



IchSagDuSagst™

Ihre Wahrheit, authentifiziert™

Powered by IchSagDuSagst

Gerichtlicher Kommunikationsbericht

Alle Nachrichten

Erstellt von IchSagDuSagst.com

1. Dieser Bericht wurde am 27. August 2026 für das Verfahren QA-DE-2026-08-27, das beim/bei der Familiengericht Berlin von Berlin, BE, DEU geführt wird, mit dem Sentiment-Analytics-KI-System von IchSagDuSagst.com erstellt. Der Hauptzweck generativer KI zur Berichtszusammenfassung besteht darin, die numerischen Analysen und bereits berechneten Analyseergebnisse für das Gericht leichter lesbar zu machen. Sie wird ausschließlich in gekennzeichneten narrativen Feldern eingesetzt und kann außerdem abgegrenztes quellenbezogenes Analysematerial zusammenfassen. Sie bestimmt oder verändert weder die zugrunde liegenden Klassifizierungen, Bewertungen oder Zählungen noch ausgewählte Quellbeweise, Diagramme oder die PDF-Zusammenstellung. Jede Instanz generativ erzeugten Berichtsinhalts wird am Ende des erzeugten Absatzes mit einem Dolchzeichen (†) gekennzeichnet. Anhang A - Methodik erläutert dieses Zeichen, nennt die verwendeten Systeme und beschreibt die Schutzmaßnahmen zur quellenbezogenen Überprüfung.
2. Die folgenden digitalen Interaktionen wurden zwischen Mr Joe Blogs und Ms Jill Blogs im Zeitraum von 09. April 2024 bis 09. Oktober 2024 analysiert:
 - a. 44 E-Mail-Nachrichten.
 - b. 0 Videoaufnahmen.
 - c. 0 aufgezeichnete Gespräche.
 - d. 0 Sprachnachrichten.
 - e. 0 Sofortnachrichten (z. B. WhatsApp, iMessage)
 - f. 0 Nachrichten, die über Co-Parenting-/Elternschafts-Apps gesendet wurden (z. B. Our Family Wizard, Talking Parents, CoParent Coordinator, AppClose)
 - g. 0 hochgeladene Screenshot-Quellbilder.
3. Die Kommunikation wurde auf die folgenden negativen Interaktionsstile hin analysiert:
 - a. Schimpfwörter: Hebt Nachrichten hervor, die Schimpfwörter, Flüche oder andere vulgäre Ausdrücke enthalten.
 - b. Bedrohungen: Identifiziert Aussagen, die drohen, jemandem Schaden zuzufügen oder bevorstehenden körperlichen Schaden andeuten.
 - c. Toxische Sprache: Bezeichnet Sprache, die allgemein feindlich, missbräuchlich oder unnötig grausam ist.
 - d. Sehr toxische Sprache: Kennzeichnet Nachrichten, die selbst nach toxischen Standards extrem feindselig oder missbräuchlich sind.
 - e. Beleidigungen: Identifiziert direkte Beschimpfungen oder herabwürdigende Bemerkungen, die an die andere Person gerichtet sind.
 - f. Identitätsangriffe: Erkennt Beleidigungen, die sich gegen Rasse, Geschlecht, Sexualität oder andere Identitätsmerkmale richten.
 - g. Verstoß gegen eine gerichtliche Anordnung: Warnt Sie, wenn eine Nachricht erscheint, die gegen eine hochgeladene gerichtliche Anordnung verstößt.
 - h. Geständnis: Findet Aussagen, in denen der Absender Tatsachen zugibt, die häusliche Gewalt oder rechtsverletzendes Verhalten belegen.

