

## **Uma análise comparativa de algoritmos de clusterização para a definição de perfis de risco de vulnerabilidades**

### **A comparative analysis of clustering algorithms for defining vulnerability risk profiles**

### **Un análisis comparativo de algoritmos de agrupamiento para definir perfiles de riesgo de vulnerabilidad**

**Madson Luiz Magno da Silva**

Mestrando em Engenharia Elétrica

Instituição: Universidade de Brasília (UNB)

Endereço: Brasília - Distrito Federal, Brasil

E-mail: madson.silva@aluno.unb.br

**João Souza Neto**

Pós-Doutor em Segurança Cibernética

Instituição: Universidade de Brasília (UNB)

Endereço: Brasília - Distrito Federal, Brasil

E-mail: neto.joao@unb.br

#### **RESUMO**

A seleção de algoritmos de aprendizado de máquina adequados é determinante para a eficácia da gestão de riscos cibernéticos baseada em dados. Este trabalho apresenta uma análise comparativa entre algoritmos de clusterização aplicados ao problema de priorização de vulnerabilidades, utilizando um *dataset* enriquecido com probabilidade de exploração (EPSS) e impacto contextual. A metodologia avaliou o desempenho dos algoritmos K-Means, Clusterização Hierárquica Aglomerativa e DBSCAN, validando a consistência dos agrupamentos através de métricas rigorosas como o Método do Cotovelo, Coeficiente de Silhueta e o *Adjusted Rand Index* (ARI). Os resultados demonstram que o algoritmo K-Means ( $k=5$ ) oferece o melhor equilíbrio entre coesão dos grupos e interpretabilidade para a segmentação de risco em larga escala, resultando em uma redução de 69% no esforço de remediação comparado a métodos tradicionais. Em contraste, o DBSCAN mostrou-se menos eficaz para a categorização geral, mas crucial para a detecção de anomalias (*outliers*) de alto risco que escapam aos perfis padrão.

**Palavras-chave:** priorização de vulnerabilidades, machine learning, clusterização, K-Means, DBSCAN, gestão de risco.

#### **ABSTRACT**

The selection of appropriate machine learning algorithms is crucial for the effectiveness of data-driven cybersecurity risk management. This work presents a comparative analysis of clustering algorithms applied to the vulnerability prioritization problem, using a dataset enriched with probability of exploitation (EPSS) and contextual impact. The methodology evaluated the performance of the K-Means, Agglomerative Hierarchical Clustering, and DBSCAN algorithms, validating the consistency of the clusters through rigorous metrics such as the Elbow Method,

Silhouette Coefficient, and the Adjusted Rand Index (ARI). The results demonstrate that the K-Means algorithm ( $k=5$ ) offers the best balance between group cohesion and interpretability for large-scale risk segmentation, resulting in a 69% reduction in remediation effort compared to traditional methods. In contrast, DBSCAN proved less effective for general categorization but crucial for detecting high-risk outliers that escape standard profiles.

**Keywords:** vulnerability prioritization, machine learning, clustering, K-Means, DBSCAN, risk management.

## RESUMEN

La selección de algoritmos de aprendizaje automático apropiados es crucial para la eficacia de la gestión de riesgos de ciberseguridad basada en datos. Este trabajo presenta un análisis comparativo de algoritmos de agrupamiento aplicados al problema de priorización de vulnerabilidades, utilizando un conjunto de datos enriquecido con probabilidad de explotación (EPSS) e impacto contextual. La metodología evaluó el rendimiento de los algoritmos K-Means, Agglomerative Hierarchical Clustering y DBSCAN, validando la consistencia de los clústeres mediante métricas rigurosas como el método Elbow, el coeficiente de silueta y el índice Rand ajustado (ARI). Los resultados demuestran que el algoritmo K-Means ( $k=5$ ) ofrece el mejor equilibrio entre cohesión de grupo e interpretabilidad para la segmentación de riesgos a gran escala, lo que resulta en una reducción del 69% en el esfuerzo de remediación en comparación con los métodos tradicionales. Por el contrario, DBSCAN resultó menos eficaz para la categorización general, pero crucial para detectar valores atípicos de alto riesgo que escapan a los perfiles estándar.

**Palabras clave:** priorización de vulnerabilidades, aprendizaje automático, agrupamiento, K-Means, DBSCAN, gestión de riesgos.

