



AUTOMAÇÃO DE PROCESSOS JURÍDICOS COM IA: UMA ARQUITETURA BASEADA EM AGENTES AUTÔNOMOS PARA DISTRIBUIÇÃO E CLASSIFICAÇÃO DE INTIMAÇÕES

SIMPEP - simpep.feb@unesp.br

UNIVERSIDADE ESTADUAL PAULISTA - UNESP - BAURU-FEB

COMISSÃO ORGANIZADORA DO SIMPEP - simpep.feb.unesp@gmail.com

UNIVERSIDADE ESTADUAL PAULISTA - UNESP - BAURU-FEB

DEPARTAMENTO DE ENGENHARIA DE PRODUÇÃO - dep.feb@unesp.br

UNIVERSIDADE ESTADUAL PAULISTA - UNESP - BAURU-FEB

Victor Hugo Aires Pimentel da Silva¹ - E-mail: victorpimente@gmail.com

José Alfredo Ferreira Costa¹ - E-mail: jafcosta@gmail.com

Esdras Daniel de Sousa Araujo Silva¹ - E-mail: esdras.sas@gmail.com

Francisco Edilvo Nunes Lima Filho¹ - E-mail: edilvolima@gmail.com

Universidade Federal do Rio Grande do Norte¹ - UFRN.

ÁREA: 03. PESQUISA OPERACIONAL

SUBÁREA: 3.7 - INTELIGÊNCIA COMPUTACIONAL

PALAVRAS-CHAVES: INTELIGÊNCIA ARTIFICIAL; NLP; MACHINE LEARNING.

1. INTRODUÇÃO

A busca por eficiência na gestão pública tem incentivado a adoção de sistemas inteligentes para automação de processos jurídicos. Nesse contexto, a integração de técnicas de Processamento de Linguagem Natural (PLN) e Aprendizado de Máquina (ML) em arquiteturas de agentes autônomos tem se mostrado uma abordagem promissora. Agentes inteligentes capazes de classificar, priorizar e resumir processos jurídicos podem contribuir significativamente para a otimização do fluxo de trabalho, reduzindo a carga operacional de servidores e aumentando a agilidade no tratamento de demandas.

A Procuradoria Fiscal do Município de Natal enfrenta desafios operacionais devido ao alto volume de peças jurídicas, muitas delas genéricas (ex.: modelos padronizados), que demanda tratamento diferenciado. A automação dessa triagem reduziria custos e agilizaria processos, liberando servidores para atividades complexas.

Este trabalho foca na classificação binária de peças jurídicas entre textos "genéricos" e "não genéricos" no âmbito da Procuradoria Fiscal do Município de Natal. A identificação automatizada de textos genéricos permitirá que agentes de IA (inteligência artificial) atuem sobre todo o processo associado e os tratem com menor intervenção humana. Para isso, foram empregadas técnicas de representação textual e avaliados diferentes algoritmos de classificação supervisionada.

Os resultados visam subsidiar a integração do modelo selecionado e adequar para integração futura no sistema multi-agente em desenvolvimento.

2. METODOLOGIA

A pesquisa foi conduzida por meio de uma abordagem aplicada, quantitativa e experimental, com foco na classificação de documentos jurídicos no contexto da Procuradoria Fiscal do Município de Natal/RN. A metodologia foi dividida em quatro etapas: (i) coleta e pré-processamento dos dados; (ii) representação vetorial dos textos; (iii) modelagem e treinamento dos algoritmos de classificação; e (iv) implementação das soluções em ambiente computacional.

2.1. Coleta e Pré-processamento dos Dados

A primeira etapa consistiu na construção do conjunto de dados. Foram inicialmente disponibilizados 23 documentos jurídicos previamente classificados como pertencentes à categoria “genérico” (classe positiva) por especialistas do setor de Apoio Fiscal da Procuradoria. Visando ampliar esse conjunto de amostras positivas, adotou-se uma estratégia de expansão supervisionada baseada em similaridade semântica. Para tal, foram geradas embeddings por um modelo BERT especializado em linguagem jurídica e ajustado para o português, foi calculada a similaridade cossenoidal entre os textos rotulados e um corpus histórico de documentos jurídicos.

