

Outra diferença significativa é que, enquanto o estudo em questão foca apenas nas dinâmicas regionais da AWS, esta dissertação explora correlações entre regiões e provedores. Isso inclui a análise de flutuações de preços spot em múltiplos contextos e períodos mais longos, abrangendo doze meses. Essa perspectiva temporal mais ampla oferece uma base mais robusta para a tomada de decisões estratégicas em ambientes computacionais complexos.

Por fim, enquanto Ekwe-Ekwe e Barker (2018) sugerem padrões de precificação baseados em médias regionais, esta dissertação incorpora análise de variabilidade temporal e correlações entre preços, demonstrando como estratégias multi-nuvem podem não apenas reduzir custos, mas também mitigar riscos associados à interrupção de instâncias spot. Assim, esta pesquisa complementa e expande significativamente os *insights* apresentados por Ekwe-Ekwe e Barker (2018), avançando o estado da arte na análise de precificação em computação em nuvem.

A tabela 1 destaca os principais pontos de cada trabalho revisado em comparação a esta dissertação.

Tabela 1 – Comparação entre os Trabalhos Relacionados e a Presente Dissertação

Aspecto	(CHITTORA; GUPTA, 2020)	(KABIR et al., 2019)	(PORTELLA et al., 2023)	(MISHRA; KE-SARWANI; YADAV, 2019)	(EKWE-EKWE; BARKER, 2018)	Esta Dissertação
Provedor(es) analisado(s)	AWS	AWS	AWS	AWS	AWS	AWS, Azure e GCP
Escopo temporal	3 meses	Não especificado	6 meses (jan/-jun 2020)	Não especificado	60 dias	12 meses
Objetivo principal	Previsão de preços spot para otimizar lances	Cálculo de densidade de probabilidade para definir preços de lances	Predição de preços em curto e longo prazo para otimização de custos	Previsão de preços de curto prazo usando transições probabilísticas	Análise da influência geográfica nos preços de instâncias spot	Analisar flutuações de preços e correlações entre regiões e provedores em ambiente multi-nuvem
Metodologia	Redes neurais LSTM empilhadas	Ajuste de curvas e padrões históricos	Métodos estatísticos e redes neurais	Transições probabilísticas baseadas em histórico	Análise estatística de padrões regionais	Análise estatística descritiva, correlação de Pearson, séries temporais
Principais contribuições	Modelo de LSTM empilhado com alta precisão preditiva	<i>Insights</i> sobre padrões diários e sazonais nos preços spot	Combinação inovadora de métodos para otimização de custos	Algoritmo eficaz para previsão de curto prazo	Identificação de regiões mais econômicas e estáveis para instâncias spot	Análise multi-nuvem de flutuações de preços, com subsídios para estratégias de alocação
Principais limitações	Restrição ao escopo temporal curto e exclusividade à AWS	Aplicabilidade limitada devido às mudanças em padrões de preços	Dependência de dados históricos e foco exclusivo na AWS	Não se adapta bem a mudanças abruptas de mercado	Curto período de análise e ausência de perspectiva multi-nuvem	Limitação de granularidade de dados em algumas regiões e ausência de dados de interrupção para Azure e GCP
Amplitude da análise	AWS, 3 meses, enfoque regional	AWS, histórico limitado	AWS, 6 meses, predição curto e longo prazo	AWS, curto prazo e histórico de transições	AWS, 60 dias, análise geográfica	AWS, Azure e GCP, 12 meses, flutuações regionais e comparação entre provedores

Fonte: Elaborada pelo autor (2025).

3 METODOLOGIA

Neste capítulo, será descrito o método aplicado para a análise dos preços históricos das instâncias spot e sob demanda nos provedores de nuvem AWS, Azure e GCP, bem como o processo de tratamento dos dados utilizados.

3.1 CONJUNTO DE DADOS

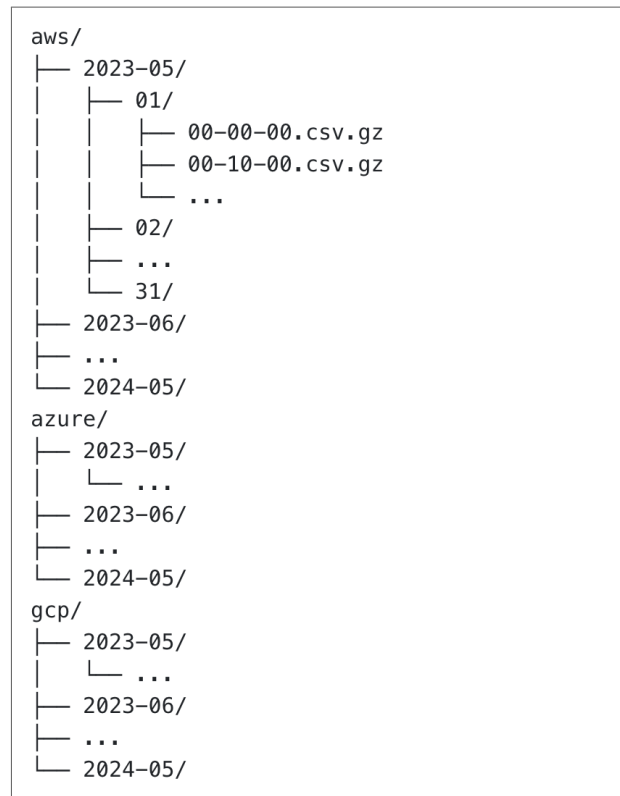
Os dados utilizados para a realização deste trabalho foram obtidos a partir do *dataset* Spotlake, que será descrito na seção 3.2. Após analisar a viabilidade do uso dos dados, foram realizadas interações por meio da ferramenta web (LEE; HWANG; LEE, 2025). No entanto, de acordo com os autores do SpotLake, não era possível obter os dados diretamente da ferramenta web devido ao grande volume. Para atender aos requisitos do experimento realizado, que envolvia um conjunto de dados de pelo menos 12 meses de coleta, abrangendo todas as regiões e famílias de máquinas da AWS, Azure e GCP, foi necessário entrar em contato com os autores do SpotLake. Foi utilizado o formulário disponibilizado na ferramenta web e também foi possível conversar com os autores durante o evento *International World Wide Web Conference 2023*. Durante essa ocasião, foram obtidas mais informações sobre os dados e também os contatos de e-mail.

Os dados de coleta de precificação das instâncias spot e sob demanda da AWS, Azure e GCP no período de Maio/2023 a Maio/2024 foram disponibilizados em arquivos *.csv.gz*. Essa coleta abrangeu um total de 12 meses. Os arquivos foram organizados em pastas separadas por provedor de nuvem, seguindo uma estrutura em árvore ilustrada na Figura 2. No segundo nível dessa estrutura, temos os diretórios referentes aos anos/meses de coleta. No terceiro nível, temos os diretórios correspondentes a cada dia do mês, e, por fim, no último nível, encontramos os arquivos *.csv.gz* contendo os dados coletados em um determinado horário, no padrão HH:mm:ss.

No caso da AWS, foram obtidos um total de **55.332** arquivos *.csv.gz*, totalizando **3.309.870** registros de precificação. Já para a Azure, foram coletados **9.186** arquivos, contendo um total de **9.769.182** registros de precificação. No caso da GCP, foram obtidos **8.913** arquivos, contendo ao todo **1.725.117** registros de precificação. A diferença no quantitativo de registros entre os provedores está relacionada diretamente a dificuldades de obtenção desses dados de

precificação, conforme descrito por Lee, Hwang e Lee (2022), bem como a quantidade de tipos de máquinas disponíveis por região e a presença geográfica desses provedores. A quantidade média de registros de preços ao longo de cada dia do período de meses analisados na AWS foi de 432 capturas (em torno de 1 captura a cada 4 minutos), na Azure 24 capturas (1 captura por hora) e no GCP 24 capturas.

Figura 2 – Árvore da estrutura dos arquivos disponibilizados.



Fonte: Elaborada pelo autor (2025)

3.2 SPOTLAKE

O SpotLake é um conjunto de dados históricos abrangendo várias instâncias spot, que também oferece um serviço de interface web para acessar esse histórico de preços. Os autores coletaram dados da AWS, Azure e GCP ao longo do tempo e realizaram uma análise para identificar as características das instâncias spot. A análise se concentrou principalmente na variação de preços e na frequência de interrupção das instâncias. O objetivo central do trabalho proposto pelos autores é de ser uma base de referência para trabalhos que investigam a precificação de instâncias spot para fins diversos (LEE; HWANG; LEE, 2022).

Além da coleta e análise de dados, os autores propuseram a centralização dessas infor-

mações com o objetivo de resolver desafios como a complexidade de acesso entre diferentes provedores, a limitação de acesso programático e a falta de dados históricos por períodos superiores a 3 meses. Além disso, a consulta desses preços varia de acordo com a região e a família de máquina em cada provedor de nuvem, o que torna ainda mais difícil obtê-los. O estudo aponta que, na época da realização, havia cerca de 547 tipos de instâncias, 17 regiões e 63 zonas de disponibilidade apenas na AWS, o que amplia significativamente o desafio de coletar informações de preços ao considerar também a Azure e a GCP.

Para verificar a precificação em um dado momento, $547 \text{ instâncias} \times 17 \text{ regiões} = 9.299$ consultas necessárias, onde havia um limite de 50 consultas únicas por conta. Diante disso, os autores Lee e Lee (1985) descrevem que conseguiram empacotar várias regiões em uma consulta, através do algoritmo bin-packing, conseguindo reduzir o número total de consultas necessárias para 2.226. Foi utilizada uma arquitetura sem servidor (*serverless*), quando aplicável, buscando diminuir os custos de operação. Além disso, foi utilizada a ferramenta SpotInfo (LEDENEV, 2021) para coleta dos dados de forma programática. A implementação atual do SpotLake mantém um conjunto de dados históricos para todas as regiões e instâncias da AWS, Azure e GCP. Os usuários podem consultar essas informações através de sua interface web (LEE; HWANG; LEE, 2025), especificando data/hora, região, tipos de instância, analisando individualmente para cada provedor de nuvem; porém, o acesso aos dados através da ferramenta web é delimitado devido ao grande volume. Os autores disponibilizam, através da ferramenta, um formulário para solicitação de acesso completo ao conjunto de dados.

Os autores mencionam que, na época, trabalhos relacionados na literatura tinham realizado análises detalhadas dos dados históricos de preços das instâncias spot, porém nenhum deles havia analisado conjuntos de dados que abrangiam a taxa de interrupção e a disponibilidade das instâncias spot. Além disso, devido à mudança no processo de precificação, conforme descrito por Ben-Yehuda et al. (2013b), os trabalhos anteriores na literatura tornaram-se obsoletos, o que aumentou ainda mais a relevância da base de dados SpotLake.

Para este estudo comparativo entre instâncias sob demanda e instâncias spot na AWS, Azure e GCP, o SpotLake foi fundamental para a fase experimental, entretanto algumas limitações foram encontradas, tais como a não disponibilização das informações de quantitativo de vCPU, RAM e Sistemas Operacionais das famílias de máquinas e a disponibilização da taxa de interrupção apenas para a AWS.

O serviço está atualmente disponível publicamente de acordo com Lee, Hwang e Lee (2025), permitindo que pesquisadores e desenvolvedores acessem o conjunto de dados que

pode ser utilizado para prever a disponibilidade de instâncias spot e planejar estratégias de custo-efetividade em suas aplicações.

3.3 TRATAMENTO E TRANSFORMAÇÃO DOS DADOS

Com o objetivo de melhorar a performance da análise dos dados, o conteúdo de cada arquivo .csv.gz foi extraído de maneira programática por meio de um script Python desenvolvido especificamente para essa tarefa. Este script possibilitou a transferência dos dados para uma base de dados orientada a documentos (MongoDB) que possui sua estrutura conforme a Figura 3.



Figura 3 – Estrutura da base MongoDB do experimento.

Fonte: Elaborada pelo autor (2025)

A coleção *instances_type* teve por objetivo armazenar os tipos de instâncias que possui similaridade entre os provedores de nuvem conforme discutido na seção 3.3.2 e conta com 7 documentos armazenados. A seguir um exemplo de um dos documentos dessa coleção.

Código Fonte 1 – Exemplo de documento da coleção *instances_type*

```

1 {
2   "_id": {
3     "$oid": "6660db5a0d93b05f576a0084"
4   },
5   "en": "General Purpose",
6   "pt_br": "Proposito Geral"
7 }

```

A coleção *regions* teve por objetivo armazenar o agrupamento das regiões entre os provedores de nuvem, conforme discutido na seção 3.3.1, levando em consideração a equivalência geográfica e proximidade. A coleção conta com 22 documentos armazenados, que representam as regiões agrupadas. A seguir, um exemplo de um dos documentos dessa coleção referente à região central da Europa.

Código Fonte 2 – Exemplo de documento da coleção *regions*

```
1 {  
2   "_id": {  
3     "$oid": "65c57ac89ab1f6404c0dac3d"  
4   },  
5   "alias": "eu_central",  
6   "aws_name": [  
7     "eu-central-1",  
8     "eu-central-2"  
9   ],  
10  "azure_name": [  
11    "centraluseuap",  
12    "dewestcentral",  
13    "plcentral"  
14  ],  
15  "gcp_name": [  
16    "europe-west3",  
17    "europe-west4"  
18  ]  
19 }
```

Por fim, a coleção *pricing* teve por objetivo armazenar o registro de precificação diário para cada provedor de nuvem, associando-o ao respectivo tipo de instância e região através do *instanceTypeId* e *regionId*. Para cada provedor, foram realizadas capturas ao longo do dia em intervalos de tempo variados, de acordo com Lee, Hwang e Lee (2022) e a coleção conta com 14.804.169 documentos armazenados. A seguir, um exemplo de um dos documentos dessa coleção referente ao tipo de instância p2.8xlarge da AWS na região

Código Fonte 3 – Exemplo de documento da coleção *pricing*

```
1 {  
2   "_id": { "$oid": "666269e353cdc49e1ae457db" },  
3   "provider": "aws",
```

```

4  "name": "p2.8xlarge",
5  "date": "2023-10-03",
6  "instanceTypeId": {"$oid": "6660d9fa0d93b05f576a0083"},
7  "regionId": {"$oid": "65c579df9ab1f6404c0dac38"},
8  "daily_on_demand_prices": [13.744, 13.744, 13.744, ..., 13.744],
9  "daily_spot_prices": [6.7478, 6.7661, 6.7243, ..., 6.7661]
10 }

```

É importante salientar que os dados brutos (RAW) obtidos do SpotLake não apresentaram um padrão uniforme de tabulação, nomes e ordem das colunas. Adicionalmente, para alguns tipos de máquinas, o valor para uma coluna específica estava ausente, o que requereu uma atenção especial durante o processo de preparação dos dados para análise. Este tratamento dos dados foi crucial para assegurar a integridade e a precisão dos resultados obtidos nas análises que serão discutidas no Capítulo 4.

3.3.1 Agrupamento das regiões

De acordo com as informações disponibilizadas em AWS Regions (2025), Azure Regions (2025) e GCP Regions (2025), cada um desses provedores de nuvem possui datacenters distribuídos ao redor do mundo, organizados em regiões e zonas de disponibilidade (*AZ - Availability Zones*). Estrategicamente, um provedor de nuvem pode ou não ter presença em determinado continente ou país, ou até mesmo possuir mais de uma zona de disponibilidade no mesmo país. Diante disso, com o objetivo de facilitar a comparação entre os provedores de nuvem, este trabalho adotou um agrupamento das regiões equivalentes entre a AWS, Azure e GCP, atribuindo a cada grupo de regiões um *alias* para representá-las. O critério para este agrupamento foi a proximidade geográfica entre os *datacenters* de cada provedor de nuvem, especialmente quando não existia uma equivalência direta. Tal informação foi extraída a partir das páginas web de cada provedor de nuvem e armazenada na coleção *regions* da base MongoDB. Este agrupamento está detalhado na tabela 2.

Conforme o documento de exemplo 2, a seguir está o detalhamento do *schema* adotado:

- **_id**: identificador do documento na coleção *regions*;
- **alias**: nome do grupo criado para representar a união das regiões similares de cada provedor de nuvem;

- **aws_name:** regiões da AWS com equivalência e proximidade geográfica com os demais provedores;
- **azure_name:** regiões da Azure com equivalência e proximidade geográfica com os demais provedores;
- **gcp_name:** regiões da GCP com equivalência e proximidade geográfica com os demais provedores

Tabela 2 – Agrupamento das regiões

Alias	Região AWS	Região Azure	Região GCP
af_south	af-south-1	zanorth, zawest	africa-south1
ap_east	ap-east-1	apeast	asia-east2
ap_northeast_1	ap-northeast-1	chnorth, jaeast	asia-east1, asia-northeast1
ap_northeast_2	ap-northeast-2	chwest, krcentral, krsouth	asia-northeast3
ap_northeast_3	ap-northeast-3	jawest	asia-northeast2
ap_south_1	ap-south-1	incentral	asia-south1
ap_south_2	ap-south-2	insouth	asia-south2
ap_southeast_1	ap-southeast-1	apsoutheast, aucentral, sesouth	asia-southeast1, australia-southeast1
ap_southeast_2	ap-southeast-2	aucentral2, aueast, ausoutheast	australia-southeast2
ca_central	ca-central-1	cacental, caeast	northamerica-northeast1, northamerica-northeast2
eu_central	eu-central-1, eu-central-2	centraluseap, dewestcentral, plcentral	europa-west3, europa-west4
eu_north	eu-north-1	denorth, noeast, nowest, secentral	europa-north1
eu_west_1	eu-west-1	eastus2euap, eunorth, euwest, frsouth	europa-west1, europa-west6
eu_west_2	eu-west-2	uknorth, uksouth, uksouth2, ukwest	europa-west2, me-central2, me-west1
eu_west_3	eu-west-3	frcentral	europa-west10, europa-west9
me_central	me-central-1	aenorth	europa-central2
me_south	me-south-1	aecentral	asia-southeast2, me-central1
sa_east	sa-east-1	brsouth, brsoutheast	southamerica-east1, southamerica-west1
us_east_1	us-east-1	useast	us-east1, us-east4
us_east_2	us-east-2	eastusslv, escentral, useast2, usnorthcentral	us-east5
us_west_1	us-west-1	uswest2	us-west2, us-west4
us_west_2	us-west-2	uscentral, ussouthcentral, uswest, uswest3, uswestcentral	us-central1, us-south1, us-west1, us-west3

Fonte: Elaborada pelo autor (2025)

3.3.2 Agrupamento dos tipos de instâncias

Assim como para as regiões, os provedores de nuvem AWS, Azure e GCP possuem uma grande variedade de tipos de instâncias com diferentes configurações de vCPU, memória RAM, arquitetura de processador, entre outras características. Esses tipos de instância são agrupados nos provedores através de famílias de máquinas, que representam categorias de instâncias com perfis de desempenho e preços semelhantes (AWS Instances Type, 2025), (Azure Instances Type, 2025), (GCP Instances Type, 2025).

Com o intuito de viabilizar a análise comparativa entre os provedores, este trabalho também definiu um agrupamento dos tipos de instâncias em famílias. Esse agrupamento foi realizado

com base nas especificações técnicas disponibilizadas pelos provedores, como número de vCPUs, memória RAM, arquitetura de processador e uso de recursos, buscando identificar perfis de instâncias equivalentes entre os provedores. Foi necessária a realização dessa etapa, pois o *dataset* SpotLake não continha tais informações no momento da realização deste experimento. É importante enfatizar que a lista de instâncias analisadas foi a que continha no conjunto de dados e não se trata de uma lista exaustiva se comparada à disponibilizada por cada provedor de nuvem. Portanto, o quantitativo de máquinas utilizadas neste experimento, considerando os dados disponibilizados e o agrupamento aplicado, consta nas seções a seguir.

3.3.2.1 *Propósito Geral*

Para máquinas consideradas de propósito geral, foram agrupadas na AWS 219 diferentes tipos de instâncias, na Azure 53 tipos de instância e no GCP 128 tipos de instâncias. O detalhamento consta na tabela 3 (vide também o apêndice A).

3.3.2.2 *Otimizado para Memória*

Para máquinas consideradas otimizadas para memória, foram agrupadas na AWS 234 diferentes tipos de instâncias, na Azure 102 tipos de instância e no GCP 23 tipos de instâncias. Em razão do tamanho da tabela, o detalhamento consta no apêndice A, tabela 7.

3.3.2.3 *Otimizado para Computação*

Para máquinas consideradas otimizadas para computação, foram agrupadas na AWS 175 diferentes tipos de instâncias, na Azure 62 tipos de instância e no GCP 88 tipos de instâncias. Em razão do tamanho da tabela, o detalhamento consta no apêndice A, tabela 8.

3.3.2.4 *Computação Acelerada*

Para máquinas consideradas de computação acelerada, foram agrupadas na AWS 70 diferentes tipos de instâncias, na Azure 1049 tipos de instância e no GCP 18 tipos de instâncias. Em razão do tamanho da tabela, o detalhamento consta no apêndice A, tabela 9.

3.3.2.5 *Otimizado para Armazenamento*

Para máquinas consideradas otimizadas para armazenamento, foram agrupadas na AWS 65 diferentes tipos de instâncias, na Azure 12 tipos de instância e no GCP 2 tipos de instâncias. Em razão do tamanho da tabela, o detalhamento consta no apêndice A, tabela 10.