- i. Zwangsausübendes Verhalten: Sucht nach Beschreibungen, in denen jemand das Leben einer anderen Person eng kontrolliert.
- j. Finanzielle Misshandlung: Weist auf Formulierungen hin, die Geld, Zugang zu Mitteln oder finanzielle Freiheit einschränken.
- k. Elterliche Kritik: Hebt Angriffe auf die Erziehungsfähigkeiten oder Entscheidungen der anderen Person hervor.
- l. Körperbeschämung: Oberflächenbemerkungen, die den Körper oder das Aussehen einer Person verspotten oder herabsetzen.
- m. Mobbing: Kennzeichnet Versuche, durch aggressive Sprache einzuschüchtern, zu nötigen oder zu dominieren.
- n. Nicht verifizierte Missbrauchsvorwürfe: Identifiziert Anschuldigungen von Missbrauch oder Straftaten, für die im vom analysierenden Nutzer hochgeladenen Dokumentenbündel keine gerichtliche Feststellung der Schuld vorliegt.
- o. Unerwünschte sexuelle Annäherungen: Erkennt sexuelle Annäherungen, Vorschläge, explizite Bemerkungen oder intime Kommentare und nutzt nahe Zurückweisung, Grenzen oder fehlende Gegenseitigkeit, um unerwünschte Annäherungen von klar willkommenen oder erwiderten intimen Äußerungen zu unterscheiden.
- p. Rechtliche Drohungen: Erkennt Drohungen, Anwälte, Gerichte oder die Polizei als Druckmittel einzubeziehen.
- q. Zwang zur Selbstverletzung: Hebt Aussagen hervor, die Selbstverletzung fördern, Druck ausüben oder vorschlagen.
- r. Angedrohte Selbstverletzung: Erkennt Aussagen, in denen Selbstverletzung oder Suizid angedroht werden, um Druck oder Kontrolle auszuüben.
- s. Stalking: Kennzeichnet Eingeständnisse, heimlich jemandem zu folgen, dessen Bewegungen zu überwachen oder zu verfolgen.
- t. Offizielle Mediationsverweigerung: Notiert explizite Verweigerungen zur Teilnahme an Mediation oder geführtem Dialog.
- u. Erpressung: Erkennt bedingte Bedrohungen, die Compliance oder Zugeständnisse fordern.
- v. Schaminduktion: Kennzeichnet Sprache, die darauf abzielt, zu erniedrigen oder die andere Person wertlos erscheinen zu lassen.
- w. Moralische Überlegenheit: Hebt Botschaften hervor, die eine ethische Überlegenheit behaupten, um jemanden herabzusetzen.
- x. Punktestand: Erkenntnisse über vergangene Fehler, die zur Kontrolle oder Forderung von Wiedergutmachung verwendet werden.
- y. Abwertung: Erkennt Du-Botschaften, die Wert, Kompetenz, Verlässlichkeit, Fürsorge oder Bedeutung der anderen Person mindern.
- z. Rachsüchtige Absicht: Kennzeichnet Äußerungen des Wunsches nach Vergeltung oder Rache.
- aa. Ausnutzung von Verwundbarkeit: Identifiziert Manipulation, die die Geheimnisse oder Wunden einer Person als Waffe einsetzt.
- ab. Schuldzuweisungen: Hebt Versuche hervor, durch Schuld oder emotionale Verpflichtungen Kontrolle auszuüben.
- ac. Performative Entschuldigung: Erkennt hohle Entschuldigungen, die zur Steuerung der Wahrnehmung angeboten werden, anstatt den Schaden zu beheben.
- ad. Ausdruck von Groll: Kennzeichnet anhaltende Bitterkeit, die verwendet wird, um zu bestrafen oder zu beschämen.
- ae. Urteilende Einordnung: Identifiziert pauschale moralische Verurteilungen des Charakters der anderen Person.
- af. Ultimatum: Kennzeichnet bedingte Forderungen, bei denen die Zusammenarbeit von der Einhaltung abhängt.
- ag. Emotionale Entwertung: Hebt Aussagen hervor, die die Gefühle des Empfängers abtun, minimieren oder verspotten.
- ah. Waffenisierte Zuschreibung von psychischer Gesundheit: Kennzeichnet Aussagen, die psychische Erkrankungen oder psychische Instabilität zuschreiben, um die Glaubwürdigkeit des Empfängers zu

