



INSTITUTO FEDERAL
São Paulo
Campus São Paulo

REGRASP

Detecção de Irregularidades em Contratações Públicas: Uma Abordagem com *Machine Learning* Baseada em Consenso entre Algoritmos

Detection of Irregularities in Public Procurement: A Machine Learning-Based Approach Using Algorithm Consensus

Raphael Hendrigo de Souza Gonçalves

Especialista em Gestão Pública Municipal
Universidade Federal de São João Del-Rei
<https://orcid.org/0009-0009-3635-6395>
raphael.goncalves@tcmsp.tc.br

Wendel Marcos dos Santos

Doutor em Administração
Instituto Federal de Educação, Ciência e Tecnologia de São Paulo
<https://orcid.org/0000-0002-1336-3623>
wendel.santos@ifsp.edu.br

Luiz Fernando Postingel Quirino

Especialista em Gestão de Riscos e Cibersegurança
Instituto Federal de Educação, Ciência e Tecnologia de São Paulo
<https://orcid.org/0000-0001-8481-4033>
luiz.quirino@ifsp.edu.br

Histórico do artigo

Recebido em: 02.04.2026
Revisões solicitadas: 23.04.2026
Versão reformulada recebida: 25.04.2026
Aprovado em: 08.06.2026
Publicado em: 08.06.2026

RESUMO

Este estudo insere-se no contexto da crescente demanda por transparência e eficiência na governança de contratações públicas. O objetivo foi identificar indícios de irregularidades por meio de técnicas de *Machine Learning* não supervisionado, com abordagem quantitativa baseada em mineração de dados do Portal Nacional de Contratações Públicas (PNCP), utilizando *K-Means*, *DBSCAN*, *Local Outlier Factor* e *Isolation Forest* combinados por análise de consenso. Os resultados evidenciaram diferentes níveis de sensibilidade entre os modelos: *K-Means* apresentou a maior taxa de anomalias, *Isolation Forest* e *LOF* comportamento mais conservador, e o *DBSCAN* superior qualidade de agrupamento. A análise de consenso identificou 76 processos sinalizados por ao menos três métodos simultaneamente, representando 3,47% do total analisado. Concluiu-se que a abordagem multialgorítmica contribui de forma objetiva e replicável para o fortalecimento dos mecanismos de auditoria pública, oferecendo ferramenta de triagem inteligente que apoia a priorização de casos por órgãos de controle e fiscalização.

Palavras-chave: Governança Pública; Contratações Públicas; *Machine Learning*; Detecção de Anomalias; Consenso entre Algoritmos.

ABSTRACT

This study addresses the growing demand for transparency and efficiency in public procurement governance. The objective was to identify irregularities through unsupervised Machine Learning techniques, using a quantitative data mining approach based on data from the Brazilian National Public Procurement Portal (PNCP),

applying K-Means, DBSCAN, Local Outlier Factor, and Isolation Forest combined through consensus analysis. The results revealed different sensitivity levels across models: K-Means showed the highest anomaly rate, Isolation Forest and LOF displayed more conservative behavior, and DBSCAN demonstrated superior clustering quality. The consensus analysis identified 76 procurement processes flagged by at least three methods simultaneously, representing 3.47% of the total analyzed. It was concluded that the proposed multi-algorithm approach contributes objectively and replicably to strengthening public audit mechanisms, offering an intelligent screening tool that supports case prioritization by oversight and regulatory bodies.

Keywords: Public Governance; Public Procurement; Machine Learning; Anomaly Detection; Algorithm Consensus.

1. Introdução

A crescente demanda por transparência, integridade e *accountability* na gestão pública tem impulsionado o desenvolvimento de mecanismos mais sofisticados de monitoramento e controle das contratações governamentais (Boege & Marques, 2024; Da Silva Pereira & Dorneles, 2025). Em um cenário caracterizado por elevado volume transacional, heterogeneidade de procedimentos e assimetria informacional entre agentes públicos e privados, a identificação de indícios de irregularidade ao longo de todo o ciclo licitatório torna-se um desafio crítico para a governança pública contemporânea (Da Silva & De Oliveira, 2024; Garbaccio & Ramos, 2025).

A licitação, enquanto processo administrativo destinado à seleção da proposta mais vantajosa ao interesse público, está fundamentada em princípios como legalidade, impessoalidade, moralidade, publicidade, isonomia e probidade administrativa (Filho, 2025; Martins et al., 2025). Ainda assim, a violação desses princípios pode ocorrer em diferentes modalidades de contratação, manifestando-se por meio de práticas como direcionamento de edital, sobrepreço, conluio entre fornecedores, fracionamento indevido de despesas e uso inadequado de hipóteses legais de exceção (Júnior et al., 2023; Da Silva & De Oliveira, 2024). Tais irregularidades não apenas geram perdas econômicas, mas também comprometem a legitimidade institucional e a eficiência alocativa do setor público (Nery et al., 2024; Reis et al., 2025).

A promulgação da Lei nº 14.133/2021 (Nova Lei de Licitações e Contratos Administrativos) representou um avanço na modernização do arcabouço normativo brasileiro, ao introduzir instrumentos de governança, planejamento e gestão de riscos nas contratações públicas (Filho, 2025; Souza Silva et al., 2025). No entanto, a complexidade operacional dos processos e a multiplicidade de atores envolvidos mantêm espaço para a ocorrência de desvios, inclusive em modalidades competitivas que, em tese, deveriam mitigá-los (Nery et al., 2024; Reis et al., 2025).

As consequências das irregularidades em contratações públicas são amplamente documentadas na literatura, indicando que uma parcela significativa dos recursos públicos pode ser comprometida por práticas ilícitas ou ineficiências (Bezerra et al., 2022; Dzeco & Simione, 2024; Boege & Marques, 2024; Nery et al., 2024). Além do impacto financeiro, tais práticas afetam a qualidade dos serviços públicos, distorcem a concorrência e reduzem a confiança nas instituições (Bezerra et al., 2022; Reis et al., 2025). No contexto brasileiro, órgãos de controle têm identificado padrões recorrentes de irregularidades, evidenciando limitações dos métodos tradicionais de auditoria diante do crescente volume e complexidade dos dados (Bezerra et al., 2022; Rodrigues et al., 2026).

Diante desse cenário, a aplicação de técnicas de Ciência de Dados emerge como uma alternativa promissora para o fortalecimento da fiscalização e da auditoria pública. Em particular, métodos de Aprendizado Não Supervisionado mostram-se adequados para a detecção de anomalias em bases de dados licitatórios, uma vez que não dependem de conjuntos previamente rotulados: uma limitação recorrente no domínio de fraudes públicas

(Abreu et al., 2024; Pavão et al., 2025). Algoritmos como o *Isolation Forest* (IF), o *DBSCAN* (*Density-Based Spatial Clustering of Applications with Noise*), o *Local Outlier Factor* (LOF) e o *K-Means* permitem identificar padrões atípicos sob diferentes perspectivas (tais como isolamento estatístico, densidade, vizinhança local e distância a centróides) possibilitando uma análise complementar e mais robusta dos dados (Dalmaijer et al., 2022; Khetarpaul et al., 2022).

Nesse contexto, este estudo parte da hipótese de que a combinação de algoritmos de aprendizado não supervisionado, por meio de uma abordagem baseada em consenso, aumenta a robustez e a confiabilidade na detecção de indícios de irregularidades em contratações públicas. Assim, este artigo tem como objetivo principal aplicar técnicas de *Machine Learning* para identificar indícios de irregularidades em contratações públicas. Como objetivo complementar, busca-se analisar o desempenho de diferentes algoritmos por meio de uma abordagem baseada em consenso, contribuindo para o fortalecimento do controle social e da transparência administrativa.