## 1 INTRODUÇÃO

A gestão de vulnerabilidades transformou-se em um requisito fundamental para lidar com o volume massivo de falhas reportadas anualmente. Com mais de 29.000 novas CVEs (*Common Vulnerabilities and Exposures*) registradas apenas em 2023 (NVD, 2024), métodos tradicionais de priorização baseados em regras estáticas tornaram-se insuficientes. O padrão da indústria, o *Common Vulnerability Scoring System* (CVSS), tem sido alvo de críticas consistentes na literatura acadêmica. Howland (2023) demonstra que o CVSS carece de justificção estatística robusta para suas métricas, falhando em correlacionar severidade técnica com a probabilidade real de ataque. Similarmente, Allodi e Massacci (2014) evidenciaram que a grande maioria das vulnerabilidades classificadas como "Críticas" pelo CVSS nunca são exploradas em ambientes reais, gerando desperdício de recursos e fadiga de alertas.

Para mitigar essas limitações, a adoção de métricas dinâmicas como o *Exploit Prediction Scoring System* (EPSS) introduziu a dimensão de probabilidade ao cálculo de risco (Jacobs et al., 2021). No entanto, a disponibilidade de dados multidimensionais impõe um novo desafio: qual técnica algorítmica é mais adequada para processar esses dados e gerar decisões de segurança confiáveis?

Embora existam estudos de *Machine Learning* (ML) em segurança, como a predição supervisionada de vulnerabilidades (Lele, 2022) e a clusterização de dados de honeynets (Kashtalian; Sochor, 2021), há uma carência de análises comparativas diretas entre algoritmos de partição e densidade focadas especificamente na priorização estratégica baseada em EPSS. A literatura tende a focar na detecção de intrusão ou na classificação binária de exploits, deixando uma lacuna sobre qual estrutura matemática melhor se adapta à distribuição de risco em ambientes corporativos.

Inspirado em estudos comparativos de outras áreas (Bansal, 2023), este artigo aborda essa lacuna conduzindo uma análise comparativa entre três abordagens de aprendizado não supervisionado: particionamento (K-Means), hierárquica e baseada em densidade (DBSCAN). O objetivo principal é avaliar a eficácia desses métodos na segmentação de vulnerabilidades em perfis de risco acionáveis, considerando o equilíbrio entre a robustez estatística dos agrupamentos e a interpretabilidade necessária para a tomada de decisão estratégica em cibersegurança.

## 1.1 OBJETIVOS

O objetivo geral deste trabalho é realizar uma análise comparativa entre três abordagens de aprendizado não supervisionado, particionamento (K-Means), hierárquica e baseada em densidade (DBSCAN), para a segmentação de vulnerabilidades em perfis de risco acionáveis. Para atingir o objetivo principal, foram definidos os seguintes objetivos específicos:

- a) enriquecer um *dataset* de vulnerabilidades corporativas com as métricas EPSS (probabilidade) e NCS (impacto contextual);
- b) aplicar os algoritmos selecionados para identificar o número ideal de agrupamentos (*clusters*) que representem a estrutura de risco da organização;
- c) avaliar o desempenho técnico dos algoritmos utilizando as métricas de Silhueta, Davies-Bouldin e o *Adjusted Rand Index* (ARI);

- d) mensurar a eficiência operacional da abordagem proposta em comparação ao método tradicional (CVSS), visando a redução do volume de alertas prioritários.

## **2 REFERENCIAL TEÓRICO E TRABALHOS RELACIONADOS**

Para fundamentar a análise comparativa proposta, esta seção apresenta as métricas de risco utilizadas e os algoritmos de clusterização avaliados, seguida pela discussão dos trabalhos correlatos que motivaram este estudo.

### **2.1 MÉTRICAS DE PRIORIZAÇÃO DE VULNERABILIDADES**

A avaliação de risco baseia-se historicamente na combinação de impacto e probabilidade, transcendendo a visão unidimensional de severidade técnica. Cumpre ressaltar que existem diversas métricas e frameworks de priorização de vulnerabilidades disponíveis na literatura e no mercado. Contudo, para os fins deste estudo, optou-se pela utilização do CVSS e do EPSS, visto que figuram entre as métricas mais consolidadas, abertas e utilizadas globalmente pela comunidade de segurança, permitindo a reprodutibilidade dos métodos aplicados.