- untergraben oder dessen Bedenken abzutun.
- ai. Gaslighting: Realitätsuntergrabende Sprache (Gaslighting-Indikator): erkennt ausdrückliche Versuche, das Gedächtnis, die Wahrnehmung oder die Fähigkeit einer Person, Realität zu unterscheiden, zu diskreditieren.
 - aj. Emotionale Verschuldung: Kennzeichnet Druck basierend auf impliziter emotionaler Verschuldung wie 'nach allem, was ich für dich getan habe'.
 - ak. Rückzug als Bestrafung: Kennzeichnet Drohungen, die Kommunikation oder Zusammenarbeit abubrechen, um den Empfänger unter Druck zu setzen.
 - al. Kritik (Gottman Institut): Fasst Momente zusammen, in denen die Person den Charakter anstelle des Verhaltens angreift.
 - am. Defensivität (Gottman Institut): Kennzeichnet Antworten, die Schuld zurückweisen, Verantwortung abstreiten, Ausreden machen oder gegenbeschweren.
 - an. Verachtung (Gottman Institut): Kennzeichnet Kommunikation, die von Spott, Verachtung oder Respektlosigkeit geprägt ist.
 - ao. Mauern (Gottman Institut): Identifiziert Rückzugstaktiken wie Stille, kurze Antworten oder das Verlassen des Gesprächs.
 - ap. Schlagen: Kennzeichnet Clips, in denen jemand eine andere Person mit der Hand, Faust oder dem Unterarm schlägt.
 - aq. Treten: Erkennt Tritte oder stoßende Fußbewegungen gegen eine andere Person.
 - ar. An den Haaren ziehen: Hebt Szenen hervor, in denen jemand eine andere Person an den Haaren packt oder zieht.
 - as. Zu Boden zwingen: Markiert Aufnahmen, in denen jemand eine andere Person gewaltsam zu Boden drückt oder wirft.
 - at. Einschränkung: Hebt Aufnahmen hervor, in denen eine Person eine andere greift, festhält oder zieht.
 - au. Ersticken: Identifiziert Szenen, in denen jemand den Hals oder die Atemwege einer anderen Person einschränkt.
 - av. Sexualisierter Kontakt: Erkennt sexualisierte Berührungen, Begrapschen, erzwungene Küsse oder intimen Körperkontakt im Bild.
 - aw. Missbräuchliche Gesten: Weist auf bedrohliche Handgesten hin, die dazu dienen, einzuschüchtern oder zu erniedrigen.
 - ax. Blockieren: Erkennt Versuche, den Weg einer Person zu blockieren oder sie am Verlassen zu hindern.
 - ay. In die Enge treiben oder Bewegung blockieren: Hebt Situationen hervor, in denen jemand eine andere Person einkesselt oder ihre Bewegung eng kontrolliert.
 - az. Telefon- oder Gerätezugriff verhindern: Erkennt Versuche, einer anderen Person das Telefon oder ein anderes Gerät wegzunehmen oder dessen Nutzung zu verhindern.
 - ba. Beißen, Kratzen oder Spucken: Erkennt aggressives Beißen, Kratzen oder Spucken gegenüber einer anderen Person.
 - bb. Drücken, Schubsen, Stolpern oder Ähnliches: Kennzeichnet Bewegungen, die jemanden schubsen, stolpern oder anderweitig aus dem Gleichgewicht bringen.
 - bc. Gegenstände werfen: Hebt Objekte hervor, die auf oder in der Nähe von jemandem geworfen werden, um Angst zu erzeugen oder Schaden zuzufügen.
 - bd. Flüssigkeiten oder Substanzen werfen: Erkennt das absichtliche Schütten, Spritzen oder Werfen von Flüssigkeiten oder anderen Substanzen auf eine Person.
 - be. Zerstörerische Handlungen: Hebt Situationen hervor, in denen jemand aus Wut Gegenstände beschädigt, schlägt oder zerstört.
 - bf. Drohung mit einer Waffe (Pistolen, Messer usw.): Identifiziert Momente, in denen eine Waffe gegen eine andere Person gerichtet wird.
 - bg. Andere körperliche Auseinandersetzungen: Beinhaltet jede andere körperliche Auseinandersetzung, die offensichtliches Unbehagen oder Schmerz verursacht.
4. Die Kommunikation wurde außerdem auf die folgenden positiven Interaktionsstile hin analysiert:
- a. Proaktives Co-Parenting: Erkennt kooperative Vorschläge, die das Co-Parenting auf Kurs halten.
 - b. Anfragen zur Mediation: Identifiziert höfliche Einladungen zur Lösung von Streitigkeiten durch Mediation.

- c. Deeskalation: Hebt Bemühungen hervor, einen Konflikt zu beruhigen, Sorgen zu validieren oder die Situation zu entspannen.
- d. Liebeserklärungen & Hingabe: Zeigt romantische Bekundungen von Liebe, Hingabe oder Verehrung.
- e. Reue: Lenkt die Aufmerksamkeit auf aufrichtige Ausdrucksformen des Bedauerns für verursachten Schaden.
- f. Um Verzeihung bitten: Hinweise auf verletzliche Bitten um Verzeihung oder Rückkehr.
- g. Vergebung gewähren: Zeigt, wenn jemand eindeutig Vergebung an die andere Partei gewährt.
- h. Empathie: Identifiziert Sprache, die mit den Gefühlen oder dem Schmerz einer anderen Person in Resonanz steht.
- i. Brückenbau: Hebt Einladungen hervor, Kompromisse zu finden, zusammenzuarbeiten oder einen Mittelweg zu suchen.
- j. Gesunde Grenzsetzung: Erkennt nicht-zwanghafte Aussagen, die persönliche, Kontakt- oder Beziehungsgrenzen setzen.
- k. Gnade / Barmherzigkeit: Hebt Momente der Freundlichkeit hervor, die auch dann angeboten werden, wenn sie nicht erforderlich sind.

