



DISSERTAÇÃO DE MESTRADO PROFISSIONAL

**PRIORIZAÇÃO INTELIGENTE DE VULNERABILIDADES:
UMA ABORDAGEM BASEADA EM PERFIS DE RISCO
DESCOBERTOS COM MACHINE LEARNING**

Madson Luiz Magno da Silva

Prof. Dr. João Souza Neto

Programa de Pós-Graduação Profissional em Engenharia
Elétrica

DEPARTAMENTO DE ENGENHARIA ELÉTRICA

UNIVERSIDADE DE BRASÍLIA
FACULDADE DE TECNOLOGIA
DEPARTAMENTO DE ENGENHARIA ELÉTRICA

**PRIORIZAÇÃO INTELIGENTE DE VULNERABILIDADES: UMA ABORDAGEM
BASEADA EM PERFIS DE RISCO DESCOBERTOS COM MACHINE LEARNING**

*INTELLIGENT VULNERABILITY PRIORITIZATION: A RISK PROFILE-BASED APPROACH
DISCOVERED WITH MACHINE LEARNING*

Madson Luiz Magno da Silva

Orientador: Prof. Dr. João Souza Neto, FT/UnB

PUBLICAÇÃO: PPEE.MP.119

BRASÍLIA-DF

2026

UNIVERSIDADE DE BRASÍLIA
Faculdade de Tecnologia

DISSERTAÇÃO DE MESTRADO PROFISSIONAL
**PRIORIZAÇÃO INTELIGENTE DE VULNERABILIDADES:
UMA ABORDAGEM BASEADA EM PERFIS DE RISCO
DESCOBERTOS COM MACHINE LEARNING**

Madson Luiz Magno da Silva
Prof. Dr. João Souza Neto

*Dissertação de Mestrado Profissional submetida ao Departamento de
Engenharia*

*Elétrica como requisito parcial para obtenção
do grau de Mestre em Engenharia Elétrica*

Banca Examinadora

Prof. Dr. João Souza Neto, FT/UnB
Presidente

Prof. Dr. Joao Jose Costa Gondim, FT/UnB
Examinador Interno

Prof. Dr. João Luiz Pereira Marciano, Câmara
dos Deputados (CD)
Examinador Externo

Prof. Dr. Geraldo Pereira Rocha Filho
Suplente

Ficha catalográfica elaborada automaticamente,
com os dados fornecidos pelo(a) autor(a)

MS586p

Magno da Silva, Madson Luiz
PRIORIZAÇÃO INTELIGENTE DE VULNERABILIDADES: UMA
ABORDAGEM BASEADA EM PERFIS DE RISCO DESCOBERTOS COM MACHINE
LEARNING / Madson Luiz Magno da Silva; orientador João Souza
Neto. Brasília, 2026.
69 p.

Dissertação(Mestrado em Engenharia Elétrica) Universidade
de Brasília, 2026.

1. Priorização de Vulnerabilidades. 2. Machine Learning.
3. Inteligência Artificial Generativa. 4. Inteligência
Artificial Generativa. 5. Dados Sintéticos. I. Souza Neto,
João, orient. II. Título.

“Como, então, viveremos?”

Os Arrais

Agradecimentos

Agradeço primeiramente a Deus pela sabedoria e saúde durante toda esta jornada.

À minha família, pelo apoio incondicional, amor e paciência nos momentos de ausência necessários para a conclusão desta pesquisa.

Ao meu orientador, Prof. Dr. João Souza Neto, pela orientação precisa, disponibilidade e por compartilhar seu vasto conhecimento acadêmico e profissional, essenciais para a concretização deste trabalho.

Aos colegas do Programa de Pós-Graduação, pelas valiosas discussões e trocas de experiências.

A todos que, direta ou indiretamente, contribuíram para a realização desta dissertação.

RESUMO

A gestão de vulnerabilidades em infraestruturas críticas governamentais enfrenta o desafio de priorizar de forma eficiente as ameaças. O uso isolado de métricas estáticas, como o *Common Vulnerability Scoring System* (CVSS), gera fadiga de alertas devido ao alto volume de classificações de criticidade máxima que não são exploradas na prática. Esta dissertação propõe um modelo bidimensional de priorização de riscos que integra a Probabilidade baseada em Exploração (EPSS) e a Nota de Criticidade de Sistema (NCS), utilizando aprendizado de máquina não supervisionado. A seleção algorítmica foi fundamentada em estudo comparativo prévio, aplicando-se o *K-Means* para segmentar as vulnerabilidades em perfis de risco e o DBSCAN para a detecção de anomalias espaciais. Para atestar a escalabilidade e a robustez do modelo sem comprometer a Segurança Operacional (OPSEC) do Estado, a pesquisa incluiu testes de estresse em larga escala. Empregou-se Inteligência Artificial Generativa, por meio de redes adversárias (CTGAN), para criar uma base sintética com 101.000 registros. O procedimento preservou a fidelidade estatística do ambiente real, validada rigorosamente pelo teste de *Kolmogorov-Smirnov*. Os resultados em ambiente controlado indicaram uma redução de 69% no esforço de remediação imediata em comparação ao método tradicional. Nos testes de estresse, o *K-Means* manteve a eficiência computacional no particionamento de grandes volumes de dados, enquanto o DBSCAN identificou com precisão os cenários extremos injetados artificialmente. Conclui-se que o modelo proposto oferece uma solução escalável, reproduzível e segura para o direcionamento de recursos em segurança cibernética, figurando como um instrumento de antecipação aplicável ao Subsistema de Inteligência de Segurança Pública (SISP).

Palavras-chave: Priorização de Vulnerabilidades, Machine Learning, Inteligência Artificial Generativa, Segurança Pública, Dados Sintéticos.

ABSTRACT

Vulnerability management in government critical infrastructures faces the challenge of efficiently prioritizing threats. The isolated use of static metrics, such as the *Common Vulnerability Scoring System* (CVSS), generates alert fatigue due to the high volume of maximum criticality classifications that are not exploited in practice. This dissertation proposes a two-dimensional risk prioritization model that integrates the *Exploit Prediction Scoring System* (EPSS) and the System Criticality Score (NCS) and uses unsupervised machine learning. The algorithmic selection was based on a previous comparative study that applied *K-Means* to segment vulnerabilities into risk profiles and DBSCAN for spatial anomaly detection. To attest the model's scalability and robustness without compromising the State's Operational Security (OPSEC), the research included large-scale stress tests. Generative Artificial Intelligence, using an adversarial network (CTGAN), was employed to create a synthetic database of 101,000 records. The procedure preserved the statistical fidelity of the real environment, rigorously validated by the *Kolmogorov-Smirnov* test. Results in a controlled environment indicated a 69% reduction in immediate remediation effort compared to the traditional method. In the stress tests, *K-Means* maintained computational efficiency in partitioning large volumes of data, while DBSCAN accurately identified the artificially injected extreme scenarios. We conclude that the proposed model offers a scalable, reproducible, and secure solution for directing cybersecurity resources and serves as an anticipation instrument applicable to the Public Security Intelligence Subsystem (SISP).

Keywords: Vulnerability Prioritization, Machine Learning, Generative Artificial Intelligence, Public Security, Synthetic Data.

SUMÁRIO

1	INTRODUÇÃO	1
1.1	O PROBLEMA DA GESTÃO DE VULNERABILIDADES.....	1
1.2	A LACUNA: A INSUFICIÊNCIA DO CVSS	2
1.3	PROPOSTA: UMA ABORDAGEM BASEADA EM RISCO REAL	2
1.4	OBJETIVOS DA PESQUISA	3
1.5	ESTRUTURA DO TRABALHO.....	4
2	REFERENCIAL TEÓRICO	5
2.1	GESTÃO DE RISCOS: O CONCEITO FUNDAMENTAL	5
2.2	RISCO EM SEGURANÇA DA INFORMAÇÃO (SI).....	5
2.2.1	AMEAÇA, VULNERABILIDADE E IMPACTO	6
2.3	GESTÃO DE VULNERABILIDADES	6
2.4	TAXONOMIA DE VULNERABILIDADES: CWE E CVE	7
2.4.1	CWE: A FRAQUEZA (A CAUSA RAIZ)	7
2.4.2	CVE: A VULNERABILIDADE (A INSTÂNCIA)	7
2.5	O MÉTODO TRADICIONAL DE PRIORIZAÇÃO: CVSS.....	8
2.6	A LACUNA DE PRIORIZAÇÃO: A CRÍTICA AO CVSS	9
2.7	SOLUÇÕES MODERNAS PARA A PRIORIZAÇÃO	10
2.7.1	A DIMENSÃO DA PROBABILIDADE: EPSS.....	10
2.7.2	A DIMENSÃO DA DECISÃO: SSVC.....	11
2.8	A DIMENSÃO DO IMPACTO: O CONTEXTO DE NEGÓCIO (PPSI/NCS) ...	11
2.9	A PRIORIZAÇÃO DE VULNERABILIDADES COMO ESTRATÉGIA DE INTE- LIGÊNCIA DE SEGURANÇA PÚBLICA (SISP)	12
2.10	A FERRAMENTA DE ANÁLISE: MACHINE LEARNING NÃO SUPERVISIO- NADO.....	13
2.11	ALGORITMOS DE CLUSTERIZAÇÃO PARA DESCOBERTA DE PERFIS	14
2.11.1	K-MEANS: AGRUPAMENTO BASEADO EM CENTROIDE	14
2.11.2	DBSCAN: AGRUPAMENTO BASEADO EM DENSIDADE E DETECÇÃO DE ANOMALIAS.....	15
2.12	O USO DE DADOS SINTÉTICOS E A SEGURANÇA OPERACIONAL.....	15
3	TRABALHOS CORRELATOS	17
3.1	MACHINE LEARNING PARA DETECÇÃO DE ANOMALIAS EM SEGURANÇA	17
3.2	APLICAÇÕES DE CLUSTERING EM CIBERSEGURANÇA	18
3.3	ABORDAGENS HÍBRIDAS DE PRIORIZAÇÃO DE RISCO	18

3.4	CONTEXTO DA PESQUISA NO ÂMBITO LOCAL (PPEE/UNB)	19
3.5	GERAÇÃO DE DADOS SINTÉTICOS PARA AVALIAÇÃO DE MODELOS	20
3.6	POSICIONAMENTO DA DISSERTAÇÃO E SÍNTESE DOS TRABALHOS CORRELATOS	21
4	METODOLOGIA DA PESQUISA	22
4.1	FASE 1: DESCOBERTA DE ATIVOS (RECONHECIMENTO)	23
4.2	FASE 2: FILTRAGEM DE ATIVOS	23
4.3	FASE 3: JUSTIFICATIVA DA SIMULAÇÃO DE VULNERABILIDADES (CWE vs. CVE)	24
4.4	FASE 4: ASSOCIAÇÃO DE VULNERABILIDADES (CVEs)	25
4.5	FASE 5: BALANCEAMENTO DO DATASET	25
4.6	FASE 6: SIMULAÇÃO DE IMPACTO (NCS)	25
4.7	FASE 7: ENRIQUECIMENTO DE PROBABILIDADE (EPSS)	26
4.8	FASE 8: PRÉ-PROCESSAMENTO (PADRONIZAÇÃO)	26
4.9	FASE 9: GERAÇÃO DE DADOS SINTÉTICOS EM LARGA ESCALA	27
4.10	FASE 10: INJEÇÃO DE ANOMALIAS (EDGE CASES)	27
5	RESULTADOS E DISCUSSÃO	29
5.1	ANÁLISE EXPLORATÓRIA DE DADOS (EDA)	29
5.1.1	ANÁLISE DE CORRELAÇÃO DAS MÉTRICAS DE RISCO (EPSS vs. NCS)	29
5.1.2	VISUALIZAÇÃO DA MATRIZ DE RISCO (DISPERSÃO)	30
5.2	SELEÇÃO DO NÚMERO DE CLUSTERS (K)	31
5.2.1	MÉTODO DO COTOVELO (ELBOW METHOD)	32
5.2.2	ANÁLISE ESTRUTURAL (DENDROGRAMA)	33
5.3	DESCOBERTA E CARACTERIZAÇÃO DOS PERFIS DE RISCO	34
5.3.1	PERFIL 1: RISCO CRÍTICO (CLUSTER 0)	35
5.3.2	PERFIL 2: AMEAÇA DE VOLUME (CLUSTER 1)	36
5.3.3	PERFIL 3: RISCO MÉDIO (CLUSTER 2)	36
5.3.4	PERFIL 4: GIGANTE ADORMECIDO (CLUSTER 3)	37
5.3.5	PERFIL 5: RISCO MÍNIMO (CLUSTER 4)	37
5.4	VALIDAÇÃO COMPLEMENTAR E DETECÇÃO DE ANOMALIAS (DBSCAN)	38
5.5	ANÁLISE DE IMPACTO QUANTITATIVO: A REDUÇÃO DA “FADIGA DE ALERTAS”	39
5.6	VISUALIZAÇÃO CONSOLIDADA DA MATRIZ DE RISCO	41
5.7	VALIDAÇÃO EM LARGA ESCALA E TESTES DE ESTRESSE	42
5.7.1	VALIDAÇÃO ESTATÍSTICA DOS DADOS SINTÉTICOS	42
5.7.2	DESEMPENHO ALGORÍTMICO NO CENÁRIO DE BIG DATA E ANOMALIAS	43
6	CONCLUSÃO E TRABALHOS FUTUROS	46

6.1	CONTRIBUIÇÕES DA PESQUISA	46
6.2	LIMITAÇÕES DA PESQUISA	47
6.3	TRABALHOS FUTUROS	48
REFERÊNCIAS BIBLIOGRÁFICAS.....		50

LISTA DE FIGURAS

4.1	Fluxograma do pipeline de engenharia de dados e de modelagem.	23
5.1	Matriz de correlação de Pearson entre EPSS e NCS (PPSI) normalizadas.	30
5.2	Matriz de Risco: Dispersão entre as EPSS e as NCS Normalizadas.	31
5.3	Análise do Método do Cotovelo (WCSS) para determinar o k ideal.	32
5.4	Análise estrutural dos agrupamentos via Dendrograma de Clusterização Hierárquica.	33
5.5	Mapa de calor (Heatmap) dos valores médios dos centroides (Perfis de Risco).	35
5.6	Comparativo de Priorização: Modelo Tradicional (CVSS) vs. Modelo Proposto (Perfis de Risco K-Means).	40
5.7	Matriz de Risco Consolidada: Perfis de Risco (K-Means, por cor), Centroides (círculos amarelos) e Anomalias (DBSCAN, marcadores 'X').	41
5.8	Preservação multivariada das matrizes de correlação no ambiente gerado por IA.	43
5.9	Comparação da densidade de probabilidade entre o ambiente real e o sintético.	43
5.10	Comparativo de mitigação da fadiga de alertas sob estresse massivo (101.000 instâncias).	44
5.11	Matriz de Risco consolidada no cenário de Big Data governamental e de anomalias extremas.	45

LISTA DE TABELAS

2.1	Resumo dos eixos de validação com dados sintéticos em cibersegurança.	16
3.1	Síntese dos trabalhos correlatos sobre dados sintéticos em cibersegurança.	21
4.1	Composição da base de dados para testes de estresse em larga escala.	28
5.1	Análise dos centroides dos 5 perfis de risco identificados pelo K-Means.	34
5.2	Exemplos de Anomalias (Outliers) Estratégicas Identificadas pelo DBSCAN.	38
5.3	Comparativo estatístico e métricas de fidelidade (Mundo Real vs. Sintético).	42

1 INTRODUÇÃO

A sociedade contemporânea vivencia uma profunda Transformação Digital, um fenômeno que transcende a mera adoção de tecnologia e reconfigura fundamentalmente as operações de organizações públicas e privadas em escala global (1). A digitalização de serviços, impulsionada pela busca por eficiência, alcance e inovação, estabeleceu o *software* como o alicerce central sobre o qual se constroem as interações econômicas, sociais e governamentais. Esta crescente dependência de sistemas de informação e infraestruturas tecnológicas, embora traga benefícios inegáveis, introduz simultaneamente uma complexidade sem precedentes.

Esta digitalização massiva resulta em uma expansão exponencial da superfície de ataque cibernético (2). Cada novo sistema conectado, cada aplicação *web* e cada linha de código implementada representam um novo vetor potencial de ameaças. Neste cenário hiperconectado, a Segurança da Informação (SI) deixa de ser uma função de suporte técnico para se tornar um pilar estratégico e indispensável à continuidade e à soberania de qualquer entidade (3, 4, 5).

O crescimento do mundo cibernético é indissociável do aumento das ameaças. Ataques de *ransomware*, exfiltração de dados e interrupção de serviços críticos tornaram-se ocorrências diárias, gerando prejuízos financeiros e reputacionais significativos (6, 7). O cerne deste problema reside na existência de falhas em *softwares*; falhas que, quando exploráveis, são catalogadas como vulnerabilidades (8).

1.1 O PROBLEMA DA GESTÃO DE VULNERABILIDADES

Diante deste panorama, a Gestão de Vulnerabilidades (GV) emerge como um processo crítico da Segurança da Informação. A GV compreende o ciclo contínuo de identificação, classificação, priorização, remediação e mitigação de falhas de segurança (9). Contudo, a velocidade com que novas vulnerabilidades são descobertas, ultrapassando dezenas de milhares de registros anuais no catálogo *Common Vulnerabilities and Exposures* (CVE), tornou impraticável o processo de correção manual ou reativo (10).

O desafio prático enfrentado pelas equipes de segurança não é a detecção de vulnerabilidades, pois ferramentas automatizadas, como o OpenVAS ou o OWASP ZAP, são eficazes para identificar fraquezas (CWEs) (11, 12). O verdadeiro desafio é a priorização. As equipes são inundadas por um volume de alertas que excede em muito sua capacidade de remediação. Este fenômeno, conhecido na literatura como "fadiga de alertas" (*alert fatigue*) (13), força as organizações a tomar

decisões difíceis sobre o que corrigir primeiro.

1.2 A LACUNA: A INSUFICIÊNCIA DO CVSS

Historicamente, a principal ferramenta para esta priorização é o *Common Vulnerability Scoring System* (CVSS) (14). O CVSS fornece uma pontuação de 0 a 10 que avalia a severidade técnica de uma vulnerabilidade, analisando métricas como o vetor de ataque, a complexidade e o impacto teórico na confidencialidade, integridade e disponibilidade. Embora essencial para a padronização, o uso exclusivo do CVSS para priorização é amplamente reconhecido como uma lacuna metodológica.

Estudos demonstram que o CVSS é um mau preditor da exploração real de vulnerabilidades (15). O sistema foca na gravidade teórica da falha, ignorando o contexto operacional do ativo afetado e a probabilidade real de um ataque ocorrer. O resultado é uma distorção da priorização: vulnerabilidades com pontuação "Crítica"(9.0-10.0) jamais podem ser exploradas, enquanto falhas "Médias" podem estar sendo ativamente usadas em campanhas de ataque (16, 17, 18, 19). Em suma, o CVSS responde se uma vulnerabilidade é perigosa, mas não se é provável ou impactante para um negócio específico.

1.3 PROPOSTA: UMA ABORDAGEM BASEADA EM RISCO REAL

Esta dissertação propõe uma abordagem alternativa para a priorização, fundamentada em uma definição clássica de risco, onde o risco cibernético é uma função de duas dimensões principais: a probabilidade de um evento de ameaça e o impacto que esse evento teria na organização (20, 21).