Tabela 3 – Agrupamento de instâncias de propósito geral

AWS	Azure	GCP
m5a.8xlarge, m5d.large, m5dn.large, m6gd.4xlarge, m5zn.12xlarge, m6in.4xlarge, m7gd.xlarge, m6idn.metal, m2.2xlarge, m4.large, m5ad.large, m7a.48xlarge, m7i.metal-48xl, m5.8xlarge, m6gd.8xlarge, m6in.8xlarge, m6i.32xlarge, m5zn.large, t2.small, t4g.small, t3a.nano, m5ad.2xlarge, m4.4xlarge, m6id.12xlarge, m6gd.16xlarge, m7i-flex.xlarge, m5dn.8xlarge, m7i.24xlarge, m6gd.xlarge, m6i.metal, m7a.large, m6a.xlarge, m7gd.2xlarge, m7i-flex.8xlarge, t3.large, m5d.metal, m6idn.32xlarge, m6id.4xlarge, m5ad.16xlarge, m5a.xlarge, m4.xlarge, m7i.12xlarge, m5ad.4xlarge, t4g.nano, m7g.16xlarge, t3.micro, m5a.24xlarge, m7gd.16xlarge, m6g.metal, m4.10xlarge, m7a.32xlarge, t4g.xlarge, m1.large, m6in.xlarge, m7g.large, m5ad.12xlarge, m5.xlarge, m5n.xlarge, m7g.2xlarge, m6in.24xlarge, m7gd.large, m5a.16xlarge, m7i-flex.large, m6in.2xlarge, m6g.2xlarge, m5n.16xlarge, t4g.2xlarge, m6in.12xlarge, m5n.8xlarge, m5n.large, m7a.24xlarge, m5.24xlarge, m6id.xlarge, m6in.16xlarge, m6i.16xlarge, m7a.medium, m7gd.4xlarge, m5n.4xlarge, m6in.metal, m5zn.2xlarge, m6a.16xlarge, m6gd.metal, m6idn.8xlarge, m6id.8xlarge, m7g.xlarge, m7a.4xlarge, t3.2xlarge, m6a.large, m6a.48xlarge, m6idn.large, t2.medium, m7i.large, m3.medium, m5d.16xlarge, m5d.8xlarge, m7i.16xlarge, t3a.small, m6a.4xlarge, m6a.12xlarge, m5dn.2xlarge, m6in.large, m6g.16xlarge, m6a.32xlarge, m3.2xlarge, m5dn.16xlarge, m5dn.metal, m5ad.24xlarge, m7a.metal-48xl, m5a.4xlarge, m7i.48xlarge, m7a.2xlarge, m7gd.metal, m2.xlarge, m7i.4xlarge, m7a.xlarge, m5.metal, m5n.12xlarge, m6idn.16xlarge, m7g.medium, m6idn.4xlarge, t4g.micro, m7a.12xlarge, m5zn.metal, m6idn.24xlarge, m3.xlarge, m4.16xlarge, m7a.16xlarge, t2.large, m7i-flex.4xlarge, m5ad.8xlarge, m5zn.xlarge, m5dn.4xlarge, t4g.medium, m5.2xlarge, m6id.large, t2.xlarge, m6gd.12xlarge, m5dn.12xlarge, m7g.8xlarge, m5zn.6xlarge, m6a.2xlarge, m6gd.large, m6i.large, m6g.12xlarge, m5d.4xlarge, m5.4xlarge, t3a.large, m6a.metal, m5dn.xlarge, m6idn.12xlarge, m7gd.medium, m7g.4xlarge, m6gd.2xlarge, m6a.8xlarge, m6i.24xlarge, m7g.12xlarge, m7gd.8xlarge, t3a.2xlarge, m5zn.3xlarge, m6id.32xlarge, m5a.large, m6g.4xlarge, m5n.24xlarge, m7a.8xlarge, m6i.8xlarge, m6idn.xlarge, m6i.xlarge, t4g.large, m5d.2xlarge, m5n.metal, m4.2xlarge, t3.nano, m6i.12xlarge, m3.large, m5dn.24xlarge, t3a.micro, m6g.8xlarge, m7g.metal, m1.medium, t1.micro, m5a.12xlarge, m5n.2xlarge, m6g.xlarge, m6id.2xlarge, t3.xlarge, m6g.large, m6idn.2xlarge, m6in.32xlarge, m7i-flex.2xlarge, m7gd.12xlarge, m7i.2xlarge, m5a.2xlarge, m1.small, m6id.24xlarge, t3.medium, m5.16xlarge, m7i.metal-24xl, m7i.xlarge, m5ad.xlarge, m5d.xlarge, m5d.24xlarge, m1.xlarge, m6a.24xlarge, t2.2xlarge, m6id.16xlarge, m6i.2xlarge, m7i.8xlarge, m5d.12xlarge, t2.micro, m6gd.medium, m5.12xlarge, m5.large, m6g.medium, m6i.4xlarge, m6id.metal, t3a.xlarge, t3a.medium, m2.4xlarge, t3.small	basic_a3, basic_a4, basic_a0, basic_a1, basic_a2, _fxmdstypel, standard_d1_v2, standard_d2_v2, standard_d3_v2, standard_d4_v2, standard_d5_v2, standard_d2_v3, standard_d4_v3, standard_d8_v3, standard_d16_v3, standard_d32_v3, standard_d2_v4, standard_d4_v4, standard_d8_v4, standard_d16_v4, standard_d32_v4, standard_d2_v5, standard_d4_v5, standard_d8_v5, standard_d16_v5, standard_d32_v5, standard_ds1_v2, standard_ds3_v2, standard_ds5_v2, standard_b1ms, standard_b2s, standard_b2ms, standard_b4ms, standard_b8ms, standard_a1_v2, standard_a2_v2, standard_a4_v2, standard_a8_v2, standard_a2m_v2, standard_a4m_v2, standard_a8m_v2, standard_e2s_v3, standard_e4s_v3, standard_e8s_v3, standard_e16s_v3, standard_e32s_v3, standard_e64s_v3	t2a-standard-2, n1-highmem-2, n2d-highcpu-64, g2-standard-24, n2d-standard-16, n1-standard-2, n2d-highmem-2, n2-highmem-64, t2d-standard-16, n1-highcpu-96, n2d-highmem-16, n2-highmem-48, e2-highmem-8, e2-highcpu-32, n2d-highcpu-96, n1-standard-96, n1-highcpu-32, n1-standard-8, n2-highcpu-2, n2-standard-16, n2-standard-4, n1-ultramem-160, n2d-highcpu-128, n2d-standard-64, n2d-standard-8, n1-highmem-64, n2-standard-32, n1-highmem-8, e2-standard-2, n1-highmem-4, g2-standard-48, n1-highcpu-16, g2-standard-16, g2-standard-96, n2-highmem-32, n2-highmem-2, t2a-standard-4, g2-standard-12, n2d-highcpu-16, e2-highcpu-16, n2d-highmem-96, n2-highmem-128, n1-standard-1, n2-standard-2, n2-highcpu-64, n2-standard-96, n1-highcpu-8, e2-highcpu-2, e2-micro, n2d-highcpu-2, e2-highmem-2, n2d-standard-80, t2a-standard-8, t2a-standard-32, t2d-standard-8, n2-highcpu-4, n2d-highmem-32, g1-small, n1-ultramem-40, n2d-highmem-48, n2d-highmem-8, n1-highcpu-4, g2-standard-4, n2d-highcpu-4, n2d-highmem-64, n2-standard-64, t2a-standard-1, e2-standard-8, n2d-standard-48, n1-highmem-96, n2d-standard-96, n2d-highcpu-80, n2d-highcpu-224, n1-highcpu-2, n1-ultramem-80, n2d-highcpu-48, e2-highcpu-8, t2d-standard-48, t2d-standard-60, n2d-highcpu-32, n1-standard-32, e2-highmem-4, n2-highmem-80, t2d-standard-2, n2-highcpu-16, f1-micro, n2d-standard-2, n1-highmem-16, e2-small, n2-highmem-4, n1-standard-16, n1-megamem-96, n2-standard-8, e2-standard-32, n2-highmem-16, n2d-highmem-80, n2-highmem-8, g2-standard-8, t2a-standard-48, n2-highcpu-8, n2-highcpu-48, g2-standard-32, n2-standard-128, n2d-standard-32, e2-standard-16, n2d-standard-224, n2d-standard-4, e2-medium, t2a-standard-16, t2d-standard-4, n2-highcpu-32, n2-highcpu-80, n1-standard-4, n2-standard-80, n1-highcpu-64, n2d-highcpu-8, n2-standard-48, e2-highmem-16, n2d-standard-128, n2d-highmem-4, e2-highcpu-4, n1-standard-64, n1-highmem-32, n2-highcpu-96, t2d-standard-1, n2-highmem-96, t2d-standard-32, e2-standard-4

Fonte: Elaborada pelo autor (2025)

Durante a execução do experimento, foi identificado que algumas famílias de máquinas não estavam disponíveis em todas as regiões equivalentes analisadas para determinados provedores. Além disso, em alguns casos não houve coleta de preços no SpotLake. Essa situação resulta do processo de agrupamento adotado sobre os dados históricos do SpotLake, que visou padronizar a comparação entre regiões e famílias de máquinas com características similares, além de refletir a ausência de dados no próprio *dataset*. Esse agrupamento permitiu consolidar um volume expressivo de informações, totalizando mais de 14 milhões de registros de preços processados, garantindo uma análise abrangente. Embora algumas famílias específicas de máquinas tenham sido excluídas devido à falta de dados consistentes entre os provedores, o conjunto de dados resultante manteve ampla representatividade e relevância para os objetivos do estudo, preservando a integridade e a comparabilidade das análises realizadas.

3.4 AMBIENTE E FERRAMENTAS ADOTADAS PARA EXECUÇÃO DO EXPERIMENTO

Para realizar a extração, tratamento, transformação e carga dos dados brutos (*RAW*) provenientes do SpotLake, bem como a execução das análises exploratórias e estatísticas propostas neste trabalho, foi necessário o uso de uma infraestrutura computacional robusta, capaz de lidar com a volumetria expressiva dos dados envolvidos.

Optou-se pelo uso de máquinas virtuais EC2 da AWS, especificamente da família **x2gd.xlarge**, que oferecem recursos adequados ao tipo de carga de trabalho executada. Estas instâncias disponibilizam 4 vCPUs *ARM-based*, 64 GB de memória RAM e 500 GB de armazenamento SSD associado a um EBS (*Elastic Block Store*), garantindo alto desempenho no processamento e persistência dos dados. O sistema operacional adotado foi Ubuntu 20.04 LTS, bem como uma base de dados MongoDB executando em contêiner Docker.

3.4.1 Ferramentas de desenvolvimento e bibliotecas utilizadas

O desenvolvimento dos scripts responsáveis por todas as etapas do experimento foi realizado utilizando a linguagem de programação Python, na versão 3.9, devido à sua robustez e ao ecossistema de bibliotecas voltado para a manipulação e análise de grandes volumes de dados. Abaixo estão detalhadas as principais bibliotecas utilizadas:

1. **Pandas**: Biblioteca central para manipulação e análise de dados tabulares. Com suas es-

truturas de dados eficientes, como os *DataFrames*, o Pandas permitiu realizar a limpeza, filtragem e agrupamento dos conjuntos de dados de forma programática e estruturada.

2. **NumPy**: Amplamente utilizada na computação científica em Python, o NumPy foi essencial para operações matemáticas e manipulação eficiente de arrays multidimensionais, oferecendo o suporte necessário para cálculos vetorizados e numéricos.
3. **Matplotlib e Seaborn**: Para a geração de gráficos e visualizações, foram utilizadas as bibliotecas Matplotlib e Seaborn. O Matplotlib forneceu a flexibilidade necessária para criar visualizações customizáveis, enquanto o Seaborn foi utilizado para simplificar a criação de gráficos estatísticos.
4. **Apache Parquet**: Para otimizar o processamento e armazenamento intermediário dos dados, foi adotado o formato Apache Parquet. Trata-se de um formato de arquivo aberto, projetado para armazenar dados em colunas, permitindo compactação eficiente e acesso rápido, o que o torna ideal para tarefas analíticas.

3.4.2 Estratégias adotadas para otimização de desempenho

Apesar do ambiente computacional robusto e das ferramentas selecionadas, o alto volume de dados (aproximadamente 48 GB) apresentou desafios de desempenho, especialmente em análises que demandam processamento intensivo, como a geração de matrizes de correlação e *heatmaps* para múltiplas regiões e tipos de instâncias. Para mitigar tais dificuldades, foram implementadas as seguintes estratégias de otimização:

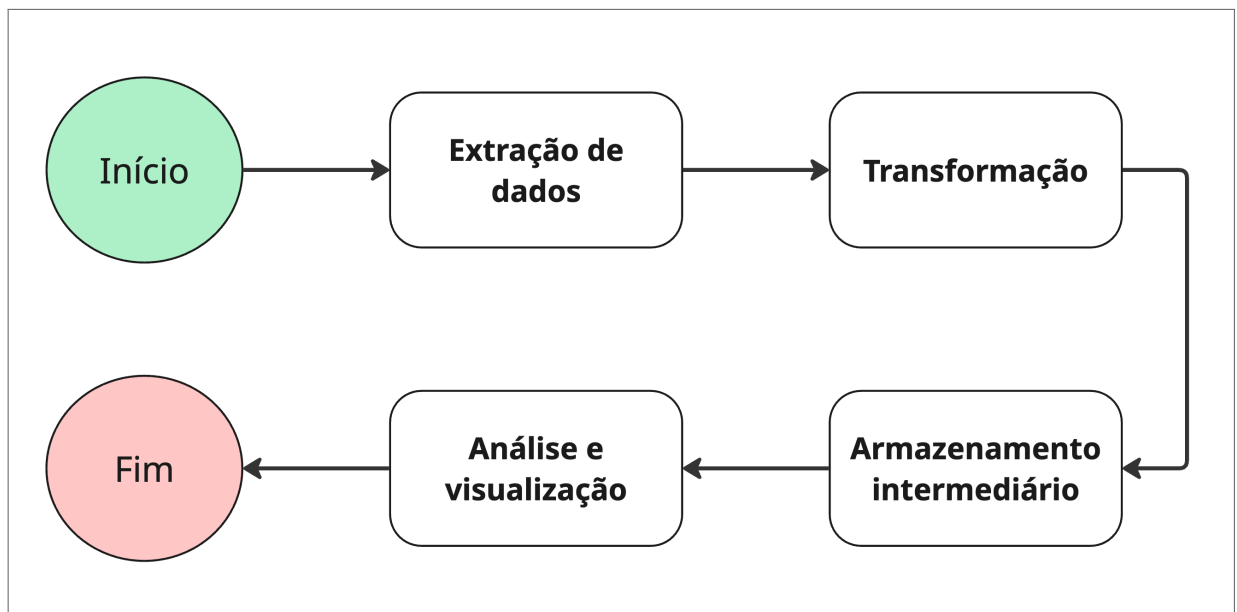
1. **Processamento em *chunks***: Ao invés de carregar todo o volume de dados na memória em uma única etapa, o processamento foi particionado em *chunks* menores, de tamanho configurável, de modo que cada fração dos dados fosse processada de forma independente. Este método permitiu reduzir o consumo de memória, mantendo a integridade dos resultados.
2. **Uso de cache intermediário com Apache Parquet**: Após o processamento de cada *chunk*, os resultados intermediários foram armazenados em arquivos Parquet. Esse formato não apenas reduziu o tempo de leitura e escrita, mas também permitiu que os *DataFrames* estruturados fossem reutilizados em análises subsequentes sem a necessi-

dade de processar novamente os dados brutos. Como resultado, a etapa de análise final pôde ser executada de maneira significativamente mais rápida e eficiente.

3. **Combinação de resultados:** Após o processamento de cada *chunk*, os resultados parciais foram combinados em um *DataFrame* unificado, possibilitando a geração de gráficos e métricas finais com desempenho otimizado. Esse método reduziu o tempo total de execução, evitando sobrecarga de memória e permitindo escalabilidade no processamento.

3.5 FLUXO GERAL DO EXPERIMENTO

O fluxo de execução do experimento pode ser resumido nas seguintes etapas:



Fonte: Elaborada pelo autor (2025)

1. **Extração:** Leitura programática dos dados brutos *.csv.gz* do SpotLake e armazenamento nas *collections* do MongoDB.
2. **Transformação:** Processamento dos dados em *chunks*, com aplicação de filtros e cálculos estatísticos.
3. **Armazenamento intermediário:** Escrita dos resultados parciais no formato Parquet para uso em análises posteriores.
4. **Análise e visualização:** Análise estatística e geração de gráficos com base nos dados estruturados.

3.6 MÉTODOS ESTATÍSTICOS APLICADOS

A análise dos preços das instâncias spot e sob demanda nos provedores AWS, Azure e GCP exigiu a aplicação de métodos estatísticos adequados para identificar padrões, variabilidades e correlações presentes nos dados. As técnicas estatísticas selecionadas foram escolhidas para fornecer uma visão detalhada das flutuações de preço, distribuição dos valores e relações entre os provedores ao longo do tempo e entre diferentes tipos de instâncias.

3.6.1 Análise Descritiva

A análise descritiva foi a primeira etapa e teve como objetivo resumir e compreender o comportamento dos preços das instâncias spot e sob demanda. Nesta fase, foram utilizadas as seguintes métricas:

- **Média aritmética:** Para avaliar o preço médio ao longo do tempo.
- **Mediana:** Indicador robusto para detectar valores centrais, minimizando a influência de outliers.
- **Desvio padrão:** Para medir a dispersão dos preços em relação à média.
- **Mínimo e máximo:** Valores limites identificados em cada conjunto de dados.

Essas estatísticas foram aplicadas separadamente para os preços spot e sob demanda, permitindo uma comparação direta entre os modelos de precificação.

3.6.2 Distribuição de Preços

Para analisar a distribuição dos preços das instâncias, foram construídos gráficos específicos, como:

- **Histogramas:** Para visualizar a frequência dos valores em intervalos específicos.
- **Gráficos de densidade:** Utilizando a técnica de Kernel Density Estimation (KDE), facilitando a análise visual das tendências de distribuição.

Essa abordagem permitiu identificar se os preços seguiam distribuições normais, assimétricas ou apresentavam picos específicos, revelando padrões ocultos em diferentes regiões e tipos de instâncias.

3.6.3 Boxplots e Visualização de Outliers

Boxplots foram gerados para comparar a variação dos preços entre meses, regiões e provedores. Eles permitiram observar:

- **Mediana e intervalos interquartis (IQR):** Indicando a dispersão central dos preços.
- **Outliers:** Preços fora do padrão, que podem sinalizar instabilidades ou condições de mercado específicas.

Essas visualizações foram particularmente úteis para avaliar a estabilidade dos preços sob demanda e spot entre os provedores.

3.6.4 Correlação e Análise de Dependência

A análise de correlação foi aplicada para investigar a relação entre os preços das instâncias spot nos diferentes provedores e entre regiões equivalentes. As métricas utilizadas incluem:

- **Coefficiente de Correlação de Pearson:** Para medir a relação linear entre os preços, com valores variando entre -1 e 1.
- **Matriz de Correlação:** Visualizada através de *heatmaps* gerados com as bibliotecas Seaborn e Matplotlib.

Essa análise forneceu *insights* sobre o nível de dependência ou independência dos preços entre os provedores, ajudando a identificar possíveis padrões sincronizados de precificação.

3.6.5 Séries Temporais

Para compreender a variação dos preços ao longo do tempo, foram aplicadas técnicas de análise de séries temporais:

- **Gráficos de linhas:** Para visualizar a evolução dos preços mensalmente.
- **Médias móveis (*Moving Average*):** Aplicadas para suavizar flutuações diárias e destacar tendências de longo prazo.

A análise das séries temporais permitiu observar tendências sazonais, como aumentos ou quedas de preço em determinados períodos, além de variações entre as regiões e tipos de instâncias.

Com a aplicação das técnicas descritas ao longo deste capítulo, foi possível estruturar e preparar os dados necessários para a realização das análises propostas. A combinação de ferramentas computacionais e métodos estatísticos apropriados permitiu o processamento eficiente da grande volumetria de dados obtidos do SpotLake. Estratégias como o uso de *chunks* para processamento em paralelo e o armazenamento intermediário em arquivos Parquet foram essenciais para otimizar o desempenho dos experimentos e garantir a integridade dos resultados. No próximo capítulo, serão apresentados os resultados das análises realizadas, destacando a variação dos preços das instâncias spot e sob demanda, as correlações entre os provedores e as tendências temporais observadas para diferentes tipos de instâncias e regiões. Essas análises visam responder às perguntas de pesquisa e fornecer *insights* relevantes para o entendimento do comportamento de precificação de máquinas spot e sob demanda em ambientes multi-nuvem.

4 ANÁLISES E RESULTADOS

Neste capítulo, são apresentados os resultados obtidos a partir da análise de precificação de instâncias spot nos provedores de nuvem AWS, Azure e *Google Cloud Platform* (GCP). O objetivo foi compreender as flutuações de preços, identificar padrões regionais e temporais, e avaliar as correlações entre os provedores em diferentes contextos geográficos e tipos de instância. Para isso, foram utilizados dados históricos do *dataset* SpotLake, abrangendo um período de doze meses, de maio de 2023 a maio de 2024. A análise foi estruturada para oferecer uma visão abrangente e detalhada sobre as dinâmicas de precificação em ambientes multi-nuvem.

Primeiramente, foi realizada uma análise descritiva dos preços, observando as tendências temporais e as variações de preços por provedor, região e tipo de instância. Esses resultados foram representados por meio de gráficos de séries temporais e boxplots, que destacam tanto os padrões gerais quanto os desvios provocados por *outliers*.

Em seguida, foi explorada a distribuição dos preços utilizando histogramas e gráficos de densidade. Essa abordagem permitiu identificar padrões específicos em cada provedor, destacando as diferenças regionais e as características únicas dos modelos de precificação dinâmicos.

Para ampliar a compreensão, foi realizada uma comparação direta entre os provedores, investigando quais oferecem os preços mais competitivos ou estáveis para instâncias equivalentes em diferentes regiões. Essa análise buscou fornecer *insights* práticos para empresas que pretendem adotar estratégias multi-nuvem.

Por fim, foram analisadas as correlações entre os provedores e as regiões, utilizando mapas de calor para evidenciar o grau de interdependência nos preços das instâncias spot. Essa etapa foi crucial para entender como as políticas de precificação de cada provedor podem influenciar o comportamento do mercado.

Essas análises são complementadas por discussões práticas que conectam os resultados com os objetivos da pesquisa, destacando sua aplicabilidade para o mercado e seu valor acadêmico.

4.1 ACESSO AOS GRÁFICOS GERADOS

Com o objetivo de permitir análises complementares, todos os gráficos gerados a partir deste trabalho estão disponíveis no repositório digital associado¹. Esses gráficos incluem:

¹ Google Drive Repository.

- Boxplots de preços spot e sob demanda para todas as regiões e famílias de máquinas.
- Histogramas e gráficos de densidade (KDE) para todas as combinações de provedores, regiões e famílias.
- Mapas de calor representando a distribuição dos preços ao longo do tempo.
- Gráficos de linhas da variação percentual de preços spot e sob demanda para todas as regiões e famílias de máquinas.
- Arquivos .csv contendo a matriz de correlação de Pearson e valores de medidas estatísticas e densidade.

Esse material complementa as análises apresentadas no Capítulo 4, oferecendo uma visão mais abrangente dos padrões de precificação em diferentes contextos.

4.2 ANÁLISE DESCRITIVA DOS DADOS

Para analisar e comparar os preços de instâncias, a mediana foi escolhida como a medida representativa para os preços mensais em vez da média. A principal justificativa para essa escolha está na natureza da distribuição dos preços spot, que tende a apresentar uma variabilidade significativa e, frequentemente, *outliers* ou picos momentâneos de preço. Esses *outliers* podem ocorrer devido a picos pontuais de demanda, políticas de gerenciamento de recursos dos provedores ou eventos específicos que causam aumentos temporários nos preços. Tais variações abruptas impactam de maneira desproporcional a média dos preços, distorcendo a representação do custo típico de uma instância spot ao longo do mês.

A mediana, por outro lado, é uma medida de tendência central que é mais resistente a *outliers* e melhor reflete o comportamento típico dos preços durante um período. Em vez de ser afetada pelos valores extremos, a mediana capta o ponto médio da distribuição dos preços, proporcionando uma visão mais estável e representativa do custo ao longo de cada mês. Em outras palavras, ao adotar a mediana, conseguimos minimizar o impacto de variações extremas, enfatizando o valor central da distribuição de preços e permitindo comparações mais confiáveis entre meses, regiões e tipos de instância.

A representação da distribuição de preços ao longo de cada mês por meio de boxplots complementa essa análise. Os gráficos de boxplot ilustram visualmente a amplitude e a dispersão dos preços, destacando a mediana, o intervalo interquartil (que representa a maior parte

dos preços observados), além de eventuais valores extremos que se destacam da distribuição central. Tais gráficos permitem observar como os preços spot e sob demanda variam mês a mês, proporcionando a compreensão da justificativa para a escolha da mediana. Além disso, as linhas de média adicionadas aos boxplots permitem contrastar essa medida com a mediana e reforçam a importância de utilizar a mediana como uma medida representativa em contextos onde a distribuição dos preços apresenta alta variação.

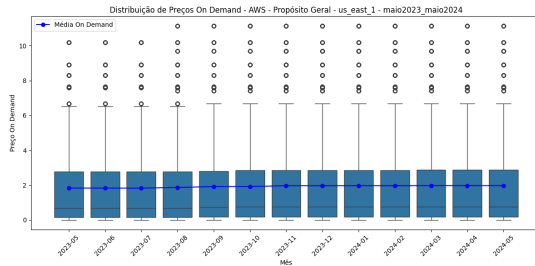
Os gráficos de boxplot apresentados na figura 4 ilustram a variação de preços ao longo de doze meses (maio de 2023 a maio de 2024) para a região **us_east_1** nos provedores AWS, Azure e GCP, abrangendo tanto instâncias sob demanda quanto instâncias spot. A região **us_east_1** foi selecionada por ser amplamente utilizada em diversas cargas de trabalho e por representar um cenário relevante para análise devido à sua alta demanda. Para manter a clareza e a objetividade, optamos por apresentar os dados apenas para instâncias da família de propósito geral. Foi adotado para o eixo X cada mês do período analisado e no eixo Y o preço em dólar (USD) tanto para o modelo spot quanto sob demanda.

Devido às diferenças nos modelos de precificação dos provedores, as escalas dos gráficos podem variar entre eles e entre os tipos sob demanda e spot, o que reflete as características específicas de cada provedor e suas estratégias de precificação. É importante observar essas variações ao interpretar os gráficos, pois elas reforçam a análise das dinâmicas de preços descritas no texto.

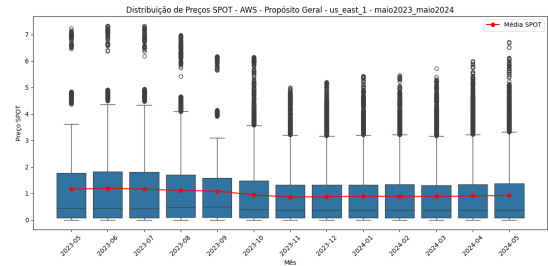
A partir dos gráficos de boxplot apresentados, é possível observar importantes aspectos estatísticos sobre os preços das instâncias spot e sob demanda dos provedores AWS, Azure e GCP na região **us_east_1**, no período de maio de 2023 a maio de 2024. Esses gráficos revelam a presença de uma quantidade significativa de *outliers*, que representam flutuações abruptas de preço. Essas flutuações são frequentemente associadas a picos temporários de demanda ou ajustes dinâmicos realizados pelos provedores. O GCP, por exemplo, apresenta uma maior dispersão nos preços das instâncias spot, indicando uma política de precificação mais volátil em comparação aos outros provedores. Em contrapartida, o Azure se destaca pela estabilidade em seus preços sob demanda, com intervalos interquartis mais estreitos e menor presença de valores extremos.