Analyse

5. Die analysierten Interaktionen zeigen wiederholt anhaltende Konflikte zwischen dem Kontoinhaber Mr Joe Blogs und seiner Gegenpartei Ms Jill Blogs in Bezug auf das Kind Johnny, insbesondere Streit über Umgangs-/Zugangsregelungen, konkrete Situationen im Alltag (u. a. Verhalten während Mahlzeiten) sowie darüber, wie sich die elterlichen Auseinandersetzungen auf Johnny auswirken. In den früheren Phasen dominieren eskalierende, vorwurfsvoll-ausgerichtete Kommunikation mit Drohungen, persönlichen Beleidigungen/abwertenden Aussagen und wechselseitigen konditionierten Konsequenzandrohungen (u. a. im Kontext von rechtlichen Schritten oder Umgangsrechten). Später verlagert sich der Ton erkennbar hin zu Konfliktmanagement: Beide Seiten sprechen stärker über das „Schritt zurücknehmen“, das Erstellen eines gemeinsamen, praktikablen Zeitplans und das Aufsetzen respektvollerer Absprachen im Rahmen von Mediation. Gegen Ende der in-scope Kommunikation treten neben den weiterhin bestehenden organisatorischen Streitpunkten zusätzliche Themen auf, insbesondere Fragen zur Sexualität bzw. zum Coming-out von Jill und die Sorge, wie darüber mit Johnny kommuniziert werden soll, ergänzt durch Belastungs-/Finanzsorgen; zugleich wird die Angemessenheit von Mediation erneut hinterfragt und die Option anwaltlicher/legalen Beratung wird erwogen. Insgesamt entwickelt sich die Kommunikation über die Zeit von hochkonfrontativer Eskalation hin zu stärker strukturierter, lösungsorientierter Sprache, wobei Konfliktfelder (Umgangsrechte, Kommunikation zwischen den Eltern) bestehen bleiben und sich um psychosoziale/kommunikative Aspekte erweitern.[†]

6. Die folgenden positiven Kommunikationsstile wurden identifiziert:

Kommunikationstyp	Mr Joe Blogs	Ms Jill Blogs
Reue	0	0
Brückenbau	0	2
Liebeserklärungen & Hingabe	0	0
Gesunde Grenzsetzung	0	0
Um Verzeihung bitten	0	0
Vergebung gewähren	1	0
Empathie	1	0
Deeskalation	1	1
Anfragen zur Mediation	1	0
Gnade / Barmherzigkeit	0	0
Proaktives Co-Parenting	0	2

7. Die folgenden negativen Kommunikationsstile wurden identifiziert:

Kommunikationstyp	Mr Joe Blogs	Ms Jill Blogs
Hat den Konflikt ausgelöst	2	1
Moralische Überlegenheit	0	0
Ausdruck von Groll	4	0
Waffenisierte Zuschreibung von psychischer Gesundheit	0	0
Finanzielle Misshandlung	0	0
Rückzug als Bestrafung	0	2
Elterliche Kritik	1	0
Rechtliche Drohungen	1	4
Performative Entschuldigung	0	0
Zwangsausübendes Verhalten	0	1
Identitätsangriffe	0	0
Sehr toxische Sprache	0	1

Ultimatum	0	3
Emotionale Verschuldung	0	0
Schimpfwörter	3	1
Gaslighting	0	0
Schaminduktion	0	1
Geständnis	0	0
Urteilende Einordnung	2	0
Zwang zur Selbstverletzung	1	0
Schulduweisungen	0	0
Angedrohte Selbstverletzung	0	0
Körperbeschämung	0	1
Stalking	0	0
Offizielle Mediationsverweigerung	0	2
Erpressung	0	1
Ausnutzung von Verwundbarkeit	0	0
Unerwünschte sexuelle Annäherungen	1	0
Rachsüchtige Absicht	0	0
Verachtung (Gottman Institut)	3	0
Beleidigungen	4	3
Verstoß gegen eine gerichtliche Anordnung	0	0
Emotionale Entwertung	0	0
Punkttestand	0	0
Toxische Sprache	6	3
Mobbing	1	0
Bedrohungen	1	0
Defensivität (Gottman Institut)	0	0
Nicht verifizierte Missbrauchsvorwürfe	0	1
Kritik (Gottman Institut)	0	0
Abwertung	5	4
Mauern (Gottman Institut)	0	0

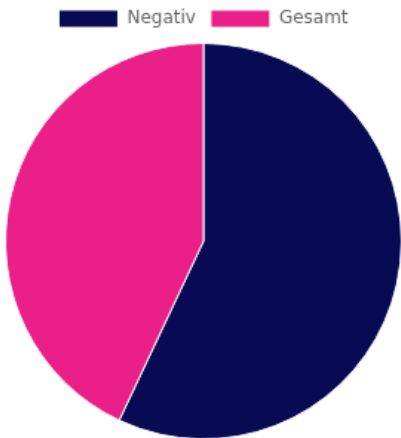
8. Die folgenden negativen Handlungen wurden in Video-Beweismaterial beobachtet:

Handlungstyp	Mr Joe Blogs	Ms Jill Blogs
Zu Boden zwingen	0	0
Flüssigkeiten oder Substanzen werfen	0	0
Gegenstände werfen	0	0
Missbräuchliche Gesten	0	0
In die Enge treiben oder Bewegung blockieren	0	0
Telefon- oder Gerätezugriff verhindern	0	0
Beißen, Kratzen oder Spucken	0	0
Schlagen	0	0
Andere körperliche Auseinandersetzungen	0	0
Zerstörerische Handlungen	0	0
Einschränkung	0	0

