

SISTEMA DE RECOMENDAÇÃO JURISPRUDENCIAL BASEADO EM INTELIGÊNCIA ARTIFICIAL PARA O TRIBUNAL DE JUSTIÇA DO ESTADO DE SERGIPE

SISTEMA DE RECOMENDAÇÃO JURISPRUDENCIAL BASEADO EM INTELIGÊNCIA ARTIFICIAL PARA O TRIBUNAL DE JUSTIÇA DO ESTADO DE SERGIPE

Relatório técnico apresentado pelo(a) mestrando(a) Max Ricardo Borges Ribeiro ao Mestrado Profissional em Administração Pública em Rede, sob orientação do(a) docente Glauco de Figueiredo Carneiro, como parte dos requisitos para obtenção do título de Mestre em Administração Pública.

Contextualização e Justificativa 04

Fundamentação Teórica 05

Definição dos Dados 07

Arquitetura da Solução Técnica 09

Benefícios e Impactos Esperados 11

Desafios e Lições Aprendidas 13

Recomendações e Próximos Passos 15

Considerações Finais 13

Responsáveis pela proposta de intervenção 17

Referências 19

CONTEXTUALIZAÇÃO E JUSTIFICATIVA

O Tribunal de Justiça de Sergipe (TJSE), assim como outros tribunais estaduais, enfrenta um cenário de elevada judicialização, com demandas crescentes e um número limitado de recursos humanos e tecnológicos. Nos juizados cíveis, essa realidade é agravada pela natureza repetitiva de grande parte dos processos, tais quais questões de consumo, telefonia, serviços bancários e planos de saúde, cuja resolução poderia ser acelerada caso houvesse fácil acesso a decisões anteriores similares.

No Tribunal de Justiça do Estado de Sergipe (TJSE), a estrutura de uma vara judicial é tipicamente organizada em dois setores principais: a Secretaria, responsável pela movimentação e trâmite processual, e o Gabinete, focado na elaboração intelectual dos atos judiciais e decisões. A Secretaria é composta por técnicos judiciários e um chefe/diretor de secretaria ou escrivão, enquanto os Gabinetes são formados pelo juiz da vara e seus assessores.

Historicamente, diversas automações foram aplicadas na movimentação processual das Secretarias, buscando otimizar o fluxo de trabalho e o tratamento de grandes volumes de documentos. No entanto, em virtude desse avanço, o aumento da celeridade nas Secretarias tem resultado, em muitos casos, em um crescimento do estoque de processos nos Gabinetes, que demandam análise e decisão, exacerbando a carga de trabalho de magistrados e assessores.

Nesse cenário, o Produto Técnico, um sistema de recomendação jurisprudencial baseado em inteligência artificial, surge como uma solução complementar fundamental. Seu objetivo primordial é melhorar a eficiência e a produtividade no Gabinete, utilizando processamento de linguagem natural e técnicas de similaridade vetorial para identificar acórdãos e precedentes relevantes a partir das petições iniciais entrantes. Isso visa liberar o tempo dos magistrados para se concentrarem em casos mais complexos, além de promover maior consistência e uniformização das decisões judiciais, atendendo a uma necessidade prioritária identificada pelo TJSE.

A morosidade decorrente da necessidade de análise manual de jurisprudência compromete a eficiência, além de propiciar disparidades decisórias entre varas e juízes que enfrentam situações jurídicas idênticas. Tal variabilidade gera insegurança jurídica e incentiva recursos desnecessários às Turmas Recursais.

Nesse contexto, a aplicação de inteligência artificial (IA) não se limita a uma solução tecnológica, mas representa uma transformação profunda no modo como o Judiciário lida com demandas repetitivas. O sistema proposto responde a desafios críticos ao:

- Automatizar a identificação de jurisprudência relevante com base em similaridade semântica;
- Favorecer a uniformização das decisões judiciais ao integrar acórdãos similares diretamente ao processo decisório de primeira instância;
- Liberar tempo dos magistrados para focar em casos mais complexos, onde a análise humana é indispensável;
- Promover maior previsibilidade e segurança jurídica para os jurisdicionados.

O produto que está sendo desenvolvido e implantado no Tribunal de Justiça do Estado de Sergipe (TJSE) é uma solução de análise de similaridade de documentos jurídicos baseada em inteligência artificial. Essa tecnologia visa otimizar o processamento das petições iniciais nos Juizados Especiais Cíveis, permitindo que os documentos sejam comparados automaticamente com acórdãos anteriores, auxiliando na identificação de precedentes e na padronização das decisões judiciais.

FUNDAMENTAÇÃO TEÓRICA

A base teórica da solução encontra-se na interseção entre o Direito e a ciência da computação, mais especificamente na área de inteligência artificial aplicada ao processamento de linguagem natural (PLN). O PLN permite a compreensão computacional de textos, essencial para analisar documentos jurídicos em linguagem natural, que apresentam estrutura, vocabulário e formalismo próprios.

O Processamento de Linguagem Natural (PLN) tem sido amplamente explorado na área da inteligência artificial para tarefas complexas de análise de textos (JURAFSKY; MARTIN, 2021). No contexto jurídico, onde a precisão e a coerência são fundamentais, o uso de técnicas de PLN permite a extração e organização de informações de maneira mais eficiente. A análise de similaridade vetorial de documentos jurídicos é uma aplicação que busca comparar textos com base em sua estrutura semântica, proporcionando suporte à pesquisa jurisprudencial, identificação de precedentes e categorização de casos. Nesse sentido, modelos baseados em embeddings, como os gerados pelo BERT, têm demonstrado desempenho superior ao lidar com nuances linguísticas e contextos específicos dos textos legais (DEVLIN et al., 2019).

O BERT revolucionou a forma como os modelos de PLN capturam relações entre palavras em frases, empregando uma abordagem baseada em aprendizado profundo para gerar representações semânticas de textos. Diferente de métodos tradicionais, como o TF-IDF ou a análise de tópicos, que dependem de características explícitas dos documentos, os embeddings do BERT fornecem vetores que encapsulam significados contextuais. Isso é particularmente útil na área jurídica, onde a interpretação de textos depende não apenas das palavras isoladas, mas também da maneira como elas se relacionam dentro de um corpus normativo e jurisprudencial. Com essa abordagem, é possível medir a similaridade entre documentos com maior precisão e relevância.

A aplicação da análise de similaridade vetorial com embeddings gerados pelo BERT se dá por meio da transformação dos documentos jurídicos em vetores multidimensionais. Uma vez representados numericamente, esses textos podem ser comparados utilizando métricas como a distância de cosseno ou técnicas de aprendizado supervisionado para agrupamento (MIKOLOV et al., 2013). Dessa forma, sistemas de inteligência artificial podem sugerir decisões jurídicas semelhantes, identificar padrões argumentativos e até auxiliar na automatização de tarefas jurídicas repetitivas. Além disso, essa abordagem permite a integração de grandes volumes de informação sem comprometer a acurácia das análises, possibilitando uma melhoria na eficiência dos profissionais da área.