A partir dessa análise, foram identificados 606 textos com maior similaridade semântica. Importante ressaltar que os textos selecionados, apesar de semanticamente próximos, não consistem em duplicatas exatas dos originais, podendo apresentar variações estruturais sutis, como alterações de pontuação.

Os 606 textos foram então submetidos à validação manual por um especialista jurídico, resultando em um conjunto final com 350 amostras da classe positiva (genérico) e 256 amostras da classe negativa (não genérico).

2.2. Representação Vetorial dos Textos

Embora os embeddings tenham sido utilizados na etapa de expansão semântica, para a modelagem supervisionada optou-se pela aplicação da técnica TF-IDF (Term Frequency-Inverse Document Frequency), visando maior interpretabilidade dos vetores gerados. Que transforma os textos em representações vetoriais esparsas baseadas na importância relativa dos termos. Foi utilizada a implementação do TfidfVectorizer da biblioteca Scikit-learn

2.3. Modelagem e Treinamento dos Algoritmos de Classificação

Para a tarefa de classificação binária dos textos jurídicos, foram avaliados três algoritmos supervisionados: Random Forest, Support Vector Machine (SVM) e Árvore de Decisão. Cada modelo foi integrado a um pipeline composto pelas etapas de vetorização TF-IDF e ajuste do classificador, garantindo replicabilidade e isolamento do fluxo de dados em cada experimento.

A partir da avaliação do desempenho dos modelos foi conduzida por meio de validação cruzada estratificada com 5 folds (StratifiedKFold), de modo a preservar a proporção entre as classes em cada subdivisão dos dados. Para cada iteração o pipeline era ajustado apenas com

os dados de treinamento, garantindo ausência de vazamento de dados (data leakage).

2.4. Implementação Das Soluções Em Ambiente Computacional

O pipeline foi desenvolvido em Python, utilizando bibliotecas como Scikit-learn, Pandas, NumPy e Transformers (HuggingFace). O experimento foi documentado e versionado para reprodutibilidade futura e integração com a arquitetura baseada em agentes autônomos que está sendo desenvolvida.

3. RESULTADOS

As métricas para classificação analisadas incluíram: acurácia, precisão, recall, F1-score e AUC (Área sob a Curva ROC), com médias e desvios padrão calculados ao longo das 5 partições. A tabela 1 apresenta os resultados obtidos para cada classificador.

Tabela 1 - Resultado da classificação usando validação cruzada

	Acurácia	Precisão	Recall	F1-Score	AUC
Random Forrest	90.82% (±0.96%)	93.74% (±3.40%)	90.29% (±3.55%)	91.85% (±0.86%)	0.97 (±0.01)
SVM	92.30% (±2.24%)	94.50% (±2.78%)	92.00% (±2.23%)	93.20% (±1.96%)	0.92 (±0.02)
Decision Tree	87.87% (±1.09%)	90.42% (±1.79%)	88.29% (±2.62%)	89.30% (±1.05%)	0.88 (±0.01)

Fonte: autoria própria (2025).

Os resultados demonstram que todos os modelos apresentaram desempenho satisfatório, com destaque para o SVM, que obteve a maior acurácia, precisão, recall e F1-score entre os modelos testados. O Random Forest apresentou o maior valor de AUC, indicando elevada capacidade discriminativa. No entanto, esse desempenho em AUC deve ser interpretado com cautela, uma vez que o Random Forest foi o único modelo, dentre os três avaliados, que forneceu diretamente probabilidades de classe, o que favorece seu desempenho nessa métrica específica. A Árvore de Decisão, embora com métricas ligeiramente inferiores, demonstrou desempenho competitivo. Esses resultados fornecem subsídios para a escolha de modelos conforme critérios de desempenho, interpretabilidade ou eficiência nas aplicações futuras.

