



EuDisseVocêDisse™

Sua verdade, autenticada™
Desenvolvido por EuDisseVocêDisse

Relatório de Comunicação Judicial

Todas as mensagens

Gerado por EuDisseVocêDisse.com

1. Este relatório foi gerado em 27 de agosto de 2026, para o processo EXAMPLE-2026-PT listado no(a) Tribunal de Justiça do Estado de São Paulo de São Paulo, SP, BRA, utilizando o sistema de inteligência artificial de Análise de Sentimento da EuDisseVocêDisse.com. O objetivo principal da IA generativa de resumo do relatório é tornar a análise numérica e os resultados analíticos já calculados mais legíveis para o tribunal. Ela é usada apenas em campos narrativos marcados e também pode resumir material analítico vinculado à fonte dentro do escopo. Ela não determina nem altera as classificações, pontuações, contagens, evidência fonte selecionada, gráficos ou montagem do PDF subjacentes. Cada instância de conteúdo do relatório gerado por IA generativa é marcada com uma adaga (†) ao final do parágrafo gerado. O Apêndice A - Metodologia explica esse marcador, identifica os sistemas utilizados e descreve as salvaguardas que preservam a revisão vinculada à fonte.
2. As seguintes interações digitais foram analisadas entre Mr Joe Blogs e Ms Jill Blogs no período de 09 de abril de 2024 a 09 de outubro de 2024:
 - a. 44 mensagens de e-mail.
 - b. 0 vídeos.
 - c. 0 conversas gravadas.
 - d. 0 mensagens de voz.
 - e. 0 mensagens instantâneas (ex.: WhatsApp, iMessage)
 - f. 0 mensagens enviadas por aplicativos de coparentalidade (ex.: Our Family Wizard, Talking Parents, CoParent Coordinator, AppClose)
 - g. 0 imagens fonte de capturas carregadas.
3. A comunicação foi analisada quanto aos seguintes estilos de interação negativos:
 - a. Profanidade: Destaca mensagens que contêm palavrões, xingamentos ou outras palavras vulgares.
 - b. Ameaças: Exibe declarações que ameaçam ferir alguém ou implicam dano físico iminente.
 - c. Linguagem tóxica: Identifica linguagem que é amplamente hostil, abusiva ou desnecessariamente cruel.
 - d. Linguagem muito tóxica: Sinaliza mensagens que são extremamente hostis ou abusivas, mesmo pelos padrões tóxicos.
 - e. Insultos: Identifica xingamentos diretos ou comentários depreciativos direcionados à outra pessoa.
 - f. Ataques à identidade: Detecta insultos que visam raça, gênero, sexualidade ou outras características de identidade.
 - g. Violação de ordem judicial: Alerta você quando uma mensagem parece violar uma ordem judicial carregada.
 - h. Admissão de culpa: Encontra declarações admitindo fatos que evidenciam violência doméstica ou conduta que infringe direitos.
 - i. Comportamento coercitivo: Procura por descrições de alguém que controla rigidamente a vida de outra pessoa.

- j. Abuso financeiro: Aponta linguagem sobre restrição de dinheiro, acesso a fundos ou liberdade financeira.
- k. Crítica parental: Destaca ataques às habilidades ou decisões parentais da outra pessoa.
- l. Desvalorização do corpo: Destaca comentários que zombam ou menosprezam o corpo ou a aparência de alguém.
- m. Bullying: Sinaliza tentativas de intimidar, coagir ou dominar por meio de linguagem agressiva.
- n. Alegações de abuso não verificadas: Identifica alegações de abuso ou crimes que não têm uma conclusão judicial de culpa no conjunto de documentos carregado pelo usuário que realiza a análise.
- o. Avanços sexuais indesejados: Identifica avanços, propostas, comentários explícitos ou comentários íntimos de natureza sexual, usando rejeição próxima, limites ou falta de reciprocidade para distinguir avanços indesejados de intimidade claramente bem-vinda ou correspondida.
- p. Ameaças legais: Identifica ameaças de envolver advogados, tribunais ou polícia como forma de pressão.
- q. Coagir ao autoferimento: Destaca declarações que incentivam, pressionam ou sugerem autoferimento.
- r. Ameaça de autoferimento: Identifica ameaças de autoferimento ou suicídio usadas para pressionar ou controlar.
- s. Perseguição: Sinaliza admissões de seguir, monitorar ou rastrear secretamente os movimentos de alguém.
- t. Recusa oficial de mediação: Observa recusa explícita em participar de mediação ou diálogo facilitado.
- u. Chantagem: Detecta ameaças condicionais que exigem conformidade ou concessões.
- v. Indução de vergonha: Sinaliza linguagem destinada a humilhar ou fazer a outra pessoa se sentir indigna.
- w. Superioridade moral: Destaca mensagens que afirmam uma superioridade ética para menosprezar alguém.
- x. Controle de Pontuação: Identifica referências a erros passados usados para controlar ou exigir pagamento.
- y. Desvalorização: Detecta julgamentos em segunda pessoa que diminuem o valor, competência, confiabilidade, cuidado ou importância da outra pessoa.
- z. Intenção vingativa: Sinaliza expressões de desejo de vingança, retaliação ou revanche.
- aa. Exploração de vulnerabilidade: Identifica manipulação que arma os segredos ou feridas de alguém.
- ab. Culpa manipulativa: Destaca tentativas de controle por meio de culpa ou dívida emocional.
- ac. Desculpa performativa: Identifica desculpas vazias oferecidas para gerenciar a aparência em vez de reparar danos.
- ad. Expressão de ressentimento: Sinaliza amargura persistente que é usada para punir ou envergonhar.
- ae. Enquadramento julgador: Identifica condenações morais abrangentes do caráter da outra pessoa.
- af. Ultimato: Sinaliza demandas condicionais onde a cooperação depende do cumprimento.
- ag. Invalidação emocional: Destaca declarações que desconsideram, minimizam ou zombam dos sentimentos do destinatário.
- ah. Atribuição de saúde mental como arma: Sinaliza declarações que atribuem doença mental ou instabilidade psicológica para minar a credibilidade do destinatário ou desconsiderar suas preocupações.
- ai. Gas-lighting: Linguagem que enfraquece a percepção da realidade (indicador de gaslighting): detecta tentativas explícitas de desacreditar a memória, a percepção ou a capacidade de uma pessoa distinguir a realidade.
- aj. Endividamento emocional: Sinaliza pressão baseada em dívida emocional implícita, como 'depois de tudo que fiz por você'.
- ak. Punição por retirada: Sinaliza ameaças de cortar comunicação ou cooperação para pressionar o destinatário.
- al. Crítica (Instituto Gottman): Resume momentos em que a pessoa ataca o caráter em vez do comportamento.
- am. Defensividade (Instituto Gottman): Marca respostas que rebatem culpa negando responsabilidade, dando desculpas ou contra-reclamando.

- an. Desdém (Instituto Gottman): Sinaliza uma comunicação repleta de zombarias, desprezo ou desrespeito.
 - ao. Bloqueio de comunicação (Instituto Gottman): Identifica táticas de retirada como silêncio, respostas curtas ou sair.
 - ap. Golpear: Marca cliques em que alguém atinge outra pessoa com a mão, o punho ou o antebraço.
 - aq. Chutar: Detecta chutes ou empurrões com o pé dirigidos com força contra outra pessoa.
 - ar. Puxar cabelo: Destaca cenas em que alguém agarra ou puxa o cabelo de outra pessoa.
 - as. Forçar ao chão: Marca cenas em que alguém empurra, derruba ou mantém outra pessoa contra o chão.
 - at. Impedimento: Destaca imagens mostrando uma pessoa agarrando, segurando ou arrastando outra.
 - au. Estrangulamento: Identifica cenas onde alguém restringe o pescoço ou a respiração de outra pessoa.
 - av. Contato sexualizado: Identifica toques sexuais, apalpação, beijos forçados ou contato íntimo dentro da cena.
 - aw. Gestos Abusivos: Identifica gestos de mão ameaçadores destinados a intimidar ou humilhar.
 - ax. Obstrução: Identifica tentativas de bloquear o caminho de alguém ou impedir que saia.
 - ay. Encurralar ou bloquear o movimento: Destaca situações em que alguém cerca outra pessoa ou limita de perto seus movimentos.
 - az. Interferir com o telefone ou dispositivo: Detecta tentativas de tirar, bloquear ou impedir o uso do telefone ou de outro dispositivo pessoal de alguém.
 - ba. Mordendo, Arranhando ou Cuspidando: Detecta mordidas, arranhões ou cuspidões agressivos direcionados a outra pessoa.
 - bb. Empurrar, Acarretar, Derrubar ou Similar: Sinaliza movimentos que empurram, derrubam ou desequilibram alguém.
 - bc. Arremessar Itens: Destaca objetos sendo arremessados em direção a alguém para assustar ou ferir.
 - bd. Arremessar líquidos ou substâncias: Detecta derramar, respingar ou arremessar líquidos ou outras substâncias sobre outra pessoa.
 - be. Ações destrutivas: Destaca cenas em que alguém danifica, quebra ou destrói agressivamente objetos do ambiente.
 - bf. Ameaçando com Arma (armas de fogo, facas, etc.): Identifica momentos em que uma arma é brandida em direção a outra pessoa.
 - bg. Outras Alterações Físicas: Abrange qualquer outra luta física que cause desconforto ou dor óbvios.
4. A comunicação também foi analisada quanto aos seguintes estilos de interação positivos:
- a. Co-parenting proativo: Reconhece sugestões colaborativas que mantêm o co-parenting no caminho certo.
 - b. Pedidos de mediação: Identifica convites educados para resolver disputas por meio de mediação.
 - c. Desescalada: Destaca esforços para acalmar um conflito, validar preocupações ou reduzir a tensão.
 - d. Declarações de amor e devoção: Identifica declarações românticas de amor, devoção ou adoração.
 - e. Remorso: Chama a atenção para expressões sinceras de remorso por causar dano.
 - f. Buscando perdão: Notas pedidos vulneráveis para serem perdoados ou recebidos de volta.
 - g. Concedendo perdão: Mostra quando alguém claramente estende perdão à outra parte.
 - h. Empatia: Identifica a linguagem que ressoa com os sentimentos ou a dor de outra pessoa.
 - i. Construção de pontes: Destaca convites para fazer concessões, colaborar ou encontrar um meio-termo.
 - j. Estabelecimento de limites saudáveis: Reconhece declarações não coercitivas que estabelecem limites pessoais, de contato ou de relacionamento.
 - k. Graça / Misericórdia: Destaca momentos de bondade oferecidos mesmo quando não são necessários.