O trabalho está estruturado da seguinte forma: a Seção 2 apresenta o referencial teórico; a Seção 3 descreve os procedimentos metodológicos; a Seção 4 discute os resultados obtidos; e a Seção 5 traz as conclusões e recomendações.

2. Referencial Teórico

2.1. Licitações Públicas e Irregularidades Nas Contratações

As licitações públicas constituem o principal mecanismo por meio do qual a administração pública seleciona, de forma isonômica e transparente, fornecedores de bens, serviços e obras. Seu fundamento jurídico está previsto na Constituição Federal de 1988, que estabelece a obrigatoriedade do processo competitivo como regra geral das contratações governamentais (Filho, 2025; Martins et al., 2025). Esse processo é orientado por princípios como legalidade, impessoalidade, moralidade, publicidade, eficiência e probidade administrativa, os quais visam assegurar a prevalência do interesse público sobre interesses privados (Boege & Marques, 2024; Reis et al., 2025).

A Lei nº 14.133/2021 promoveu a modernização do regime jurídico das contratações públicas no Brasil, ao consolidar normas anteriormente dispersas e introduzir instrumentos voltados à governança, ao planejamento e à gestão de riscos (Boege & Marques, 2024; Souza Silva et al., 2025;). Entre suas inovações, destacam-se a ampliação das modalidades licitatórias e o fortalecimento de práticas de *compliance* e controle interno, refletindo uma abordagem mais estruturada e preventiva da gestão contratual (Da Silva & De Oliveira, 2024).

Apesar desses avanços normativos, a literatura evidencia que irregularidades persistem em diferentes etapas e modalidades de contratação pública, não se restringindo às hipóteses de contratação direta (Bezerra et al., 2022; Nery et al., 2024). Tais irregularidades incluem, entre outras, direcionamento de editais, sobrepreço, conluio entre fornecedores, fracionamento indevido de despesas e uso inadequado de exceções legais, frequentemente associadas a falhas de planejamento, fragilidade dos mecanismos de controle e limitações na capacidade técnica dos agentes envolvidos (Junior et al., 2023; Da Silva & De Oliveira, 2024).

No contexto brasileiro, evidências empíricas oriundas de auditorias conduzidas por órgãos como a Controladoria-Geral da União (CGU) e o Tribunal de Contas da União (TCU) apontam a recorrência desses padrões, inclusive em situações em que a irregularidade decorre de práticas estruturais (Bezerra et al., 2022; Souto et al., 2025). Esses achados reforçam a necessidade de abordagens mais sistemáticas e baseadas em dados para a identificação de indícios de irregularidade, especialmente diante do volume e da complexidade crescentes das informações relacionadas às contratações públicas.

2.2 Mineração de Dados e *Machine Learning*

A Mineração de Dados refere-se ao processo de extração de padrões, tendências e conhecimento relevante a partir de grandes volumes de dados, estruturados ou não, por meio da integração de técnicas estatísticas, *Machine Learning* e gestão de dados (Bazzaz Abkenar et al., 2021). No setor público, sua aplicação tem se expandido no apoio a atividades de auditoria e controle, permitindo a análise sistemática de registros administrativos em escala incompatível com abordagens tradicionais baseadas em inspeção manual (Abreu et al., 2024; Garbaccio & Ramos, 2025).

O *Machine Learning*, por sua vez, constitui um subcampo da Inteligência Artificial voltado ao desenvolvimento de modelos capazes de identificar padrões e realizar inferências a partir de dados (Arão, 2024). Esses modelos são usualmente classificados em três paradigmas: supervisionado, não supervisionado e por reforço (Teixeira et al., 2019; De Almeida & Chinelato, 2025). No contexto da auditoria pública, destaca-se o aprendizado não supervisionado, uma vez que a indisponibilidade de bases rotuladas (decorrente da subnotificação de irregularidades e da limitada consolidação de decisões definitivas) representa uma restrição prática recorrente (Garbaccio & Ramos, 2025; Souto et al., 2025).

Diante dessa limitação, métodos não supervisionados permitem identificar observações que se desviam dos padrões predominantes nos dados, sinalizando potenciais indícios de irregularidade para análise posterior (Junior et al., 2023; Pavão et al., 2025). Essa abordagem é particularmente relevante no contexto das contratações públicas, caracterizadas por elevado volume de dados, heterogeneidade de procedimentos e dinamicidade dos padrões de comportamento, o que dificulta a aplicação de regras fixas de detecção (Silva et al., 2023; Da Silva Pereira & Dorneles, 2025).

2.2.1 *K-Means*

O *K-Means* é um dos algoritmos de agrupamento (*clustering*) mais amplamente utilizados na literatura de *Machine Learning* (Sinaga & Yang, 2020; Khetarpaul et al., 2022). Seu funcionamento consiste em particionar um conjunto de observações em *k clusters*, de modo que cada observação seja atribuída ao *cluster* cujo centróide (ponto médio dos membros do grupo) apresenta a menor distância euclidiana (Dalmaijer et al., 2022; Khetarpaul et al., 2022). O processo é iterativo: os centróides são recalculados a cada ciclo até que a convergência seja atingida, minimizando a variância intragrupo (Umargono, Suseno & Gunawan, 2020; Sari, Al-Khowarizmi & Batubara, 2021).

No contexto da análise de processos licitatórios, o *K-Means* permite identificar grupos de contratos com características similares (como faixas de valor, perfis de fornecedores ou frequência de contratação) e, a partir daí, detectar registros que não se enquadram adequadamente em nenhum dos agrupamentos formados ou que se concentram em clusters estatisticamente atípicos (Sinaga & Yang, 2020). Embora o *K-Means* não seja um algoritmo de detecção de anomalias em sentido estrito, sua aplicação combinada com métricas de distância ao centróide permite sinalizar observações de *outliers* com boa eficiência computacional (Khetarpaul et al., 2022).

2.2.2 *Isolation Forest*

O *Isolation Forest* (IF) é um algoritmo não supervisionado desenvolvido especificamente para a detecção de anomalias em grandes volumes de dados (Liu et al., 2008). Seu princípio de funcionamento difere dos métodos baseados em densidade ou distância: em vez de modelar o comportamento normal dos dados para depois identificar desvios, o IF parte

da premissa de que anomalias são pontos raros e estruturalmente distintos, portanto mais fáceis de isolar (Bai et al., 2025).

O algoritmo constrói um conjunto de árvores de decisão aleatórias (denominadas *isolation trees*) realizando partições binárias sucessivas e aleatórias do espaço de atributos (Shao et al., 2025). Observações anômalas tendem a ser isoladas em poucos passos, resultando em caminhos mais curtos nas árvores, enquanto observações regulares requerem um número maior de partições para serem separadas (Liu et al., 2008). O escore de anomalia de cada observação é calculado com base no comprimento médio do caminho de isolamento ao longo de todas as árvores construídas (Bai et al., 2025). O IF apresenta desempenho computacional superior a métodos como Local Outlier Factor (LOF) e *DBSCAN* em conjuntos de dados de alta dimensionalidade, sendo amplamente adotado em aplicações de detecção de fraudes financeiras e auditoria de dados (Shao et al., 2025).