Ao longo do período analisado, as **tendências temporais reveladas** pelos gráficos mostram que a média e a mediana dos preços seguem comportamentos similares em muitos momentos, mas divergem especialmente nas instâncias spot do GCP e da AWS, onde a presença de *outliers* tem maior impacto na média. Um exemplo disso é a leve tendência de queda

AWS

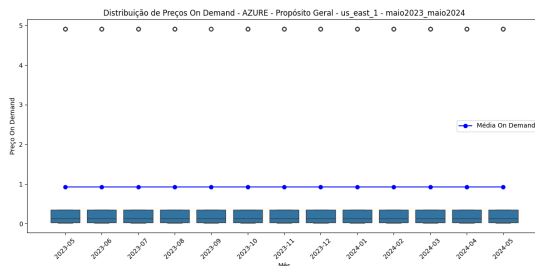


(a) Instâncias Sob Demanda

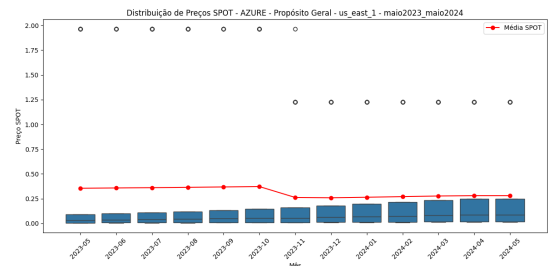


(b) Instâncias Spot

Azure

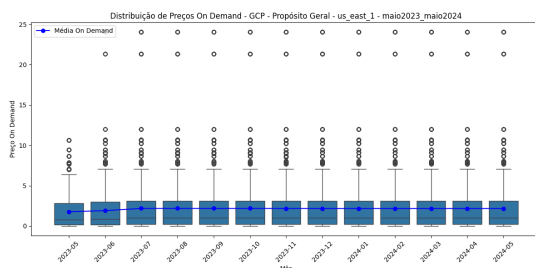


(c) Instâncias Sob Demanda

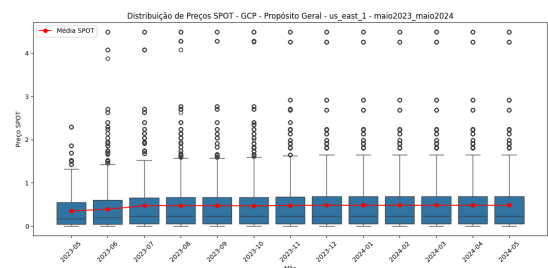


(d) Instâncias Spot

GCP



(e) Instâncias Sob Demanda



(f) Instâncias Spot

Figura 4 – Distribuição de preços sob demanda e spot para os provedores AWS, Azure e GCP na região *us_east_1*, considerando máquinas de propósito geral no período de maio de 2023 a maio de 2024 – Observe as variações ao interpretar os gráficos, pois elas reforçam a análise das dinâmicas de preços descritas no texto.

Fonte: Elaborada pelo autor (2025)

observada nos preços spot da AWS ao longo do período, o que pode refletir ajustes sazonais ou mudanças na demanda. Já os preços do Azure e do GCP mantiveram maior estabilidade, principalmente nas instâncias sob demanda, reforçando a previsibilidade oferecida por esses provedores.

Outra análise interessante se refere à **comparação entre os preços spot e sob demanda**. Como esperado, os preços sob demanda são consistentemente mais altos, devido ao modelo de precificação baseado na garantia de disponibilidade. Contudo, as instâncias spot apresentam uma variabilidade significativamente maior, evidenciando a dinâmica de mercado dessa modalidade. Essa característica destaca a importância de avaliar o balanço entre custo e previsibilidade ao selecionar instâncias para diferentes tipos de cargas de trabalho.

Além disso, as diferenças entre os provedores revelam **estratégias distintas de precificação**. A AWS, por exemplo, apresentou uma maior quantidade de *outliers* em seus preços spot, o que sugere maior sensibilidade a picos de demanda. O Azure demonstrou estabilidade em ambos os tipos de instâncias, enquanto o GCP exibiu maior dispersão nos preços, refletindo sua maior variabilidade de mercado. Esses padrões reforçam a necessidade de considerar as particularidades de cada provedor ao planejar cargas de trabalho em ambientes multi-nuvem.

Essas análises trazem implicações práticas importantes. Para cargas de trabalho que exigem previsibilidade de custos, o Azure se mostra uma opção atrativa. Por outro lado, o GCP pode oferecer oportunidades de economia em situações onde a variabilidade de preços seja tolerável. A escolha da região **us_east_1** como objeto de análise reflete sua relevância como uma das mais utilizadas para diferentes cargas de trabalho, mas as observações feitas aqui podem variar para outras regiões. Dessa forma, estratégias multi-nuvem devem considerar não apenas os custos, mas também a dinâmica de precificação e a variabilidade de cada provedor.

Com base nesses gráficos, conclui-se que a mediana é uma medida mais adequada para representar os preços das instâncias, pois é menos influenciada pelos *outliers* que distorcem a média. Ainda assim, a análise detalhada das distribuições e das características específicas de cada provedor oferece insights valiosos para o planejamento estratégico em ambientes de nuvem pública.

É importante ressaltar que as conclusões apresentadas são baseadas na análise dos gráficos da região **us_east_1** e para a família de máquinas de propósito geral. Essas observações podem variar significativamente entre outras regiões e agrupamentos de famílias de máquinas, dadas as diferenças nas políticas de precificação e dinâmicas de demanda específicas de cada cenário. Considerando a grande quantidade de dados analisados, foram gerados um total de

406 gráficos de boxplot, abrangendo diversas combinações de regiões e famílias de máquinas. Contudo, pela limitação de espaço, não é viável comentar detalhadamente todos esses gráficos nesta seção. Para permitir uma análise mais abrangente, todos os boxplots gerados estão disponíveis no repositório mencionado em 4.1.

4.2.1 Histogramas e Distribuições de Preços

Os histogramas e gráficos de densidade (KDE) foram gerados para ilustrar a distribuição dos preços sob demanda e spot para a região **us_east_1** com instâncias da família de Propósito Geral. Esses gráficos ajudam a identificar os padrões de precificação e variações observadas no período de análise, de maio de 2023 a maio de 2024.

Abaixo estão as principais observações para os provedores AWS, Azure e GCP:

- **AWS:** Os preços sob demanda apresentam uma concentração em faixas mais baixas (predominantemente abaixo de \$2 USD), enquanto os preços spot são ainda menores, com picos de densidade próximos a \$0.5 USD. No entanto, os preços spot mostram maior dispersão, refletindo a flexibilidade e a volatilidade da precificação dinâmica desse provedor.
- **Azure:** Os preços sob demanda mantêm uma distribuição mais uniforme e estável, com densidades concentradas em torno de \$1 USD. Os preços spot, por outro lado, apresentam menor variabilidade, mas destacam valores muito baixos com raros *outliers*, indicando uma política de preços previsível para ambos os tipos de instância.
- **GCP:** O GCP apresentou a maior dispersão de preços entre os provedores. Os preços sob demanda se estendem para valores mais altos, enquanto os preços spot exibem alta densidade para valores baixos, mas com variação significativa, indicando uma política de precificação mais agressiva e dinâmica.

A Figura 5 apresenta os histogramas e gráficos de densidade para os três provedores mencionados, evidenciando os padrões descritos acima.

A análise dos histogramas reforça os padrões descritos, mostrando que:

- Os preços sob demanda possuem uma distribuição mais concentrada em valores específicos, principalmente para o Azure.

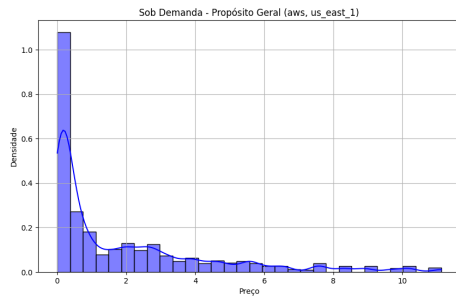
- Os preços spot, apesar de oferecerem vantagens econômicas significativas, apresentam uma dispersão maior, especialmente no GCP, sugerindo que sua volatilidade deve ser considerada no planejamento de workloads.

Os histogramas e gráficos de densidade apresentados para a região *us_east_1*, focados na família de propósito geral, refletem claramente as diferenças nas políticas de precificação dos provedores.

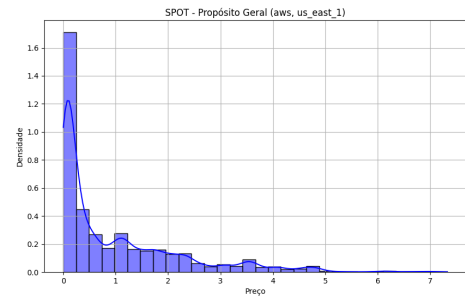
- **AWS:** Nota-se uma maior dispersão nos preços spot, indicando que esse provedor apresenta maior sensibilidade a flutuações na demanda. Tal comportamento pode ser resultado de ajustes dinâmicos frequentes em resposta a variações de capacidade ociosa.
- **Azure:** As distribuições de preços sob demanda exibem menor variabilidade, com histogramas que apresentam picos bem definidos e concentrados. Isso sugere uma política de preços previsível, adequada para cargas de trabalho críticas que demandam estabilidade.
- **GCP:** A distribuição de preços spot no GCP é a mais ampla entre os três provedores, evidenciando uma maior volatilidade. Essa característica pode oferecer oportunidades de economia em situações onde a variação de custos seja tolerável.

A análise visual dos histogramas valida ainda mais a escolha da mediana como métrica principal para representar os preços. Enquanto a média é fortemente impactada pelos valores extremos evidentes nas distribuições de spot, a mediana reflete de forma mais estável o custo típico de uma instância.

AWS

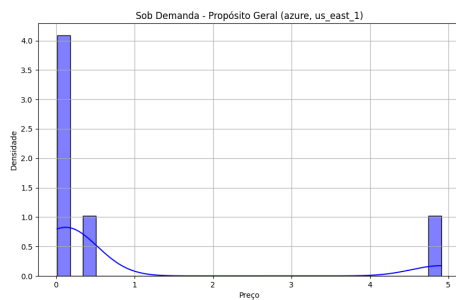


(a) Instâncias Sob Demanda

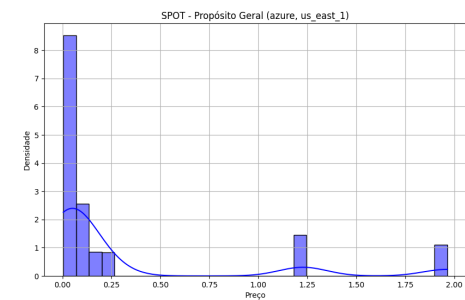


(b) Instâncias Spot

Azure

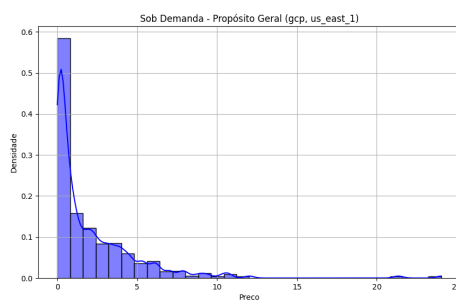


(c) Instâncias Sob Demanda

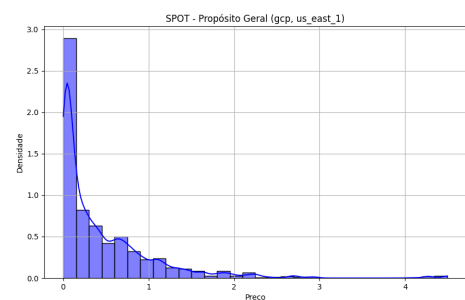


(d) Instâncias Spot

GCP



(e) Instâncias Sob Demanda



(f) Instâncias Spot

Figura 5 – Histogramas e KDEs de preços sob demanda e spot para os provedores AWS, Azure e GCP na região *us_east_1*, considerando máquinas de Propósito Geral no período de maio de 2023 a maio de 2024.

Fonte: Elaborada pelo autor (2025)

4.3 COMPARAÇÃO DOS PREÇOS SOB DEMANDA E SPOT ENTRE OS PROVEDORES

Os mapas de calor apresentados a seguir fornecem uma análise consolidada dos preços sob demanda e spot por região e tipo de instância nos provedores AWS, Azure e GCP. Esses mapas foram elaborados para capturar padrões de precificação ao longo do período analisado, considerando diferentes regiões e agrupamentos de instâncias. A visualização permite identificar variações de preços que podem estar relacionadas a estratégias específicas de alocação de recursos por parte dos provedores ou demandas sazonais em determinadas regiões.

Inicialmente, foram realizadas análises individualizadas de cada provedor, explorando os dados para identificar características particulares de suas estratégias de precificação. Em seguida, foi realizada uma análise comparativa entre os provedores, com o objetivo de compreender o potencial de redução de custos ao adotar uma abordagem multi-nuvem. Esta abordagem busca considerar os fatores geográficos e tipos de instância, permitindo inferir como a combinação de serviços de diferentes provedores pode otimizar a relação custo-benefício em cenários específicos.

Os mapas de calor foram adotados por serem representações visuais da mediana dos preços, que foram escolhidas devido à sua robustez frente a variações extremas, como discutido anteriormente. Eles oferecem uma perspectiva detalhada, possibilitando uma análise mais clara das diferenças regionais e por tipo de instância, além de fornecer subsídios para estratégias de migração ou alocação multi-nuvem.

4.3.1 Análise dos preços sob demanda e spot: AWS

Os mapas de calor apresentados para o provedor AWS oferecem uma visão detalhada da mediana dos preços sob demanda e spot por região e tipo de instância ao longo do período analisado. Cada gráfico reflete padrões distintos de precificação, com variações significativas entre as regiões e famílias de máquinas. Essa análise destaca como os preços são influenciados por fatores específicos, como a demanda local, o tipo de instância e a estratégia de precificação do provedor.

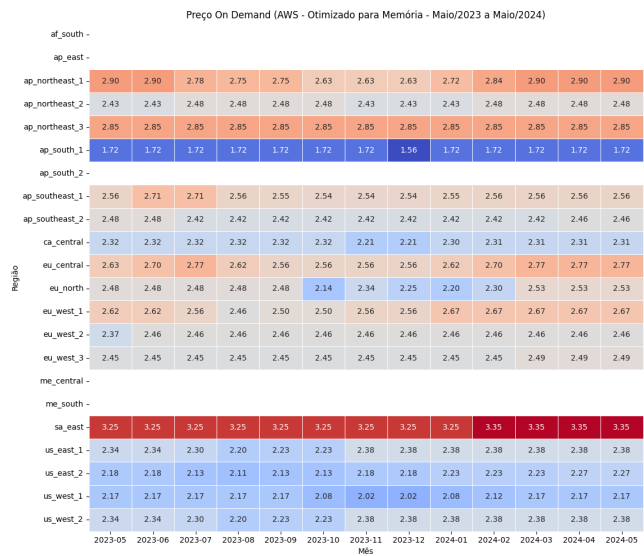
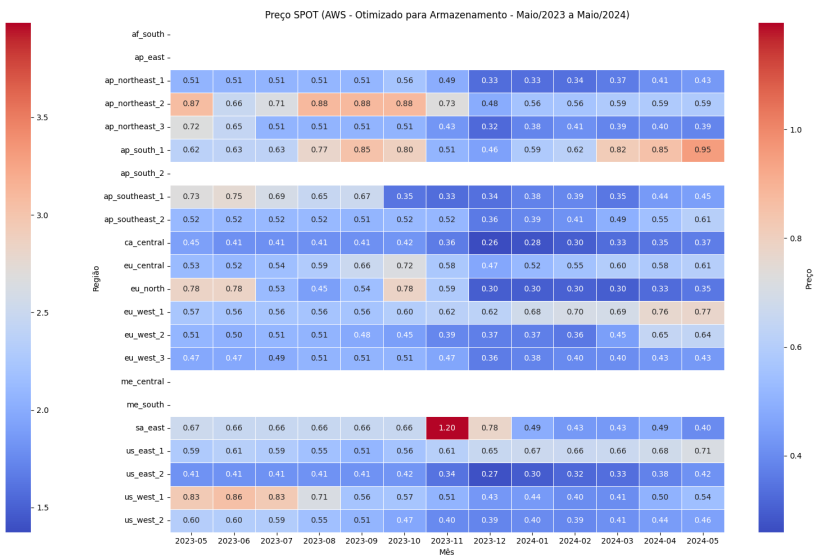
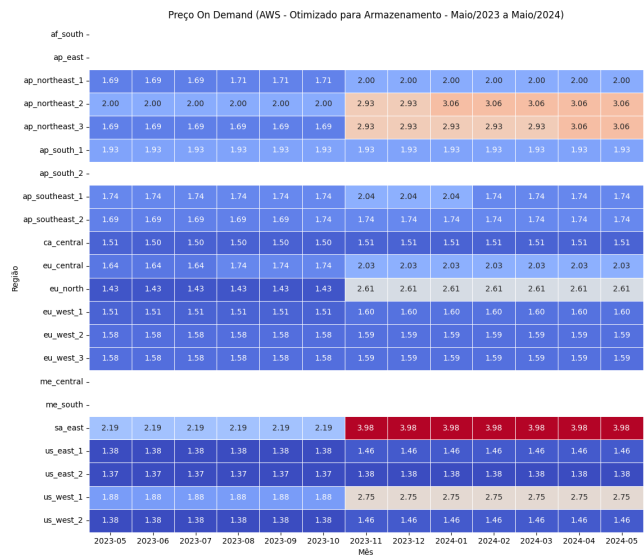
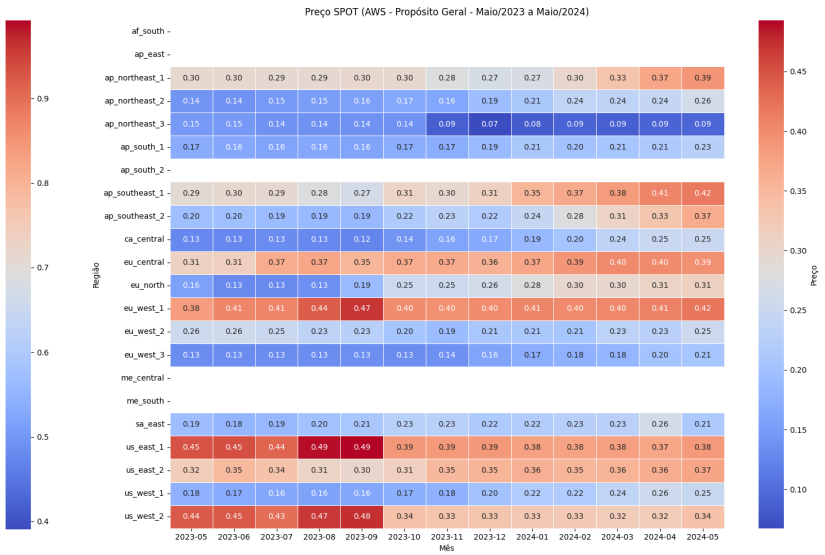
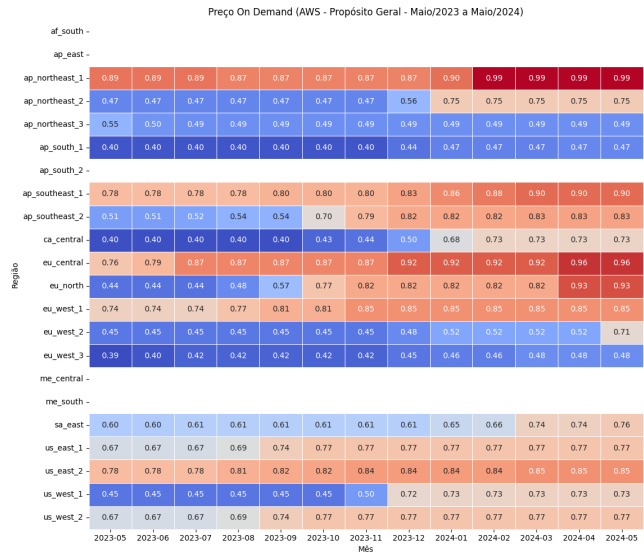
Observa-se que algumas regiões, como **us-east-1** e **eu-west-1**, apresentam uma maior estabilidade de preços em comparação com outras, como **ap-south-1** e **sa-east-1**, que exibem maior volatilidade nos preços spot. Essa diferença pode estar relacionada a variações de demanda local ou limitações de capacidade nas regiões menos estáveis. Além disso, percebe-se

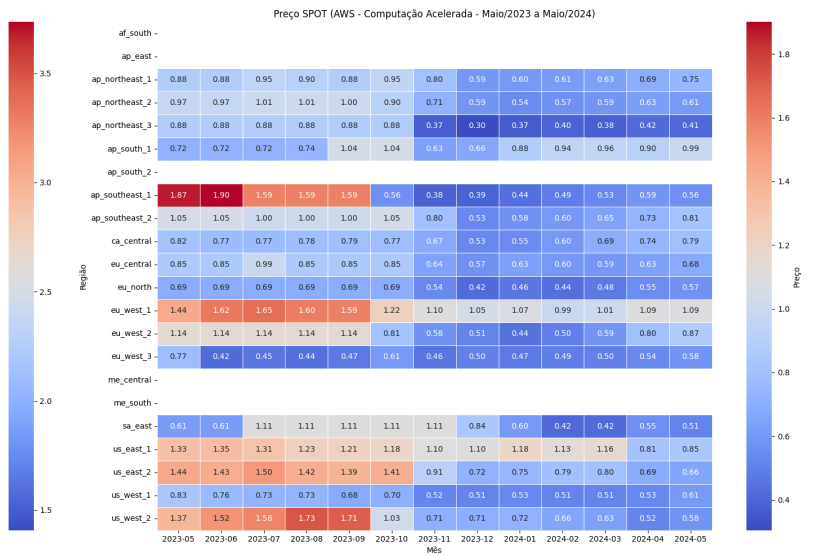
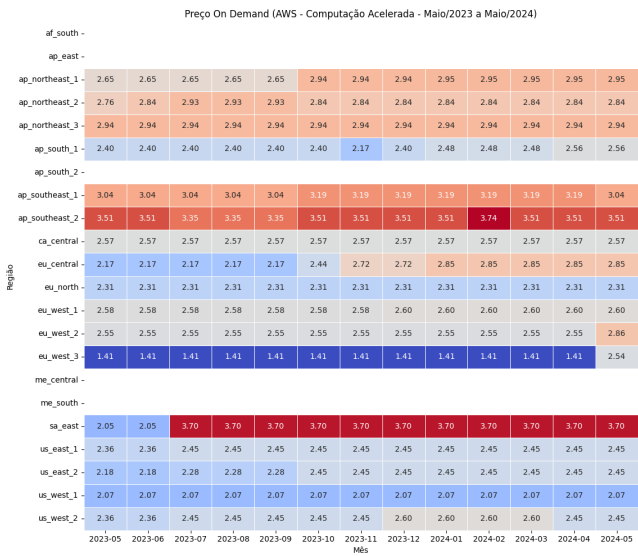
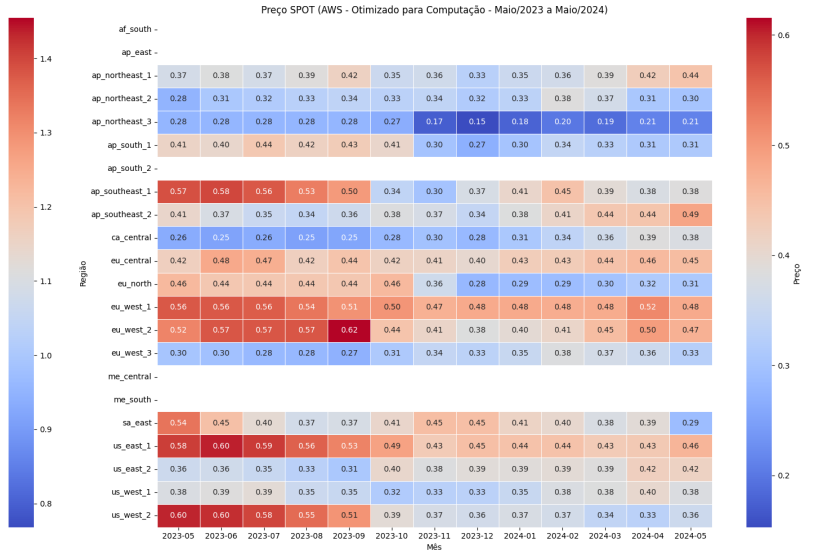
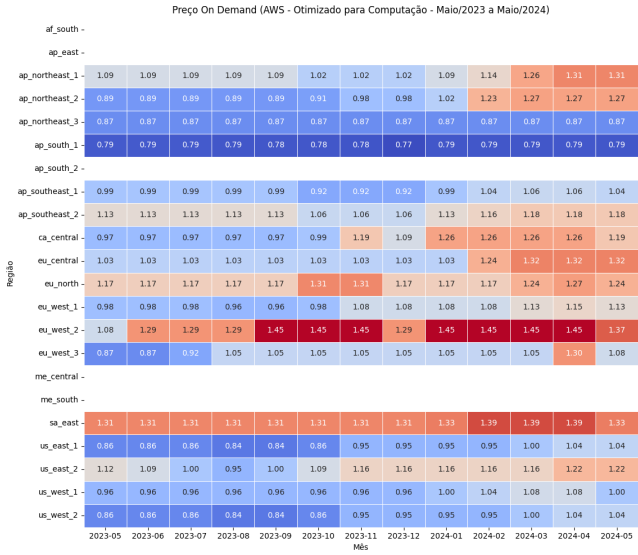
que os preços sob demanda, de maneira geral, apresentam maior uniformidade em comparação aos preços spot, que são mais suscetíveis a flutuações abruptas.

Uma característica notável nos mapas de calor é a ausência de registros de preços para determinadas regiões e famílias de máquinas. Tal fato pode ser explicado pela indisponibilidade desses dados no *dataset* SpotLake durante o período analisado.