A fundamentação teórica para o uso dessas técnicas em documentos jurídicos passa pelo reconhecimento da complexidade e especificidade do vocabulário legal. A linguagem jurídica é frequentemente caracterizada por expressões ambíguas e construções formais que exigem um tratamento diferenciado na modelagem computacional. O treinamento de modelos baseados em BERT para esse tipo de aplicação requer ajustes finos e curadoria de dados representativos do domínio jurídico. Estudos nessa área têm mostrado que adaptações específicas do modelo podem aprimorar sua capacidade de captar relações entre conceitos legais, tornando a análise de similaridade vetorial uma ferramenta robusta para aprimorar a pesquisa jurídica assistida por inteligência artificial (DEVLIN et al., 2019).

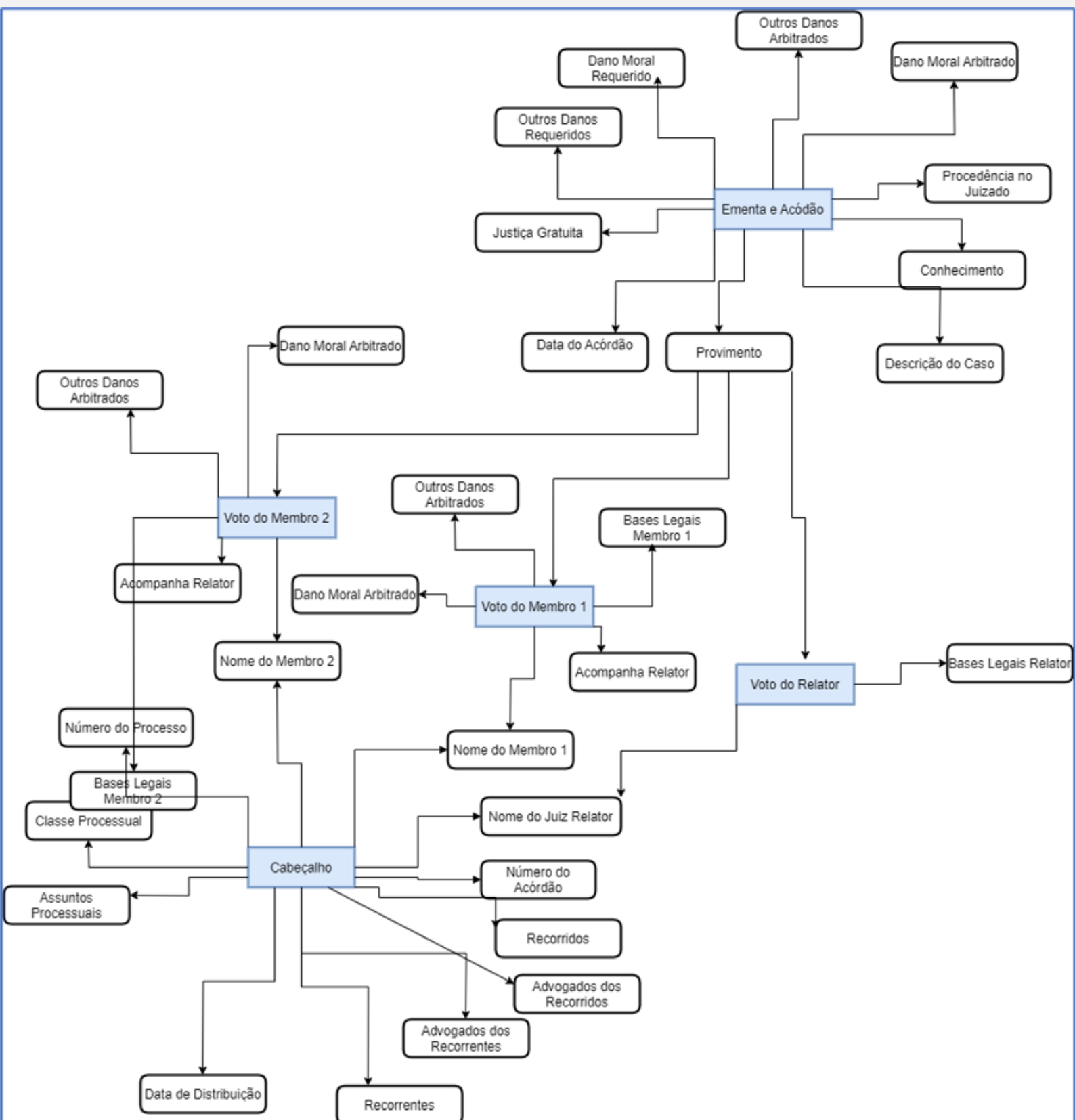
DEFINIÇÃO DOS DADOS

A definição dos dados a serem extraídos dos acórdãos judiciais para a análise das petições iniciais que ingressam nos gabinetes dos Juizados Especiais Cíveis do Tribunal de Justiça do Estado de Sergipe (TJSE) foi realizada com base em reuniões estratégicas com especialistas das áreas negociais. Esse processo colaborativo teve como objetivo garantir que a estruturação dos dados e o treinamento da inteligência artificial fossem conduzidos de forma assertiva, possibilitando maior precisão na interpretação e categorização das informações jurídicas.

Durante as reuniões, especialistas do setor jurídico e tecnológico analisaram o fluxo processual das petições iniciais, identificando os elementos essenciais que poderiam contribuir para a padronização e otimização do tratamento desses documentos. Entre os principais dados considerados estão: identificação das partes envolvidas, descrição da causa de pedir, fundamentação jurídica, pedidos formulados, dispositivos legais citados e o desfecho do acórdão em situações correlatas. A inclusão desses elementos visa estabelecer um modelo de análise que possa correlacionar novos casos com decisões preexistentes, contribuindo para maior eficiência na triagem e encaminhamento dos processos. Além da definição dos dados essenciais, as discussões também abordaram a estruturação das informações para o treinamento da

inteligência artificial. Para que os algoritmos de aprendizado profundo, especialmente os baseados em Processamento de Linguagem Natural (PLN), possam operar com precisão, os documentos precisam ser organizados em um formato que facilite a extração de padrões. Assim, foi estabelecida a necessidade de pré-processamento textual, segmentação de informações e enriquecimento semântico dos dados, permitindo que o modelo compreenda as relações jurídicas de forma contextualizada e comparativa.

A partir dessas reuniões estratégicas, foi elaborado um mapa mental (figura 1) contendo todos os dados identificados como essenciais para a análise das petições iniciais. Esse mapa serviu como uma representação visual da estrutura da informação, facilitando a compreensão dos vínculos entre os diferentes elementos jurídicos e auxiliando na organização do processo de treinamento da inteligência artificial. Para garantir a precisão e a utilidade dos dados definidos, o mapa mental foi validado pelas áreas finalísticas, assegurando que sua aplicação estivesse alinhada com as necessidades reais dos gabinetes dos Juizados Especiais Cíveis do TJSE. Esse ciclo de validação e ajustes permitiu refinamentos importantes na abordagem adotada, consolidando a fundamentação teórica e prática da solução desenvolvida.



ARQUITETURA DA SOLUÇÃO TÉCNICA

A implementação da solução de análise de petições iniciais nos Juizados Especiais Cíveis do TJSE exige uma arquitetura robusta e estruturada, que combine técnicas avançadas de Processamento de Linguagem Natural (PLN) e inteligência artificial para garantir a precisão na identificação e correlação de documentos jurídicos. Para isso, foi definida uma infraestrutura que contempla desde a escolha e configuração de modelos como o BERT até a utilização de sistemas de armazenamento vetorial, como Milvus, para otimizar a busca por similaridade textual. Além disso, diferentes datasets foram integrados para validar a abordagem, permitindo que a solução seja ajustada progressivamente com base em testes práticos. A organização sequencial das atividades detalhadas a seguir visa assegurar que cada etapa do desenvolvimento contribua para uma implementação eficiente, confiável e alinhada às necessidades do TJSE.