#### **2.1.1 CVSS (Common Vulnerability Scoring System)**

O CVSS é o padrão da indústria para avaliar a severidade de vulnerabilidades. Em sua versão v3.1 (FIRST, 2019), o CVSS foca nas características intrínsecas da falha (Vetor Base), resultando em uma pontuação de 0 a 10. Contudo, como apontam Allodi e Massacci (2014), o CVSS mede a severidade técnica, não o risco de exploração. A dependência exclusiva dessa métrica gera um alto índice de falsos positivos na priorização, pois ignora o contexto de ameaça ativa.

#### **2.1.2 EPSS (Exploit Prediction Scoring System)**

Proposto por Jacobs et al. (2021), o EPSS utiliza aprendizado de máquina treinado em dados de threat intelligence para estimar a probabilidade ( $P \in [0,1]$ ) de uma CVE ser explorada

nos próximos 30 dias. Esta métrica introduz a dimensão dinâmica de ameaça que falta ao CVSS, permitindo que as organizações foquem na probabilidade real de um ataque ocorrer, em vez de apenas na teórica severidade da falha.

## 2.2 ALGORITMOS DE CLUSTERIZAÇÃO

A clusterização visa agrupar objetos de dados de tal forma que objetos no mesmo grupo (cluster) sejam mais similares entre si do que com aqueles em outros grupos. Este trabalho compara três abordagens clássicas:

### 2.2.1 K-Means (Particionamento)

É um algoritmo de particionamento que agrupa os dados em um número pré-definido ( $k$ ) de clusters. Seu objetivo é minimizar a inércia, ou a soma das distâncias quadráticas entre as amostras e o centroide de seu cluster (MacQueen, 1967). Devido à sua simplicidade e escalabilidade, é amplamente utilizado em diversas áreas, incluindo cibersegurança, para segmentar comportamentos e identificar grupos de atividades maliciosas.

### 2.2.2 Clusterização hierárquica aglomerativa

Diferente do particionamento direto, este método constrói uma hierarquia de clusters de baixo para cima. Inicialmente, cada ponto de dado é seu próprio cluster, e o algoritmo funde iterativamente os pares mais próximos até que apenas um cluster (ou o número desejado) permaneça. O resultado, visualizado como um dendrograma, é fundamental para a análise exploratória da estrutura subjacente dos dados, permitindo validar a existência natural de agrupamentos sem a necessidade de pré-definir  $k$  (Bansal, 2023).

### 2.2.3 DBSCAN (Baseado em Densidade)

O Density-Based Spatial Clustering of Applications with Noise agrupa pontos que estão densamente compactados, marcando como outliers (ruído) os pontos que residem sozinhos em

regiões de baixa densidade (Ester et al., 1996). Sua principal vantagem sobre os métodos anteriores é a capacidade de encontrar clusters com formas arbitrárias e sua robustez a ruído, tornando-o ideal para a detecção de anomalias em segurança, onde ataques podem se manifestar como comportamentos atípicos e isolados.

## 2.3 TRABALHOS RELACIONADOS E LACUNA DE PESQUISA

A aplicação de Machine Learning em cibersegurança tem sido amplamente explorada, embora frequentemente focada em detecção de intrusão ou classificação supervisionada. Kashtalian e Sochor (2021), por exemplo, utilizaram K-Means para agrupar dados de honeynets, demonstrando a eficácia de métodos não supervisionados na identificação de padrões de ataque e segmentação de comportamento malicioso.

No campo da análise comparativa de algoritmos, Bansal (2023) conduziu um estudo comparando K-Means, Clusterização Hierárquica e DBSCAN para a segmentação de clientes. O autor concluiu que, enquanto o K-Means oferece eficiência para grandes bases, a abordagem hierárquica fornece melhor entendimento estrutural. Similarmente, Ester et al. (1996), ao proporem o DBSCAN, evidenciaram sua superioridade na detecção de anomalias espaciais, uma característica desejável para identificar riscos atípicos.

Entretanto, observa-se uma lacuna na literatura quanto à aplicação comparativa direta desses métodos para a priorização estratégica de vulnerabilidades. A maioria dos estudos atuais foca na predição de exploits (supervisionado) ou na análise técnica de redes. Diferente dos trabalhos citados, esta pesquisa não busca apenas agrupar ameaças técnicas ou clientes, mas sim definir perfis de risco de negócio, combinando a probabilidade de exploração (EPSS) com o impacto contextual (PPSI). A contribuição deste trabalho reside na avaliação de qual filosofia algorítmica (partição, hierarquia ou densidade) melhor traduz essa matriz de risco multidimensional em decisões operacionais acionáveis.