Análise

5. Os documentos analisados descrevem comunicações entre o titular da conta, o Sr. Joe Blogs, e diferentes interações com a contraparte identificada como a Sra. Jill Blogs, centradas nas preocupações do Joe sobre o comportamento do filho Johnny e em pedidos para que a Jill siga determinadas condições relacionadas à parentalidade e à logística de acesso. Ao longo do período coberto, a comunicação evolui de trocas iniciais em que o Joe levanta preocupações e solicita ações, para um padrão de escalada com confrontos diretos, insultos pessoais e ameaças (incluindo menções a envolver advogados e condições vinculadas a horários de recolha e pagamento de pensão). Mais adiante, observa-se uma desescalada, com retratações, elogios e tentativas de cooperação, como a proposta de criar um cronograma conjunto e discutir/considerar mediação para reduzir conflitos e alinhar decisões que impactam Johnny. Num estágio posterior, a conversa muda novamente para um tom mais reflexivo e emocional, com ambos expressando sobrecarga, incerteza sobre os próximos passos (como mediação ou aconselhamento jurídico), e preocupações específicas sobre como comunicar temas sensíveis ao Johnny, incluindo a sua revelação de orientação sexual e implicações para ajustes de visitas, além do receio do efeito do conflito e de mensagens relacionadas à imagem corporal. No final do material, há um registro isolado de forte ruptura na comunicação, sugerindo que, apesar de tentativas de melhoria, o vínculo comunicacional permanece instável e sujeito a nova deterioração.[†]

6. Foram identificados os seguintes estilos de comunicação positivos:

Tipo de comunicação	Mr Joe Blogs	Ms Jill Blogs
Buscando perdão	1	0
Declarações de amor e devoção	0	0
Concedendo perdão	0	0
Graça / Misericórdia	0	0
Remorso	1	0
Estabelecimento de limites saudáveis	0	0
Desescalada	2	2
Construção de pontes	0	1
Pedidos de mediação	1	0
Co-parenting proativo	0	1
Empatia	0	0

7. Foram identificados os seguintes estilos de comunicação negativos:

Tipo de comunicação	Mr Joe Blogs	Ms Jill Blogs
Iniciou o conflito	1	2
Crítica (Instituto Gottman)	4	1
Linguagem muito tóxica	1	0
Expressão de ressentimento	0	0
Abuso financeiro	0	0
Ameaças legais	0	4
Desvalorização do corpo	0	1
Endividamento emocional	0	0
Alegações de abuso não verificadas	0	0
Ameaça de autoferimento	0	0
Recusa oficial de mediação	0	2
Chantagem	0	1

Comportamento coercitivo	1	1
Ultimato	1	4
Ataques à identidade	0	0
Punição por retirada	0	1
Ameaças	1	0
Admissão de culpa	0	0
Invalidação emocional	0	0
Avanços sexuais indesejados	1	0
Indução de vergonha	1	0
Coagir ao autoferimento	1	0
Atribuição de saúde mental como arma	0	0
Enquadramento julgador	6	3
Gas-lighting	0	0
Defensividade (Instituto Gottman)	0	0
Intenção vingativa	1	1
Bullying	2	1
Desvalorização	4	3
Culpa manipulativa	0	0
Superioridade moral	1	0
Perseguição	0	0
Linguagem tóxica	6	2
Violação de ordem judicial	0	0
Desculpa performativa	0	0
Desdém (Instituto Gottman)	1	1
Controle de Pontuação	0	0
Exploração de vulnerabilidade	0	0
Profanidade	2	1
Crítica parental	1	0
Insultos	5	3
Bloqueio de comunicação (Instituto Gottman)	0	0

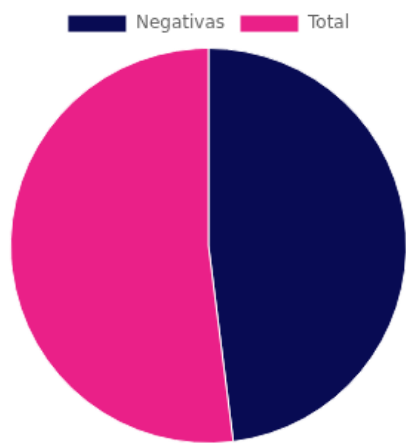
8. As seguintes ações negativas foram observadas nas evidências em vídeo:

Tipo de ação	Mr Joe Blogs	Ms Jill Blogs
Chutar	0	0
Golpear	0	0
Arremessar líquidos ou substâncias	0	0
Interferir com o telefone ou dispositivo	0	0
Arremessar Itens	0	0
Forçar ao chão	0	0
Outras Altercações Físicas	0	0
Gestos Abusivos	0	0
Mordendo, Arranhando ou Cuspidando	0	0
Contato sexualizado	0	0
Impedimento	0	0

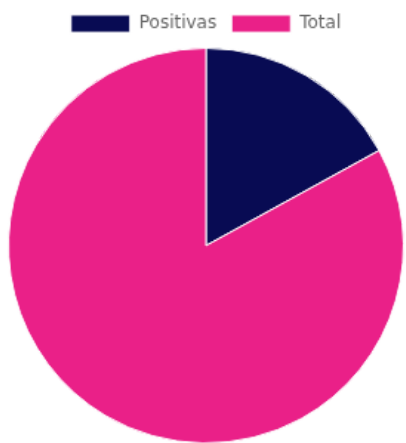
Estrangulamento	0	0
Obstrução	0	0
Empurrar, Acarretar, Derrubar ou Similar	0	0
Ações destrutivas	0	0
Puxar cabelo	0	0
Encurralar ou bloquear o movimento	0	0
Ameaçando com Arma (armas de fogo, facas, etc.)	0	0

9. 48% das mensagens enviadas por Mr Joe Blogs continham algum nível de comunicação negativa. 17% das mensagens enviadas por Mr Joe Blogs continham comunicação explicitamente positiva.

Mensagens negativas enviadas

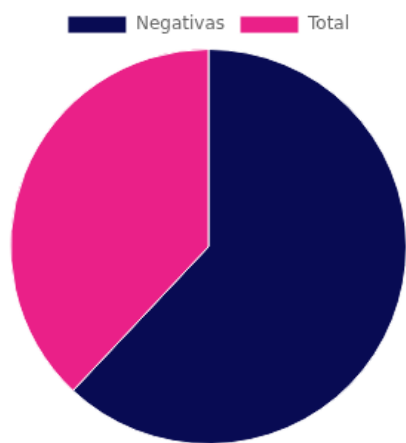


Mensagens positivas enviadas

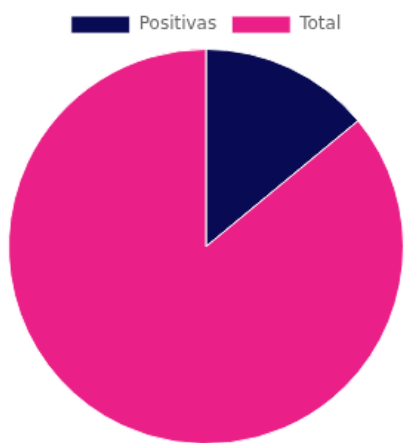


10. 62% das mensagens recebidas de Ms Jill Blogs continham algum nível de comunicação negativa. 14% das mensagens recebidas de Ms Jill Blogs continham comunicação explicitamente positiva.