Adicionalmente, observa-se que as escalas dos gráficos foram ajustadas para cada família de máquina, de modo a preservar a legibilidade das variações de preço. Essa decisão permitiu destacar as diferenças intrínsecas entre os preços das instâncias sob demanda e spot, mas requer atenção ao comparar diretamente gráficos de diferentes famílias ou regiões.

No que diz respeito às famílias de máquinas, as instâncias otimizadas para computação (*Compute Optimized*) e memória (*Memory Optimized*) apresentaram maiores variações de preço spot em relação às instâncias de propósito geral (*General Purpose*). Isso indica que cargas de trabalho específicas podem estar mais sujeitas a estratégias dinâmicas de precificação da AWS. Por outro lado, as instâncias de propósito geral tendem a manter uma precificação mais previsível e alinhada entre as regiões, tornando-se uma escolha potencialmente mais vantajosa para cenários de custo previsível.





4.3.2 Análise dos preços sob demanda e spot: Azure

Os mapas de calor da Azure mostram as medianas dos preços sob demanda e spot em diversas regiões e tipos de instância ao longo do período analisado. Essa abordagem permite identificar flutuações e padrões específicos de precificação, considerando as particularidades regionais e das famílias de máquinas.

Nota-se que algumas regiões não possuem registros de preços no *dataset* SpotLake, refletindo diferenças de disponibilidade regional ou lacunas nos dados coletados. Em algumas famílias de máquinas, a ausência de dados também limita a formação de conclusões, o que é esperado dada a diversidade de opções disponíveis.

A precificação sob demanda apresenta padrões relativamente estáveis em regiões amplamente utilizadas, como **eastus** e **westeurope**, mas há maior variação em regiões menos concorridas, como **brazilsouth**. Isso sugere que a demanda local e as estratégias de alocação influenciam a competitividade regional.

Para instâncias spot, os mapas destacam uma maior amplitude de variação de preços, refletindo a dinâmica desse modelo de precificação. Regiões como **eastus2** e **uksouth** mostram maior volatilidade, enquanto **canadacentral** apresenta preços mais consistentes ao longo do período.

A escala dos gráficos sob demanda e spot não é uniforme devido às diferenças nas variações de preço entre os modelos e famílias de máquina. Essa escolha busca preservar a legibilidade e permitir uma análise mais clara das tendências observadas.

Conclui-se que a Azure adota estratégias de precificação sob demanda focadas em competitividade regional, enquanto os preços spot são mais sensíveis à oferta e demanda. A análise reforça a importância de considerar regiões e tipos de instância ao planejar cargas de trabalho, especialmente quando se busca equilibrar custos e estabilidade de preços.

4.3.3 Análise dos preços sob demanda e spot: GCP

Os mapas de calor apresentados para o *Google Cloud Platform* (GCP) mostram a mediana dos preços sob demanda e spot em diferentes regiões e tipos de instância ao longo do período analisado. Essa análise detalha as estratégias de precificação específicas do GCP, destacando as características regionais e de famílias de máquinas que influenciam o comportamento de preços.

Uma observação inicial é a ausência de dados para algumas regiões e tipos de instância no *dataset* SpotLake, o que pode refletir tanto a disponibilidade limitada dessas opções no período analisado quanto lacunas nos registros coletados. Essa limitação deve ser considerada ao interpretar os resultados, pois influencia a abrangência das conclusões. Adicionalmente, como nas análises dos outros provedores, a escala dos mapas foi ajustada individualmente para preservar a legibilidade e capturar as variações específicas de cada modelo de precificação.

Os preços sob demanda no GCP apresentam padrões geralmente estáveis em regiões centrais, como **us-central1** e **europa-west1**, que são amplamente utilizadas para cargas de trabalho diversas. Por outro lado, regiões periféricas, como **southamerica-east1**, exibem maior variação nos preços, o que pode estar associado a uma menor concorrência ou estratégias específicas para atrair cargas de trabalho para essas localidades.

Em relação às instâncias spot, os mapas de calor indicam uma variabilidade maior em regiões de alta demanda, como **asia-southeast1** e **us-west1**, onde os preços refletem a dinâmica da oferta ociosa e das flutuações de demanda. Regiões com menor concorrência, como **australia-southeast1**, apresentam preços mais consistentes, demonstrando um comportamento mais previsível para instâncias spot.

As famílias de máquinas no GCP seguem padrões similares aos observados nos outros provedores, com instâncias otimizadas para memória e computação exibindo maiores flutuações nos preços spot. As instâncias de propósito geral mantêm preços mais estáveis, tornando-se uma opção atrativa para cenários que priorizam previsibilidade de custos.

Conclui-se que o GCP adota estratégias de precificação que buscam equilíbrio entre competitividade e eficiência regional, especialmente para instâncias sob demanda. As variações nos preços spot reforçam a importância de considerar regiões específicas ao planejar cargas de trabalho que busquem custos reduzidos, especialmente em estratégias de migração ou alocação multi-nuvem.

4.3.4 Análise comparativa dos mapas de calor entre os provedores, famílias de máquinas e regiões

A análise comparativa dos mapas de calor entre AWS, Azure e GCP revela diferenças significativas nas estratégias de precificação adotadas por cada provedor, oferecendo uma visão consolidada das dinâmicas de mercado e das oportunidades para otimização de custos em ambientes multi-nuvem. Comparando as medianas de preços em diversas regiões e famílias de máquinas, observam-se padrões que refletem como cada provedor ajusta seus recursos para atender a demandas específicas.

Nas regiões de maior utilização, como **us-east-1** (AWS), **eastus** (Azure) e **us-central1** (GCP), os preços sob demanda apresentaram maior consistência, indicando que a alta concorrência regional contribui para uma precificação mais estável e previsível. Em contrapartida, regiões como **sa-east-1** e **brazilsouth** exibiram maior volatilidade, sugerindo que a menor concorrência e a variabilidade na demanda local afetam os preços de forma mais acentuada.

Os preços spot, como esperado, mostraram maior variabilidade entre os provedores. Na AWS, regiões como **ap-south-1** registraram picos momentâneos de preços, indicando maior sensibilidade a flutuações de demanda. Por outro lado, o GCP apresentou preços spot mais consistentes em regiões como **europa-west1**, refletindo uma abordagem conservadora em sua gestão de capacidade ociosa. A Azure, por sua vez, destacou-se pela volatilidade em regiões como **eastus2**, onde a alta demanda parece exercer maior influência nos preços dinâmicos.

Entre as famílias de máquinas, as instâncias otimizadas para memória (*Memory Optimized*) na Azure exibiram preços spot proporcionalmente mais elevados, reforçando a alta demanda por esse tipo de instância em cargas intensivas, como bancos de dados em memória. Em contrapartida, as instâncias de propósito geral (*General Purpose*) mostraram-se mais competitivas no GCP, particularmente em regiões como **us-central1**, evidenciando esforços para atrair cargas de trabalho genéricas.

A análise também evidenciou que, em regiões como **us-west-1** e **europa-west-3**, os preços spot na AWS foram mais previsíveis, enquanto **ap-southeast-1** apresentou maior variabilidade. Essa diferença sugere que a localização geográfica influencia diretamente as estratégias de precificação, com regiões de maior capacidade ociosa oferecendo oportunidades mais consistentes para economias.

De modo geral, a AWS e a Azure tendem a adotar estratégias regionais específicas para instâncias otimizadas, enquanto o GCP apresenta uma abordagem mais homogênea na pre-

cificação spot. Essa distinção reforça a importância de uma abordagem multi-nuvem, que pode explorar as características de cada provedor para otimizar custos. Por exemplo, cargas de trabalho que exigem estabilidade de preços podem se beneficiar de instâncias de propósito geral no GCP, enquanto demandas específicas por recursos otimizados podem encontrar maior vantagem na AWS ou Azure, dependendo da região.

4.4 VALIDAÇÃO ESTATÍSTICA E ANÁLISE DE CORRELAÇÃO

Para complementar as análises descritivas e visuais realizadas até o momento, esta seção introduz uma validação estatística focada na análise de correlação entre preços de instâncias spot e sob demanda. O objetivo principal é identificar relações estatísticas entre os preços oferecidos pelos provedores AWS, Azure e GCP em uma mesma região, bem como investigar como os preços variam entre diferentes regiões e tipos de instância para um mesmo provedor.

O coeficiente de correlação de Pearson foi escolhido como a métrica principal para essa análise, uma vez que ele mede a força e a direção de uma relação linear entre duas variáveis.

A correlação de Pearson é uma medida estatística que avalia a força e a direção do relacionamento linear entre duas variáveis. Seu valor varia entre -1 e 1, sendo calculado pela seguinte fórmula:

$$r = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2 \sum_{i=1}^n (y_i - \bar{y})^2}} \quad (4.1)$$

Onde:

- r é o coeficiente de correlação de Pearson;
- x_i e y_i representam os valores das variáveis X e Y ;
- $\bar{x}$ e $\bar{y}$ são as médias das variáveis X e Y ;
- n é o número total de observações.

Os valores de r podem ser interpretados da seguinte forma:

- $r = 1$: Correlação perfeitamente positiva, indicando que, à medida que uma variável aumenta, a outra também aumenta de forma proporcional.
- $r = -1$: Correlação perfeitamente negativa, indicando que, à medida que uma variável aumenta, a outra diminui de forma proporcional.

- $r = 0$: Nenhuma correlação linear, indicando que não há relação linear entre as variáveis.
- $0 < r < 1$: Correlação positiva parcial, onde valores mais próximos de 1 indicam um relacionamento linear forte.
- $-1 < r < 0$: Correlação negativa parcial, onde valores mais próximos de -1 indicam um relacionamento linear forte, mas inverso.

Os resultados apresentados a seguir fornecem *insights* adicionais sobre a interdependência entre os preços, auxiliando na interpretação das dinâmicas de mercado e estratégias de precificação.

4.4.1 Análise das Correlações entre Provedores

Os cálculos de correlação de Pearson entre os preços spot e sob demanda dos três provedores (AWS, Azure, GCP) revelaram padrões interessantes sobre a relação entre os preços nas diferentes regiões e famílias de máquinas.

Por exemplo, na região *ca_central*, observou-se uma forte correlação positiva entre Azure e GCP, com valores de 0.9593 para preços spot e 0.9797 para preços sob demanda na família *Propósito Geral*. Esses valores indicam uma forte relação linear positiva entre os preços dos dois provedores. No caso dos preços spot, a correlação de 0.9593 sugere que, quando os preços de um provedor aumentam, os preços do outro também tendem a aumentar de forma semelhante. Para os preços sob demanda, o valor ainda mais alto de 0.9797 sugere que os preços dos dois provedores seguem praticamente o mesmo padrão de comportamento nessa região.

Em contrapartida, foram identificadas correlações negativas em algumas combinações de provedores, indicando diferenças mais significativas nas estratégias de precificação. Um exemplo claro é a região *us_west_2* para a família *Otimizado para Memória*, onde a correlação entre AWS e GCP foi de -0.7205 para preços spot e -0.7573 para preços sob demanda. Esses valores refletem uma relação linear negativa, o que significa que aumentos nos preços de um provedor estão associados a quedas nos preços do outro. Essa discrepância pode apontar para estratégias distintas de precificação, como diferenciação para atrair diferentes perfis de clientes ou gerenciamento de capacidade ociosa.

Além disso, a família *Computação Acelerada* apresentou uma alta correlação entre AWS e Azure na região *us_west_1*, com um valor de 0.9449 para preços sob demanda. Esse valor

indica uma relação linear forte e positiva, sugerindo que os preços sob demanda de ambas as plataformas variam de forma muito semelhante nessa região. Esse comportamento pode ser atribuído a uma maior previsibilidade nos preços dessa categoria, possivelmente devido a uma demanda estável ou estratégias de alinhamento competitivo entre os provedores.

Outro exemplo relevante está na região *eu_west_3*, onde a correlação de 0.8735 para os preços spot entre AWS e Azure na família *Computação Acelerada* aponta para uma forte relação linear positiva. Esse alinhamento sugere que ambos os provedores respondem de maneira semelhante a flutuações de mercado nessa região e família de instâncias. No entanto, não foi possível calcular uma correlação válida para os preços sob demanda entre esses provedores, devido à ausência de variação nos dados.

Os resultados apresentados na Tabela 4 destacam as correlações mais significativas identificadas entre os preços spot e sob demanda nas regiões selecionadas. A análise completa, incluindo todas as famílias de instâncias e regiões, está detalhada no Apêndice B.

Tabela 4 – Correlação de Pearson entre provedores nas regiões destacadas

Família	Região	Par de Provedores	Correlação Spot	Correlação Sob Demanda
Propósito Geral	ca_central	azure-gcp	0.9593	0.9797
Otimizado para Memória	us_west_2	aws-gcp	-0.7205	-0.7573
Propósito Geral	us_west_2	azure-gcp	0.7195	0.9662
Computação Acelerada	us_west_1	aws-azure	0.2532	0.9449
Computação Acelerada	eu_west_3	aws-azure	0.8735	-
Otimizado para Memória	ap_southeast_1	aws-gcp	0.5969	0.6690
Otimizado para Computação	us_east_1	aws-gcp	-0.9650	0.8882
Propósito Geral	ap_south_1	azure-gcp	0.7996	0.9741

Fonte: Elaborada pelo autor (2025). Dados completos no Apêndice B.

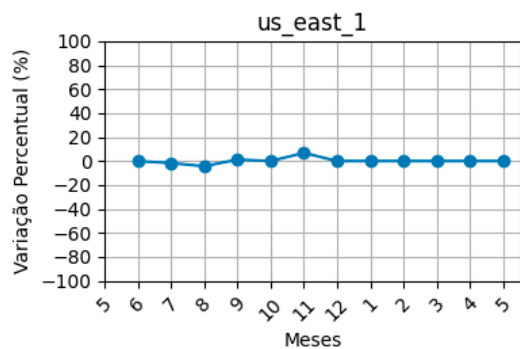
4.5 AVALIAÇÃO DE PADRÕES TEMPORAIS E SAZONAIS

A análise de padrões temporais e sazonais nos preços das instâncias spot e sob demanda revela dinâmicas importantes sobre o comportamento de preços em diferentes regiões e tipos de instância. Os gráficos selecionados a seguir foram escolhidos para representar as variações mais significativas observadas entre os provedores AWS, Azure e GCP, destacando regiões e famílias de máquinas relevantes para a análise.

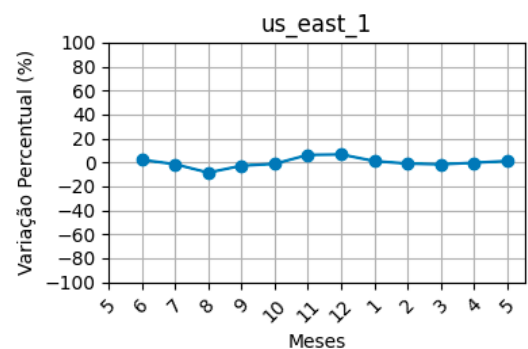
A análise dos padrões temporais também permite correlacionar as flutuações de preços com eventos sazonais globais ou regionais. Por exemplo, o período entre novembro e janeiro, marcado por eventos como *Black Friday* e Natal, frequentemente resulta em aumento de demanda por infraestrutura de nuvem em regiões como *us_east* e *eu_west*, especialmente para

instâncias sob demanda. Em contrapartida, meses como fevereiro e março tendem a apresentar uma redução sazonal na demanda em algumas regiões devido à menor atividade empresarial após as festividades de final de ano. Adicionalmente, entre junho e agosto, observam-se padrões sazonais influenciados pelas férias de verão no hemisfério norte, afetando tanto os preços sob demanda quanto spot de maneira diferenciada, dependendo da região e do provedor.

Os gráficos da **AWS** nas figuras 6 e 7 evidenciam a estabilidade de preços sob demanda na região *sa_east* para instâncias otimizadas para memória (*Memory Optimized*). Em contrapartida, demonstra maior volatilidade dos preços spot nesta região, se comparada à região *us_east_1*. Essa comparação reflete como as políticas regionais e a demanda influenciam diretamente o comportamento de preços em diferentes modelos de precificação e famílias de máquina.

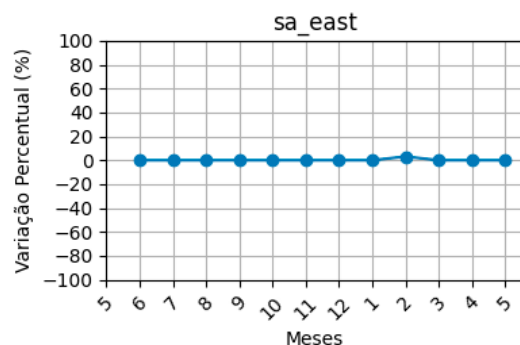


(a) Variação de preços **sob demanda** na AWS para a região *us_east_1*, instâncias Otimizadas para Memória.

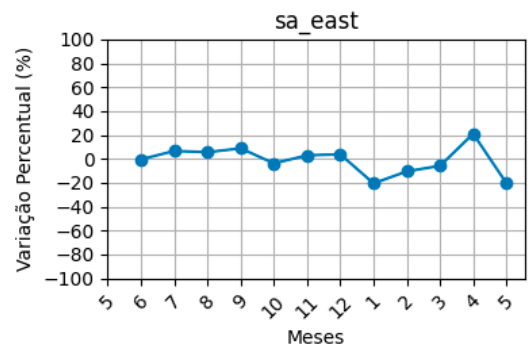


(b) Variação de preços **spot** na AWS para a região *us_east_1*, instâncias Otimizadas para Memória.

Figura 6 – Variação de preços sob demanda e spot na AWS para a região *us_east_1*, instâncias Otimizadas para Memória – Período: 05/2023 a 05/2024.



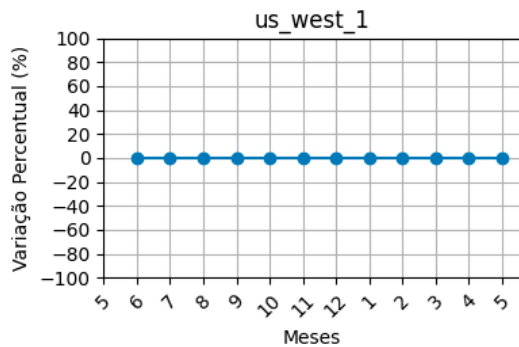
(a) Variação de preços **sob demanda** na AWS para a região *sa_east*, instâncias Otimizadas para Memória.



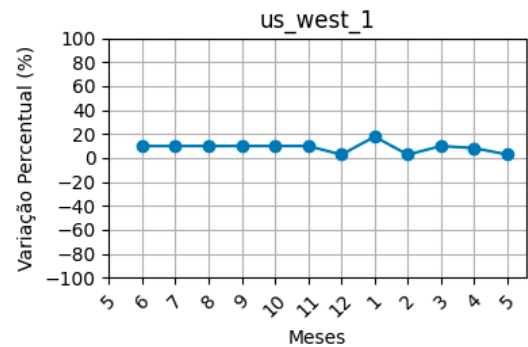
(b) Variação de preços **spot** na AWS para a região *sa_east*, instâncias Otimizadas para Memória.

Figura 7 – Variação de preços sob demanda e spot na AWS para a região *sa_east*, instâncias Otimizadas para Memória – Período: 05/2023 a 05/2024.

Para a **Azure**, os gráficos das figuras 8 e 9 mostram padrões relativamente estáveis nos preços sob demanda na região *us_west_1* para instâncias de propósito geral (*General Purpose*), enquanto a região *ca_central* apresenta variações mais acentuadas em ambos os modelos de precificação para a mesma família de máquinas. Essas discrepâncias indicam que regiões com menor concorrência tendem a apresentar maior volatilidade, enquanto regiões amplamente utilizadas exibem maior previsibilidade.

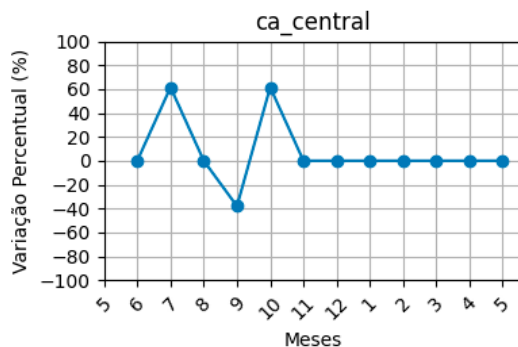


(a) Variação de preços **sob demanda** na Azure para a região *us_west_1*, instâncias de Propósito Geral.

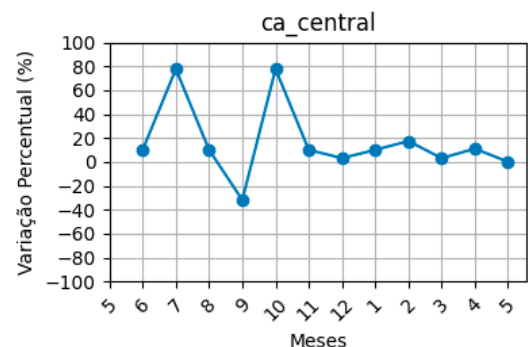


(b) Variação de preços **spot** na Azure para a região *us_west_1*, instâncias de Propósito Geral.

Figura 8 – Variação de preços sob demanda e spot na Azure para a região *us_west_1*, instâncias de Propósito Geral – Período: 05/2023 a 05/2024.



(a) Variação de preços **Sob Demanda** na Azure para a região *ca_central*, instâncias de Propósito Geral.

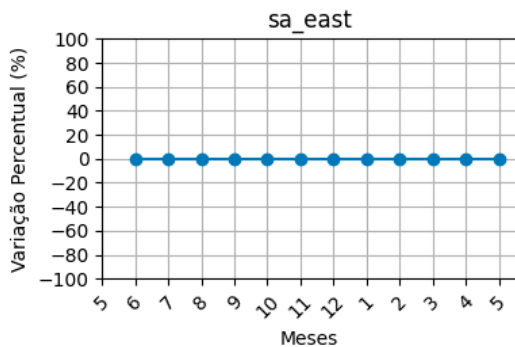


(b) Variação de preços **spot** na Azure para a região *ca_central*, instâncias de Propósito Geral.

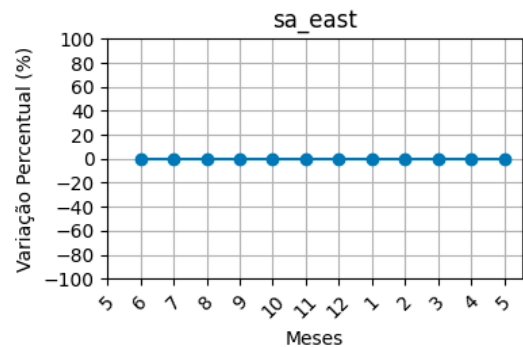
Figura 9 – Variação de preços sob demanda e spot na Azure para a região *ca_central*, instâncias de Propósito Geral – Período: 05/2023 a 05/2024.

No caso do **GCP**, a região *sa_east* para instâncias otimizadas para computação (*Compute Optimized*) destaca-se pela estabilidade de preços em ambos os modelos de precificação, conforme demonstrado na figura 10, enquanto a região *us_west_2* apresenta uma volatilidade moderada nos preços para a mesma família de máquina, conforme a figura 11. Esses padrões

sugerem que o GCP adota estratégias mais homogêneas de precificação, embora ainda existam diferenças notáveis em regiões específicas.

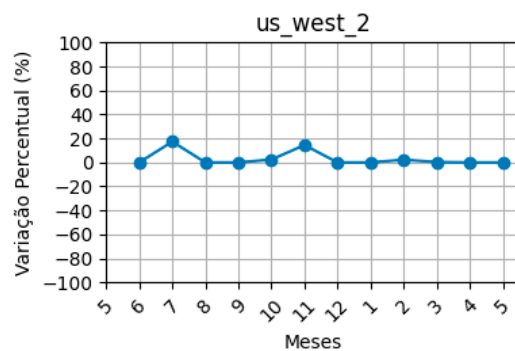


(a) Variação de preços **sob demanda** na GCP para a região *sa_east*, instâncias Otimizadas para Computação.

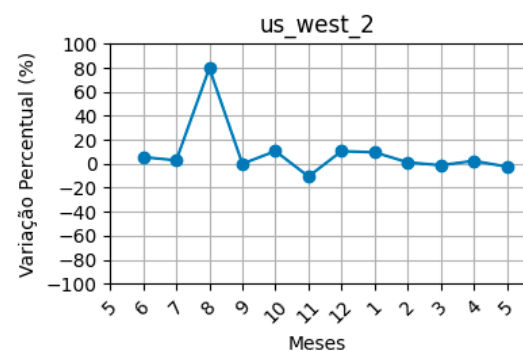


(b) Variação de preços **spot** na GCP para a região *sa_east*, instâncias Otimizadas para Computação.

Figura 10 – Variação de preços sob demanda e spot na GCP para a região *sa_east*, instâncias Otimizadas para Computação – Período: 05/2023 a 05/2024.



(a) Variação de preços **sob demanda** na GCP para a região *us_west_2*, instâncias Otimizadas para Computação.



(b) Variação de preços **spot** na GCP para a região *us_west_2*, instâncias Otimizadas para Computação.

