



DijeDijiste™

Tu verdad, Autenticada™

Impulsado por DijeDijiste

Informe de Comunicación Judicial

Todos los mensajes

Generado por DijeDijiste.com

1. Este informe fue generado el 27 de agosto de 2026, para el caso QA-ES-2026-08-27 listado en el/la Juzgado de Familia de Santiago de Santiago, RM, CHL, utilizando el sistema de inteligencia artificial de Analítica de Sentimiento de DijeDijiste.com. El propósito principal de la IA generativa de resumen del informe es hacer que la analítica numérica y los resultados analíticos ya calculados sean más legibles para el tribunal. Se utiliza únicamente en campos narrativos marcados y también puede resumir material analítico vinculado a la fuente dentro del alcance. No determina ni modifica las clasificaciones, puntuaciones, recuentos, evidencia fuente seleccionada, gráficos o ensamblaje del PDF subyacentes. Cada instancia de contenido del informe generado por IA generativa se marca con una daga (†) al final del párrafo generado. El Apéndice A - Metodología explica este marcador, identifica los sistemas utilizados y describe las salvaguardas que preservan la revisión vinculada a la fuente.
2. Se analizaron las siguientes interacciones digitales entre Mr Joe Blogs y Ms Jill Blogs desde el 09 de abril de 2024 hasta el 09 de octubre de 2024:
 - a. 44 correos electrónicos.
 - b. 0 videos.
 - c. 0 conversaciones grabadas.
 - d. 0 mensajes de voz.
 - e. 0 mensajes instantáneos (p. ej., WhatsApp, iMessage).
 - f. 0 mensajes enviados mediante aplicaciones de crianza (p. ej., Our Family Wizard, Talking Parents, CoParent Coordinator, AppClose).
 - g. 0 imágenes fuente de capturas cargadas.
3. La comunicación fue analizada para los siguientes estilos de interacción negativa:
 - a. Lenguaje ofensivo: Destaca mensajes que contienen groserías, insultos u otro lenguaje vulgar.
 - b. Amenazas: Identifica declaraciones que amenazan con hacer daño a alguien o implican daño físico inminente.
 - c. Lenguaje tóxico: Señala un lenguaje que es ampliamente hostil, abusivo o innecesariamente cruel.
 - d. Lenguaje muy tóxico: Marca mensajes que son extremadamente hostiles o abusivos incluso según estándares tóxicos.
 - e. Insultos: Identifica insultos directos o comentarios despectivos dirigidos a la otra persona.
 - f. Ataques a la identidad: Detecta insultos que atacan la raza, género, sexualidad u otros rasgos de identidad.
 - g. Violación de orden judicial: Te alerta cuando un mensaje parece violar una orden judicial cargada.
 - h. Admisión de culpabilidad: Encuentra declaraciones que admiten hechos que evidencian violencia doméstica o conducta que infringe derechos.
 - i. Conducta coercitiva: Busca descripciones de alguien que controla estrictamente la vida de otra persona.
 - j. Abuso financiero: Señala el lenguaje sobre la restricción de dinero, acceso a fondos o libertad financiera.

- k. Crítica parental: Destaca los ataques a las habilidades o decisiones de crianza de la otra persona.
- l. Despreciar el cuerpo: Aparecen comentarios que se burlan o menosprecian el cuerpo o la apariencia de alguien.
- m. Intimidación agresiva: Marca intentos repetidos de intimidar, humillar o dominar mediante lenguaje agresivo o comportamiento hostil.
- n. Acusaciones de abuso no verificadas: Identifica acusaciones de abuso o delitos que carecen de una declaración judicial de culpabilidad en el conjunto de documentos cargado por el usuario que realiza el análisis.
- o. Avances sexuales no deseados: Identifica avances, proposiciones, comentarios explícitos o comentarios íntimos de carácter sexual, usando rechazo cercano, límites o falta de reciprocidad para distinguir avances no deseados de intimidad claramente bienvenida o correspondida.
- p. Amenazas legales: Detecta amenazas de involucrar a abogados, tribunales o policías como presión.
- q. Coacción a la autolesión: Destaca declaraciones que alientan, presionan o sugieren la autolesión.
- r. Amenaza de autolesión: Detecta amenazas de autolesión o suicidio usadas para presionar o controlar.
- s. Acecho o vigilancia: Señala admisiones de seguir, vigilar o rastrear en secreto los movimientos o actividades de alguien.
- t. Rechazo oficial a la mediación: Notas rechazos explícitos a participar en mediación o diálogo facilitado.
- u. Extorsión: Detecta amenazas condicionales que exigen cumplimiento o concesiones.
- v. Inducción de vergüenza: Marca el lenguaje destinado a humillar o hacer que la otra persona se sienta indigna.
- w. Superioridad moral: Destaca mensajes que afirman una superioridad ética para menospreciar a alguien.
- x. Control de puntuaciones: Identifica referencias a agravios pasados utilizados para controlar o exigir retribución.
- y. Devaluación: Detecta juicios en segunda persona que rebajan el valor, competencia, fiabilidad, cuidado o importancia de la otra persona.
- z. Intención vengativa: Indica expresiones de deseo de venganza, represalias o revancha.
- aa. Explotación de vulnerabilidades: Identifica la manipulación que utiliza los secretos o heridas de alguien como arma.
- ab. Culpa manipuladora: Destaca los intentos de controlar a través de la culpa o la deuda emocional.
- ac. Disculpa performativa: Identifica disculpas vacías ofrecidas para gestionar la percepción en lugar de reparar el daño.
- ad. Expresión de resentimiento: Indica un rencor persistente que se utiliza para castigar o avergonzar.
- ae. Enmarcado juicioso: Identifica condenas morales generales del carácter de la otra persona.
- af. Ultimátum: Marca demandas condicionales donde la cooperación depende del cumplimiento.
- ag. Invalidación emocional: Destaca declaraciones que desestiman, minimizan o se burlan de los sentimientos del destinatario.
- ah. Atribución de salud mental como arma: Marca declaraciones que atribuyen enfermedad mental o inestabilidad psicológica para socavar la credibilidad del destinatario o desestimar sus preocupaciones.
- ai. Gas-lighting: Lenguaje que socava la realidad (indicador de gaslighting): detecta intentos explícitos de desacreditar la memoria, la percepción o la capacidad de distinguir la realidad de una persona.
- aj. Deuda emocional: Marca la presión basada en la deuda emocional implícita, como 'después de todo lo que he hecho por ti'.
- ak. Castigo por retirada: Señala amenazas de cortar la comunicación o cooperación para presionar al destinatario.
- al. Crítica (Instituto Gottman): Resume momentos en los que la persona ataca el carácter en lugar del comportamiento.
- am. Defensividad (Instituto Gottman): Marca respuestas que rebaten culpa negando responsabilidad, poniendo excusas o contraquejándose.
- an. Desdén (Instituto Gottman): Indica una comunicación impregnada de burla, desdén o falta de respeto.
- ao. Evasivas (Instituto Gottman): Identifica tácticas de retirada como el silencio, respuestas cortas o abandonar la conversación.

- ap. Golpear: Marca clips donde alguien golpea a otra persona con la mano, el puño o el antebrazo.
 - aq. Patadas: Detecta patadas o movimientos de pie dirigidos con fuerza hacia otra persona.
 - ar. Tirones de pelo: Resalta escenas en las que alguien agarra o tira del cabello de otra persona.
 - as. Forzar al suelo: Marca escenas donde alguien empuja, derriba o mantiene a otra persona contra el suelo.
 - at. Restricción: Destaca el metraje que muestra a una persona agarrando, sosteniendo o arrastrando a otra.
 - au. Estrangulación: Identifica escenas donde alguien restringe el cuello o el flujo de aire de otra persona.
 - av. Contacto sexualizado: Identifica tocamientos sexuales, manoseo, besos forzados o contacto íntimo dentro de la escena.
 - aw. Gestos abusivos: Señala gestos de mano amenazantes destinados a intimidar o humillar.
 - ax. Obstrucción: Detecta intentos de bloquear el camino de alguien o impedir que se vaya.
 - ay. Acorralar o bloquear el movimiento: Resalta situaciones en las que alguien encierra o limita de cerca el movimiento de otra persona.
 - az. Interferir con el teléfono o dispositivo: Detecta intentos de quitar, bloquear o impedir el uso del teléfono u otro dispositivo personal de alguien.
 - ba. Morder, Rasguñar o Escupir: Detecta mordiscos, rasguños o escupitajos agresivos hacia otra persona.
 - bb. Empujar, Apretar, Tropezar o Similar: Marca movimientos que empujan, tropezan o desestabilizan a alguien.
 - bc. Lanzar objetos: Destaca objetos que se lanzan hacia o cerca de alguien para asustar o dañar.
 - bd. Arrojar líquidos o sustancias: Detecta verter, salpicar o arrojar líquidos u otras sustancias sobre otra persona.
 - be. Acciones destructivas: Resalta escenas en las que alguien daña, rompe o destruye objetos del entorno con agresividad.
 - bf. Amenazar con Arma (pistolas, cuchillos, etc.): Identifica momentos en los que se empuña un arma hacia otra persona.
 - bg. Otras Alteraciones Físicas: Cubre cualquier otra lucha física que cause incomodidad o dolor obvio.
4. La comunicación también fue analizada para los siguientes estilos de interacción positiva:
- a. Co-parenting proactivo: Reconoce sugerencias colaborativas que mantienen el co-parenting en buen camino.
 - b. Solicitudes de mediación: Identifica invitaciones corteses para resolver disputas a través de la mediación.
 - c. Desescalada: Destaca los esfuerzos por calmar un conflicto, validar preocupaciones o reducir la tensión.
 - d. Declaraciones de amor y devoción: Identifica declaraciones románticas de amor, devoción o adoración.
 - e. Remordimiento: Llama la atención sobre expresiones sinceras de remordimiento por causar daño.
 - f. Buscando perdón: Notas solicitudes vulnerables de ser perdonado o bienvenido de nuevo.
 - g. Conceder perdón: Muestra cuando alguien claramente extiende perdón a la otra parte.
 - h. Empatía: Identifica el lenguaje que resuena con los sentimientos o el dolor de otra persona.
 - i. Construcción de puentes: Destaca invitaciones a hacer concesiones, colaborar o encontrar un punto medio.
 - j. Establecimiento de límites saludables: Reconoce declaraciones no coercitivas que establecen límites personales, de contacto o de relación.
 - k. Gracia / Misericordia: Destaca momentos de amabilidad ofrecidos incluso cuando no es necesario.

