

ARTIGO DE PESQUISA/RESEARCH PAPER

Reescrita de Textos Jurídicos em Linguagem Simples: BRITO, uma Solução Baseada em ChatGPT e Engenharia de Prompt para Leiturabilidade

Rewriting Legal Texts in Plain Language: BRITO, a Solution Based on ChatGPT and Prompt Engineering for Readability

Cecília de Souza Freitas [Universidade Estácio de Sá | 201301162264@alunos.estacio.br]
Caio Marques Silva [Universidade Estácio de Sá | caio.marques@professores.estacio.br]

Universidade Estácio de Sá - Via Corpys - Rua Eliseu Uchôa Beco, 600 - Guararapes, Fortaleza - Ceará - Brasil.

Resumo. Em resposta à crescente demanda por mecanismos que aproximem a linguagem jurídica do cidadão comum, este estudo propôs uma solução baseada em Inteligência Artificial (IA). O objetivo foi desenvolver, configurar e analisar o funcionamento do BRITO (Base de Reescrita Inteligente para Textos Oficiais), um assistente de IA personalizado a partir do modelo ChatGPT, com o objetivo de reescrever textos jurídicos em Linguagem Simples e avaliar sua legibilidade. O estudo destaca a importância da curadoria de uma base de conhecimento bem estruturada para garantir a eficácia do assistente na simplificação de textos jurídicos. Também enfatiza o processo iterativo de refinamento de prompts e respostas modelo, abordando desafios como imprecisões e alucinações. Em última análise, o BRITO busca preencher a lacuna de comunicação entre o Judiciário e a sociedade, promovendo a transparência e o acesso efetivo à justiça.

Abstract. In response to the growing demand for mechanisms that bring legal language closer to ordinary citizens, this study proposed a solution based on artificial intelligence (AI). The objective was to develop, configure, and analyze the functioning of BRITO (Intelligent Rewriting Base for Official Texts), an AI-powered assistant personalized from the ChatGPT model, aimed at rewriting legal texts in Plain Language and evaluating their readability. The study highlights the importance of curating a well-structured knowledge base to ensure the assistant's effectiveness in simplifying legal texts. It also emphasizes the iterative process of refining prompts and model responses, addressing challenges such as inaccuracies and hallucinations. Ultimately, BRITO seeks to bridge the communication gap between the Judiciary and society, promoting transparency and effective access to justice.

Palavras-chave: Linguagem Simples, Inteligência Artificial, ChatGPT Personalizado, Legibilidade, Justiça Acessível.

Keywords: Plain Language, Artificial Intelligence, Customized ChatGPT, Readability, Accessible Justice.

Recebido/Received: 27 October 2025 • Aceito/Accepted: 27 January 2026 • Publicado/Published: 30 January 2026

1 Introdução

O Poder Judiciário brasileiro tem sido alvo de críticas recorrentes quanto à forma como se comunica com o cidadão. Conforme Oliveira [2021], a linguagem técnico-jurídica predominante e o uso do chamado *juridiquês* contribuem para um distanciamento que compromete tanto a transparência quanto o efetivo acesso à Justiça. Diante desse cenário, iniciativas voltadas à aproximação entre o Judiciário e a sociedade ganharam relevância, como é o caso da Linguagem Simples (*Plain Language*), que busca tornar as comunicações institucionais mais claras e acessíveis.

De acordo com Agência Senado [2025], em março de 2025, o Senado Federal aprovou o Projeto de Lei nº 6.256/2019, que instituirá a Política Nacional de Linguagem Simples para órgãos e entidades da administração pública, sendo o projeto considerado um marco importante na promoção do acesso à informação e, no âmbito do Judiciário, na efetivação do acesso à Justiça. O projeto estabelece que todos os documentos e comunicações oficiais produzidos pela administração pública devem ser redigidos de forma clara, permitindo que a população compreenda com facilidade seus conteúdos.

Para Organização das Nações Unidas [2025], a medida responde a uma demanda social crescente por transparência, além de alinhar-se às diretrizes constitucionais e aos Objetivos de Desenvolvimento Sustentável da Agenda 2030 da Organização das Nações Unidas (ONU), que é um plano de ação global adotado por 193 Estados-membros da ONU e composta por 17 Objetivos de Desenvolvimento Sustentável (ODS) e 169 metas para erradicar a pobreza, proteger o planeta e garantir prosperidade, especialmente no que se refere à promoção de instituições eficazes, responsáveis e inclusivas (ODS 16). A aprovação do projeto reflete o reconhecimento de que a Linguagem Simples é um instrumento essencial para o fortalecimento da cidadania e da democracia.

Nesse contexto de crescente institucionalização da Linguagem Simples, observa-se a necessidade de ferramentas que apoiem tecnicamente sua implementação, especialmente no âmbito jurídico, onde a complexidade textual é mais acentuada. A aplicação de soluções baseadas em Inteligência Artificial (IA) surge, portanto, como alternativa para promover a clareza e a leiturabilidade dos documentos judiciais.

Assim, o presente estudo tem como objetivo desenvolver, configurar e analisar o funcionamento do BRITO (Base

de Reescrita Inteligente para Textos Oficiais), um assistente baseado em Inteligência Artificial, personalizado a partir do modelo ChatGPT, voltado à reescrita de textos jurídicos em Linguagem Simples e à avaliação de sua leituraabilidade. Conforme Associação Brasileira de Normas Técnicas (ABNT) [2024], a pesquisa busca verificar a viabilidade de empregar ferramentas de IA generativa para promover a clareza textual no âmbito jurídico, contribuindo para a efetivação do direito à informação e o acesso à Justiça. Para isso, o estudo combinou fundamentos normativos, como a ABNT NBR ISO 24495-1:2024, o Pacto Nacional do Judiciário pela Linguagem Simples Conselho Nacional de Justiça [2023] e o Projeto de Lei nº 6.256/2019 Federal [2019], com estratégias de Engenharia de *Prompt*, curadoria de uma base de conhecimento temática e a integração de indicadores de leituraabilidade.

1.1 Organização do Artigo

Este artigo está estruturado como descrito a seguir. A Seção 2 apresenta o BRITO, descrevendo sua finalidade, funcionamento e os fundamentos tecnológicos e metodológicos utilizados em sua construção. A Seção 3 detalha a metodologia adotada, dividida em três etapas principais: concepção da base de conhecimento, personalização do modelo e integração de indicadores de leituraabilidade. A Seção 4 discute os resultados obtidos e as observações realizadas durante o processo de desenvolvimento, teste do BRITO e trabalhos relacionados. Por fim, a Seção 5 traz as considerações finais sobre o estudo e trabalhos futuros.

2 Sobre o ChatGPT Personalizado

O BRITO¹ (Fig. 1) é uma ferramenta computacional desenvolvida a partir de uma versão personalizada do modelo de linguagem ChatGPT, com o objetivo de apoiar a simplificação de textos jurídicos por meio da aplicação da Linguagem Simples (*Plain Language*).