Figura 11 – Variação de preços sob demanda e spot na GCP para a região *us_west_2*, instâncias Otimizadas para Computação – Período: 05/2023 a 05/2024.

Ao analisar as respostas às variações sazonais, é evidente que os provedores adotam estratégias distintas. A AWS, por exemplo, apresenta uma maior amplitude de variação nos preços spot em regiões como *ap-south-1*, possivelmente buscando maximizar a utilização de recursos ociosos. Já a Azure exibe oscilações menos frequentes em regiões como *eastus*, sugerindo uma estratégia de precificação mais estável para atender a cargas de trabalho críticas. O GCP, por sua vez, mantém uma abordagem ainda mais conservadora, com menor volatilidade nos preços spot em regiões amplamente utilizadas, como *us-central1*. Essas diferenças refletem as prioridades de cada provedor, sendo importantes ao considerar estratégias de migração ou planejamento multi-nuvem.

Os gráficos restantes estão disponíveis para consulta no repositório público mencionado em 4.1, permitindo uma análise detalhada e abrangente para todos os provedores, regiões e famílias de máquinas.

Os padrões temporais observados nos gráficos oferecem valiosos *insights* para estratégias de planejamento em nuvem. Identificar sazonalidades permite às empresas ajustar a alocação de recursos em períodos de alta demanda, escalonando instâncias em regiões mais estáveis e de menor custo. Além disso, em períodos de baixa demanda, o uso de instâncias spot em regiões com preços mais consistentes, como *us-west-1* e *europa-west-3*, pode gerar economias significativas. Essas estratégias são particularmente relevantes para organizações que operam globalmente e desejam explorar abordagens multi-nuvem para otimizar a relação custo-benefício de suas operações.

As análises apresentadas nesta seção oferecem implicações práticas para a tomada de decisão em ambientes multi-nuvem. A análise detalhada das variações de preços sob demanda e spot em diferentes regiões e tipos de instâncias permite que organizações tomem decisões mais informadas ao planejar a alocação de suas cargas de trabalho. Por exemplo, para cargas de trabalho críticas que exigem estabilidade e previsibilidade de custos, as instâncias sob demanda em regiões com preços historicamente consistentes, como **us-east-1** na AWS ou **eastus** na Azure, podem ser uma escolha estratégica.

Por outro lado, as empresas que buscam reduzir custos operacionais podem explorar regiões e provedores que apresentem padrões de preços spot mais previsíveis, como as instâncias otimizadas para computação no GCP na região *europa-west1*. Essa abordagem é particularmente interessante para cargas de trabalho tolerantes a interrupções, como processamento em lote ou análises não críticas, que podem ser ajustadas para aproveitar momentos de preços reduzidos.

Além disso, os resultados sugerem que uma estratégia multi-nuvem, que combina o uso de instâncias sob demanda e spot de diferentes provedores, pode otimizar a relação custo-benefício. Por exemplo, uma organização pode optar por instâncias sob demanda da Azure para cargas de trabalho estáveis em regiões amplamente utilizadas e, ao mesmo tempo, aproveitar instâncias spot da AWS em regiões com menor volatilidade de preços para cargas dinâmicas.

Por fim, a análise evidencia que a escolha da região e do tipo de instância deve ser guiada não apenas pelo preço médio, mas também pela variabilidade e previsibilidade dos custos ao longo do tempo. Ao considerar esses fatores em conjunto, as organizações podem criar estratégias de alocação mais resilientes e econômicas, maximizando os benefícios dos ambientes

multi-nuvem.

4.6 DISCUSSÃO E APLICABILIDADE PRÁTICA

A análise detalhada dos preços de instâncias spot e sob demanda revelou importantes padrões de variação entre provedores, regiões e tipos de instância. Nesta seção, discutimos como esses resultados podem ser aplicados na prática, explorando as regiões mais vantajosas para cada provedor, estratégias para alocação eficiente de recursos em ambientes multi-nuvem e a relevância dessas abordagens para empresas que buscam otimizar custos e aumentar a resiliência de suas operações. Além disso, apresentamos casos de uso hipotéticos que exemplificam a aplicabilidade dos *insights* obtidos, conectando-os a possíveis cenários do mercado e à pesquisa acadêmica. Essa discussão visa não apenas consolidar as implicações práticas da análise realizada, mas também fornecer um ponto de partida para trabalhos futuros no campo de precificação dinâmica e alocação inteligente de instâncias em nuvens públicas.

4.6.1 Regiões mais vantajosas para cada provedor

Os resultados das análises realizadas ao longo deste capítulo destacam as regiões mais vantajosas para cada provedor, considerando as características de preços spot e sob demanda.

No caso da AWS, a região *us-east-1* demonstrou consistentemente os menores valores médios e maior previsibilidade nos preços sob demanda, tornando-se atrativa para cargas críticas e contínuas. Já as regiões *ap-southeast-1* e *eu-west-1* apresentaram maior competitividade nos preços spot, sendo indicadas para cargas de trabalho menos sensíveis à disponibilidade.

Para a Azure, a estabilidade de preços foi uma característica predominante em regiões como *us-west-2* e *eu-north*, que apresentaram baixos desvios padrões tanto para preços spot quanto sob demanda. Essa estabilidade reforça a adequação da Azure para cenários em que a previsibilidade de custos é essencial, como na execução de aplicações críticas e contratos de longo prazo.

No GCP, as regiões *us-central1* e *asia-southeast1* se destacaram pela ampla variação de preços spot, indicando oportunidades de economia para empresas dispostas a explorar modelos de escalabilidade dinâmica. No entanto, essa volatilidade também exige monitoramento constante e integração com políticas de alocação automatizadas para maximizar os benefícios de custo.

4.6.2 Aplicabilidade prática: Estratégias para alocação multi-nuvem

Os resultados obtidos oferecem importantes direcionamentos para empresas que operam em ambientes multi-nuvem. A seguir, destacamos algumas estratégias práticas baseadas nos *insights* apresentados:

- **Explorar a complementaridade de provedores:** A análise revelou que cada provedor possui características distintas em suas políticas de precificação. Por exemplo, a AWS oferece flexibilidade em regiões de alta demanda (*us-east-1*), enquanto o GCP se destaca em cenários que demandam elasticidade e preços dinâmicos (*us-central1*). Uma abordagem que combine esses pontos fortes pode oferecer redução de custos e maior disponibilidade.
- **Aproveitar as vantagens regionais:** Empresas que planejam cargas de trabalho intensivas em regiões específicas podem optar por provedores que apresentam menor custo e maior estabilidade. Por exemplo, a Azure é uma escolha confiável em regiões como *eu-west-2*, enquanto a AWS e o GCP oferecem melhor custo-benefício para cargas de trabalho temporárias em regiões como *ap-southeast-1*.
- **Monitorar sazonalidades e eventos de mercado:** Os dados mostraram que períodos sazonais e eventos específicos, como *Black Friday*, afetam significativamente a dinâmica de preços. Isso sugere que empresas podem planejar cargas de trabalho menos urgentes para períodos de menor demanda, otimizando seus custos.
- **Adotar ferramentas automatizadas de alocação:** Dada a volatilidade dos preços spot, principalmente em provedores como o GCP, o uso de plataformas de gerenciamento automatizado pode ser crucial. Ferramentas que integram algoritmos de decisão para alocação de cargas de trabalho com base em custo e desempenho são altamente recomendadas.

4.6.3 Casos de Uso: Aplicação Prática dos Resultados

Para demonstrar a aplicabilidade prática dos resultados obtidos na análise de preços, apresentamos dois casos de uso hipotéticos que exploram cenários realistas de alocação de instâncias em nuvens públicas. Esses exemplos destacam os custos estimados para diferentes

estratégias de alocação, utilizando os dados analisados nesta dissertação. As estratégias são apresentadas em ordem decrescente de custo, sendo: (i) todas as instâncias alocadas sob demanda em um único provedor, (ii) todas as instâncias spot alocadas em um único provedor, e (iii) todas as instâncias spot distribuídas em um ambiente multi-nuvem. Essa hierarquia evidencia as oportunidades de economia proporcionadas por tais abordagens.

4.6.3.1 Cenário 1: Processamento de Dados em Batch

Imagine uma empresa que precisa executar uma carga de trabalho intensiva de processamento de dados em *batch*, consumindo 100 instâncias de propósito geral por um período de 48 horas. Três abordagens são consideradas:

- **Opção 1: Sob Demanda em um Único Provedor (AWS)** A empresa utiliza 100 instâncias sob demanda na região **us_east_1** da AWS, onde o custo médio de uma instância sob demanda é \$1,9268/hora. O custo total seria:

$$Custo = 100 \times 1,9268 \times 48 = \mathbf{\$9.249,60}$$

- **Opção 2: Spot em um Único Provedor (AWS)** Ao optar por instâncias spot na mesma região, o custo médio seria \$0,986/hora. Assim, o custo total seria:

$$Custo = 100 \times 0,986 \times 48 = \mathbf{\$4.732,80}$$

- **Opção 3: Multi-Nuvem (AWS, Azure, GCP)** Distribuindo as instâncias entre os provedores mais baratos:

- AWS: 40 instâncias na **us_east_1** (\$0,986/hora);
- Azure: 30 instâncias na **us_west_2** (\$0,2088/hora);
- GCP: 30 instâncias na **eu_west_3** (\$0,5111/hora).

O custo total seria:

$$Custo = (40 \times 0,986 \times 48) + (30 \times 0,2088 \times 48) + (30 \times 0,5111 \times 48)$$

$$Custo = \$1.894,08 + \$300,67 + \$735,97 = \mathbf{\$2.930,72}$$

Conclusão: A abordagem multi-nuvem oferece uma economia de aproximadamente **68%** em relação à abordagem sob demanda e de **38%** em relação à abordagem spot de um único provedor.

4.6.3.2 Cenário 2: Aplicação Web Escalável

Uma *startup* precisa manter 50 instâncias otimizadas para memória durante um mês (30 dias) para hospedar uma aplicação web escalável. As opções analisadas são:

- **Opção 1: Sob Demanda em um Único Provedor (GCP)** Utilizando 50 instâncias sob demanda na região **us_east_1**, onde o custo médio é \$13,7756/hora, o custo total seria:

$$Custo = 50 \times 13,7756 \times 24 \times 30 = \mathbf{\$495.940,80}$$

- **Opção 2: Spot em um Único Provedor (GCP)** Utilizando 50 instâncias spot na mesma região, com custo médio de \$3,4309/hora, o custo total seria:

$$Custo = 50 \times 3,4309 \times 24 \times 30 = \mathbf{\$123.509,04}$$

- **Opção 3: Multi-Nuvem (AWS, Azure, GCP)** Distribuindo as instâncias entre os provedores mais econômicos:
 - AWS: 20 instâncias na **eu_west_3** (\$1,0444/hora);
 - Azure: 15 instâncias na **us_west_1** (\$2,5325/hora);
 - GCP: 15 instâncias na **eu_west_3** (\$3,4089/hora).

O custo total seria:

$$Custo = (20 \times 1,0444 \times 24 \times 30) + (15 \times 2,5325 \times 24 \times 30) + (15 \times 3,4089 \times 24 \times 30)$$

$$Custo = \$15.600,96 + \$27.369,00 + \$36.781,92 = \mathbf{\$79.751,88}$$

Conclusão: A abordagem multi-nuvem oferece uma economia de aproximadamente **84%** em relação à abordagem sob demanda e de **35%** em relação à abordagem spot de um único provedor.

Implicações Práticas: Esses casos de uso destacam o potencial de economia ao adotar uma estratégia multi-nuvem, utilizando dados históricos e análises de preço como suporte para a alocação eficiente de recursos. Além disso, essas observações servem como base para futuras pesquisas sobre mecanismos dinâmicos de alocação de instâncias em ambientes multi-nuvem.

4.6.4 Relevância para a pesquisa acadêmica e o mercado

Os resultados obtidos neste estudo possuem várias implicações práticas e acadêmicas. Para o mercado, as conclusões auxiliam na tomada de decisão para a escolha de provedores e regiões em estratégias multi-nuvem. As diferenças identificadas entre os preços spot e sob demanda reforçam a importância de análises detalhadas para a otimização de custos, especialmente em empresas que operam com margens reduzidas e grande volume de operações.

No contexto acadêmico, os resultados expandem a compreensão sobre as dinâmicas de precificação dos provedores de nuvem em cenários multi-nuvem. Estudos futuros podem explorar modelos preditivos para prever preços spot, bem como a predição da revogação dessas instâncias, possibilitando uma migração entre regiões e/ou provedores considerando a relação custo-benefício.

De forma prática, a abordagem apresentada neste capítulo oferece uma base sólida para a implementação de políticas de otimização de custos em empresas que buscam explorar os benefícios da nuvem pública. A escolha criteriosa das regiões e provedores, aliada ao uso de ferramentas automatizadas, permite alinhar os objetivos de custo e desempenho, contribuindo para estratégias mais eficazes e sustentáveis em ambientes multi-nuvem.

4.6.5 Diferenciais da Pesquisa

Esta pesquisa apresenta diversos diferenciais em relação aos trabalhos anteriores na área, oferecendo contribuições relevantes tanto para a academia quanto para o mercado. Primeiramente, a análise considera simultaneamente os três principais provedores de serviços em nuvem — AWS, Azure e GCP —, o que permite uma visão comparativa abrangente das dinâmicas de precificação de instâncias spot e sob demanda em um contexto multi-nuvem. Essa abordagem contrasta com a maioria dos trabalhos relacionados, que frequentemente se limitam à análise de um único provedor, principalmente a AWS.

Outro aspecto é o uso de dados históricos que abrangem 12 meses, permitindo identificar tendências sazonais e padrões temporais na precificação, algo raramente explorado de forma detalhada. Além disso, a pesquisa utilizou uma metodologia de agrupamento de regiões e tipos de instâncias, garantindo a comparabilidade entre provedores e refletindo as características específicas de cada cenário geográfico.

O volume expressivo de dados analisados, totalizando mais de 14 milhões de registros e mais

de 1.000 gráficos gerados, é outro diferencial significativo, proporcionando uma base sólida para as conclusões apresentadas. Adicionalmente, o trabalho enfatiza a aplicabilidade prática dos resultados, sugerindo estratégias de alocação multi-nuvem para empresas e demonstrando cenários reais de economia com o uso de instâncias spot.

Por fim, este estudo se conecta diretamente com desafios atuais da computação em nuvem, ao propor como trabalhos futuros um mecanismo dinâmico de acompanhamento de preços e alocação de instâncias, baseado em dados históricos e análise em tempo real.

4.7 INTEGRAÇÃO DOS RESULTADOS COM OS OBJETIVOS DA PESQUISA

As análises realizadas ao longo deste capítulo permitiram explorar detalhadamente como os preços spot variam entre os provedores AWS, Azure e GCP em diferentes regiões geográficas e tipos de instância. A identificação de padrões e correlações nos preços evidenciou que as estratégias de precificação são altamente influenciadas pelo contexto regional e pelas características das instâncias.

Em regiões amplamente utilizadas, como *us-east-1*, os preços mostraram maior estabilidade, com correlações mais consistentes entre os provedores. Por outro lado, regiões menos concorridas, como *ap-northeast-1*, apresentaram maior volatilidade, refletindo dinâmicas locais de demanda e oferta. Essas variações destacam a importância de considerar tanto a localização geográfica quanto o tipo de instância ao planejar estratégias de alocação multi-nuvem.

A análise temporal revelou que a volatilidade dos preços spot influencia diretamente a correlação entre os provedores. Em períodos de alta estabilidade, como os meses iniciais do ano, as correlações foram mais evidentes, especialmente para instâncias de propósito geral. Eventos sazonais, como a *Black Friday*, também impactaram os padrões de precificação, causando flutuações significativas em algumas famílias de instâncias, como as otimizadas para memória e computação. Esses resultados indicam que os preços spot são mais dinâmicos e sensíveis às condições de mercado do que os preços sob demanda, que mantêm maior previsibilidade.

Além disso, as características das instâncias mostraram-se determinantes para a correlação de preços. Instâncias de propósito geral destacaram-se por sua estabilidade, enquanto as otimizadas para computação e memória apresentaram maior volatilidade devido à natureza específica de suas cargas de trabalho. Esses padrões reforçam a necessidade de uma análise cuidadosa ao selecionar instâncias para diferentes cenários operacionais.

Os resultados apresentados neste capítulo não apenas respondem às perguntas de pesquisa,

mas também fornecem *insights* práticos para empresas que operam em ambientes multi-nuvem. A compreensão das dinâmicas de precificação permite otimizar custos e melhorar a eficiência das estratégias de alocação, contribuindo para o avanço do estado da arte em pesquisas sobre computação em nuvem.

4.7.1 Sumarização da Análise

A Tabela 5 resume os principais aspectos metodológicos e quantitativos da pesquisa, evidenciando a abrangência e profundidade da análise realizada.

Tabela 5 – Resumo Quantitativo e Metodológico da Pesquisa

Descrição	Valor
Total de arquivos analisados (.csv.gz)	73.431
AWS	55.332
Azure	9.186
GCP	8.913
Total de registros de preços analisados	14.804.169
AWS	3.309.870
Azure	9.769.182
GCP	1.725.117
Provedores de nuvem analisados	AWS, Azure, GCP
Agrupamentos de Regiões geográficas	22
Agrupamento de Famílias de máquinas	5
Total de tipos de instâncias analisados	557
Período de análise	12 meses (maio/2023 a maio/2024)
Total de gráficos gerados	> 1.000
Total de correlações calculadas	1.386
Total de pares regiões-provedor analisadas	66
Volume total de dados armazenados (em GB)	> 50
Tempo médio para análise de um arquivo	2.3 segundos
Tecnologias utilizadas	MongoDB, Python, Parquet
Estratégia de visualização	Boxplots, heatmaps, histogramas
Instâncias utilizadas na análise computacional	AWS x2gd.xlarge (4 vCPUs, 64 GB RAM)
Horas de utilização de máquina na análise computacional	~4.015 horas

Fonte: Elaborada pelo autor (2025)

5 CONCLUSÃO E TRABALHOS FUTUROS

Este trabalho analisou a precificação de instâncias spot e sob demanda nos principais provedores de nuvem — AWS, Azure e *Google Cloud Platform* (GCP) —, considerando um cenário multi-nuvem a partir de dados históricos do *dataset* SpotLake. A pesquisa foi motivada pelo impacto financeiro associado à escolha de instâncias e pela crescente adoção de estratégias multi-nuvem em ambientes organizacionais.

As análises realizadas permitiram responder às questões de pesquisa propostas. Observou-se que as estratégias de precificação variam significativamente entre provedores e regiões, e que fatores sazonais e características específicas das instâncias exercem influência relevante nas dinâmicas de preços. Identificou-se também que, embora existam correlações em determinados contextos, elas não são uniformes, o que evidencia estratégias de precificação independentes entre os provedores.

Os resultados obtidos fornecem subsídios para a formulação de estratégias de alocação de recursos em ambientes multi-nuvem, permitindo a identificação de oportunidades de otimização de custos a partir da análise de flutuações regionais e comportamentos de precificação. A integração de dados de AWS, Azure e GCP em uma mesma análise, abrangendo um período de doze meses, constitui uma contribuição relevante para a literatura sobre gestão de custos em computação em nuvem.

5.1 PRINCIPAIS LIMITAÇÕES

Embora os resultados obtidos sejam significativos, algumas limitações devem ser reconhecidas. A primeira limitação está relacionada ao período de análise. Apesar de abrangente, o intervalo de doze meses pode não capturar completamente padrões sazonais de longo prazo ou mudanças estruturais nos modelos de precificação dos provedores. Estudos futuros com períodos de análise mais extensos poderiam fornecer uma visão mais completa dessas dinâmicas.

Outra limitação importante é a ausência de dados detalhados sobre taxas de interrupção para instâncias spot nos provedores Azure e GCP. Esses dados são cruciais para avaliar a relação custo-benefício de instâncias spot em cargas de trabalho críticas, como bancos de dados ou sistemas de processamento em tempo real. Além disso, a análise focou exclusivamente nos três maiores provedores globais de nuvem, excluindo *players* regionais ou emergentes que poderiam

oferecer alternativas competitivas em certos cenários.

Por fim, a análise foi limitada a métricas relacionadas aos preços e às correlações entre eles de acordo com o *dataset*. Métricas adicionais, como custos de transferência de dados, poderiam enriquecer significativamente as análises e oferecer uma visão mais holística das escolhas de alocação em nuvem.

5.2 TRABALHOS FUTUROS

Com base nas limitações identificadas e nas oportunidades levantadas, sugerem-se os seguintes trabalhos futuros:

- **Estudos de longo prazo:** Ampliar o período de análise para capturar sazonalidades de longo prazo e identificar tendências estruturais nos modelos de precificação dos provedores.
- **Incorporação de métricas adicionais:** Considerar dados sobre taxas de interrupção e custos de transferência de dados, permitindo uma análise mais abrangente.
- **Desenvolvimento de ferramentas dinâmicas:** O desenvolvimento de ferramentas dinâmicas baseadas em aprendizado de máquina representa uma direção promissora para trabalhos futuros. Essas ferramentas poderiam combinar dados históricos e em tempo real para prever flutuações de preços spot, permitindo decisões mais ágeis e otimizadas. Por exemplo, um algoritmo preditivo poderia recomendar a melhor combinação de instâncias em diferentes regiões e provedores, ajustando automaticamente a alocação de recursos conforme as mudanças no mercado.
- **Ampliação do escopo:** Incluir provedores regionais e emergentes, como Alibaba Cloud e Oracle Cloud, para explorar dinâmicas de precificação em mercados menos consolidados.
- **Estudos sobre impacto ambiental:** Avaliar como estratégias multi-nuvem baseadas em instâncias spot podem contribuir para a sustentabilidade, otimizando o uso de recursos computacionais ociosos e reduzindo o consumo de energia.
- **Simulações de cenários práticos:** Desenvolver estudos de caso que integrem a análise de preços com aplicações reais, simulando cenários de economia e otimização para diferentes cargas de trabalho.

5.3 CONTRIBUIÇÕES ACADÊMICAS E IMPLICAÇÕES TEÓRICAS

Além das contribuições práticas, este trabalho também oferece contribuições para a literatura na área de computação em nuvem, preenchendo lacunas relevantes e complementando estudos anteriores. A análise integrada dos três maiores provedores de nuvem pública — AWS, Azure e GCP — no contexto de precificação de instâncias spot é um diferencial desta pesquisa. Enquanto a maioria dos estudos foca em um único provedor ou em análises isoladas, este trabalho adota uma perspectiva comparativa e abrangente, explorando as dinâmicas de precificação em um contexto multi-nuvem.

Uma das lacunas preenchidas está relacionada à falta de estudos que analisam a correlação de preços entre os provedores de nuvem. Este trabalho não apenas identificou padrões de correlação — positivos, negativos ou nulos — entre os preços de instâncias spot, mas também explorou como essas correlações podem ser utilizadas para otimizar a alocação de recursos computacionais. Isso complementa estudos anteriores que abordavam a precificação de instâncias de forma independente, sem considerar a interação entre provedores.

Além disso, a pesquisa introduziu uma metodologia para análise de dados históricos de preços em larga escala, utilizando o *dataset* SpotLake. Com isso, foi possível avaliar variações temporais e sazonais, algo que é frequentemente ausente na literatura. Esse enfoque temporal complementa trabalhos que se concentram em análises pontuais de precificação, oferecendo uma perspectiva mais dinâmica e realista das flutuações do mercado.

Outro ponto de destaque é a inclusão de diferentes categorias de instâncias — como propósito geral, otimizado para computação e otimizado para memória — em uma análise comparativa entre regiões e provedores. Isso amplia a aplicabilidade dos resultados e atende a diferentes cenários de uso na prática.

Ao identificar regiões e categorias de instâncias mais vantajosas, o trabalho também contribui para uma compreensão mais aprofundada das dinâmicas de mercado em regiões geográficas específicas. Por exemplo, enquanto estudos anteriores tendem a focar nas regiões mais populares, esta pesquisa avaliou agrupamentos regionais mais amplos, incluindo regiões emergentes e menos estudadas.

Por fim, a análise integrada dos dados históricos e a proposta de estratégias para otimização em ambientes multi-nuvem complementam uma literatura que, muitas vezes, aborda a computação em nuvem limitada a um único provedor. Os resultados obtidos nesta pesquisa não apenas preenchem essas lacunas, mas também abrem novos caminhos para a investigação

futura, ampliando as fronteiras do conhecimento acadêmico sobre precificação e estratégias multi-nuvem.

5.4 CONSIDERAÇÕES FINAIS

Este trabalho representa um avanço significativo na análise de precificação de instâncias spot em um contexto multi-nuvem. Além de oferecer uma visão abrangente sobre as dinâmicas de preços entre os três maiores provedores de nuvem, ele destaca a relevância de estratégias baseadas em dados para a otimização de custos em ambientes computacionais cada vez mais complexos.

Os resultados obtidos fornecem subsídios para empresas e pesquisadores que desejam explorar o potencial das instâncias spot, apontando caminhos para economias significativas e maior flexibilidade operacional. Ao mesmo tempo, as limitações e os desafios identificados abrem espaço para uma agenda de pesquisa promissora, que pode não apenas aprofundar o conhecimento sobre precificação em nuvem, mas também contribuir para a sustentabilidade e a eficiência no uso de recursos computacionais.