2.2.3 Local Outlier Factor (LOF)

O *Local Outlier Factor* (LOF) é um algoritmo baseado em densidade local, proposto por Breunig et al. (2000) que avalia o grau de anomalia de cada observação em relação à densidade de sua vizinhança (Basheer et al., 2026). Para cada ponto, o LOF calcula a razão entre a densidade local média dos seus k vizinhos mais próximos e sua própria densidade local: valores significativamente superiores a 1 indicam que o ponto está em uma região de baixa densidade comparada à sua vizinhança, caracterizando-o como *outlier* (Shahrestani & Sanislav, 2024).

A principal vantagem do LOF em relação a métodos globais, como o K-Means, é sua capacidade de detectar anomalias contextuais: observações que seriam consideradas normais em uma escala global, mas que se destacam como atípicas dentro de um contexto local específico (Basheer et al., 2026). No âmbito das licitações públicas, isso é especialmente relevante: um contrato de valor elevado pode ser perfeitamente regular em determinado segmento de mercado, mas altamente atípico quando comparado aos contratos de mesmo objeto celebrados por um município de pequeno porte (Da Silva Pereira & Dorneles, 2025; Souza Silva et al., 2025).

2.2.4 DBSCAN

O *DBSCAN* (*Density-Based Spatial Clustering of Applications with Noise*) é um algoritmo de agrupamento baseado em densidade, proposto por Ester et al. (1996), que identifica clusters como regiões de alta densidade separadas por regiões de baixa densidade (Singh et al., 2022). Diferentemente do *K-Means*, o *DBSCAN* não exige a definição prévia do número de clusters e é capaz de identificar agrupamentos de formas arbitrárias, sem a restrição de geometrias convexas (Singh et al., 2022).

O algoritmo classifica cada ponto do conjunto de dados em uma de três categorias: ponto central (*core point*), situado em uma região densa; ponto de borda (*border point*), localizado na periferia de um cluster; ou ruído (*noise point*), isolado em região de baixa densidade (Hanafi & Saadatfar, 2022). Os pontos classificados como ruído correspondem, na prática, às anomalias detectadas pelo algoritmo (Singh et al., 2022). No contexto deste trabalho, o *DBSCAN* é particularmente útil para identificar processos de contratação que não se agrupam com nenhum conjunto coerente de registros, sinalizando potenciais irregularidades estruturais que merecem investigação aprofundada (Júnior et al., 2023).

3. Procedimentos Metodológicos

Esta pesquisa caracteriza-se como quantitativa e exploratória, com abordagem baseada em mineração de dados aplicada ao domínio das contratações públicas. O estudo

utiliza dados secundários obtidos de fonte oficial, submetidos a etapas sequenciais de coleta, pré-processamento e modelagem por algoritmos não supervisionados de detecção de anomalias. Por se tratar de dados públicos, abertos e anonimizados provenientes do Portal Nacional de Contratações Públicas (PNCP), esta pesquisa não envolveu identificação de indivíduos, estando dispensada de apreciação por comitê de ética, conforme o Art. 1º, Parágrafo único, incisos II e III, da Resolução CNS nº 510/2016 (Brasil, 2016), em consonância com a Lei nº 12.527/2011 (Brasil, 2011).

Cabe destacar que, por se tratar de uma abordagem baseada em aprendizado não supervisionado, não há utilização de dados previamente rotulados ou definição de grupos de controle, limitação recorrente no domínio de fraudes públicas, decorrente da subnotificação de irregularidades e da limitada consolidação de decisões definitivas (Garbaccio & Ramos, 2025; Souto et al., 2025). Nesse contexto, a análise de consenso entre os quatro algoritmos atua como mecanismo de robustez na sinalização de anomalias, reduzindo a incidência de falsos positivos ao priorizar registros identificados por métodos com premissas distintas (Aggarwal, 2016).

3.1 Coleta de Dados

Os dados foram coletados por meio da API do Portal Nacional de Contratações Públicas (PNCP), repositório eletrônico oficial instituído pelo art. 174 da Lei nº 14.133/2021, responsável pela divulgação centralizada e obrigatória dos atos relativos às contratações públicas brasileiras (Da Silva Pereira et al., 2024). A coleta foi realizada por meio de requisições à API REST do PNCP, especificamente ao *endpoint* de consulta de processos de licitação. O retorno foi fornecido em formato *JSON* e convertido em estrutura tabular (*dataframe*) por meio da biblioteca Pandas (McKinney, 2011) resultando em um conjunto inicial de 2.967 registros e 52 variáveis. A base analisada corresponde à totalidade dos registros disponíveis no recorte temporal e temático considerado, caracterizando-se como um censo dos dados acessíveis. Dessa forma, não se aplicam procedimentos de amostragem probabilística, sendo a suficiência da base determinada pela abrangência e completude dos dados disponíveis (Marconi & Lakatos, 2017).

3.2 Ambiente de Implementação e Análise

Toda a implementação foi desenvolvida na linguagem *Python*, amplamente adotada na comunidade científica de Ciência de Dados por sua versatilidade e pelo rico ecossistema de bibliotecas especializadas (Sinaga & Yang, 2020). A coleta e manipulação dos dados foram realizadas com as bibliotecas *Requests* e *Pandas* (McKinney, 2011; Dalmaijer, Nord & Astle, 2022); o pré-processamento e a redução de dimensionalidade via PCA contaram com o suporte do *Scikit-learn*, que também forneceu as implementações dos algoritmos *K-Means*, *DBSCAN*, *LOF* e *Isolation Forest* (Umargono, Suseno & Gunawan, 2020); e a visualização dos resultados foi produzida com *Matplotlib*. O ambiente de desenvolvimento utilizado foi o *Jupyter Notebook*, que permite a execução interativa do código e a documentação integrada das etapas analíticas, favorecendo a transparência e a reprodutibilidade da pesquisa (Khetarpaul et al., 2022).

3.3 Pré-Processamento

O pré-processamento seguiu um protocolo sistemático de limpeza e redução dimensional, organizado em cinco etapas. Primeiramente, foram removidos registros duplicados (decorrentes de instabilidades na conexão com a API), reduzindo o conjunto para 2.188 observações únicas. Em seguida, foram excluídas as colunas com ausência total de valores: (orgaoSubRogado, justificativaPresencial, unidadeSubRogada, linkProcessoEletronico)

e aquelas com representatividade inferior a dois registros não nulos, por não contribuírem informativamente para a modelagem.

Na sequência, procedeu-se à análise de redundância entre variáveis de identificação e controle, eliminando campos que representavam a mesma informação em formatos distintos (Umargono, Suseno & Gunawan, 2020). A variável `valorTotalHomologado` foi removida em favor de `valorTotalEstimado`, dada a correlação de 0,97 entre ambas e a maior completude desta última. As colunas `modalidadeId` e `modalidadeNome` foram descartadas por apresentarem valor constante em todo o conjunto: consequência natural do recorte analítico adotado.

Por fim, as variáveis categóricas foram transformadas por meio de *One-Hot Encoding*, técnica que representa cada categoria como uma coluna binária, evitando a imposição artificial de ordenação entre categorias (Rodríguez et al., 2018). Após a análise de correlação, variáveis com coeficiente igual ou superior a 0,90 foram removidas com o objetivo de mitigar problemas de multicolinearidade (Kyriazos & Poga, 2023). Como resultado, obteve-se um conjunto final composto por 2.188 observações e 41 variáveis. As etapas de pré-processamento são apresentadas de forma sintética na Tabela 1.