Para operacionalizar este modelo de risco bidimensional, a metodologia proposta combina duas métricas de ponta, substituindo a pontuação única do CVSS por uma análise de portfólio:

1. **Probabilidade de Exploração:** Quantificada pelo *Exploit Prediction Scoring System* (EPSS). O EPSS é um modelo de *Machine Learning* mantido pelo FIRST.org que calcula a probabilidade, em um intervalo de 0% a 100%, de uma vulnerabilidade (identificada pelo seu CVE) ser explorada ativamente nos próximos 30 dias (15).
2. **Impacto no Negócio:** Quantificado pela Nota de Criticidade de Sistema (NCS). Esta métrica é uma derivação do *Framework* de Privacidade e Segurança da Informação (PPSI) do governo brasileiro, que, por sua vez, baseia-se no Acórdão 1.889/2020 do Tribunal de Con-

tas da União (TCU) (22, 23). A NCS avalia a criticidade de um ativo (sistema) com base no seu papel no negócio, considerando critérios como danos financeiros, reputacionais e à missão da organização.

A hipótese central deste trabalho é que, ao tratar vulnerabilidades como pontos de dados em um espaço bidimensional (Probabilidade x Impacto), é possível aplicar técnicas de *Machine Learning* não supervisionadas, especificamente algoritmos de *clustering* como o k-Means e o DBSCAN, para descobrir "perfis de risco" (*risk profiles*) naturais e acionáveis (24, 25, 26, 27). Esta abordagem visa transformar a "fadiga de alertas" em uma estratégia de remediação focada, permitindo que as equipes de segurança concentrem seus recursos limitados nas ameaças que representam o maior risco real e contextualizado para a organização.

A metodologia demonstrou resultados positivos em ambientes controlados, e a seleção dos algoritmos foi fundamentada em estudo comparativo prévio publicado na literatura (19). Contudo, a validação de uma ferramenta analítica voltada para infraestruturas críticas requer a verificação de sua estabilidade em cenários de alto volume de dados. A literatura da área aponta uma escassez de bases de dados públicas que detalhem o impacto de negócios em órgãos do governo. Essa restrição ocorre devido à necessidade de manter a confidencialidade e a segurança operacional das informações, o que dificulta a reprodutibilidade das pesquisas. Para suprir essa lacuna metodológica, o presente trabalho realizou testes de estresse adicionais. O método consistiu na incorporação de registros com alta probabilidade de exploração, provenientes do catálogo CISA KEV, e na geração de dados sintéticos por meio de redes adversárias generativas (28, 29). Essa etapa avaliou a escalabilidade dos algoritmos e o comportamento do modelo diante de casos extremos. O procedimento buscou preservar o rigor analítico sem expor dados governamentais reais.

1.4 OBJETIVOS DA PESQUISA

O objetivo geral desta dissertação é desenvolver e avaliar uma metodologia de priorização inteligente de vulnerabilidades baseada em perfis de risco, utilizando métricas de probabilidade de exploração (EPSS) e impacto de negócio (NCS) analisadas por meio de *Machine Learning* não supervisionado.

Para alcançar este objetivo geral, definem-se os seguintes objetivos específicos:

1. Realizar um levantamento bibliográfico sobre a gestão de riscos, a gestão de vulnerabilidades, as limitações do CVSS e as abordagens modernas de priorização, como o EPSS e o SSVC.

2. Detalhar a metodologia de coleta e tratamento de dados, incluindo a identificação de ativos reais, a associação de CVEs e a simulação da criticidade de ativos (NCS) para criar um *dataset* de pesquisa viável.
3. Aplicar algoritmos de clusterização (k-Means) para segmentar as vulnerabilidades em grupos distintos e aplicar algoritmos de densidade (DBSCAN) para identificar anomalias (*outliers*).
4. Analisar e caracterizar os *clusters* resultantes para definir perfis de risco qualitativos (ex.: "Risco Crítico", "Gigantes Adormecidos") que orientem diferentes estratégias de remediação.
5. Validar quantitativamente a eficácia da metodologia proposta, comparando o volume de alertas gerado pelo modelo de perfis de risco com o gerado pela priorização tradicional baseada apenas no CVSS.
6. Validar a escalabilidade e a resiliência do modelo proposto por meio de testes de estresse, com o uso de catálogos de vulnerabilidades ativamente exploradas e a geração de dados sintéticos por redes adversárias generativas, verificando o comportamento dos algoritmos em cenários com alto volume de dados e casos anômalos.

1.5 ESTRUTURA DO TRABALHO

Este trabalho está organizado em seis capítulos. O Capítulo 1, esta introdução, apresenta o problema, a motivação e os objetivos da pesquisa. O Capítulo 2 estabelece o referencial teórico, detalhando os conceitos fundamentais de risco e segurança, bem como as métricas de priorização. O Capítulo 3 apresenta os trabalhos correlatos e situa esta pesquisa no contexto de outras soluções de *Machine Learning* para cibersegurança. O Capítulo 4 detalha a metodologia de pesquisa e o *pipeline* de engenharia de dados. O Capítulo 5 apresenta e discute os resultados obtidos, incluindo a caracterização dos perfis de risco e a validação do impacto do modelo. A estrutura da dissertação inclui um espaço para a validação complementar da metodologia. O documento detalha os testes de estresse aplicados ao modelo, em ambientes sintéticos e de larga escala. Por fim, o Capítulo 6 conclui o estudo, resumindo as contribuições, apontando as limitações e sugerindo trabalhos futuros.

2 REFERENCIAL TEÓRICO

Este capítulo estabelece os fundamentos conceituais necessários à compreensão da metodologia proposta. Inicia-se com as definições formais de risco, conforme as normas internacionais, afunilando-se para o domínio específico do risco em segurança da informação. Em seguida, detalha os componentes do risco — Ameaça, Vulnerabilidade e Impacto —, apresenta o processo cíclico de Gestão de Vulnerabilidades e, por fim, estabelece os conceitos fundamentais (CWE e CVE) que alicerçam a coleta de dados e a análise desta dissertação.

2.1 GESTÃO DE RISCOS: O CONCEITO FUNDAMENTAL

No contexto organizacional, a gestão de riscos é um processo essencial para a tomada de decisão estratégica. A norma ABNT NBR ISO 31000:2018 define risco como o "efeito da incerteza nos objetivos"(30). Esta definição é crucial, pois desvincula o risco de uma conotação puramente negativa; a incerteza pode gerar tanto ameaças quanto oportunidades. A gestão de riscos, portanto, não visa eliminar todo e qualquer risco, mas sim gerenciar essas incertezas de forma estruturada para aumentar a probabilidade de que os objetivos organizacionais sejam alcançados.

O processo de gestão de riscos, conforme detalhado na ISO 31000 e em *frameworks* consolidados, é iterativo e compreende as etapas de identificação, análise, avaliação e tratamento (mitigação, aceitação, transferência ou evitação) dos riscos. Esta dissertação foca em uma subárea crítica deste processo: o risco cibernético.

2.2 RISCO EM SEGURANÇA DA INFORMAÇÃO (SI)

A Segurança da Informação (SI) é a disciplina dedicada a proteger os ativos de informação contra um espectro de ameaças, visando assegurar a continuidade dos negócios, minimizar danos e maximizar o retorno sobre os investimentos (5). A SI fundamenta-se em três pilares clássicos, conhecidos como a Tríade CIA: Confidencialidade (garantia de que a informação é acessível apenas por entidades autorizadas), Integridade (salvaguarda da exatidão e completude da informação) e Disponibilidade (garantia de que a informação esteja acessível quando solicitada) (4).

No domínio da SI, o risco é definido de forma mais específica. Publicações como a ISO/IEC

27005:2022 (20) e o NIST SP 800-30 (21) convergem para um modelo em que o risco é uma função de três componentes principais: Ameaça, Vulnerabilidade e Impacto.

2.2.1 Ameaça, Vulnerabilidade e Impacto

A compreensão da interação entre estes três componentes é a base para qualquer avaliação de risco cibernético. Uma ameaça é qualquer circunstância ou evento com potencial para causar dano a um ativo de informação. As ameaças podem ser naturais (ex.: desastres) ou humanas, sendo estas últimas divididas em acidentais (ex.: erro operacional) ou intencionais (ex.: um ator malicioso, ou *hacker*).

Uma vulnerabilidade é uma fraqueza em um ativo ou em seus controles de segurança que pode ser explorada por uma ou mais ameaças (9). É crucial notar que uma vulnerabilidade, por si só, não representa risco. Ela é uma condição latente, uma falha em um *software* ou uma configuração inadequada, que só se torna um problema na presença de uma ameaça que saiba explorá-la.

O impacto refere-se ao resultado adverso ou ao dano causado à organização caso uma ameaça explore com sucesso uma vulnerabilidade. No contexto desta dissertação, o impacto não é medido tecnicamente (ex.: "o servidor foi comprometido"), mas sim em termos de negócio (ex.: "o comprometimento do servidor resultou em perda financeira, dano reputacional ou interrupção da missão crítica da organização") (5).

O risco cibernético, portanto, só se materializa quando uma ameaça (o "quem/o quê") explora uma vulnerabilidade (o "onde/como") e causa impacto (o "prejuízo").

2.3 GESTÃO DE VULNERABILIDADES

A Gestão de Vulnerabilidades (GV) é o processo tático e operacional da Gestão de Riscos de SI que se concentra especificamente no componente "vulnerabilidade". Dado que as organizações têm controle limitado sobre as ameaças externas, mas controle total sobre seus próprios sistemas, a GV é considerada uma das práticas de defesa mais eficazes (9).

O ciclo de vida da Gestão de Vulnerabilidades é frequentemente descrito em quatro fases, conforme popularizado por *frameworks* como o do NIST (National Institute of Standards and Technology):

- **Descobrir (Discover):** A primeira fase consiste em realizar um inventário de ativos e utilizar ferramentas de varredura (*scanners*) para identificar vulnerabilidades conhecidas nesses

ativos.

- **Priorizar (Prioritize/Assess):** Esta é a fase de análise. As vulnerabilidades descobertas são avaliadas com base na severidade e no contexto de negócio para determinar a urgência da correção.
- **Corrigir (Remediate):** A fase de tratamento, na qual as equipes de tecnologia aplicam correções (*patches*) ou mitigações (controles compensatórios) necessárias.
- **Monitorar (Verify/Monitor):** A fase final verifica se a correção foi aplicada com sucesso e reinicia o ciclo, buscando continuamente novas vulnerabilidades.

Como destacado na Introdução (Capítulo 1), a fase de "Priorizar" é o principal gargalo deste ciclo e o foco desta pesquisa. Para priorizar, é necessário, primeiro, classificar o que foi encontrado, o que exige uma taxonomia padronizada.

2.4 TAXONOMIA DE VULNERABILIDADES: CWE E CVE

A comunicação eficaz sobre falhas de segurança depende de um vocabulário comum. A organização MITRE, com o apoio da comunidade global de segurança, mantém dois dos catálogos mais importantes dessa taxonomia: o CWE e o CVE. A confusão entre esses dois conceitos é comum, mas a distinção entre eles é fundamental para a metodologia desta dissertação.

2.4.1 CWE: A Fraqueza (A Causa Raiz)

O *Common Weakness Enumeration* (CWE) é um dicionário de tipos de fraquezas de *software* e *hardware* (31). O CWE não descreve uma vulnerabilidade em um produto específico, mas sim a classe de erro de programação ou de projeto que a leva. Por exemplo, "CWE-89: Neutralização Imprópria de Elementos Especiais usados em um Comando SQL ('SQL Injection')" é a descrição da fraqueza genérica. Ferramentas de análise estática e dinâmica, como o OWASP ZAP, são projetadas para identificar essas classes de fraquezas no código (11).

2.4.2 CVE: A Vulnerabilidade (A Instância)

O *Common Vulnerabilities and Exposures* (CVE) é o padrão global para identificar e catalogar instâncias específicas de vulnerabilidades em produtos de *software* e *hardware* (8). Um CVE representa a materialização de um ou mais CWEs em um produto real.

Por exemplo, "CVE-2017-5638" é uma vulnerabilidade específica de execução remota de código no Apache Struts. A causa raiz (o tipo de fraqueza) dessa vulnerabilidade se enquadra em uma ou mais categorias da CWE. O CVE é o identificador único que permite que equipes de segurança, pesquisadores e ferramentas de ameaças (como o EPSS) se refiram inequivocamente à mesma falha.

A metodologia desta dissertação foi diretamente influenciada por essa taxonomia. Como as ferramentas de varredura (ex.: ZAP) reportam CWEs e as ferramentas de inteligência de ameaças (ex.: EPSS) monitoram CVEs, um "cruzamento" direto dos dados não é trivial e exigiu a simulação de um *dataset* associando CVEs reais a ativos, conforme detalhado no Capítulo 4.

2.5 O MÉTODO TRADICIONAL DE PRIORIZAÇÃO: CVSS

Historicamente, o padrão de indústria para a priorização de vulnerabilidades tem sido o *Common Vulnerability Scoring System* (CVSS) (14). Mantido pelo FIRST.org, o CVSS é um *framework* aberto projetado para comunicar as características e a severidade técnica de vulnerabilidades de *software*. Sua finalidade é fornecer uma pontuação numérica que auxilie as organizações a avaliar a urgência da remediação.

A pontuação CVSS é composta por três grupos de métricas:

- **Métricas Base (Base Score):** Representam as qualidades intrínsecas e estáticas da vulnerabilidade, que não mudam com o tempo ou entre ambientes. Este grupo avalia a Explorabilidade (Vetores de Ataque, Complexidade do Ataque, Privilégios Requeridos, Interação do Usuário) e o Impacto (Confidencialidade, Integridade, Disponibilidade) (32). Esta é a pontuação (de 0,0 a 10,0) mais comumente usada para priorização.
- **Métricas Temporais (Temporal Score):** Ajustam a pontuação base de acordo com fatores que mudam ao longo do tempo, como a disponibilidade de um código de exploração (*Exploit Code Maturity*) ou de uma correção oficial (*Remediation Level*) (32).
- **Métricas Ambientais (Environmental Score):** Permitem que uma organização ajuste a pontuação ao seu contexto de negócio específico, modificando a importância dos pilares de impacto (Confidencialidade, Integridade, Disponibilidade) conforme o ativo afetado (32).

Na prática, devido à complexidade de avaliar métricas temporais e ambientais em larga escala, a vasta maioria das organizações utiliza exclusivamente a Pontuação Base (CVSS Base Score) como seu principal (e, muitas vezes, único) mecanismo de priorização (15). As vulnerabilidades

são então categorizadas em níveis de severidade Baixo, Médio, Alto e Crítico, e as equipes de segurança são instruídas a corrigir primeiro as de maior severidade.

2.6 A LACUNA DE PRIORIZAÇÃO: A CRÍTICA AO CVSS

Embora o CVSS seja fundamental para padronizar a descrição da severidade, seu uso como única ferramenta de priorização de risco constitui uma lacuna metodológica amplamente documentada na literatura (17, 18, 15). O problema central é que o CVSS é um "mau preditor" (*poor predictor*) da exploração real de vulnerabilidades (15).

A pontuação base do CVSS é estática e concentra-se na gravidade teórica da falha, ignorando duas das variáveis mais importantes do risco real: a probabilidade de exploração e o impacto contextual no negócio (19). Conforme apontado por Spring et al. (2021) (16), uma vulnerabilidade "Crítica" em um servidor de desenvolvimento interno representa um risco muito menor do que uma "Média" em um servidor *web* público que processa dados de clientes.

Estudos empíricos confirmam essa distorção. Jacobs et al. (2020) (15) observaram que apenas cerca de 5% de todas as vulnerabilidades publicadas são, de fato, exploradas ativamente no mundo real. No entanto, a priorização baseada no CVSS frequentemente classifica uma fração muito maior de vulnerabilidades como "Críticas" ou "Altas". No *dataset* desta pesquisa, por exemplo (conforme detalhado no Capítulo 5), 79,5% das vulnerabilidades foram classificadas nestas duas categorias superiores.

Este fenômeno é conhecido como "fadiga de alertas" (*alert fatigue*) (13, 16). As equipes de segurança são inundadas por um volume impraticável de alertas de alta prioridade, o que torna impossível remediar todos eles. Na prática, quando tudo é prioritário, nada é prioridade. Zeng et al. (2022) (17) e Yoon et al. (2023) (18) reforçam que o CVSS não foi projetado para avaliar o risco isoladamente, pois não associa de forma realista as vulnerabilidades à probabilidade de exploração.

Conclui-se, portanto, que uma metodologia de priorização eficaz não pode depender apenas da severidade estática. Ela deve, obrigatoriamente, incorporar uma medida de Probabilidade (o que será explorado) e uma medida de Impacto (o que importa para a organização).

2.7 SOLUÇÕES MODERNAS PARA A PRIORIZAÇÃO

Em resposta direta às limitações do CVSS, a comunidade de segurança desenvolveu novos *frameworks* que se alinham melhor ao conceito de risco (Probabilidade vs. Impacto). Duas abordagens se destacam como o estado da arte: o EPSS, focado na probabilidade, e o SSVC, focado no processo de decisão contextual.

2.7.1 A Dimensão da Probabilidade: EPSS

O *Exploit Prediction Scoring System* (EPSS) é a principal solução *data-driven* para quantificar a dimensão da Probabilidade do risco (33, 19). Mantido pelo mesmo FIRST.org que hospeda o CVSS, o EPSS é uma iniciativa de um Grupo de Interesse Especial (SIG) que reúne especialistas da indústria, da academia e do governo (15).

Diferentemente do CVSS, o EPSS não é um julgamento qualitativo, mas sim o resultado de um modelo de *Machine Learning* (especificamente, um XGBoost) treinado diariamente (15). Este modelo analisa um vasto conjunto de características (*features*) para calcular sua pontuação, incluindo:

- **Características da CVE:** idade da vulnerabilidade e métricas CVSS.
- **Disponibilidade de Exploit:** Verificações em fontes públicas como Metasploit, Exploit-DB e GitHub (15).
- **Discussão Pública:** Menções da CVE em mídias sociais (como o Twitter) e em ferramentas de segurança ofensiva (15).
- **Telemetria de Exploração (Mundo Real):** O componente mais crítico. O EPSS recebe dados de telemetria de uma vasta rede de parceiros (incluindo fornecedores de segurança como Fortinet, AlienVault, etc.) sobre tentativas reais de exploração observadas em seus sistemas de detecção de intrusão (IDS/IPS) e em *honeypots* (15).

O resultado é uma pontuação única, de 0% a 100%, que representa a probabilidade de uma vulnerabilidade (identificada pelo seu CVE) ser ativamente explorada nos próximos 30 dias (15).

A validação do modelo (EPSS v3) demonstra sua eficácia superior. Jacobs et al. (2023) (15) mostram que o modelo apresenta um desempenho 82% superior ao de seu predecessor. Na prática, uma estratégia que prioriza vulnerabilidades com EPSS acima de 8,8% (esforço de 7,3% do total) alcança a mesma cobertura de vulnerabilidades exploradas (82%) que uma estratégia CVSS

"Crítico ou Alto"(esforço de 58,1% do total) (15). O EPSS, portanto, é a métrica de probabilidade utilizada nesta dissertação.

2.7.2 A Dimensão da Decisão: SSVC

Uma abordagem alternativa influente é a *Stakeholder-Specific Vulnerability Categorization* (SSVC), desenvolvida pelo Software Engineering Institute (SEI) da Carnegie Mellon University e adotada pela CISA (16).

O SSVC não produz uma pontuação numérica, mas sim uma árvore de decisão qualitativa. O *framework* reconhece que a priorização não é um cálculo único, mas um processo de decisão que depende do contexto do *stakeholder* (a organização) (19). O analista de segurança utiliza a árvore de decisão do SSVC para avaliar uma vulnerabilidade com base em vários fatores, como o "Estado da Exploração"(se já existe prova de conceito), o "Impacto Técnico"e o "Contexto do Ativo".

O resultado do SSVC não é uma nota, mas uma ordem de ação imediata: *Act* (Agir), *Attend* (Atender), *Track* (Rastrear) ou *Defer* (Adiar) (19). Embora altamente eficaz para priorização contextualizada, o SSVC é um processo inerentemente manual e qualitativo, que exige análise caso a caso.

A metodologia desta dissertação se inspira em ambas as abordagens: utiliza a métrica de probabilidade quantitativa do EPSS, mas busca contextualizar o impacto proposta pelo SSVC por meio da métrica quantitativa NCS.