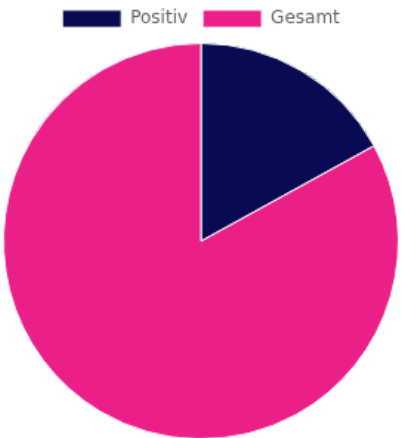
Blockieren	0	0
Sexualisierter Kontakt	0	0
Treten	0	0
Drohung mit einer Waffe (Pistolen, Messer usw.)	0	0
Erstickern	0	0
An den Haaren ziehen	0	0
Drücken, Schubsen, Stolpern oder Ähnliches	0	0

9. 57% der von Mr Joe Blogs gesendeten Nachrichten enthielten ein gewisses Maß an negativer Kommunikation. 17% der von Mr Joe Blogs gesendeten Nachrichten enthielten ausdrücklich positive Kommunikation.

Negative gesendete Nachrichten

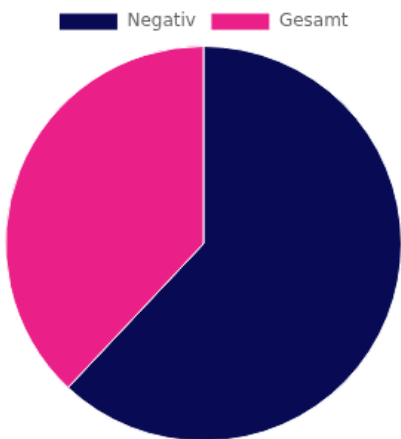


Positive gesendete Nachrichten

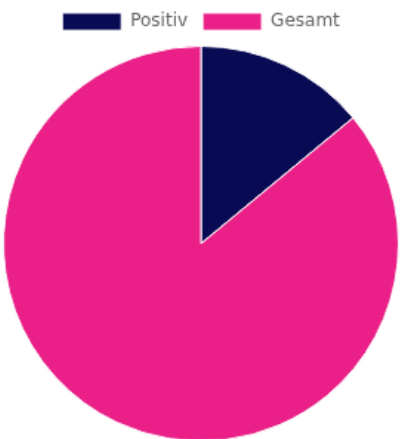


10. 62% der von Ms Jill Blogs empfangenen Nachrichten enthielten ein gewisses Maß an negativer Kommunikation. 14% der von Ms Jill Blogs empfangenen Nachrichten enthielten ausdrücklich positive Kommunikation.