## 3 METODOLOGIA

Esta pesquisa classifica-se como de cunho exploratório e descritivo, adotando uma abordagem quantitativa para a análise dos dados. O experimento foi desenhado para realizar uma

análise comparativa das abordagens de clusterização sobre um mesmo *dataset*, permitindo a caracterização dos perfis de risco resultantes. O fluxo metodológico compreende a construção e o enriquecimento dos dados, o pré-processamento e a definição das métricas de avaliação.

### 3.1 CONSTRUÇÃO E ENRIQUECIMENTO DO DATASET

A base de dados foi construída a partir de um relatório de vulnerabilidades simulado, representativo de um ambiente corporativo real. Para viabilizar a análise de risco multidimensional, os dados brutos foram enriquecidos com duas métricas fundamentais:

- a) impacto contextual (NCS): Para superar a generalidade do CVSS, calculou-se a Nota de Criticidade de Sistema (NCS) baseada no Framework PPSI (MGI, 2024). A métrica pondera o impacto no negócio (I) e a vulnerabilidade inerente (V) do ativo, conforme a Equação 1:  $NCS = (2 \times I) + V$ . Essa métrica segue o modelo de avaliação proposto pelo Tribunal de Contas da União (TCU, 2020);
- b) probabilidade de exploração (EPSS): O dataset foi enriquecido com os scores EPSS (pe), obtidos da base da FIRST.org (FIRST, 2024). O EPSS representa a probabilidade estimada de uma vulnerabilidade ser ativamente explorada nos próximos 30 dias ( $0 \leq pe \leq 1$ ). Vulnerabilidades sem registro correspondente receberam valor 0.

#### 3.1.1 Balanceamento e pré-processamento

Para garantir a relevância estatística e evitar que ativos com um número excessivo de vulnerabilidades dominassem a formação dos *clusters*, foi aplicada uma técnica de balanceamento por amostragem aleatória. Para cada ativo contendo mais de 50 vulnerabilidades, definiu-se um limite superior variável (sorteado entre 50 e 250), selecionando-se uma amostra aleatória (*random\_state* = 42) sem ordenação prévia por severidade. Ativos com 50 ou menos vulnerabilidades foram mantidos integralmente.

Antes da aplicação dos algoritmos, os vetores de características  $x = [EPSS, NCS]$  foram submetidos à padronização (*StandardScaler*), transformando os dados para média zero e desvio padrão unitário. Esta etapa é crucial para algoritmos baseados em distância, como o K-Means, evitando que a escala das variáveis distorça o cálculo da inércia (James *et al.*, 2021).

### 3.2 MÉTRICAS DE AVALIAÇÃO

Para garantir uma análise comparativa abrangente, foram utilizadas métricas de avaliação para determinar o número ideal de *clusters*, medir a qualidade intrínseca dos agrupamentos e quantificar a similaridade entre as partições geradas. Para evitar a segmentação excessiva do texto, as métricas aplicadas estão descritas a seguir:

- a) determinação de k (Método do Cotovelo): O número ideal de *clusters* foi investigado utilizando o Método do Cotovelo (*Elbow Method*), que analisa a variação da inércia (WCSS - *Within-Cluster Sum of Squares*). O método busca identificar o ponto de inflexão onde a adição de novos *clusters* deixa de resultar em uma redução significativa da variância intra-cluster (MacQueen, 1967);
- b) Coeficiente de Silhueta (*Silhouette Score*): Mede quão similar um objeto é ao seu próprio *cluster* em comparação com os outros, avaliando a coesão e separação sem a necessidade de um gabarito (*ground truth*). Os valores variam de -1 a 1, onde valores mais altos indicam *clusters* densos e bem definidos (Arbelaitz *et al.*, 2013);
- c) Índice Davies-Bouldin: Avalia a qualidade da clusterização baseando-se na razão entre a dispersão intra-cluster e a separação inter-cluster. Valores mais baixos indicam uma melhor partição, com *clusters* mais compactos e mais distantes entre si (Komaragiri, 2022);
- d) métrica de similaridade (*Adjusted Rand Index* - ARI): Utilizada para quantificar o grau de concordância entre os agrupamentos gerados pelos diferentes algoritmos (ex: K-Means vs. Hierárquico). O ARI mede a similaridade entre duas partições de dados, corrigindo a probabilidade de agrupamento ao acaso. Seus valores variam de -1 a 1, onde 1 indica concordância perfeita e 0 representa uma similaridade aleatória (Yoon, 2023).