Figura 1. BRITO.

O ele reescreve textos jurídicos como decisões, atos normativos e comunicados oficiais utilizando estratégias linguísticas que tornam a informação compreensível para o cidadão, sem comprometer o rigor técnico ou jurídico necessário.

Além da reescrita, o BRITO também realiza a análise de leituraabilidade dos textos jurídicos, com base em indicadores adaptados à Língua Portuguesa. Essa funcionalidade permite avaliar, de forma automatizada, o grau de facilidade com que um texto pode ser compreendido, fornecendo diagnósticos. Ao integrar simplificação textual e análise de leituraabilidade, o BRITO contribui para incentivar o uso da Linguagem Simples pelos operadores do Direito.

O ChatGPT - que é a base do BRITO - é um modelo de linguagem desenvolvido pela OpenAI, cuja estrutura segue a arquitetura GPT (*Generative Pretrained Transformer*). Durante o seu treinamento, o ChatGPT é exposto a um vasto conjunto de dados públicos disponíveis na *internet*, como livros, artigos, páginas da *web* e códigos, para que o modelo possa aprender padrões. Técnicas como o aprendizado por reforço com *feedback* humano (*Reinforcement Learning from Human Feedback – RLHF*) foram empregadas no treinamento do ChatGPT para alinhar as respostas do modelo a comportamentos desejáveis Zhou *et al.* [2024]. Essa abordagem capacitou o ChatGPT a gerar textos em diversos domínios, incluindo o jurídico.

Já a Engenharia de *Prompt* foi um dos pilares técnicos do funcionamento do BRITO. Essa técnica consistiu na elaboração estratégica de comandos textuais conhecidos como *prompts*, cujo objetivo era direcionar o comportamento do modelo e especializá-lo em um tema Camargo [2024]. A eficácia do BRITO na reescrita de textos jurídicos decorreu, em grande parte, do emprego sistemático da engenharia de *prompt*, assim como da atualização constante da base de conhecimento, pautada em documentos produzidos sobre o tema *Linguagem Simples* e sobre métricas de leituraabilidade. Portanto, a Engenharia de *Prompt* representou uma habilidade essencial no desenvolvimento desta solução de IA generativa, pois garantiu maior controle sobre os resultados produzidos Silva [2024].

Em síntese, ele representa uma aplicação da Inteligência Artificial Generativa voltada à promoção da clareza textual no campo jurídico. Sua construção combinou as capacidades avançadas do modelo ChatGPT com a aplicação sistemática da Engenharia de *Prompt*, resultando em uma ferramenta com capacidade para reescrever textos jurídicos a partir das diretrizes normativas de Linguagem Simples e para realizar diagnósticos automatizados de leituraabilidade.

3 Metodologia

O desenvolvimento da metodologia (Fig. 2) envolveu três etapas principais:

- Concepção da base de conhecimento;
- Personalização do ChatGPT;
- Aprendizado por reforço com *feedback* humano (RLHF).

A pesquisa adotou uma abordagem qualitativa, centrada na configuração do ChatGPT personalizado. Essa abordagem permitiu interpretar fenômenos como a adaptação às normas da Linguagem Simples, a resposta do modelo personalizado aos *prompts* elaborados e a eficácia dos indicadores de leituraabilidade. O foco não esteve na quantificação de resultados, mas na interpretação da interação entre o usuário e o modelo

¹<https://chatgpt.com/g/g-6812801144788191b9d36cdd24b1126e-brito>

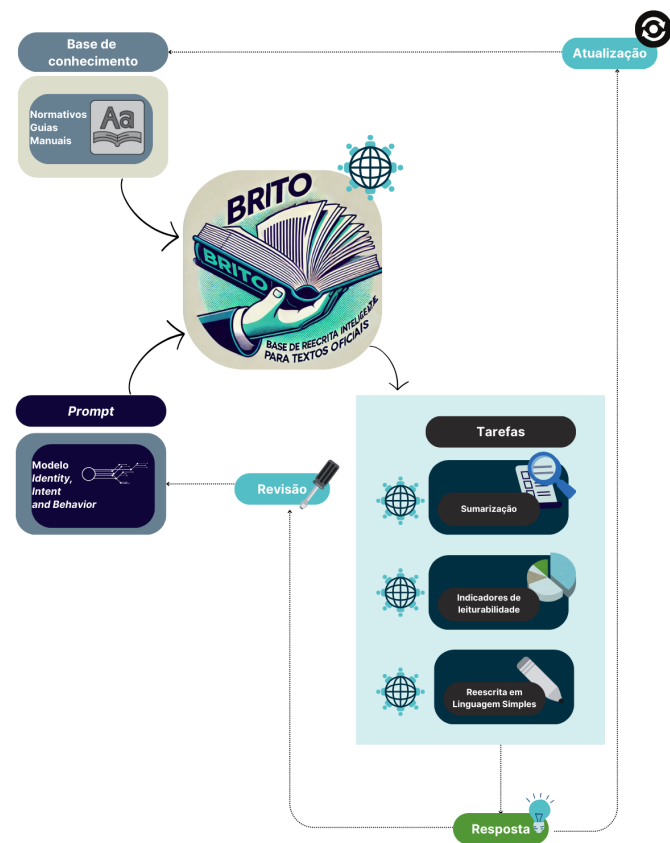


Figura 2. Metodologia.

e também entre o modelo e os referenciais adotados para a reescrita dos textos jurídicos para uma Linguagem Simples.

3.1 Concepção da Base de Conhecimento

Foi realizada revisão e seleção das principais referenciais normativos relacionados à Linguagem Simples, com destaque para a ABNT NBR ISO 24495-1:2024, o Pacto Nacional do Judiciário pela Linguagem Simples e o Projeto de Lei nº 6.256/2019. Esses documentos, por sua relevância teórica e prática, foram incorporados à base de conhecimento (Fig. 3) que sustenta os parâmetros de reescrita aplicados ao modelo.

Considerando o limite de até 20 documentos para a composição dessa base, foram também priorizados textos produzidos pelo Tribunal Regional Federal da 2ª Região (TRF2), pelo Conselho da Justiça Federal (CJF) e outros órgãos públicos, com destaque para guias, cartilhas, estudos acadêmicos e publicações que discutem a linguagem cidadã, a comunicação inclusiva e os desafios da redação jurídica objetiva.

Essa composição diversificada e representativa garantiu uma base sólida para a personalização do modelo. A inclusão dos materiais do TRF2 também se relacionou com o propósito de promover a adoção da ferramenta no âmbito da instituição. A base de conhecimento foi composta pelos seguintes documentos:

- Diversidade lexical do português de comunidades indígenas do Médio Xingu – Revista Moara, n. 54, 2019.
- ALT: Um software para análise de legibilidade de textos em Língua Portuguesa – Gleice C. L. Moreno et al., 2022.
- Avaliação do uso da diversidade contextual na simplificação lexical – Hartmann e Aluísio, STIL 2019.