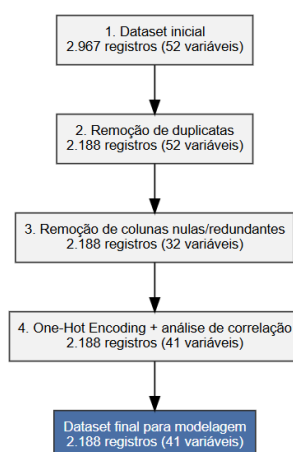
Tabela 1: Síntese do pré-processamento

Etapa	Registros	Variáveis
Dataset inicial	2.967	52
Remoção de duplicatas	2.188	52
Remoção de colunas nulas/redundantes	2.188	32
One-Hot Encoding + análise de correlação	2.188	41

Fonte: Autoria própria.

As etapas de pré-processamento evidenciam as transformações sucessivas do conjunto de dados ao longo das fases de limpeza e preparação. Complementarmente, a Figura 1 apresenta uma representação gráfica desse processo, permitindo uma compreensão mais clara do fluxo sequencial adotado na preparação dos dados.

Figura 1: Fluxograma das etapas de pré-processamento dos dados, desde o conjunto inicial até a definição do *dataset* final para modelagem



Fonte: Autoria própria.

3.4 Algoritmos para Detecção de Anomalias

A seleção dos algoritmos baseou-se no princípio de complementaridade metodológica, operacionalizado por três critérios. Primeiro, a diversidade de paradigmas de detecção: foram escolhidos algoritmos que operam sob premissas distintas: *K-Means*, baseado em distância ao centroide (Sinaga & Yang, 2020; Khetarpaul et al., 2022); *DBSCAN*, baseado em densidade global (Ester et al., 1996; Hanafi & Saadatfar, 2022); *Local Outlier Factor* (LOF), baseado em

densidade local (Breunig et al., 2000; Basheer et al., 2026); e *Isolation Forest*, baseado em isolamento estatístico (Liu et al., 2008; Shao et al., 2025), minimizando a dependência de um único pressuposto sobre a estrutura dos dados. Segundo, a adequação ao contexto de dados tabulares administrativos de alta dimensionalidade, para o qual *Isolation Forest* e *LOF* apresentam desempenho consistente na literatura (Shahrestani & Sanislav, 2024; Shao et al., 2025), enquanto *K-Means* e *DBSCAN* oferecem interpretabilidade por meio de agrupamentos coerentes. Terceiro, a viabilidade computacional: todos estão disponíveis na biblioteca *scikit-learn* com implementações eficientes e amplamente validadas, favorecendo reprodutibilidade e escalabilidade. Algoritmos baseados em redes neurais, como *autoencoders*, foram descartados por demandarem maior volume de dados rotulados para calibração ou apresentarem menor interpretabilidade em contextos de auditoria pública (Garbaccio & Ramos, 2025).

A adoção dessa abordagem combinada (*ensemble*) permite reduzir a incidência de falsos positivos e aumentar a robustez das detecções, ao integrar múltiplos critérios de identificação de anomalias (Garbaccio & Ramos, 2025; Souto et al., 2025). Para viabilizar a análise e visualização dos resultados, aplicou-se a Análise de Componentes Principais (PCA), técnica de redução de dimensionalidade que preserva as direções de maior variância do conjunto original (Li et al., 2025). Os resultados obtidos e a análise comparativa entre os algoritmos são apresentados na seção seguinte.

4 Resultados e Discussão

A aplicação dos quatro algoritmos não supervisionados ao conjunto de 2.188 processos de dispensa de licitação resultou em diferentes volumes e padrões de anomalias detectadas, conforme sintetizado na Tabela 2.

Tabela 2: Anomalias detectadas por algoritmo

Algoritmo	Anomalias detectadas	Proporção (%)
K-Means	499	22,81%
DBSCAN	317	14,49%
LOF (Local Outlier Factor)	219	10,01%
Isolation Forest	217	9,92%

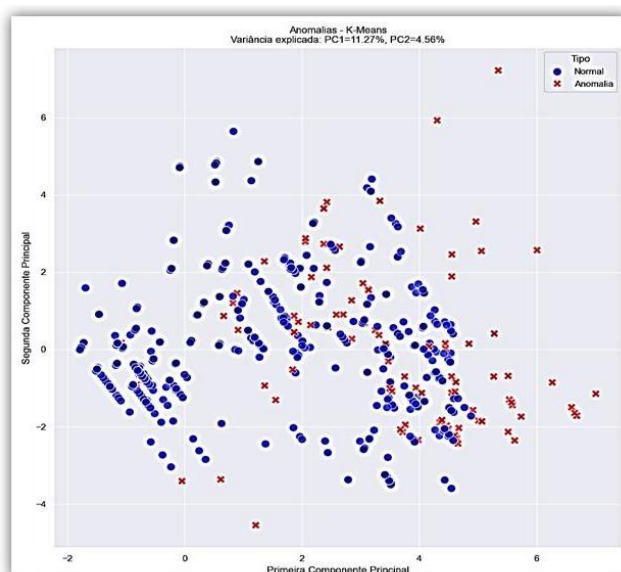
Fonte: Autoria própria.

O *Isolation Forest* e o *LOF* apresentaram proporções mais restritas de detecção, sinalizando aproximadamente 10% dos registros como anômalos, enquanto o *K-Means* registrou a maior taxa (22,81%), comportamento consistente com sua sensibilidade à distância em relação aos centróides. As subseções a seguir detalham os resultados de cada algoritmo, acompanhados de sua representação visual no espaço das duas primeiras componentes principais.

4.1 Anomalias identificadas pelo algoritmo *K-MEANS*

O algoritmo *K-Means* identificou 499 observações anômalas, correspondendo a 22,81% do conjunto analisado. A Figura 2 apresenta a projeção dos dados no espaço das duas primeiras componentes principais, permitindo a visualização da distribuição relativa das observações.

Figura 2: Projeção das observações no espaço das duas primeiras componentes principais (PCA), destacando as anomalias identificadas pelo algoritmo K-Means



Fonte: Autoria própria.

Observa-se que as anomalias (em vermelho) tendem a se distribuir de forma mais dispersa e afastada das regiões de maior concentração de observações normais (em azul), indicando a presença de registros com características significativamente distintas do padrão predominante. Esse comportamento é consistente com a lógica do K-Means, que se baseia na distância em relação aos centróides dos clusters, sendo mais sensível à identificação de desvios globais no espaço de atributos (Sinaga & Yang, 2020; Khetarpaul et al., 2022).

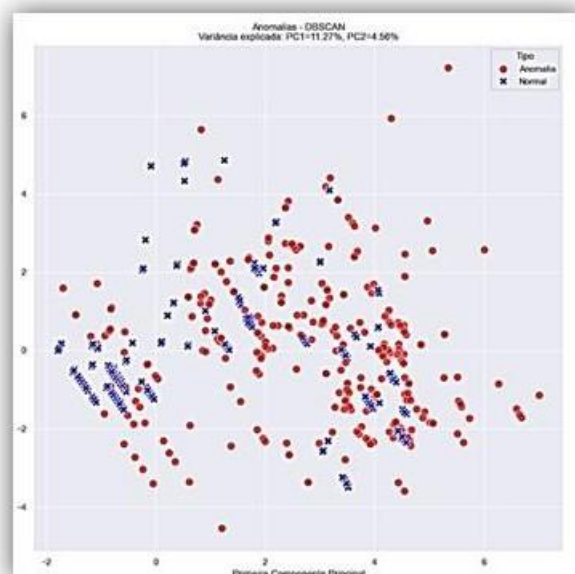
Ressalta-se que as componentes principais representam combinações lineares das variáveis originais, sendo utilizadas neste estudo exclusivamente para fins de visualização (Li et al., 2025). Dessa forma, não se busca interpretar individualmente o significado de cada componente, mas sim analisar a distribuição relativa das observações no espaço reduzido. Adicionalmente, a proporção de variância explicada pelas duas primeiras componentes (PC1 = 11,27% e PC2 = 4,56%) reforça o caráter exploratório da representação gráfica.