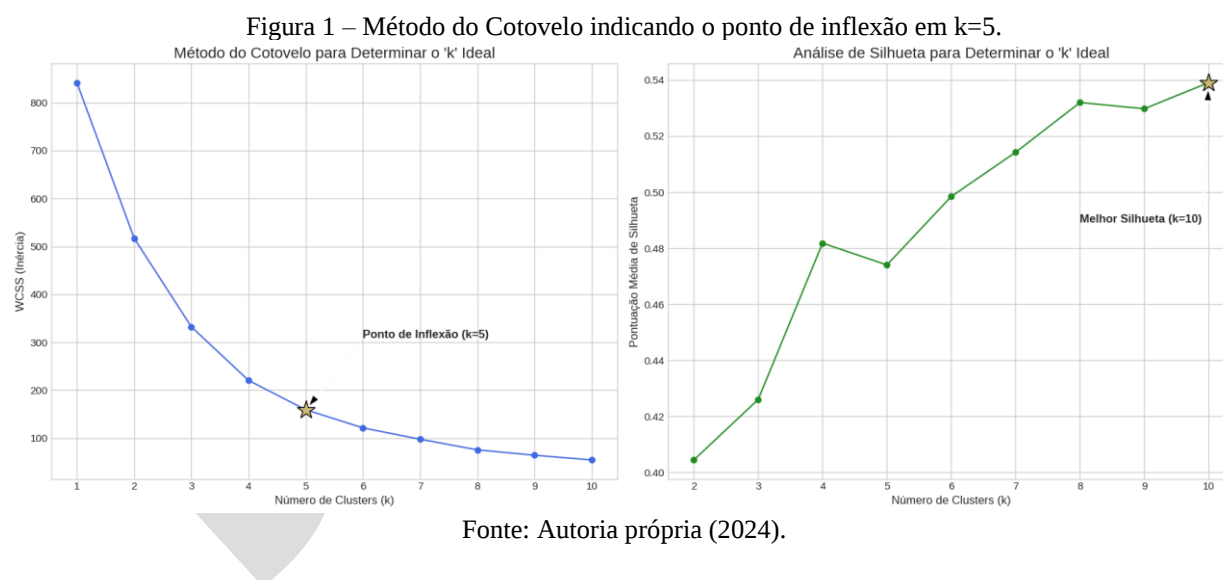
## 4 RESULTADOS E DISCUSSÃO

Nesta seção, são apresentados os resultados da aplicação das metodologias de clusterização. A análise segue uma lógica sequencial: inicialmente, determina-se o número ideal de agrupamentos (*k*), seguido pela comparação técnica entre os algoritmos e, por fim, a tradução desses agrupamentos em perfis de risco acionáveis e a análise de eficiência operacional.

#### 4.1 DETERMINAÇÃO DO NÚMERO DE CLUSTERS (K)

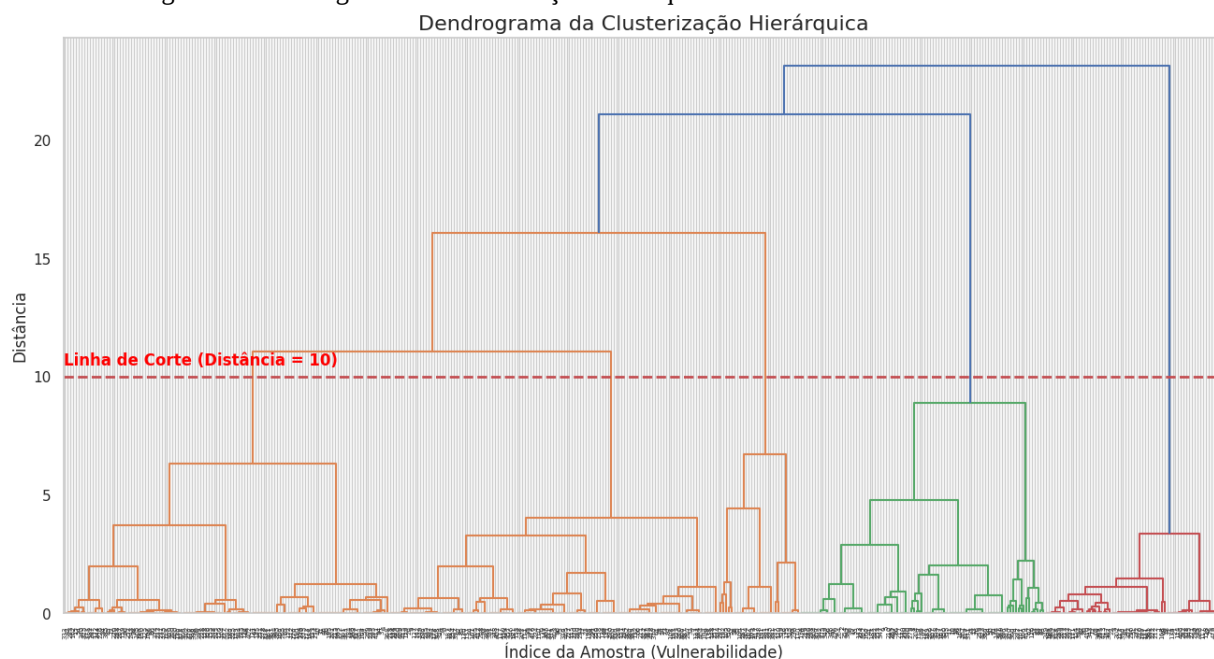
Para determinar o número ideal de perfis de risco, utilizou-se uma abordagem combinada entre o Método do Cotovelo (*Elbow Method*) e a análise estrutural via dendrograma.