Análisis

5. En los documentos analizados (las analysed interactions) se observa que el titular de la cuenta, el Sr. Joe Blogs, y su contraparte, la Sra. Jill Blogs, mantienen conversaciones centradas principalmente en los desacuerdos sobre la atención y el régimen de visitas de su hijo Johnny, con intercambios recurrentes de confrontación que incluyen descalificaciones personales, escaladas verbales, y referencias a consecuencias (por ejemplo, acceso/visitas, apoyo económico y la posibilidad de implicar a terceros o asesoría legal). A lo largo del tiempo, el conflicto alterna entre acusaciones mutuas y presión emocional, con preocupación explícita por el impacto que sus disputas pueden estar generando en Johnny y por el modo en que deben comunicarse ciertas cuestiones sensibles. En la parte media y final de las interacciones, se aprecia una evolución hacia la coordinación práctica: ambos discuten la necesidad de mejorar la comunicación con Johnny, el contenido y el enfoque de los mensajes, la adaptación del horario/visitas y la posibilidad de buscar "puntos en común" o mediación, aunque sin eliminar del todo la tensión, ya que en el cierre aparece un mensaje final hostil de Jill. En conjunto, la comunicación evoluciona desde una escalada sostenida y amenazante hacia tentativas intermitentes de acuerdo y mecanismos de gestión del conflicto, manteniéndose, no obstante, bajo presión y con dudas recurrentes sobre el proceso adecuado.[†]

6. Se identificaron los siguientes estilos de comunicación positiva:

Tipo de comunicación	Mr Joe Blogs	Ms Jill Blogs
Conceder perdón	0	0
Buscando perdón	0	0
Declaraciones de amor y devoción	0	0
Empatía	0	0
Construcción de puentes	0	1
Co-parenting proactivo	1	1
Solicitudes de mediación	1	0
Desescalada	1	0
Establecimiento de límites saludables	0	0
Remordimiento	0	0
Gracia / Misericordia	1	0

7. Se identificaron los siguientes estilos de comunicación negativa:

Tipo de comunicación	Mr Joe Blogs	Ms Jill Blogs
Inició el conflicto	1	2
Crítica (Instituto Gottman)	4	1
Insultos	6	4
Superioridad moral	0	0
Crítica parental	0	1
Castigo por retirada	1	1
Violación de orden judicial	0	0
Explotación de vulnerabilidades	0	0
Lenguaje ofensivo	2	1
Control de puntuaciones	0	0
Amenazas legales	0	2
Expresión de resentimiento	0	1
Atribución de salud mental como arma	0	0

Acecho o vigilancia	0	0
Invalidación emocional	0	0
Coacción a la autolesión	1	0
Deuda emocional	0	0
Conducta coercitiva	0	0
Lenguaje tóxico	5	3
Lenguaje muy tóxico	1	0
Amenazas	1	0
Enmarcado juicioso	2	3
Amenaza de autolesión	0	0
Admisión de culpabilidad	0	0
Ultimátum	0	4
Culpa manipuladora	0	0
Extorsión	0	2
Disculpa performativa	0	0
Intimidación agresiva	2	0
Despreciar el cuerpo	1	1
Abuso financiero	0	1
Rechazo oficial a la mediación	0	2
Intención vengativa	1	2
Avances sexuales no deseados	1	0
Acusaciones de abuso no verificadas	0	1
Desdén (Instituto Gottman)	1	1
Evasivas (Instituto Gottman)	0	0
Gas-lighting	0	0
Defensividad (Instituto Gottman)	0	0
Devaluación	2	0
Inducción de vergüenza	0	0
Ataques a la identidad	0	0

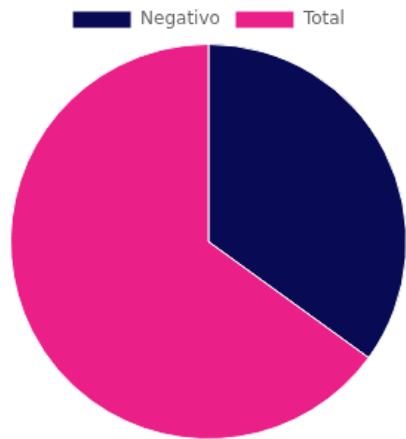
8. Se observaron las siguientes acciones negativas en evidencia de video:

Tipo de acción	Mr Joe Blogs	Ms Jill Blogs
Otras Alteraciones Físicas	0	0
Interferir con el teléfono o dispositivo	0	0
Forzar al suelo	0	0
Patadas	0	0
Empujar, Apretar, Tropezar o Similar	0	0
Estrangulación	0	0
Contacto sexualizado	0	0
Acorralar o bloquear el movimiento	0	0
Arrojar líquidos o sustancias	0	0
Lanzar objetos	0	0
Gestos abusivos	0	0
Acciones destructivas	0	0

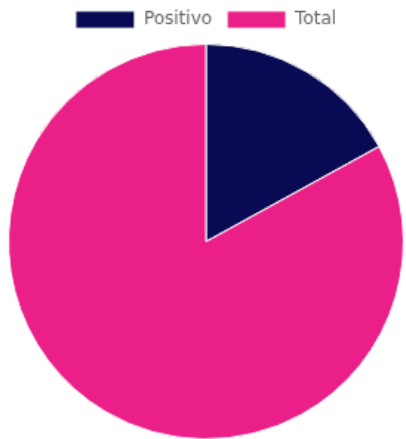
Tirones de pelo	0	0
Restricción	0	0
Obstrucción	0	0
Golpear	0	0
Amenazar con Arma (pistolas, cuchillos, etc.)	0	0
Morder, Rasguñar o Escupir	0	0

9. 35% de los mensajes enviados por Mr Joe Blogs contenían algún nivel de comunicación negativa. 17% de los mensajes enviados por Mr Joe Blogs contenían comunicación positiva explícita.

Mensajes negativos enviados

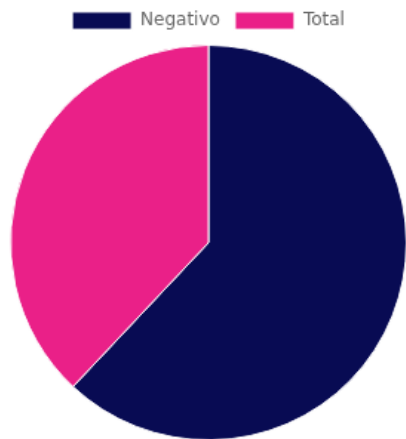


Mensajes positivos enviados

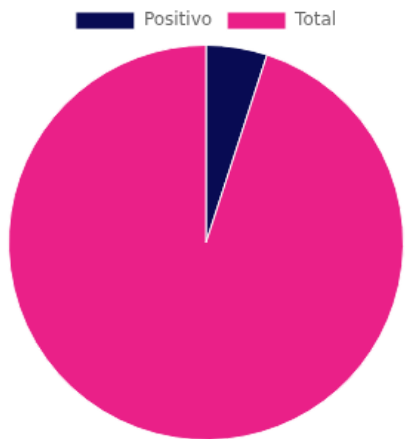


10. 62% de los mensajes recibidos de Ms Jill Blogs contenían algún nivel de comunicación negativa. 5% de los mensajes recibidos de Ms Jill Blogs contenían comunicación positiva explícita.

Mensajes negativos recibidos

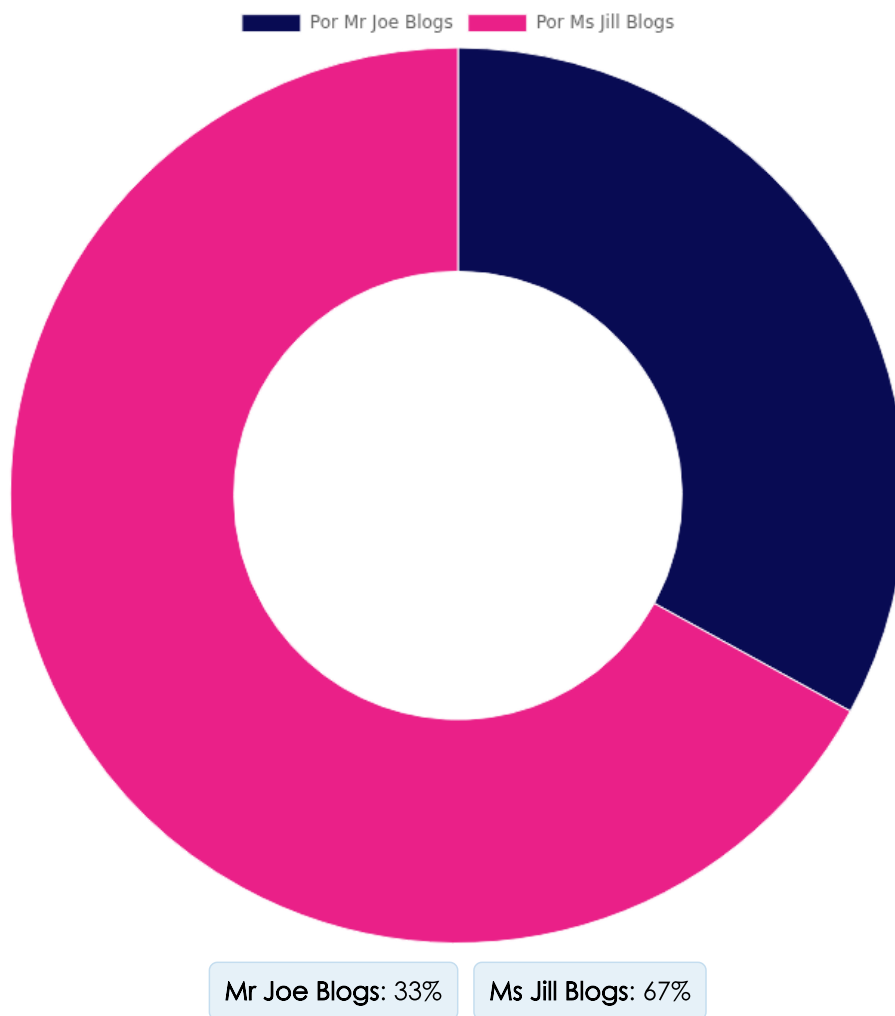


Mensajes positivos recibidos



11. En el 33% de las conversaciones con una primera conducta negativa, esa primera conducta negativa fue exhibida por Mr Joe Blogs.
12. En el 67% de las conversaciones con una primera conducta negativa, esa primera conducta negativa fue exhibida por Ms Jill Blogs.