PASSOS DE IMPLEMENTAÇÃO

Extração e Armazenamento dos Acórdãos Estruturados

A primeira etapa do projeto envolveu a extração, transformação e carga (ETL) de acórdãos estruturados em HTML a partir da base de Business Intelligence (BI). Esse processo se apresenta na Figura 2 e foi fundamental para consolidar um conjunto robusto de dados jurídicos que serviria de referência para a análise e treinamento da inteligência artificial.

Inicialmente, foram identificados e selecionados os acórdãos armazenados na base de BI, priorizando documentos que apresentassem estruturação adequada e informações relevantes para a classificação e comparação das petições iniciais. A extração desses documentos foi realizada por meio de scripts automatizados que capturaram os conteúdos em formato HTML, preservando a formatação original e os elementos estruturais essenciais.

Após a extração, os dados passaram por uma fase de transformação, onde foram limpos, padronizados e convertidos em texto estruturado. Esse processo envolveu a remoção de códigos desnecessários, identificação de seções específicas dos acórdãos e adequação das informações para posterior análise semântica. Com isso, foi possível assegurar que os dados estivessem organizados de maneira eficiente para serem utilizados pelos modelos de Processamento de Linguagem Natural (PLN).

Por fim, os acórdãos foram transferidos para um banco de dados PostgreSQL, criado exclusivamente para este projeto. Esse banco foi configurado para suportar consultas eficientes e permitir a indexação das informações jurídicas, facilitando o acesso e a recuperação dos textos armazenados. No total, mais de 3.000 acórdãos foram processados e inseridos no banco, garantindo uma base sólida para as próximas etapas da implementação da solução de inteligência artificial.

Essa estruturação inicial proporcionou um ambiente de dados confiável e bem definido, essencial para que os modelos de IA possam operar com precisão na análise e comparação de documentos jurídicos.



➤ PASSOS

Configuração do Modelo BERT

O primeiro passo envolveu a utilização do modelo BERT. Esse modelo foi usado como base para geração de embeddings semânticos, fundamentais para a análise e compreensão dos textos jurídicos. O processo inclui a instalação do modelo, verificação da compatibilidade com o ambiente de execução e configuração inicial para recebimento de consultas.

Configuração da Solução Milvus

O próximo passo foi configurar a solução de armazenamento vetorial utilizando Milvus, conforme as orientações disponíveis na URL mencionada. Milvus foi utilizado para armazenar as embeddings geradas pelo BERT e permitir buscas eficientes de similaridade entre textos jurídicos, possibilitando uma melhor correlação entre petições e acórdãos passados

Criação da Base Vetorial

Após a geração dos embeddings, é necessário armazená-los em uma base vetorial. Para isso, é utilizada a plataforma Milvus, que é especializada no armazenamento e recuperação de vetores de alta dimensão. A criação dessa base permite que futuras buscas de similaridade sejam realizadas de forma eficiente, garantindo rapidez na comparação de novas perguntas com dados históricos.

Criação do Repositório GIT

Para garantir o versionamento e a colaboração eficiente, foi criado um projeto no Git para servir como repositório central do desenvolvimento. Esse repositório armazenará os códigos, scripts de testes e documentação do sistema.

Revisão de Parâmetros

Após os testes iniciais, os parâmetros do servidor BERT foram revisados e ajustados conforme necessário. Isso incluiu otimizações para melhor utilização da memória, ajuste de parâmetros e configuração de serviços auxiliares para suportar a demanda de análise.

Ambiente Python para STS

Foi montado um ambiente Python no servidor para suportar o uso de Semantic Textual Similarity (STS), permitindo que o sistema avalie a similaridade semântica entre os textos processados. Essa configuração garantiu a compatibilidade com as bibliotecas necessárias para a análise dos acórdãos.

Carregamento da Solução QA

O primeiro passo consiste na leitura e processamento do arquivo CSV que contém pares de perguntas e respostas. Esse arquivo serve como base para testar a capacidade do modelo BERT de compreender e gerar representações vetoriais dos textos. O carregamento dos dados envolve a verificação da estrutura do arquivo, remoção de inconsistências e conversão das perguntas em um formato adequado para análise.

Testes da Solução QA

Com os datasets devidamente carregados, a solução foi estada para verificar a capacidade de resposta às perguntas jurídicas. O desempenho do modelo foi avaliado por meio de métricas como precisão, relevância e tempo de resposta, garantindo que ele seja eficaz na análise de documentos jurídicos.

Carga do Dataset de 100 registros(teste)

Para validar o funcionamento inicial do sistema de Question Answering (QA), foi utilizado um conjunto de 100 registros. Esse dataset permite realizar testes preliminares para avaliar a precisão das respostas geradas e identificar possíveis ajustes necessários na parametrização do modelo.

Carga do Dataset de 1000 registros

Após os testes iniciais, o próximo passo foi utilizar um dataset mais robusto, contendo 1000 registros, conforme sugerido na URL específica. Esse aumento no volume de dados visa testar a escalabilidade da solução e sua capacidade de processar informações com diferentes níveis de complexidade jurídica.

➤ PASSOS

Encode do Dataset ITD/STJ no Servidor BERT

Nesta etapa, foi realizado o teste de encoding do dataset ITD/STJ diretamente no servidor BERT. Isso permitirá verificar se o modelo consegue processar os textos desse conjunto de dados corretamente, garantindo que as embeddings sejam geradas de maneira precisa e relevante.

Com o ambiente STS configurado, foi executado um script para testar o encoding do dataset ITD/STJ usando essa técnica. Esse teste visa verificar como o modelo interpreta e compara os textos jurídicos, refinando a abordagem de análise de similaridade.

Encode dos Relatórios da Base de Acórdãos do TJSE

Após a limpeza dos dados, será realizado o encoding do texto dos relatórios dos acórdãos do TJSE. Essa etapa transformará os documentos jurídicos em vetores semânticos que poderão ser comparados utilizando métricas de similaridade vetorial.

Inserção dos embeddings gerados na base Milvus

Com a base vetorial configurada, os embeddings das perguntas são inseridos no Milvus. Esse processo de armazenamento garante que os dados possam ser utilizados em consultas futuras, permitindo que o sistema de Question Answering (QA) encontre respostas relevantes com base na similaridade dos vetores.