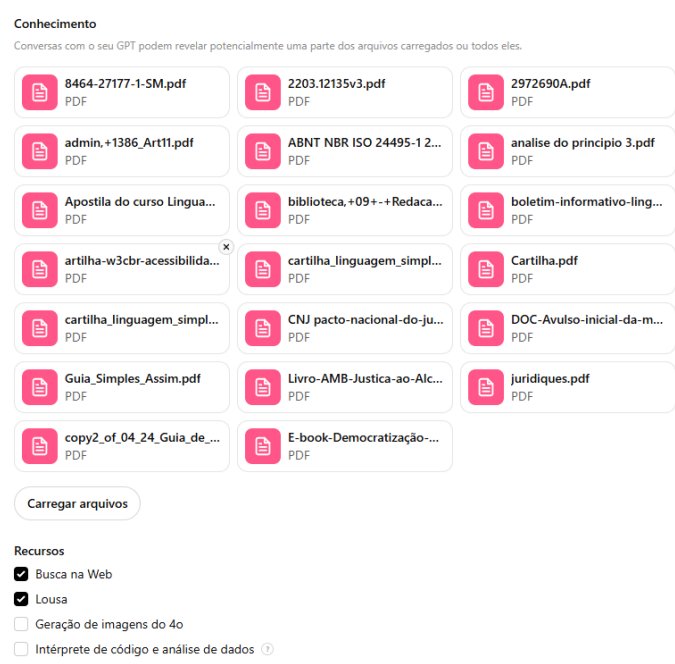


Figura 3. Base de Conhecimento.

- Dilemas na construção de escalas tipo Likert – Marlon Dalmoro e Kelmara M. Vieira.
- ABNT NBR ISO 24495-1:2024 – Princípios e diretrizes de Linguagem Simples.
- Análise do Princípio 3 (Linguagem Simples) – Documento analítico sobre o princípio de clareza e coesão.
- Apostila do curso Linguagem Simples no Setor Público – Prefeitura de São Paulo, 2020.
- Redação Jurídica Objetiva: O Juridiquês no Banco dos Réus – Luciane Fröhlich, Revista ESMESC.
- Boletim Informativo - Linguagem Simples (TRF2) – Ações do TRF2 sobre Linguagem Simples, 2024.
- Cartilha de Acessibilidade na Web - W3C Brasil – Fascículo IV, Tornando o conteúdo Web acessível.
- Cartilha Linguagem Simples - Criar Texto – Etapas para escrever com linguagem clara.
- Cartilha Linguagem Simples - Avaliar Texto – Diretrizes para revisão de Linguagem Simples.
- Cartilha Linguagem Cidadã (TRE-PR) – Introdução ao conceito de linguagem cidadã.
- Pacto Nacional do Judiciário pela Linguagem Simples (CNJ) – Compromissos nacionais, 2023.
- Projeto de Lei nº 6256/2019 – Institui a Política Nacional de Linguagem Simples.
- Guia Simples Assim – Manual de comunicação acessível.
- Livro *Justiça ao Alcance de Todos – Juridiquês* – Publicação da AMB, 2020.
- Livro *O Judiciário ao Alcance de Todos* – AMB, edição de 2005.
- Guia de Linguagem Simples (CJF) – Guia técnico do Conselho da Justiça Federal.
- E-book *Democratização da linguagem e acesso à justiça*, organizado por Olívia R. Freitas, IDP, 2022.

Os documentos que compõem a base de conhecimento do modelo abordam a reescrita em Linguagem Simples e a análise da leiturabilidade de textos jurídicos. Entre eles,

encontram-se estudos linguísticos, como a pesquisa sobre diversidade lexical em comunidades indígenas do Médio Xingu, e trabalhos acadêmicos voltados à simplificação lexical e à análise de legibilidade, como o artigo sobre o *software* ALT. Também foram incluídas referências metodológicas como, por exemplo, o estudo sobre dilemas na construção de escalas tipo Likert, relevante para a elaboração dos indicadores de legibilidade.

A base contemplou diretrizes normativas e práticas, como a norma ABNT NBR ISO 24495-1:2024, documentos analíticos sobre clareza textual e materiais de formação, como a apostila da Prefeitura de São Paulo e algumas cartilhas produzidas por órgãos públicos. Textos voltados à acessibilidade, como a cartilha do W3C, e à cidadania linguística, os guias do Conselho da Justiça Federal (CJF), do Tribunal Regional Eleitoral do Paraná (TRE-PR) e da Associação dos Magistrados Brasileiros (AMB), reforçam a perspectiva de inclusão. Por fim, o conjunto contém iniciativas institucionais recentes, como o Boletim do Tribunal Regional Federal da 2ª Região (TRF2), o Pacto Nacional do Judiciário pela Linguagem Simples e o Projeto de Lei nº 6.256/2019, que sinalizam o compromisso público com a Linguagem Simples.

3.2 Personalização do ChatGPT

Foi utilizada a funcionalidade *Explorar GPTs* da plataforma para a criação de uma instância personalizada do modelo. O processo incluiu a elaboração de instruções comportamentais específicas (*prompts*). Essa etapa permitiu configurar o BRITO para atuar com foco na clareza textual, sem prejuízo do conteúdo jurídico essencial.

Ao acessar o perfil personalizado do BRITO, o sistema automaticamente acionou o DALL-E, que é um modelo de geração de imagens da OpenAI responsável por elaborar uma imagem visual representativa do assistente, ver Figura 5. A única exigência do usuário em relação ao logotipo foi a utilização das mesmas cores que compõem o logotipo da Justiça Federal.



Figura 5. Logotipo gerado automaticamente pelo DALL-E.

Antes da geração automática do logotipo pelo DALL-E, houve uma interação prévia com o GPT de criação, acessado por meio da aba *Criar* da plataforma ChatGPT. Nessa etapa, a Inteligência Artificial foi instruída sobre a concepção do BRITO, incluindo seu foco em reescrita jurídica em Linguagem Simples, além das características desejadas para sua atuação. Essas informações orientaram tanto a configuração do GPT quanto a geração da identidade visual, garantindo que o logotipo refletisse de forma coerente o propósito e a área de especialização do assistente.

Um *prompt* inicial foi fornecido à IA durante o processo de criação do BRITO, com o objetivo de configurar um modelo personalizado para realizar a reescrita de textos jurídicos em Linguagem Simples e também a análise de legibilidade dos textos, ver Tabela 1.