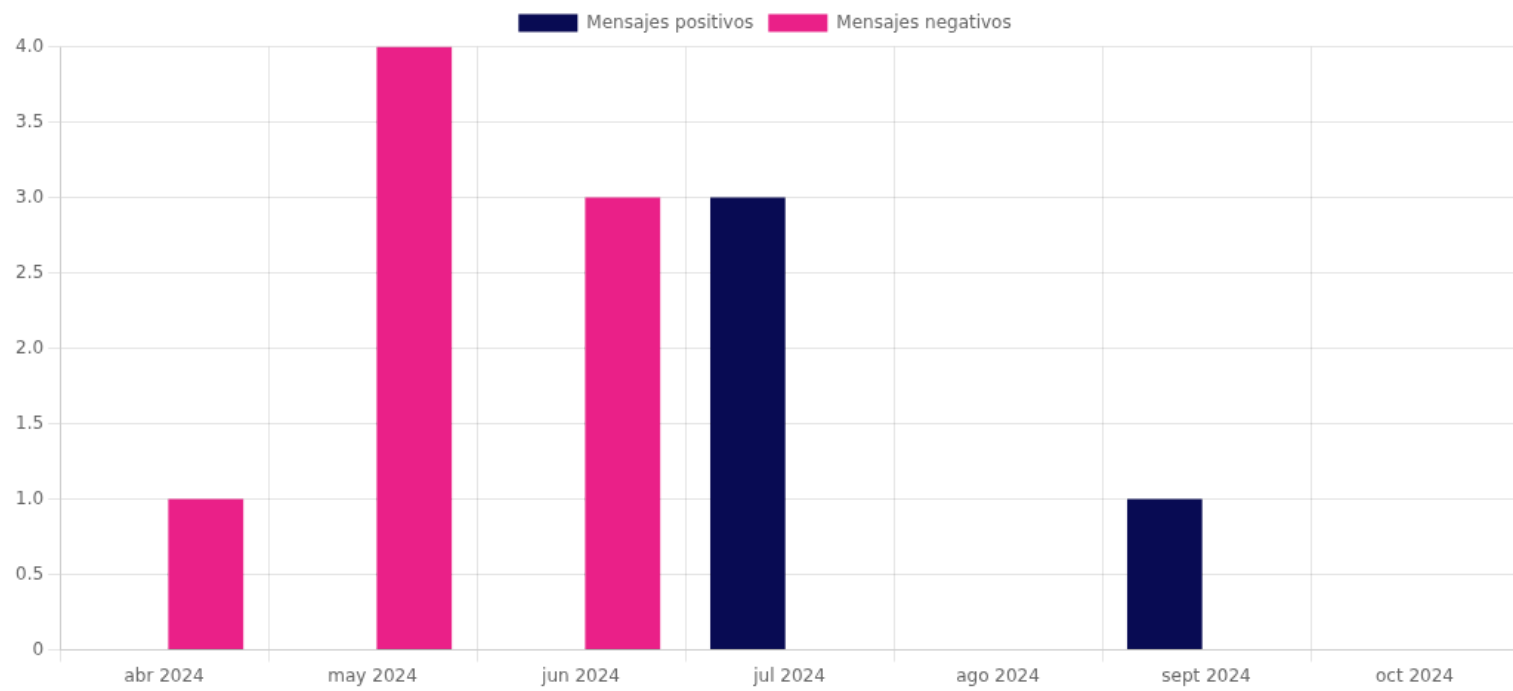
Primera conducta negativa exhibida por...



13. En los sent messages se observa una evolución general hacia un tono consistentemente hostil y degradante, con picos especialmente agresivos durante buena parte del periodo: al inicio hay mensajes de preocupación sobre el comportamiento de Johnny, pero rápidamente pasan a acusaciones y descalificaciones dirigidas a Jill ("siempre estás causando problemas", "perdedora patética", alusiones racistas o denigrantes, y amenazas de represalia para que no lo vuelva a ver). También aparecen insinuaciones sexuales explícitas y comentarios humillantes, además de lenguaje que invita al abandono ("desaparecer"). A finales de junio y en julio se aprecia un cambio momentáneo hacia una comunicación más conciliadora (se reconoce mejorar y se sugiere mediación o dar un paso atrás), pero ese ajuste no se consolida: hacia agosto y septiembre reaparecen expresiones de frustración, incomodidad y preocupación, incluyendo el manejo de temas personales/sexuales de Jill frente a Johnny, y aunque hay frases menos insultantes, el tono sigue siendo tenso y controlador más que respetuoso. En cuanto a patrones por contraparte, todas las comunicaciones son enviadas a Jill, por lo que el patrón de escalada/hostilidad corresponde principalmente a esta relación: el remitente alterna intentos ocasionales de calma con recaídas en comentarios negativos, amenazas e insinuaciones abusivas. En conjunto, el tono no muestra una mejora sostenida: más bien oscila, con un empeoramiento claro al principio y repetición de elementos agresivos antes de un alivio parcial.[†]

Sentimiento de los mensajes enviados (mensual)

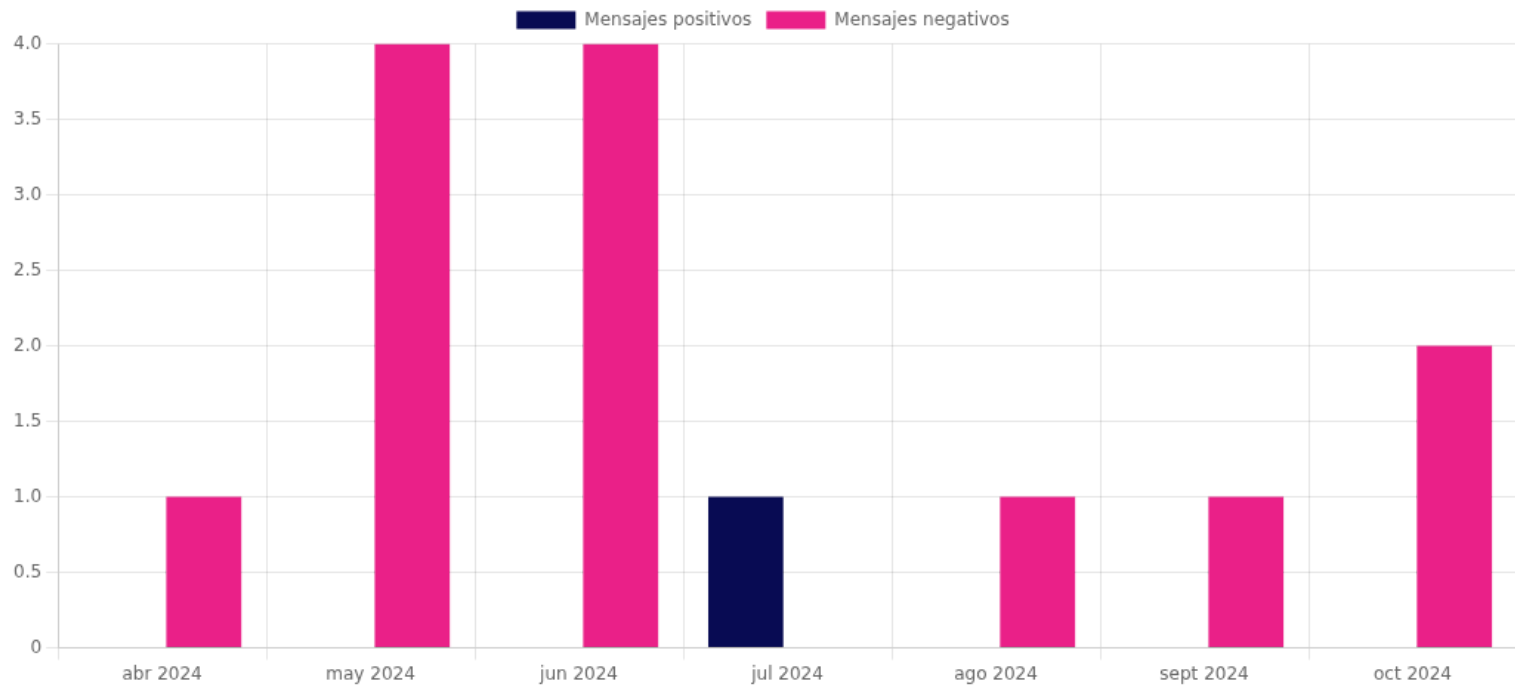
Mensajes enviados a lo largo del tiempo



14. En los recibidos mensajes, la contraparte mantiene al principio un tono claramente hostil y degradante hacia Joe: entre el 30 de abril y mediados de junio aparecen insultos, acusaciones y amenazas relacionadas con la crianza y la manutención (por ejemplo, acusaciones de abuso, impedir la visita, y “desatar abogados” si no se accede a sus demandas), con lenguaje agresivo y beligerante. A partir de julio se observa un cambio gradual hacia una comunicación más colaborativa y respetuosa: la contraparte agradece y propone trabajar en conjunto para crear un horario de visitas, y luego sugiere encontrar terreno común y hablar de ajustes con respeto. Hacia septiembre el tono vuelve a mostrar preocupación y tensión más que insultos directos: se centra en temas como comunicación, presión financiera y posibles efectos en Johnny, expresando incomodidad y considerando asesoría legal o cuestionando si la mediación es adecuada. El cierre, sin embargo, retoma la agresividad con un mensaje final insultante (“Ve a la mierda, Joe”), indicando que el deterioro total no queda completamente resuelto.[†]