Criação do índice para permitir buscas eficientes

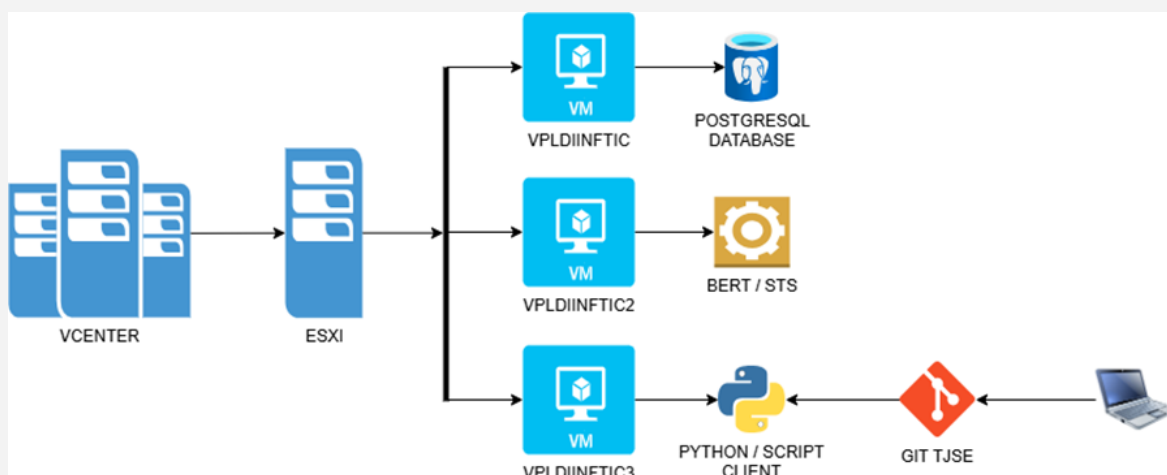
Após a leitura do arquivo, as perguntas são extraídas e organizadas para envio ao servidor BERT. Esse processo permite que os dados sejam padronizados antes de serem processados pelo modelo. Além disso, a segmentação das perguntas facilita a análise do desempenho do sistema em diferentes contextos e categorias de questionamentos.

Extração da lista de perguntas para envio ao servidor BERT

Para otimizar o tempo de consulta e melhorar a precisão das buscas no Milvus, é criado um índice que organiza os embeddings de maneira estruturada. O índice melhora a eficiência das buscas ao permitir que o sistema localize rapidamente os vetores mais semelhantes a uma consulta específica, acelerando o processo de recuperação de respostas.

Geração de embeddings das perguntas pelo modelo BERT

Com as perguntas prontas, elas são enviadas ao servidor BERT, que gera embeddings para cada uma delas. Embeddings são representações matemáticas do significado de cada texto, permitindo que comparações de similaridade sejam feitas de maneira eficiente. O BERT processa as perguntas levando em consideração seu contexto e estrutura linguística, resultando em vetores que encapsulam a semântica das frases.



Fluxo de Implantação

➤ PASSOS

Ao fim da configuração do ambiente, a infraestrutura pode ser logicamente visualizada conforme Figura 3.

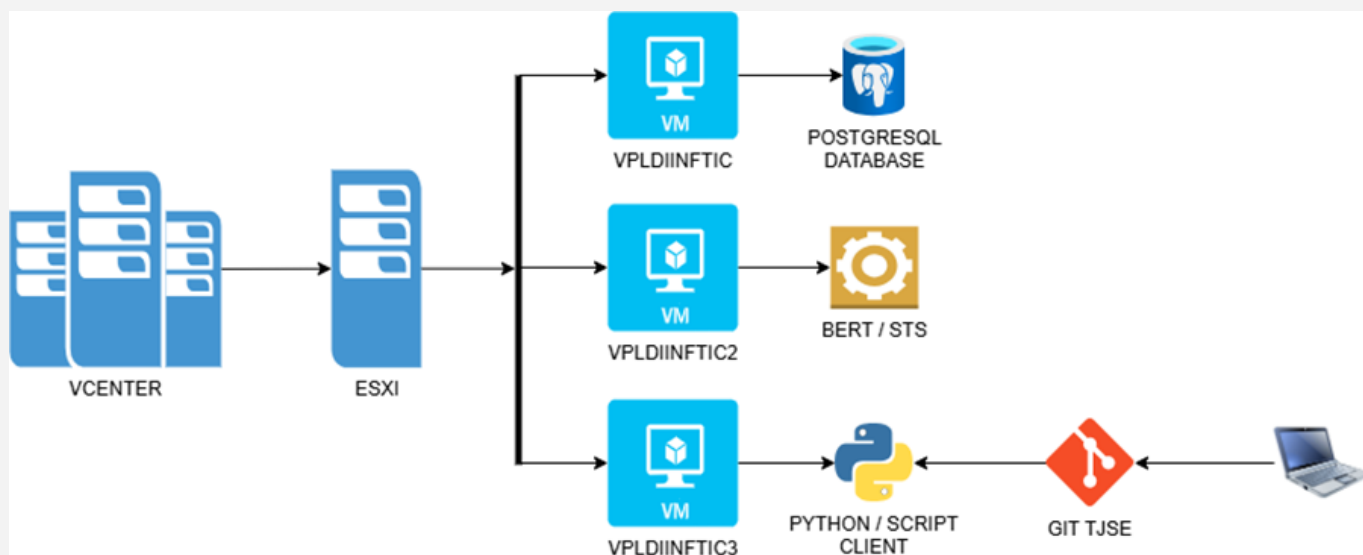


Figura 3: Infraestrutura da Solução

Além disso, o fluxo de informações entre os servidores de embeddings, o banco de dados vetorial e o banco de dados transacional pode ser verificado na Figura 4.

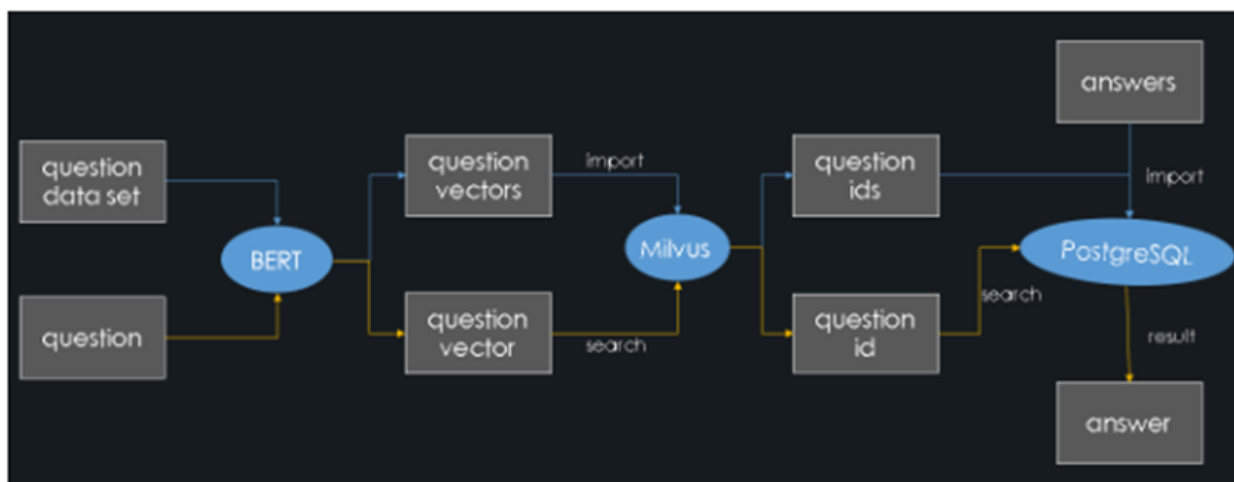


Figura 4: Fluxo da Ferramenta

BENEFÍCIOS E IMPACTOS ESPERADOS

Benefícios Diretos

Redução do retrabalho: menos tempo gasto na formulação de sentenças para casos repetitivos.

Melhor qualidade das decisões: fundamentação mais robusta e aderente aos precedentes.

Padronização do entendimento jurídico: diminuição de divergências em casos semelhantes.