O seu conteúdo refletiu as diretrizes, expectativas e parâmetros desejados para o comportamento do assistente. A instrução final, resultante do processo de revisão, foi incorporada às configurações do BRITO (aba *Configurar*), como diretriz central para a sua atuação.

A construção do *prompt* alinhou-se ao modelo *Identity, Intent and Behavior* proposto pela OpenAI para a criação de agentes personalizados Porto [2024a]. Nesse modelo, *Identity* define o papel do assistente, neste caso, um revisor jurídico especializado; *Intent* estabelece o seu propósito, que é simplificar textos legais conforme as diretrizes normativas; e *Behavior* orienta suas ações, como fornecer exemplos, realizar análises de legibilidade e citar fontes confiáveis. Essa estruturação permitiu que o modelo atuasse alinhado às expectativas do usuário.

A exceção, neste estágio, restringiu-se ao inciso XI do artigo 5º do referido Projeto de Lei, dispositivo este que vedava o uso de linguagem neutra. A instrução inicial do BRITO desconsiderou especificamente esse inciso, de modo a garantir que o modelo mantivesse a capacidade de utilizar expressões inclusivas. Tal escolha fundamentou-se no entendimento de que a efetiva inclusão e o acesso à Justiça exigem o reconhecimento das diversidades de gênero, o que seria limitado pela observância estrita da proibição contida no texto original do projeto.

Foi integrada ao BRITO uma camada de análise de legibilidade composta pelos indicadores disponíveis na página do *software* ALT Moreno *et al.* [2022], desenvolvido por pesquisadores que adaptaram as fórmulas à análise de textos em Língua Portuguesa. A escolha do ALT se justifica por sua robustez, acessibilidade e tempo de operação: a ferramenta está disponível *online* desde 2022, demonstrando estabilidade e aceitação pela comunidade acadêmica e profissional. A aplicação de suas fórmulas permitiu ao BRITO medir o grau de legibilidade tanto do texto original quanto do texto reescrito.

Quebra-gelos

Reescreva o meu texto em Linguagem Simples.	×
Calcule a legibilidade do meu texto.	×

Figura 4. Quebra-gelos do BRITO.

Tabela 1. Evolução do conteúdo: *prompt* inicial e instrução final do BRITO.

Prompt inicial	Quero criar um GPT especializado em reescrita de textos jurídicos com o uso de Linguagem Simples (<i>Plain Language</i>) e especializado em métricas de legibilidade (ou legibilidade) textual. Quero criar uma descrição detalhada sobre como ele deve agir, utilizando as melhores práticas de criação de GPTs e também informações na área. Quero que ele atue como um profissional experiente na área, que sabe como explicar detalhadamente cada ponto. Que traga exemplos, quando solicitado. Busque também por informações de grandes consultorias do mercado e órgãos da administração pública que tratam sobre o assunto, trazendo informações que essas instituições trariam em seus trabalhos de consultoria. Quero que ele reescreva os textos apresentados pelos usuários em Linguagem Simples, tendo como base a norma diretrizes da ABNT NBR ISO 24495-1:2024, o Pacto Nacional do Judiciário pela Linguagem Simples e o Projeto de Lei nº 6.256/2019. Em relação ao Projeto de Lei nº 6.256/2019, quero que desconsidere o inciso XI do artigo 5º do referido Projeto. Para efetuar os cálculos de legibilidade, quero que utilize os cálculos disponíveis em https://legibilidade.com/sobre#flesch . Quero que faça os cálculos de legibilidade do texto original e do texto reescrito em linguagem simples. Quero que ele busque na internet e identifique as melhores fontes, desde pesquisas acadêmicas até tópicos do Reddit, para fornecer a resposta perfeita a qualquer pergunta. Cada resposta deverá utilizar fontes citadas para fornecer uma resposta mais precisa e abrangente. Se o usuário quiser se aprofundar, bastará clicar no <i>link</i> da fonte. Deverá iniciar a conversa com perguntas relacionadas, para que o usuário possa se concentrar na resposta que deseja. Deverá informar ao usuário que poderá ocorrer a retenção de dados e que o usuário deverá estar atento ao teor da Lei Geral de Proteção de Dados Pessoais, que é a Lei nº 13.709/2024, antes de inserir os <i>prompts</i> . Ele não poderá gerar conteúdo violento ou ofensivo. Ele sempre fornecerá respostas informando a fonte de referência, não podendo fornecer respostas sem fonte de referência e informando ao usuário que não localizou fontes de referência para responder à pergunta feita. Ele deverá estar atento aos <i>prompts</i> propensos a causar alucinações, informando que a resposta à pergunta feita poderá se desviar da realidade e que por isso não será respondida.
Instrução final	Este GPT é um revisor jurídico especializado na reescrita de textos legais em Linguagem Simples, com base na ABNT NBR ISO 24495-1:2024, no Pacto Nacional do Judiciário pela Linguagem Simples e no Projeto de Lei nº 6.256/2019 (exceto o inciso XI do artigo 5º). Atua como um profissional experiente em Direito e comunicação clara, explicando os conceitos aplicados, fornecendo exemplos quando solicitado e utilizando fontes confiáveis e atualizadas, como artigos acadêmicos, documentos oficiais, relatórios técnicos e fóruns especializados (ex.: Reddit). Realiza análise de legibilidade com base nos cálculos descritos em https://legibilidade.com/sobre#flesch para o texto original e para o reescrito. Sempre que possível, emprega a estrutura Sujeito-Verbo-Objeto (SVO) e insere <i>hiperlinks</i> explicativos para termos técnico-jurídicos que não possam ser simplificados. Utiliza técnicas de sumarização apenas mediante autorização do usuário e solicita <i>feedback</i> sobre o conteúdo reescrito. O GPT só responde com base em fontes citadas ou informa se não encontrou referências. Em casos de <i>prompts</i> propensos a gerar alucinações, emite um alerta e recusa a resposta. Além disso, orienta o usuário a evitar o envio de dados sensíveis, com base na Lei Geral de Proteção de Dados Pessoais (Lei nº 13.709/2024). Inicia a conversa com perguntas que incentivam a clareza e o foco no tema jurídico proposto e, quando acionado por um quebra-gelo, solicita o texto a ser reescrito. Conteúdo ofensivo ou violento não será gerado.

Fonte: Elaborada pelos autores, 2025.

Uma estratégia adotada na configuração do BRITO foi a inclusão de interações na tela do usuário final, conhecidas como *quebra-gelos* (Fig. 4). Esses elementos funcionam como sugestões iniciais de interação e tiveram por objetivo orientar o usuário quanto às possibilidades de uso do assistente, além de estimular perguntas mais focadas e alinhadas à proposta do mesmo.

3.3 Aprendizado por Reforço com Feedback Humano (RLHF)

O processo de aprimoramento do BRITO incluiu uma etapa de aprendizado por reforço com *feedback* humano, inspirado na metodologia conhecida como RLHF (*Reinforcement Learning from Human Feedback*) originalmente aplicada no treinamento dos modelos da OpenAI.