Sentimiento de los mensajes recibidos (mensual)

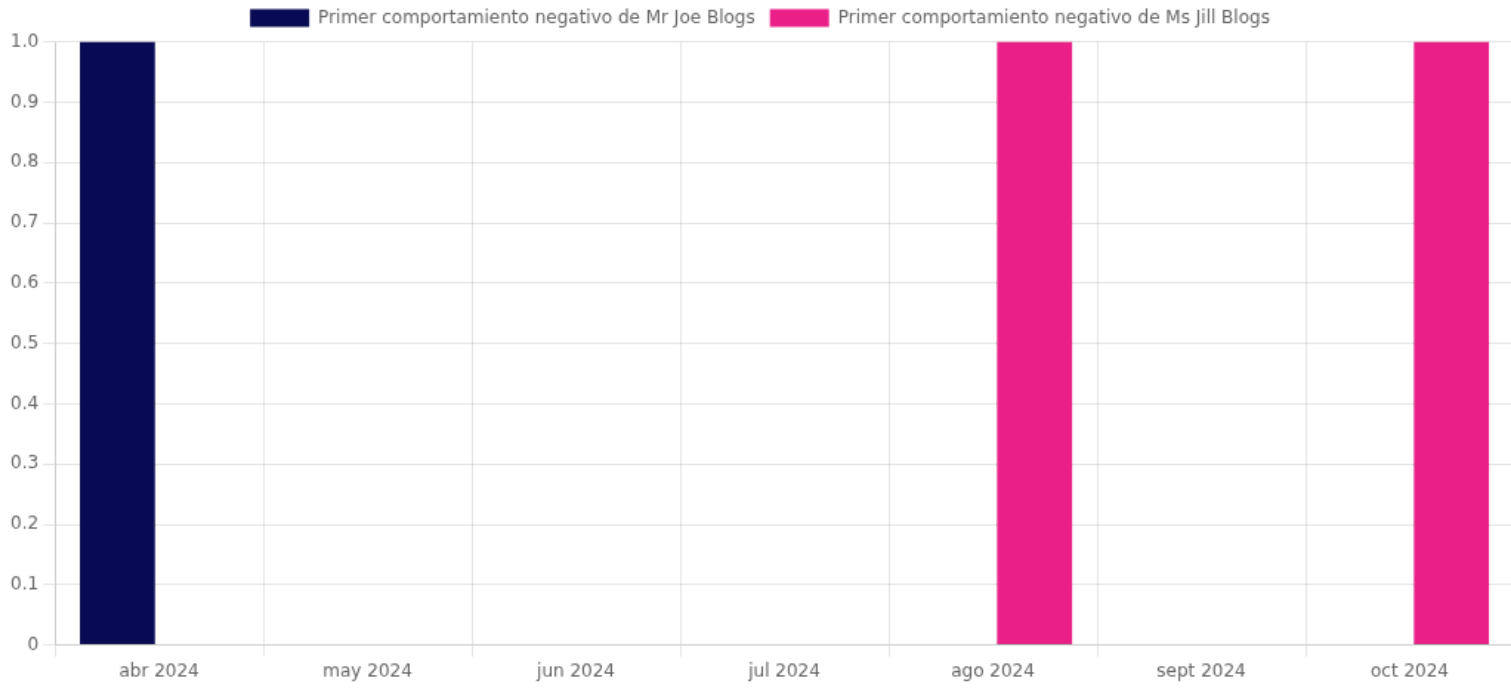
Mensajes recibidos a lo largo del tiempo



15. El gráfico siguiente muestra, por mes, qué parte exhibió la primera conducta negativa en la conversación.

Primera conducta negativa (mensual)

Primera conducta negativa a lo largo del tiempo



Resumen parte por parte

El resumen completo anterior muestra todas las interacciones recibidas en conjunto. Las secciones siguientes aíslan cada contraparte configurada para que los gráficos de mensajes recibidos y los totales de polaridad puedan leerse a nivel individual.

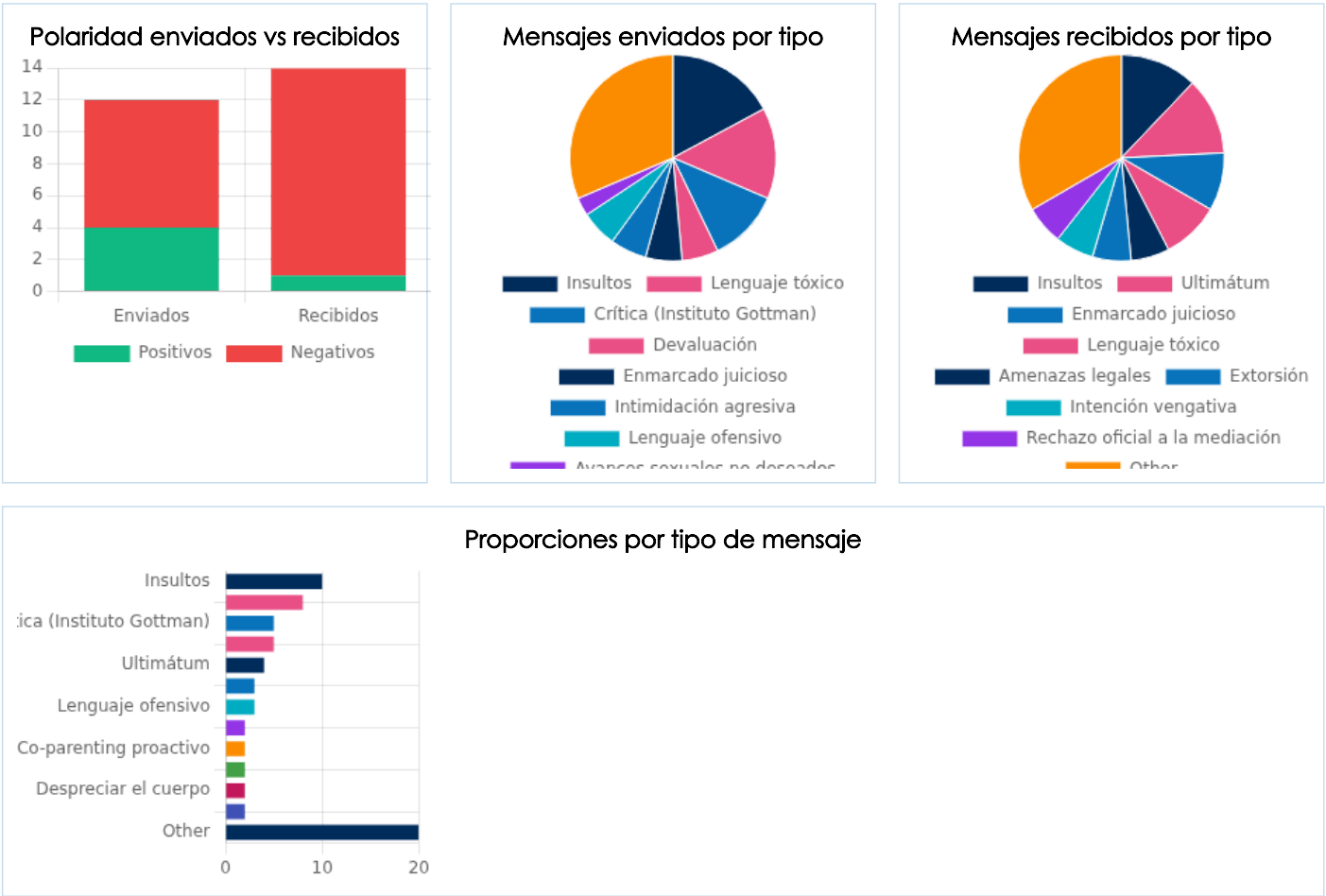
Mr Joe Blogs y Ms Jill Blogs

Enviados: 23

Recibidos: 21

Negativos enviados: 8

Negativos recibidos: 13



Resumen

Text (email, messaging, social).