Esses resultados indicam que o K-Means é capaz de capturar padrões anômalos associados a desvios estruturais mais amplos nos dados. No entanto, essa característica também pode resultar na identificação de um número mais elevado de observações como anômalas, o que reforça a importância da utilização de abordagens complementares para validação dos achados (Garbaccio & Ramos, 2025; Souto et al., 2025).

4.2 Anomalias identificadas pelo algoritmo DBSCAN

Diferentemente do K-Means, o DBSCAN baseia-se na densidade dos dados, sendo capaz de identificar regiões de baixa densidade como anômalas (Ester et al., 1996; Singh et al., 2022). Na Figura 3, observa-se que as anomalias (em vermelho) tendem a se concentrar nas regiões periféricas do espaço projetado, afastadas das áreas de maior densidade de observações normais (em azul). Esse padrão indica que o algoritmo é particularmente eficaz na identificação de registros que não pertencem a nenhum agrupamento coeso, capturando estruturas mais dispersas nos dados (Hanafi & Saadatfar, 2022).

Figura 3: Projeção das observações no espaço das duas primeiras componentes principais (PCA), destacando as anomalias identificadas pelo algoritmo DBSCAN em regiões de baixa densidade



Fonte: Autoria própria.

Ressalta-se que as componentes principais representam combinações lineares das variáveis originais, sendo utilizadas neste estudo exclusivamente para fins de visualização (Li et al., 2025). Assim, a interpretação dos resultados concentra-se na distribuição relativa das observações no espaço reduzido. A proporção de variância explicada pelas duas primeiras componentes (PC1 = 11,27% e PC2 = 4,56%) reforça o caráter exploratório da representação gráfica.

Os resultados sugerem que o DBSCAN apresenta maior capacidade de identificar anomalias estruturais associadas à baixa densidade, ao mesmo tempo em que mantém maior coesão nos agrupamentos principais, o que contribui para sua superioridade em termos de métricas de qualidade de clusterização (Sari et al., 2021; Hanafi & Saadatfar, 2022). O DBSCAN apresentou o melhor desempenho entre os modelos baseados em agrupamento (Tabela 3). Diferentemente do K-Means, o algoritmo identificou anomalias em regiões de baixa densidade, capturando registros dispersos que não pertencem a nenhum cluster coeso — padrão visível na Figura 3, onde os pontos anômalos concentram-se nas periferias do espaço PCA.

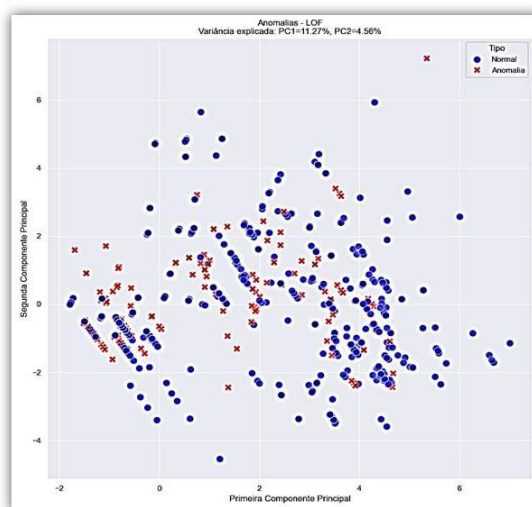
4.3 Anomalias identificadas pelo algoritmo *LOCAL OUTLIER FACTOR (LOF)*

O algoritmo Local Outlier Factor (LOF) identificou 219 observações anômalas, correspondendo a 10,01% do conjunto analisado. Diferentemente dos métodos baseados em padrões globais, o LOF opera com base na densidade relativa de cada observação em relação à sua vizinhança local, sendo capaz de detectar anomalias contextuais (Breunig et al., 2000; Basheer et al., 2026).

Conforme ilustrado na Figura 4, observa-se que as anomalias (em vermelho) frequentemente se situam próximas a regiões densas de observações normais (em azul), indicando que tais registros não se destacam por estarem isolados no espaço global, mas por apresentarem comportamento atípico em relação ao seu contexto local (Shahrestani & Sanislav, 2024). Essa característica distingue o LOF dos métodos globais e o torna

especialmente sensível a irregularidades contextuais, como contratos atípicos dentro de um mesmo perfil de órgão ou município (Basheer et al., 2026; Da Silva Pereira & Dorneles, 2025).

Figura 4: Projeção das observações no espaço das duas primeiras componentes principais (PCA), destacando as anomalias identificadas pelo algoritmo Local Outlier Factor (LOF) com base na densidade relativa local



Fonte: Autoria própria.

Ressalta-se que as componentes principais representam combinações lineares das variáveis originais, sendo utilizadas neste estudo exclusivamente para fins de visualização (Li et al., 2025). Assim, a análise concentra-se na distribuição relativa das observações no espaço reduzido. A proporção de variância explicada pelas duas primeiras componentes (PC1 = 11,27% e PC2 = 4,56%) reforça o caráter exploratório da representação gráfica.

Esses resultados indicam que o LOF complementa os demais algoritmos ao capturar anomalias que não seriam identificadas por abordagens baseadas exclusivamente em desvios globais ou baixa densidade (Shahrestani & Sanislav, 2024), contribuindo para uma análise mais abrangente quando integrado à estratégia de consenso adotada no estudo (Aggarwal, 2016).

4.4 Anomalias identificadas pelo algoritmo *Isolation Forest*

O algoritmo *Isolation Forest* identificou 217 observações anômalas, correspondendo a 9,92% do conjunto analisado, sendo a menor proporção entre os modelos avaliados. Esse resultado reflete o caráter mais conservador do método, que se baseia no isolamento das observações por meio de partições aleatórias no espaço de atributos (Liu et al., 2008; Shao et al., 2025).

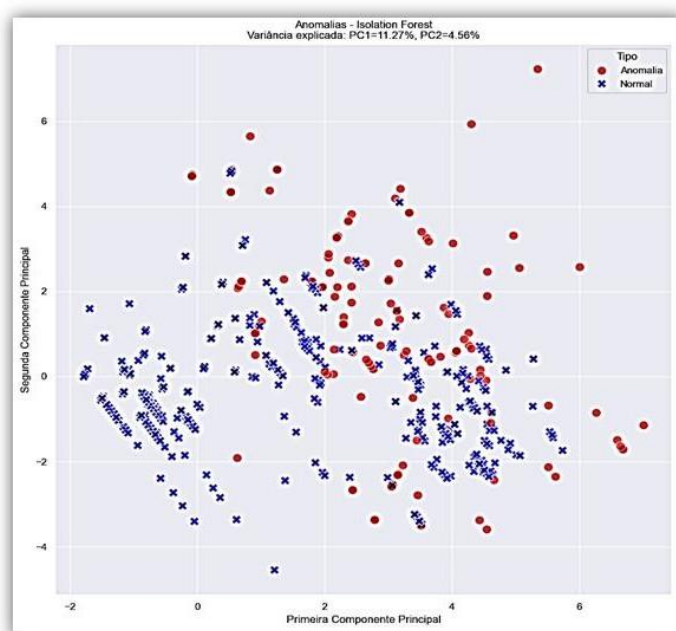
Conforme ilustrado na Figura 5, observa-se que as anomalias (em vermelho) tendem a se concentrar em regiões específicas do espaço projetado, distintas das áreas ocupadas pela maioria das observações normais (em azul). Esse padrão sugere que o algoritmo foi capaz de identificar registros estruturalmente isolados, ou seja, com características que os tornam mais facilmente separáveis no espaço multidimensional original (Bai et al., 2025).