Negative empfangene Nachrichten

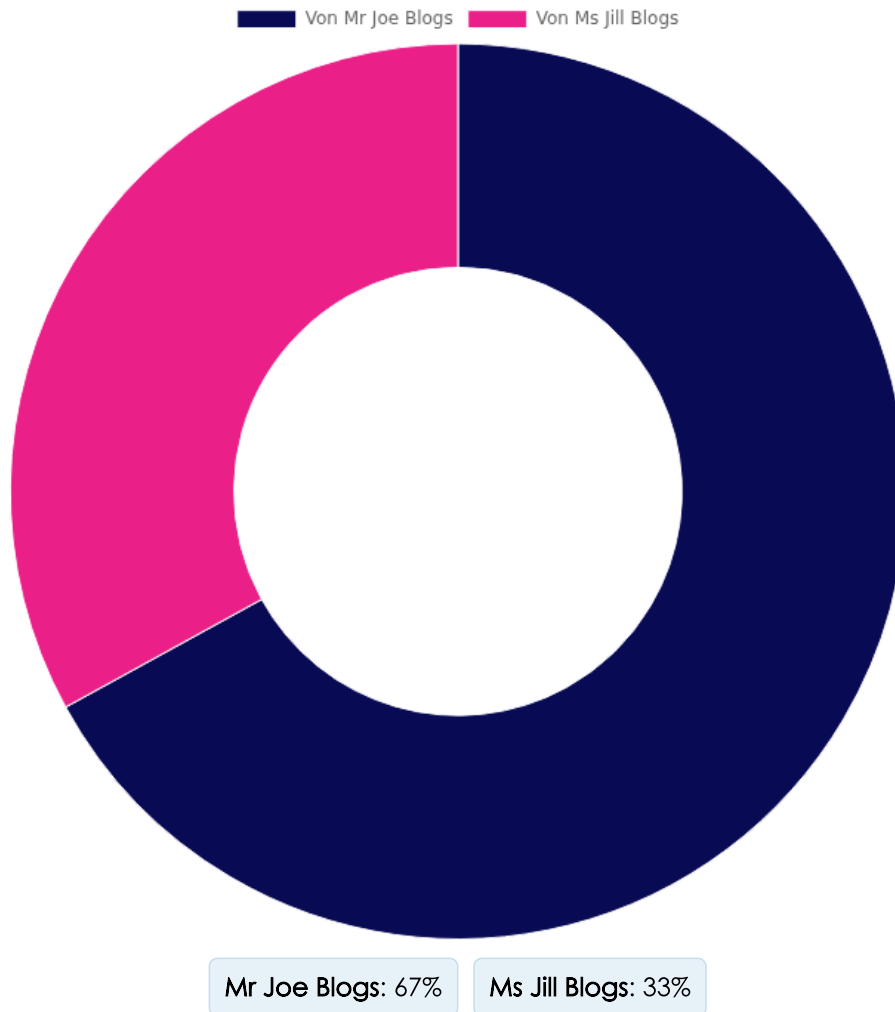


Positive empfangene Nachrichten



11. In 67% der Gespräche mit einer ersten negativen Verhaltensweise wurde diese erste negative Verhaltensweise von Mr Joe Blogs gezeigt.
12. In 33% der Gespräche mit einer ersten negativen Verhaltensweise wurde diese erste negative Verhaltensweise von Ms Jill Blogs gezeigt.

Erste negative Verhaltensweise gezeigt von...

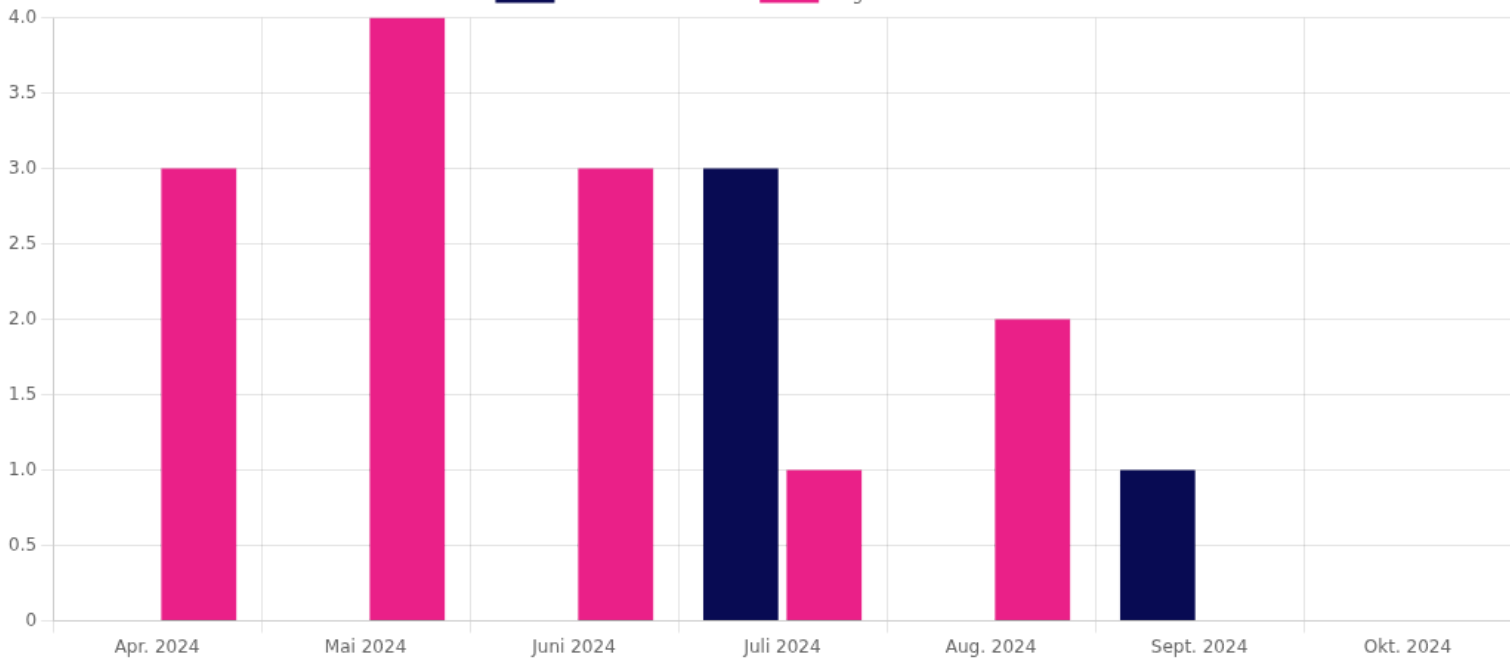


13. In den sent messages zeigt sich insgesamt eine deutliche Zuspitzung und zwischenzeitliche Eskalation des Tons: Früh im Zeitraum (April) beschreiben die Nachrichten das Verhalten von Johnny als stark unangemessen und fordern wiederholt ein Eingreifen, wobei bereits direkte Beschimpfungen und vulgäre/abwertende Formulierungen auftauchen. Im Mai setzt sich die Negativität fort bzw. wird schlimmer: Jill wird in aggressiven, erniedrigenden Aussagen adressiert („erbmärlicher Verlierer“, abwertende rassistische Bemerkungen) und es folgen Drohungen, Jill den Kontakt zu Johnny zu entziehen. Auch Mitte/Ende Juni bleiben die Nachrichten stark negativ und beleidigend, bevor sie Anfang Juli zwar kurz freundlicher wirken („zurückgenommen“, Anerkennung als Elternteil), aber ohne dass sich das Muster dauerhaft stabilisiert. Ab September verlagert sich der Ton teilweise in Richtung sachlicher und weniger konfrontativ (Fokus auf Kommunikation, Mediation, Auswirkungen auf Johnny), dennoch bleibt die Korrespondenz insgesamt wechselhaft und nicht durchgehend „gesünder“—ein klarer Trend hin zu respektvoller, durchgehend höflicherem Umgang ist nur teilweise erkennbar.[†]