2.8 A DIMENSÃO DO IMPACTO: O CONTEXTO DE NEGÓCIO (PPSI/NCS)

Enquanto o EPSS fornece uma solução robusta para a dimensão da Probabilidade, uma priorização de risco eficaz exige uma segunda dimensão: o Impacto (21). Conforme criticado na Seção 2.6, as métricas de impacto do CVSS são genéricas e agnósticas em relação ao ambiente. A metodologia desta dissertação busca preencher essa lacuna por meio de uma métrica de impacto contextualizada, derivada do ambiente de aplicação: o setor público brasileiro.

A métrica de impacto selecionada é a Nota de Criticidade de Sistema (NCS), definida no *framework* de Privacidade e Segurança da Informação (PPSI) (22). O PPSI, mantido pelo Ministério da Gestão e da Inovação em Serviços Públicos (MGI), é a estrutura de governança de segurança e privacidade que orienta os órgãos do Sistema de Administração dos Recursos de Tecnologia da Informação (SISP) do Governo Federal.

A relevância desta métrica é evidenciada em diversos trabalhos acadêmicos no contexto brasileiro, muitos dos quais originados no Programa de Pós-Graduação Profissional em Engenharia Elétrica (PPEE) da Universidade de Brasília (UnB) e no Encontro de Administração da Justiça (ENAJUS). Esses estudos analisam a gestão de riscos de negócio, de privacidade e de segurança cibernética especificamente no âmbito do Poder Judiciário brasileiro (5, 34, 35, 36).

A própria origem da métrica NCS fundamenta-se em uma necessidade de governança pública, tendo sido baseada no modelo de avaliação de criticidade de sistemas estabelecido pelo Tribunal de Contas da União (TCU) no Acórdão 1.889/2020 (23). O modelo calcula a nota de criticidade (NC) de um ativo (sistema) por meio da fórmula $NC = (2 \times I) + V$, em que 'I' representa o impacto para o negócio (com base em critérios como dano financeiro, reputacional e à missão) e 'V' representa a vulnerabilidade do ativo (com base nos controles de segurança existentes) (19).

Portanto, a NCS não é uma métrica técnica genérica, mas sim uma avaliação de impacto de negócios contextualizada, alinhada às diretrizes de governança do setor público brasileiro. A sua utilização como dimensão de "Impacto" nesta dissertação permite alinhar diretamente a priorização de vulnerabilidades ao risco real do negócio da organização.

2.9 A PRIORIZAÇÃO DE VULNERABILIDADES COMO ESTRATÉGIA DE INTELIGÊNCIA DE SEGURANÇA PÚBLICA (SISP)

A transformação da segurança da informação em um pilar estratégico estatal exige a integração das práticas de defesa cibernética ao Subsistema de Inteligência de Segurança Pública (SISP). A complexidade das ameaças digitais impõe aos órgãos de segurança a necessidade de adotar ferramentas de Inteligência Estratégica Cibernética, capazes de mapear superfícies de ataque e antecipar incidentes de forma proativa (37, 38).

A gestão e a priorização de vulnerabilidades operam como instrumentos de inteligência acionáveis. Ao empregar técnicas de Inteligência de Fontes Abertas (OSINT) para a descoberta de ativos e o cruzamento de métricas probabilísticas de exploração, as agências de segurança conseguem identificar alvos prioritários antes da concretização de um ataque (37). Além de proteger a Segurança Operacional (OpSec) das infraestruturas críticas, a focalização dos esforços mitigatórios nos perfis de risco mais elevados reduz significativamente o Retorno sobre o Investimento (ROI) do adversário, desarticulando a eficiência de campanhas criminosas em larga escala (37, 39). A otimização dos recursos humanos e tecnológicos, por meio da priorização inteligente, converge diretamente com as diretrizes de reestruturação da governança cibernética no âmbito do Sistema Brasileiro de Inteligência (38).

2.10 A FERRAMENTA DE ANÁLISE: MACHINE LEARNING NÃO SUPERVISIONADO

A combinação das métricas de Probabilidade (EPSS) e de Impacto (NCS) resulta em um *dataset* bidimensional. Analisar manualmente milhares de pontos de dados neste espaço para identificar padrões é inviável. A solução para esta análise em larga escala reside na aplicação de técnicas de Ciência de Dados (*Data Science*).

Em um contexto de Segurança da Informação, o volume de dados gerados por *logs*, alertas de *scanners* e *feeds* de inteligência de ameaças caracteriza um problema de *Big Data* (6, 2). A Ciência de Dados é o campo interdisciplinar que aplica métodos científicos, processos e algoritmos para extrair conhecimento e *insights* a partir de dados estruturados e não estruturados. A principal ferramenta para a extração de padrões é o *Machine Learning* (ML) (40).

O *Machine Learning* é um subcampo da Inteligência Artificial que se concentra no desenvolvimento de algoritmos que permitem aos computadores aprender a partir de dados, sem serem explicitamente programados para cada tarefa. Os métodos de ML são comumente divididos em duas categorias principais:

1. **Aprendizado Supervisionado (*Supervised Learning*):** O algoritmo é treinado com um *dataset* rotulado, em que a resposta correta já é conhecida (ex.: classificar *e-mails* como 'spam' ou 'não-spam'). O objetivo é aprender uma função que mapeie entradas em saídas.
2. **Aprendizado Não Supervisionado (*Unsupervised Learning*):** O algoritmo recebe um *dataset* "não rotulado" e deve encontrar estruturas ou padrões latentes nos dados por conta própria. O objetivo não é prever uma saída, mas sim descobrir agrupamentos ou anomalias (40, 41).

Para o problema desta dissertação, "Quais perfis de risco existem em meu ambiente?", o aprendizado supervisionado não é adequado, pois os perfis (rótulos) não são conhecidos *a priori*. A abordagem correta é o aprendizado não supervisionado, especificamente a tarefa de *clustering* (agrupamento), que visa descobrir esses perfis de forma emergente a partir da estrutura dos próprios dados (6).

2.11 ALGORITMOS DE CLUSTERIZAÇÃO PARA DESCOBERTA DE PERFIS

A clusterização é a tarefa de agrupar um conjunto de objetos (neste caso, vulnerabilidades) de modo que os objetos de um mesmo grupo (chamado de *cluster*) sejam mais similares entre si do que com os de outros grupos. Esta pesquisa utiliza dois dos algoritmos de *clustering* mais fundamentais e amplamente validados na literatura: K-Means e DBSCAN (25).

2.11.1 K-Means: Agrupamento Baseado em Centróide

O K-Means é, talvez, o algoritmo de *clustering* particional mais conhecido e utilizado, proposto originalmente por MacQueen (1967) (42). Sua lógica baseia-se em centróides: o algoritmo tenta particionar n observações em k *clusters*, em que cada observação pertence ao *cluster* cujo centróide (a média) é o mais próximo.

O processo é iterativo:

1. **Inicialização:** k centróides são inicializados aleatoriamente no espaço dos dados.
2. **Atribuição:** Cada ponto de dados é atribuído ao *cluster* cujo centróide é o mais próximo. A "proximidade" é classicamente medida pela distância euclidiana (24).
3. **Atualização:** Os centróides são recalculados como a média de todos os pontos de dados atribuídos ao respectivo *cluster*.

Estes passos são repetidos até que os centróides se estabilizem (convergência). O K-Means visa minimizar a soma das distâncias quadráticas intra-cluster (WCSS), também conhecida como "inércia".

Uma consideração metodológica crítica do K-Means, decorrente do uso da distância euclidiana, é sua sensibilidade a diferentes escalas de dados. Se uma *feature* (ex.: NCS, variando de 0 a 100) tiver uma escala muito maior do que outra (ex.: EPSS, de 0 a 1), ela dominará o cálculo da distância e distorcerá os *clusters*. Por esta razão, a padronização dos dados (ex.: usando `StandardScaler` para que todas as *features* tenham média 0 e desvio padrão 1) é uma etapa de pré-processamento obrigatória, conforme detalhado no Capítulo 4 (19).

A determinação do número ideal de *clusters* (k) é um desafio comum, frequentemente resolvido por métodos heurísticos como o Método do Cotovelo (*Elbow Method*) e a análise da Pontuação de Silhueta (*Silhouette Score*), ambos utilizados nesta dissertação (43, 44). O uso de K-Means

para agrupar dados de cibersegurança e identificar padrões de ameaça é uma abordagem validada na literatura (27, 26, 3).

2.11.2 DBSCAN: Agrupamento Baseado em Densidade e Detecção de Anomalias

O *Density-Based Spatial Clustering of Applications with Noise* (DBSCAN), proposto por Ester et al. (1996) (45), oferece uma abordagem fundamentalmente diferente. Em vez de particionar todos os dados com base em centroides, o DBSCAN agrupa pontos que estão "densamente conectados" e marca como *outliers* (ruído) os pontos isolados em regiões de baixa densidade.

O algoritmo é definido por dois parâmetros:

1. **Epsilon (ϵ):** O raio da vizinhança. Define a distância máxima para que dois pontos sejam considerados vizinhos.
2. **MinPts (Minimum Samples):** O número mínimo de pontos necessários dentro de uma vizinhança ϵ para que um ponto seja considerado um *core point* (ponto central de um *cluster*).

A principal vantagem do DBSCAN é sua capacidade de encontrar *clusters* de formas arbitrárias (não apenas esféricas, como o K-Means) e, crucialmente para esta dissertação, de identificar ruído (*outliers*) (25).

No contexto da priorização de vulnerabilidades, um *outlier* não é um erro estatístico, mas sim um risco anômalo, uma vulnerabilidade com uma combinação tão única de probabilidade e impacto que não se encaixa em nenhum perfil padrão. Conforme demonstrado por Fuhnwi et al. (2023) (25) e Goswami (2024) (40), o DBSCAN é uma ferramenta de primeira linha para a detecção de anomalias em dados de segurança.

Nesta dissertação, o K-Means é usado como ferramenta principal para identificar os perfis de risco (*clusters* principais), enquanto o DBSCAN é utilizado como ferramenta complementar, essencial para validar a estrutura de densidade dos *clusters* e isolar as anomalias estratégicas (*outliers*).

2.12 O USO DE DADOS SINTÉTICOS E A SEGURANÇA OPERACIONAL

A literatura científica aponta a escassez de bases de dados reais e atualizadas para o treinamento de modelos de aprendizado de máquina em segurança da informação (28). Organizações, especial-

mente no setor público, evitam compartilhar informações detalhadas sobre suas vulnerabilidades e os impactos de negócio correspondentes. A divulgação desses dados expõe a infraestrutura crítica e viola os princípios da Segurança Operacional (OPSEC), dificultando a pesquisa e a reprodutibilidade dos estudos na área (46).

Para contornar essa restrição sem comprometer a validade do experimento, pesquisadores recorrem à geração de dados sintéticos. O uso de Inteligência Artificial Generativa, com destaque para as Redes Adversárias Generativas (GANs), permite a criação de bases de dados que preservam a distribuição estatística do ambiente original. Esse método garante alta fidelidade matemática aos dados reais, mantendo o anonimato e a privacidade das informações institucionais (47, 48).

Além de proteger a origem dos dados, a abordagem sintética permite realizar testes de estresse. Modelos de clusterização necessitam de validação em cenários com grande volume de dados para atestar sua escalabilidade (49). O uso de dados gerados artificialmente facilita o balanceamento de classes e a injeção deliberada de casos extremos, testando os limites dos algoritmos na detecção de anomalias (28). A Tabela 2.1 resume os eixos que justificam a adoção dessa técnica na presente pesquisa.

Tabela 2.1: Resumo dos eixos de validação com dados sintéticos em cibersegurança.

Fonte: elaborada pelo autor (2025).

Eixo de Validação	Descrição Teórica	Aplicação na Pesquisa
Confidencialidade (OPSEC)	Ocultação das falhas reais do ambiente produtivo.	Preservar a segurança e a privacidade do órgão público analisado.
Reprodutibilidade	Geração de uma base de dados comparável e pública.	Permitir a validação do método estatístico por outros pesquisadores.
Teste de estresse	Injeção de grande volume de registros e de casos extremos.	Avaliar a escalabilidade do K-Means e a detecção de anomalias pelo DBSCAN.

3 TRABALHOS CORRELATOS

Enquanto o Capítulo 2 estabeleceu os fundamentos conceituais do risco, da gestão de vulnerabilidades e das métricas de priorização, este capítulo situa a presente dissertação no contexto da literatura científica existente. A metodologia aqui proposta, a aplicação de *clustering* não supervisionado a um *dataset* bidimensional de Probabilidade (EPSS) e Impacto (NCS), representa uma síntese de várias linhas de pesquisa em cibersegurança.

A análise dos trabalhos correlatos é apresentada em seis seções. Primeiramente, revisam-se as aplicações de *Machine Learning* para detecção de anomalias em segurança, estabelecendo o *framework* geral. Em segundo lugar, detalham-se as aplicações específicas dos algoritmos de *clustering* (K-Means e DBSCAN) em problemas de cibersegurança. Terceiro, são analisadas abordagens híbridas concorrentes que, assim como este trabalho, buscam resolver o problema da priorização de risco. Quarto, a pesquisa é contextualizada no âmbito da produção científica local do Programa de Pós-Graduação em Engenharia Elétrica (PPEE) da Universidade de Brasília (UnB) e do seu domínio de aplicação, o setor público brasileiro. Quinto, introduz-se a geração de dados sintéticos para avaliação de modelos. Por fim, a sexta seção apresenta o posicionamento da dissertação e a síntese dos trabalhos correlatos.

3.1 MACHINE LEARNING PARA DETECÇÃO DE ANOMALIAS EM SEGURANÇA

A transição dos métodos de segurança reativos, baseados em assinaturas, para métodos proativos, baseados em comportamento, é uma resposta direta à "fadiga de alertas"(16) e ao volume de dados de ameaças (2). A literatura recente converge em favor do uso de *Machine Learning* (ML) como ferramenta primária para esta tarefa.

Manoharan & Sarker (2022) (6) argumentam que o ML revoluciona a detecção de ameaças ao permitir que os sistemas "detectem anomalias, analisem padrões de comportamento e prevejam ameaças potenciais"(6). Mohammed (2024) (7) reforça essa visão ao detalhar a transformação das operações do Centro de Operações de Segurança (SOC) por meio do ML. O autor destaca o *User and Entity Behavior Analytics* (UEBA), um sistema que utiliza aprendizado não supervisionado (como *clustering* e análise de anomalias) para modelar o comportamento "normal" de usuários e dispositivos e, assim, identificar desvios que sinalizam uma ameaça, como um ataque interno ou uma conta comprometida (7).

Estes trabalhos validam a premissa de que a aplicação de ML não supervisionado para encontrar padrões e anomalias é uma abordagem de ponta para trazer inteligência ao caos de dados de segurança.

3.2 APLICAÇÕES DE CLUSTERING EM CIBERSEGURANÇA

Dentro do aprendizado não supervisionado, os algoritmos de *clustering* são frequentemente empregados para agrupar dados de segurança e identificar padrões de ataque, perfis de ameaça ou, como no caso desta dissertação, perfis de risco.

O K-Means, devido à sua eficiência e interpretabilidade, é uma escolha comum. Kashtalian & Sochor (2021) (27), por exemplo, aplicaram o K-Means a dados de *honeynets* para determinar "padrões de ataque" com base na atividade do invasor ao longo do tempo. De forma similar, Cao et al. (2025) (26) propõem um "K-Means com peso dinâmico" para avaliar a situação de cibersegurança global, agrupando países com base em seus índices de segurança. Outros, como Khaoula & Mohamed (2023) (24), buscam otimizar o algoritmo para dados de segurança, propondo o "K-means-dist", que combina as distâncias euclidianas e de correlação para melhorar a detecção de intrusões no *dataset* CIC-IDS2018.

O DBSCAN também é validado na literatura, especialmente por sua capacidade de detectar anomalias. Goswami (2024) (40) o identifica como uma técnica central para detecção de anomalias em tempo real, pois seu modelo baseado em densidade pode isolar eventos que não se conformam a nenhum *cluster* normal. Fuhnwi et al. (2023) (25) realizaram um estudo empírico comparando diretamente o K-Means (representativo) com o DBSCAN (baseado em densidade) para detecção de *outliers*, concluindo que, embora ambos sejam eficazes, o K-Means (especificamente o k-means++) apresentou maior robustez em suas métricas.

Esta dissertação alinha-se a estes trabalhos ao empregar o K-Means como método principal de perfilagem (assim como Kashtalian e Cao) e o DBSCAN como método complementar de validação e detecção de anomalias (assim como Goswami e Fuhnwi).

3.3 ABORDAGENS HÍBRIDAS DE PRIORIZAÇÃO DE RISCO

A aplicação de modelos de risco bidimensionais (Probabilidade x Impacto) para a priorização de vulnerabilidades é uma área de pesquisa emergente, com várias propostas concorrentes que visam superar o CVSS.

Zeng et al. (2022) (17) propuseram o LICALITY, um sistema que, como este trabalho, avalia o risco com base em "Likelihood and Criticality" (Probabilidade e Criticidade). No entanto, sua metodologia difere na derivação das métricas: o LICALITY utiliza *Natural Language Processing* (NLP) na descrição textual da vulnerabilidade para inferir a probabilidade e as próprias métricas de impacto do CVSS para definir a criticidade.

Yoon et al. (2023) (18) também propõem uma nova métrica, a *Real-world Risk Score* (RWS), que critica o CVSS por ignorar a "exploitability no mundo real". Sua abordagem também utiliza o EPSS como uma de suas entradas (Probabilidade de Exploração), mas a combina com uma métrica de *Exploit Code Maturity* (ECM), que avalia o nível de "maturidade" (ex.: prova de conceito, arma funcional) do código de *exploit* disponível publicamente.

Finalmente, Krishnan et al. (2022) (50) utilizam uma abordagem de *clustering* para *profiling*, mas com um objetivo diferente. Em vez de criar perfis de vulnerabilidades, eles criam perfis de dispositivos IoT com base em comportamentos de rede esperados definidos pelo *Manufacturer Usage Description* (MUD). Utilizando *clustering* hierárquico, eles detectam anomalias comportamentais que indicam um dispositivo comprometido.

Esta dissertação se diferencia dessas abordagens ao propor uma combinação única de métricas: utiliza a probabilidade *data-driven* do EPSS (como em Yoon et al.), mas combina-a com uma métrica de impacto de negócio contextualizada (NCS/PPSI), em vez de métricas técnicas (CVSS, como em LICALITY) ou de maturidade de *exploit* (ECM, como em Yoon et al.).

3.4 CONTEXTO DA PESQUISA NO ÂMBITO LOCAL (PPEE/UNB)

Esta dissertação se insere em uma linha de pesquisa consolidada no Programa de Pós-Graduação Profissional em Engenharia Elétrica (PPEE) da Universidade de Brasília (UnB), que foca na aplicação de *frameworks* de segurança e gestão de risco no contexto do setor público e do Poder Judiciário brasileiro.

Trabalhos anteriores do programa, como os de Arruda (35), Ribeiro (36) e Alves (5), exploraram em profundidade os desafios da Gestão de Riscos Cibernéticos, da Governança de Riscos de TI, do Tratamento de Incidentes (GRSO) e da mitigação de riscos de negócio no Judiciário. Estes trabalhos estabelecem a relevância de *frameworks* como o NIST e a ISO, e analisam os desafios específicos de privacidade (LGPD) e segurança cibernética neste domínio (34).

A presente pesquisa contribui para essa linha de investigação ao introduzir uma metodologia quantitativa e automatizada (ML) para um problema clássico de gestão de risco: a priorização de vulnerabilidades. Ao adotar a métrica NCS, derivada do *framework* PPSI e do Acórdão do

TCU (22, 23), documentos centrais para o contexto explorado por esses trabalhos correlatos, esta dissertação oferece uma solução que não é apenas academicamente robusta, mas também diretamente aplicável e relevante para o domínio prioritário de aplicação do PPEE.