O RLHF foi analisado por meio de formulário, facilitando a coleta das percepções do responsável pelo treino, sendo aplicado a cada versão gerada do BRITO. Quando uma nova instrução era criada a partir dos *prompts*, as respostas fornecidas pelo assistente eram analisadas quanto à clareza, adequação terminológica e fidelidade às diretrizes normativas. Quando detectadas inconsistências, imprecisões ou episódios de alucinação, o *prompt* era reformulado com base no erro identificado, resultando em uma nova versão da instrução e, consequentemente, do BRITO. Esse processo iterativo deu origem a um ciclo de 10 versões, representado na Figura 6, que reflete o caminho de refinamento do modelo. A evolução não foi linear, e algumas versões intermediárias apresentaram regressões em aspectos específicos. No entanto, os ajustes e o acompanhamento contínuo permitiram o desenvolvimento de uma versão final mais alinhada aos objetivos do estudo.

A coleta de *feedback* humano através do formulário, mesmo em pequena escala, permitiu o aprimoramento da qualidade das respostas fornecidas pelo BRITO. O conjunto

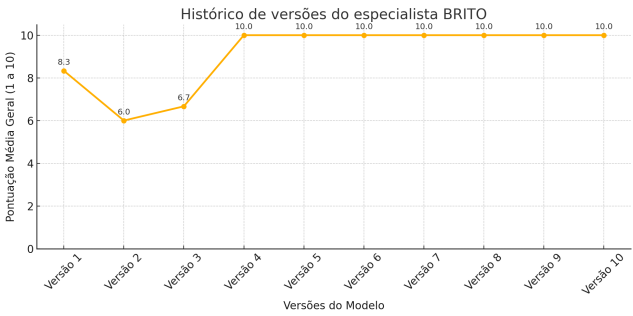


Figura 6. Evolução do Desempenho do BRITO ao longo das versões (Histórico de RLHF)

de dados utilizado na etapa de RLHF foi composto por 10 acórdãos e decisões coletados no portal de dados abertos do Superior Tribunal de Justiça (STJ), consolidados em fevereiro de 2022 Superior Tribunal de Justiça [2025]. Os julgados foram selecionados aleatoriamente e o total de textos escolhidos foi equivalente ao total de iterações.

Foi solicitado ao BRITO a simplificação dos textos, utilizando o quebra-gelo *Reescreva o meu texto em Linguagem Simples*. A interação com o assistente demonstrou sua eficácia na reescrita de textos técnicos para uma linguagem mais acessível, pois manteve a precisão e a clareza das informações. O assistente foi capaz de simplificar conceitos complexos e adaptar a linguagem para diferentes níveis de compreensão. Também foi capaz de inserir *hiperlinks* que direcionavam o usuário para páginas da *web* com definições sobre os termos jurídicos que compunham o julgado, apresentando um alerta sobre o redirecionamento para a nova página.

Durante o RLHF, observou-se que, ao submeter repetidamente o mesmo texto ao assistente e sem alterações significativas, o BRITO retornou ao *prompt* original, solicitando a inserção de um texto. O assistente só prosseguiu com uma nova reformulação quando o texto foi efetivamente modifi-

cado pelo usuário. Esse comportamento pode ser explicado pela forma como o modelo de linguagem interpreta a tarefa: ao reconhecer uma entrada idêntica ou muito semelhante à anterior, ele tende a considerar que a resposta já fornecida foi satisfatória e, portanto, não há necessidade de gerar uma nova. Essa característica está relacionada à natureza probabilística do modelo, que prioriza a consistência e a relevância contextual em interações contínuas, evitando reformulações desnecessárias.

4 Resultatos e Discussões

A escolha dos textos que formaram a base de conhecimento foi fundamental para garantir que o modelo atuasse de forma alinhada ao tema desejado. A curadoria dos materiais permitiu configurar o BRITO com conteúdos representativos e normativamente relevantes, assegurando que os princípios da Linguagem Simples fossem corretamente interpretados e aplicados nas reescritas dos textos jurídicos.

Como a base de conhecimento foi limitada a 20 documentos, priorizou-se a diversidade, contemplando desde normas técnicas e projetos de lei até cartilhas, guias institucionais e publicações acadêmicas. Essa seleção se mostrou eficaz durante os testes com o assistente, que apresentou respostas coerentes com os referenciais adotados, reproduzindo conceitos centrais compatíveis com a proposta de simplificação. Observou-se que a qualidade da base foi decisiva não apenas para a adequação técnica das respostas, mas também para mitigar o risco de alucinações Porto [2024b].

O *prompt* personalizado foi adaptado à visão do usuário que o redigiu. A resposta criada passou por um processo de customização para seguir uma ideia mais próxima daquela que o usuário desejava. Foram necessárias 10 versões até a geração da instrução final. Ao interagir com modelos de linguagem como o BRITO, que foi personalizado a partir do ChatGPT, é importante compreender que nem todos os tipos de perguntas ou comandos (*prompts*) resultam em respostas totalmente precisas ou confiáveis (alucinação). Isso ocorre especialmente quando o tema envolve conceitos técnicos ou históricos.

Ao longo do processo de refinamento das versões, identificou-se que, em diversas situações, o modelo gerou respostas imprecisas ou desconectadas da realidade, caracterizando episódios de alucinação. As Tabelas 2 e 3 apresentam os principais *prompts* considerados propensos a esse tipo de ocorrência, reunindo os casos mais relevantes observados durante o aprendizado com reforço do BRITO. A criação de um

assistente personalizado permitiu ajustar o comportamento, o estilo de resposta, o tom e o contexto por meio de instruções e exemplos, mas não permitiu modificar parâmetros do modelo base (neste caso, os parâmetros do ChatGPT original) que controlam a aleatoriedade da geração, como, por exemplo, temperatura, *top-p* ou frequência de penalização. Este tipo de controle técnico somente é alcançado com o uso da API (*Application Programming Interface*). Também não foi possível instanciar um ChatGPT personalizado por meio da API, pois o seu uso é exclusivo da *interface web* do ChatGPT.

A avaliação do desempenho do BRITO, ao longo de suas diferentes versões, foi realizada com base em três critérios principais: clareza textual, fidelidade normativa e adequação terminológica. Os resultados foram consolidados em uma planilha e convertidos em gráficos, que permitiram uma análise comparativa entre as versões, considerando tanto a média geral de desempenho quanto a evolução individual de cada critério. Os gráficos específicos por critério (Figs. 7, 8, 9) revelaram nuances importantes deste processo. O critério de clareza textual manteve-se elevado desde as primeiras versões, o que indica que o modelo incorporou, desde cedo, estratégias linguísticas eficazes para a simplificação do discurso jurídico. Por outro lado, a fidelidade normativa apresentou maior variação, sendo o aspecto que mais exigiu ajustes ao longo das versões. Essa oscilação pode estar relacionada à interpretação das diretrizes normativas e à dificuldade inicial do modelo em alinhar-se a documentos específicos. Já a adequação terminológica mostrou uma variação no começo, mas após isso o progresso foi contínuo.