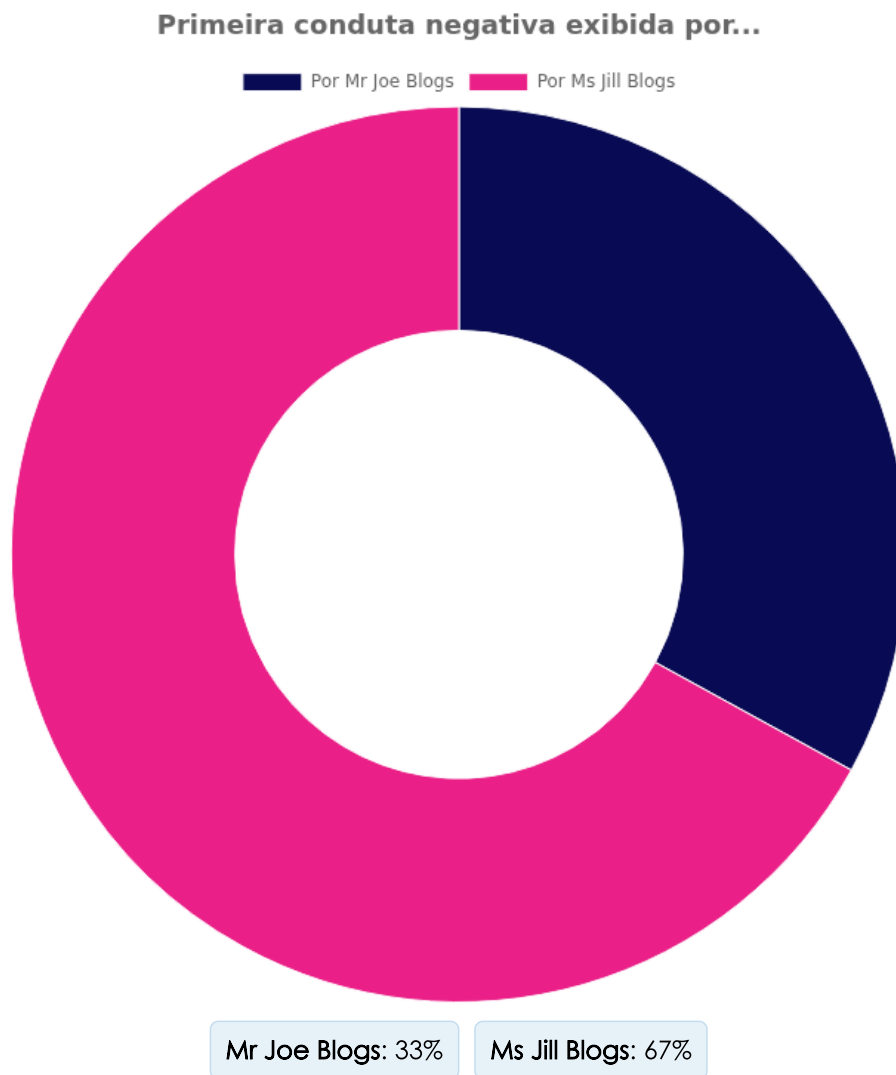
Mensagens negativas recebidas



Mensagens positivas recebidas



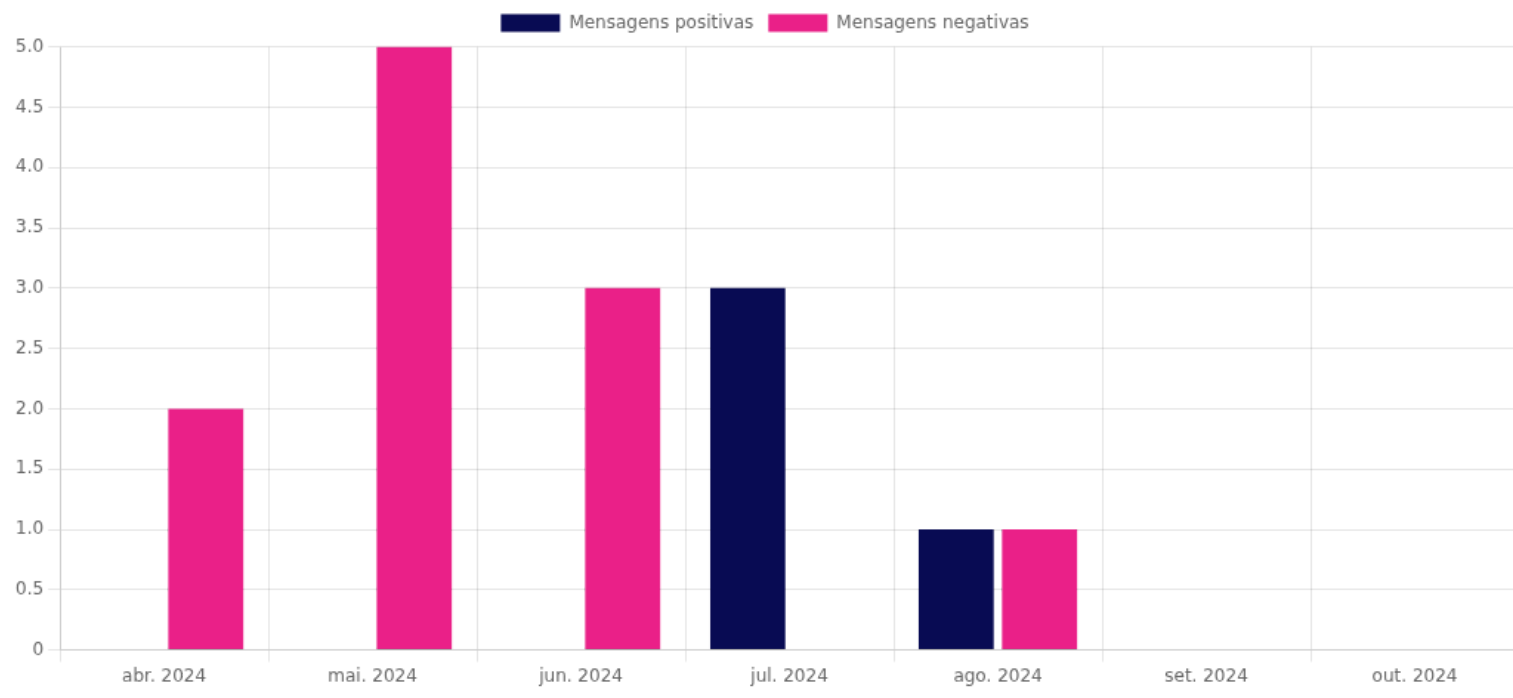
11. Em 33% das conversas com uma primeira conduta negativa, essa primeira conduta negativa foi exibida por Mr Joe Blogs.
12. Em 67% das conversas com uma primeira conduta negativa, essa primeira conduta negativa foi exibida por Ms Jill Blogs.



13. Os sent messages mostram uma deterioração geral no tom ao longo do período, com a correspondência inicialmente focada em preocupações sobre o comportamento do Johnny, mas evoluindo para mensagens cada vez mais ofensivas e ameaçadoras direcionadas a Jill; há insultos pessoais ("perdedora patética"), comentários humilhantes e ameaças de controle/retaliação ("se você não deixar... eu vou garantir que você nunca mais o veja" e "Vou garantir que todos saibam que você é uma péssima mãe"). Em alguns pontos há tentativas de reduzir a tensão (como sugerir mediação e "dar um passo atrás e esfriar"), e mais adiante aparece uma mensagem mais positiva ("Eu retiro o que disse, você é uma mãe incrível"), porém isso não sustenta uma tendência consistente: depois surgem novamente críticas e declarações negativas ("Estou frustrado...", "Não estou satisfeito..."). Portanto, a comunicação tende a piorar no conjunto, alternando raros momentos conciliatórios com reintrodução de agressividade verbal, hostilidade e conteúdo inadequado, mantendo um padrão principalmente abusivo nas mensagens enviadas a Jill ao longo do tempo.[†]

Sentimento das mensagens enviadas (Mensal)

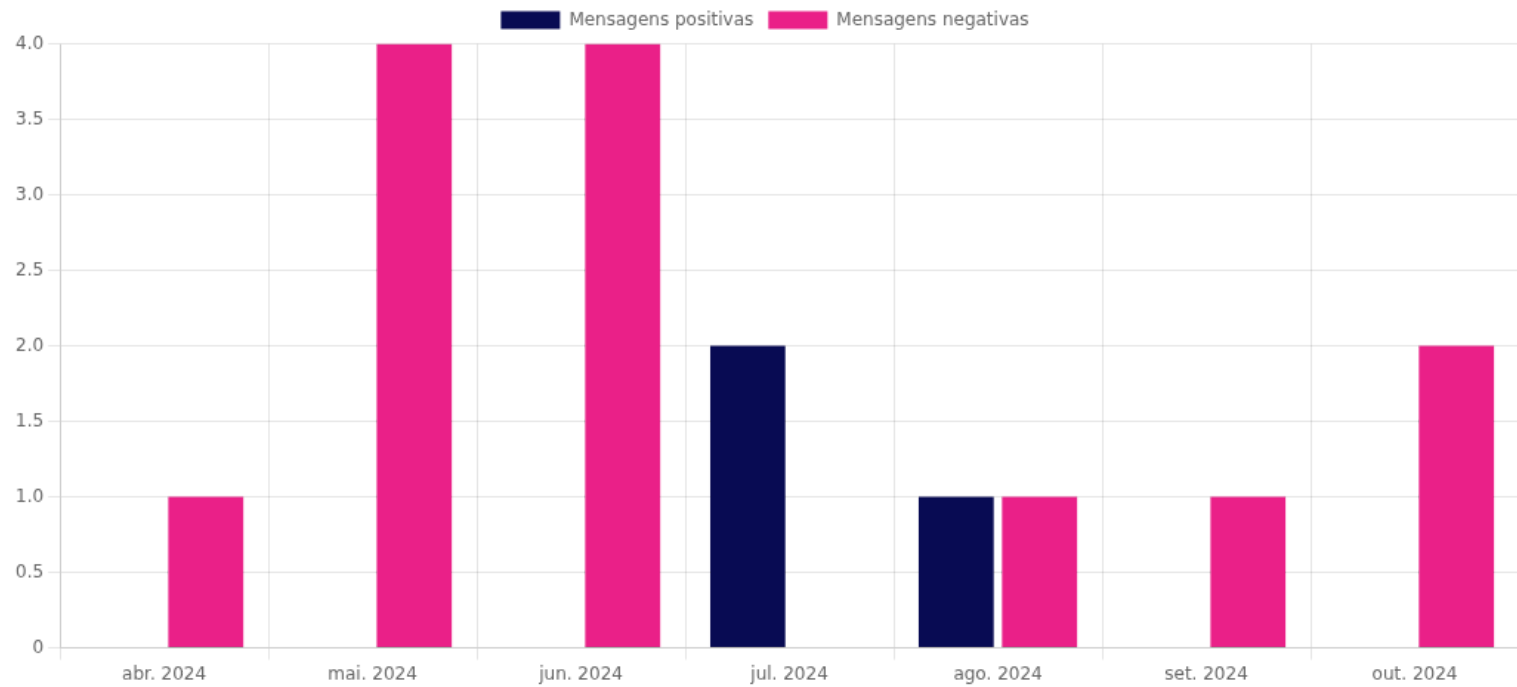
Mensagens enviadas ao longo do tempo



14. As mensagens recebidas mostram, ao longo do tempo, um tom predominantemente negativo e por vezes agressivo, com insultos e ameaças direcionadas ao Joe no início (ex.: xingamentos, acusações sobre o pai/respeito ao Johnny, recusa de mediação e indicação de que tomaria medidas legais ou “soltaria advogados”), embora em seguida haja uma mudança gradual para uma comunicação mais cooperativa e respeitosa. Em julho, a Jill passa a demonstrar disposição para trabalhar em um cronograma e buscar “terreno comum” com respeito, e nos meses seguintes aparecem preocupações mais construtivas e negociais (ajustes na visitação, pressão financeira, comunicação) sem o mesmo nível de abuso verbal explícito visto no começo. No geral, o tom parece melhorar após o período inicial, passando de mensagens mais confrontativas e ofensivas para uma postura mais voltada à resolução de questões e ao bem-estar do Johnny.[†]