Impactos Organizacionais

Cultura digital e inovação: servidores passam a enxergar a tecnologia como aliada, não substituta.

Eficiência institucional: maior número de casos resolvidos com menor esforço.

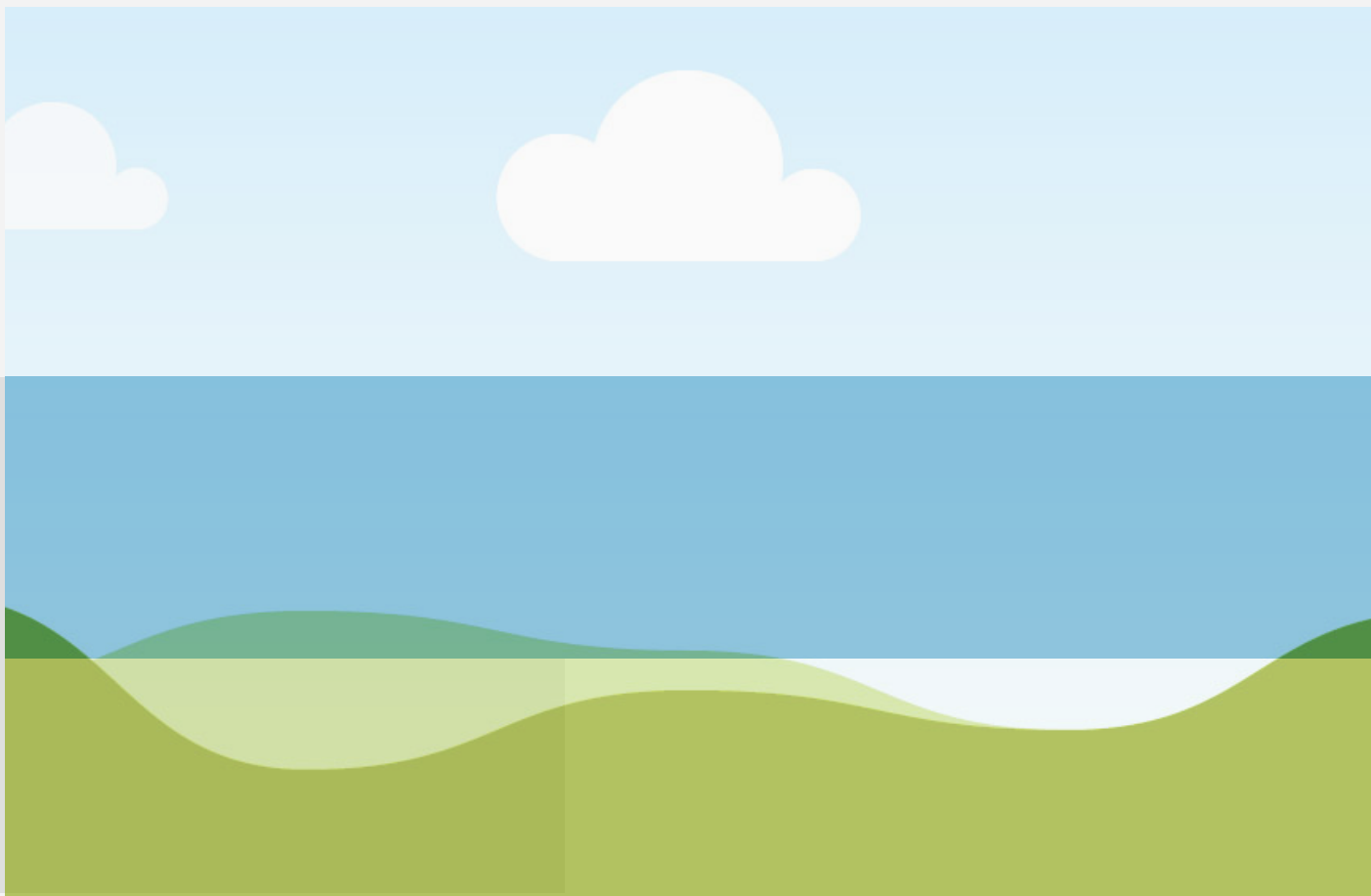
Base de dados centralizada e reutilizável: os dados processados poderão ser utilizados em outras ferramentas, como dashboards de jurisprudência.

Impactos Sociais

Decisões mais céleres e previsíveis: maior confiança por parte dos cidadãos.

Redução de recursos desnecessários: mais agilidade também nas Turmas Recursais.

Fomento à confiança institucional: o TJSE se posiciona como protagonista da inovação judiciária no Brasil.



DESAFIOS E LIÇÕES APRENDIDAS

Desafios Técnicos

Diversidade nos formatos de documentos: dificuldades em padronizar o texto extraído dos sistemas processuais.

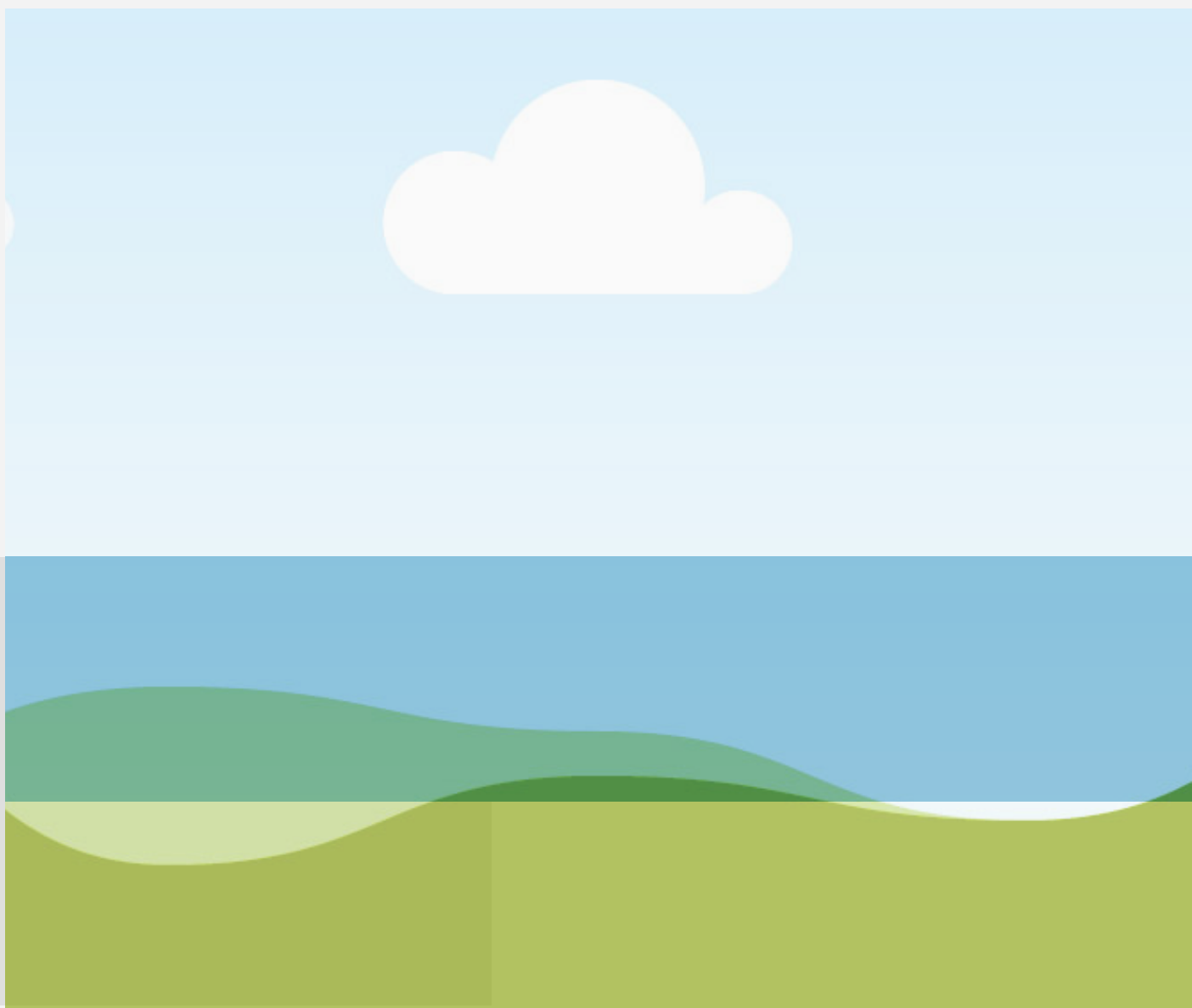
Ausência de GPU restringe a aplicação de modelos PLN.
Domínio jurídico específico: necessidade de ajustes finos nos modelos para expressões e estruturas jurídicas locais.