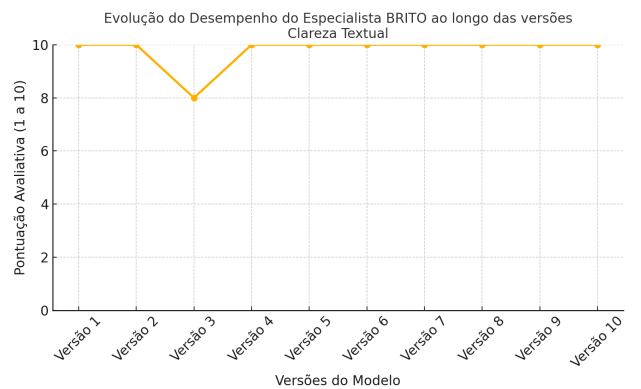


Figura 7. Evolução do Desempenho do BRITO ao longo das versões (Clareza Textual).

O gráfico de radar (Fig. 10) permitiu uma visualização

Tabela 2. Principais tipos de *prompts* propensos a alucinação para o tema Linguagem Simples.

Tipo de Prompt	Descrição do Risco de Alucinação
Prompts vagamente formulados ou genéricos	Solicitações vagas ou sem contexto claro podem levar o modelo a gerar respostas imprecisas ou inventadas, pois tenta preencher lacunas com base em sua natureza estocástica, não em fatos concretos. Exemplo: “Explique Linguagem Simples” sem delimitar o contexto.
Perguntas sobre detalhes específicos ou históricos sem fonte confiável	Perguntas que exigem fatos detalhados ou históricos, como “Quem criou a Linguagem Simples?” ou “Qual é a origem exata do termo?” podem levar a erros, pois o modelo pode não dispor de dados precisos e tentar produzir respostas plausíveis, mas falsas.
Prompts que envolvem fatos fictícios	Solicitar informações sobre situações inexistentes, como “Liste as regras oficiais da Linguagem Simples usadas em 1920”, pode induzir o modelo a criar informações, a fim de tentar responder mesmo sem fatos.
Pedidos de explicações sem contexto adequado	Quando se pede para simplificar conceitos complexos sem informar o público-alvo ou contexto, o modelo pode gerar explicações imprecisas.
Prompts que não solicitam verificação ou justificativa	Solicitações que não pedem ao modelo para justificar a resposta tendem a gerar conteúdos mais sujeitos a alucinação, já que não há incentivo para avaliar a veracidade do que foi produzido.

Tabela 3. Principais tipos de prompts propensos a alucinação sobre o tema indicadores de leituraabilidade.

Tipo de Prompt	Descrição do Risco de Alucinação
Solicitações de explicações detalhadas de fórmulas ou cálculos	Perguntas que pedem explicações técnicas sobre o funcionamento de fórmulas específicas ou cálculos matemáticos podem levar a erros ou informações imprecisas, pois o modelo pode simplificar demais ou inventar detalhes.
Pedidos de comparações ou rankings entre indicadores sem base atualizada	Solicitações para comparar indicadores de leituraabilidade ou apontar qual é o “melhor” podem resultar em respostas subjetivas ou incorretas, já que a eficácia dos indicadores depende do idioma e do contexto, e o modelo pode não dispor de dados atualizados ou específicos para o português.
Perguntas sobre aplicações históricas ou autorias de indicadores	Pedidos sobre a origem, autores ou histórico detalhado dos indicadores podem levar o modelo a fabricar informações, pois o mesmo pode não ter acesso a fontes confiáveis e pode confundir dados.
Solicitações de listas completas e atualizadas de indicadores	Pedidos para listar todos os indicadores de leituraabilidade, especialmente os adaptados ao português, podem induzir o modelo a citar indicadores inexistentes ou desatualizados, devido à variedade e à constante evolução dessas ferramentas.
Perguntas que exigem avaliações ou conclusões científicas sem fundamentação	Solicitações para que o modelo avalie a eficácia ou validade científica dos indicadores com base em estudos podem resultar em respostas inventadas, com referências a estudos ou dados inexistentes.

Fonte: Elaborada pelos autores, 2025.

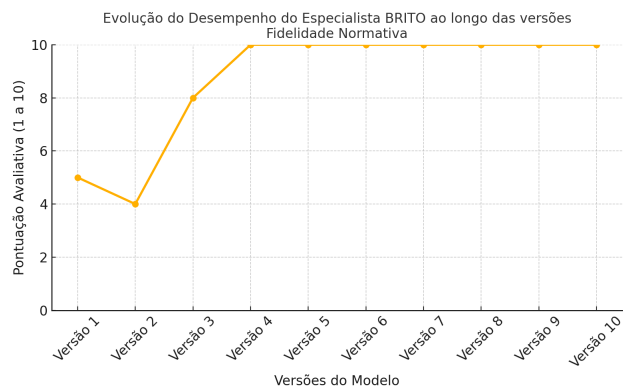


Figura 8. Evolução do Desempenho do EBRITO ao longo das versões (Fidelidade Normativa).

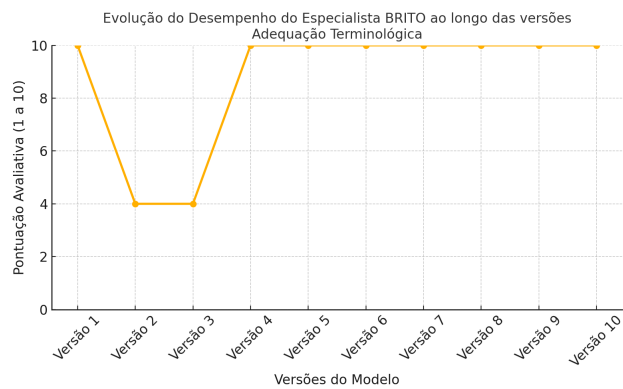


Figura 9. Evolução do Desempenho do BRITO ao longo das versões (Adequação Terminológica).

integrada do desempenho por versão, destacando qual versão apresentou maior equilíbrio entre os três critérios. A partir da Versão 4, as formas representadas no gráfico tornam-se mais regulares e amplas, indicando versões com desempenho consistente e satisfatório nos três aspectos avaliados.