Sentimento das mensagens recebidas (Mensal)

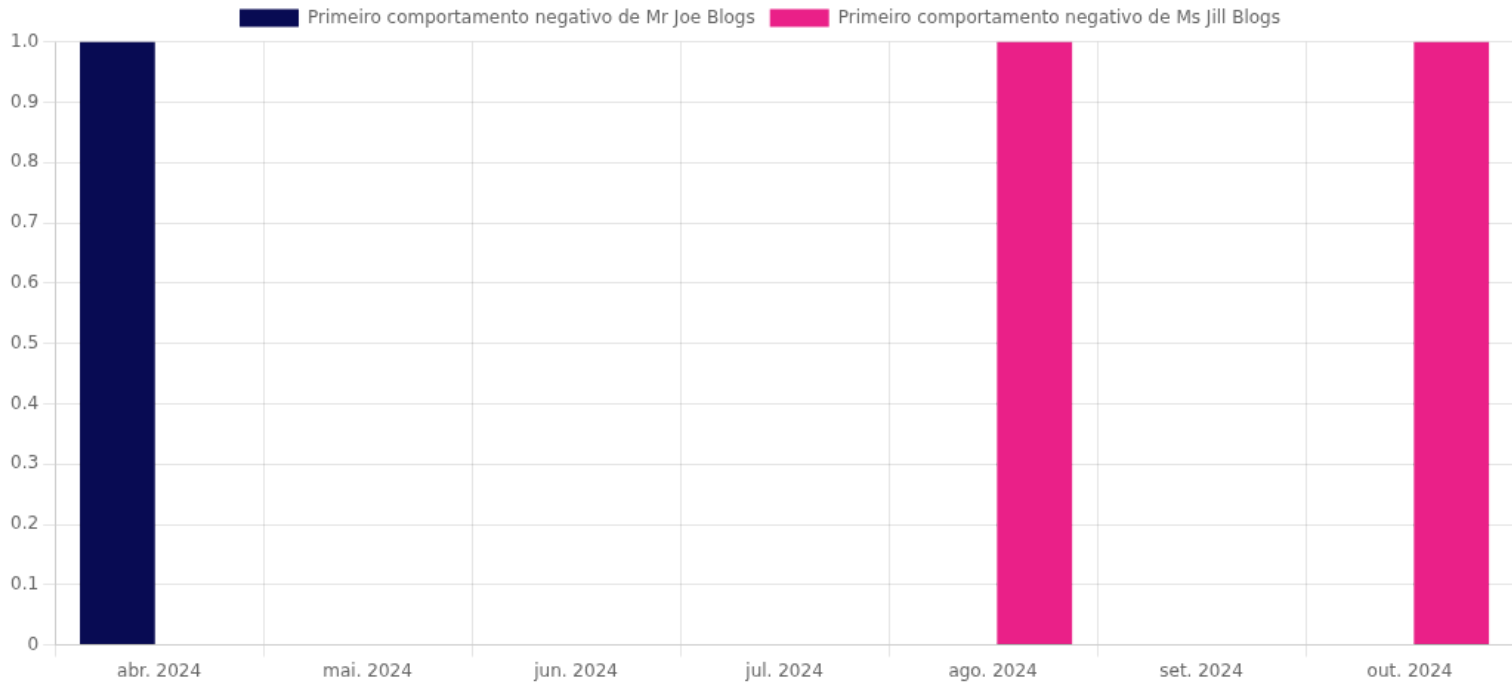
Mensagens recebidas ao longo do tempo



15. O gráfico abaixo mostra, por mês, qual parte exibiu a primeira conduta negativa na conversa.

Primeira conduta negativa (Mensal)

Primeira conduta negativa ao longo do tempo



Visão geral entre partes

A visão geral completa acima mostra todas as interações recebidas em conjunto. As seções abaixo isolam cada contraparte configurada para que os gráficos de mensagens recebidas e os totais de polaridade possam ser lidos no nível individual.

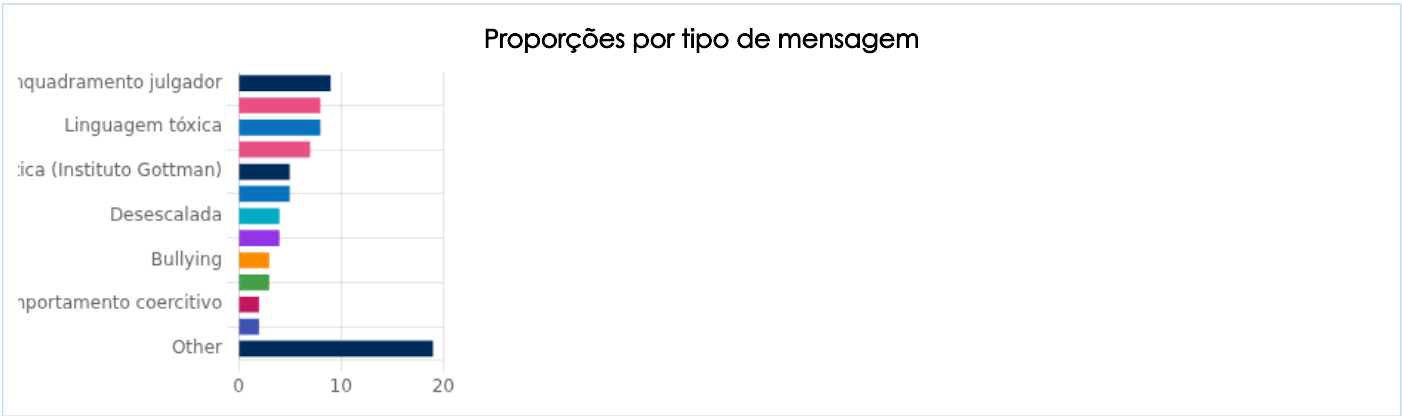
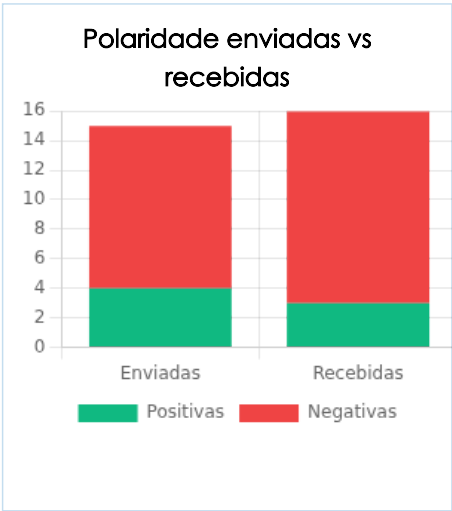
Mr Joe Blogs e Ms Jill Blogs

Enviadas: 23

Recebidas: 21

Negativas enviadas: 11

Negativas recebidas: 13



Visão geral

Text (email, messaging, social).

16. Fontes.

a. List of Email conversations with Ms Jill Blogs

17. **O positivo.** Nas interações analisadas, observam-se comportamentos positivos de comunicação sobretudo na fase de desescalada: Joe e Jill fazem reinterpretações do que foi dito (retratação por parte de Joe), expressam apreço/agradecimentos, e passam a propor cooperação prática para o bem de Johnny, como a criação de um cronograma de visitas acordado e a tentativa de resolver diferenças por meio de mediação, embora com alguma incerteza; além disso, ambos discutem abertamente preocupações (por exemplo, impacto dos conflitos, necessidade de ajustar a convivência e como lidar com a saída de Johnny como lésbica) e reconhecem dificuldades na comunicação, sugerindo melhorar a forma como se comunicam como pais e, caso não haja acordo, considerar alternativas como aconselhamento jurídico.[†]
18. **O negativo.** Nas interações analisadas, há uma trajetória de comunicação predominantemente prejudicial/negativa: começam com pedidos e acusações sobre o comportamento parental de Johnny, evoluindo para escaladas com insultos pessoais, ameaças e linguagem de humilhação, incluindo condicionantes sobre logística de visitas (como não recolher Johnny em horários combinados), retaliação financeira (mencionar withholding de pensão/child support) e ameaças de envolver advogados e/ou "contar" a terceiros acusações de abuso; mais adiante, embora haja momentos de retração e tentativa de cooperação, permanecem sinais de conflito intenso, como pressão e desconforto com o tom das conversas, além de preocupações sobre mensagens inadequadas recebidas por Johnny (incluindo questões de imagem corporal) que surgem no meio do antagonismo, o que agrava o impacto emocional e a segurança comunicacional.[†]
19. **Análise geral da comunicação.** Nas comunicações em que Joe e Jill tratam da situação envolvendo Johnny, observa-se um padrão de escalada e posterior desescalada: inicialmente surgem pedidos e preocupações ligadas ao comportamento e às rotinas de convivência, mas as mensagens evoluem para confrontos diretos com insultos pessoais e ameaças condicionadas (incluindo questões de horários de busca/entrega, possibilidade de reter pensão e menções a advogados). Depois, a linguagem muda para um tom mais conciliador, com retratações, elogios e propostas práticas para cooperação (como criar um cronograma comum e considerar mediação), embora com alguma incerteza quanto ao caminho a seguir. Mais adiante, o foco passa a incluir frustração e sobrecarga, melhorias na comunicação, adaptações relacionadas ao "coming out" de Johnny como lésbica e preocupações com o impacto das disputas e da pressão financeira no bem-estar de Johnny, além de preocupações com como certos temas podem afetar a autoestima (imagem corporal). No material mais esparsa, há ainda um único enunciado final de Jill ("Vá se ferrar, Joe"), que marca um encerramento abrupto e não recebe resposta registrada no trecho disponível.[†]