3.5 GERAÇÃO DE DADOS SINTÉTICOS PARA AVALIAÇÃO DE MODELOS

A aplicação de dados sintéticos para suprir a escassez de amostras reais e testar os limites matemáticos de algoritmos é uma abordagem estabelecida na literatura. Trabalhos correlatos utilizam Inteligência Artificial Generativa para criar cenários de estresse e balancear classes em bases de dados de segurança.

Ammara, Ding e Tutschku (51) realizaram uma análise comparativa de arquiteturas generativas para tráfego de rede. Os autores demonstraram que os modelos baseados em redes adversárias generativas, especificamente o modelo CTGAN, apresentam superioridade na preservação das propriedades estatísticas (fidelidade) e no desempenho dos algoritmos de detecção (utilidade).

No contexto das métricas de validação de dados tabulares sintéticos, a pesquisa científica requer a avaliação da semelhança estatística de cada atributo. Kim et al. (52) utilizaram a ferramenta SDMetrics para calcular a fidelidade das variáveis individuais em ambientes de detecção de anomalias. A metodologia aplicada pelos autores emprega o teste de Kolmogorov-Smirnov (KS) para avaliar as variáveis contínuas e a Distância de Variação Total (TVD) para as variáveis categóricas. Essa padronização de testes permite atestar a proximidade estrutural da base artificialmente gerada em relação aos dados originais.

A Tabela 3.1 sintetiza os principais estudos recentes que utilizam dados sintéticos e testes de estresse em segurança da informação, alinhando as metodologias de validação da presente dissertação ao estado da arte.

Tabela 3.1: Síntese dos trabalhos correlatos sobre dados sintéticos em cibersegurança.

Fonte: elaborada pelo autor (2025).

Autor (Ano)	Foco do Estudo	Modelo Gerativo	Métrica de Validação
Ammara et al. (2024)	Framework de seleção arquitetural para geração de tráfego de rede.	CTGAN, Copula-GAN	Correlação e Fidelidade Estrutural
Kim et al. (2025)	Geração de dados com preservação de privacidade para sistemas de detecção.	Tabular Diffusion	KS e TVD
Cossetin Neto et al. (2025) (53)	Aprimoramento da detecção de ataques persistentes avançados.	GAN-Transformers	Distância Wasserstein
Nawaz et al. (2025) (54)	Aumento de dados para mitigar falsos negativos em classes minoritárias.	CTGAN	F1-Score e Recall em dados anômalos

3.6 POSICIONAMENTO DA DISSERTAÇÃO E SÍNTESE DOS TRABALHOS CORRELATOS

Os trabalhos correlatos validam as escolhas metodológicas desta dissertação e destacam sua contribuição singular. Enquanto a literatura demonstra consenso quanto à insuficiência do CVSS, as soluções propostas divergem. Abordagens como LICALITY (17) e RWS (18) focam em métricas técnicas, enquanto trabalhos de *clustering* em cibersegurança, como os de Kashtalian et al. (27) e Cao et al. (26), focam em padrões de ataque.

Esta dissertação sintetiza essas frentes: aplica os métodos de *clustering* (validados por Kashtalian, Fuhnwi (25), etc.) à métrica de probabilidade (EPSS, de Jacobs et al. (15)) e à métrica de impacto (NCS), esta última contextualizada para o setor público brasileiro (conforme discutido por Arruda (35), Alves (5), etc.). A eficácia dessa abordagem foi formalmente validada na literatura científica com a publicação dos resultados na Revista Tecnológica uniFATEC (55). Adicionalmente, conforme demonstrado por da Silva et al. (2025) (56) na identificação de ativos críticos (as "Joias da Coroa"), a correta valoração do ativo é um passo metodológico essencial que embasa a dimensão de impacto do modelo. A contribuição não está na invenção de um novo algoritmo, mas na aplicação de uma combinação única de métricas e métodos para criar perfis de risco estratégicos e acionáveis.

4 METODOLOGIA DA PESQUISA

Este capítulo detalha a arquitetura da pesquisa e a sequência de procedimentos técnicos adotados para atingir os objetivos propostos. A metodologia foi estruturada como uma abordagem mista, combinando procedimentos quantitativos e qualitativos para a construção e análise de um modelo de priorização de vulnerabilidades (57, 58).

Quanto à sua natureza, esta pesquisa classifica-se como aplicada, pois visa gerar conhecimento para aplicação prática na resolução de um problema específico e real: a priorização de vulnerabilidades em ambientes corporativos. No que tange aos objetivos, a pesquisa é de caráter exploratório-descritivo, buscando proporcionar uma nova compreensão do problema ao investigar o uso de métricas contextuais de risco (EPSS e NCS) e de técnicas de *Machine Learning* (clusterização) (58).

O fluxo metodológico, ilustrado na Figura 4.1, foi executado em um *pipeline* de engenharia de dados composto por dez fases, que vão desde a descoberta de ativos reais até a injeção de anomalias. O código fonte, os conjuntos de dados e os artefatos de pesquisa desenvolvidos neste trabalho estão disponíveis publicamente no repositório do GitHub (<https://github.com/MadsonMagno/Mestrado-UnB-DataScience>), garantindo a reprodutibilidade dos experimentos e fomentando futuras pesquisas na área.

É imperativo destacar que a metodologia proposta foca estritamente na fase de categorização e priorização. Ao final, no momento de definir a ação corretiva a ser tomada (aceitar, mitigar, transferir ou evitar), deve-se considerar ativamente o apetite ao risco predefinido pela organização dentro de seu ciclo de tratamento de risco geral. A gradação de criticidade sugerida pelo modelo atua como um balizador quantitativo, mas a decisão executiva final dependerá sempre do apetite ao risco corporativo.

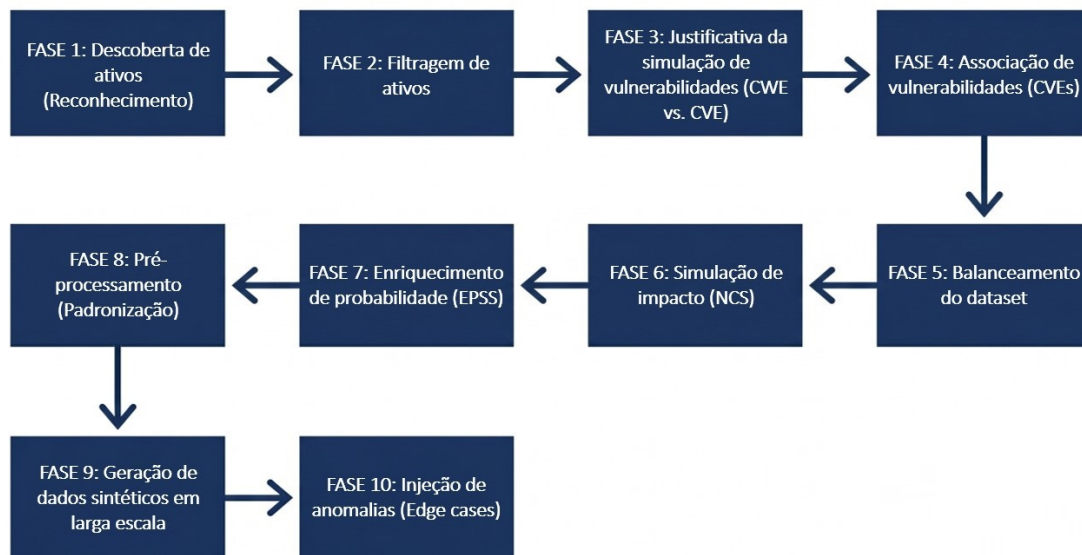


Figura 4.1: Fluxograma do pipeline de engenharia de dados e de modelagem.

4.1 FASE 1: DESCOBERTA DE ATIVOS (RECONHECIMENTO)

Diferentemente de uma simulação puramente teórica, esta pesquisa partiu da descoberta de ativos digitais reais pertencentes a uma organização pública. O objetivo desta fase de reconhecimento (*reconnaissance*) foi construir uma lista de alvos (sistemas *web*) que servisse de base para um cenário autêntico de gestão de vulnerabilidades.

Para esta tarefa, foram utilizados *scripts* de *Open Source Intelligence* (OSINT) e ferramentas de enumeração de subdomínios, conforme documentado nos *notebooks* da pesquisa (59). Ferramentas como DNSDumpster e Sublist3r foram executadas para consultar fontes públicas e identificar subdomínios associados ao domínio principal da organização-alvo (ex.: `mj.gov.br`). O resultado desta fase foi um conjunto bruto de URLs de ativos potencialmente expostos na *Internet*.

4.2 FASE 2: FILTRAGEM DE ATIVOS

A lista bruta de ativos identificada na Fase 1 não era diretamente utilizável. Muitos subdomínios poderiam estar inativos ou protegidos por mecanismos de autenticação que impedem a análise de vulnerabilidades por *scanners* automatizados.

Portanto, foi realizada uma fase de filtragem, documentada nos *notebooks* da pesquisa (59). Esta

filtragem consistiu em duas etapas: primeiro, um teste de conectividade (HTTP GET) foi realizado para verificar quais URLs estavam ativas e respondiam; segundo, foi realizada uma análise (no caso desta pesquisa, manual e exploratória) para remover sistemas que exigiam autenticação (*login*) imediata.

O resultado foi um *dataset* de ativos reais, públicos e acessíveis, que serviu como escopo de ativos para a análise de risco.

4.3 FASE 3: JUSTIFICATIVA DA SIMULAÇÃO DE VULNERABILIDADES (CWE VS. CVE)

A intenção original da pesquisa era realizar varreduras de vulnerabilidade diretamente nos ativos filtrados da Fase 2. Para isso, foram empregadas ferramentas padrão de mercado, como o OWASP ZAP (conforme *script scan_zap.ipynb*) e o OpenVAS (59).

Os relatórios gerados (*2025-06-29-ZAP-Report-.json* e *report1_openvas.csv*) (60, 61) foram bem-sucedidos em identificar diversas falhas de segurança. Contudo, esta etapa revelou uma incompatibilidade metodológica fundamental que inviabilizou o prosseguimento direto:

1. **Scanners (ZAP/OpenVAS):** Reportam Fraquezas (CWEs - *Common Weakness Enumeration*). Por exemplo, o ZAP identifica a ausência do cabeçalho de segurança "X-Content-Type-Options" como "CWE-693: X-Content-Type-Options Header Missing".
2. **Métrica de Probabilidade (EPSS):** É indexada por vulnerabilidades (CVEs - *Common Vulnerabilities and Exposures*). O EPSS fornece a probabilidade de exploração de um CVE específico (ex.: "CVE-2021-44228").

Não existe uma relação direta ou unívoca que permita converter automaticamente um CWE genérico (a classe de erro) em um CVE específico (a instância de erro em um *software*). Como a métrica de Probabilidade (EPSS) é um pilar central desta dissertação, e ela exige um CVE como entrada, os dados brutos dos *scanners* (baseados em CWE) se mostraram incompatíveis com os objetivos da pesquisa. Esta limitação técnica exigiu uma mudança metodológica para a simulação de vulnerabilidades.

4.4 FASE 4: ASSOCIAÇÃO DE VULNERABILIDADES (CVES)

Para contornar a incompatibilidade entre CWE e CVE, adotou-se uma abordagem de simulação de vulnerabilidades. Em vez de escanear os ativos (Fase 2) em busca de fraquezas, esta fase associou-os a uma lista de vulnerabilidades (CVEs) reais e publicamente conhecidas.

O processo, executado em um *notebook* próprio (59), consistiu em mapear os ativos filtrados para um conjunto de CVEs reais. Estes CVEs foram obtidos de bancos de dados oficiais, como o *National Vulnerability Database* (NVD) e o catálogo do MITRE. Esta associação simulou o *dataset* que uma organização madura possuiria em seu sistema de Gestão de Vulnerabilidades: uma lista de seus ativos e as vulnerabilidades conhecidas (CVEs) em cada um deles. O resultado foi o *dataset* `dataset_vulnerabilidades_simulado.csv` (62).

4.5 FASE 5: BALANCEAMENTO DO DATASET

Uma análise exploratória do *dataset* simulado revelou um problema comum em ambientes reais: o viés de amostragem. Alguns ativos, particularmente sistemas mais antigos ou complexos, apresentavam um número desproporcionalmente elevado de vulnerabilidades (centenas), enquanto outros tinham poucas. Se utilizado diretamente, este desequilíbrio faria com que os ativos "super-representados" dominassem o cálculo da distância no K-Means, distorcendo os perfis de risco.

Para corrigir isso, foi aplicada uma técnica de amostragem aleatória (balanceamento), conforme executado no *notebook* da pesquisa (59). Um limite aleatório de vulnerabilidades por ativo (entre 50 e 250) foi imposto, garantindo que nenhum ativo individual pudesse enviesar o modelo. O resultado deste processo foi o *dataset* `dataset_vulnerabilidades_balanceado_final.csv` (62).

4.6 FASE 6: SIMULAÇÃO DE IMPACTO (NCS)

Com o *dataset* de vulnerabilidades balanceado, a próxima etapa foi atribuir a segunda dimensão do risco: o "Impacto no Negócio". Como o acesso aos dados de criticidade reais dos ativos (Fase 1) não estava disponível, esta dimensão também foi simulada de forma controlada, conforme documentado no *notebook* da pesquisa (59).

Seguindo a metodologia do PPSI/TCU (Capítulo 2, Seção 2.8), os ativos foram divididos em perfis de criticidade (ex.: Alta, Média, Baixa) e tiveram seus valores de Impacto (I) e Vulnerabilidade

(V) simulados para gerar a Nota de Criticidade (NC) final, $NC = (2 \times I) + V$. O objetivo foi criar um cenário equilibrado, com uma distribuição realista da criticidade dos ativos. Esta nota de impacto foi então associada a todas as vulnerabilidades pertencentes a aquele ativo, gerando o *dataset* `relatorio_final_criticidade_ponderado.csv` (62).

4.7 FASE 7: ENRIQUECIMENTO DE PROBABILIDADE (EPSS)

A etapa seguinte consistiu em adicionar a primeira dimensão do risco: a "Probabilidade". Este processo foi documentado em um *notebook* da pesquisa (59).

O *script* iterou sobre cada vulnerabilidade (identificada pelo respectivo CVE) no *dataset* balanceado (Fase 5) e consultou o *feed* oficial de dados do EPSS v3 (mantido pelo FIRST.org) (63). A pontuação de probabilidade correspondente (valor entre 0 e 1) foi extraída e adicionada ao *dataset*. Vulnerabilidades que não possuíam um *score* EPSS (ex.: CVEs muito recentes ou muito antigos) receberam o valor 0, indicando probabilidade de exploração nula ou desconhecida.

O resultado desta fase foi um *dataset* (62) agora completo com as duas dimensões de risco: a Probabilidade (EPSS) e o Impacto (NCS/PPSI).

4.8 FASE 8: PRÉ-PROCESSAMENTO (PADRONIZAÇÃO)

A etapa final da engenharia de dados, executada em um *notebook* da pesquisa (59), consistiu no pré-processamento para a modelagem de *Machine Learning*.

Conforme estabelecido no Referencial Teórico (Seção 2.11.1), algoritmos baseados em distância, como o K-Means, são altamente sensíveis à escala das variáveis. No *dataset* da Fase 7, as duas dimensões de risco possuíam escalas drasticamente diferentes (EPSS variando de 0 a 1; NCS variando de 0,78 a 2,27, conforme a simulação).

Para garantir que ambas as dimensões tivessem peso igual no cálculo das distâncias euclidianas, foi aplicado o `StandardScaler` da biblioteca Scikit-learn (64). Este transformador removeu a média e escalou os dados para a variância unitária (média 0 e desvio padrão 1).

O resultado deste *pipeline* de 8 fases foi o *dataset* final `dataset_artigo.csv` (62). Este arquivo, contendo 421 vulnerabilidades com suas métricas de probabilidade e impacto normalizadas, serviu como a entrada principal para o *notebook* de modelagem (59), cujos resultados são detalhados no próximo capítulo. Antes da execução dos testes de estresse (Fases 9 e 10),

os algoritmos selecionados (K-Means e DBSCAN) foram primeiramente aplicados a este *dataset* padronizado original, permitindo a validação da eficácia do modelo em um cenário base.

4.9 FASE 9: GERAÇÃO DE DADOS SINTÉTICOS EM LARGA ESCALA

Para atestar a validade estrutural do modelo de agrupamento e sua resiliência estatística, o *pipeline* metodológico foi expandido com a aplicação de um teste de estresse em massa. Empregou-se a biblioteca *Synthetic Data Vault* (SDV) operando com a arquitetura de redes neurais generativas condicionais para dados tabulares (CTGAN).

O *dataset* padronizado, contendo 421 registros reais, serviu como base de treinamento para a rede neural. O modelo generativo aprendeu as matrizes de correlação latentes e as distribuições probabilísticas bidimensionais (Impacto e Probabilidade) do ambiente original. Após o ajuste dos parâmetros, a arquitetura gerou 100.000 amostras puramente sintéticas. O procedimento permitiu criar um cenário hipotético de *Big Data* governamental, preservando a fidelidade geométrica dos dados sem expor infraestruturas críticas reais.

4.10 FASE 10: INJEÇÃO DE ANOMALIAS (EDGE CASES)

A validação da eficácia do algoritmo DBSCAN como filtro espacial exigiu a simulação de cenários atípicos que ocorrem com baixa frequência na natureza, mas representam alto risco estratégico. Adotou-se a técnica de *Attack Augmentation*, inserindo artificialmente 1.000 casos extremos na massa sintética gerada.

O primeiro grupo de anomalias consistiu em 500 registros configurados com severidade técnica máxima (CVSS 10.0) e probabilidade de exploração quase nula (EPSS próximo de 0.0001), simulando o ruído responsável pela fadiga de alertas. O segundo grupo incorporou 500 registros com baixa severidade técnica (CVSS 0.0), mas com probabilidade de exploração extrema (EPSS 0.99), o que representa ataques furtivos altamente prováveis.

A base final submetida à validação estrutural perfaz 101.000 registros. A Tabela 4.1 resume a composição do ambiente computacional desenvolvido para o teste de estresse.

Tabela 4.1: Composição da base de dados para testes de estresse em larga escala.

Fonte: elaborada pelo autor (2026).

Origem dos Dados	Volume	Objetivo Metodológico
Dataset Original (Fases 1 a 8)	421	Treinamento de arquitetura generativa (SDV/CTGAN).
Massa Sintética (CTGAN)	100.000	Avaliação de escalabilidade computacional (Big Data).
Anomalias (Alto CVSS / Baixo EPSS)	500	Teste de mitigação de falsos positivos com k-Means.
Anomalias (Baixo CVSS / Alto EPSS)	500	Teste de detecção de anomalias espaciais com DBSCAN.
Volume Total do Experimento	101.000	Validação matemática e estrutural em cenários extremos.

A fundamentação para a seleção dos algoritmos K-Means e DBSCAN decorre de uma análise comparativa prévia validada na literatura científica (19). O estudo avaliou o desempenho de algoritmos de particionamento (K-Means), métodos hierárquicos aglomerativos (Ward) e abordagens baseadas em densidade (DBSCAN) na segmentação de matrizes bidimensionais de risco. A investigação comprovou que o K-Means (com $k=5$) apresenta o maior Coeficiente de Silhueta para a clusterização global, oferecendo o melhor equilíbrio entre coesão estatística e interpretabilidade para a tomada de decisão gerencial. Em contrapartida, o método DBSCAN demonstrou superioridade matemática (evidenciada pelo baixo Índice Davies-Bouldin) na detecção estrita de núcleos densos e isolamento de anomalias espaciais (19). O grau de concordância entre as partições, atestado por um Adjusted Rand Index (ARI) próximo a 0,90, validou a robustez estrutural da matriz e justificou a aplicação conjunta do K-Means para o estabelecimento de políticas gerais e do DBSCAN para a detecção de instâncias anômalas (*outliers*).