Ressalta-se que as componentes principais representam combinações lineares das variáveis originais, sendo utilizadas neste estudo exclusivamente para fins de visualização (Li et al., 2025). Assim, a análise concentra-se na distribuição relativa das observações no espaço reduzido. A proporção de variância explicada pelas duas primeiras componentes (PC1 = 11,27% e PC2 = 4,56%) reforça o caráter exploratório da representação gráfica.

Os resultados indicam que o *Isolation Forest* apresenta maior rigor na identificação de anomalias, priorizando casos mais evidentes de isolamento estrutural (Shahrestani & Sanislav, 2024). Essa característica contribui para a redução de falsos positivos e complementa as

demais abordagens utilizadas no estudo, reforçando a robustez da estratégia de detecção baseada em consenso (Aggarwal, 2016).

Figura 5: Projeção das observações em PCA com anomalias identificadas pelo Isolation Forest



Fonte: Autoria própria.

4.5 Desempenho Comparativo Dos Modelos

A Tabela 3 consolida as métricas de avaliação dos quatro algoritmos.

Tabela 3: Métricas comparativas de desempenho

Modelo	Silhouette	Calinski-Harabasz	Davies-Bouldin	Anomaly Ratio	Avg Score
K-Means	0,4590	229,32	2,5637	0,2281	—
DBSCAN	0,8922	10.845,02	0,1586	0,1444	—
LOF (Local Outlier Factor)	—	—	—	0,1001	31,40
Isolation Forest	—	—	—	0,0992	—

Fonte: Autoria própria.

Os traços (—) indicam métricas não aplicáveis ao algoritmo em questão. Os resultados apresentados na Tabela 3 evidenciam diferenças significativas no comportamento dos modelos avaliados. O DBSCAN apresentou o melhor desempenho entre os algoritmos de agrupamento, com valores elevados de Silhouette (0,8922) e Calinski-Harabasz (10.845,02), além de baixo índice de Davies-Bouldin (0,1586), indicando alta coesão interna e boa separação entre os clusters formados (Hanafi & Saadatfar, 2022; Sari et al., 2021).

Em contraste, o K-Means apresentou desempenho inferior nessas métricas (Silhouette = 0,4590; Davies-Bouldin = 2,5637), sugerindo menor capacidade de capturar a estrutura natural dos dados (Dalmaijer et al., 2022). Esse comportamento está alinhado à maior proporção de anomalias identificadas por esse modelo (22,81%), indicando maior sensibilidade a desvios globais (Sinaga & Yang, 2020).

No caso do LOF e do Isolation Forest, a ausência de métricas de clustering é compensada pela análise do Anomaly Ratio e, no caso do LOF, pelo Avg Score, que indica maior intensidade média das anomalias detectadas (Basheer et al., 2026). Ambos os modelos apresentaram proporções semelhantes de anomalias (10,01% e 9,92%, respectivamente), sugerindo um comportamento mais conservador em comparação aos métodos de agrupamento (Shahrestani & Sanislav, 2024).

De forma geral, observa-se que os diferentes algoritmos apresentam perfis distintos de detecção, reforçando a importância da abordagem combinada adotada neste estudo, na qual a análise de consenso permite integrar essas diferentes perspectivas para uma identificação mais robusta de anomalias (Aggarwal, 2016; Garbaccio & Ramos, 2025).

4.6 Análise de Consenso

Para aumentar a confiabilidade das sinalizações, foi realizada análise de consenso entre os quatro modelos, classificando cada observação conforme o número de algoritmos que a identificou simultaneamente como anômala (Aggarwal, 2016), conforme Tabela 4.

Tabela 4: Distribuição de anomalias por nível de consenso

Nível de consenso	Nº de algoritmos	Registros	Proporção (%)
Consenso total	4	11	0,50%
Consenso forte	3	65	2,97%
Consenso moderado	2	156	7,13%
Consenso fraco	1	314	14,35%

Fonte: Autoria própria.

Os 11 registros identificados por todos os quatro algoritmos representam os casos de maior prioridade investigativa, dado o alto grau de convergência entre métodos com premissas distintas. Somados aos 65 casos de consenso forte, totalizam 76 processos — 3,47% do conjunto analisado — que merecem atenção prioritária dos órgãos de controle. A predominância de casos com consenso fraco (314 registros) confirma que cada algoritmo captura dimensões complementares de anomalia, reforçando a pertinência da abordagem multimetodológica adotada.

Os resultados obtidos evidenciam que a aplicação combinada de algoritmos não supervisionados é capaz de identificar, de forma sistemática e escalável, processos de dispensa de licitação com perfil estatisticamente atípico. Os 76 casos de consenso total e forte (nos quais ao menos três dos quatro algoritmos convergiram na sinalização de anomalia) representam o subconjunto de maior relevância investigativa, pois a concordância entre métodos de premissas distintas reduz substancialmente a probabilidade de falsos positivos (Aggarwal, 2016).

Do ponto de vista prático, esses 76 processos constituem uma lista prioritária objetiva para triagem por órgãos de controle como o TCU e a CGU, substituindo a seleção manual por um critério analítico reproduzível. Essa capacidade é especialmente relevante no contexto das dispensas emergenciais previstas no Art. 75, inciso VIII, da Lei nº 14.133/2021, modalidade historicamente associada a irregularidades decorrentes de ausência de planejamento ou de criação artificial de urgência (Bezerra et al., 2022; Da Silva & De Oliveira, 2024). A automação da triagem permite que os auditores concentrem esforços nos casos de maior risco, aumentando a eficiência do controle sem ampliação proporcional de recursos.

Embora não sejam aplicáveis medidas tradicionais de incerteza, como intervalos de confiança, no contexto de aprendizado não supervisionado com dados não rotulados (Garbaccio & Ramos, 2025; Souto et al., 2025), a robustez dos resultados foi avaliada por meio da convergência entre diferentes algoritmos. Nesse sentido, a análise de consenso entre os modelos funciona como um mecanismo de validação indireta, indicando maior confiabilidade nos casos identificados simultaneamente por múltiplas abordagens (Aggarwal, 2016).

Cabe ressaltar que as anomalias detectadas constituem indícios, não confirmações de irregularidade, uma vez que a detecção de anomalias identifica desvios estatísticos que demandam investigação complementar para validação substantiva (Júnior et al. 2023; Souto

et al., 2025). Nesse sentido, a abordagem proposta posiciona-se como uma camada inicial de triagem inteligente, complementar (e não substituta) aos métodos tradicionais de auditoria.