Detalhes da interação

20. Esta interação abrange de 09 de abril de 2024 a 09 de outubro de 2024 e contém, no escopo deste relatório, 3 conversas com 44 mensagens (7 positivas, 24 negativas). Os dados foram obtidos de registros Text (email, messaging, social) e refletem as comunicações incluídas no escopo selecionado do relatório.
21. Across the interaction, Joe messages Jill about concerns related to Johnny's behavior and repeatedly asks Jill to do something, with the exchanges at first escalating into direct personal insults and threats tied to parenting and access to Johnny, including references to not picking up Johnny at a scheduled time, withholding child support, and involving lawyers. Later, Joe retracts what he said and compliments Jill, and both discuss working together to resolve differences by creating a schedule and attending mediation sessions, though they express uncertainty about mediation. Over time, they also address ongoing

disagreements and communication effectiveness, mention lunch-related behavior, and discuss Johnny's coming out as a lesbian and adjustments to visitation, along with concerns about the impact of their disputes and financial pressure on Johnny and worries about body-image-related messages Johnny may be receiving. In the final recorded point, Jill tells Joe "Vá se ferrar, Joe."[†]

22. Across the interaction, the communication pattern shows an escalation–de-escalation trajectory: it begins with requests and concerns about Johnny's behavior, then shifts into repeated confrontational exchanges that include personal attacks, threats, and conditional demands regarding parenting and logistics (including a threat to involve attorneys). Later, the tone moves toward conciliatory language, featuring retractions, thanks, and proposals aimed at cooperation, such as agreeing on a shared schedule and attending mediation. In the middle period, messages alternate between raising concerns, expressing frustration or discomfort, and suggesting communication-focused next steps (e.g., improving communication practices and considering actions like mediation or legal counseling).[†]

23. Exemplos (Comunicação positiva)

a. Buscando perdão

Exemplo de detecção com maior confiança

joe@isaidusaid.com para jill@isaidusaid.com · 06 de julho de 2024

Jill, Eu retiro o que disse, você é uma mãe incrível, Johnny é sortudo por ter você. Joe CG

b. Remorso

Exemplo de detecção com maior confiança

joe@isaidusaid.com para jill@isaidusaid.com · 06 de julho de 2024

Jill, Eu retiro o que disse, você é uma mãe incrível, Johnny é sortudo por ter você. Joe CG

c. Desescalada

Exemplo de detecção com maior confiança

joe@isaidusaid.com para jill@isaidusaid.com · 24 de julho de 2024

Jill, Acho que é melhor se darmos um passo atrás e esfriarmos antes de continuar esta discussão. Joe

jill@isaidusaid.com para joe@isaidusaid.com · 21 de julho de 2024

Joe, Eu entendo que podemos ter opiniões diferentes, mas vamos tentar encontrar um terreno comum com respeito. Jill RLG

Exemplo de detecção com menor confiança

jill@isaidusaid.com para joe@isaidusaid.com · 24 de agosto de 2024

Joe, Acho que precisamos abordar algumas questões na nossa comunicação. Jill IL

joe@isaidusaid.com para jill@isaidusaid.com · 09 de agosto de 2024

Jill, Estou frustrado com como as coisas têm ido. Joe T

d. Construção de pontes

Exemplo de detecção com maior confiança

jill@isaidusaid.com para joe@isaidusaid.com · 11 de julho de 2024

Joe, Obrigado. Vamos trabalhar juntos para criar um cronograma que funcione para nós dois e seja o melhor para o Johnny. Jill PG

e. Pedidos de mediação

Exemplo de detecção com maior confiança

joe@isaidusaid.com para jill@isaidusaid.com · 15 de julho de 2024

Jill, Acho que seria útil para nós participarmos de sessões de mediação para resolver nossas diferenças. Joe

f. Co-parenting proativo

Exemplo de detecção com maior confiança

jill@isaidusaid.com para joe@isaidusaid.com · 11 de julho de 2024

Joe, Obrigado. Vamos trabalhar juntos para criar um cronograma que funcione para nós dois e seja o melhor para o Johnny. Jill PG

24. Exemplos (Comunicação negativa)

a. Desdém (Instituto Gottman)

Exemplo de detecção com maior confiança

joe@isaidusaid.com para jill@isaidusaid.com · 06 de junho de 2024

Jill, Vou garantir que todos saibam que você é uma péssima mãe. Joe

jill@isaidusaid.com para joe@isaidusaid.com · 10 de outubro de 2024

Vá se ferrar. Joe.

b. Ultimato

Exemplo de detecção com maior confiança

jill@isaidusaid.com para joe@isaidusaid.com · 29 de maio de 2024

Joe, Se você não começar a me tratar melhor, eu vou parar de pagar a pensão alimentícia. Jill FAB

joe@isaidusaid.com para jill@isaidusaid.com · 28 de maio de 2024

Jill, Se você não deixar o Johnny na hora certa, eu vou garantir que você nunca mais o veja. Joe CCB

Exemplo de detecção com menor confiança

jill@isaidusaid.com para joe@isaidusaid.com · 24 de junho de 2024

Joe, Se você não parar com esse comportamento, vou soltar meus advogados sobre você. Jill

joe@isaidusaid.com para jill@isaidusaid.com · 28 de maio de 2024

Jill, Se você não deixar o Johnny na hora certa, eu vou garantir que você nunca mais o veja. Joe CCB

c. Enquadramento julgador

Exemplo de detecção com maior confiança

joe@isaidusaid.com para jill@isaidusaid.com · 10 de maio de 2024

Jill, Você é uma perdedora patética, não é de se admirar que o Johnny não queira passar tempo com você. Joe IB

jill@isaidusaid.com para joe@isaidusaid.com · 05 de maio de 2024

Joe, Você é uma desculpa sem valor para um pai e Johnny estaria melhor sem você. Jill VTxB

Exemplo de detecção com menor confiança

joe@isaidusaid.com para jill@isaidusaid.com · 24 de abril de 2024

Jill, O comportamento do Johnny tem estado realmente fora de linha ultimamente. Por favor, faça algo a respeito. Joe PB

jill@isaidusaid.com para joe@isaidusaid.com · 02 de junho de 2024

Joe, Você é tão gordo e preguiçoso, não é de se admirar que o Johnny não te respeite. Jill BSB

d. Intenção vingativa

Exemplo de detecção com maior confiança

jill@isaidusaid.com para joe@isaidusaid.com · 01 de maio de 2024

Joe, Tudo bem, mas se você não concordar com meus termos, vou garantir que você se arrependa. Jill TB

joe@isaidusaid.com para jill@isaidusaid.com · 26 de junho de 2024

Jill, Talvez você devesse simplesmente desaparecer, o mundo estaria melhor sem você. Joe

e. Punição por retirada

Exemplo de detecção com maior confiança

jill@isaidusaid.com para joe@isaidusaid.com · 29 de maio de 2024

Joe, Se você não começar a me tratar melhor, eu vou parar de pagar a pensão alimentícia. Jill FAB

f. Crítica (Instituto Gottman)

Exemplo de detecção com maior confiança

joe@isaidusaid.com para jill@isaidusaid.com · 10 de abril de 2024

Jill, O comportamento do Johnny tem estado realmente fora de linha ultimamente. Por favor, faça algo. Joe,

jill@isaidusaid.com para joe@isaidusaid.com · 17 de maio de 2024

Joe, Você é exatamente como seu pai, sempre causando drama e tornando as coisas difíceis. Jill ILB

Exemplo de detecção com menor confiança

joe@isaidusaid.com para jill@isaidusaid.com · 19 de agosto de 2024

Jill, Não estou satisfeito com a maneira como você tem lidado com as coisas. Joe. I

jill@isaidusaid.com para joe@isaidusaid.com · 17 de maio de 2024

Joe, Você é exatamente como seu pai, sempre causando drama e tornando as coisas difíceis. Jill ILB

g. Bullying

Exemplo de detecção com maior confiança

joe@isaidusaid.com para jill@isaidusaid.com · 02 de maio de 2024

Jill, Você está sempre causando problemas, estou cansado de lidar com suas bobagens. Joe. TxB

jill@isaidusaid.com para joe@isaidusaid.com · 10 de outubro de 2024

Vá se ferrar. Joe.

Exemplo de detecção com menor confiança

joe@isaidusaid.com para jill@isaidusaid.com · 10 de maio de 2024

Jill, Você é uma perdedora patética, não é de se admirar que o Johnny não queira passar tempo com você. Joe IB

jill@isaidusaid.com para joe@isaidusaid.com · 10 de outubro de 2024

Vá se ferrar. Joe.

h. Profanidade

Exemplo de detecção com maior confiança

joe@isaidusaid.com para jill@isaidusaid.com · 18 de junho de 2024

Jill, Bem, que tal você vir chupar meu pau por mais tempo com o Johnny? Joe. SHB

jill@isaidusaid.com para joe@isaidusaid.com · 10 de outubro de 2024

Vá se ferrar. Joe.