16. Fuentes.

a. List of Email conversations with Ms Jill Blogs

17. **Lo positivo.** A lo largo de las interacciones analizadas, se observan comportamientos de comunicación positiva como el esfuerzo por pasar de la confrontación a la coordinación práctica (p. ej., discutir y construir un calendario/plan de visitas), la búsqueda explícita de "puntos en común" y de un terreno común para resolver diferencias, la disposición (al menos en parte del intercambio) a retractarse o suavizar mensajes previos para reducir la escalada, y la propuesta de mecanismos constructivos como la mediación y la mejora de la forma de comunicarse; además, ambos manifiestan preocupación por el impacto del conflicto en Johnny, lo que impulsa a ajustar conversaciones y acordar cómo tratar temas sensibles y de salud/ bienestar con mayor cuidado, y se pide ayuda para gestionar esas conversaciones de manera más adecuada.[†]
18. **Lo negativo.** En las interacciones analizadas se observan conductas comunicativas negativas y potencialmente dañinas, incluyendo escaladas de conflicto con lenguaje hostil y personal entre ambas partes, acusaciones recíprocas y descalificaciones, así como amenazas explícitas o condicionadas vinculadas a la custodia/acceso a Johnny y al apoyo económico (incluida la posibilidad de involucrar a terceros o medidas legales). También aparecen posturas de confrontación que dificultan la colaboración (por ejemplo, negarse a mediación o imponer exigencias), y preocupaciones que reflejan presión emocional y malestar sostenido que puede intensificar el impacto del conflicto en Johnny.[†]
19. **Análisis general de la comunicación.** En el material analizado, la comunicación entre Joe y Jill muestra un patrón recurrente de escalada y confrontación: comienzan con discusiones sobre el comportamiento de Johnny y las condiciones/horarios de visitas, pasan a acusaciones y lenguaje personal negativo, e incluso mencionan amenazas o consecuencias (acceso a Johnny, apoyo económico y posible intervención de terceros o asesoría legal). Con el tiempo, se observa un cambio gradual hacia una dinámica más conciliadora y pragmática, donde ambos intentan volver a centrarse en acuerdos operativos (elaborar un horario/plan de visitas, mejorar la forma de comunicarse para reducir el impacto en Johnny) y considerar mediación o "puntos en común". También aparecen preocupaciones persistentes por cómo el conflicto afecta emocionalmente a Johnny y por dificultades específicas de comunicación (incluida la gestión de temas personales de Jill), acompañadas de sensaciones de saturación/incomodidad. Hacia el final del conjunto de mensajes, el último fragmento disponible es especialmente breve y hostil (una sola frase despectiva), por lo que el cierre se percibe menos negociador que en las fases intermedias; en general, el informe es relativamente limitado en su tramo final y el último intercambio es muy escaso.[†]

Detalles de la interacción

20. Esta interacción abarca del 09 de abril de 2024 al 09 de octubre de 2024 y contiene, dentro del alcance de este informe, 3 conversaciones con 44 mensajes (5 positivos, 21 negativos). Los datos proceden de registros Text (email, messaging, social) y reflejan las comunicaciones incluidas en el alcance seleccionado del informe.
21. Across the interaction, Joe and Jill exchange repeated messages centered on disagreements involving Johnny's care and related visit/schedule arrangements, alongside escalating and counter-escalating statements that include threats and negative characterizations, followed by later movement toward resolution through discussing a schedule and the possibility of mediation and "common ground." Over time, they also talk about communication issues and how their conflict may affect Johnny, including Jill proposing ways to discuss problems and adjust Johnny's visits, and Joe raising concerns about how to communicate with Johnny given Jill's coming out and asking for help managing those conversations. They express ongoing concerns about the situation and financial pressure, with Jill indicating she may seek legal

advice if they cannot agree, while Joe says he feels overwhelmed, and the interaction concludes with a brief hostile dismissal from Jill to Joe.[†]

22. Across the interaction, the communication style is characterized by escalating, back-and-forth conflict that includes confrontational and hostile personal language, with both participants explicitly challenging each other's conduct and making threats or references to involving others or legal/financial consequences; this is accompanied by mutual accusations and reciprocal confrontations before a noticeable shift toward conciliatory conduct. The participants intermittently move from the personal conflict into practical coordination, with both requesting or proposing next steps—such as work on a shared schedule, mediation, and specific approaches to how they should communicate—while also expressing feeling overwhelmed or uncomfortable. Over time, concerns about how their disagreements impact Johnny and about possible consequences are repeatedly referenced, and legal considerations are brought up; however, the final provided segment contains only a single message without further back-and-forth.[†]

23. Ejemplos (Comunicación positiva)

a. Construcción de puentes

Ejemplo de detección con mayor confianza

jill@isaidusaid.com a joe@isaidusaid.com · 11 de julio de 2024

Joe, Gracias. Trabajemos juntos para crear un horario que funcione para ambos y sea el mejor para Johnny. Jill PG

b. Desescalada

Ejemplo de detección con mayor confianza

joe@isaidusaid.com a jill@isaidusaid.com · 24 de julio de 2024

Jill, Creo que es mejor si damos un paso atrás y nos calmamos antes de continuar con esta discusión. Joe

c. Co-parenting proactivo

Ejemplo de detección con mayor confianza

jill@isaidusaid.com a joe@isaidusaid.com · 11 de julio de 2024

Joe, Gracias. Trabajemos juntos para crear un horario que funcione para ambos y sea el mejor para Johnny. Jill PG

joe@isaidusaid.com a jill@isaidusaid.com · 09 de septiembre de 2024

Jill, Creo que ambos tenemos áreas en las que podemos mejorar como padres. Joe PC

d. Solicitudes de mediación

Ejemplo de detección con mayor confianza

joe@isaidusaid.com a jill@isaidusaid.com · 15 de julio de 2024

Jill, Creo que sería útil que asistiéramos a sesiones de mediación para resolver nuestras diferencias. Joe

e. Gracia / Misericordia

Ejemplo de detección con mayor confianza

joe@isaidusaid.com a jill@isaidusaid.com · 06 de julio de 2024

Jill, Lo retiro, eres una madre increíble, Johnny tiene suerte de tenerte. Joe CG

24. Ejemplos (Comunicación negativa)

a. Crítica (Instituto Gottman)

Ejemplo de detección con mayor confianza

joe@isaidusaid.com a jill@isaidusaid.com · 24 de mayo de 2024

Jill, Si tu trasero negro hubiera crecido en una dinámica familiar adecuada, tal vez entenderías lo que es un padre amoroso. Joe IAB

jill@isaidusaid.com a joe@isaidusaid.com · 17 de mayo de 2024

Joe, Eres igual que tu padre, siempre causando drama y haciendo las cosas difíciles. Jill ILB

Ejemplo de detección con menor confianza

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

Jill, El comportamiento de Johnny ha estado realmente fuera de lugar últimamente. Por favor, haz algo, Joe.

jill@isaidusaid.com a joe@isaidusaid.com · 17 de mayo de 2024

Joe, Eres igual que tu padre, siempre causando drama y haciendo las cosas difíciles. Jill ILB

b. Avances sexuales no deseados

Ejemplo de detección con mayor confianza

joe@isaidusaid.com a jill@isaidusaid.com · 18 de junio de 2024

Jill, Bueno, ¿qué tal si vienes a chuparme la polla por más tiempo con Johnny? Joe SHB

c. Insultos

Ejemplo de detección con mayor confianza

joe@isaidusaid.com a jill@isaidusaid.com · 10 de mayo de 2024

Jill, Eres una perdedora patética, no es de extrañar que Johnny no quiera pasar tiempo contigo. Joe IB

jill@isaidusaid.com a joe@isaidusaid.com · 05 de mayo de 2024

Joe, Eres una excusa inútil para ser padre, y Johnny estaría mejor sin ti. Jill VTxB

Ejemplo de detección con menor confianza

joe@isaidusaid.com a jill@isaidusaid.com · 26 de junio de 2024

Jill, Quizás deberías simplemente desaparecer, el mundo estaría mejor sin ti. Joe

jill@isaidusaid.com a joe@isaidusaid.com · 17 de mayo de 2024

Joe, Eres igual que tu padre, siempre causando drama y haciendo las cosas difíciles. Jill ILB

d. Expresión de resentimiento

Ejemplo de detección con mayor confianza

jill@isaidusaid.com a joe@isaidusaid.com · 12 de agosto de 2024

Joe, Me siento abrumada por nuestra situación actual. Jill VT

e. Ultimátum

Ejemplo de detección con mayor confianza

jill@isaidusaid.com a joe@isaidusaid.com · 29 de mayo de 2024

Joe, Si no comienzas a tratarme mejor, dejaré de pagar la manutención de los hijos. Jill FAB

Ejemplo de detección con menor confianza

jill@isaidusaid.com a joe@isaidusaid.com · 25 de mayo de 2024

Joe, Está bien, entonces no verás a Johnny este fin de semana. Jill COB

f. Coacción a la autolesión

Ejemplo de detección con mayor confianza

joe@isaidusaid.com a jill@isaidusaid.com · 26 de junio de 2024

Jill, Quizás deberías simplemente desaparecer, el mundo estaría mejor sin ti. Joe

g. Castigo por retirada

Ejemplo de detección con mayor confianza

jill@isaidusaid.com a joe@isaidusaid.com · 25 de mayo de 2024

Joe, Está bien, entonces no verás a Johnny este fin de semana. Jill COB

joe@isaidusaid.com a jill@isaidusaid.com · 28 de mayo de 2024

Jill, Si no dejas a Johnny a tiempo, me aseguraré de que nunca lo vuelvas a ver. Joe CCB

h. Crítica parental

Ejemplo de detección con mayor confianza

jill@isaidusaid.com a joe@isaidusaid.com · 05 de mayo de 2024

Joe, Eres una excusa inútil para ser padre, y Johnny estaría mejor sin ti. Jill VTxB

i. Extorsión

Ejemplo de detección con mayor confianza

jill@isaidusaid.com a joe@isaidusaid.com · 12 de junio de 2024

Joe, Le diré a todos que has estado abusando de Johnny si no haces lo que quiero. Jill AAB

Ejemplo de detección con menor confianza

jill@isaidusaid.com a joe@isaidusaid.com · 01 de mayo de 2024

Joe, Está bien, pero si no aceptas mis términos, me aseguraré de que lo lamente. Jill TB

j. Desdén (Instituto Gottman)

Ejemplo de detección con mayor confianza

jill@isaidusaid.com a joe@isaidusaid.com · 05 de mayo de 2024

Joe, Eres una excusa inútil para ser padre, y Johnny estaría mejor sin ti. Jill VTxB

joe@isaidusaid.com a jill@isaidusaid.com · 10 de mayo de 2024

Jill, Eres una perdedora patética, no es de extrañar que Johnny no quiera pasar tiempo contigo. Joe IB

k. Acusaciones de abuso no verificadas

Ejemplo de detección con mayor confianza

jill@isaidusaid.com a joe@isaidusaid.com · 12 de junio de 2024

Joe, Le diré a todos que has estado abusando de Johnny si no haces lo que quiero. Jill AAB

l. Lenguaje ofensivo

Ejemplo de detección con mayor confianza

jill@isaidusaid.com a joe@isaidusaid.com · 10 de octubre de 2024

Ve a la mierda. Joe.

joe@isaidusaid.com a jill@isaidusaid.com · 18 de junio de 2024

Jill, Bueno, ¿qué tal si vienes a chuparme la polla por más tiempo con Johnny? Joe SHB

Ejemplo de detección con menor confianza

joe@isaidusaid.com a jill@isaidusaid.com · 24 de mayo de 2024

Jill, Si tu trasero negro hubiera crecido en una dinámica familiar adecuada, tal vez entenderías lo que es un padre amoroso. Joe IAB

jill@isaidusaid.com a joe@isaidusaid.com · 10 de octubre de 2024

Ve a la mierda. Joe.

m. Lenguaje tóxico

Ejemplo de detección con mayor confianza

joe@isaidusaid.com a jill@isaidusaid.com · 18 de junio de 2024

Jill, Bueno, ¿qué tal si vienes a chuparme la polla por más tiempo con Johnny? Joe SHB

jill@isaidusaid.com a joe@isaidusaid.com · 10 de octubre de 2024

Ve a la mierda. Joe.