A Figura 1 apresenta a curva da inércia (*Within-Cluster Sum of Squares - WCSS*) para  $k$  variando de 2 a 10. Observa-se um ponto de inflexão claro em  $k = 5$ . Conforme detalhado na análise experimental, a redução da inércia ao passar de  $k = 4$  para  $k = 5$  foi significativa (queda de 225,35 para 159,09), enquanto a adição de *clusters* subsequentes resultou em ganhos marginais decrescentes (James *et al.*, 2021).



Embora a Análise de Silhueta tenha sugerido matematicamente um  $k$  maior (pico em  $k = 10$ ), optou-se por  $k = 5$  para manter a interpretabilidade do modelo de negócio. Esta decisão foi corroborada pela inspeção visual do dendrograma da Clusterização Hierárquica (Figura 2), onde um corte na distância euclidiana de 10 particiona os dados naturalmente em cinco grandes grupos estruturais.

Figura 2 – Dendrograma da Clusterização Hierárquica validando a estrutura de 5 clusters.



Fonte: Autoria própria (2024).

## 4.2 COMPARAÇÃO DE ALGORITMOS

Com  $k = 5$  definido, comparou-se o desempenho dos algoritmos K-Means, Clusterização Hierárquica e DBSCAN. A Tabela 1 resume as métricas de qualidade intrínseca obtidas.

Tabela 1. Métricas de Qualidade dos Agrupamentos

| Algoritmo             | Pontuação de Silhueta ( $\uparrow$ ) | Índice Davies-Bouldin ( $\downarrow$ ) |
|-----------------------|--------------------------------------|----------------------------------------|
| K-Means ( $k=5$ )     | <b>0.4740</b>                        | 0.7175                                 |
| Hierárquico (Ward)    | 0.4662                               | 0.7222                                 |
| DBSCAN (sem outliers) | 0.4145                               | <b>0.4954</b>                          |

Fonte: Autoria própria (2024).

O K-Means apresentou o maior Coeficiente de Silhueta (0.4740), indicando *clusters* ligeiramente mais coesos e bem definidos em comparação à abordagem hierárquica. O DBSCAN obteve o melhor Índice Davies-Bouldin (0.4954), o que reflete a alta compactação dos núcleos densos que ele identificou; no entanto, sua incapacidade de classificar a totalidade dos dados (gerando muitos ruídos ou agrupando dados dispersos em um único *cluster* grande) limitou sua utilidade como método primário de segmentação.

Para validar a robustez da partição gerada pelo K-Means, calculou-se o *Adjusted Rand*

*Index* (ARI) em relação à Clusterização Hierárquica. O resultado de 0.8997 indica uma concordância quase perfeita entre os dois métodos, sugerindo que a estrutura de risco de 5 perfis é intrínseca aos dados e não um artefato do algoritmo escolhido.