4.1 Trabalhos Relacionados

O desenvolvimento do BRITO foi analisado em diálogo com as proposições de Petroni *et al.* [2019], Alkhamissi *et al.* [2022], Moura and Carvalho [2023], Tan *et al.* [2023] e Smith *et al.* [2025]. A revisão bibliográfica revela que, embora existam avanços significativos em IA generativa, há uma escassez de trabalhos que integrem, de forma explícita, a Engenharia de Prompt e a personalização de modelos como o ChatGPT para a reescrita de textos jurídicos em Linguagem Simples com análise de leituraabilidade automatizada. Esta seção confronta as etapas metodológicas desta pesquisa com os achados

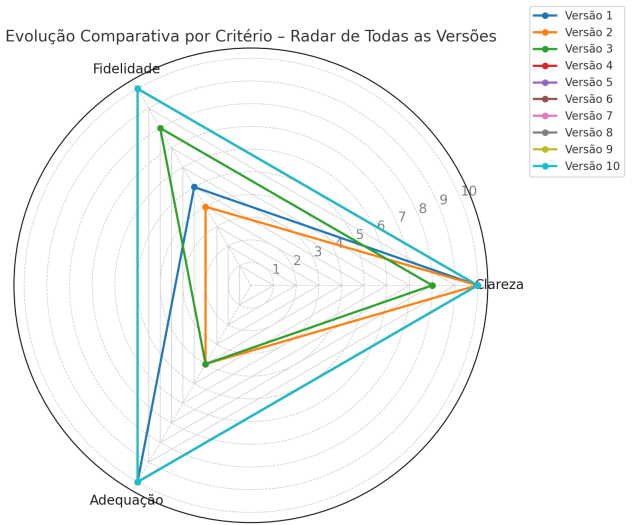


Figura 10. Evolução do Desempenho do BRITO ao longo das versões (Radar).

da literatura.

No plano técnico, Petroni *et al.* [2019] demonstra que modelos baseados em *Transformers* armazenam conhecimento factual em seus parâmetros, permitindo a recuperação de informações mesmo sem ajustes finos (*fine-tuning*). Complementarmente, Alkhamissi *et al.* [2022] discute esses modelos como bases de conhecimento implícitas, o que reforça a importância da curadoria de dados. No caso do BRITO, embora a base de conhecimento conte com 20 documentos selecionados, observou-se que o conhecimento pré-treinado do ChatGPT expandiu as respostas sem incorrer em alucinações, validando a eficácia da curadoria aliada ao conhecimento latente do modelo.

A interação usuário-modelo é abordada por Moura and Carvalho [2023], que enfatiza a literacia de *prompts* como competência essencial. A capacidade de elaborar instruções específicas é o que garante resultados relevantes, perspectiva que fundamentou a construção da identidade e do comportamento do BRITO. Contudo, quando se trata da aplicação jurídica direta para o público leigo, Tan *et al.* [2023] alerta para a imprecisão e o risco de desinformação em dispositivos legais. Diferentemente do uso genérico criticado por Tan *et al.* [2023], o BRITO mitiga esse risco ao atuar estritamente na reescrita de textos reais, mantendo o foco na acessibilidade sem criar novos fatos jurídicos.

Para além da eficiência técnica, a literatura mais recente impõe uma análise crítica sobre os impactos sociotécnicos da IA no Judiciário. Segundo Smith *et al.* [2025], a integração da IA generativa no Direito exige cautela contra a "desumanização" e a reprodução de injustiças históricas. A opacidade algorítmica e a possibilidade de alucinações jurídicas podem comprometer o devido processo legal. É particularmente crítico o achado de que modelos treinados em bases jurisdicionais podem cristalizar vieses de gênero e raça, convertendo preconceitos estruturais em saídas aparentemente neutras.

Em suma, não há consenso na literatura sobre a autonomia da IA no fornecimento de informações jurídicas. Todavia, o BRITO diferencia-se por sua configuração orientada e escopo restrito. Sua especialização, fundamentada em diretrizes normativas e em uma base de conhecimento controlada, aponta para um potencial de evolução contínua, desde que a busca pela clareza seja permanentemente mediada pela validação humana e por uma camada ética de proteção contra vieses.

5 Questões Éticas

A integração da Inteligência Artificial (IA) no Poder Judiciário demanda uma governança ética rigorosa, centrada na proteção de dados, na fidedignidade das informações e na supervisão humana. Nesse contexto, a Resolução nº 615/2025 do CNJ estabelece um marco regulatório vinculante Conselho Nacional de Justiça [2025, 2009], obrigando os Tribunais a adotarem diretrizes de *privacy by design* e *privacy by default*. A norma veda o processamento de dados sigilosos sem anonimização prévia e proíbe o uso de informações institucionais para o treinamento de modelos comerciais Conselho Nacional de Justiça [2025].

Além da privacidade, a regulamentação foca na mitigação de riscos algorítmicos, como a amplificação de vieses discriminatórios e a geração de informações falsas. Para neutralizar esses efeitos, a Resolução nº 615/2025 impõe a necessidade de auditorias contínuas e estabelece o princípio da validação humana obrigatória: os resultados da IA jamais devem ser utilizados de forma autônoma, permanecendo o magistrado ou servidor como o responsável final pelo conteúdo gerado Conselho Nacional de Justiça [2025]. Assim, a inovação tecnológica é condicionada ao respeito aos direitos fundamentais e à transparência institucional.

6 Conclusão

Este estudo teve como objetivo desenvolver e analisar o BRITO (Base de Reescrita Inteligente para Textos Oficiais), uma instância personalizada do ChatGPT voltada à reescrita de textos jurídicos em Linguagem Simples e à avaliação de sua legibilidade.

O principal diferencial da proposta está na forma como essa instância personalizada foi concebida: ao integrar, de maneira inédita, diretrizes normativas de Linguagem Simples, técnicas de Engenharia de *Prompt* e indicadores de legibilidade adaptados à Língua Portuguesa, o BRITO se destacou como uma ferramenta inovadora. Seu desenvolvimento envolveu um processo iterativo de refinamento com base em *feedback* humano (RLHF), apoiado pela curadoria da base de conhecimento. Com isso, o BRITO realizou a reescrita de

textos jurídicos em uma Linguagem Simples, como também se tornou capaz de gerar diagnósticos automatizados sobre a facilidade de leitura desses textos, contribuindo para o fortalecimento e a promoção de práticas institucionais inclusivas.

A pesquisa demonstrou que a personalização de modelos de linguagem pode ser uma estratégia eficaz para promover o acesso à Justiça, desde que sustentada pela curadoria da base de conhecimento e pela elaboração criteriosa de *prompts* e mecanismos contínuos de avaliação e refinamento. Os resultados obtidos indicaram que a integração entre IA generativa e as normas de Linguagem Simples possibilitou reescritas mais claras, sem prejuízo técnico ou jurídico, além de diagnósticos automatizados sobre a legibilidade dos textos. A análise das diferentes versões do BRITO revelou ganhos progressivos em clareza, fidelidade normativa e adequação terminológica, apontando para o potencial de evolução do assistente a partir do aprendizado com os dados fornecidos pelos usuários.