5 RESULTADOS E DISCUSSÃO

Este capítulo apresenta os resultados obtidos a partir da aplicação da metodologia detalhada no capítulo 4. A análise inicia-se com a Análise Exploratória de Dados (EDA), que valida as premissas da pesquisa, e segue para a determinação do número ideal de *clusters* (k). Posteriormente, são detalhados os perfis de risco identificados pela modelagem, a análise de anomalias (*outliers*), o impacto quantitativo da metodologia na priorização, a visualização consolidada da matriz de risco e a validação em larga escala com testes de estresse.

5.1 ANÁLISE EXPLORATÓRIA DE DADOS (EDA)

Antes de aplicar qualquer modelo de *Machine Learning*, é fundamental realizar uma Análise Exploratória de Dados (EDA). Esta prática, consolidada por Tukey (1977) (65), permite investigar as características principais do *dataset* `dataset_artigo.csv` (62), validar as hipóteses da pesquisa e identificar padrões que guiarão a modelagem.

Nesta pesquisa, a EDA (documentada no *notebook* `projeto_cienciadedados.ipynb` (59)) teve dois objetivos centrais:

1. Validar que as métricas de "Probabilidade"(EPSS) e "Impacto"(NCS) são estatisticamente independentes.
2. Confirmar visualmente a existência de agrupamentos (*clusters*) no *dataset* que justifiquem o uso de algoritmos de clusterização.

5.1.1 Análise de Correlação das Métricas de Risco (EPSS vs. NCS)

A primeira e mais importante hipótese a ser validada é a de que as duas dimensões de risco escolhidas, Probabilidade (EPSS) e Impacto (NCS), são independentes. Se as métricas tivessem alta correlação (positiva ou negativa), o modelo bidimensional seria redundante, pois uma delas poderia ser usada para prever a outra. O ideal é que as métricas sejam ortogonais, ou seja, que meçam facetas distintas do risco.

Para verificar esta independência, foi calculada a Matriz de Correlação de Pearson entre as duas variáveis normalizadas, conforme apresentado na Figura 5.1.

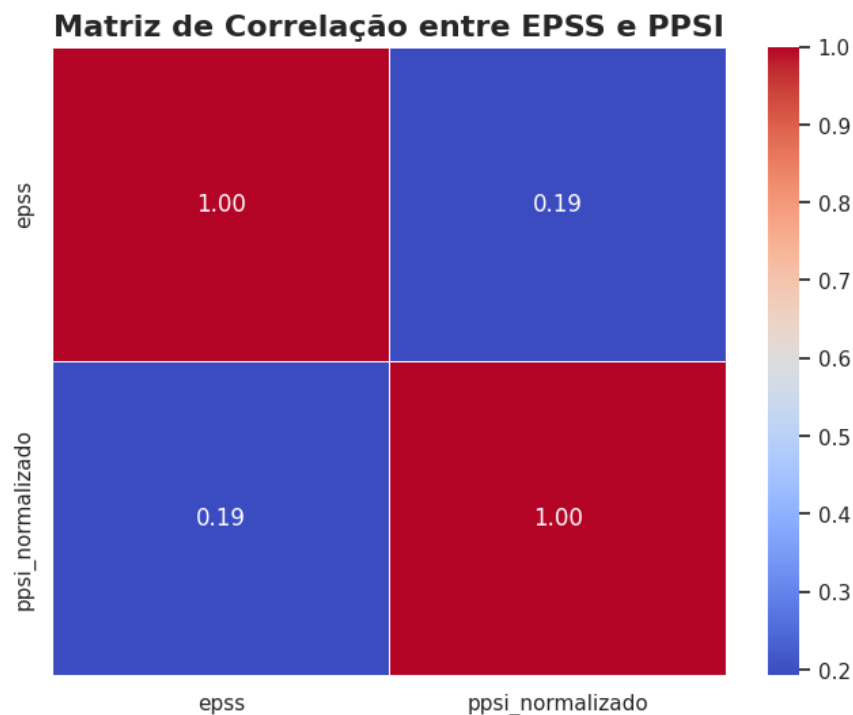


Figura 5.1: Matriz de correlação de Pearson entre EPSS e NCS (PPSI) normalizadas.

Fonte: elaborado pelo autor (2025).

O resultado, documentado no artigo (19), revela um coeficiente de correlação de 0,19. Este valor, muito próximo de zero, comprova empiricamente a baixa correlação linear entre as variáveis (19). Na prática, isso confirma que a probabilidade de uma vulnerabilidade ser explorada (EPSS) não possui relação direta com a criticidade do ativo em que ela reside (NCS). Esta independência estatística é fundamental, pois confirma a premissa de que o modelo bidimensional captura duas facetas complementares e necessárias do risco real.

5.1.2 Visualização da Matriz de Risco (Dispersão)

Confirmada a independência das métricas, a segunda etapa da EDA consistiu em investigar visualmente a distribuição dos dados para verificar se havia agrupamentos naturais. Se os pontos de dados (vulnerabilidades) estivessem distribuídos uniformemente e aleatoriamente no espaço, algoritmos de *clustering* teriam dificuldade em encontrar padrões significativos.

Para esta análise, foi gerado um gráfico de dispersão (*scatterplot*), ou "Matriz de Risco", posicionando cada uma das 421 vulnerabilidades do *dataset* de acordo com suas pontuações normalizadas de EPSS (eixo X) e NCS (eixo Y), conforme a Figura 5.2.

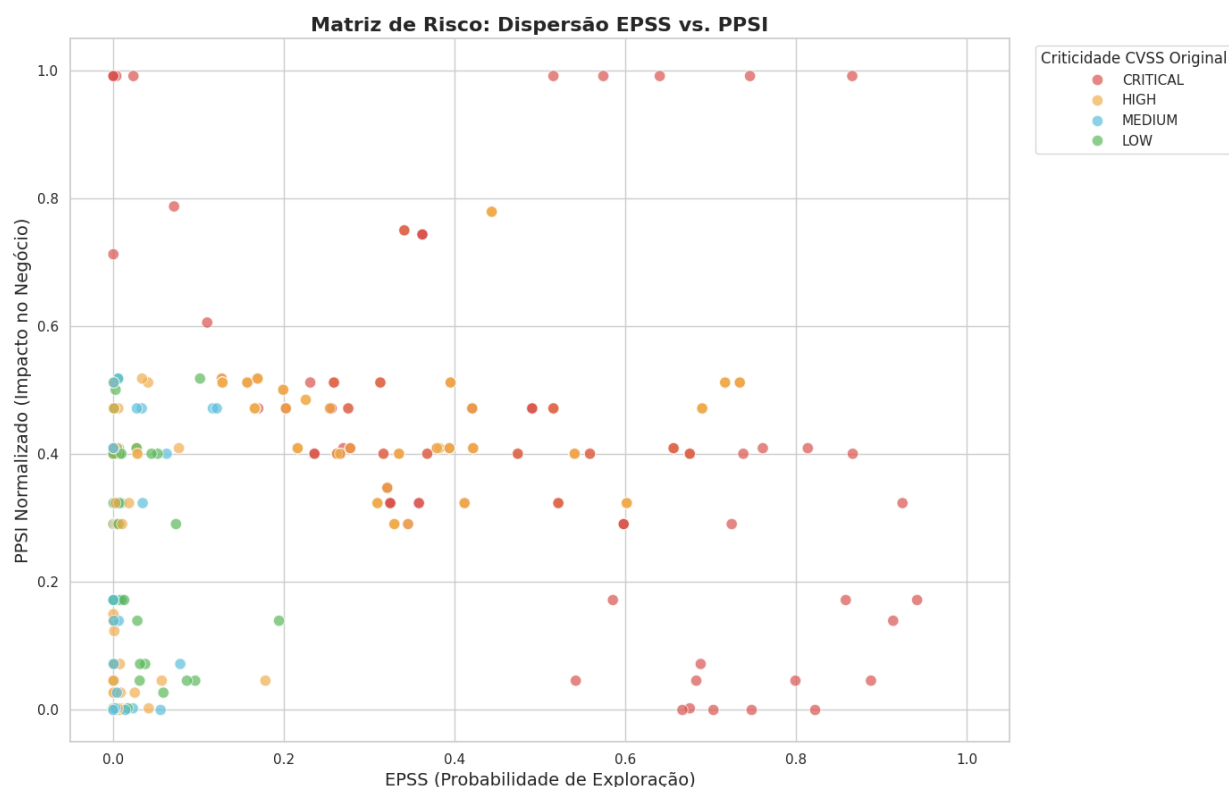


Figura 5.2: Matriz de Risco: Dispersão entre as EPSS e as NCS Normalizadas.

Fonte: elaborado pelo autor (2025).

A inspeção visual da Figura 5.2 evidencia claramente que as vulnerabilidades não se distribuem de forma uniforme (19). Ao contrário, elas formam concentrações (grupos) em regiões distintas do gráfico. Nota-se uma grande concentração de pontos no quadrante inferior esquerdo (baixo EPSS, baixo NCS), representando o "ruído" de vulnerabilidades de baixo risco. Em contrapartida, os quadrantes de maior risco (superior esquerdo, inferior direito e superior direito) são mais esparsos, mas apresentam agrupamentos visuais nítidos (19).

Esta distribuição desigual e a presença de "vazios" nas concentrações de dados sugerem fortemente a existência de uma estrutura de agrupamento. Isso justifica plenamente a aplicação de algoritmos de *clustering* para a identificação e caracterização formal desses perfis de risco latentes.

5.2 SELEÇÃO DO NÚMERO DE CLUSTERS (K)

A metodologia desta pesquisa utiliza o K-Means como principal algoritmo de *clustering*. Conforme detalhado no Capítulo 2, uma limitação do K-Means é a necessidade de pré-especificar o número de *clusters* (k) (42). Uma escolha inadequada de k pode levar a agrupamentos sub-

representados (poucos *clusters*) ou super-representados (*clusters* demais, dividindo grupos que deveriam ser únicos).

Para determinar o valor de k de forma objetiva, foram utilizadas duas técnicas heurísticas complementares, ambas executadas no *notebook* da pesquisa (59): uma análise quantitativa da variância (Método do Cotovelo) e uma análise estrutural (Dendrograma).

5.2.1 Método do Cotovelo (Elbow Method)

O Método do Cotovelo é uma das técnicas mais comuns para a seleção de k . Ele consiste em executar o algoritmo K-Means para uma faixa de valores de k (neste caso, de 1 a 10) e plotar a "inércia" (WCSS - *Within-Cluster Sum of Squares*) para cada valor de k . A inércia mede a soma das distâncias quadráticas de cada ponto ao centroide, ou seja, o quão compactos são os *clusters*.

O objetivo é identificar o "ponto de cotovelo" no gráfico: o ponto em que a adição de um novo *cluster* deixa de reduzir de forma expressiva a inércia, indicando o ponto de retornos decrescentes. O resultado desta análise é apresentado na Figura 5.3.

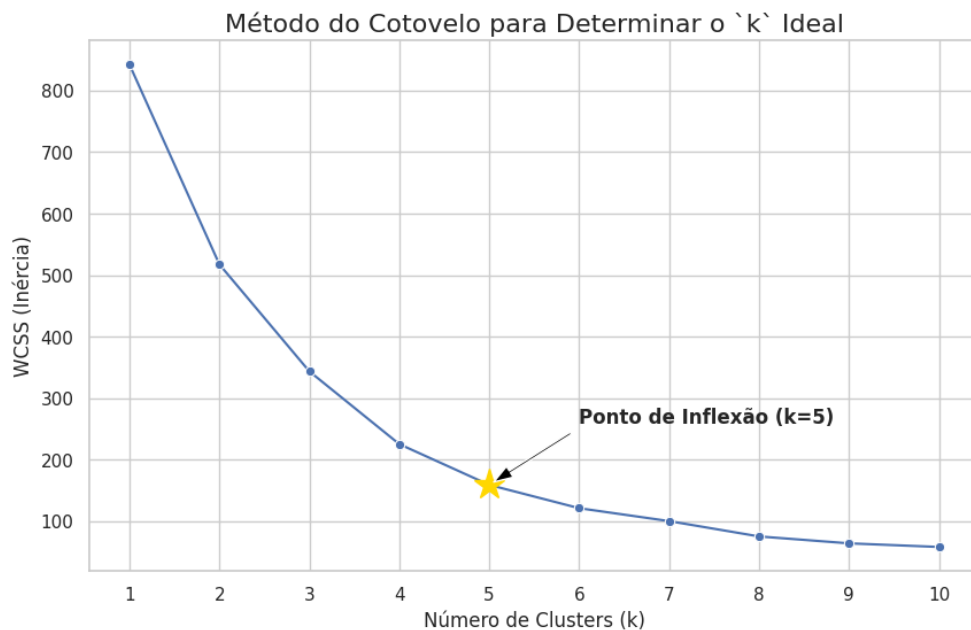


Figura 5.3: Análise do Método do Cotovelo (WCSS) para determinar o k ideal.

Fonte: Elaborado pelo autor (2025).

A análise da Figura 5.3 revela uma queda abrupta na inércia, de $k = 1$ a $k = 4$. Conforme documentado no artigo (19), embora a transição de $k = 4$ para $k = 5$ ainda promova uma redução significativa, a redução obtida ao passar de $k = 5$ para $k = 6$ é consideravelmente menor. Este ponto em $k = 5$ é interpretado como o "cotovelo", representando o melhor equilíbrio entre a

coesão intra-cluster (baixa inércia) e a complexidade do modelo (número de *clusters*) (19).

5.2.2 Análise Estrutural (Dendrograma)

Para complementar a análise quantitativa do Método do Cotovelo, realizou-se uma análise estrutural por meio do *Clustering* Hierárquico Aglomerativo. Este método não requer a pré-definição de k , mas constrói uma árvore de *clusters* (Dendrograma) que mostra como os pontos de dados são agrupados em grupos com base na distância.

A Figura 5.4 apresenta o dendrograma gerado para o *dataset*. A inspeção visual desta estrutura permite validar a escolha de k .

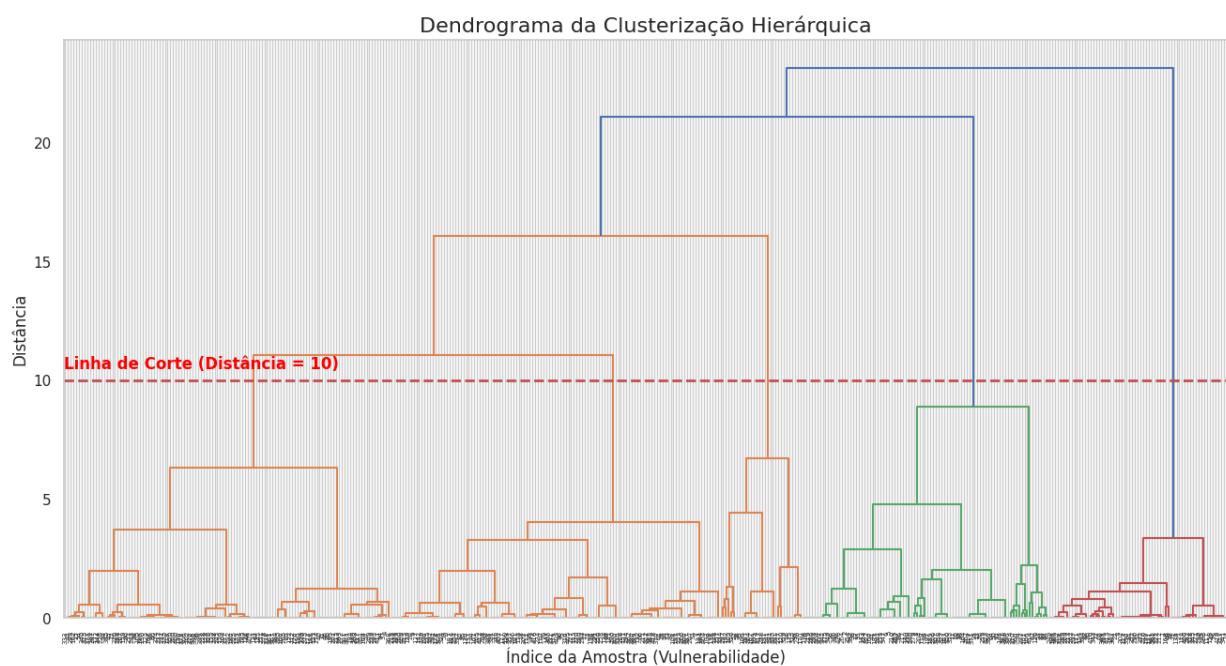


Figura 5.4: Análise estrutural dos agrupamentos via Dendrograma de Clusterização Hierárquica.

Fonte: Elaborado pelo autor (2025).

Como demonstra a Figura 5.4, ao traçar uma linha de corte horizontal na maior distância vertical que não intercepta um agrupamento já consolidado (indicada no gráfico original do artigo (19)), a estrutura de dados é naturalmente dividida em cinco agrupamentos estruturalmente distintos.

A convergência entre a análise de variância (Método do Cotovelo) e a análise estrutural (Dendrograma), ambas indicando $k = 5$, forneceu uma forte justificativa metodológica para a adoção deste valor na modelagem final com o K-Means.

5.3 DESCOBERTA E CARACTERIZAÇÃO DOS PERFIS DE RISCO

Após a validação do número ideal de agrupamentos ($k = 5$), o algoritmo K-Means foi executado sobre o *dataset* normalizado `dataset_artigo.csv` (62). A principal contribuição desta metodologia reside na análise dos centroides de cada um dos 5 *clusters* resultantes. O centroide representa o "centro de gravidade" de um *cluster*, ou seja, o valor médio das *features* (EPSS e NCS) de todas as vulnerabilidades que compõem aquele grupo (19).

A análise desses centroides, documentada no *notebook* da pesquisa (59), permite transcender a análise técnica individual de cada falha e identificar "perfis de risco" estratégicos. A Tabela 5.1 apresenta os valores médios (não normalizados, para facilitar a interpretação) das métricas EPSS (Probabilidade) e NCS (Impacto) para cada *cluster* identificado.

Tabela 5.1: Análise dos centroides dos 5 perfis de risco identificados pelo K-Means.

Fonte: Elaborado pelo autor (2025).

Cluster ID	EPSS (Média)	NCS (Média)	Perfil de Risco Proposto
0	0,33	0,86	Risco Crítico
1	0,67	0,35	Ameaça de Volume
2	0,34	0,41	Risco Médio (Padrão)
3	0,07	0,44	Gigante Adormecido
4	0,02	0,07	Risco Mínimo (Ruído)

Para facilitar a interpretação visual desses perfis, a Figura 5.5 apresenta os mesmos dados da Tabela 5.1 em um mapa de calor (*heatmap*). Neste gráfico, as cores mais quentes (vermelho escuro) indicam valores mais elevados, permitindo a identificação imediata das características dominantes de cada perfil de risco.