4.Dificuldade para processar bases de treinamento.

Impactos Organizacionais

A personalização do modelo para o contexto jurídico local é tão importante quanto a tecnologia em si.

A escuta ativa dos usuários melhora tanto o design da solução quanto sua aceitação.



RECOMENDAÇÕES E PRÓXIMOS PASSOS

Recomendações Imediatas

Testar novos modelos de PLN em infraestrutura com GPU(NVIDIA L40S) adquirida;

Implantar dashboards gerenciais para visualizar métricas de uso, feedbacks e impacto das recomendações.

Estabelecer um canal contínuo de comunicação com os usuários para registrar sugestões e erros.

Oportunidades de Expansão

Incorporação de jurisprudência do STJ e TJ de outros estados, permitindo uma visão mais ampla de precedentes.

Criação de um sistema de extração de teses jurídicas, com anotação automática dos fundamentos legais.

Previsão de desfecho com base em histórico processual, auxiliando na conciliação e instrução.

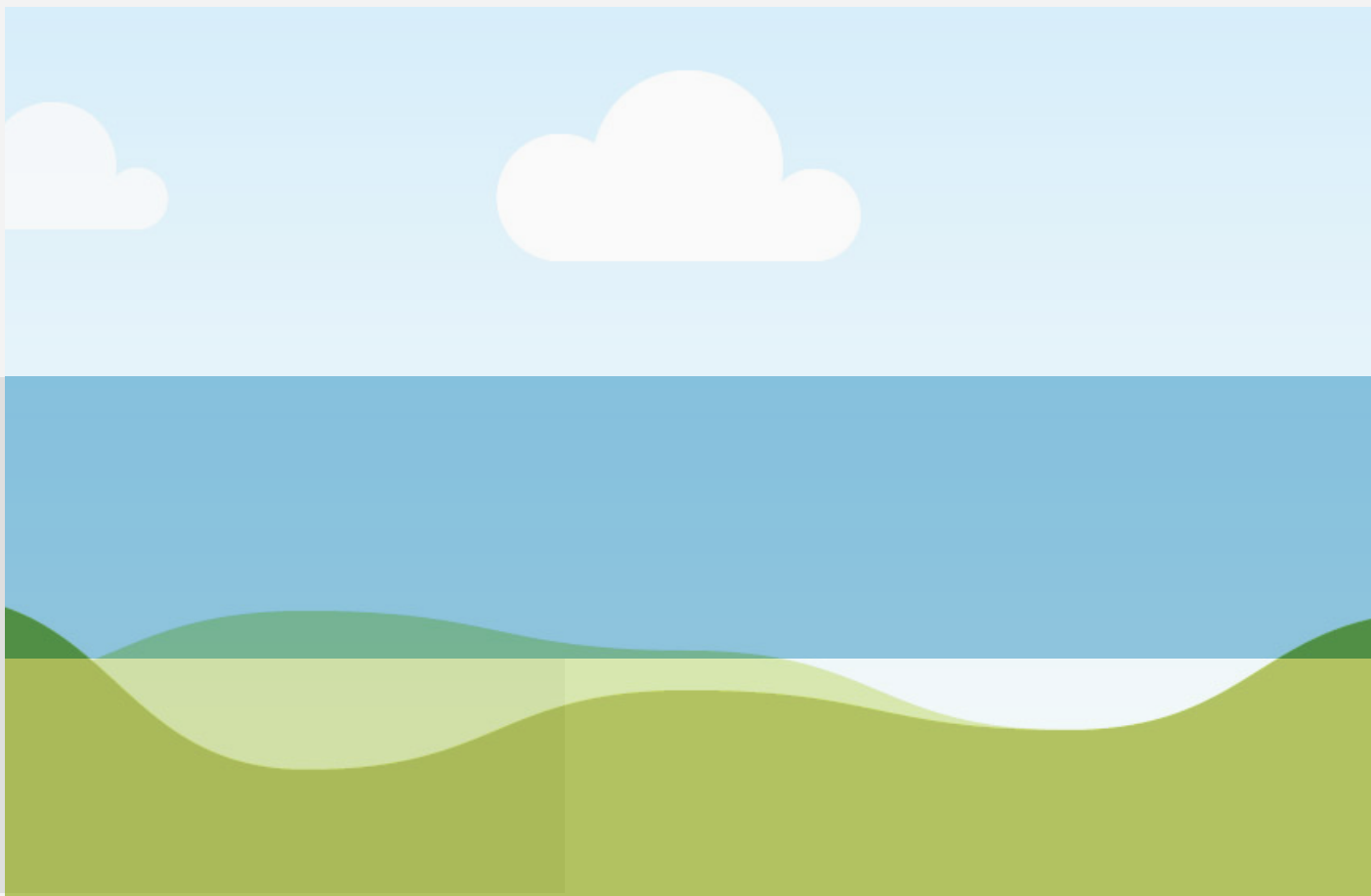
Disponibilização da solução para colaboração de outros Tribunais.

Propostas de Pesquisas Futuras

Aprimoramento com LLMs para explicações em linguagem natural.

Aprendizado por reforço com feedback do usuário: algoritmos que se adaptam com base na aceitação ou rejeição dos resultados.

Estudos quantitativos e qualitativos sobre o impacto do sistema na produtividade e qualidade das decisões.