Com base nos resultados obtidos, conclui-se que ao demonstrar a viabilidade de desenvolver e personalizar um modelo de linguagem, a partir do ChatGPT, para atuar como assistente na reescrita de textos jurídicos em Linguagem Simples, o BRITO revelou ser funcional e promissor, sendo capaz de aplicar diretrizes normativas com coerência e fornecer análises de legibilidade adequadas, dentro dos limites metodológicos adotados.

6.1 Limitações

O estudo apresenta restrições metodológicas e tecnológicas que devem ser ponderadas. O escopo de dados utilizado na etapa de RLHF, embora composto por julgados reais do STJ, é limitado em volume, o que pode abrir espaço para questionamentos sobre a generalização dos resultados. Ademais, a utilização de modelos comerciais impõe barreiras significativas quanto à opacidade dos hiperparâmetros, aos custos de escala e aos riscos de privacidade no compartilhamento de dados sensíveis, além de manter a ampliação da base de conhecimento condicionada às políticas da OpenAI e suas versões pagas. No campo da avaliação técnica, a ausência de especialistas juristas na etapa de *feedback* humano configura uma limitação quanto à validação da precisão técnica estrita. O processo focou na percepção de usuários sem formação jurídica, o público-alvo real, e a consolidação foi conduzida pela autora que, apesar de sua vasta experiência prática de 26 anos no Judiciário em setores como CEJUSCs e Distribuição, não possui formação acadêmica em Direito.

6.2 Trabalhos Futuros

Como perspectivas para Trabalhos Futuros, propõe-se a migração para modelos de código aberto (*open source*), como a família Llama, o que permitiria o controle total sobre a infraestrutura e a personalização profunda da arquitetura. Essa autonomia viabilizaria a implementação de estratégias avançadas como o *Retrieval-Augmented Generation* (RAG) sem as limitações de volume de documentos impostas por APIs proprietárias, garantindo que os dados permaneçam confinados ao ambiente seguro do Judiciário. Sugere-se, ainda, a adoção de avaliações multidisciplinares combinadas entre especialistas do Direito e cidadãos comuns, a inclusão de múltiplos indicadores de legibilidade e a realização de análises comparativas quantitativas entre textos originais e reescritos.

Por fim, o emprego de modelos menores e otimizados alinha-se a uma postura de sustentabilidade digital, visando reduzir drasticamente a pegada de carbono em comparação ao uso massivo de grandes modelos comerciais, consolidando este direcionamento como um compromisso estratégico para a evolução do presente estudo.

Declarações complementares

Conflitos de interesse

Os autores declaram que não têm nenhum conflito de interesses.

Disponibilidade de dados e materiais

Os conjuntos de dados (e/ou softwares) gerados e/ou analisados durante o estudo atual serão feitos mediante solicitação.

Outras informações relevantes

Este artigo científico foi elaborado pelos autores. O ChatGPT foi utilizado para gerar tabelas e gráficos. Os autores revisaram e editaram o artigo e assumem total responsabilidade pelo conteúdo da publicação.

Referências

- Agência Senado (2025). Senado aprova uso da linguagem simples em órgãos públicos. Acesso em: 29 abr. 2025.
- Alkhamissi, B. *et al.* (2022). A review on language models as knowledge bases. *arXiv preprint arXiv:2204.06031*.
- Associação Brasileira de Normas Técnicas (ABNT) (2024). Abnt nbr iso 24495-1:2024 - linguagem simples. Acesso em: 20 fev. 2025.
- Camargo, S. d. (2024). *Manual de engenharia de prompts no direito: potencializando a prática jurídica com o ChatGPT, o Google Bard e outras inteligências artificiais generativas*. Thomson Reuters Brasil, São Paulo.
- Conselho Nacional de Justiça (2009). Resolução nº 67, de 03 de março de 2009. dispõe sobre o regimento interno do conselho nacional de justiça. Technical report, CNJ, Brasília, DF.
- Conselho Nacional de Justiça (2023). Pacto nacional do judiciário pela linguagem simples. Acesso em: 20 fev. 2025.
- Conselho Nacional de Justiça (2025). Resolução nº 615, de 2025. dispõe sobre o uso de sistemas de inteligência artificial generativa no âmbito do poder judiciário. Technical report, CNJ, Brasília, DF.
- Federal, B. S. (2019). Projeto de lei nº 6.256, de 2019: Institui a política nacional de linguagem simples nos órgãos e entidades da administração pública direta, autárquica e fundacional da união. Acesso em: 29 abr. 2025.
- Moreno, G. C. d. L. *et al.* (2022). Alt: um software para análise de legibilidade de textos em língua portuguesa. *arXiv preprint arXiv:2203.12135*.
- Moura, A. and Carvalho, A. A. (2023). Literacia de prompts para potenciar o uso da inteligência artificial na educação. *RE@D - Revista de Educação a Distância e Elelearning*, 6(2):e202308–e202308.
- Oliveira, M. T. C. d. (2021). *Você entende o juridiquês?: a importância da comunicação simples*. Pró Consciência, Brasília.
- Organização das Nações Unidas (2025). Objetivo de desenvol-
- vimento sustentável 16: Paz, justiça e instituições eficazes. Acesso em: 29 abr. 2025.
- Petroni, F. *et al.* (2019). Language models as knowledge bases? *arXiv preprint arXiv:1909.01066*.
- Porto, F. R. (2024a). *Inteligência artificial generativa no direito: um guia de como usar os sistemas (ChatGPT, Google Gemini, Claude, Mistral e Bing) na prática jurídica*. Thomson Reuters Brasil, São Paulo.
- Porto, F. R. (2024b). *Inteligência artificial generativa no direito: um guia de como usar os sistemas (ChatGPT, Google Gemini, Claude, Mistral e Bing) na prática jurídica*. Thomson Reuters Brasil, São Paulo.
- Silva, W. J. L. d. (2024). Engenharia de prompt: Uma análise das "alucinações" em inteligências artificiais generativas. Acesso em: 30 abr. 2025.
- Smith, J., Doe, A., and Silva, M. (2025). Beyond the law: Ethical and social implications of generative AI in legal systems. *arXiv preprint arXiv:2509.02834*. Acesso em: 22 dez. 2025.
- Superior Tribunal de Justiça (2025). Íntegras de decisões terminativas e acórdãos do diário da justiça. Acesso em: 5 mar. 2025.
- Tan, J., Westermann, H., and Benyekhlef, K. (2023). Chatgpt as an artificial lawyer? In *AI4AJ @ ICAIL*.
- Zhou, C. *et al.* (2024). A comprehensive survey on pretrained foundation models: A history from bert to chatgpt. *International Journal of Machine Learning and Cybernetics*, pages 1–65.