Dessa forma, este trabalho não apenas contribui para o avanço do conhecimento acadêmico, mas também oferece ferramentas práticas para empresas que operam em um mercado altamente competitivo. Acredita-se que os resultados aqui apresentados poderão servir como base para estudos futuros, inspirando novas abordagens e soluções no campo da computação em nuvem.

REFERÊNCIAS

- A. W. S. Blogs. *New Amazon EC2 Spot Pricing Model*. 2018. <<https://aws.amazon.com/blogs/compute/new-amazon-ec2-spot-pricing/>>. Accessed on January 10, 2025.
- ABUALKISHIK, A.; ALWAN, A.; GULZAR, Y. Disaster recovery in cloud computing systems: An overview. *International Journal of Advanced Computer Science and Applications*, v. 11, p. 702, 10 2020.
- ALSHAREEF, H. N. Current development, challenges, and future trends in cloud computing: A survey. *International Journal of Advanced Computer Science and Applications*, The Science and Information Organization, v. 14, n. 3, 2023. Disponível em: <<http://dx.doi.org/10.14569/IJACSA.2023.0140337>>.
- AWS. *Amazon EC2 Pricing*. 2025. <<https://aws.amazon.com/pt/ec2/pricing/>>. Accessed on January 10, 2025.
- AWS Instances Type. *AWS Instances Type*. 2025. <<https://aws.amazon.com/pt/ec2/instance-types>>. Accessed on January 10, 2025.
- AWS Regions. *AWS Global Infrastructure*. 2025. <https://aws.amazon.com/pt/about-aws/global-infrastructure/regions_az/>. Accessed on January 10, 2025.
- Azure Instances Type. *Azure Instances Type*. 2025. <<https://learn.microsoft.com/pt-br/azure/virtual-machines/sizes/overview?tabs=breakdownseries%2Cgeneralisizelist%2Ccomputesizelist%2Cmemorysizelist%2Cstoragesizelist%2Cgpuisizelist%2Cfpgasizelist%2Chpcisizelist#vm-size-and-series-naming>>. Accessed on January 10, 2025.
- Azure Regions. *Azure Global Infrastructure*. 2025. <<https://azure.microsoft.com/pt-br/explore/global-infrastructure/geographies/>>. Accessed on January 10, 2025.
- BAUGHMAN, M.; CATON, S.; HAAS, C.; CHARD, R.; WOLSKI, R.; FOSTER, I.; CHARD, K. Deconstructing the 2017 changes to aws spot market pricing. In: *Proceedings of the 10th Workshop on Scientific Cloud Computing*. [S.l.]: Association for Computing Machinery, 2019. p. 19–26.
- BEN-YEHUDA, O. A.; BEN-YEHUDA, M.; SCHUSTER, A.; TSAFRIR, D. Deconstructing Amazon EC2 spot instance pricing. *ACM Transactions on Economics and Computation (TEAC)*, ACM New York, NY, USA, v. 1, n. 3, p. 1–20, 2013.
- BEN-YEHUDA, O. A.; BEN-YEHUDA, M.; SCHUSTER, A.; TSAFRIR, D. Deconstructing amazon ec2 spot instance pricing. v. 1, n. 3, 2013. ISSN 2167-8375.
- BERENBERG, A.; CALDER, B. Deployment archetypes for cloud applications. *ACM Comput. Surv.*, Association for Computing Machinery, New York, NY, USA, v. 55, n. 3, 2022. ISSN 0360-0300. Disponível em: <<https://doi.org/10.1145/3498336>>.
- CHARD, R.; CHARD, K.; WOLSKI, R.; MADDURI, R.; NG, B.; BUBENDORFER, K.; FOSTER, I. Cost-aware cloud profiling, prediction, and provisioning as a service. *IEEE Cloud Computing*, v. 4, p. 48–59, 07 2017.

- CHITTORA, V.; GUPTA, C. P. Dynamic spot price forecasting using stacked lstm networks. In: *2020 3rd International Conference on Intelligent Sustainable Systems (ICISS)*. [S.l.: s.n.], 2020. p. 1080–1085.
- EKWE-EKWE, N.; BARKER, A. Location, location, location: Exploring amazon ec2 spot instance pricing across geographical regions. In: *2018 18th IEEE/ACM International Symposium on Cluster, Cloud and Grid Computing (CCGRID)*. [S.l.: s.n.], 2018. p. 370–373.
- Flexera Software. *2024 State of the Cloud Report*. 2024. Disponível em <<https://www.flexera.com/>>. Disponível em: <<https://www.flexera.com/>>.
- GCP Instances Type. *GCP Instances Type*. 2025. <<https://cloud.google.com/compute/docs/machine-resource>>. Accessed on January 10, 2025.
- GCP Regions. *GCP Global Infrastructure*. 2025. <<https://cloud.google.com/about/locations>>. Accessed on January 10, 2025.
- Google Cloud. *GCP Spot Interruption*. 2025. <<https://cloud.google.com/compute/docs/instances/spot>>. Accessed on January 10, 2025.
- Google Cloud Platform. *Google Cloud Compute Engine VM Instance Pricing*. 2025. <<https://cloud.google.com/compute/vm-instance-pricing#pricing>>. Accessed on January 10, 2025.
- KABIR, H. M. D.; KHOSRAVI, A.; HOSEN, M. A.; NAHAVANDI, S.; BUYYA, R. Probability density for amazon spot instance price. In: *2019 IEEE International Conference on Industrial Technology (ICIT)*. [S.l.: s.n.], 2019. p. 887–892.
- KAUR, T.; KAMBOJ, S. Descriptive analysis of the cloud computing services and deployment models. In: *2023 International Conference for Advancement in Technology (ICONAT)*. [S.l.: s.n.], 2023. p. 1–6.
- LEDENEV, A. *Spotinfo: Open-source spot instance advisor cli tool*. 2021. <<https://github.com/alexei-led/spotinfo>>.
- LEE, C. C.; LEE, D. T. A simple on-line bin-packing algorithm. *J. ACM*, 1985.
- LEE, S.; HWANG, J.; LEE, K. Spotlake: Diverse spot instance dataset archive service. In: *2022 IEEE International Symposium on Workload Characterization (IISWC)*. [S.l.: s.n.], 2022. p. 242–255.
- LEE, S.; HWANG, J.; LEE, K. *SpotLake Web*. 2025. <<https://spotlake.ddps.cloud/>>. Accessed on January 10, 2025.
- LEE, Y. C.; LIAN, B. Cloud bursting scheduler for cost efficiency. In: *2017 IEEE 10th International Conference on Cloud Computing (CLOUD)*. [S.l.: s.n.], 2017. p. 774–777.
- LI, F.; WU, G.; LU, J.; JIN, M.; AN, H.; LIN, J. Smartcmp: A cloud cost optimization governance practice of smart cloud management platform. In: *2022 IEEE 7th International Conference on Smart Cloud (SmartCloud)*. [S.l.: s.n.], 2022. p. 171–176.
- LI, Z.; ZHANG, H.; O'BRIEN, L.; JIANG, S.; ZHOU, Y.; KIHIL, M.; RANJAN, R. Spot pricing in the cloud ecosystem: A comparative investigation. *Journal of Systems and Software*, v. 114, p. 1–19, 2016. ISSN 0164-1212. Disponível em: <<https://www.sciencedirect.com/science/article/pii/S0164121215002332>>.

Microsoft. *Azure Spot Removal Policies*. 2025. <<https://learn.microsoft.com/pt-br/azure/virtual-machines/spot-vms>>. Accessed on January 10, 2025.

MISHRA, A. K.; KESARWANI, A.; YADAV, D. K. Short term price prediction for preemptible vm instances in cloud computing. In: *2019 IEEE 5th International Conference for Convergence in Technology (I2CT)*. [S.l.: s.n.], 2019. p. 1–9.

MUSHTAQ, M.; AKRAM, U.; KHAN, I.; KHAN, S.; SHAHZAD, A.; ULAH, A. Cloud computing environment and security challenges: A review. *International Journal of Advanced Computer Science and Applications*, v. 8, p. 183–195, 10 2017.

PORTELLA, G. J.; NAKANO, E.; RODRIGUES, G. N.; BOUKERCHE, A.; MELO, A. C. M. A. A novel statistical and neural network combined approach for the cloud spot market. *IEEE Transactions on Cloud Computing*, v. 11, n. 1, p. 278–290, 2023.

WU, X.; YU, H.; CASALE, G.; GAO, G. Towards Cost-Optimal Policies for DAGs to Utilize IaaS Clouds with Online Learning . *IEEE Transactions on Services Computing*, IEEE Computer Society, n. 01, p. 1–18, 2021. ISSN 1939-1374. Disponível em: <<https://doi.ieeecomputersociety.org/10.1109/TSC.2025.3536305>>.

ZHANG, P.; WANG, X.; LI, C.; WANG, L.; XIE, J. A dynamic pricing model for virtual machines in cloud environments. In: *2020 IEEE International Conference on Smart Cloud (SmartCloud)*. [S.l.: s.n.], 2020. p. 13–18.

ZHANG, X.; HUANG, Z.; WU, C.; LI, Z.; LAU, F. C. M. Dynamic vm scaling: Provisioning and pricing through an online auction. *IEEE Transactions on Cloud Computing*, v. 9, n. 1, p. 131–144, 2021.

APÊNDICE A – TABELAS DE AGRUPAMENTO DE INSTÂNCIAS

Tabela 6 – Agrupamento de instâncias de propósito geral

AWS	Azure	GCP
m5a.8xlarge, m5d.large, m5dn.large, m6gd.4xlarge, m5zn.12xlarge, m6in.4xlarge, m7gd.xlarge, m6idn.metal, m2.2xlarge, m4.large, m5ad.large, m7a.48xlarge, m7i.metal-48xl, m5.8xlarge, m6gd.8xlarge, m6in.8xlarge, m6i.32xlarge, m5zn.large, t2.small, t4g.small, t3a.nano, m5ad.2xlarge, m4.4xlarge, m6id.12xlarge, m6gd.16xlarge, m7i-flex.xlarge, m5dn.8xlarge, m7i.24xlarge, m6gd.xlarge, m6i.metal, m7a.large, m6a.xlarge, m7gd.2xlarge, m7i-flex.8xlarge, t3.large, m5d.metal, m6idn.32xlarge, m6id.4xlarge, m5ad.16xlarge, m5a.xlarge, m4.xlarge, m7i.12xlarge, m5ad.4xlarge, t4g.nano, m7g.16xlarge, t3.micro, m5a.24xlarge, m7gd.16xlarge, m6g.metal, m4.10xlarge, m7a.32xlarge, t4g.xlarge, m1.large, m6in.xlarge, m7g.large, m5ad.12xlarge, m5.xlarge, m5n.xlarge, m7g.2xlarge, m6in.24xlarge, m7gd.large, m5a.16xlarge, m7i-flex.large, m6in.2xlarge, m6g.2xlarge, m5n.16xlarge, t4g.2xlarge, m6in.12xlarge, m5n.8xlarge, m5n.large, m7a.24xlarge, m5.24xlarge, m6id.xlarge, m6in.16xlarge, m6i.16xlarge, m7a.medium, m7gd.4xlarge, m5n.4xlarge, m6in.metal, m5zn.2xlarge, m6a.16xlarge, m6gd.metal, m6idn.8xlarge, m6id.8xlarge, m7g.xlarge, m7a.4xlarge, t3.2xlarge, m6a.large, m6a.48xlarge, m6idn.large, t2.medium, m7i.large, m3.medium, m5d.16xlarge, m5d.8xlarge, m7i.16xlarge, t3a.small, m6a.4xlarge, m6a.12xlarge, m5dn.2xlarge, m6in.large, m6g.16xlarge, m6a.32xlarge, m3.2xlarge, m5dn.16xlarge, m5dn.metal, m5ad.24xlarge, m7a.metal-48xl, m5a.4xlarge, m7i.48xlarge, m7a.2xlarge, m7gd.metal, m2.xlarge, m7i.4xlarge, m7a.xlarge, m5.metal, m5n.12xlarge, m6idn.16xlarge, m7g.medium, m6idn.4xlarge, t4g.micro, m7a.12xlarge, m5zn.metal, m6idn.24xlarge, m3.xlarge, m4.16xlarge, m7a.16xlarge, t2.large, m7i-flex.4xlarge, m5ad.8xlarge, m5zn.xlarge, m5dn.4xlarge, t4g.medium, m5.2xlarge, m6id.large, t2.xlarge, m6gd.12xlarge, m5dn.12xlarge, m7g.8xlarge, m5zn.6xlarge, m6a.2xlarge, m6gd.large, m6i.large, m6g.12xlarge, m5d.4xlarge, m5.4xlarge, t3a.large, m6a.metal, m5dn.xlarge, m6idn.12xlarge, m7gd.medium, m7g.4xlarge, m6gd.2xlarge, m6a.8xlarge, m6i.24xlarge, m7g.12xlarge, m7gd.8xlarge, t3a.2xlarge, m5zn.3xlarge, m6id.32xlarge, m5a.large, m6g.4xlarge, m5n.24xlarge, m7a.8xlarge, m6i.8xlarge, m6idn.xlarge, m6i.xlarge, t4g.large, m5d.2xlarge, m5n.metal, m4.2xlarge, t3.nano, m6i.12xlarge, m3.large, m5dn.24xlarge, t3a.micro, m6g.8xlarge, m7g.metal, m1.medium, t1.micro, m5a.12xlarge, m5n.2xlarge, m6g.xlarge, m6id.2xlarge, t3.xlarge, m6g.large, m6idn.2xlarge, m6in.32xlarge, m7i-flex.2xlarge, m7gd.12xlarge, m7i.2xlarge, m5a.2xlarge, m1.small, m6id.24xlarge, t3.medium, m5.16xlarge, m7i.metal-24xl, m7i.xlarge, m5ad.xlarge, m5d.xlarge, m5d.24xlarge, m1.xlarge, m6a.24xlarge, t2.2xlarge, m6id.16xlarge, m6i.2xlarge, m7i.8xlarge, m5d.12xlarge, t2.micro, m6gd.medium, m5.12xlarge, m5.large, m6g.medium, m6i.4xlarge, m6id.metal, t3a.xlarge, t3a.medium, m2.4xlarge, t3.small	basic_a3, basic_a4, basic_a0, basic_a1, basic_a2, _fxmdstypel, standard_d1_v2, standard_d2_v2, standard_d3_v2, standard_d4_v2, standard_d5_v2, standard_d4_v3, standard_d8_v3, standard_d32_v3, standard_d32_v4, standard_d2_v4, standard_d8_v4, standard_d32_v4, standard_d2_v5, standard_d8_v5, standard_d32_v5, standard_ds1_v2, standard_ds3_v2, standard_ds5_v2, standard_b1s, standard_b1ms, standard_b2s, standard_b2ms, standard_b4ms, standard_b8ms, standard_a1_v2, standard_a2_v2, standard_a4_v2, standard_a8_v2, standard_a2m_v2, standard_a4m_v2, standard_a8m_v2, standard_e4s_v3, standard_e16s_v3, standard_e64s_v3	t2a-standard-2, n1-highmem-2, n2d-highcpu-64, g2-standard-24, n2d-standard-16, n1-standard-2, n2d-highmem-2, n2-highmem-64, t2d-standard-16, n1-highcpu-96, n2d-highmem-16, n2-highmem-48, e2-highmem-8, e2-highcpu-32, n2d-highcpu-96, n1-standard-96, n1-highcpu-32, n1-standard-8, n2-highcpu-2, n2-standard-16, n2-standard-4, n1-ultramem-160, n2d-highcpu-128, n2d-standard-64, n2d-standard-8, n1-highmem-64, n2-standard-32, n1-highmem-8, e2-standard-2, n1-highmem-4, g2-standard-48, n1-highcpu-16, g2-standard-16, g2-standard-96, n2-highmem-32, n2-highmem-2, t2a-standard-4, g2-standard-12, n2d-highcpu-16, e2-highcpu-16, n2d-highmem-96, n2-highmem-128, n1-standard-1, n2-standard-2, n2-highcpu-64, n2-standard-96, n1-highcpu-8, e2-highcpu-2, e2-micro, n2d-highcpu-2, e2-highmem-2, n2d-standard-80, t2a-standard-8, t2a-standard-32, t2d-standard-8, n2-highcpu-4, n2d-highmem-32, g1-small, n1-ultramem-40, n2d-highmem-48, n2d-highmem-8, n1-highcpu-4, g2-standard-4, n2d-highcpu-4, n2d-highmem-64, n2-standard-64, t2a-standard-1, e2-standard-8, n2d-standard-48, n1-highmem-96, n2d-standard-96, n2d-highcpu-80, n2d-highcpu-224, n1-highcpu-2, n1-ultramem-80, n2d-highcpu-48, e2-highcpu-8, t2d-standard-48, t2d-standard-60, n2d-highcpu-32, n1-standard-32, e2-highmem-4, n2-highmem-80, t2d-standard-2, n2-highcpu-16, f1-micro, n2d-standard-2, n1-highmem-16, e2-small, n2-highmem-4, n1-standard-16, n1-megamem-96, n2-standard-8, e2-standard-32, n2-highmem-16, n2d-highmem-80, n2-highmem-8, g2-standard-8, t2a-standard-48, n2-highcpu-8, n2-highcpu-48, g2-standard-32, n2-standard-128, n2d-standard-32, e2-standard-16, n2d-standard-224, n2d-standard-4, e2-medium, t2a-standard-16, t2d-standard-4, n2-highcpu-32, n2-highcpu-80, n1-standard-4, n2-standard-80, n1-highcpu-64, n2d-highcpu-8, n2-standard-48, e2-highmem-16, n2d-standard-128, n2d-highmem-4, e2-highcpu-4, n1-standard-64, n1-highmem-32, n2-highcpu-96, t2d-standard-1, n2-highmem-96, t2d-standard-32, e2-standard-4

Fonte: Elaborada pelo autor (2025)

AWS			Azure		GCP	
r3.2xlarge, x1e.2xlarge, x2gd.xlarge, r5b.8xlarge, r6in.large, r7a.xlarge, r7a.large, r7iz.8xlarge, r6idn.8xlarge, r4.8xlarge, x2idn.16xlarge, r6idn.xlarge, x2iezn.2xlarge, r6idn.32xlarge, r7iz.large, r7a.12xlarge, r5n.xlarge, r6a.48xlarge, r6id.12xlarge, r7a.2xlarge, x2iedn.4xlarge, r7g.metal, r5.2xlarge, r4.4xlarge, r5b.metal, r6id.24xlarge, x1e.8xlarge, r5.12xlarge, x2gd.16xlarge, r7gd.4xlarge, r71.12xlarge, r7iz.12xlarge, r6id.2xlarge, r6in.24xlarge, r6g.4xlarge, r7g.4xlarge, r5a.24xlarge, r6a.16xlarge, r4.16xlarge, r5d.large, r7a.32xlarge, r71.24xlarge, r6g.12xlarge, x2idn.metal, r3.large, r5.24xlarge, x2idn.32xlarge, r5a.8xlarge, r5b.xlarge, r5ad.16xlarge, r5a.12xlarge, r6idn.12xlarge, r6a.24xlarge, r7gd.16xlarge, r6id.16xlarge, r5d.24xlarge, r5n.24xlarge, r71.2xlarge, r3.8xlarge, r5ad.8xlarge, r6id.metal, r7g.medium, r5a.large, r5dn.8xlarge, r7a.4xlarge, r3.4xlarge, r5.4xlarge, x2iedn.2xlarge, r5a.16xlarge, r5dn.4xlarge, r7gd.12xlarge, r7g.16xlarge, r5.metal, r7gd.8xlarge,	r5dn.24xlarge, r7iz.2xlarge, r6a.2xlarge, r6in.4xlarge, r6i.32xlarge, r7iz.32xlarge, r7g.8xlarge, r6g.metal, r5.16xlarge, r5b.24xlarge, r5a.4xlarge, r5b.2xlarge, r6idn.metal, x2idn.24xlarge, r6idn.large, r6idn.12xlarge, r5n.2xlarge, r7g.12xlarge, r6i.metal, r7iz.metal-16xl, x1e.4xlarge, r6in.xlarge, r7g.large, r5d.12xlarge, r4.xlarge, z1d.6xlarge, x2gd.8xlarge, r6a.8xlarge, r7gd.2xlarge, r5ad.24xlarge, r7i.metal-48xl, z1d.12xlarge, x1e.xlarge, r5n.metal, r5dn.12xlarge, r5dn.2xlarge, r6id.8xlarge, r5n.4xlarge, r5b.4xlarge, r6i.16xlarge, x1e.32xlarge, r7gd.medium, r6a.4xlarge, r7iz.16xlarge, r5n.large, r6i.24xlarge, r7a.48xlarge, r7gd.metal, z1d.metal, r6gd.8xlarge, z1d.large, r5d.2xlarge, r5dn.large, r5.8xlarge, r5b.16xlarge, r6gd.large, r5a.2xlarge, r6a.xlarge, r7i.xlarge, r7gd.large, r6i.4xlarge, x2gd.2xlarge, r5n.12xlarge, r6gd.xlarge, r5d.xlarge, r6idn.4xlarge, r6id.large, r7i.4xlarge, x2iedn.metal, x1e.16xlarge, r5ad.large, r7iz.xlarge, r7iz.4xlarge, z1d.xlarge, r6gd.4xlarge, r6g.8xlarge, r6g.16xlarge, x2iedn.16xlarge, r4.2xlarge, r5a.xlarge,	standard_eb2s_v5, standard_eb8s_v5, standard_eb32s_v5, standard_ce2s_v4, standard_ce8s_v4, standard_ce32s_v4, standard_m2ms, standard_m8ms, standard_m32ms, standard_e2s_v4, standard_e8s_v4, standard_e32s_v4, standard_eb2s_v3, standard_eb8s_v3, standard_eb32s_v3, standard_ce2s_v3, standard_ce8s_v3, standard_ce32s_v3, standard_m2s_v2, standard_m8s_v2, standard_m32s_v2, standard_e2as_v4, standard_e8as_v4, standard_e32as_v4, standard_eb2as_v4, standard_eb8as_v4, standard_eb32as_v4, standard_ce2as_v5, standard_ce8as_v5, standard_ce32as_v5, standard_m2ms_v5, standard_m8ms_v5, standard_m32ms_v5, standard_e2ds_v5, standard_e8ds_v5, standard_e32ds_v5, standard_eb2ds_v5, standard_eb8ds_v5, standard_eb32ds_v5, standard_ce2ds_v5, standard_ce8ds_v5, standard_ce32ds_v5, standard_m2ds_v5, standard_m8ds_v5, standard_m32ds_v5, standard_e2ts_v5, standard_e8ts_v5, standard_e32ts_v5, standard_eb2ts_v5, standard_eb8ts_v5, standard_eb32ts_v5,	standard_eb4s_v5, standard_eb16s_v5, standard_eb64s_v5, standard_ce4s_v4, standard_ce16s_v4, standard_ce64s_v4, standard_m4ms, standard_m16ms, standard_m64ms, standard_e4s_v4, standard_e16s_v4, standard_e64s_v4, standard_eb4s_v3, standard_eb16s_v3, standard_eb64s_v3, standard_ce4s_v3, standard_ce16s_v3, standard_ce64s_v3, standard_m4s_v2, standard_m16s_v2, standard_m64s_v2, standard_e4as_v4, standard_e16as_v4, standard_e64as_v4, standard_eb4as_v4, standard_eb16as_v4, standard_eb64as_v4, standard_ce4as_v5, standard_ce16as_v5, standard_ce64as_v5, standard_m4ms_v5, standard_m16ms_v5, standard_m64ms_v5, standard_e4ds_v5, standard_e16ds_v5, standard_e64ds_v5, standard_eb4ds_v5, standard_eb16ds_v5, standard_eb64ds_v5, standard_ce4ds_v5, standard_ce16ds_v5, standard_ce64ds_v5, standard_m4ds_v5, standard_m16ds_v5, standard_m64ds_v5, standard_e4ts_v5, standard_e16ts_v5, standard_e64ts_v5, standard_eb4ts_v5, standard_eb16ts_v5, standard_eb64ts_v5,	m3-ultramem-64, m1-ultramem-40, m3-ultramem-128, m1-megamem-96, ultramem-32, x4-megamem-1440-metal, x4-megamem-1920-metal, x4-megamem-1440-metal, m2-ultramem-208, ultramem-416, m2-hypermem-416, m2-ultramem-416, m2-hypermem-416	m1-ultramem-160, m3-megamem-128, m3-megamem-64, m1-ultramem-80, m3-megamem-960-metal, x4-megamem-960-metal, x4-megamem-960-metal, m2-ultramem-208, m2-megamem-416, m2-ultramem-208, m2-megamem-416	

Fonte: Elaborada pelo autor (2025)