5. Conclusão

Esta pesquisa demonstrou que técnicas de Machine Learning não supervisionado constituem uma abordagem viável, eficiente e replicável para a triagem automatizada de indícios de irregularidades em contratações públicas. A partir de 2.188 registros coletados do PNCP e submetidos a quatro algoritmos de detecção de anomalias, a análise de consenso resultou na identificação de 76 processos com alta prioridade investigativa, subconjunto robusto e objetivamente fundamentado para atuação dos órgãos de controle. Esses achados consolidam a viabilidade do *framework* multialgorítmico proposto como ferramenta de triagem inteligente, complementar aos métodos tradicionais de auditoria.

Como principal contribuição, este estudo avança em relação à literatura ao propor uma abordagem multialgorítmica baseada em análise de consenso, capaz de integrar diferentes perspectivas de detecção de anomalias — globais, locais e estruturais — em um único *framework* analítico. Essa estratégia permite reduzir a dependência de um único modelo e aumentar a robustez dos resultados, oferecendo um critério mais objetivo para priorização investigativa em contextos de dados não rotulados.

A análise de consenso entre os modelos revelou 76 processos sinalizados por ao menos três algoritmos simultaneamente, constituindo um subconjunto prioritário para investigação pelos órgãos de controle. Esses casos não representam confirmações de fraude, mas sim indícios objetivos, derivados de padrões estatísticos consistentes, que justificam aprofundamento investigativo.

Do ponto de vista da gestão pública, a abordagem proposta alinha-se diretamente aos princípios de transparência, eficiência e controle social previstos na Lei nº 14.133/2021 (Boege & Marques, 2024; Souza Silva et al., 2025), ao oferecer aos órgãos fiscalizadores uma ferramenta de triagem inteligente que contribui para a alocação mais eficiente dos recursos humanos em atividades de auditoria. A automação da detecção inicial de anomalias permite que os auditores concentrem seus esforços analíticos nos casos de maior risco potencial, aumentando a efetividade dos mecanismos de controle sem ampliação proporcional de custos.

Como limitações, destaca-se que a qualidade dos resultados está diretamente condicionada à completude e consistência dos dados disponibilizados pelo PNCP, bem como ao fato de que desvios estatísticos não correspondem automaticamente a irregularidades jurídicas, demandando validação por meio de análise documental e jurídica complementar (Silva et al., 2023; Nery et al., 2024; Rodrigues et al., 2026), além do fato da inexistência (até o momento) de validação externa com dados rotulados.

Para pesquisas futuras, recomenda-se a ampliação do recorte temporal e temático dos dados, a incorporação de variáveis relacionadas ao histórico dos fornecedores e a exploração de abordagens supervisionadas, à medida que decisões dos Tribunais de Contas possam ser utilizadas como base de rotulagem.

Referências

- Abreu, B. M., Pereira, T. H. S., & Gomes-Jr, L. (2024). Detecção de Fraudes em Licitações Públicas: Uma Comparação de Modelos de Detecção de Anomalias. *Anais Da XIX Escola Regional De Banco De Dados*, 81–90.
- Aggarwal, C. C. (2016). *Outlier analysis*. Springer International Publishing AG
- Arão, C. (2024). Por trás da inteligência artificial: uma análise das bases epistemológicas do *Machine Learning*. *Trans/Form/Ação*, 47(3).
- Bai, M., Yang, B., Wu, Z., & Ma, J. (2025). Adaptive Hybrid Model with CatBoost for Predicting Electric Market Competition. 2025 *International Conference on Communication, Computer, and Information Technology (IC3IT)*, 1–6.
- Basheer, M. Y. I., Ali, A. M., Hamid, N. H. A., Nordin, S., Osman, R., Yusoff, N., & Gu, X. (2026). Adaptive local outlier factor. *Evolving Systems*, 17(2). <https://doi.org/10.1007/s12530-026-09808-y>
- Bazzaz Abkenar, S., Kshani, M., Mahdipour, E. and Jameii, S. (2021), Big data analytics meets social media: a systematic review of techniques, open issues, and future directions, *Telematics and Informatics*, Vol. 57, p. 101517, [doi: 10.1016/j.tele.2020.101517](https://doi.org/10.1016/j.tele.2020.101517)
- Bezerra, A. J. S., Da Silva, F. a. S., Da Silva, F. a. S., Arenas, M. V. D. S., & Arenas, M. V. D. S. (2022). Principais irregularidades em obras públicas no Estado de Rondônia apontadas pelo TCE-RO no período de 2018 a 2021. *Brazilian Journal of Development*, 8(4), 26025–26040. <https://doi.org/10.34117/bjdv8n4-219>
- Boege, M. G. M., & Marques, S. D. R. B. (2024). Compliance e integridade na nova Lei de Licitações: análise do diálogo competitivo como instrumento de eficiência e transparência nas contratações públicas. *Academia De Direito*, 6, 2995–3014. <https://doi.org/10.24302/acaddir.v6.5083>
- BRASIL. Conselho Nacional de Saúde. Resolução nº 510, de 7 de abril de 2016. Brasília: Ministério da Saúde, 2016. Disponível em: https://bvsmis.saude.gov.br/bvs/saudelegis/cns/2016/res0510_07_04_2016.html. Acesso em: 25 abr. 2026.
- BRASIL. Lei nº 12.527, de 18 de novembro de 2011. Presidência da República, 2011. Disponível em: https://www.planalto.gov.br/ccivil_03/_ato2011-2014/2011/lei/l12527.htm. Acesso em: 25 abr. 2026.
- BRASIL. Lei nº 14.133/2021. Disponível em: https://www.planalto.gov.br/ccivil_03/_ato2019-2022/2021/lei/l14133.htm. Acesso em: 18 mar. 2026.
- Breunig, M. M., Kriegel, H., Ng, R. T., & Sander, J. (2000). LOF. *Proceedings of the 2000 ACM SIGMOD International Conference on Management of Data*, 93–104. <https://doi.org/10.1145/342009.335388>
- Dalmaijer, E. S., Nord, C. L., & Astle, D. E. (2022). Statistical power for cluster analysis. *BMC Bioinformatics*, <https://doi.org/10.1186/s12859-022-04675-1>
- Da Silva, J. V. S., & De Oliveira, R. D. (2024). Fraude em licitações: uma perspectiva da lei nº 14.133/2021. *COGNITIONIS Scientific Journal*, 7(2), e500. <https://doi.org/10.38087/2595.8801.500>
- Da Silva Pereira, R., Dantas, J. R. a. F., Santos, J. B. D., De Abreu Mascarenhas, A., & Ney, O. F. (2024). Do portal nacional de contratações públicas (PNCP) e disposições transitórias e finais. *International Journal of Scientific Management and Tourism*, 10(4), e1005.
- Da Silva Pereira, B., & Dorneles, C. F. (2025). Uma Metodologia para Coleta e Mapeamento de Dados de Licitações Públicas dos Portais da Transparência dos Municípios de Santa Catarina. *Anais Da Escola Regional De Banco De Dados (ERBD)*, 50–59. <https://doi.org/10.5753/erbd.2025.7267>
- De Almeida, J., & Chinelato, C. (2025). Identification of electromyographic signals using machine learning techniques and low-cost technologies. *Revista Para Graduandos/Instituto Federal De Educação, Ciência E Tecnologia De São Paulo - Campus São Paulo - REGRASP*, 10(3), 18-31. <https://doi.org/10.47734/regrasp.v10.03.p18-31>
- Dzeco, E. A., & Simione, A. A. (2024). Processo De Procurement Na Conjuntura De Crise Da Pandemia Da Covid-19 E A Transparência Na Contratação Pública Em Moçambique. *Estudo & Debate*, 31(2).
- Ester M, Kriegel H-P, Sander J, Xu X. (1996) A density-based algorithm for discovering clusters in large spatial databases with noise KDD'96: Proceedings of the Second International Conference on Knowledge Discovery and Data Mining. Pages 226 – 231 <https://dl.acm.org/doi/10.5555/3001460.3001507>
- Filho, J. R. C. (2025). Princípios norteadores das licitações públicas à luz da lei nº 14.133/2021. *LUMEN ET VIRTUS*, 16(55), e11292. <https://doi.org/10.56238/levv16n55-109>
- Garbaccio, G. L., & Ramos, D. F. (2025). Utilização da inteligência artificial na gestão pública para identificação de irregularidades em licitações e contratos: experiências brasileira e portuguesa. <http://dx.doi.org/10.26668/revistajur.2316-753X.v3i83.7741>
- Hanafi, N., & Saadatfar, H. (2022). A fast DBSCAN algorithm for big data based on efficient density calculation. *Expert Systems With Applications*, 203, 117501.
- Júnior, D. P. G., De Sousa Filho, G. F., & Cabral, L. D. a. F. (2023). Classificação de fraudes em licitações