4. CONCLUSÃO

A metodologia adotada, combinando curadoria especializada, expansão supervisionada via similaridade semântica proporcionou uma base confiável para treinamento dos modelos e validação dos resultados. A validação cruzada é crucial para medir a robustez dos modelos para que os agentes de IA atuem com precisão e previsibilidade em ambientes institucionais sensíveis, como a gestão de processos jurídicos.

Após a classificação automática dos textos jurídicos como “genéricos” por meio de modelos de aprendizado de máquina, os agentes de IA assumem o processamento subsequente dessas peças, atuando na geração de resumos, minutas e outras tarefas de apoio à atuação jurídica. A identificação automatizada funciona como um filtro preliminar que direciona o fluxo de documentos para tratamento otimizado, permitindo que os agentes autônomos atuem com foco e eficiência na produção de conteúdo jurídico com base nos dados classificados.

Concluimos que os modelos avaliados são adequados para a tarefa de classificação binária de textos jurídicos. O classificador SVM apresentou o melhor equilíbrio entre precisão e recall, o que é especialmente relevante para aplicações em sistemas autônomos que devem evitar tanto a omissão de casos relevantes quanto classificações equivocadas. O Random Forest, por sua vez, destacou-se na métrica de AUC, favorecido por sua capacidade nativa de estimar probabilidades, o que pode ser explorado por agentes que tomam decisões baseadas em níveis de confiança.

5. REFERÊNCIAS

Classificação com Máquina de Vetores de Suporte (SVM) - Análise Macro. [s.d.].

Disponível em: <https://analisemacro.com.br/econometria-e-machine-learning/classificacao-com-maquina-de-vetores-de-suporte-svm/>. Acesso em: 10 jun. 2025.

COMPARAÇÃO ENTRE AS TÉCNICAS DE REGRESSÃO LOGÍSTICA, ÁRVORE DE DECISÃO, BAGGING E RANDOM FOREST APLICADAS A UM ESTUDO DE - Estatística e Ciência de Dados. [s.d.]. Disponível em: https://coordest.ufpr.br/wp-content/uploads/2018/12/TCC_DanielEricson.pdf. Acesso em: 10 jun. 2025.

FERREIRA, M. H. P. Classificação de peças processuais jurídicas: Inteligência Artificial no Direito. 2018. 78 f. Trabalho de Conclusão de Curso (Bacharelado em Engenharia de Software) – Faculdade UnB Gama, Universidade de Brasília, Brasília, 2018.

GOOGLE. Accuracy, Precision, Recall, and F1. Google Machine Learning Crash Course. [s.d.]. Disponível em: <https://developers.google.com/machine-learning/crash-course/classification/accuracy-precision-recall?hl=pt-br>. Acesso em: 23 maio 2025.

NOGUTI, M. Y. Comparison of Natural Language Processing Algorithms Applied to Small Supervised Datasets in the Legal Domain. 2022. Dissertação (Mestrado em Informática) - Setor de Ciências Exatas, Universidade Federal do Paraná, Curitiba, 2022.

Disponível em: <https://www.inf.ufpr.br/lesoliveira/download/MarianaMSC.pdf>. Acesso em: [data de acesso].

PEDREGOSA, F. et al. Scikit-learn: Machine Learning in Python. **Journal of Machine Learning Research**, v. 12, p. 2825-2830, 2011.

REIMERS, N.; GUREVYCH, I. Sentence-BERT: Sentence Embeddings using Siamese BERT-Networks. In: PROCEEDINGS OF THE 2019 CONFERENCE ON EMPIRICAL METHODS IN NATURAL LANGUAGE PROCESSING AND THE 9TH INTERNATIONAL JOINT CONFERENCE ON NATURAL LANGUAGE PROCESSING (EMNLP-IJCNLP), 2019, Hong Kong, China. **Association for Computational Linguistics**, 2019. p. 3982-3992.

SOUZA, F.; NOGUEIRA, R.; LOTUFO, R. BERTimbau: Pretrained BERT Models for Brazilian Portuguese. In: CERRI, R.; PRATI, R. C. (Ed.). **Intelligent Systems**. Cham: Springer International Publishing, 2020. p. 403-417. (Lecture Notes in Computer Science, v. 12319).