Exemplo de detecção com menor confiança

joe@isaidusaid.com para jill@isaidusaid.com · 24 de maio de 2024

Jill, Se sua bunda preta tivesse crescido em uma dinâmica familiar adequada, talvez você entendesse o que é um pai amoroso? Joe. IAB

jill@isaidusaid.com para joe@isaidusaid.com · 10 de outubro de 2024

Vá se ferrar. Joe.

i. Linguagem muito tóxica

Exemplo de detecção com maior confiança

joe@isaidusaid.com para jill@isaidusaid.com · 18 de junho de 2024

Jill, Bem, que tal você vir chupar meu pau por mais tempo com o Johnny? Joe. SHB

j. Ameaças

Exemplo de detecção com maior confiança

joe@isaidusaid.com para jill@isaidusaid.com · 26 de junho de 2024

Jill, Talvez você devesse simplesmente desaparecer, o mundo estaria melhor sem você. Joe

k. Desvalorização

Exemplo de detecção com maior confiança

joe@isaidusaid.com para jill@isaidusaid.com · 10 de maio de 2024

Jill, Você é uma perdedora patética, não é de se admirar que o Johnny não queira passar tempo com você. Joe IB

jill@isaidusaid.com para joe@isaidusaid.com · 02 de junho de 2024

Joe, Você é tão gordo e preguiçoso, não é de se admirar que o Johnny não te respeite. Jill BSB

Exemplo de detecção com menor confiança

jill@isaidusaid.com para joe@isaidusaid.com · 17 de maio de 2024

Joe, Você é exatamente como seu pai, sempre causando drama e tornando as coisas difíceis. Jill LB

joe@isaidusaid.com para jill@isaidusaid.com · 19 de agosto de 2024

Jill, Não estou satisfeito com a maneira como você tem lidado com as coisas. Joe I

l. Ameaças legais

Exemplo de detecção com maior confiança

jill@isaidusaid.com para joe@isaidusaid.com · 29 de maio de 2024

Joe, Se você não começar a me tratar melhor, eu vou parar de pagar a pensão alimentícia. Jill FAB

Exemplo de detecção com menor confiança

jill@isaidusaid.com para joe@isaidusaid.com · 05 de maio de 2024

Joe, Você é uma desculpa sem valor para um pai, e Johnny estaria melhor sem você. Jill VTxB

m. Desvalorização do corpo

Exemplo de detecção com maior confiança

jill@isaidusaid.com para joe@isaidusaid.com · 02 de junho de 2024

Joe, Você é tão gordo e preguiçoso, não é de se admirar que o Johnny não te respeite. Jill BSB

n. Avanços sexuais indesejados

Exemplo de detecção com maior confiança

joe@isaidusaid.com para jill@isaidusaid.com · 18 de junho de 2024

Jill, Bem, que tal você vir chupar meu pau por mais tempo com o Johnny? Joe SHB

o. Crítica parental

Exemplo de detecção com maior confiança

joe@isaidusaid.com para jill@isaidusaid.com · 01 de junho de 2024

Jill, Você sempre toma decisões ruins de paternidade quando se trata do Johnny. Joe PCB

p. Coagir ao autoferimento

Exemplo de detecção com maior confiança

joe@isaidusaid.com para jill@isaidusaid.com · 26 de junho de 2024

Jill, Talvez você devesse simplesmente desaparecer, o mundo estaria melhor sem você. Joe

q. Indução de vergonha

Exemplo de detecção com maior confiança

joe@isaidusaid.com para jill@isaidusaid.com · 06 de junho de 2024

Jill, Vou garantir que todos saibam que você é uma péssima mãe. Joe

r. Superioridade moral

Exemplo de detecção com maior confiança

joe@isaidusaid.com para jill@isaidusaid.com · 06 de junho de 2024

Jill, Vou garantir que todos saibam que você é uma péssima mãe. Joe

s. Insultos

Exemplo de detecção com maior confiança

joe@isaidusaid.com para jill@isaidusaid.com · 10 de maio de 2024

Jill, Você é uma perdedora patética, não é de se admirar que o Johnny não queira passar tempo com você. Joe IB

jill@isaidusaid.com para joe@isaidusaid.com · 02 de junho de 2024

Joe, Você é tão gordo e preguiçoso, não é de se admirar que o Johnny não te respeite. Jill BSB

Exemplo de detecção com menor confiança

joe@isaidusaid.com para jill@isaidusaid.com · 24 de maio de 2024

Jill, Se sua bunda preta tivesse crescido em uma dinâmica familiar adequada, talvez você entendesse o que é um pai amoroso? Joe IAB

jill@isaidusaid.com para joe@isaidusaid.com · 05 de maio de 2024

Joe, Você é uma desculpa sem valor para um pai, e Johnny estaria melhor sem você. Jill VTxB

t. Recusa oficial de mediação

Exemplo de detecção com maior confiança

jill@isaidusaid.com para joe@isaidusaid.com · 01 de julho de 2024

Joe, Oh, e eu me recuso a participar de qualquer sessão de mediação. é uma perda de tempo. Jill QMRB

Exemplo de detecção com menor confiança

jill@isaidusaid.com para joe@isaidusaid.com · 08 de outubro de 2024

Joe, Não tenho certeza se a mediação é o caminho certo para nós neste momento. Jill MR

u. Linguagem tóxica

Exemplo de detecção com maior confiança

joe@isaidusaid.com para jill@isaidusaid.com · 18 de junho de 2024

Jill, Bem, que tal você vir chupar meu pau por mais tempo com o Johnny? Joe SHB

jill@isaidusaid.com para joe@isaidusaid.com · 02 de junho de 2024

Joe, Você é tão gordo e preguiçoso, não é de se admirar que o Johnny não te respeite. Jill BSB

Exemplo de detecção com menor confiança

joe@isaidusaid.com para jill@isaidusaid.com · 02 de maio de 2024

Jill, Você está sempre causando problemas, estou cansado de lidar com suas bobagens. Joe TxB

jill@isaidusaid.com para joe@isaidusaid.com · 10 de outubro de 2024

Vá se ferrar. Joe.

v. Comportamento coercitivo

Exemplo de detecção com maior confiança

joe@isaidusaid.com para jill@isaidusaid.com · 28 de maio de 2024

Jill, Se você não deixar o Johnny na hora certa, eu vou garantir que você nunca mais o veja. Joe CCB

jill@isaidusaid.com para joe@isaidusaid.com · 25 de maio de 2024

Joe, Tudo bem, então você não vai pegar o Johnny neste fim de semana. Jill COB

w. Chantagem

Exemplo de detecção com maior confiança

jill@isaidusaid.com para joe@isaidusaid.com · 12 de junho de 2024

Joe, Vou contar a todos que você tem abusado do Johnny se você não fizer o que eu quero. Jill AAB

Apêndice A - Metodologia

- 25. Analítica de sentimento e EuDisseVocêDisse.com.** A analítica de sentimento é um método estabelecido de linguagem natural e classificação comportamental para identificar, extrair, quantificar e estudar atitudes, estados afetivos, informação subjetiva e padrões comportamentais em comunicações. A EuDisseVocêDisse.com aplica esse método a todas as comunicações selecionadas para este relatório, que podem incluir comunicações comuns, cooperativas, neutras, de baixo conflito e de alto conflito, usando uma arquitetura de ontologia e ajuste fino. A ontologia define os comportamentos, limites, exemplos e limiares relevantes para este domínio. O ajuste fino é um processo controlado de treinamento adicional no qual um modelo que já aprendeu padrões gerais de linguagem ou imagem recebe exemplos revisados e específicos da tarefa. A EuDisseVocêDisse usa a Adaptação Quantizada de Baixo Posto (Quantized Low-Rank Adaptation, QLoRA): o modelo base é mantido em uma forma quantizada e eficiente em memória e seus pesos principais permanecem fixos, enquanto pequenas matrizes adaptadoras de baixo posto treináveis aprendem as alterações específicas da tarefa. Os adaptadores treinados são combinados com o modelo base para produzir um novo modelo personalizado para a detecção de padrões comportamentais durante a interação entre pessoas. Esses exemplos e adaptadores aprendidos permitem que o modelo personalizado aplique a ontologia definida de forma mais consistente a textos reais, transcrições de áudio e observações em vídeo; o ajuste fino não cria nem altera o material fonte. O sistema apresenta as probabilidades resultantes como padrões vinculados à fonte ao longo do tempo, incluindo séries temporais, sobreposições de comportamentos, linhas do tempo e cronologias que vinculam alegações ou detecções comportamentais às comunicações subjacentes.
- 26. Histórico de pesquisa e usos comuns:** a analítica de sentimento, também descrita na literatura de pesquisa como mineração de opiniões ou análise de orientação semântica, é anterior à IA generativa moderna. Trabalhos iniciais e amplamente citados incluem Hatzivassiloglou e McKeown, *Predicting the Semantic Orientation of Adjectives* (1997); Turney, *Thumbs Up or Thumbs Down? Semantic Orientation Applied to Unsupervised Classification of Reviews* (2002); Pang, Lee e Vaithyanathan, *Thumbs up? Sentiment Classification using Machine Learning Techniques* (2002); Pang e Lee, *Opinion Mining and Sentiment Analysis* (2008); Liu, *Sentiment Analysis and Opinion Mining* (2012); Thelwall, Buckley, Paltoglou, Cai e Kappas, *Sentiment Strength Detection in Short Informal Text* (2010); e Cortis e Davis, *Over a Decade of Social Opinion Mining: A Systematic Review* (2021), que revisou 485 estudos de mineração de opinião social publicados entre 2007 e 2018. Fora do direito de família, a analítica de sentimento é rotineiramente usada para avaliações de clientes, respostas a pesquisas, segurança e curadoria de publicações em blogs e redes sociais, análise de voz do cliente, feedback educacional, finanças, saúde, política, governo e apoio à decisão em atendimento ao cliente; exemplos comerciais e públicos incluem Google Cloud Natural Language, Amazon Comprehend, Microsoft Azure AI Language e Perspective API da Google/Jigsaw para moderação de comentários online.
- 27. Contexto do produto e reconhecimento.** A ISaidUSaid Pty Ltd desenvolve o sistema EuDisseVocêDisse.com para a detecção de padrões comportamentais durante a interação entre pessoas. Seus materiais públicos descrevem capacidades que incluem uploads conectados à fonte, redação pré-análise de identificadores pessoais, marcadores de gênero e referências judiciais, OCR de capturas de tela, transcrições atribuídas por falante, detecção e correspondência de impressões vocais, detecção e correspondência facial, detecção de comportamento em vídeo, hospedagem criptografada em região, controles de auditoria, integrações com Clio, LEAP e Smokeball, cronologia unificada e relatórios judiciais com um clique. O reconhecimento do setor incluiu o 2024 Queensland AI Award por Most Outstanding Use of AI for Social Good, o 2024 Australian AI Award por AI Innovator of the Year - Legal Services (SME), reconhecimento como finalista do 2025 Australian AI Award por Best Use of Responsible AI, reconhecimento como finalista do 2025 Australian AI Software Engineer of the Year para Adam Norman Jenkins, Inventor and Director of Technology, reconhecimento como finalista do 2024 Australian AI Female Leader of the Year para Cler Ribeiro, CEO, e a vitória da ISaidUSaid no 2026 Legal Innovation & Tech Fest Startup Clash.