- públicas através do agrupamento de empresas em conluíus. *Anais Do Latin American Symposium on Digital Government (LASDiGov)*, 13–24. <https://doi.org/10.5753/wcge.2023.229519>
- Khetarpaul, S., Sharma, D., Jose, J. I., & Saragur, M. (2022). Real-Time detection and visualization of traffic conditions by mining Twitter data. In *Lecture notes in computer science* (pp. 141–152). https://doi.org/10.1007/978-3-031-15512-3_11
- Kyriazos, T., Poga, M. (2023) Dealing with Multicollinearity in Factor Analysis: The Problem, Detections, and Solutions. *Open Journal of Statistics*, 13, 404-424.
- Liu, F. T., Ting, K. M., & Zhou, Z. (2008). Isolation Forest. *2008 Eighth IEEE International Conference on Data Mining*, 413–422. <https://doi.org/10.1109/icdm.2008.17>
- Li, Y., Li, S., Xiong, X., Hao, Z., Liu, Y., & Mu, W. (2025). A prediction method for wine customer consumption behavior using data enhancement and attention-based CNN-LSTM. *British Food Journal*. Vol. ahead-of-print No. ahead-of-print. <https://doi.org/10.1108/BFJ-07-2025-0899>
- Marconi, M. A.; Lakatos, E. M. Fundamentos de metodologia científica. 8. ed. São Paulo: Atlas, 2017.
- Martins, R. S. N., Fêlix, M. P., & Ianuzzi, A. M. M. (2025). Princípios fundamentais das licitações públicas no Brasil. *Caderno Pedagógico*, 22(13), e21216.
- McKinney, W. (2011). Pandas: A Foundational Python Library for Data Analysis and Statistics. *Python for High Performance and Scientific Computing*, 14, 1-9.
- Nery, E. L. M., De Souza Nunes, L., & De Barbuda, A. S. (2024). ATOS DE IMPROBIDADE ADMINISTRATIVA EM LICITAÇÕES NO DIREITO BRASILEIRO. *Revista Jurídica Do Nordeste Mineiro*, 3(3). <https://doi.org/10.61164/rjnm.v3i3.2222>
- Pavão, J. N., Brandão, D., & Belloze, K. (2025). Large Language Models para detecção de conluíus em licitações. *Anais Do XL Simpósio Brasileiro De Bancos De Dados*, 893–899.
- Reis, I. O., Castro, R. I. F., & Piffer, D. M. (2025). Lei Da Transparência E Legitimidade Em Licitações De Obras Públicas. *Revista Gestão E Conhecimento*, 19(1), e386.
- Rodrigues, H. D. S., Santos, R. S. D., & De Santana Neto, H. G. (2026). Responsabilidade Do Agente Público E Do Contratado Nas Licitações: Análise Da Jurisprudência Após A Lindb (ART. 28). *Revista Ibero-Americana De Humanidades, Ciências E Educação*, 12(2), 1–9.
- Rodríguez, P., Bautista, M. A., Gonzàlez, J., & Escalera, S. (2018). Beyond one-hot encoding: Lower dimensional target embedding. *Image and Vision Computing*, 75, 21–31. <https://doi.org/10.1016/j.imavis.2018.04.004>
- Sari, I. P., Al-Khowarizmi, A., & Batubara, I. H. (2021). Cluster Analysis Using K-Means Algorithm and Fuzzy C- Means Clustering For Grouping Students' Abilities In Online Learning Process. *Journal of Computer Science, Information Technology and Telecommunication Engineering*, 2(1), 139–144.
- Shao, M., Shao, H., Wang, X., Gao, Y., & Liu, B. (2025). Interpretable anomaly detection using extended isolation forest with adaptive thresholds. *Structural Health Monitoring*. <https://doi.org/10.1177/14759217251339607>
- Shahrestani, S., & Sanislav, I. (2024). How does dimensionality influence outlier detection effectiveness in multivariate geochemical data? insights from LOF and IF methods. *Earth Science Informatics*, 18(1). <https://doi.org/10.1007/s12145-024-01611-0>
- Silva, M. O., Costa, L. L., Bezerra, G., Gomide, L. D., Hott, H. R., Oliveira, G. P., Brandão, M. A., Lacerda, A., & Pappa, G. (2023). Análise de Sobrepreço em Itens de Licitações Públicas. *Anais Do Latin American Symposium on Digital Government (LASDiGov)*, 118–129.
- Sinaga, K. P., & Yang, M. (2020). Unsupervised K-Means clustering algorithm. *IEEE Access*, 8, 80716–80727. <https://doi.org/10.1109/access.2020.2988796>
- Singh, H. V., Girdhar, A., & Dahiya, S. (2022). A Literature survey based on DBSCAN algorithms. *2022 6th International Conference on Intelligent Computing and Control Systems (ICICCS)*, 751–758. <https://doi.org/10.1109/iciccs53718.2022.9788440>
- Teixeira, C. E., Bonin, G., & Campo, A. (2019). Aplicação De Técnicas De Machine Learning Para Controle De Uma Mão Robótica. *Revista Para Graduandos/Instituto Federal De Educação, Ciência E Tecnologia De São Paulo - Campus São Paulo - REGRASP*, 4(2), 116-128
- Umargono, E., Suseno, J. E., & Gunawan, S. K. V. (2020). K-Means Clustering Optimization Using the Elbow Method and Early Centroid Determination Based on Mean and Median Formula. *Intelligent Automation & Soft Computing*. <https://doi.org/10.2991/assehr.k.201010.019>
- Souto, A.L, Gomes, V. C., & Riveros, J. L. T. (2025). Inteligência Artificial em Auditoria de Licitações. *Síntese.*, 2(1). <https://doi.org/10.70690/f9axe22>
- Souza Silva, E. M., Cabrinha, N. M. G., Da Silva Thompson, N. P., & Costa, L. C. P. (2025). A nova lei de licitações e contratos (lei no 14.133/2021): impactos no planejamento, eficiência e transparência dos processos licitatórios nos municípios da administração pública do estado do Amazonas. *Revista Políticas Públicas & Cidades*, 14(3), e1946. <https://doi.org/10.23900/2359-1552v14n3-41-2025>.