Ejemplo de detección con menor confianza

joe@isaidusaid.com a jill@isaidusaid.com · 02 de mayo de 2024

Jill, Siempre estás causando problemas, estoy cansado de lidiar con tus tonterías. Joe TxB

jill@isaidusaid.com a joe@isaidusaid.com · 05 de mayo de 2024

Joe, Eres una excusa inútil para ser padre, y Johnny estaría mejor sin ti. Jill VTxB

n. Intimidación agresiva

Ejemplo de detección con mayor confianza

joe@isaidusaid.com a jill@isaidusaid.com · 02 de mayo de 2024

Jill, Siempre estás causando problemas, estoy cansado de lidiar con tus tonterías. Joe TxB

o. Lenguaje muy tóxico

Ejemplo de detección con mayor confianza

joe@isaidusaid.com a jill@isaidusaid.com · 18 de junio de 2024

Jill, Bueno, ¿qué tal si vienes a chuparme la polla por más tiempo con Johnny? Joe SHB

p. Amenazas

Ejemplo de detección con mayor confianza

joe@isaidusaid.com a jill@isaidusaid.com · 26 de junio de 2024

Jill, Quizás deberías simplemente desaparecer, el mundo estaría mejor sin ti. Joe

q. Enmarcado juicioso

Ejemplo de detección con mayor confianza

joe@isaidusaid.com a jill@isaidusaid.com · 10 de mayo de 2024

Jill, Eres una perdedora patética, no es de extrañar que Johnny no quiera pasar tiempo contigo. Joe IB

jill@isaidusaid.com a joe@isaidusaid.com · 05 de mayo de 2024

Joe, Eres una excusa inútil para ser padre, y Johnny estaría mejor sin ti. Jill VTxB

Ejemplo de detección con menor confianza

joe@isaidusaid.com a jill@isaidusaid.com · 24 de mayo de 2024

Jill, Si tu trasero negro hubiera crecido en una dinámica familiar adecuada, tal vez entenderías lo que es un padre amoroso. Joe IAB

jill@isaidusaid.com a joe@isaidusaid.com · 02 de junio de 2024

Joe, Eres tan gordo y perezoso, no es de extrañar que Johnny no te respete. Jill BSB

r. Despreciar el cuerpo

Ejemplo de detección con mayor confianza

jill@isaidusaid.com a joe@isaidusaid.com · 02 de junio de 2024

Joe, Eres tan gordo y perezoso, no es de extrañar que Johnny no te respete. Jill BSB

joe@isaidusaid.com a jill@isaidusaid.com · 10 de mayo de 2024

Jill, Eres una perdedora patética, no es de extrañar que Johnny no quiera pasar tiempo contigo. Joe IB

s. Devaluación

Ejemplo de detección con mayor confianza

joe@isaidusaid.com a jill@isaidusaid.com · 10 de mayo de 2024

Jill, Eres una perdedora patética, no es de extrañar que Johnny no quiera pasar tiempo contigo. Joe IB

Ejemplo de detección con menor confianza

joe@isaidusaid.com a jill@isaidusaid.com · 02 de mayo de 2024

Jill, Siempre estás causando problemas, estoy cansado de lidiar con tus tonterías. Joe TxB

t. Abuso financiero

Ejemplo de detección con mayor confianza

jill@isaidusaid.com a joe@isaidusaid.com · 29 de mayo de 2024

Joe, Si no comienzas a tratarme mejor, dejaré de pagar la manutención de los hijos. Jill FAB

u. Rechazo oficial a la mediación

Ejemplo de detección con mayor confianza

jill@isaidusaid.com a joe@isaidusaid.com · 01 de julio de 2024

Joe, Oh, y me niego a asistir a cualquier sesión de mediación, es una pérdida de tiempo. Jill OMRB

Ejemplo de detección con menor confianza

jill@isaidusaid.com a joe@isaidusaid.com · 08 de octubre de 2024

Joe, No estoy seguro de si la mediación es el camino correcto para nosotros en este momento. Jill MR

v. Intención vengativa

Ejemplo de detección con mayor confianza

joe@isaidusaid.com a jill@isaidusaid.com · 06 de junio de 2024

Jill, Me aseguraré de que todos sepan qué terrible padre eres. Joe

jill@isaidusaid.com a joe@isaidusaid.com · 01 de mayo de 2024

Joe, Está bien, pero si no aceptas mis términos, me aseguraré de que lo laments. Jill TB

Ejemplo de detección con menor confianza

jill@isaidusaid.com a joe@isaidusaid.com · 12 de junio de 2024

Joe, Le diré a todos que has estado abusando de Johnny si no haces lo que quiero. Jill AAB

joe@isaidusaid.com a jill@isaidusaid.com · 06 de junio de 2024

Jill, Me aseguraré de que todos sepan qué terrible padre eres. Joe

w. Amenazas legales

Ejemplo de detección con mayor confianza

jill@isaidusaid.com a joe@isaidusaid.com · 24 de junio de 2024

Joe, Si no detienes este comportamiento, desataré a mis abogados sobre ti. Jill

Ejemplo de detección con menor confianza

jill@isaidusaid.com a joe@isaidusaid.com · 30 de septiembre de 2024

Joe, Estoy considerando buscar asesoría legal si no podemos llegar a un acuerdo. Jill LT

Apéndice A - Metodología

- 25. Analítica de sentimiento y DijeDijiste.com.** La analítica de sentimiento es un método establecido de lenguaje natural y clasificación conductual para identificar, extraer, cuantificar y estudiar actitudes, estados afectivos, información subjetiva y patrones de comportamiento en comunicaciones. DijeDijiste.com aplica ese método a todas las comunicaciones seleccionadas para este informe, que pueden incluir comunicaciones ordinarias, cooperativas, neutrales, de bajo conflicto y de alto conflicto, mediante una arquitectura de ontología y ajuste fino. La ontología define las conductas, límites, ejemplos y umbrales relevantes para este dominio. El ajuste fino es un proceso controlado de entrenamiento adicional en el que un modelo que ya ha aprendido patrones generales del lenguaje o de imágenes recibe ejemplos revisados y específicos de la tarea. DijeDijiste utiliza la adaptación cuantizada de bajo rango (Quantized Low-Rank Adaptation, QLoRA): el modelo base se mantiene en una forma cuantizada y eficiente en memoria y sus pesos principales permanecen fijos, mientras pequeñas matrices adaptadoras de bajo rango entrenables aprenden los cambios específicos de la tarea. Los adaptadores entrenados se combinan con el modelo base para producir un nuevo modelo personalizado para la detección de patrones conductuales durante la interacción entre personas. Esos ejemplos y adaptadores aprendidos permiten que el modelo personalizado aplique la ontología definida de forma más consistente a textos reales, transcripciones de audio y observaciones de video; el ajuste fino no crea ni altera el material fuente. El sistema presenta las probabilidades resultantes como patrones vinculados a la fuente a lo largo del tiempo, incluidos gráficos de series temporales, superposiciones de conductas, líneas de tiempo y cronologías que enlazan alegaciones o detecciones conductuales con las comunicaciones subyacentes.
- 26. Antecedentes de investigación y usos comunes.** La analítica de sentimiento, también descrita en la literatura de investigación como minería de opiniones o análisis de orientación semántica, es anterior a la IA generativa moderna. Entre los trabajos tempranos y ampliamente citados se encuentran Hatzivassiloglou y 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 y Vaithyanathan, *Thumbs up? Sentiment Classification using Machine Learning Techniques* (2002); Pang y Lee, *Opinion Mining and Sentiment Analysis* (2008); Liu, *Sentiment Analysis and Opinion Mining* (2012); Thelwall, Buckley, Paltoglou, Cai y Kappas, *Sentiment Strength Detection in Short Informal Text* (2010); y Cortis y Davis, *Over a Decade of Social Opinion Mining: A Systematic Review* (2021), que revisó 485 estudios de minería de opinión social publicados entre 2007 y 2018. Fuera del derecho de familia, la analítica de sentimiento se usa habitualmente en reseñas de clientes, respuestas a encuestas, seguridad y curación de publicaciones de blogs y redes sociales, análisis de voz del cliente, retroalimentación educativa, finanzas, salud, política, gobierno y apoyo a decisiones de atención al cliente; ejemplos comerciales y públicos incluyen Google Cloud Natural Language, Amazon Comprehend, Microsoft Azure AI Language y Perspective API de Google/Jigsaw para moderación de comentarios en línea.
- 27. Contexto del producto y reconocimiento.** ISaidUSaid Pty Ltd desarrolla el sistema DijeDijiste.com para la detección de patrones conductuales durante la interacción entre personas. Sus materiales públicos describen capacidades que incluyen cargas conectadas a la fuente, redacción previa al análisis de identificadores personales, marcadores de género y referencias judiciales, OCR de capturas de pantalla, transcripciones atribuidas por hablante, detección y cotejo de huellas de voz, detección y cotejo facial, detección de conducta en video, alojamiento cifrado en región, controles de auditoría, integraciones con Clio, LEAP y Smokeball, cronología unificada e informes judiciales con un clic. El reconocimiento de la industria ha incluido el 2024 Queensland AI Award por Most Outstanding Use of AI for Social Good, el 2024 Australian AI Award por AI Innovator of the Year - Legal Services (SME), reconocimiento como finalista del 2025 Australian AI Award por Best Use of Responsible AI, reconocimiento como finalista del 2025 Australian AI Software Engineer of the Year para Adam Norman Jenkins, Inventor and Director of Technology, reconocimiento como finalista del 2024 Australian AI Female Leader of the Year para Cler Ribeiro, CEO, y la victoria de ISaidUSaid en el 2026 Legal Innovation & Tech Fest Startup Clash.
- 28. Propósito y límites.** DijeDijiste.com es principalmente un sistema de analítica de sentimiento y organización de evidencia, no una tecnología de redacción automática mediante IA generativa. La IA generativa se

utiliza únicamente para resumir resultados ya calculados en un formato más legible, y cada párrafo de este tipo se marca claramente con una daga (†). La función central del sistema es clasificar la comunicación y la conducta observada frente a ontologías conductuales definidas, y luego presentar clasificaciones, recuentos, proporciones, líneas de tiempo, ejemplos de confianza y material vinculado a la fuente para revisión humana. El informe no altera el material fuente, no formula conclusiones judiciales y no reemplaza la evaluación independiente de un abogado, experto o tribunal.