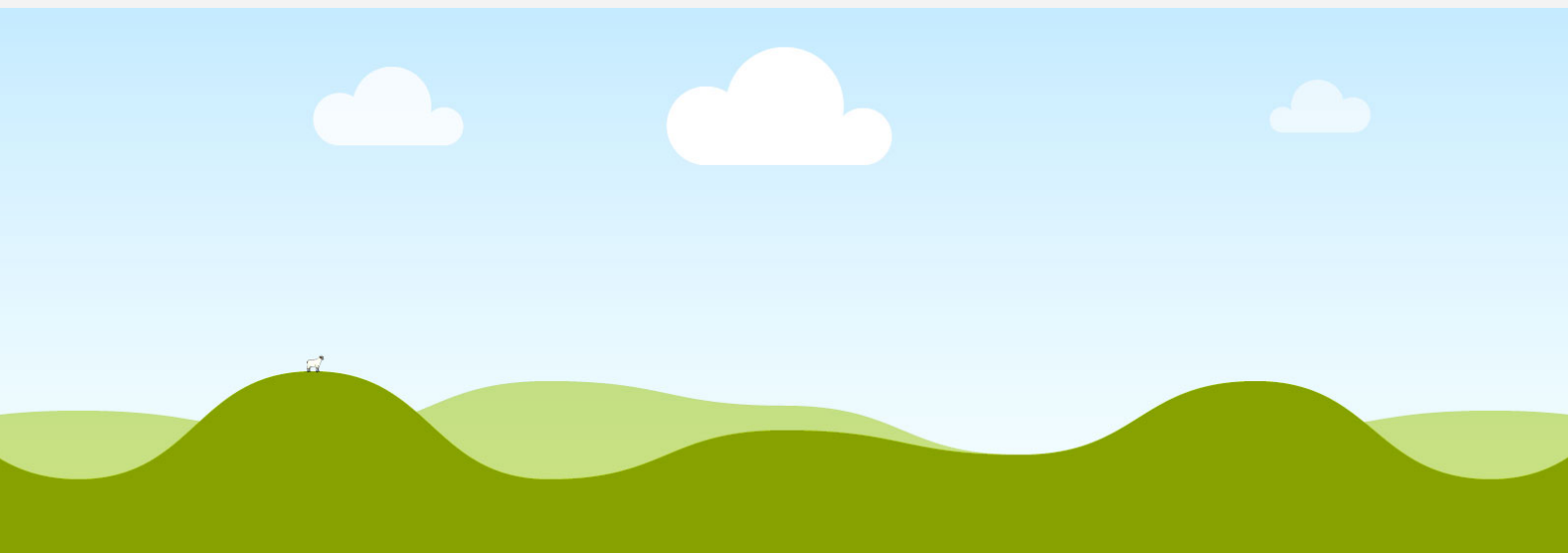
CONSIDERAÇÕES FINAIS

A ferramenta de análise de similaridade entre petições iniciais nos juizados especiais cíveis e acórdãos das turmas recursais representa um avanço significativo para a gestão processual e eficiência do sistema judiciário. Embora os resultados práticos ainda estejam pendentes da integração com o sistema judicial legado — etapa prevista nas próximas fases do projeto — o desenvolvimento e a arquitetura da solução já demonstram potencial considerável para otimizar o fluxo de trabalho nos juizados especiais, proporcionando maior previsibilidade decisória e agilidade na tramitação processual.

Durante o desenvolvimento, enfrentamos desafios técnicos relevantes, destacando-se a ausência de infraestrutura com GPUs adequadas, o que limitou consideravelmente o conjunto de modelos que pudemos testar e implementar. Somou-se a isso a escassez de modelos pré-treinados especificamente para o português jurídico brasileiro, cenário que exigiu adaptações metodológicas e abordagens alternativas para alcançar resultados satisfatórios mesmo com as limitações impostas. Estas dificuldades evidenciam a necessidade de maiores investimentos em recursos computacionais e no desenvolvimento de modelos linguísticos especializados para o domínio jurídico nacional.

Como perspectivas futuras, vislumbra-se a expansão do escopo da ferramenta para além dos juizados especiais cíveis, contemplando também varas cíveis comuns e, potencialmente, outras especialidades jurisdicionais. Adicionalmente, pretende-se compartilhar a solução com outros tribunais, fomentando a padronização de práticas e a interoperabilidade entre diferentes órgãos do Poder Judiciário. Este compartilhamento poderá não apenas maximizar o aproveitamento do recurso desenvolvido, como também contribuir para a formação de uma base de conhecimento jurídico computacionalmente tratável em escala nacional, representando um passo importante rumo à modernização e maior eficiência do sistema de justiça brasileiro. A solução, além de funcional, é escalável e replicável, podendo servir de modelo para outros tribunais.

Mais que um ganho operacional, o sistema promove um salto qualitativo, alinhando o TJSE com as melhores práticas de inovação no serviço público. A combinação de conhecimento técnico, jurídico e humano foi determinante para os resultados alcançados — e deve continuar sendo a base para os próximos passos.



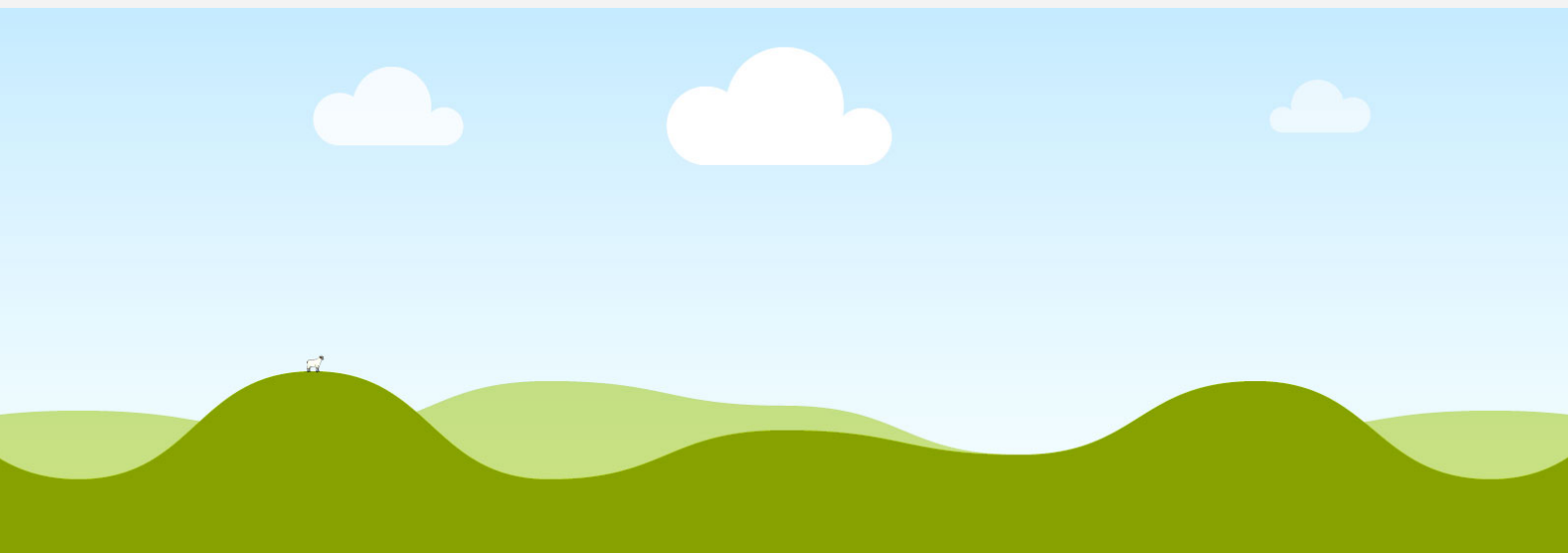
CONSIDERAÇÕES FINAIS

A ferramenta de análise de similaridade entre petições iniciais nos juizados especiais cíveis e acórdãos das turmas recursais representa um avanço significativo para a gestão processual e eficiência do sistema judiciário. Embora os resultados práticos ainda estejam pendentes da integração com o sistema judicial legado — etapa prevista nas próximas fases do projeto — o desenvolvimento e a arquitetura da solução já demonstram potencial considerável para otimizar o fluxo de trabalho nos juizados especiais, proporcionando maior previsibilidade decisória e agilidade na tramitação processual.