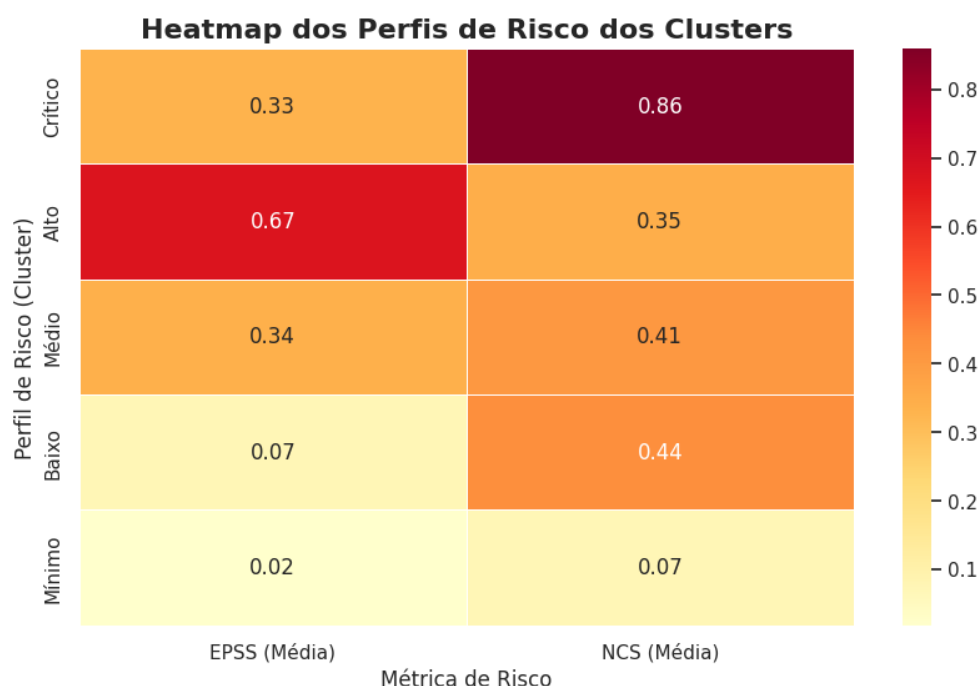


Figura 5.5: Mapa de calor (Heatmap) dos valores médios dos centroides (Perfis de Risco).

Fonte: Elaborado pelo autor (2025).

A análise combinada da Tabela 5.1 e da Figura 5.5 revela cinco perfis de risco estratégicos e acionáveis, que são detalhados exhaustivamente nas subseções seguintes.

5.3.1 Perfil 1: Risco Crítico (Cluster 0)

O Perfil "Risco Crítico"(Cluster 0) é quantitativamente definido por um Impacto (NCS) médio de 0,86 e uma Probabilidade (EPSS) média de 0,33. Este perfil representa o risco estratégico mais elevado para a organização.

A característica dominante deste *cluster* é o valor de NCS, o mais alto entre todos os perfis. Isso indica que as vulnerabilidades aqui agrupadas residem nos ativos de missão crítica (os sistemas "joias da coroa"da organização), onde uma exploração bem-sucedida resultaria em danos catastróficos, seja por meio de perdas financeiras, interrupção da missão principal ou grave dano reputacional, conforme os critérios do PPSI/TCU (22, 23).

Curiosamente, a probabilidade de exploração (EPSS 0,33) não é a mais alta do *dataset*, mas sim "moderada". Isso sugere que estas não são, necessariamente, as vulnerabilidades mais fáceis de explorar, mas que apresentam evidências suficientes de exploração no mundo real (como capturado pelo EPSS (15)) para justificar uma ação imediata. A combinação de um impacto catastrófico e de uma probabilidade de exploração moderada torna este perfil uma prioridade inequívoca.

Estratégia de Remediação: Imediata. Este perfil se alinha ao *bucket* "Act"(Agir) do *framework* SSVc (16). A remediação destas falhas deve ter precedência sobre todas as demais, pois o impacto de um incidente supera qualquer outra consideração.

5.3.2 Perfil 2: Ameaça de Volume (Cluster 1)

O Perfil "Ameaça de Volume"(Cluster 1) é definido por uma Probabilidade (EPSS) média de 0,67 e um Impacto (NCS) médio de 0,35. Este perfil é, em essência, o oposto do "Risco Crítico" em termos de características de risco.

A característica dominante aqui é a probabilidade de exploração. Com um *score* médio de 0,67, este é o *cluster* com a maior probabilidade de exploração em todo o *dataset*. Estas são as vulnerabilidades "da moda", ativamente exploradas por atacantes automatizados (*bots*) e por campanhas de *malware* em larga escala (15).

Contudo, estas falhas altamente prováveis residem em ativos de impacto de negócio moderado a baixo (NCS de 0,35). Elas representam a "linha de frente" dos incidentes de segurança; embora a exploração de uma única vulnerabilidade deste grupo não seja catastrófica, o alto volume de ataques esperado gera um ruído operacional significativo e representa a ameaça mais provável de se concretizar.

Estratégia de Remediação: Prioritária, com foco na automação. A correção destas falhas é essencial para reduzir o volume de incidentes diários e a "fadiga de alertas" da equipe de segurança. A remediação pode ser tratada logo após o Perfil Crítico, muitas vezes por meio de *patches* em massa ou de regras de *firewall* (WAF) para conter o alto volume.

5.3.3 Perfil 3: Risco Médio (Cluster 2)

O Perfil "Risco Médio"(Cluster 2) é o perfil de *baseline*, caracterizado por valores médios em ambas as dimensões: Probabilidade (EPSS) média de 0,34 e Impacto (NCS) médio de 0,41.

Estas vulnerabilidades não evidenciam o impacto extremo do Perfil Crítico nem a probabilidade de exploração generalizada do Perfil de Volume. Elas representam o risco operacional padrão que as organizações enfrentam.

Estratégia de Remediação: Padrão. Este *cluster* constitui o *backlog* de remediação, que deve ser tratado por meio dos processos contínuos de gestão de vulnerabilidades, após os perfis mais urgentes (Crítico e Alto) terem sido endereçados. Alinha-se ao *bucket* "Attend"(Atender) do SSVc (16).

5.3.4 Perfil 4: Gigante Adormecido (Cluster 3)

O Perfil "Gigante Adormecido"(Cluster 3) é, talvez, o *insight* estratégico mais importante revelado pelo *clustering*. É definido por um Impacto (NCS) médio de 0,44 e uma Probabilidade (EPSS) média de 0,07.

A característica deste perfil é a sua assimetria. Ele agrupa vulnerabilidades com baixíssima probabilidade de exploração (EPSS de 0,07), o que significa que são amplamente ignoradas pelas ferramentas de ataque automatizadas e pela comunidade pública de *exploits*. No entanto, elas residem em sistemas de impacto moderado (NCS 0,44), mais altos que os do Perfil de Volume.

Este perfil expõe a maior falha da priorização tradicional: o CVSS provavelmente classificaria muitas destas falhas como "Críticas"(devido a fatores técnicos), enquanto o EPSS (usado sozinho) as classificaria como "Baixas"(devido à baixa probabilidade). Nenhuma das métricas, por si só, captura a natureza deste risco: uma ameaça latente. Se um novo código de *exploit* para uma dessas vulnerabilidades for publicado, seu *score* EPSS aumentará drasticamente, e ela será reclassificada instantaneamente do Perfil "Gigante Adormecido"para o Perfil "Risco Crítico".

Estratégia de Remediação: Monitoramento e Planejamento. A remediação não é urgente, dada a baixa probabilidade de exploração. A estratégia correta é "Track"(Rastrear) (16): as vulnerabilidades devem ser catalogadas, e o *score* EPSS deve ser monitorado continuamente. Qualquer mudança significativa na probabilidade deve acionar um alerta para remediação imediata.

5.3.5 Perfil 5: Risco Mínimo (Cluster 4)

O Perfil "Risco Mínimo"(Cluster 4) é o grupo de menor prioridade, definido por uma Probabilidade (EPSS) média de 0,02 e um Impacto (NCS) médio de 0,07.

Ambas as dimensões do risco são as mais baixas do *dataset*. Estas vulnerabilidades representam o "ruído"de fundo da gestão de vulnerabilidades: falhas com probabilidade de exploração quase nula, localizadas em ativos de baixo ou nenhum impacto para o negócio.

Estratégia de Remediação: Aceitação de Risco ou "Menor Esforço". Corrigir estas vulnerabilidades representa uma alocação ineficiente de recursos escassos que deveriam estar focados nos Perfis Críticos e Altos. Este *cluster* é o principal candidato à aceitação formal do risco ou ao *bucket* "Defer"(Adiar) do SSVC (16), sendo tratado apenas quando todos os demais riscos tiverem sido mitigados.

5.4 VALIDAÇÃO COMPLEMENTAR E DETECÇÃO DE ANOMALIAS (DBSCAN)

Enquanto o K-Means foi eficaz na partição dos dados em perfis de risco (conforme Seção 5.3), a metodologia também empregou o algoritmo DBSCAN (*Density-Based Spatial Clustering of Applications with Noise*) (45) com um duplo propósito: (1) validar a estrutura de densidade dos agrupamentos e (2) identificar *outliers* (anomalias) (19).

O uso do DBSCAN para a detecção de anomalias é uma prática robusta, conforme validado por estudos como os de Fuhnwi et al. (2023) (25) e de Goswami (2024) (40), que destacam a capacidade do algoritmo de isolar pontos em regiões de baixa densidade que não pertencem a nenhum *cluster* principal. Em primeiro lugar, a aplicação do DBSCAN ao *dataset* (documentada no *notebook* da pesquisa (59)) identificou cinco agrupamentos principais de alta densidade, corroborando, indiretamente, a escolha de $k = 5$ para o K-Means (19).

Mais importante, o DBSCAN classificou 14 pontos de dados como "ruído" ou *outliers* (19). No contexto desta dissertação, estes pontos não são erros estatísticos, mas sim riscos anômalos: vulnerabilidades com combinações únicas de probabilidade e impacto que não se enquadram nos perfis de risco padrão. A Tabela 5.2 apresenta exemplos notáveis dessas anomalias identificadas.

Tabela 5.2: Exemplos de Anomalias (Outliers) Estratégicas Identificadas pelo DBSCAN.

Fonte: Elaborado pelo autor (2025).

CVE ID	EPSS	PPSI	Análise da Anomalia Estratégica
CVE-2003-0033	0,52	0,99	Risco extremo devido ao impacto máximo no negócio. A exploração, embora de probabilidade moderada, seria catastrófica. Representa uma ameaça latente ("bomba-relógio") que o K-Means poderia agrupar, mas que o DBSCAN corretamente isola.
CVE-2024-0195	0,92	0,32	Alvo quase certo de ataques automatizados (EPSS altíssimo). Embora o impacto direto no ativo seja baixo, pode servir como um "ponto de entrada" (<i>foothold</i>) para a rede, e sua exploração é um indicador de comprometimento com alta confiança.

A análise aprofundada desses *outliers* (Tabela 5.2) revela sua importância estratégica. O CVE-2003-0033, por exemplo, combina um impacto extremo (0,99) com uma probabilidade moderada (0,52) (19). Um modelo K-Means, focado em centroides, poderia "esconder" esta vulnerabilidade

em um único *cluster* maior. O DBSCAN, no entanto, a isola corretamente, sinalizando ao gestor de segurança que esta falha específica, embora não esteja sendo explorada em volume, representa um risco singular de "bomba-relógio" que exige um plano de mitigação dedicado. A capacidade de isolar esses casos demonstra o valor da análise de densidade não apenas para categorizar ameaças em massa, mas também para sinalizar riscos estratégicos singulares (19).

5.5 ANÁLISE DE IMPACTO QUANTITATIVO: A REDUÇÃO DA “FADIGA DE ALERTAS”

A validação final da metodologia proposta reside em sua capacidade de resolver o problema prático da "fadiga de alertas" (*alert fatigue*) (13). Conforme discutido no Capítulo 2, a priorização tradicional baseada em CVSS sobrecarrega as equipes de segurança ao classificar uma fração impraticável de vulnerabilidades como "Críticas" ou "Altas" (15).

Para quantificar o impacto da metodologia de perfis de risco, foi gerado um gráfico comparativo (Figura 5.6) que compara a distribuição de vulnerabilidades por nível de criticidade antes (usando a classificação CVSS tradicional) e depois (usando os perfis de risco do K-Means) (19).

A análise da Figura 5.6 apresenta a principal conclusão quantitativa desta dissertação. Na abordagem tradicional (em vermelho, "Antes"), a equipe de segurança enfrenta um "muro" de alertas: somando as categorias "Crítico" (126 vulnerabilidades) e "Alto" (209 vulnerabilidades), a equipe seria instruída a tratar 335 vulnerabilidades como prioridade máxima. Este volume, que corresponde a 79,5% de todo o *dataset*, é a definição exata da "fadiga de alertas": quando quase tudo é classificado como urgente, a priorização perde seu significado prático (19).

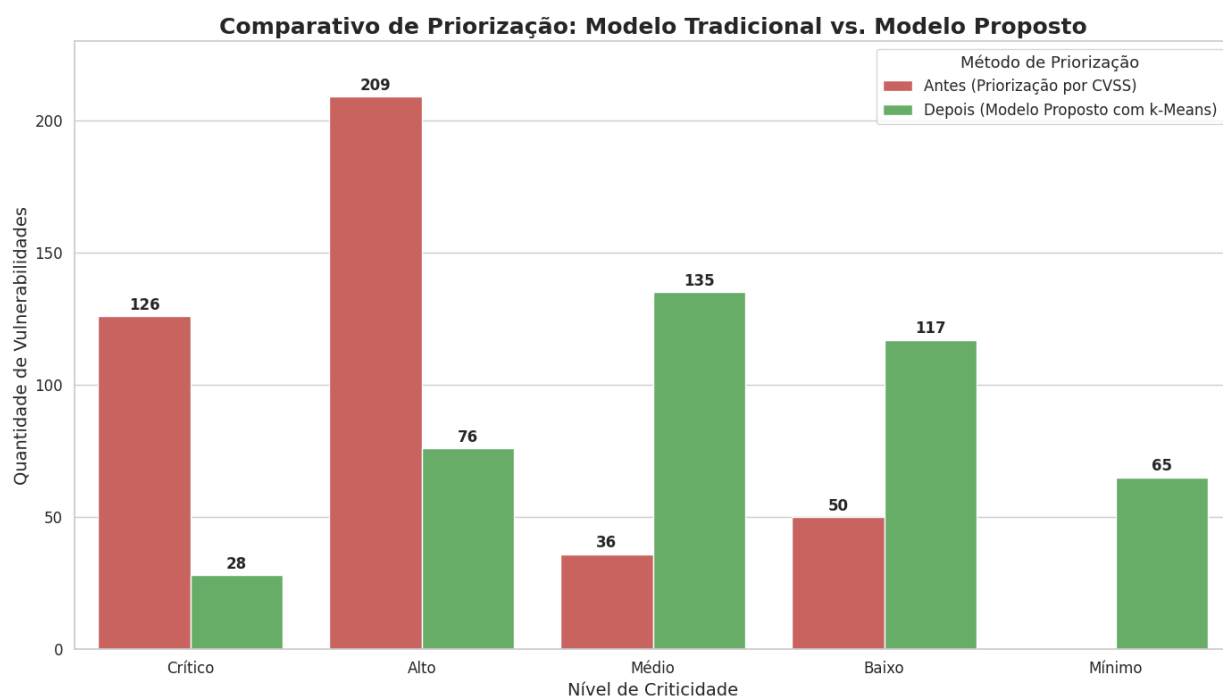


Figura 5.6: Comparativo de Priorização: Modelo Tradicional (CVSS) vs. Modelo Proposto (Perfis de Risco K-Means).

Fonte: Elaborado pelo autor (2025).

Em contrapartida, a abordagem de Perfis de Risco (em verde, "Depois") fornece um foco estratégico e gerenciável. Ao agrupar as vulnerabilidades com base no risco real (Probabilidade x Impacto), os perfis de prioridade mais alta ("Crítico" e "Alto") somam apenas 104 vulnerabilidades (28 e 76, respectivamente).

A aplicação desta metodologia resultou, portanto, em uma redução de 69% no volume de alertas de alta prioridade (de 335 para 104). Esta otimização não significa ignorar as demais 231 vulnerabilidades, mas sim reclassificá-las corretamente em perfis de menor urgência (como "Médio", "Baixo" e "Mínimo" ou "Gigantes Adormecidos" e "Risco Mínimo"), permitindo que as equipes de segurança concentrem seus recursos escassos de forma eficaz nas 104 ameaças que representam o risco mais iminente e impactante para o negócio (19).

A redução do escopo de atuação prioritária valida a viabilidade operacional do modelo em centros de comando governamentais. A abordagem permite que as equipes integrantes do Subsistema de Inteligência de Segurança Pública (SISP) superem o modelo reativo de tratamento de incidentes e adotem uma postura de antecipação. Ao restringir o volume de intervenções imediatas ao estritamente necessário, o modelo bidimensional otimiza o emprego dos ativos de segurança cibernética e garante a resiliência contínua dos serviços críticos prestados à sociedade.

5.6 VISUALIZAÇÃO CONSOLIDADA DA MATRIZ DE RISCO

Para consolidar todos os achados da modelagem, a Figura 5.7 apresenta a Matriz de Risco final. Este gráfico de dispersão combina os resultados das duas técnicas de *clustering* em uma única ferramenta de decisão estratégica.

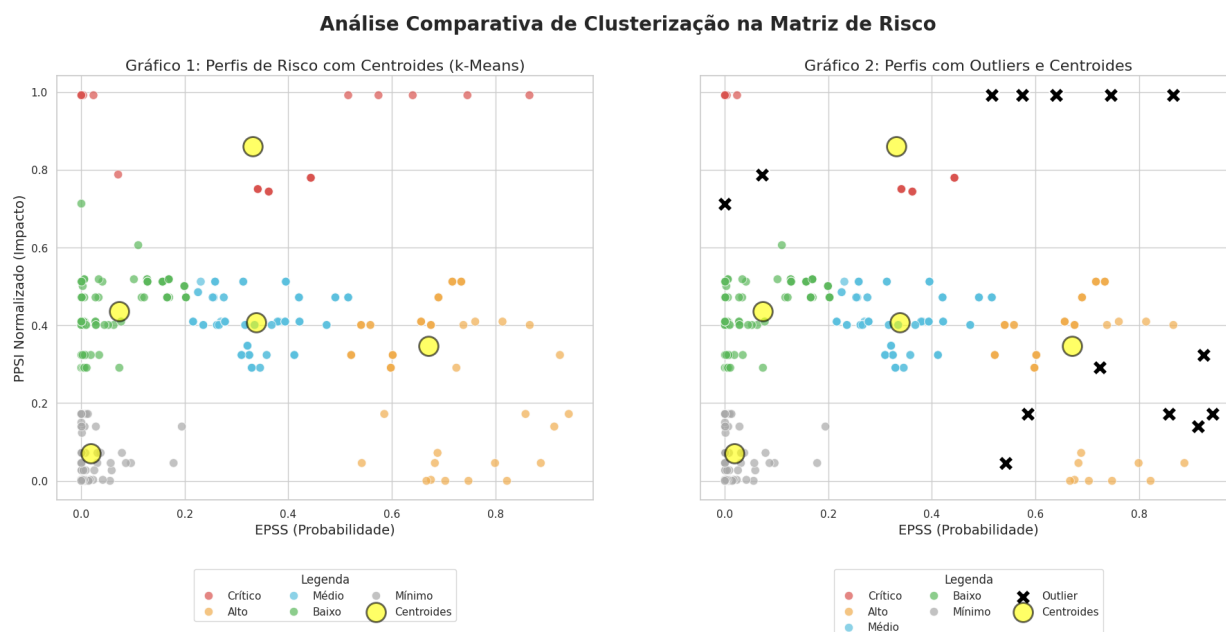


Figura 5.7: Matriz de Risco Consolidada: Perfis de Risco (K-Means, por cor), Centroides (círculos amarelos) e Anomalias (DBSCAN, marcadores 'X').

Fonte: Elaborado pelo autor (2025).

A visualização consolidada na Figura 5.7 constitui a síntese da metodologia. O "Gráfico 1"(esquerda) ilustra os cinco perfis de risco descobertos pelo K-Means, representados por cores distintas, com seus respectivos centroides (círculos amarelos), indicando o "centro de gravidade"de cada perfil.

O "Gráfico 2"(direita) sobrepõe a esta visualização os *outliers* (anomalias) identificados pelo DBSCAN, marcados com um 'X'. Esta combinação permite ao gestor de segurança uma interpretação imediata do cenário de risco. Ele pode visualizar não apenas os agrupamentos principais (os perfis de risco K-Means), mas também as ameaças atípicas e isoladas (os *outliers* do DBSCAN) que exigem investigação individualizada (19). A matriz traduz a análise complexa de dados em um mapa de priorização prático e compreensível, atendendo ao objetivo central desta pesquisa.