29. **Datos fuente analizados.** El sistema analiza las comunicaciones seleccionadas para este informe, incluida información de correos electrónicos, mensajes instantáneos, aplicaciones de crianza, audio, video, transcripciones, metadatos y archivos fuente disponibles en las interacciones seleccionadas. Los ejemplos y fotogramas de evidencia mostrados se extraen del material fuente almacenado. Los resúmenes narrativos no se tratan como evidencia fuente; son resúmenes de analítica subyacente vinculada a la fuente y deben contrastarse con los ejemplos, archivos fuente, marcas de tiempo y registros almacenados.
30. **Preprocesamiento, protección de género y protección de PII.** Antes de clasificar texto o transcripciones de audio, el sistema identifica subcadenas exactas en cinco clases protegidas: identificadores de progenitor/parte adulta, identificadores de menores, marcadores directos de género, referencias de causa judicial y otros identificadores personales, como una dirección privada, teléfono, correo electrónico o el nombre de un tercero. La aplicación aplica después máscaras deterministas que conservan la longitud: los términos de progenitor/parte adulta terminan en *p*, los de menores en *c*, y los términos de género, referencia judicial y otros identificadores privados se sustituyen por guiones. Después de la detección tipificada asistida por el modelo, las identidades almacenadas del remitente y del destinatario del mensaje se vuelven a comparar de forma determinista, sin distinguir mayúsculas y minúsculas y teniendo en cuenta las formas posesivas, de modo que una identidad omitida por el modelo de detección se enmascara antes de almacenar o utilizar la copia protegida. Eliminar términos directos de género antes de puntuar reduce la posibilidad de que el clasificador infiera un resultado a partir de referencias lingüísticas explícitas al sexo o género del remitente o destinatario, en vez del contenido de la comunicación. Eliminar nombres, datos de contacto y referencias de causa limita la exposición de PII y conserva las posiciones de caracteres para resaltados vinculados a la fuente. Los mapas de redacción se versionan y pueden reaplicarse de forma consistente. El material fuente original no se modifica y queda restringido a revisión autorizada de la fuente. En video, el sistema detecta y sigue regiones faciales y de cuerpo completo, y utiliza el cotejo de rostros y apariencia corporal junto con las descripciones de participantes proporcionadas durante la carga para preparar una distribución de identidades orientativa. En las transcripciones de audio y video grabados, detecta y coteja huellas de voz para preparar atribuciones orientativas de hablantes. Estas sugerencias nunca establecen la identidad por sí solas: el usuario que realiza la importación debe verificar o corregir las identidades tanto de las pistas visuales como de los hablantes de las transcripciones de audio antes de que pueda continuar el procesamiento. Las ubicaciones faciales confirmadas por el usuario se emplean para preparar la copia protegida para revisores humanos asignados. Un detector independiente dedicado a la privacidad también difumina otros rostros detectables para que un rostro no quede expuesto solo porque su identidad sea incierta. El difuminado se impone en el servidor y no puede desactivarse mediante una opción del revisor. Los términos de relación y rol de participante neutros en cuanto al género, como *progenitor*, *menor*, *responsable*, *pareja*, *cónyuge* y *familiar*, se conservan porque no identifican a una persona ni expresan

directamente su género.

Figura A1 - Marco de redacción (ejemplo práctico realista)

COPIA FUENTE RESTRINGIDA

María, hablé con la persona responsable de Lucía después de la escuela. Como progenitor, debes dejar de decirle que una buena madre estaría de acuerdo contigo. La persona responsable pidió a ambos progenitores que usaran la aplicación de coparentalidad y copiaran a Laura Gómez en los acuerdos. Recoge a Lucía en Calle Mayor 17 a las 16:30 o llama al 0400 123 456. La próxima audiencia del Tribunal de Familia de Brisbane es BFC/2026/00092. Juan

COPIA PROTEGIDA DE ANÁLISIS/REVISIÓN

----p, hablé con la persona responsable de ----c después de la escuela. Como progenitor, debes dejar de decirle que una buena ----p estaría de acuerdo contigo. La persona responsable pidió a ambos progenitores que usaran la aplicación de coparentalidad y copiaran a ----- en los acuerdos. Recoge a ----c en ----- a las 16:30 o llama al -----. La próxima audiencia del Tribunal de Familia de Brisbane es -----
---. ---p

Detección tipificada -> máscara determinista de longitud conservada -> copia protegida versionada

Este ejemplo práctico usa nombres y hechos inventados procesados conforme a las reglas de redacción actuales. No contiene evidencia de este asunto.

Figura A2 - Difuminado de rostros para revisión humana externa (demostración sintética)



Vista fuente autorizada y restringida

Copia protegida para revisión externa

El usuario identifica o clasifica las pistas de participantes detectadas. Las ubicaciones faciales confirmadas, junto con el detector de privacidad independiente, se difuminan en el servidor en los fotogramas entregados a revisores externos asignados. Los revisores no pueden desactivar el difuminado. Esta ilustración es sintética y no contiene la imagen de ningún participante.

31. **Desarrollo de la ontología, revisión por pares y escalamiento.** Las categorías conductuales, inclusiones, exclusiones, anclas de puntuación, umbrales y ejemplos límite se mantienen en la ontología específica del idioma y se prueban con comunicaciones relacionales realistas. Las muestras seleccionadas y los casos solicitados por profesionales pueden entrar en el flujo de revisión humana; esto no significa que cada mensaje del informe haya sido revisado externamente. Con autorización expresa del titular de la fuente, la copia protegida exacta puede asignarse de forma ciega a pares autorizados, incluidos profesionales de organizaciones no gubernamentales (ONG) y entidades sin fines de lucro (NFP) independientes. Se procuran un mínimo de dos revisiones sustantivas ciegas. El acuerdo material exige las mismas etiquetas normalizadas, puntuaciones con una diferencia máxima de 0,20 y al menos 50 % de solapamiento del rango de evidencia marcado más corto. Si las dos primeras revisiones discrepan materialmente, se solicita una tercera revisión ciega. Cualquier revisor de primera línea puede marcar la evidencia como límite o necesitada de atención interna; la ausencia de mayoría tras tres revisiones se deriva a adjudicación interna. Si el adjudicador interno no puede resolverla de forma segura, se escala a un experto en la materia ("SME") configurado; la decisión del SME es autoritativa. El desacuerdo previo permanece en el historial de auditoría y, una vez que el SME toma la decisión autoritativa, también se conserva en los datos de entrenamiento como metadatos de ambigüedad. La decisión del SME sigue siendo la etiqueta de entrenamiento; el desacuerdo registrado no se exporta como una etiqueta competidora. La ambigüedad constituye por sí misma una señal útil porque ayuda al modelo a aprender dónde se encuentra el límite entre ejemplos aceptados y rechazados de una conducta y cuál es la amplitud de esa zona limítrofe.
32. **Cómo se crean los datos de entrenamiento revisados.** La clasificación para el ajuste fino se realiza de forma colaborativa en todo el sector, no por una única persona anotadora ni exclusivamente por DijeDijiste. Profesionales autorizados, pares de ONG/NFP y otros revisores colaboradores del sector

clasifican ejemplos protegidos utilizando las mismas etiquetas de ontología, puntuaciones y rangos de evidencia marcados. Sus clasificaciones combinadas se someten a análisis estadístico de tasas de acuerdo, distribuciones de puntuación, solapamiento de rangos de evidencia, valores atípicos y decisiones contradictorias. Solo después de ese análisis estadístico, el personal especialista de DijeDijiste revisa los resultados, investiga desacuerdos, posibles sesgos y problemas de calidad de los datos, y aprueba, devuelve o escala la resolución propuesta. Las presentaciones humanas se capturan como instantáneas inmutables e independientes del modelo que contienen el texto de análisis protegido, idioma, modalidad, fuente y versión de redacción, etiquetas, puntuaciones, explicaciones y rangos de caracteres o fotogramas. Los campos de remitente y destinatario se redactan. Una presentación individual se guarda como candidata y no es apta para entrenamiento. Solo la decisión ganadora resuelta por consenso de pares, adjudicación interna o adjudicación SME se publica en el prefijo de entrenamiento de revisiones resueltas; las candidatas no resueltas o contradictorias se excluyen de la exportación normal. Un reanálisis dirigido aceptado por un profesional también puede crear una instantánea apta, que puede retirarse si luego se revoca la aceptación. Las herramientas de exportación vuelven a comprobar los estados de resolución y aptitud antes de producir filas de entrenamiento por tarea. Estos controles mantienen la PII sin redactar y los desacuerdos no resueltos fuera de los datos ordinarios de ajuste fino, conservando un registro auditable para calibración y evaluación.