#### 4.3 INTERPRETAÇÃO DOS PERFIS DE RISCO (K-MEANS)

Dada a sua superioridade no equilíbrio entre desempenho e interpretabilidade, o K-Means foi selecionado para a definição final dos perfis. A análise dos centroides (Tabela 2) revela cinco categorias de risco distintas baseadas na Probabilidade (EPSS) e Impacto (NCS).

| Tabela 2. Caracterização dos Perfis de Risco (Centroides K-Means) |             |             |                                                |
|-------------------------------------------------------------------|-------------|-------------|------------------------------------------------|
| Perfil (Cluster)                                                  | EPSS Médio  | NCS Médio   | Interpretação Estratégica                      |
| 0 - Crítico                                                       | 0.33        | <b>0.86</b> | Alto impacto de negócio; prioridade máxima.    |
| 1 - Alto                                                          | <b>0.67</b> | 0.35        | Alta probabilidade de ataque; ameaça iminente. |
| 2 - Médio                                                         | 0.34        | 0.41        | Risco padrão; gestão contínua.                 |
| 3 - Baixo                                                         | 0.07        | 0.44        | Baixa probabilidade; monitoramento.            |
| 4 - Mínimo                                                        | 0.02        | 0.07        | Risco residual; aceitável.                     |

Fonte: Autoria própria (2024).

Destaca-se a distinção entre o *cluster* "Crítico" (focado no impacto ao negócio) e o *cluster* "Alto" (focado na probabilidade de exploração). Esta separação permite estratégias de defesa diferenciadas: proteção de ativos vitais *versus* mitigação de ameaças ativas.

#### 4.4 EFICIÊNCIA OPERACIONAL E REDUÇÃO DE FADIGA

A aplicação prática destes perfis demonstrou um ganho significativo em eficiência operacional. Comparando-se a priorização tradicional (CVSS níveis *High* e *Critical*) com a abordagem proposta (*Clusters* Alto e Crítico), observou-se uma redução drástica no volume de alertas prioritários:

- a) abordagem tradicional (CVSS): 335 vulnerabilidades exigindo atenção imediata;
- b) abordagem proposta (Clusterização): 104 vulnerabilidades prioritárias.

Isso representa uma redução de esforço de 69%. A Matriz de Realocação de Risco demonstra que muitas vulnerabilidades tecnicamente severas (CVSS alto) foram reclassificadas para perfis "Médio" ou "Baixo" devido à sua baixa probabilidade real de exploração (EPSS) ou baixo impacto contextual, permitindo que as equipes foquem nos riscos reais.

#### 4.5 ANÁLISE DE ANOMALIAS COM DBSCAN

Embora o DBSCAN tenha se mostrado menos eficaz para a segmentação global, ele desempenhou um papel crucial na identificação de *outliers*. O algoritmo isolou 14 vulnerabilidades que não se enquadravam nos perfis padrão.

A análise individualizada destes pontos revelou casos de risco extremo e atípico, como:

- a) CVE-2003-0033 (EPSS 0.52, NCS 0.99): Um ativo de impacto quase máximo com probabilidade moderada, representando uma ameaça latente catastrófica que poderia ser subpriorizada em médias simples;
- b) CVE-2024-0195 (EPSS 0.92, NCS 0.32): Uma vulnerabilidade de altíssima probabilidade em um ativo de menor valor, servindo como um provável ponto de entrada para movimentação lateral.

Portanto, o DBSCAN complementa o K-Means agindo como uma rede de segurança para detectar exceções críticas que fogem à regra estatística dos centroides.

### 5 CONCLUSÃO

Este trabalho conduziu uma análise comparativa entre algoritmos de *clusterização* para a descoberta de perfis de risco de vulnerabilidades, transcendendo as limitações de métricas estáticas. A investigação confirmou que a escolha do algoritmo impacta diretamente a capacidade de traduzir dados técnicos em decisões estratégicas.

A avaliação baseada em métricas de qualidade intrínseca e estabilidade demonstrou que o K-Means ( $k = 5$ ) oferece o melhor equilíbrio para a segmentação global. Validado por um *Adjusted Rand Index* (ARI) de aproximadamente 0,90 em relação à abordagem hierárquica e por pontuações de Silhueta consistentes, o método provou ser robusto na identificação de padrões de risco bem definidos e interpretáveis.

Do ponto de vista de negócio, a aplicação dessa segmentação resultou em um ganho de eficiência tangível. Ao diferenciar vulnerabilidades de alto impacto ("Crítico") daquelas de alta probabilidade ("Alto"), a metodologia permitiu uma redução de 69% no esforço de remediação imediata, eliminando a fadiga de alertas causada por falsos positivos técnicos que saturam as equipes de segurança.