Stimmung ausgehender Nachrichten (monatlich)

Gesendete Nachrichten im Zeitverlauf

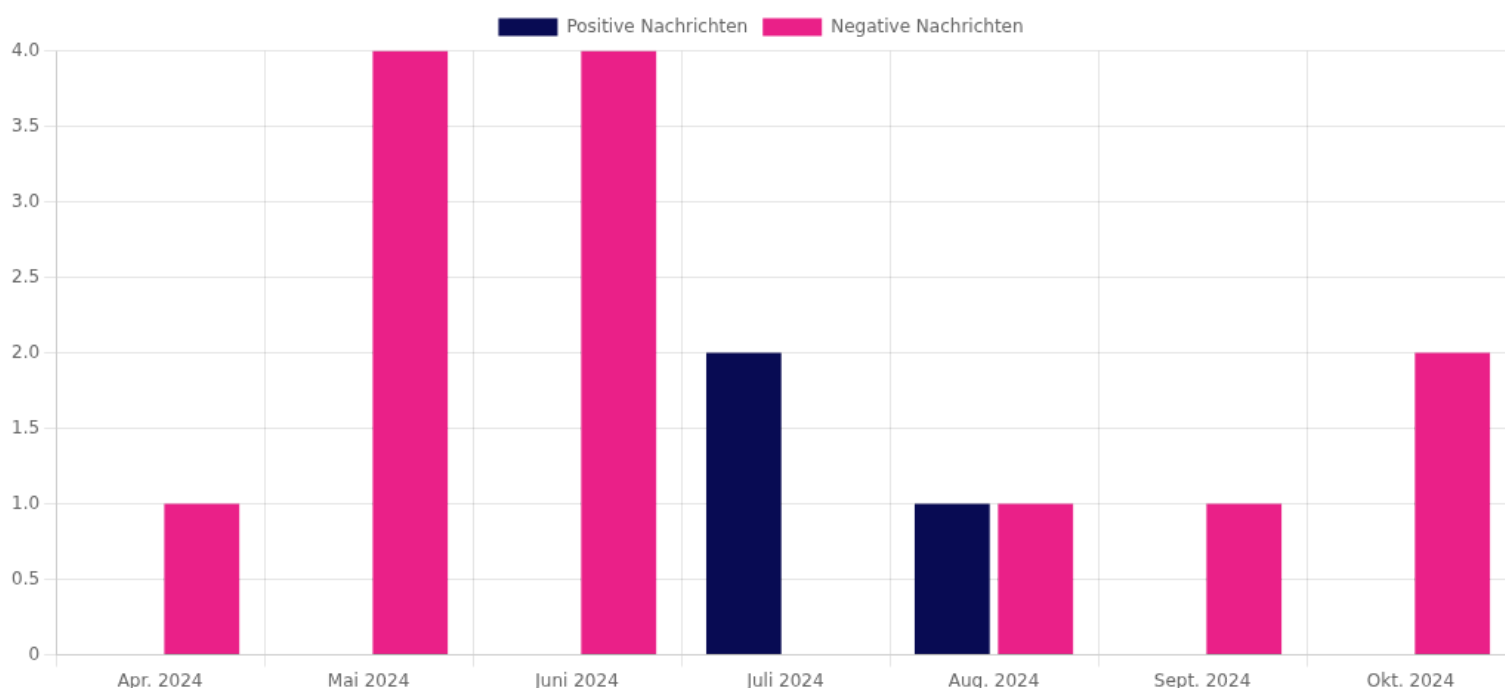
Positive Nachrichten Negative Nachrichten