5.7 VALIDAÇÃO EM LARGA ESCALA E TESTES DE ESTRESSE

A eficácia de uma ferramenta analítica voltada à gestão de infraestruturas estatais requer a comprovação de sua operabilidade em ambientes de grande escala. Esta seção descreve os achados da aplicação dos algoritmos no cenário de estresse estabelecido na fase metodológica, incluindo a validação estatística da amostra e a avaliação do desempenho dos modelos de agrupamento.

5.7.1 Validação Estatística dos Dados Sintéticos

A confiabilidade dos testes de estresse depende da fidedignidade da base de dados gerada artificialmente. Para avaliar a semelhança entre a amostra sintética (100.000 registros) e o ambiente original (421 registros), aplicou-se o teste não paramétrico de Kolmogorov-Smirnov. O teste mensura a distância máxima entre as funções de distribuição cumulativa das duas amostras, expressa pela estatística D.

Os cálculos revelaram uma aproximação geométrica expressiva entre as variáveis dimensionais de risco. A estatística D para a variável de severidade técnica (CVSS) foi de 0,189. Para a probabilidade de exploração (EPSS), a distância D foi de 0,257. Valores próximos a zero confirmam estatisticamente a sobreposição das distribuições e o sucesso da arquitetura generativa na mimetização do ambiente real.

Para consolidar a análise, extraiu-se o relatório de qualidade multivariada. A Tabela 5.3 apresenta o índice KSComplement por atributo e compara as estatísticas descritivas (médias) entre o mundo real e o cenário sintético. O escore global de fidelidade estrutural da base de dados atingiu 81,55%.

Tabela 5.3: Comparativo estatístico e métricas de fidelidade (Mundo Real vs. Sintético).

Fonte: elaborada pelo autor (2025).

Variável Avaliada	Média (Real)	Média (Sintético)	KSComplement
Severidade (CVSS)	7,13	7,15	0,80
Probabilidade (EPSS)	0,25	0,29	0,74
Criticidade (NCS)	1,68	1,71	0,85
Impacto Normalizado (PPSI)	0,53	0,52	0,77

A conservação dos tensores matemáticos comprova a estabilidade metodológica. A média do impacto no negócio (PPSI), por exemplo, apresentou uma oscilação desprezível de 0,53 para 0,52, garantindo que o algoritmo de agrupamento opere com os mesmos pressupostos lógicos em ambos os cenários.

Para atestar que a inteligência artificial não apenas replicou os dados isoladamente, mas também preservou as regras de negócio intrínsecas ao ambiente, analisou-se o comportamento multivariado. A Figura 5.8 exibe as matrizes de correlação de ambos os cenários. A manutenção do coeficiente de correlação (0,19) no ambiente sintético comprova estruturalmente que a independência entre o risco técnico (EPSS) e o impacto no negócio (NCS) foi preservada, validando a base para a clusterização bidimensional.

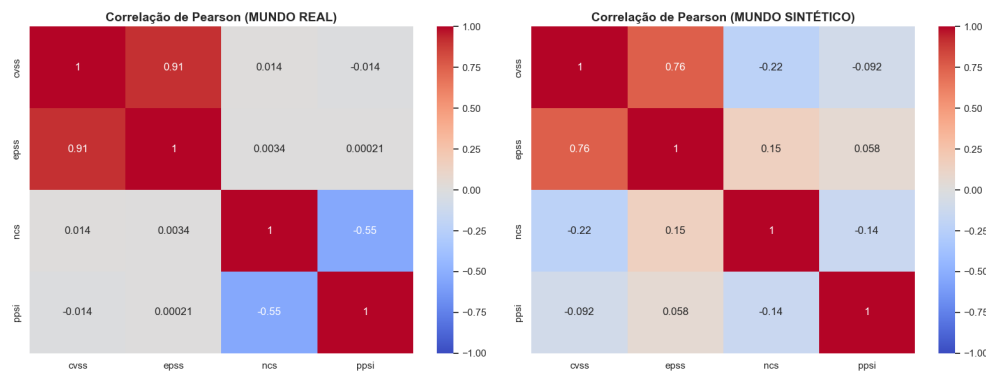


Figura 5.8: Preservação multivariada das matrizes de correlação no ambiente gerado por IA.

Fonte: Elaborada pelo autor (2025).

A comprovação visual da aderência probabilística é apresentada na Figura 5.9. A sobreposição das curvas de densidade (KDE) entre as distribuições originais e as amostras geradas pela arquitetura CTGAN evidencia a capacidade da rede neural em capturar as nuances das variáveis contínuas, sem incidência de colapso de modo (*mode collapse*).

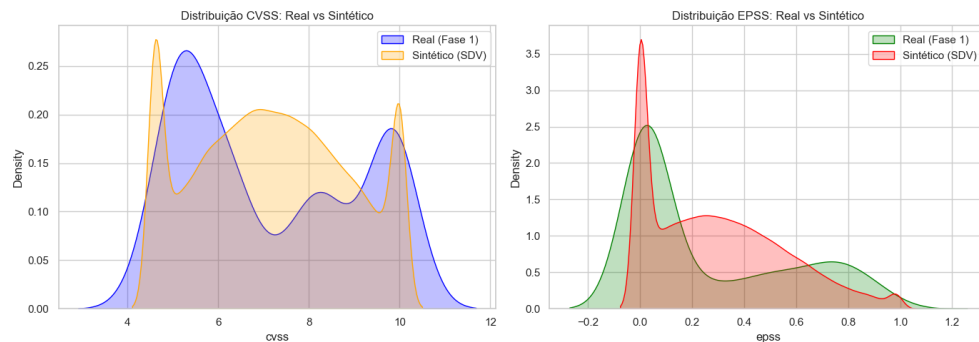


Figura 5.9: Comparação da densidade de probabilidade entre o ambiente real e o sintético.

Fonte: elaborada pelo autor (2025).

5.7.2 Desempenho Algorítmico no Cenário de Big Data e Anomalias

Com a integridade da base validada, a amostra consolidada de 101.000 registros foi submetida ao fluxo de agrupamento. O procedimento avaliou a mitigação da fadiga de alertas em proporções

corporativas e a capacidade de identificar comportamentos anômalos inseridos propositalmente no ambiente.

A execução do algoritmo K-Means processou todo o volume de dados de forma eficiente e preservou a coerência estrutural da segmentação em cinco perfis de risco. A distribuição quantitativa acompanhou a tendência observada no ambiente controlado. A matriz realocou milhares de vulnerabilidades com alta pontuação técnica (CVSS), porém, com probabilidade nula, para os quadrantes de menor prioridade. A ação comprovou a resiliência do modelo bidimensional na filtragem e redução massiva do escopo de atuação prioritária em cenários de dados volumosos.

A eficácia computacional do K-Means na mitigação da fadiga de alertas em proporções corporativas é demonstrada na Figura 5.10. O gráfico evidencia o efeito de correção de viés (*bias mitigation*), demonstrando que, mesmo diante de um volume de 101.000 instâncias contendo anomalias propositalis, o algoritmo não sofreu distorção em seus centroides, realocando de forma consistente as falhas de alto risco técnico, porém de baixa probabilidade, para perfis não prioritários.

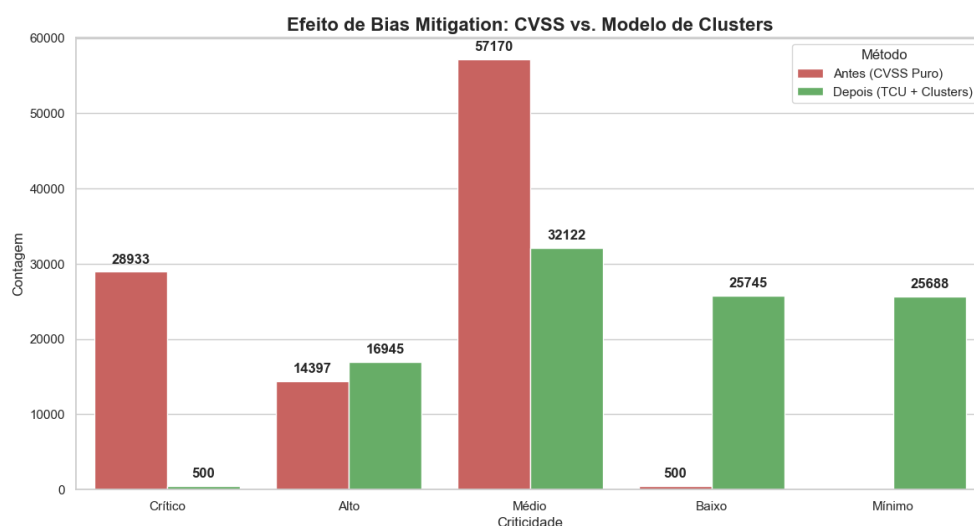


Figura 5.10: Comparativo de mitigação da fadiga de alertas sob estresse massivo (101.000 instâncias).

Fonte: elaborada pelo autor (2025).

Em paralelo, o algoritmo DBSCAN rastreou com precisão o espaço vetorial. A varredura por densidade isolou os casos anômalos injetados artificialmente (vulnerabilidades configuradas com CVSS máximo e EPSS nulo, ou com CVSS zero e EPSS absoluto), rotulando-os como ruído espacial. A análise de densidade classificou um subconjunto restrito de vulnerabilidades extremas explicitamente como “ruído” (*outliers*). Esses pontos representam perfis de risco excepcionais, como ameaças de impacto quase máximo com probabilidade moderada (ex: CVE-2003-0033), ou falhas de probabilidade de exploração altíssima em ativos de menor valor (ex: CVE-2024-0195), que comumente servem como portas de entrada para movimentação lateral.

A identificação e isolamento desses *outliers* comprovam a capacidade do modelo de capturar casos que seriam erroneamente agrupados e mascarados pela força da média em modelos de particionamento baseados unicamente em centroides. Estrategicamente, isso atesta a dupla eficácia da solução proposta: permite que a equipe de segurança utilize o K-Means para o tratamento automatizado e estruturado da massa de alertas em larga escala, enquanto a análise humana especializada é direcionada cirurgicamente apenas para as anomalias extremas apontadas como ruído discrepante pelo DBSCAN.

A visualização final do cenário de estresse é projetada na Figura 5.11. O mapeamento espacial reflete a estrutura global dos 101.000 registros. No painel à direita, destaca-se a atuação do algoritmo DBSCAN, que mapeou com sucesso as regiões de densidade isoladas e atribuiu o rótulo de anomalia espacial estritamente aos cenários anômalos inseridos artificialmente.

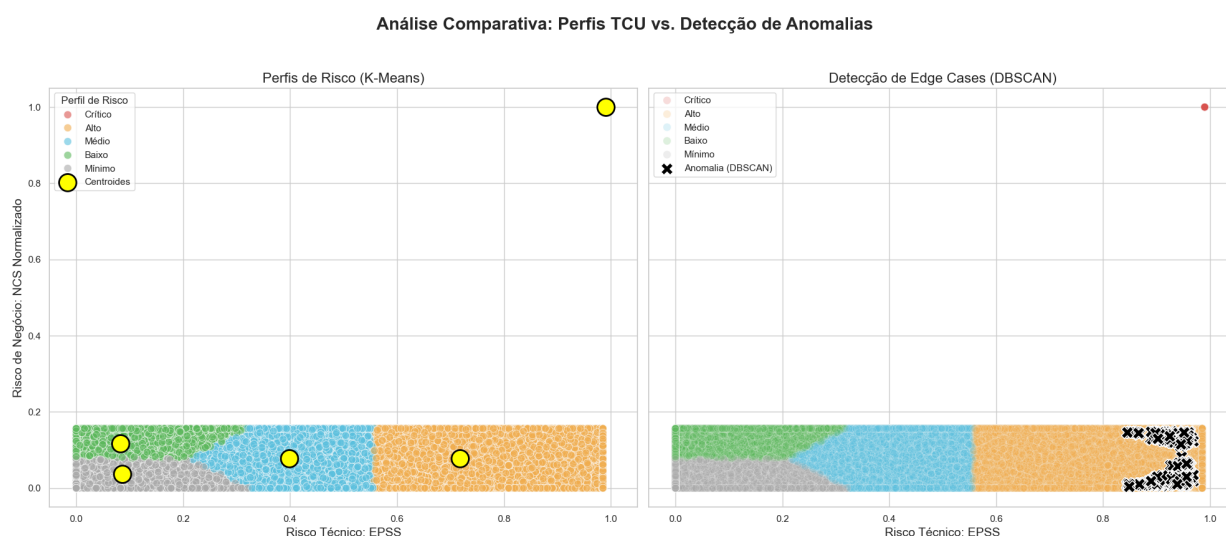


Figura 5.11: Matriz de Risco consolidada no cenário de Big Data governamental e de anomalias extremas.

Fonte: elaborada pelo autor (2025).

6 CONCLUSÃO E TRABALHOS FUTUROS

A gestão eficaz de vulnerabilidades é um pilar central da cibersegurança moderna. Contudo, a metodologia tradicional de priorização, baseada quase exclusivamente na severidade técnica (CVSS), mostrou-se insuficiente para o cenário atual. Ao classificar um volume impraticável de vulnerabilidades como "Críticas" ou "Altas", o CVSS gera uma "fadiga de alertas" que obscurece o risco real e pulveriza os recursos limitados das equipes de segurança (16, 15).

Esta dissertação propôs e validou uma metodologia alternativa de priorização inteligente, fundamentada em um modelo de risco bidimensional que reflete o contexto organizacional. A abordagem substitui a pontuação de severidade estática por um portfólio de risco que equilibra a Probabilidade de exploração (quantificada pelo EPSS (15)) e o Impacto no negócio (quantificado pela NCS, do *framework* PPSI/TCU (22, 23)).

6.1 CONTRIBUIÇÕES DA PESQUISA

A principal contribuição deste trabalho foi a demonstração empírica de que a aplicação de *Machine Learning* não supervisionado a este modelo de risco bidimensional é capaz de transformar o "ruído" da gestão de vulnerabilidades em inteligência estratégica acionável.

Ao aplicar o algoritmo K-Means, a metodologia identificou cinco perfis de risco distintos e qualitativamente relevantes (detalhados no Capítulo 5): "Risco Crítico", "Ameaça de Volume", "Risco Médio (Padrão)", "Gigante Adormecido" e "Risco Mínimo (Ruído)". Cada perfil exige uma estratégia de remediação distinta, que vai da correção imediata ao monitoramento ou à aceitação do risco.

O impacto quantitativo desta abordagem foi significativo. Conforme demonstrado na Figura 5.6, a metodologia de perfis de risco reduziu em 69,2% o volume de alertas de alta prioridade em comparação com o método CVSS. Isso representa uma otimização drástica, permitindo que as organizações foquem seus recursos nas 104 vulnerabilidades (24,5% do total) que realmente importam, em vez de se afogarem em 335 alertas (79,5% do total) (19).

Adicionalmente, a aplicação complementar do DBSCAN validou a estrutura de densidade dos agrupamentos e, mais importante, isolou 14 riscos anômalos (*outliers*), como o CVE-2003-0033 (19), demonstrando a capacidade da metodologia de sinalizar ameaças estratégicas singulares que seriam mascaradas por modelos de média, como o K-Means.

A pesquisa elevou o rigor de sua validação ao submeter a metodologia a um teste de estresse com 101.000 registros. A utilização de Inteligência Artificial Generativa, especificamente a arquitetura de redes neurais condicionais (CTGAN), demonstrou ser uma solução eficaz para simular o ambiente de vulnerabilidades de um órgão governamental. A técnica alcançou alta fidelidade estatística em relação ao mundo real, confirmada pelo teste de Kolmogorov-Smirnov. Nesse cenário de larga escala, o algoritmo K-Means atestou sua escalabilidade computacional e manteve a coerência matemática na redução de alertas prioritários. Simultaneamente, o algoritmo DBSCAN atuou com precisão cirúrgica na identificação de casos extremos inseridos artificialmente na base sintética. Essa abordagem de enriquecimento metodológico comprovou a robustez da solução em ambientes corporativos de alto volume de dados. O procedimento garantiu a reprodutibilidade científica do experimento e preservou a Segurança Operacional (OPSEC) das infraestruturas críticas do Estado, consolidando o modelo como uma ferramenta madura para adoção governamental.

Adicionalmente, a fundamentação da escolha algorítmica, por meio de estudo comparativo validado na literatura, e o alinhamento da pesquisa às diretrizes do Subsistema de Inteligência de Segurança Pública (SISP) atestam a aplicabilidade institucional da proposta. O modelo transcende a mitigação técnica de falhas e consolida-se como um ativo de Inteligência Estratégica Cibernética, capaz de otimizar o emprego de recursos humanos e antecipar ameaças nas infraestruturas críticas do Estado.

6.2 LIMITAÇÕES DA PESQUISA

A transparência acadêmica exige o reconhecimento das limitações metodológicas que definiram o escopo desta pesquisa. As limitações foram intencionais e necessárias para viabilizar a análise e basearam-se nas justificativas detalhadas no Capítulo 4.

A primeira e principal limitação foi a incompatibilidade entre CWE e CVE. As ferramentas de varredura de ativos (*scanners* como OWASP ZAP e OpenVAS) identificam fraquezas (CWEs) (60, 61), mas a métrica de probabilidade (EPSS) é indexada por vulnerabilidades (CVEs). A ausência de um mapeamento direto e confiável entre essas duas taxonomias inviabilizou o uso de um *dataset* de *scans* reais para a modelagem de probabilidade. Isso exigiu a associação simulada de CVEs reais (coletados de fontes públicas, como o NVD) aos ativos reais identificados, criando um banco de dados de Gestão de Vulnerabilidades maduro.

A segunda limitação foi a simulação do Impacto (NCS). Embora a métrica NCS seja baseada em um *framework* real e público (PPSI/TCU) (22, 23), a criticidade de cada ativo descoberto (Fase 1 da Metodologia) não era pública. A NCS foi, portanto, simulada e atribuída aos ativos para criar um cenário de risco diversificado e equilibrado, necessário para que os algoritmos de *clustering*

pudessem identificar perfis distintos.

A terceira limitação foi o equilíbrio do *dataset*. Ambientes reais frequentemente possuem ativos (ex.: sistemas legados) com um número desproporcional de vulnerabilidades. Para evitar que esses *outliers* de ativos dominassem o cálculo da distância euclidiana do K-Means e distorcessem os perfis de risco, foi aplicada uma amostragem aleatória com limite superior de vulnerabilidades por ativo.

Em suma, o *dataset* de 421 vulnerabilidades analisado é um ambiente controlado e representativo (um "estudo de caso simulado"), e os perfis de risco identificados são específicos deste cenário. Apesar destas limitações operacionais inerentes à construção do modelo para esta pesquisa, em aplicações práticas no cenário corporativo maduro, algumas das fases iniciais da metodologia, como o *footprinting* via OSINT, podem ser suprimidas. Organizações estruturadas já possuem processos de inventário que fornecem a lista de ativos, os CVEs associados e as notas de criticidade previamente calculadas, permitindo a aplicação direta das fases analíticas e de *clustering* do modelo proposto (55).

6.3 TRABALHOS FUTUROS

As contribuições e limitações desta dissertação abrem um vasto campo para pesquisas futuras, voltadas a operacionalizar e a avançar a metodologia de perfis de risco.

Uma evolução natural e imediata deste trabalho é a aplicação de *Machine Learning* supervisionado. Os cinco perfis de risco descobertos pelo K-Means (ex.: "Risco Crítico", "Gigante Adormecido") podem agora ser utilizados como "rótulos"(*labels*) para treinar um modelo de classificação (como *Random Forest* ou XGBoost). Tal modelo seria capaz de classificar automaticamente novas vulnerabilidades em tempo real, à medida que fossem descobertas, transformando esta metodologia de análise exploratória em uma ferramenta de classificação preditiva e operacional (66).