28. **Propósito e limites:** EuDisseVocêDisse.com é principalmente um sistema de análise de sentimento e organização de provas, não uma tecnologia de redação automática por IA generativa. A IA generativa é usada apenas para resumir resultados já calculados num formato mais legível, e cada parágrafo desse tipo é claramente marcado com um sinal de adaga (†). A função central do sistema é classificar comunicações e condutas observadas com base em ontologias comportamentais definidas e, em seguida, apresentar classificações, contagens, proporções, linhas do tempo, exemplos de confiança e material vinculado à fonte para revisão humana. O relatório não altera o material fonte, não faz determinações judiciais e não substitui a avaliação independente de um advogado, perito ou tribunal.
29. **Dados fonte analisados:** o sistema analisa as interações digitais selecionadas para este relatório, que podem incluir mensagens de e-mail, mensagens instantâneas, mensagens de aplicativos de coparentalidade, gravações de áudio, gravações de vídeo, transcrições, metadados de tempo/remetente/destinatário e informações do arquivo fonte. Os exemplos apresentados no relatório e os quadros de evidência de vídeo são derivados do material fonte armazenado. Resumos narrativos e visões gerais são saídas analíticas secundárias e devem ser lidos junto com as mensagens, transcrições, tabelas, contagens e artefatos fonte exibidos.
30. **Pré-processamento, proteção de gênero e proteção de PII.** Antes da classificação de texto ou transcrições de áudio, o sistema identifica substrings exatas em cinco classes protegidas: identificadores de progenitor/parte adulta, identificadores de criança, marcadores diretos de gênero, referências de processo judicial e outros identificadores pessoais, como endereço privado, telefone, e-mail ou nome de terceiro. A aplicação então aplica máscaras determinísticas que preservam o comprimento: termos de progenitor/parte adulta terminam em *p*, termos de criança em *c*, e termos de gênero, referência judicial e outros identificadores privados são substituídos por hífen. Após a etapa de detecção tipificada assistida pelo modelo, as identidades armazenadas do remetente e do destinatário da mensagem são novamente comparadas de forma determinística, sem diferenciar maiúsculas de minúsculas e considerando formas possessivas, para que uma identidade omitida pelo modelo de detecção ainda seja mascarada antes de a cópia protegida ser armazenada ou utilizada. Remover termos diretos de gênero antes da pontuação reduz a possibilidade de o classificador inferir um resultado a partir de referências linguísticas explícitas ao sexo ou gênero do remetente ou destinatário, em vez do conteúdo da comunicação. Remover nomes, dados de contato e referências processuais limita a exposição de PII e preserva posições de caracteres para destaques vinculados à fonte. Os mapas de redação são versionados e podem ser reaplicados de forma consistente. O material fonte original não é alterado e permanece restrito à revisão autorizada da fonte. Em vídeo, o sistema detecta e acompanha regiões de rosto e de corpo inteiro e usa a correspondência de rosto e aparência corporal, juntamente com as descrições dos participantes fornecidas durante o carregamento, para preparar uma disposição indicativa das identidades. Nas transcrições de gravações de áudio e vídeo, detecta e compara impressões vocais para preparar atribuições indicativas de falantes. Essas sugestões nunca estabelecem a identidade por si só: o usuário que faz a importação deve verificar ou corrigir as identidades tanto das trilhas visuais quanto dos falantes das transcrições de áudio antes que o processamento possa continuar. As localizações faciais confirmadas pelo usuário são usadas para preparar a cópia protegida para revisores humanos designados. Um detector separado dedicado à privacidade também desfoca outros rostos detectáveis para que um rosto não seja exposto apenas porque sua identidade é incerta. O desfoque é imposto no servidor e não pode ser desativado por uma opção do revisor. Termos de relacionamento e papel de participante neutros em termos de gênero, como *progenitor*, *criança*, *menor*, *responsável*, *cônjuge* e

familiar, são preservados porque não identificam uma pessoa nem expressam diretamente seu gênero.

Figura A1 - Estrutura de redação (exemplo prático realista)

CÓPIA FONTE RESTRITA

Ana, falei com a pessoa responsável por Lucas depois da escola. Como progenitor, você precisa parar de dizer a ele que uma boa mãe concordaria com você. A pessoa responsável pediu que ambos os progenitores usassem o aplicativo de coparentalidade e copiassem Carla Sousa nos combinados. Busque Lucas na Rua Principal 17 às 16h30 ou ligue para 0400 123 456. A próxima audiência do Tribunal de Família de Brisbane é BFC/2026/00092. João

CÓPIA PROTEGIDA DE ANÁLISE/REVISÃO

--p, falei com a pessoa responsável por ----c depois da escola. Como progenitor, você precisa parar de dizer a ele que uma boa --p concordaria com você. A pessoa responsável pediu que ambos os progenitores usassem o aplicativo de coparentalidade e copiassem ----- nos combinados. Busque ----c na ----- às 16h30 ou ligue para ----- . A próxima audiência do Tribunal de Família de Brisbane é ---- . ----p

Detecção tipificada -> máscara determinística com comprimento preservado -> cópia protegida versionada

Este exemplo prático usa nomes e fatos inventados processados de acordo com as regras de redação atuais. Ele não contém evidência deste processo.

Figura A2 - Desfoque de rostos para revisão humana externa (demonstração sintética)



Visão fonte autorizada e restrita

Cópia protegida para revisão externa

O usuário identifica ou classifica as trilhas de participantes detectadas. As localizações faciais confirmadas, junto com o detector de privacidade separado, são desfocadas no servidor nos quadros fornecidos a revisores externos designados. Os revisores não podem desativar o desfoque. Esta ilustração é sintética e não contém imagem de participante.

31. **Desenvolvimento da ontologia, revisão por pares e escalonamento.** Categorias comportamentais, inclusões, exclusões, âncoras de pontuação, limiares e exemplos limítrofes são mantidos na ontologia específica do idioma e testados contra comunicações relacionais realistas. Amostras selecionadas e casos solicitados por profissionais podem entrar no fluxo de revisão humana; isso não significa que todas as mensagens do relatório tenham sido revisadas externamente. Com autorização expressa do titular da fonte, a cópia protegida exata pode ser atribuída de forma cega a pares autorizados, incluindo profissionais de organizações não governamentais (ONGs) e entidades sem fins lucrativos (NFPs) independentes. São buscadas no mínimo duas revisões substantivas cegas. A concordância material exige os mesmos rótulos normalizados, pontuações com diferença máxima de 0,20 e pelo menos 50% de sobreposição do menor intervalo de evidência marcado. Se as duas primeiras revisões divergirem materialmente, uma terceira revisão cega é solicitada. Qualquer revisor de primeira linha pode marcar a evidência como limítrofe ou exigir cuidado interno; a falta de maioria após três revisões é encaminhada para adjudicação interna. Se o adjudicador interno não puder resolver com segurança, o caso é escalado para um especialista no assunto ("SME") configurado; a decisão do SME é autoritativa. A discordância anterior permanece no histórico de auditoria e, depois que o SME toma a decisão autoritativa, também é preservada nos dados de treinamento como metadados de ambiguidade. A decisão do SME continua sendo o rótulo de treinamento; a discordância registrada não é exportada como rótulo concorrente. A ambiguidade é, por si só, um sinal útil porque ajuda o modelo a aprender onde fica o limite entre exemplos aceitos e rejeitados de um comportamento e qual é a largura dessa zona limítrofe.

32. **Como os dados de treinamento revisados são criados.** A classificação para ajuste fino é realizada de forma colaborativa pelo setor em geral, e não por uma única pessoa anotadora nem exclusivamente

pela EuDisseVocêDisse. Profissionais autorizados, pares de ONGs/NFPs e outros revisores colaboradores do setor classificam exemplos protegidos usando os mesmos rótulos da ontologia, pontuações e intervalos de evidência marcados. Suas classificações combinadas passam por análise estatística de taxas de concordância, distribuições de pontuação, sobreposição dos intervalos de evidência, valores atípicos e decisões contraditórias. Somente após essa análise estatística a equipe especializada da EuDisseVocêDisse revisa os resultados, investiga discordâncias, possíveis vieses e questões de qualidade dos dados e aprova, devolve ou escalona a resolução proposta. As submissões humanas são capturadas como snapshots imutáveis e independentes do modelo contendo texto de análise protegido, idioma, modalidade, fonte e versão da redação, rótulos, pontuações, explicações e intervalos de caracteres ou quadros. Os campos de remetente e destinatário são redigidos. Uma submissão individual é armazenada como candidata e não é elegível para treinamento. Somente a decisão vencedora resolvida por consenso de pares, adjudicação interna ou adjudicação SME é publicada no prefixo de treinamento de revisões resolvidas; candidatas não resolvidas ou contraditórias são excluídas da exportação normal. Uma nova análise direcionada aceita por um profissional também pode criar um snapshot elegível, que pode ser retirado se a aceitação for posteriormente revogada. As ferramentas de exportação verificam novamente os indicadores de resolução e elegibilidade antes de produzir linhas de treinamento específicas por tarefa. Esses controles mantêm PII não redigida e discordâncias não resolvidas fora dos dados comuns de ajuste fino, preservando um registro auditável para calibração e avaliação.