AWS			Azure		GCP		
c6id.2xlarge, c6g.16xlarge, c5n.9xlarge, c6a.metal, c7g.medium, c6in.metal, c7a.12xlarge, flex.8xlarge, c4.8xlarge, c5a.24xlarge, c1.xlarge, c7i-flex.2xlarge, c7gd.4xlarge, c5d.12xlarge, c6a.16xlarge, c5a.large, c5ad.8xlarge, c1.medium, c4.4xlarge, c6gd.4xlarge, c7gn.8xlarge, c4.large, c5.large, c7gd.large, c5n.large, c6gd.8xlarge, c7g.4xlarge, c5n.18xlarge, c5ad.4xlarge, c6i.24xlarge, c7i.metal-48xl, c7i-flex.4xlarge, c7i.2xlarge, c5d.24xlarge, c5d.18xlarge, c7i-flex.large, c6a.4xlarge, c5.12xlarge, c5a.large, c5ad.2xlarge, c5d.metal, c6gd.medium, c6gn.large, c5a.12xlarge, c5a.xlarge, c6id.16xlarge, c6a.24xlarge, c6in.16xlarge, c3.4xlarge, c6g.2xlarge, c6gd.16xlarge, c7i-flex.xlarge, c5d.xlarge, c6gd.2xlarge, c6in.8xlarge, c6g.4xlarge, c7i.4xlarge, c6g.large, c6a.48xlarge, c6in.24xlarge	c7g.12xlarge, c6gn.4xlarge, c7a.8xlarge, c7g.xlarge, c7i.24xlarge, c7gn.12xlarge, c7a.metal-48xl, c7i.48xlarge, c3.large, c5a.8xlarge, c7gn.4xlarge, c4.xlarge, c5a.16xlarge, c6a.large, c6gn.12xlarge, c7a.2xlarge, c6gn.16xlarge, c7g.metal, c6id.metal, c6gd.large, c7a.24xlarge, c6id.32xlarge, c6in.32xlarge, c7i.16xlarge, c6g.12xlarge, c6i.8xlarge, c6i.12xlarge, c5n.metal, c7gd.8xlarge, c5n.2xlarge, c5d.18xlarge, c6i.4xlarge, c7gd.xlarge, c5n.4xlarge, c6a.xlarge, c6a.2xlarge, c6i.large, c6gn.xlarge, c7g.16xlarge, c7i.xlarge, c5ad.16xlarge, c5.4xlarge, c7gn.2xlarge, c7gd.medium, c5d.2xlarge, c5.2xlarge, c5ad.xlarge, c7gn.16xlarge, c6in.2xlarge, c6id.12xlarge, c6i.metal, c7gd.12xlarge, c7a.xlarge, c5a.2xlarge, c7i.8xlarge, c7a.4xlarge, c6a.32xlarge, c5.18xlarge, c6gn.8xlarge, c6in.24xlarge	c7a.medium, c7gn.large, c6gn.2xlarge, c7g.large, c4.2xlarge, c6id.24xlarge, c7i- c6a.12xlarge, c5.24xlarge, c3.2xlarge, c6a.8xlarge, c7gn.xlarge, c6g.medium, c5ad.24xlarge, c6i.xlarge, c5.9xlarge, c7gn.medium, c6in.4xlarge, c7gd.16xlarge, c6gn.medium, c5.metal, c6g.xlarge, c3.8xlarge, c5d.4xlarge, c6i.32xlarge, c6in.12xlarge, c5ad.large, c5d.9xlarge, c6id.4xlarge, c6id.large, c7i.metal-48xl, c5n.4xlarge, c6a.xlarge, c6a.2xlarge, c3.xlarge, c6in.xlarge, c7gn.metal, c5d.large, c6gd.metal, c6id.xlarge, c5n.large, c6gd.xlarge, c6i.2xlarge, c6id.8xlarge, c7a.48xlarge, c6g.8xlarge, c7i.12xlarge, c6in.large, c5a.4xlarge, c7a.16xlarge, c5ad.12xlarge, c6g.metal, c7a.32xlarge, c7g.2xlarge, c7gd.2xlarge, c6i.16xlarge,	standard_F2s_v2, standard_F8s_v2, standard_F32s_v2, standard_F64s_v2, standard_F2as_v6, standard_F8as_v6, standard_F32as_v6, standard_F64as_v6, standard_F4als_v6, standard_F16als_v6, standard_F8als_v6, standard_F2ams_v6, standard_F4ams_v6, standard_F16ams_v6, standard_F32ams_v6, standard_F64ams_v6, standard_F2s_v2, standard_F8s_v2, standard_F32s_v2, standard_F64s_v2, standard_FX2ms_v2, standard_FX8ms_v2, standard_FX12ms_v2, standard_FX24ms_v2, standard_FX32ms_v2, standard_FX64ms_v2, standard_FX96ms_v2, standard_FX4mds_v2, standard_FX8mds_v2, standard_FX16mds_v2, standard_FX24mds_v2, standard_FX32mds_v2, standard_FX48mds_v2, standard_FX64mds_v2, standard_FX96mds_v2, standard_FX12mds, standard_FX36mds	standard_F4s_v2, standard_F16s_v2, standard_F48s_v2, standard_F72s_v2, standard_F4as_v6, standard_F16as_v6, standard_F48as_v6, standard_F2als_v6, standard_F8als_v6, standard_F32als_v6, standard_F64als_v6, standard_F8ams_v6, standard_F48ams_v6, standard_F4s_v2, standard_F16s_v2, standard_F48s_v2, standard_F72s_v2, standard_FX4ms_v2, standard_FX16ms_v2, standard_FX24ms_v2, standard_FX48ms_v2, standard_FX2mds_v2, standard_FX4mds_v2, standard_FX12mds_v2, standard_FX16mds_v2, standard_FX24mds_v2, standard_FX32mds_v2, standard_FX48mds_v2, standard_FX64mds_v2, standard_FX96mds_v2, standard_FX4mds, standard_FX24mds, standard_FX48mds	c3d-highmem-8, standard-16, c3-highmem-88, standard-90-1ssd, c2d-highmem-4, highmem-360-1ssd, c3d-standard-30-1ssd, c2d-highcpu-2, c3-highcpu-60, c2d-highcpu-112, highmem-22, c3d-highcpu-30, c3-highcpu-8, c2d-standard-56, c3-highcpu-4, c3-highmem-4, c3-highmem-176, c3-standard-44-1ssd, c3-standard-8-1ssd, c3-highcpu-180, c3-standard-88-1ssd, c3-standard-176-1ssd, c2d-standard-8, c2d-highmem-32, c2d-standard-16, standard-88, c2d-highmem-8, c3d-highcpu-90, c2d-highmem-56, c2-standard-4, c3-standard-22-1ssd, c3d-highcpu-360, c2d-highmem-2, c2d-standard-112, c3-standard-4, c2d-highmem-16, c3d-highmem-30, c3-standard-8, c3-highcpu-88, c3d-highcpu-16, c2d-highcpu-32, c2d-standard-32, c3-standard-44, c3-standard-22, c3d-standard-360-1ssd, c2d-highcpu-56, c2d-standard-2, c3d-standard-8, c3d-standard-90, c3-highmem-8, c2d-highcpu-8, c3-standard-176, c3-highcpu-22, c2-standard-4, c2d-standard-4, c3d-standard-60, c2-standard-60, c3d-standard-8-1ssd, c3d-highmem-60	c3d-standard-30, c3d-highmem-90-1ssd, c3d-highmem-360, c3d-standard-4-1ssd, c3d-highmem-180, c3d-highmem-90, c3d-highmem-8-1ssd, c3d-highmem-16, c3d-highcpu-8, c3-highcpu-44, c3d-highcpu-18, c3-highmem-4, c3-highcpu-8, c3d-standard-180-1ssd, c3d-highmem-180-1ssd, c3d-standard-60-1ssd, c3d-standard-360, c3-highmem-44, c2d-highcpu-16, c3d-standard-16-1ssd, c2d-highmem-112, c3d-highmem-30-1ssd, c2-standard-30, c3d-highmem-16-1ssd, c3-highcpu-4, c3-standard-88, c3d-highcpu-4, c3-standard-16, c2d-highcpu-32, c2d-standard-44, c3-standard-56, c3d-standard-8, c3d-standard-90, c2d-highcpu-8, c3d-standard-22, c2d-standard-4, c2d-standard-4, c3d-standard-60, c2-standard-60, c3d-standard-8-1ssd, c3d-highmem-60	

Tabela 9 – Agrupamento de instâncias de Computação Acelerada

AWS			Azure		GCP
g6.24xlarge, inf1.24xlarge, g5g.16xlarge, p5.48xlarge, dl1.24xlarge, g4dn.12xlarge, g5.24xlarge, g5.4xlarge, inf2.8xlarge, g5g.2xlarge, dl2q.24xlarge, p3.2xlarge, f1.16xlarge, vt1.3xlarge, g6.48xlarge, g3.4xlarge, p3dn.24xlarge, g6.4xlarge, trn1.2xlarge, g4dn.xlarge, g4ad.xlarge, g5g.4xlarge, g6.12xlarge, g6.xlarge	g3.8xlarge, g6.2xlarge, p3.8xlarge, trn1n.32xlarge, g4dn.metal, p2.16xlarge, inf2.xlarge, g2.8xlarge, g5g.metal, vt1.24xlarge, inf1.xlarge, g3s.xlarge, g4dn.8xlarge, inf2.24xlarge, g6.8xlarge, gr6.8xlarge, vt1.6xlarge, g5.8xlarge, g4dn.2xlarge, g5.2xlarge, g4ad.16xlarge, g6.16xlarge, p2.xlarge	g5.12xlarge, g4ad.2xlarge, g4dn.4xlarge, g5.48xlarge, g5g.8xlarge, f1.4xlarge, g4ad.8xlarge, gr6.4xlarge, g5g.xlarge, inf2.48xlarge, p4d.24xlarge, g2.2xlarge, trn1.32xlarge, inf1.2xlarge, f1.2xlarge, g4ad.4xlarge, p2.8xlarge, g4dn.16xlarge, inf1.6xlarge, g5.xlarge, p3.16xlarge, g5.16xlarge, g3.16xlarge	standard_hb120-64rs_v2, standard_d16as_v4, standard_m16s, standard_e64s_v4, standard_d4d_v4, standard_dc16as_cc_v5, standard_d2als_v6, standard_e8-4s_v6, standard_e8s_v3, standard_dc48es_v5, standard_e4-2as_v4, standard_gs4-8, standard_e96s_v6, standard_ec8eds_v5, standard_np20s, standard_nc24, standard_ds14_v2, standard_d1, standard_e8d_v4, standard_e8-4as_v4, standard_d32s_v6, standard_sqlg7_amd_iaas, standard_d16plds_v5, standard_e8-2s_v3, standard_ds1_v2, standard_e8-4ads_v5, standard_b8as_v2, standard_e16-8ds_v6, standard_e64s_v6, standard_e192ids_v6, standard_b4pls_v2, standard_d64s_v4, standard_e8-2ads_v5, standard_e8ds_v5, standard_a0, standard_nv48s_v3, standard_ng8ads_v620_v1, standard_a8_v2, standard_e16bds_v5, standard_e64a_v4, standard_e8_v4, standard_l32s, standard_e16-8s_v6, standard_d32alds_v6, standard_e16ps_v5, standard_m16-4ms, standard_m32dms_v2, standard_ec128eds_v5, standard_l16s, standard_d2plds_v5, standard_e48ds_v4, standard_e48ads_v6, standard_m208ms_v2, standard_e2ads_v6, standard_m192idms_v2, standard_d2d_v5, standard_f48amds_v6, standard_b2ps_v2, standard_ds3, standard_l32s_v3, standard_e128ds_v6, standard_e48_v3, standard_e96-48as_v4, standard_d2ls_v6, standard_e8ps_v5, standard_d48s_v4, standard_e32ads_v6, standard_nv4as_v4, standard_e64-32ds_v4, standard_e64_v4, standard_e20_v5, standard_ec20as_cc_v5, standard_d32ps_v5, standard_ng32adms_v620, standard_dc96ads_v5, standard_m48ds_1_v3, standard_d2pls_v5, standard_e64ds_v4, standard_ec32ads_v5, standard_ec96as_cc_v5, standard_f2s_v2, standard_ec8as_cc_v5, standard_b8pls_v2, standard_m192is_v2, standard_e20as_v4, standard_e64_v5, standard_m32ms_v2, standard_e20s_v6, standard_d32lds_v6, standard_e64as_v5, standard_e96iads_v5, standard_e112ias_v5, standard_e16s_v3, standard_e20s_v4, standard_f4ads_v6, standard_f16amds_v6, standard_f64as_v6, standard_ec96iads_v5, standard_e8-4s_v5, standard_d16ds_v4, standard_l64as_v3, standard_d48as_v5, standard_ec32as_v5, standard_d64s_v5, standard_dc32ads_v5, standard_dc32eds_v5, standard_ds1, standard_f2ads_v6, standard_ds12-1_v2, standard_gs3, standard_dc32ds_v3, standard_e32bds_v5, standard_fx24mds, standard_f4alds_v6, standard_d64_v4, standard_nd6s, standard_d64d_v4, standard_e8-4s_v4, standard_e64-16ds_v6, standard_f8ads_v6, standard_gs4, standard_dc32as_cc_v5, standard_m176s_3_v3, standard_d48as_v6, standard_d16als_v6, standard_d2ls_v5, standard_e96as_v6, standard_d11, standard_e96d_v5, standard_e4-2ds_v5, standard_e2as_v4, standard_d64-16s_v3, standard_l80as_v3, standard_b16s_v2, standard_e64-32ds_v6, standard_d48s_v5, standard_e128-32ds_v6, standard_ec64as_v5, standard_e2_v5, standard_d48pds_v5, standard_e64bds_v5,	a2-megagpu-16g, a2-highgpu-1g, a2-ultragpu-8g, a2-ultragpu-1g, a2-highgpu-2g, a2-highgpu-8g, a2-ultragpu-4g, a2-ultragpu-2g, a2-highgpu-4g, a3-highgpu-8g, g2-standard-4, g2-standard-8, g2-standard-12, g2-standard-16, g2-standard-24, g2-standard-32, g2-standard-48, g2-standard-96	

Tabela 6 - Agrupamento de instâncias de Computação Acelerada (continuação)

AWS	Azure	GCP
-	standard_d32as_v4, standard_d8d_v5, - standard_m32s, standard_nc8ads_a10_v4, standard_ec16ads_v5, stan- dard_nc24rs_v3, standard_d14, stan- dard_m416ms_v2, standard_e64ds_v5, standard_ng32ads_v620_v1, stan- dard_nd96ams_a100_v4, standard_e16- 4as_v5, standard_e32-16s_v4, stan- dard_g3, standard_e8bds_v5, stan- dard_l48s_v2, standard_d8alds_v6, stan- dard_e32as_v4, standard_d4ads_v5, stan- dard_m48s_1_v3, standard_d32ads_v5, standard_ec2eds_v5, standard_d64ds_v5, standard_d64ads_v5, standard_e64- 32ds_v5, standard_d4alds_v6, stan- dard_e32as_v6, standard_e64ads_v5, standard_f64ams_v6, standard_e48s_v5, standard_d8as_v5, standard_d8ds_v5, standard_d48a_v4, standard_m208- 52ms_v2, standard_nv8as_v4, stan- dard_ds5_v2, standard_f1s, standard_ds11- 1_v2, standard_nc96ads_a100_v4, stan- dard_dc4as_cc_v5, standard_e64-16as_v5, standard_nd24s, standard_d16_v5, stan- dard_dc4s_v3, standard_m32ls, stan- dard_ec64es_v5, standard_d64alds_v6, standard_d8a_v4, standard_dc8as_v5, standard_d8s_v4, standard_e64- 32as_v5, standard_e8ds_v6, stan- dard_dc128eds_v5, standard_ds13-2_v2, standard_d64-32s_v3, standard_a1, stan- dard_d96as_v5_promo, standard_f2als_v6, standard_e4d_v4, standard_e48as_v5, standard_d4ls_v6, standard_fx12mds, standard_d48lds_v5, standard_e96- 24ds_v6, standard_a2, standard_m32-8ms, standard_f64als_v6, standard_d4s_v3, standard_e8s_v5, standard_d8_v3, stan- dard_m128s, standard_f16als_v6, stan- dard_b2s_v2, standard_d16as_v5_promo, standard_e32-8s_v4, standard_d12, standard_d96d_v5, standard_e8-2s_v4, standard_d8plds_v5, standard_e32-16s_v5, standard_f72als_v6, standard_dc2s, stan- dard_f16ams_v6, standard_m128-64ms, standard_l8s_v3, standard_e32-8as_v4, standard_hx176-24rs, standard_f64s_v2, standard_nv6ads_a10_v5, stan- dard_e2bds_v5, standard_d64ds_v4, stan- dard_e2as_v5, standard_dc48as_v5, stan- dard_f1alds_v6, standard_d32als_v6, stan- dard_d48ads_v5, standard_e112iads_v5, standard_dc8eds_v5, standard_e32_v5, standard_dc64as_cc_v5, stan- dard_d96as_v6, standard_f72ads_v6, stan- dard_e96as_v5, standard_e64-32as_v4, standard_e2d_v5, standard_d48alds_v6, standard_e48ads_v5, standard_b32als_v2, standard_d64d_v5, standard_e16as_v6, standard_d48_v3, standard_e48as_v6, standard_ds12_v2, standard_e32-16ds_v5, standard_e32-8ds_v5, standard_d32_v3, standard_d8ds_v6, standard_e32-8ads_v5, standard_dc4eds_v5, standard_e32- 16as_v5, standard_ec16eds_v5, stan- dard_d2s_v3, standard_d48_v5, standard_hx176-144rs, standard_dc32s_v3, standard_d128ls_v6, standard_e48d_v4, standard_f48ams_v6, standard_m128- 32ms, standard_e64ads_v6, stan- dard_dc8ds_v3, standard_d64ds_v6, standard_e8-2as_v4, standard_d96as_v4, standard_d2ads_v6, standard_nd40s_v2, standard_e64ids_v4_special, stan- dard_l80s_v3, standard_d16a_v4, stan- dard_e2d_v4, standard_m416-208ms_v2, standard_nv24, standard_e64-16s_v3,	

Tabela 6 - Agrupamento de instâncias de Computação Acelerada (continuação)

AWS	Azure	GCP
-	standard_m192ims_v2, standard_f32alds_v6, standard_l88is_v2, standard_dc96as_cc_v5, standard_e8-4s_v3, standard_d4_v4, standard_m32-16ms, standard_d64as_v5_promo, standard_d64lds_v5, standard_ec96ias_v5, standard_e192is_v6, standard_m32ts, standard_d2ds_v5, standard_d4ds_v4, standard_hb120rs_v3, standard_ec2es_v5, standard_f1, standard_hb60rs, standard_e8ads_v5, standard_e8-4ds_v5, standard_d48als_v6, standard_e96-24as_v5, standard_e2ds_v5, standard_e16ads_v5, standard_d3_v2, standard_e16bs_v5, standard_nc4as_t4_v3, standard_e96-24s_v5, standard_nv36adms_a10_v5, standard_dc96ads_cc_v5, standard_f1ads_v6, standard_e8s_v4, standard_e16-4s_v5, standard_d8as_v6, standard_d15_v2, standard_b2ats_v2, standard_e64bs_v5, standard_l32as_v3, standard_e8as_v4, standard_e8as_v5, standard_m208-52s_v2, standard_f72alds_v6, standard_e48_v5, standard_m128ds_v2, standard_f1amds_v6, standard_e8a_v4, standard_b4ls_v2, standard_pb24s, standard_f64amds_v6, standard_f1ams_v6, standard_d5_v2, standard_nv12s_v2, standard_ds2, standard_e96-48s_v6, standard_d32a_v4, standard_e16-4ads_v5, standard_l8s, standard_f4ams_v6, standard_d48ds_v5, standard_e16as_v5, standard_dc4s_v2, standard_d2d_v4, standard_e64s_v5, standard_d32-16s_v3, standard_d4, standard_f72iamds_v6, standard_f8alds_v6, standard_nc8as_t4_v3, standard_d32s_v3, standard_e2s_v3, standard_d48lds_v6, standard_e16-4s_v6, standard_f48as_v6, standard_dc32as_v5, standard_d32s_v5, standard_dc8s_v3, standard_ec96eds_v5, standard_e64ds_v6, standard_e16-8ds_v4, standard_d16s_v4, standard_d64as_v5, standard_e2_v3, standard_m64ms, standard_d4as_v6, standard_m96s_1_v3, standard_f4, standard_dc4ads_v5, standard_f4s, standard_e104i_v5, standard_d96ds_v6, standard_e4ps_v5, standard_d64pls_v5, standard_dc4es_v5, standard_b2ls_v2, standard_dc1ds_v3, standard_e8s_v6, standard_m192ids_v2, standard_ec2ads_v5, standard_dc24ds_v3, standard_d32d_v5, standard_e32_v3, standard_dc96es_v5, standard_d4plds_v5, standard_a4, standard_d32pds_v5, standard_b2as_v2, standard_hc44-32rs, standard_m192ms_v2, standard_e8ds_v4, standard_e96-24ds_v5, standard_hb120-16rs_v3, standard_hb176rs_v4, standard_e4_v3, standard_e32ds_v5, standard_e64-16s_v5, standard_d16d_v5, standard_d16ads_v6, standard_e4ds_v6, standard_ec48as_v5, standard_d8ls_v6, standard_m12s_v3, standard_nc24s_v2, standard_hb60-30rs, standard_ec4as_cc_v5, standard_hb176-144rs_v4, standard_hb120-32rs_v2, standard_e64-32ads_v5, standard_d32ds_v4, standard_m416-104s_v2, standard_ds11_v2, standard_d4_v2, standard_e32s_v4, standard_ds13-4_v2, standard_f16alds_v6, standard_d12_v2, standard_d16pls_v5, standard_nv18ads_a10_v5, standard_fx48mds, standard_m32ms, standard_d96lds_v5, standard_dc64ads_v5, standard_ec32eds_v5, standard_ds12-2_v2, standard_d2_v5, standard_e96bs_v5,	-

Tabela 6 - Agrupamento de instâncias de Computação Acelerada (continuação)

AWS	Azure	GCP
-	standard_d2_v3, standard_nd12s, standard_e64is_v3, standard_hb60-45rs, standard_e4-2ds_v4, standard_m8-2ms, standard_e16-8s_v5, standard_ec32as_cc_v5, standard_e4d_v5, standard_dc2ads_v5, standard_ec16es_v5, standard_d2lds_v5, standard_d96ls_v5, standard_ds14-4_v2, standard_m12ds_v3, standard_f4amds_v6, standard_hb120-96rs_v2, standard_l16as_v3, standard_ec48eds_v5, standard_nd96asr_v4, standard_ec8es_v5, standard_d96als_v6, standard_e64d_v5, standard_m416-208s_v2, standard_e96-24ads_v5, standard_e20ds_v6, standard_d96ads_v5, standard_dc2s_v2, standard_dc8ads_v5, standard_d64pds_v5, standard_nc32ads_a10_v4, standard_b32s_v2, standard_nc6s_v3, standard_d96lds_v6, standard_d32_v5, standard_b8s_v2, standard_dc16ads_cc_v5, standard_e2bs_v5, standard_ec20ads_cc_v5, standard_nc80adis_h100_v5, standard_d32ds_v6, standard_d16alds_v6, standard_d48d_v5, standard_d13_v2, standard_dc2es_v5, standard_e80is_v4, standard_hb120-96rs_v3, standard_nv32as_v4, standard_dc128es_v5, standard_nc12s_v2, standard_e2pds_v5, standard_l8s_v2, standard_d96as_v5, standard_d64_v5, standard_f2amds_v6, standard_nd96amsr_a100_v4, standard_ds3_v2, standard_d16_v4, standard_f48als_v6, standard_m128m, standard_e96ias_v4, standard_d96s_v6, standard_d8ps_v5, standard_d4ads_v6, standard_f48ads_v6, standard_b8als_v2, standard_d2as_v5, standard_dc4ds_v3, standard_nc40ads_h100_v5, standard_hb176-96rs_v4, standard_e8as_v6, standard_nc24rs_v2, standard_h8, standard_e16-8as_v4, standard_d48_v4, standard_e2a_v4, standard_f32als_v6, standard_e16-8as_v5, standard_dc24s_v3, standard_e8d_v5, standard_f32amds_v6, standard_g1, standard_e16d_v5, standard_e48s_v3, standard_dc96as_v5, standard_d32ads_v6, standard_e32s_v6, standard_d16_v3, standard_nv6, standard_l96s_v2, standard_d48plds_v5, standard_e64-16s_v4, standard_d16ads_v5, standard_pb12s, standard_m176s_4_v3, standard_f16s, standard_f8as_v6, standard_nc6s_v2, standard_nc6, standard_d32_v4, standard_dc48ads_cc_v5, standard_e16-4s_v3, standard_e96ds_v6, standard_e96ias_v5, standard_m24ds_v3, standard_d2_v4, standard_m192s_v2, standard_d8_v5, standard_dc8ads_cc_v5, standard_d4lds_v6, standard_d4as_v5, standard_e64-16ads_v5, standard_d48ps_v5, standard_m64ds_v2, standard_e32-8ds_v4, standard_ds4_v2, standard_d8s_v6, standard_e64is_v4_special, standard_d8pds_v5, standard_d48ls_v5, standard_e20_v4, standard_e20ps_v5, standard_e48bs_v5, standard_e8-2s_v5, standard_hb176-24rs_v4, standard_f32s_v2, standard_d8as_v4, standard_d16s_v6, standard_d128ds_v6, standard_ds15_v2, standard_l16s_v3, standard_e4ads_v6, standard_e64-16ds_v5, standard_e96ads_v5, standard_d64ls_v5, standard_dc16es_v5, standard_e64-32s_v5, standard_d4_v3, standard_e4s_v3, standard_d4s_v5,	-