Durante o desenvolvimento, enfrentamos desafios técnicos relevantes, destacando-se a ausência de infraestrutura com GPUs adequadas, o que limitou consideravelmente o conjunto de modelos que pudemos testar e implementar. Somou-se a isso a escassez de modelos pré-treinados especificamente para o português jurídico brasileiro, cenário que exigiu adaptações metodológicas e abordagens alternativas para alcançar resultados satisfatórios mesmo com as limitações impostas. Estas dificuldades evidenciam a necessidade de maiores investimentos em recursos computacionais e no desenvolvimento de modelos linguísticos especializados para o domínio jurídico nacional.

Como perspectivas futuras, vislumbra-se a expansão do escopo da ferramenta para além dos juizados especiais cíveis, contemplando também varas cíveis comuns e, potencialmente, outras especialidades jurisdicionais. Adicionalmente, pretende-se compartilhar a solução com outros tribunais, fomentando a padronização de práticas e a interoperabilidade entre diferentes órgãos do Poder Judiciário. Este compartilhamento poderá não apenas maximizar o aproveitamento do recurso desenvolvido, como também contribuir para a formação de uma base de conhecimento jurídico computacionalmente tratável em escala nacional, representando um passo importante rumo à modernização e maior eficiência do sistema de justiça brasileiro. A solução, além de funcional, é escalável e replicável, podendo servir de modelo para outros tribunais.

Mais que um ganho operacional, o sistema promove um salto qualitativo, alinhando o TJSE com as melhores práticas de inovação no serviço público. A combinação de conhecimento técnico, jurídico e humano foi determinante para os resultados alcançados — e deve continuar sendo a base para os próximos passos.



RESPONSÁVEIS PELA PROPOSTA DE INTERVENÇÃO E DATA

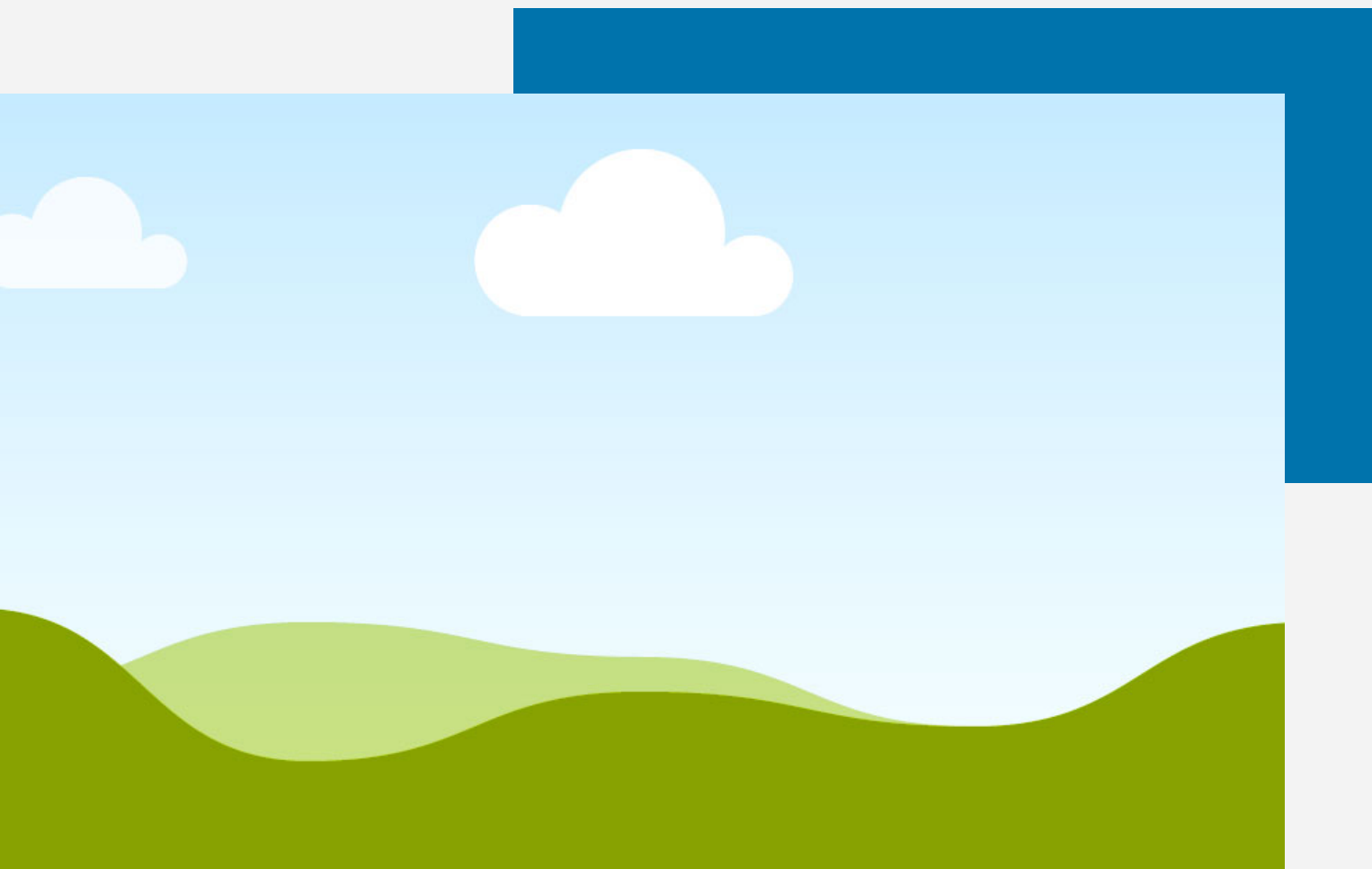
Max Ricardo Borges Riberro

Coordenador Interno do Projeto
Mestrando em Administração Pública
maxricardobr@gmail.com
Orcid: <https://orcid.org/0000-0003-4877-3937>
Lattes: <http://lattes.cnpq.br/2434224229287500>

Aracaju, 20/07/2025

Dr. Glauco de Figueiredo Carneiro

Coordenador Externo do Projeto
Orientador
glauco.carneiro@dcomp.ufs.br
Orcid: <https://orcid.org/0000-0001-6241-1612>
Lattes: <http://lattes.cnpq.br/4951846457502161>



REFERÊNCIAS

DEVLIN, J.; CHANG, M.-W.; LEE, K.; TOUTANOVA, K. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. 2019. Disponível em: <<https://arxiv.org/abs/1810.04805>>. Acesso em: 13 maio 2025.

JURAFSKY, D.; MARTIN, J. H. Speech and Language Processing. Prentice Hall, 2021.

MIKOLOV, T.; SUTSKEVER, I.; CHEN, K.; CORRADO, G.; DEAN, J. Distributed Representations of Words and Phrases and their Compositionality. NeurIPS, 2013. Disponível em: <<https://arxiv.org/abs/1310.4546>>. Acesso em: 13 maio 2025.

Discente: Max Ricardo Borges Ribeiro,

Orientador: Dr. Glauco de Figueiredo
Carneiro

Universidade Federal de Sergipe

20 de julho de 2025