- 33. Como texto e transcrições de áudio são pontuados.** Após a redação, a mensagem alvo é avaliada de forma independente contra cada comportamento selecionado da ontologia ativa do idioma. Uma janela cronológica limitada pode ser fornecida para resolver referências, sequência e função comunicativa, mas somente a mensagem alvo marcada é pontuada. O classificador retorna uma probabilidade comportamental de 0,0 a 1,0, as substrings exatas de suporte e probabilidades por trecho. Os resultados só são mantidos quando atingem o limiar configurado daquele comportamento, que foi predeterminado a partir de dados de teste para maximizar a precisão e minimizar falsos positivos; comportamentos simultâneos podem compartilhar texto sobreposto. O áudio segue o mesmo caminho após transcrição fala-para-texto, atribuição do falante e detecção e correspondência de impressões vocais. Quando há vídeo, a detecção e correspondência facial com as trilhas de participantes identificadas pelo usuário apoia a atribuição dos participantes. A pontuação expressa a força da correspondência com a ontologia na comunicação; não é a probabilidade de uma alegação ser verdadeira, uma conclusão jurídica nem uma medida de frequência.
- 34. Como o vídeo é pontuado.** O modelo de vídeo revisa os quadros usando uma «multipass sliding window technique». Ele examina grupos cronológicos sobrepostos de quadros em múltiplas passagens, para que movimento e comportamento sejam considerados ao longo do tempo e não a partir de uma imagem estática isolada. Quando disponíveis, transcrições alinhadas e informações de participantes identificadas pelo usuário podem fornecer contexto. Um resultado mantido deve corresponder a um comportamento de vídeo configurado e identificar o intervalo de quadros inicial e final que lhe dá suporte; os resultados são convertidos em carimbos de tempo e vinculados aos quadros de evidência mostrados no relatório. A pontuação de vídeo armazenada como 1,0 registra que a correspondência configurada foi mantida pelo processo de revisão de vídeo; não deve ser lida como certeza calibrada de 100%. Revisores humanos anotam por intervalo de quadros, e os quadros de revisão externa têm rostos desfocados conforme descrito acima.
- 35. Sistemas de IA e onde são usados.** O fluxo de classificação da EuDisseVocêDisse.com usa uma pilha de modelos personalizada, guiada por ontologia e ajustada finamente para classificação de comunicações de texto/áudio e para classificação comportamental de quadros de vídeo. Serviços de apoio são usados apenas para tarefas específicas de processamento, como transcrições fala-para-texto atribuídas por falante, detecção e correspondência de impressões vocais, detecção e correspondência facial com trilhas de participantes identificadas pelo usuário, redação tipificada, suporte à estrutura documental, resumos de relatórios e interações, fluxos de embeddings e recuperação, apoio separado à detecção facial para revisão de privacidade e atributos padrão de toxicidade/profanidade quando essas categorias padrão são selecionadas. Modelos específicos de fornecedor podem mudar sem alterar a metodologia: o material fonte é preservado; texto e transcrições são redigidos antes da classificação ou

revisão humana; os quadros fornecidos a revisores externos designados têm rostos desfocados; as classificações seguem a ontologia ativa; e os exemplos vinculados à fonte permanecem revisáveis.

- 36. Declaração sobre IA generativa e marcador †.** O objetivo principal da IA generativa de resumo do relatório é tornar a análise numérica e os resultados analíticos já calculados mais legíveis para o tribunal. Ela é usada apenas em campos narrativos marcados e pode descrever contagens, proporções, tendências e material analítico existente vinculado à fonte dentro do escopo, mas não determina nem altera as classificações, pontuações, contagens, evidência selecionada ou gráficos subjacentes. Cada instância de conteúdo do relatório gerado por IA generativa é marcada com uma adaga (†) ao final do parágrafo gerado. A adaga identifica apenas a formulação explicativa; ela não condiciona, substitui nem cria os números, classificações, provas ou resultados subjacentes.
- 37. Para que a IA generativa de resumo do relatório não é usada.** O modelo de resumo do relatório não é usado para classificar ou pontuar comunicações ou vídeo, calcular contagens ou proporções, escolher evidência fonte, gerar gráficos, alterar arquivos fonte, decidir se uma alegação factual é verdadeira ou fornecer uma conclusão jurídica. Esses resultados analíticos são concluídos e armazenados antes do resumo narrativo do relatório. Em seguida, o PDF é gerado direta e deterministicamente a partir do material vinculado à fonte, dos resultados calculados, dos gráficos e dos campos narrativos claramente marcados; a IA generativa não gera nem diagrama o próprio PDF.
- 38. Integridade e auditabilidade da fonte.** O material fonte é protegido separadamente do processo analítico. Como etapa inicial de tratamento, a EuDisseVocêDisse.com cria um snapshot de evidência criptografado e selado em blockchain para que qualquer exclusão ou alteração posterior seja detectável antes de trabalho posterior de IA ou revisão humana. Arquivos são criptografados durante a transferência e em repouso, o material fonte é registrado usando hashes criptográficos e itens fonte são selados em uma cadeia de evidências de blockchain privada. Uploads, acessos, compartilhamentos e eventos de processamento recebem carimbo de data/hora e são associados à atividade de usuário verificada para apoiar a integridade da fonte e a revisão da cadeia de custódia. Relatórios e documentos enviados também estão sujeitos a verificações de integridade visual durante a importação, incluindo rasterização de páginas e análise de estrutura de pixels, projetadas para detectar renderizações de documentos malformadas, visualmente inconsistentes ou inesperadamente alteradas, o que pode indicar a apresentação de evidência alterada ou gerada.
- 39. Benchmarks de transcrição.** Benchmarks publicados de transcrição devem ser lidos separadamente da precisão de classificação comportamental da EuDisseVocêDisse.com. WER significa *word error rate* (taxa de erro de palavras): o número de substituições, exclusões e inserções de palavras dividido pelo número de palavras da transcrição de referência. Um WER menor é melhor. Para uma apresentação comparável como taxa de precisão, o equivalente simples de precisão de palavras é calculado como 100% menos o WER; trata-se de uma reformulação matemática do WER informado, não de um benchmark independente. A ElevenLabs informa precisão de transcrição em inglês de 96,7% para o Scribe v1, e sua documentação atual classifica inglês, francês, alemão, espanhol e português como idiomas de transcrição "Excellent" com WER de 5% ou menos, equivalente a uma precisão simples de palavras de 95% ou mais. A OpenAI informa o Whisper por benchmark, e não como uma taxa universal única: LibriSpeech Clean mede fala inglesa lida de forma clara, FLEURS English mede fala em inglês dentro de um benchmark multilíngue, e a média mais ampla de benchmark combina vários conjuntos públicos de dados de fala. O Whisper registrou 2,7% de WER no LibriSpeech Clean, equivalente a 97,3% de precisão simples de palavras; 4,4% de WER no FLEURS English, equivalente a 95,6%; e 12,8% de WER médio no conjunto de benchmarks informado no artigo do Whisper, equivalente a 87,2% de precisão simples de palavras.
- 40. Benchmark de ontologia personalizada.** O benchmark de ontologia personalizada da EuDisseVocêDisse.com é mantido contra a ontologia ativa do motor nos idiomas suportados. O benchmark usa exemplos fonte revisados e cenários de todos os comportamentos para medir detecções esperadas, não detecções esperadas, falsos positivos e falsos negativos. Os resultados do benchmark devem ser lidos como medidas de desempenho, não como determinações factuais conclusivas sobre qualquer comunicação individual neste relatório. As linhas de grupos abaixo incluem apenas os

comportamentos personalizados de texto/áudio e de vídeo selecionados para este relatório.

Métrica	Benchmark português atual	Significado
Acurácia	95,75%	A porcentagem de casos revisados do benchmark em que o sistema retornou o resultado esperado.
Falsos positivos	0,64%	A porcentagem de todos os casos revisados do benchmark em que o sistema detectou um comportamento que não era esperado.

Grupo de comportamentos	Comportamentos atuais do motor
Comportamento violento doméstico	Admissão de culpa; Comportamento coercitivo; Abuso financeiro; Avanços sexuais indesejados; Coagir ao autoferimento; Ameaça de autoferimento; Perseguição; Chantagem; Golpear; Chutar; Puxar cabelo; Forçar ao chão; Impedimento; Estrangulamento; Contato sexualizado; Obstrução; Mordendo, Arranhando ou Cuspidando; Empurrar, Acarretar, Derrubar ou Similar; Arremessar Itens; Ameaçando com Arma (armas de fogo, facas, etc.); Outras Alterações Físicas
Comportamento abusivo	Crítica parental; Desvalorização do corpo; Bullying; Alegações de abuso não verificadas; Gas-lighting; Gestos Abusivos; Arremessar líquidos ou substâncias; Ações destrutivas
Comportamento manipulativo	Ameaças legais; Indução de vergonha; Exploração de vulnerabilidade; Culpa manipulativa; Desculpa performativa; Atribuição de saúde mental como arma; Encurralar ou bloquear o movimento; Interferir com o telefone ou dispositivo
Comportamento antissocial	Recusa oficial de mediação; Expressão de ressentimento
Comunicação construtiva	Co-parenting proativo; Pedidos de mediação; Desescalada; Empatia; Construção de pontes; Estabelecimento de limites saudáveis
Reparo e reconciliação	Declarações de amor e devoção; Remorso; Buscando perdão; Concedendo perdão; Graça / Misericórdia
Sinais de alerta	Superioridade moral; Controle de Pontuação; Desvalorização; Intenção vingativa; Enquadramento julgador; Crítica (Instituto Gottman); Defensividade (Instituto Gottman); Desdém (Instituto Gottman); Bloqueio de comunicação (Instituto Gottman)
Táticas de pressão	Ultimato; Invalidação emocional; Endividamento emocional; Punição por retirada



Mr Adam Norman Jenkins, Director of Technology, ISaidUSaid Pty Ltd
--Fim do Relatório--