No entanto, a análise também revelou que métodos de particionamento podem obscurecer ameaças singulares. Por isso, conclui-se que a arquitetura ideal para a gestão de vulnerabilidades é uma abordagem híbrida: utilizar a eficiência do K-Means para a estruturação massiva de políticas de segurança e a capacidade de detecção de densidade do DBSCAN para identificar anomalias (*outliers*) de risco extremo que fogem à regra estatística.

Como trabalhos futuros, sugere-se a integração de técnicas de Inteligência Artificial Explicável (XAI) para fornecer justificativas semânticas automáticas para cada atribuição de *cluster*, aumentando a confiança dos gestores na decisão algorítmica (Islam, 2025). Adicionalmente, a incorporação de fluxos de *Threat Intelligence* em tempo real poderá transformar este modelo de avaliação pontual em um sistema de resposta a incidentes dinâmico e adaptativo (Rodriguez, 2020).

## REFERÊNCIAS

- ALLODI, L.; MASSACCI, F. **Comparing Vulnerability Severity and Exploits Using Case-Control Studies**. *ACM Transactions on Information and System Security (TISSEC)*, v. 17, n. 1, p. 1-20, 2014.
- ARBELAITZ, O. *et al.* **An extensive comparative study of cluster validity indices**. *Pattern Recognition*, v. 46, n. 1, p. 243-256, 2013.
- 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, 2023.
- ESTER, M. *et al.* **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)*, 1996. AAAI Press, 1996. p. 226-231.
- FIRST.ORG. **Common Vulnerability Scoring System v3.1: Specification Document**. Forum of Incident Response and Security Teams, 2019. Disponível em: <https://www.first.org/cvss/v3-1/specification-document>. Acesso em: 14 out. 2025.
- FIRST.ORG. **Exploit Prediction Scoring System (EPSS) User Guide**. Forum of Incident Response and Security Teams, 2024. Disponível em: <https://www.first.org/epss/user-guide>. Acesso em: 14 out. 2025.
- HOWLAND, H. **CVSS: Ubiquitous and Broken**. *Digital Threats: Research and Practice*, v. 4, n. 1, p. 1-12, 2023.
- ISLAM, S. *et al.* **Intelligent dynamic cybersecurity risk management framework with explainability and interpretability of AI models**. *Journal of Reliable Intelligent Environments*, v. 11, p. 1-22, 2025.
- JACOBS, J. *et al.* **Exploit Prediction Scoring System (EPSS)**. *Digital Threats: Research and Practice*, v. 2, n. 3, p. 1-17, 2021.
- JAMES, G. *et al.* **An Introduction to Statistical Learning: Applications in R**. 2. ed. Springer, 2021.
- 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)*, 2021. v. 2853, p. 48-58, 2021.
- 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. 980-998, 2022.

LELE, A.; KULKARNI, V. **Machine learning based vulnerability analysis and prioritization.** *International Journal of Computer Applications*, v. 184, n. 11, p. 22-29, 2022.

MACQUEEN, J. B. **Some methods for classification and analysis of multivariate observations.** In: *PROCEEDINGS OF THE FIFTH BERKELEY SYMPOSIUM ON MATHEMATICAL STATISTICS AND PROBABILITY*, 1967. University of California Press, 1967. v. 1, p. 281-297.

MINISTÉRIO DA GESTÃO E DA INOVAÇÃO EM SERVIÇOS PÚBLICOS (MGI). **Framework do PPSI - Guia.** 2024. Disponível em: [https://www.gov.br/governodigital/pt-br/seguranca-e-protecao-de-dados/ppsi/guia\\_framework\\_psi.pdf](https://www.gov.br/governodigital/pt-br/seguranca-e-protecao-de-dados/ppsi/guia_framework_psi.pdf).

NATIONAL INSTITUTE OF STANDARDS AND TECHNOLOGY. **NVD - CVE Metrics.** 2024. Disponível em: <https://nvd.nist.gov/vuln/reporting/nvd-dashboard>. Acesso em: 14 out. 2025.

RODRIGUEZ, A.; OKAMURA, K. **Enhancing data quality in real-time threat intelligence systems using machine learning.** *Social Network Analysis and Mining*, v. 10, n. 1, p. 91, 2020.

TRIBUNAL DE CONTAS DA UNIÃO (TCU). **Acórdão 1.889/2020-TCU-Plenário.** 2020.

YOON, S.-S. *et al.* **Vulnerability assessment based on real world exploitability for prioritizing patch applications.** In: *2023 7TH CYBER SECURITY IN NETWORKING CONFERENCE (CSNet)*, 2023. p. 62-66.