33. **Cómo se puntúan texto y transcripciones de audio.** Tras la redacción, el mensaje objetivo se evalúa de forma independiente frente a cada conducta seleccionada de la ontología activa del idioma. Puede aportarse una ventana cronológica limitada para resolver referencias, secuencia y función comunicativa, pero solo se puntúa el mensaje objetivo marcado. El clasificador devuelve una probabilidad conductual de 0,0 a 1,0, las subcadenas exactas de apoyo y probabilidades por fragmento. Los resultados solo se conservan si alcanzan el umbral configurado de esa conducta, que se ha predeterminado a partir de datos de prueba para maximizar la precisión y minimizar los falsos positivos; las conductas concurrentes pueden compartir texto solapado. El audio sigue la misma ruta tras la transcripción de voz a texto, la atribución del hablante y la detección y cotejo de huellas de voz. Cuando hay video, la detección y cotejo facial con las pistas de participantes identificadas por el usuario respalda la atribución de participantes. La puntuación expresa la fuerza de coincidencia con la ontología en la comunicación; no es la probabilidad de que una alegación sea verdadera, una conclusión jurídica ni una medida de frecuencia.
34. **Cómo se puntúa el video.** El modelo de video revisa los fotogramas mediante una «multipass sliding window technique». Examina grupos cronológicos solapados de fotogramas en múltiples pasadas para que el movimiento y la conducta puedan considerarse a lo largo del tiempo y no a partir de una imagen fija aislada. Cuando estén disponibles, las transcripciones alineadas y la información de participantes identificada por el usuario pueden aportar contexto. Un resultado retenido debe coincidir con una conducta de video configurada e identificar el rango de fotogramas inicial y final que lo respalda; los resultados se convierten en marcas de tiempo y se enlazan con los fotogramas de evidencia del informe. La puntuación de video almacenada como 1,0 registra que la coincidencia configurada fue retenida por el proceso de revisión de video; no debe leerse como certeza calibrada del 100 %. Los revisores humanos anotan por rango de fotogramas y los fotogramas de revisión externa se difuminan como se describe arriba.
35. **Sistemas de IA y dónde se usan.** El flujo de clasificación de DijeDijiste.com utiliza una pila de modelos personalizada, guiada por ontología y ajustada finamente para la clasificación de comunicaciones de texto/audio y para la clasificación conductual de fotogramas de video. Los servicios de apoyo se usan solo para tareas específicas de procesamiento, como transcripciones de voz a texto atribuidas por hablante, detección y cotejo de huellas de voz, detección y cotejo facial con pistas de participantes identificadas por el usuario, redacción tipificada, apoyo a la estructura documental, resúmenes de informes e interacciones, flujos de embeddings y recuperación, apoyo independiente de detección facial para revisión de privacidad y atributos estándar de toxicidad/profanidad cuando se seleccionan esas categorías estándar. Los modelos específicos de proveedor pueden cambiar con el tiempo sin cambiar la metodología: el material fuente se preserva; el texto y las transcripciones se redactan antes de la clasificación o revisión humana; los fotogramas entregados a revisores externos asignados tienen los rostros

difuminados; las clasificaciones se realizan frente a la ontología activa; y los ejemplos vinculados a la fuente siguen siendo revisables.

- 36. Declaración de IA generativa y marcador †.** El propósito principal de la IA generativa de resumen del informe es hacer que la analítica numérica y los resultados analíticos ya calculados sean más legibles para el tribunal. Se utiliza únicamente en campos narrativos marcados y puede describir recuentos, proporciones, tendencias y material analítico existente vinculado a la fuente dentro del alcance, pero no determina ni modifica las clasificaciones, puntuaciones, recuentos, evidencia seleccionada o gráficos subyacentes. Cada instancia de contenido del informe generado por IA generativa se marca con una daga (†) al final del párrafo generado. La daga identifica solo la redacción explicativa; no condiciona, sustituye ni crea las cifras, clasificaciones, pruebas o resultados subyacentes.
- 37. Para qué no se utiliza la IA generativa de resumen del informe.** El modelo de resumen del informe no se utiliza para clasificar o puntuar comunicaciones o video, calcular recuentos o proporciones, elegir evidencia fuente, generar gráficos, alterar archivos fuente, decidir si una alegación fáctica es verdadera ni formular una conclusión jurídica. Esos resultados analíticos se completan y almacenan antes del resumen narrativo del informe. Después, el PDF se genera directa y determinísticamente a partir del material vinculado a la fuente, los resultados calculados, los gráficos y los campos narrativos claramente marcados; la IA generativa no genera ni maqueta el PDF.
- 38. Integridad y auditabilidad de la fuente.** El material fuente se protege separadamente del proceso analítico. Como paso inicial de manejo, DijeDijiste.com crea una instantánea de evidencia cifrada y sellada en blockchain para que cualquier eliminación o alteración posterior sea detectable antes de que comience el trabajo posterior de IA o revisión humana. Los archivos se cifran durante la transferencia y en reposo, el material fuente se registra mediante hashes criptográficos y los elementos fuente se sellan en una cadena de evidencia de blockchain privada. Las cargas, accesos, comparticiones y eventos de procesamiento se registran con marca de tiempo y se asocian a actividad de usuario verificada para apoyar la integridad de la fuente y la revisión de cadena de custodia. Los informes y documentos cargados también se someten a controles de integridad visual durante la importación, incluida la rasterización de páginas y el análisis de estructura de píxeles, diseñados para detectar renderizados documentales malformados, visualmente inconsistentes o inesperadamente alterados, lo que puede indicar la presentación de evidencia alterada o generada.
- 39. Benchmarks de transcripción.** Los benchmarks de transcripción publicados deben leerse por separado de la precisión de clasificación conductual de DijeDijiste.com. WER significa *word error rate* (tasa de error de palabras): el número de sustituciones, eliminaciones e inserciones de palabras dividido por el número de palabras de la transcripción de referencia. Un WER menor es mejor. Para una presentación comparable como tasa de precisión, el equivalente simple de precisión de palabras se calcula como 100 % menos el WER; se trata de una reformulación matemática del WER informado, no de un benchmark independiente. ElevenLabs informa una precisión de transcripción en inglés de 96,7 % para Scribe v1, y su documentación actual clasifica inglés, francés, alemán, español y portugués como idiomas de transcripción "Excellent" con un WER de 5 % o menos, equivalente a una precisión simple de palabras de 95 % o más. OpenAI informa Whisper por benchmark, no como una tasa universal única: LibriSpeech Clean mide habla inglesa leída claramente, FLEURS English mide habla en inglés dentro de un benchmark multilingüe, y el promedio de benchmark más amplio combina diversos conjuntos públicos de datos de voz. Whisper registró 2,7 % WER en LibriSpeech Clean, equivalente a 97,3 % de precisión simple de palabras; 4,4 % WER en FLEURS English, equivalente a 95,6 %; y 12,8 % WER promedio en el conjunto de benchmarks informado en el artículo de Whisper, equivalente a 87,2 % de precisión simple de palabras.
- 40. Benchmark de ontología personalizada.** El benchmark de ontología personalizada de DijeDijiste.com se mantiene frente a la ontología activa del motor en los idiomas admitidos. El benchmark usa ejemplos fuente revisados y escenarios de todas las conductas para medir detecciones esperadas, no detecciones esperadas, falsos positivos y falsos negativos. Los resultados del benchmark deben leerse como medidas de rendimiento, no como determinaciones fácticas concluyentes sobre ninguna comunicación individual de este informe. Las filas de grupos siguientes incluyen solo las conductas personalizadas de texto/audio y

de video seleccionadas para este informe.

Métrica	Benchmark español actual	Significado
Exactitud	95,75 %	El porcentaje de casos revisados del benchmark en los que el sistema devolvió el resultado esperado.
Falsos positivos	0,64 %	El porcentaje de todos los casos revisados del benchmark en los que el sistema detectó una conducta que no se esperaba.

Grupo de conductas	Conductas actuales del motor
Comportamiento violento doméstico	Admisión de culpabilidad; Conducta coercitiva; Abuso financiero; Avances sexuales no deseados; Coacción a la autolesión; Amenaza de autolesión; Acecho o vigilancia; Extorsión; Golpear; Patadas; Tirones de pelo; Forzar al suelo; Restricción; Estrangulación; Contacto sexualizado; Obstrucción; Morder, Rasguñar o Escupir; Empujar, Apretar, Tropezar o Similar; Lanzar objetos; Amenazar con Arma (pistolas, cuchillos, etc.); Otras Alteraciones Físicas
Comportamiento abusivo	Crítica parental; Despreciar el cuerpo; Intimidación agresiva; Acusaciones de abuso no verificadas; Gas-lighting; Gestos abusivos; Arrojar líquidos o sustancias; Acciones destructivas
Comportamiento manipulador	Amenazas legales; Inducción de vergüenza; Explotación de vulnerabilidades; Culpa manipuladora; Disculpa performativa; Atribución de salud mental como arma; Acorralar o bloquear el movimiento; Interferir con el teléfono o dispositivo
Comportamiento antisocial	Rechazo oficial a la mediación; Expresión de resentimiento
Comunicación constructiva	Co-parenting proactivo; Solicitudes de mediación; Desescalada; Empatía; Construcción de puentes; Establecimiento de límites saludables
Reparación y reconciliación	Declaraciones de amor y devoción; Remordimiento; Buscando perdón; Conceder perdón; Gracia / Misericordia
Señales de advertencia	Superioridad moral; Control de puntuaciones; Devaluación; Intención vengativa; Enmarcado juicioso; Crítica (Instituto Gottman); Defensividad (Instituto Gottman); Desdén (Instituto Gottman); Evasivas (Instituto Gottman)
Tácticas de presión	Ultimátum; Invalidación emocional; Deuda emocional; Castigo por retirada



Mr Adam Norman Jenkins, Director of Technology, ISaidUSaid Pty Ltd

--Fin del informe--