Uma segunda linha de pesquisa seria a exploração de algoritmos de *clustering* mais avançados para a descoberta de perfis. Esta dissertação utilizou implementações clássicas do K-Means e DBSCAN. A literatura recente, no entanto, sugere melhorias passíveis de teste. Cao et al. (2025) (26), por exemplo, propõem um "K-Means com peso dinâmico" que poderia refinar a separação entre *clusters*. De forma ainda mais avançada, El Maouaki et al. (2024) (67) demonstraram resultados promissores com algoritmos de *Quantum Machine Learning* (QML), como o QCSWAP K-means, para agrupar dados de cibersegurança (especialmente o catálogo CISA KEV). Testar se estes algoritmos quânticos produzem perfis de risco mais coesos do que os do K-Means clássico seria uma contribuição de ponta.

Finalmente, o maior desafio para trabalhos futuros é a aplicação desta metodologia em um ambiente organizacional real. Isso exigiria superar as limitações aqui identificadas, envolvendo (1) a colaboração com as equipes de negócio para obter os *scores* reais de NCS dos ativos e (2) o desenvolvimento de um mapeamento (provavelmente probabilístico) entre CWEs (dos *scanners*) e CVEs (do EPSS) para automatizar o *pipeline* de dados. A validação do modelo de perfis de risco com dados reais, ruidosos e em larga escala é o teste definitivo de sua eficácia e valor estratégico para a gestão de vulnerabilidades.

REFERÊNCIAS BIBLIOGRÁFICAS

- 1 SOUSA, F. L. d. *Transformação Digital no Contexto da Inteligência de Estado: Análise e Mitigação das Vulnerabilidades do Documento Digital*. Dissertação (Dissertação de Mestrado Profissional) — Universidade de Brasília (UnB), Brasília-DF, Brasil, 2023.
- 2 OFOEGBU, K. D. O.; OSUNDARE, O. S.; IKE, C. S.; FAKEYEDE, O. G.; IGE, A. B. Data-driven cyber threat intelligence: Leveraging behavioral analytics for proactive defense mechanisms. *Computer Science & IT Research Journal*, v. 4, n. 3, p. 502–524, 2023.
- 3 VAL, O. O.; KOLADE, T. M.; GBADEBO, M. O.; SELESI-AINA, O.; OLATEJU, O. O.; OLANIYI, O. O. Strengthening cybersecurity measures for the defense of critical infrastructure in the united states. *SSRN Electronic Journal / Asian Journal of Research in Computer Science*, v. 17, n. 11, p. 25–45, 2024.
- 4 SILVA, T. C. d. *3SW: Um Modelo Adaptado do CIS Controls v8.1 para Mitigação de Ameaças a Servidores Web*. Dissertação (Dissertação de Mestrado Profissional) — Universidade de Brasília (UnB), Brasília-DF, Brasil, 2025.
- 5 ALVES, R. S. *Proposta de uma Linha de Base de Controles de Segurança da Informação para Mitigação dos Riscos de Negligência do Poder Judiciário Brasileiro*. Dissertação (Dissertação de Mestrado Profissional) — Universidade de Brasília (UnB), Brasília-DF, Brasil, 2024.
- 6 MANOHARAN, A.; SARKER, M. Revolutionizing cybersecurity: Unleashing the power of artificial intelligence and machine learning for next-generation threat detection. *International Research Journal of Modernization in Engineering Technology and Science (IRJMETs)*, v. 04, n. 12, p. 2151–2164, 2022.
- 7 MOHAMMED, A. Transforming soc operations: Harnessing the power of ai and ml for enhanced threat detection. *International Journal of Research Culture Society*, v. 08, n. Special Issue 35, 2024.
- 8 The MITRE Corporation. *CVE - Common Vulnerabilities and Exposures*. 2024. <<https://cve.mitre.org/>>. Acessado em: 10/11/2025.
- 9 KOMARAGIRI, V. B.; EDWARD, A. Ai-driven vulnerability management and automated threat mitigation. *International Journal of Scientific Research and Management (IJSRM)*, v. 10, n. 10, p. EC–2022–980–998, 2022.
- 10 JACOBS, J.; ROMANOSKY, S.; SUCIU, O.; EDWARDS, B.; SARABI, A. Enhancing vulnerability prioritization: Data-driven exploit predictions with community-driven insights. In: *2023 IEEE European Symposium on Security and Privacy Workshops (EuroS&PW)*. [S.l.]: IEEE, 2023. p. 194–206.
- 11 OWASP. *ZAP Scanning Report*. [S.l.], 2025. Relatório de varredura gerado em 29/06/2025. Disponível em: <<https://www.zaproxy.org/>>.

- 12 Greenbone Networks. *OpenVAS Vulnerability Assessment Report*. [S.l.], 2025. Relatório de varredura de vulnerabilidades. Disponível em: <<https://www.openvas.org/>>.
- 13 TARIQ, S.; CHHETRI, M. B.; NEPAL, S.; PARIS, C. Alert fatigue in security operations centres: Research challenges and opportunities. *ACM Computing Surveys*, ACM New York, NY, v. 57, n. 9, p. 1–38, 2025.
- 14 FIRST.org, Inc. *Common Vulnerability Scoring System v3.1: Specification Document*. 2019. <<https://www.first.org/cvss/v3.1/specification-document>>. Acessado em: 10/11/2025.
- 15 JACOBS, J.; ROMANOSKY, S.; SUCIU, O.; EDWARDS, B.; SARABI, A. Enhancing vulnerability prioritization: Data-driven exploit predictions with community-driven insights. In: *2023 IEEE European Symposium on Security and Privacy Workshops (EuroS&PW)*. [S.l.: s.n.], 2023.
- 16 SPRING, J. M.; HATLEBACK, E.; HOUSEHOLDER, A. D.; MANION, A.; SHICK, D. *Prioritizing Vulnerability Response: A Stakeholder-Specific Vulnerability Categorization (Version 2.0)*. [S.l.], 2021. CMU/SEI-2021-TR-006.
- 17 ZENG, Z.; YANG, Z.; HUANG, D.; CHUNG, C.-J. Licality-likelihood and criticality: Vulnerability risk prioritization through logical reasoning and deep learning. *IEEE Transactions on Network and Service Management*, v. 19, n. 2, p. 1746–1760, 2022.
- 18 YOON, S.-S.; KIM, D.-Y.; KIM, G.-G.; EUOM, I.-C. Vulnerability assessment based on real world exploitability for prioritizing patch applications. In: *2023 7th Cyber Security in Networking Conference (CSNet)*. [S.l.: s.n.], 2023.
- 19 SILVA, M. L. M. d.; NETO, J. S. Uma análise comparativa de algoritmos de clusterização para a definição de perfis de risco de vulnerabilidades. *REVISTA DELOS*, v. 19, n. 79, p. e9206, abr. 2026. Disponível em: <<https://ojs.revistadelos.com/ojs/index.php/delos/article/view/9206>>.
- 20 International Organization for Standardization (ISO). *ISO/IEC 27005:2022 - Information security, cybersecurity and privacy protection – Guidance on managing information security risks*. [S.l.], 2022.
- 21 FORCE, J. T. *NIST Special Publication 800-30 Rev. 1: Guide for Conducting Risk Assessments*. Gaithersburg, MD, USA, 2012.
- 22 Ministério da Gestão e da Inovação em Serviços Públicos (MGI). *Guia Operacional - Framework do PPSI*. 2024. <https://www.gov.br/governodigital/pt-br/seguranca-e-protecao-de-dados/ppsi/guia_framework_psi.pdf>. Acessado em: 10/11/2025.
- 23 Tribunal de Contas da União (TCU). *Acórdão 1.889/2020-TCU-Plenário*. 2020. <<https://pesquisa.apps.tcu.gov.br/#/documento/acordao-completo/1889-2020/PLENARIO/GSI/>>. Acessado em: 10/11/2025.
- 24 KHAOULA, R.; MOHAMED, M. K-means-dist: A novel approach for enhanced cybersecurity clustering using combined distance metrics. In: *2023 10th International Conference on Wireless Networks and Mobile Communications (WINCOM)*. [S.l.: s.n.], 2023.

- 25 FUHNWI, G. S.; AGBAJE, J. O.; OSHINUBI, K.; PETER, O. J. An empirical study on anomaly detection using density-based and representative-based clustering algorithms. *Journal of the Nigerian Society of Physical Sciences*, v. 5, 2023.
- 26 CAO, X.; ZHANG, J.; ZHENG, P.; BU, S.; CHENG, K. Global cybersecurity situation assessment combining k-means clustering and cybersecurity index. In: *2025 5th International Symposium on Computer Technology and Information Science (ISCTIS)*. [S.l.: s.n.], 2025.
- 27 KASHTALIAN, A.; SOCHOR, T. K-means clustering of honeynet data with unsupervised representation learning. In: *Proceedings of the 2nd International Workshop on Intelligent Information Technologies and Systems of Information Security (IntelITSIS 2021)*. [S.l.: s.n.], 2021.
- 28 NETO, A. C.; PREGARDIER, R. C.; SANTOS, C. R. P. dos; FÜLBER-GARCIA, V.; SILVA, L. A. L. Aprimorando a detecção de apTs com geração de dados sintéticos baseada em gan-transformers. In: *Anais do Simpósio Brasileiro de Cibersegurança (SBSeg 2025)*. Foz do Iguaçu: SBC, 2025. p. 131–146.
- 29 SARMIN, F. J.; SARKAR, A. R.; WANG, Y.; MOHAMMED, N. Synthetic data: revisiting the privacy-utility trade-off. *International Journal of Information Security*, v. 24, n. 4, p. 156, 2025.
- 30 Associação Brasileira de Normas Técnicas (ABNT). *ABNT NBR ISO 31000:2018 - Gestão de riscos â€” Diretrizes*. Rio de Janeiro, RJ, Brasil, 2018.
- 31 The MITRE Corporation. *CWE - Common Weakness Enumeration*. 2024. <<https://cwe.mitre.org/>>. Acessado em: 10/11/2025.
- 32 VALADARES, D. C. G.; PERKUSICH, A.; SANTOS, D. F. d. S. Measuring security with a score system. In: *2024 11th International Conference on Future Internet of Things and Cloud (FiCloud)*. [S.l.: s.n.], 2024.
- 33 FIRST.org, Inc. *EPSS - Frequently Asked Questions (FAQ)*. 2024. <<https://www.first.org/epss/faq>>. Acessado em: 10/11/2025.
- 34 ZOTTMANN, C. E. M.; GEORG, M. A. C.; ALVES, R. S.; SILVA, M. A. da; NUNES, R. R. Proposta de metodologia para avaliação de riscos de privacidade para “rgãos do poder judiciário no brasil. In: *Anais do Encontro de Administração da Justiça (ENAJUS)*. Brasília-DF, Brasil: [s.n.], 2023.
- 35 FERREIRA, L. V. A. *O papel da Auditoria Interna na Gestão de Riscos Cibernéticos em Instituições Financeiras Brasileiras - Estudo sob a perspectiva das trãs linhas*. Dissertação (Dissertação de Mestrado Profissional) — Universidade de Brasília (UnB), Brasília-DF, Brasil, 2024.
- 36 CARDOSO, L. H. F. *Gestão do Risco Cibernético â€” Implantação ADS-B no âmbito do SISCEAB por Meio do Método de Gerenciamento de Riscos â€” Segurança Operacional (GRSO)*. Dissertação (Dissertação de Mestrado Profissional) — Universidade de Brasília (UnB), Brasília-DF, Brasil, 2023.

- 37 SOUSA, R. B. G. de. Inteligência estratégica no espaço cibernético na segurança pública. *Revista de Inteligência de Segurança Pública (RISP)*, Rio de Janeiro, v. 7, n. 1, p. 20–31, 2024.
- 38 SILVA, F. A. da. A reestruturação do sistema brasileiro de inteligência (sisbin) e suas implicações para o subsistema de inteligência de segurança pública (sisp). *Revista de Inteligência de Segurança Pública (RISP)*, Rio de Janeiro, v. 7, n. 1, p. 32–44, 2024.
- 39 SANTOS, D. C. dos. *A priorização dos incidentes cibernéticos como estratégia de tratamento e resposta a incidentes cibernéticos. Ensaio Acadêmico*. Rio de Janeiro: Escola Superior de Guerra (ESG), 2025.
- 40 GOSWAMI, M. J. Ai-based anomaly detection for real-time cybersecurity. *International Journal of Research and Review Techniques (IJRRT)*, v. 3, n. 1, p. 45–53, 2024.
- 41 MAMUN, A. A.; HOSSAIN, M. S.; RISHAD, S. M. S. I.; RAHMAN, M. M.; TISHA, S. A.; SHAKIL, F.; CHOUDHURY, M. Z. M. E.; DAS, A. C.; DAS, R.; SULTANA, S. Machine learning for stock market security measurement: A comparative analysis of supervised, unsupervised, and deep learning models. *THE AMERICAN JOURNAL OF ENGINEERING AND TECHNOLOGY*, v. 06, n. 11, p. 63–76, 2024.
- 42 MACQUEEN, J. B. Some methods for classification and analysis of multivariate observations. In: *Proceedings of the Fifth Berkeley Symposium on Mathematical Statistics and Probability, Volume 1: Statistics*. [S.l.]: University of California Press, 1967. p. 281–297.
- 43 BANSAL, A. Optimizing customer segmentation for enhanced recommendation systems through comparative analysis of k-means, hierarchical clustering, and dbSCAN algorithms. *International Journal of Core Engineering & Management*, v. 7, n. 6, p. 59–71, 2023.
- 44 MOZUMDER, M. A. S.; MAHMUD, F.; SHAK, M. S.; SULTANA, N.; RODRIGUES, G. N.; RAFI, M. A.; FARAZI, M. Z. R.; RAZAUL, K. M.; HASAN, M.; ISLAM, M. R. Optimizing customer segmentation in the banking sector: A comparative analysis of machine learning algorithms. *Journal of Computer Science and Technology Studies*, v. 6, n. 4, p. 31–42, 2024.
- 45 ESTER, M.; KRIEGEL, H.-P.; SANDER, J.; XU, X. A density-based algorithm for discovering clusters in large spatial databases with noise. In: *Proceedings of the Second International Conference on Knowledge Discovery and Data Mining (KDD'96)*. [S.l.]: AAAI Press, 1996. p. 226–231.
- 46 EMAM, K. E.; MOSQUERA, L.; HESSA, R. Synthetic data, synthetic trust: navigating data challenges in the digital revolution. *Journal of Medical Internet Research*, v. 22, n. 11, p. e21303, 2020.
- 47 AL-KHAFAJI, A. et al. Generative ai driven synthetic attack augmentation for enhanced apt detection. *IEEE Access*, v. 12, p. 1–15, 2024.
- 48 LI, J. et al. Generative adversarial local density-based unsupervised anomaly detection. *PLOS One*, v. 17, n. 3, p. e0264101, 2022.

- 49 PANT, Y. R.; LEIGH, L.; RUEDA, J. F. Improving k-means clustering: A comparative study of parallelized version of modified k-means algorithm for clustering of satellite images. *Algorithms*, v. 18, n. 8, p. 532, 2025.
- 50 KRISHNAN, P.; JAIN, K.; BUYYA, R.; VIJAYAKUMAR, P.; NAYYAR, A.; BILAL, M.; SONG, H. Mud-based behavioral profiling security framework for software-defined iot networks. *IEEE Internet of Things Journal*, v. 9, n. 9, p. 6611–6622, 2022.
- 51 AMMARA, D. A.; DING, J.; TUTSCHKU, K. Architectural selection framework for synthetic network traffic: Quantifying the fidelity-utility trade-off. *IEEE Access*, v. 12, p. 1–18, 2024.
- 52 KIM, J. et al. Novel synthetic dataset generation method with privacy-preserving for intrusion detection system. *Applied Sciences*, v. 15, n. 19, p. 10609, 2025.
- 53 NETO, A. C.; PREGARDIER, R. C.; SANTOS, C. R. dos; FULBER-GARCIA, V.; SILVA, L. A. Aprimorando a detecção de apts com geração de dados sintéticos baseada em gan-transformers. In: SBC. *Simpósio Brasileiro de Segurança da Informação e de Sistemas Computacionais (SBSeg)*. [S.l.], 2025. p. 131–146.
- 54 NAWAZ, M.; TAHIRA, S.; YASMIN, A. Generative ai driven synthetic attack augmentation for enhanced intrusion detection using an imbalanced dataset. Preprints, 2025.
- 55 SILVA, M. L. M. d.; NETO, J. S.; NUNES, R. R.; GARCIA, H. Priorização inteligente de vulnerabilidades: uma abordagem baseada em perfis de risco descobertos com machine learning. *REVISTA TECNOLÓGICA - uniFATEC*, v. 17, n. 1, p. 1–21, 2026.
- 56 SILVA, E. G. da; GEORG, M. A. C.; JÚNIOR, L. A. R.; FERREIRA, L. R.; NUNES, R. R. Tomada de decisão em segurança cibernética: Modelo para identificação de joias da coroa. *Revista Ibérica de Sistemas e Tecnologias de Informação*, Associação Ibérica de Sistemas e Tecnologias de Informacao, n. E77, p. 60–73, 2025.
- 57 GIL, A. C. *Como elaborar projetos de pesquisa*. 4. ed. São Paulo: Atlas, 2002. ISBN 9788522431694.
- 58 MARCONI, M. de A.; LAKATOS, E. M. *Fundamentos de metodologia científica*. 8. ed. São Paulo: Atlas, 2017. ISBN 9788597010763.
- 59 SILVA, M. L. M. d. *Notebooks Jupyter da Pesquisa: Pipeline de Coleta, Engenharia de Dados e Modelagem*. 2025. Arquivos da execução prática da dissertação, incluindo: teste_urls_ativas_mj.ipynb, scan_zap.ipynb, gerar_dataset.ipynb, dataset_avaliacao_ppsi.ipynb, dataset_enriquecido_epss.ipynb, dataset_vuln_epss_ppsi_normalizar.ipynb e projeto_cienciadedados.ipynb.
- 60 OWASP ZAP. *Relatório de Varredura de Vulnerabilidades (JSON)*. 2025. Arquivo 2025-06-29-ZAP-Report-.json gerado para esta pesquisa, documentando CWEs em ativos reais.
- 61 OpenVAS. *Relatório de Varredura de Vulnerabilidades (CSV)*. 2025. Arquivo report1_openvas.csv gerado para esta pesquisa, documentando vulnerabilidades em ativos reais.

- 62 SILVA, M. L. M. d. *Dataset Final da Pesquisa: Vulnerabilidades, EPSS e NCS*. 2025. Dataset final utilizado na modelagem, `dataset_artigo.csv`, e seus arquivos de origem (ex: `dataset_vulnerabilidades_balanceado_final.csv`, `relatorio_final_criticidade_ponderado.csv`, `dataset_completo_ml_balanceado.csv`, etc.).
- 63 FIRST.org, Inc. *EPSS Data Feeds*. 2025. <https://www.first.org/epss/data_stats>. Acessado em: 10/09/2025.
- 64 PEDREGOSA, F. e. a. Scikit-learn: Machine learning in python. *Journal of Machine Learning Research*, v. 12, p. 2825–2830, 2011.
- 65 TUKEY, J. W. *Exploratory Data Analysis*. Reading, MA: Addison-Wesley, 1977. ISBN 9780201076165.
- 66 ALTHAR, R. R.; SAMANTA, D.; KAUR, M.; SINGH, D.; LEE, H. N. Automated risk management based software security vulnerabilities management. *IEEE Access*, v. 10, p. 90597–90608, 2022.
- 67 MAOUAKI, W. E.; INNAN, N.; MARCHISIO, A.; SAID, T.; BENNAI, M.; SHAFIQUE, M. Quantum clustering for cybersecurity. In: *2024 IEEE International Conference on Quantum Computing and Engineering (QCE)*. [S.l.: s.n.], 2024.