Tabela 6 - Agrupamento de instâncias de Computação Acelerada (continuação)

AWS	Azure	GCP
-	standard_d96ls_v6, standard_e20ds_v4, - standard_dc8_v2, standard_m416s_8_v2, standard_ec48es_v5, standard_e32-8s_v5, standard_d8lds_v6, standard_e32s_v3, standard_b2pts_v2, standard_nv16as_v4, standard_l48as_v3, standard_d48s_v6, standard_e32a_v4, standard_e96-48s_v5, standard_e16-8s_v3, standard_f1als_v6, standard_d8als_v6, standard_nd40rs_v2, standard_b16ps_v2, standard_e96- 48ds_v6, standard_e80ids_v4, standard_hb120-64rs_v3, standard_gs5, standard_gs2, standard_d48ls_v6, stan- dard_ng16ads_v620_v1, standard_nc12, standard_dc16ads_v5, standard_e4pds_v5, standard_d32ls_v5, standard_hb176- 48rs_v4, standard_d32ds_v5, stan- dard_m16ms, standard_l64s_v3, stan- dard_e16ads_v6, standard_d32lds_v5, standard_nv12ads_a10_v5, stan- dard_f32as_v6, standard_e2ps_v5, standard_e96-48ads_v5, stan- dard_e104id_v5, standard_ec128es_v5, standard_l64s_v2, standard_m128dms_v2, standard_e4s_v6, standard_ds13_v2, stan- dard_f72as_v6, standard_l96ias_v3, standard_d8ads_v5, standard_e8- 2s_v6, standard_e4-2s_v3, stan- dard_m24s_v3, standard_d32as_v6, standard_d48as_v5_promo, stan- dard_e16s_v6, standard_d4lds_v5, stan- dard_e32ds_v4, standard_d4s_v6, stan- dard_ec96ads_cc_v5, standard_e32d_v4, standard_e16a_v4, standard_e8_v5, stan- dard_b32ls_v2, standard_ec4ads_cc_v5, standard_hc44-16rs, standard_d4a_v4, standard_e64as_v6, standard_m8-4ms, standard_e8-4ds_v6, standard_m64- 16ms, standard_b16pls_v2, stan- dard_e64_v3, standard_e2s_v4, stan- dard_d4_v5, standard_l8as_v3, stan- dard_dc64eds_v5, standard_e16-8ads_v5, standard_f8als_v6, standard_e4as_v6, standard_ec4ads_v5, standard_h16mr, standard_e8-4ds_v4, standard_b16ls_v2, standard_m64s, standard_dc8es_v5, standard_e48as_v4, standard_e2s_v6, standard_e16_v4, standard_h16m, standard_e20s_v5, standard_a6, stan- dard_dc32es_v5, standard_d2_v2, standard_f48s_v2, standard_d16d_v4, standard_f1as_v6, standard_e4_v4, standard_e64s_v3, standard_ds4, stan- dard_ds2_v2, standard_nd96isr_h100_v5, standard_dc48ads_v5, standard_d16s_v3, standard_nc24ads_a100_v4, stan- dard_e96ds_v5, standard_d4ahs_v4, standard_ec32ads_cc_v5, stan- dard_e4ads_v5, standard_m64-32ms, standard_d64a_v4, standard_hb120rs_v2, standard_d2pds_v5, standard_m128, standard_d15i_v2, standard_d2ps_v5, standard_e4-2s_v5, standard_ec4as_v5, standard_m176ds_3_v3, stan- dard_sqlg7_amd_nvme, standard_pb6s, standard_d13, standard_e16pds_v5, standard_f2ams_v6, standard_e4ds_v5, standard_e4-2s_v4, standard_f4s_v2, standard_d48d_v4, standard_f16s_v2, standard_d4s_v4, standard_m16-8ms, standard_d96ads_v6, standard_e48ds_v5, standard_sqlg7_nvme, standard_e48s_v6, standard_d32-8s_v3, standard_f48alds_v6, standard_nv36ads_a10_v5, stan- dard_ds14, standard_d48pls_v5, stan- dard_d64_v3, standard_dc2eds_v5, standard_d16as_v5, standard_e16ds_v4, standard_e64as_v4, standard_m128ms, standard_ec16as_cc_v5,	

Tabela 6 - Agrupamento de instâncias de Computação Acelerada (continuação)

AWS	Azure	GCP
-	standard_dc1s_v2, standard_gs4-4, standard_ec20as_v5, standard_e20ads_v5, standard_e4-2as_v5, standard_e16d_v4, standard_e16-8ds_v5, standard_e2as_v6, standard_e32ps_v5, standard_b4s_v2, standard_d11_v2, standard_d32d_v4, standard_f4as_v6, standard_d32ls_v6, standard_e20pds_v5, standard_ds13, standard_m96s_2_v3, standard_dc16s_v3, standard_e48a_v4, standard_d8ls_v5, standard_fx36mids, standard_f2, standard_a3, standard_b32as_v2, standard_m208-104s_v2, standard_e8bs_v5, standard_dc64es_v5, standard_dc96eds_v5, standard_e48_v4, standard_d16s_v5, standard_e64i_v3, standard_e32_v4, standard_m416s_v2, standard_e32ds_v6, standard_b2ts_v2, standard_b16als_v2, standard_e4-2s_v6, standard_dc16ds_v3, standard_np40s, standard_g5, standard_dc32ads_cc_v5, standard_d3, standard_a2m_v2, standard_dc1s_v3, standard_d8s_v5, standard_f64alds_v6, standard_e4_v5, standard_f8ams_v6, standard_gs1, standard_ds12, standard_e8-2ds_v4, standard_d128s_v6, standard_l16s_v2, standard_e32-16s_v6, standard_e4-2ds_v6, standard_e96_v5, standard_ec2as_v5, standard_d4ds_v5, standard_d128lds_v6, standard_nc12s_v3, standard_e48d_v5, standard_ec4eds_v5, standard_e32-8ds_v6, standard_d2ds_v6, standard_l112ias_v3, standard_e96bds_v5, standard_dc16eds_v5, standard_nd96is_h100_v5, standard_ec16as_v5, standard_ec4es_v5, standard_ec96es_v5, standard_e112ibs_v5, standard_e8-2as_v5, standard_e2ds_v4, standard_e104ids_v5, standard_e32-8s_v6, standard_e16-4ds_v5, standard_d4ds_v6, standard_d16pds_v5, standard_e96-48ds_v5, standard_e20ads_v6, standard_e20a_v4, standard_dc64as_v5, standard_e4s_v4, standard_d2, standard_e4s_v5, standard_d2s_v4, standard_m8ms, standard_nc16as_t4_v3, standard_ec64eds_v5, standard_m64s_v2, standard_d8as_v5_promo, standard_a7, standard_f64ads_v6, standard_m208-104ms_v2, standard_d2as_v6, standard_d16ls_v6, standard_d64as_v4, standard_f2as_v6, standard_e64id_v4_special, standard_e2ds_v6, standard_b2als_v2, standard_d2ads_v5, standard_d8_v4, standard_e2ads_v5, standard_f8, standard_e16as_v4, standard_e96-48as_v5, standard_d8ads_v6, standard_nv12, standard_ec8ads_cc_v5, standard_dc4s, standard_d64ads_v6, standard_f4als_v6, standard_ec96as_v5, standard_ec48ads_v5, standard_e64-32s_v4, standard_nc48ads_a100_v4, standard_d4pds_v5, standard_d16as_v6, standard_dc48s_v3, standard_d8lds_v5, standard_d32s_v4, standard_b2pls_v2, standard_d4d_v5, standard_nv24s_v2, standard_e16s_v4, standard_ec8as_v5, standard_d2ds_v4, standard_hb120-32rs_v3, standard_e32bs_v5, standard_e20as_v5, standard_gs5-16, standard_m416is_v2, standard_hx176rs, standard_e32-16as_v4, standard_e2s_v5, standard_d32plds_v5, standard_e64-16as_v4, standard_d16lds_v5, standard_e8-4as_v5, standard_d16ds_v6, standard_d64lds_v6, standard_e4a_v4, standard_e32-16s_v3,	-

Tabela 6 - Agrupamento de instâncias de Computação Acelerada (continuação)

AWS	Azure	GCP
-	standard_ec16ads_cc_v5, standard_d16ls_v5, standard_dc16as_v5, standard_dc2as_v5, standard_e32-16ads_v5, standard_l48s_v3, standard_ds15i_v2, standard_e16s_v5, standard_m96ds_1_v3, standard_h8m, standard_nc16ads_a10_v4, standard_e20as_v6, standard_ec64ads_cc_v5, standard_dc4ads_cc_v5, standard_d4als_v6, standard_e20ds_v5, standard_d64ls_v6, standard_f32ads_v6, standard_d16lds_v6, standard_b8ls_v2, standard_e64d_v4, standard_e32-8as_v5, standard_nd24rs, standard_e104is_v5, standard_d64plds_v5, standard_d4pls_v5, standard_d14_v2, standard_nv72ads_a10_v5, standard_e128-64ds_v6, standard_dc48ds_v3, standard_d64als_v6, standard_m128ms_v2, standard_ds14-8_v2, standard_f8s, standard_b8ps_v2, standard_e4ds_v4, standard_d4as_v5_promo, standard_e48ds_v6, standard_e48bds_v5, standard_d96s_v5, standard_f8amds_v6, standard_d8d_v4, standard_d2a_v4, standard_e112ibds_v5, standard_e16-4as_v4, standard_m128s_v2, standard_d8ds_v4, standard_g2, standard_nd96asr_a100_v4, standard_e4-2ads_v5, standard_e96as_v4, standard_ec48ads_cc_v5, standard_e4as_v5, standard_b16as_v2, standard_h16, standard_m64ls, standard_dc48as_cc_v5, standard_m64ms_v2, standard_ds11, standard_d48ds_v6, standard_np10s, standard_ec64ads_v5, standard_e20d_v5, standard_d64s_v6, standard_l32s_v2, standard_d2s_v5, standard_e96ads_v6, standard_d96a_v4, standard_m416-104ms_v2, standard_e32pds_v5, standard_d48as_v4, standard_e64-16ds_v4, standard_e48s_v4, standard_e16ds_v5, standard_e128-32s_v6, standard_f16as_v6, standard_dc2s_v3, standard_nc24s_v3, standard_d16ds_v5, standard_m208s_v2, standard_e32-16ds_v4, standard_ec48as_cc_v5, standard_d2as_v4, standard_e16_v5, standard_nc24r, standard_h16r, standard_e16-8s_v4, standard_d8s_v3, standard_d48ads_v6, standard_e32-16ds_v6, standard_e4bds_v5, standard_nv12s_v3, standard_e16-4ds_v4, standard_gs5-8, standard_m176ds_4_v3, standard_a5, standard_d64as_v6, standard_nc64as_t4_v3, standard_ec20ads_v5, standard_nv24s_v3, standard_hx176-96rs, standard_e32ads_v5, standard_nv6s_v2, standard_d48ds_v4, standard_dc8as_cc_v5, standard_l80s_v2, standard_a4m_v2, standard_e128-64s_v6, standard_b4ps_v2, standard_d96ds_v5, standard_hb60-15rs, standard_ec32es_v5, standard_e32as_v5, standard_d2s_v6, standard_d64ps_v5, standard_b4as_v2, standard_d2lds_v6, standard_m96ds_2_v3, standard_e8ads_v6, standard_e16-4s_v4, standard_e16-4ds_v6, standard_d32as_v5, standard_hx176-48rs, standard_hb120-16rs_v2, standard_ec64as_cc_v5, standard_d1_v2, standard_b4als_v2, standard_dc64ads_cc_v5, standard_e64i_v4_special, standard_d4as_v4, standard_m64, standard_hc44rs, standard_f32ams_v6, standard_e64-32s_v3, standard_e96-24s_v6, standard_f72s_v2, standard_e96s_v5, standard_d96alds_v6, standard_e16_v3, standard_a4_v2, standard_dc2ds_v3, standard_d4ps_v5, standard_e32-8s_v3,	-

Tabela 6 - Agrupamento de instâncias de Computação Acelerada (continuação)

AWS	Azure	GCP
-	standard_e20s_v3, standard_f16, standard_f2alids_v6, standard_d2alids_v6, standard_e2_v4, standard_l4s, standard_a8m_v2, standard_f72iams_v6, standard_a1_v2, standard_dc48eds_v5, standard_d64s_v3, standard_e96-24as_v4, standard_d48s_v3, standard_d8pls_v5, standard_e20_v3, standard_g4, standard_e4as_v4, standard_e64-16s_v6, standard_d32as_v5_promo, standard_e96a_v4, standard_e8-2ds_v6, standard_e8-2ds_v5, standard_ec96ads_v5, standard_ec8ads_v5, standard_e32s_v5, standard_f8s_v2, standard_d32pls_v5, standard_a2_v2, standard_d96_v5, standard_d16ps_v5, standard_d2as_v5_promo, standard_f2s, standard_e32d_v5, standard_e128s_v6, standard_e4bs_v5, standard_d4ls_v5, standard_m64dms_v2, standard_fx4mds, standard_e16ds_v6, standard_e8pds_v5, standard_e64-32s_v6, standard_e20d_v4, standard_dc4as_v5, standard_e8_v3, standard_f16ads_v6	-

Fonte: Elaborada pelo autor (2025)

Tabela 10 – Agrupamento de instâncias otimizadas para armazenamento

AWS	Azure	GCP
d3.4xlarge, i4i.2xlarge, im4gn.4xlarge, is4gen.4xlarge, i3.metal, h1.8xlarge, im4gn.large, i4i.xlarge, i4g.8xlarge, d2.2xlarge, d3en.12xlarge, i4g.large, is4gen.medium, i3en.24xlarge, i3en.3xlarge, d3.xlarge, is4gen.8xlarge, i3en.12xlarge, i4i.24xlarge, is4gen.xlarge, d3en.2xlarge, is4gen.large, d2.8xlarge, i2.4xlarge, d2.4xlarge, i4g.2xlarge, h1.4xlarge, i4i.16xlarge, i3.4xlarge, d3en.8xlarge, i4i.metal, i3.16xlarge, i3.2xlarge, d3en.xlarge, im4gn.8xlarge, i4i.32xlarge, d3.8xlarge, im4gn.2xlarge, d2.xlarge, i2.xlarge, i3en.large, i4i.large, i4i.12xlarge, i3en.6xlarge, is4gen.2xlarge, i2.2xlarge, im4gn.16xlarge, d3en.6xlarge, im4gn.xlarge, i2.8xlarge, d3en.4xlarge, i4g.4xlarge, h1.2xlarge, i3.xlarge, i4i.4xlarge, i4g.xlarge, i4i.8xlarge, i3.8xlarge, h1.16xlarge, i3.large, d3.2xlarge, i4g.16xlarge, i3en.xlarge, i3en.2xlarge, i3en.metal	standard_L8s_v2, standard_L16s_v2, standard_L32s_v2, standard_L48s_v2, standard_L64s_v2, standard_L80s_v26, standard_L8as_v3, standard_L16as_v3, standard_L32as_v3, standard_L48as_v3, standard_L64as_v3, standard_L80as_v3	z3-highmem-88, z3-highmem-176

Fonte: Elaborada pelo autor (2025)

APÊNDICE B – TABELAS DE CORRELAÇÃO DE PEARSON ENTRE PROVEDORES, REGIÕES E TIPOS DE INSTÂNCIA

Tabela 11 – Correlação de Pearson entre provedores por região - Computação Acelerada

Família	Região	Par de Provedores	Correlação Spot	Correlação Sob Demanda
Computação Acelerada	ap_south_1	aws-azure	0.2641	-0.0804
Computação Acelerada	eu_north	aws-azure	-0.8714	0.2549
Computação Acelerada	eu_west_3	aws-azure	0.8735	0.0462
Computação Acelerada	eu_west_2	aws-azure	-0.0127	0.2566
Computação Acelerada	eu_west_2	aws-gcp	-0.5969	-0.0418
Computação Acelerada	eu_west_2	azure-gcp	0.1173	-0.035
Computação Acelerada	eu_west_1	aws-azure	0.2926	0.8567
Computação Acelerada	ap_northeast_2	aws-azure	0.8838	-0.9283
Computação Acelerada	ap_northeast_2	aws-gcp	-0.6263	
Computação Acelerada	ap_northeast_2	azure-gcp	-0.8007	
Computação Acelerada	ap_northeast_1	aws-azure	-0.5115	0.6296
Computação Acelerada	ap_northeast_1	aws-gcp	0.1197	-0.4795
Computação Acelerada	ap_northeast_1	azure-gcp	-0.0149	0.0456
Computação Acelerada	ca_central	aws-azure	0.4202	0.0189
Computação Acelerada	sa_east	aws-azure	0.5532	0.535
Computação Acelerada	ap_southeast_1	aws-azure	0.8544	0.4699
Computação Acelerada	ap_southeast_1	aws-gcp	-0.6886	-0.2401
Computação Acelerada	ap_southeast_1	azure-gcp	-0.8944	-0.1144
Computação Acelerada	ap_southeast_2	aws-azure	0.8752	0.6545
Computação Acelerada	eu_central	aws-azure	0.2878	0.9456
Computação Acelerada	eu_central	aws-gcp	-0.4476	-0.3095
Computação Acelerada	eu_central	azure-gcp	-0.7378	-0.2963
Computação Acelerada	us_east_1	aws-azure	-0.2122	0.6501
Computação Acelerada	us_east_1	aws-gcp	0.941	0.115
Computação Acelerada	us_east_1	azure-gcp	-0.2941	0.064
Computação Acelerada	us_east_2	aws-azure	-0.5542	0.7454
Computação Acelerada	us_east_2	aws-gcp	0.337	0.0363
Computação Acelerada	us_east_2	azure-gcp	0.0625	0.0459
Computação Acelerada	us_west_1	aws-azure	0.2532	0.9449
Computação Acelerada	us_west_1	aws-gcp	-0.5063	
Computação Acelerada	us_west_1	azure-gcp	-0.8539	
Computação Acelerada	us_west_2	aws-azure	-0.7215	0.5025
Computação Acelerada	us_west_2	aws-gcp	-0.1083	-0.7885
Computação Acelerada	us_west_2	azure-gcp	0.314	-0.4307

Fonte: Elaborada pelo autor (2025)

Tabela 12 – Correlação de Pearson entre provedores por região - Propósito Geral

Família	Região	Par de Provedores	Correlação Spot	Correlação Sob Demanda
Propósito Geral	ap_south_1	aws-azure	-0.1233	0.3246
Propósito Geral	ap_south_1	aws-gcp	-0.2326	0.2867
Propósito Geral	ap_south_1	azure-gcp	0.7996	0.9741
Propósito Geral	eu_north	aws-gcp	-0.0258	0.8365
Propósito Geral	eu_west_3	aws-azure	-0.6464	-0.3595
Propósito Geral	eu_west_3	aws-gcp	-0.2119	0.4262
Propósito Geral	eu_west_3	azure-gcp	0.4089	-0.2685
Propósito Geral	eu_west_2	aws-azure	-0.7004	-0.7002
Propósito Geral	eu_west_2	aws-gcp	-0.6017	0.5073
Propósito Geral	eu_west_2	azure-gcp	0.5561	-0.7184
Propósito Geral	eu_west_1	aws-azure	0.4776	0.0149
Propósito Geral	eu_west_1	aws-gcp	-0.6912	0.6692
Propósito Geral	eu_west_1	azure-gcp	-0.0089	-0.1104
Propósito Geral	ap_northeast_3	aws-azure	-0.7691	0.1628
Propósito Geral	ap_northeast_3	aws-gcp	0.0362	-0.9995
Propósito Geral	ap_northeast_3	azure-gcp	-0.0439	-0.1639
Propósito Geral	ap_northeast_2	aws-azure	-0.2975	0.0727
Propósito Geral	ap_northeast_2	aws-gcp	-0.5898	0.4069
Propósito Geral	ap_northeast_2	azure-gcp	0.5802	-0.0276
Propósito Geral	ap_northeast_1	aws-azure	0.8902	0.3761
Propósito Geral	ap_northeast_1	aws-gcp	-0.4967	0.0902
Propósito Geral	ap_northeast_1	azure-gcp	-0.5302	0.0094
Propósito Geral	ca_central	aws-azure	0.2381	0.4716
Propósito Geral	ca_central	aws-gcp	0.2322	0.4628
Propósito Geral	sa_east	aws-azure	-0.2129	0.3023
Propósito Geral	sa_east	aws-gcp	-0.8109	0.4197
Propósito Geral	sa_east	azure-gcp	0.6515	0.9698
Propósito Geral	ap_southeast_2	aws-azure	0.5224	0.1889
Propósito Geral	ap_southeast_2	aws-gcp	0.2898	0.6082
Propósito Geral	ap_southeast_2	azure-gcp	0.697	-0.0384
Propósito Geral	eu_central	aws-azure	-0.8702	0.0633
Propósito Geral	eu_central	aws-gcp	-0.655	0.8483
Propósito Geral	eu_central	azure-gcp	0.9174	-0.1007
Propósito Geral	us_east_2	aws-azure	-0.1268	0.6323
Propósito Geral	us_east_2	aws-gcp	-0.3295	0.6176
Propósito Geral	us_east_2	azure-gcp	0.788	0.9767
Propósito Geral	us_west_1	aws-azure	0.8127	0.2786
Propósito Geral	us_west_1	aws-gcp	-0.7412	0.4566
Propósito Geral	us_west_1	azure-gcp	-0.4429	0.4984
Propósito Geral	us_west_2	aws-azure	-0.1052	0.6638
Propósito Geral	us_west_2	aws-gcp	-0.6873	0.7589
Propósito Geral	us_west_2	azure-gcp	0.7195	0.9662
Propósito Geral	ap_east	azure-gcp	0.634	-0.0699
Propósito Geral	ap_south_2	azure-gcp	0.2755	-0.0204

Fonte: Elaborada pelo autor (2025)

Tabela 13 – Correlação de Pearson entre provedores por região - Otimizada para Computação

Família	Região	Par de Provedores	Correlação Spot	Correlação Sob Demanda
Otimizado para Computação	ap_south_1	aws-gcp	0.6808	0.2441
Otimizado para Computação	eu_west_2	aws-gcp	0.526	-0.7155
Otimizado para Computação	eu_west_1	aws-gcp	-0.956	0.8897
Otimizado para Computação	ap_northeast_3	aws-gcp	-0.4419	0.0
Otimizado para Computação	ap_northeast_2	aws-gcp		0.0526
Otimizado para Computação	ap_northeast_1	aws-gcp	-0.562	0.5102
Otimizado para Computação	ca_central	aws-gcp	-0.5595	0.2454
Otimizado para Computação	sa_east	aws-gcp	0.1736	-0.0665
Otimizado para Computação	ap_southeast_1	aws-gcp	-0.9614	0.4223
Otimizado para Computação	eu_central	aws-gcp	-0.6487	0.5794
Otimizado para Computação	us_east_1	aws-gcp	-0.965	0.8882
Otimizado para Computação	us_east_2	aws-gcp	-0.5553	0.5334
Otimizado para Computação	us_west_1	aws-gcp	-0.2506	0.5071
Otimizado para Computação	us_west_2	aws-gcp	-0.9287	0.8864

Fonte: Elaborada pelo autor (2025)

Tabela 14 – Correlação de Pearson entre provedores por região - Otimizada para Memória

Família	Região	Par de Provedores	Correlação Spot	Correlação Sob Demanda
Otimizado para Memória	ap_south_1	aws-gcp	-0.1398	0.4702
Otimizado para Memória	eu_west_3	aws-gcp	0.9498	-0.5823
Otimizado para Memória	eu_west_2	aws-gcp	0.0592	-0.0515
Otimizado para Memória	eu_west_1	aws-gcp	-0.1044	-0.6572
Otimizado para Memória	ap_northeast_3	aws-gcp	0.3882	
Otimizado para Memória	ap_northeast_2	aws-gcp	0.008	-0.0851
Otimizado para Memória	ap_northeast_1	aws-gcp	0.9376	0.5159
Otimizado para Memória	ca_central	aws-gcp	0.5586	0.3957
Otimizado para Memória	sa_east	aws-gcp	0.3822	-0.5212
Otimizado para Memória	ap_southeast_1	aws-gcp	0.5969	0.669
Otimizado para Memória	ap_southeast_2	aws-gcp	-0.0896	0.2055
Otimizado para Memória	eu_central	aws-gcp	0.7481	0.602
Otimizado para Memória	us_east_1	aws-gcp	-0.3835	-0.371
Otimizado para Memória	us_west_1	aws-gcp	-0.3926	0.08
Otimizado para Memória	us_west_2	aws-gcp	-0.7205	-0.7573

Fonte: Elaborada pelo autor (2025)