14. In den „received messages“ fällt zunächst eine klar eskalierende Dynamik auf: In den frühen Monaten richten sich die Nachrichten der Absenderin mehrfach mit stark abwertender und drohender Tonlage gegen Joe (z.B. Beschimpfungen, Druckausübung über Unterhalt und die Ankündigung, rechtliche Schritte bzw. „Anwälte“ einzusetzen). Danach wird der Ton zwar stellenweise „geordnet“ bzw. konfliktnäher („lass uns zusammenarbeiten“, „versuchen, respektvoll einen gemeinsamen Nenner zu finden“), gleichzeitig bleiben aber auch wieder vereinzelte aggressive Inhalte präsent (u.a. Beleidigungen und das Ablehnen von Mediation als „Zeitverschwendung“). Ab August/September äußert die Absenderin mehr Sorge und Kommunikationsprobleme, spricht häufiger konkrete Themen (Belastung, Nachrichten an Johnny, Anpassungen zum Besuchsrecht) an und wirkt dabei weniger in Form pauschaler Beschimpfungen als in der ersten Phase; dennoch gibt es keinen durchgehend ruhigeren, konsistent höflichen Verlauf, weil die Kommunikation insgesamt von wiederkehrender Konfrontation und Drohkulisse geprägt bleibt. Insgesamt lässt sich damit sagen: Es gibt einzelne Phasen, in denen die Tonalität etwas „gesitteter“ wirkt, aber der Gesamtrend ist nicht eindeutig besser—die Kommunikation bleibt phasenweise aggressiv und negativ.[†]

Stimmung eingehender Nachrichten (monatlich)

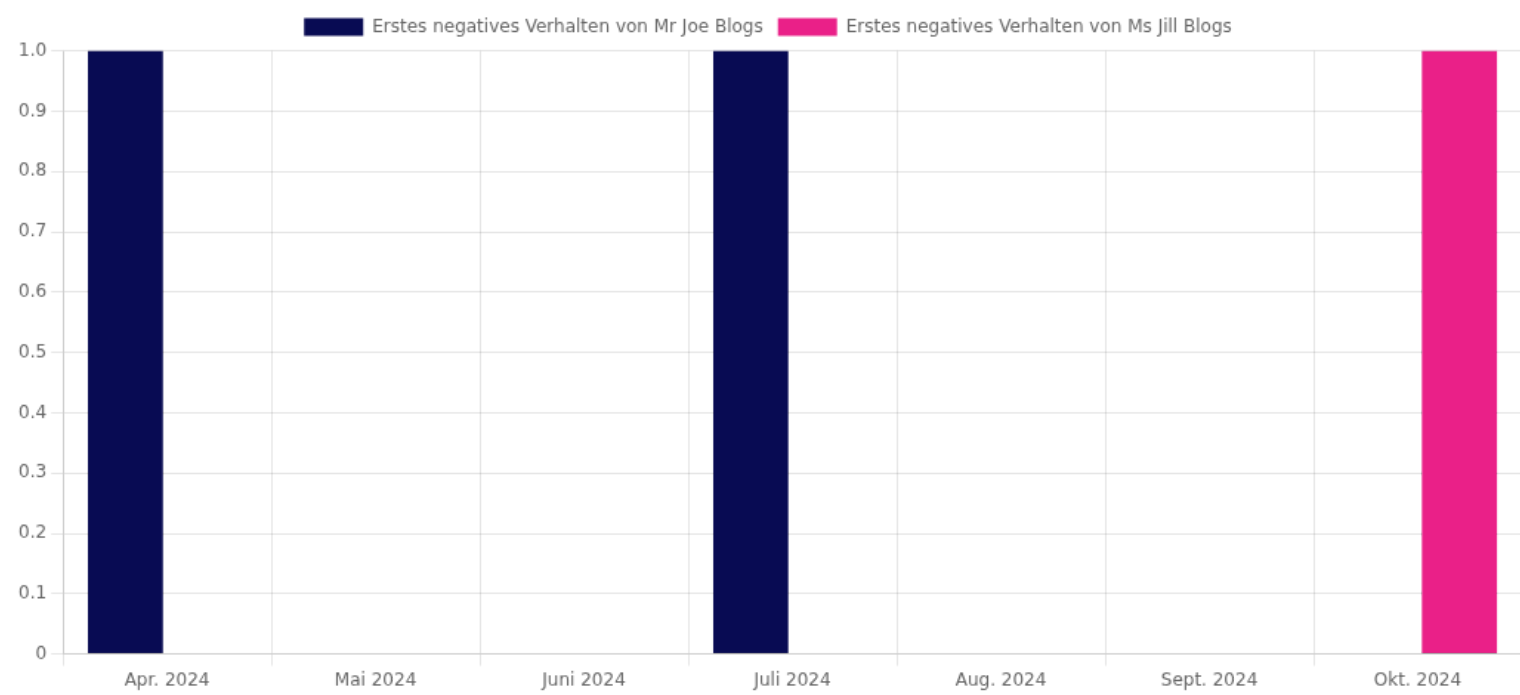
Empfangene Nachrichten im Zeitverlauf



15. Die folgende Grafik zeigt monatsweise, welche Partei die erste negative Verhaltensweise im Gespräch gezeigt hat.

Erste negative Verhaltensweise (monatlich)

Erste negative Verhaltensweise im Zeitverlauf



Partei-zu-Partei-Überblick

Der vollständige Überblick oben zeigt alle empfangenen Interaktionen zusammen. Die folgenden Abschnitte isolieren jede konfigurierte Gegenpartei, damit die Diagramme zu empfangenen Nachrichten und die Polaritätswerte auf Ebene der einzelnen Person gelesen werden können.

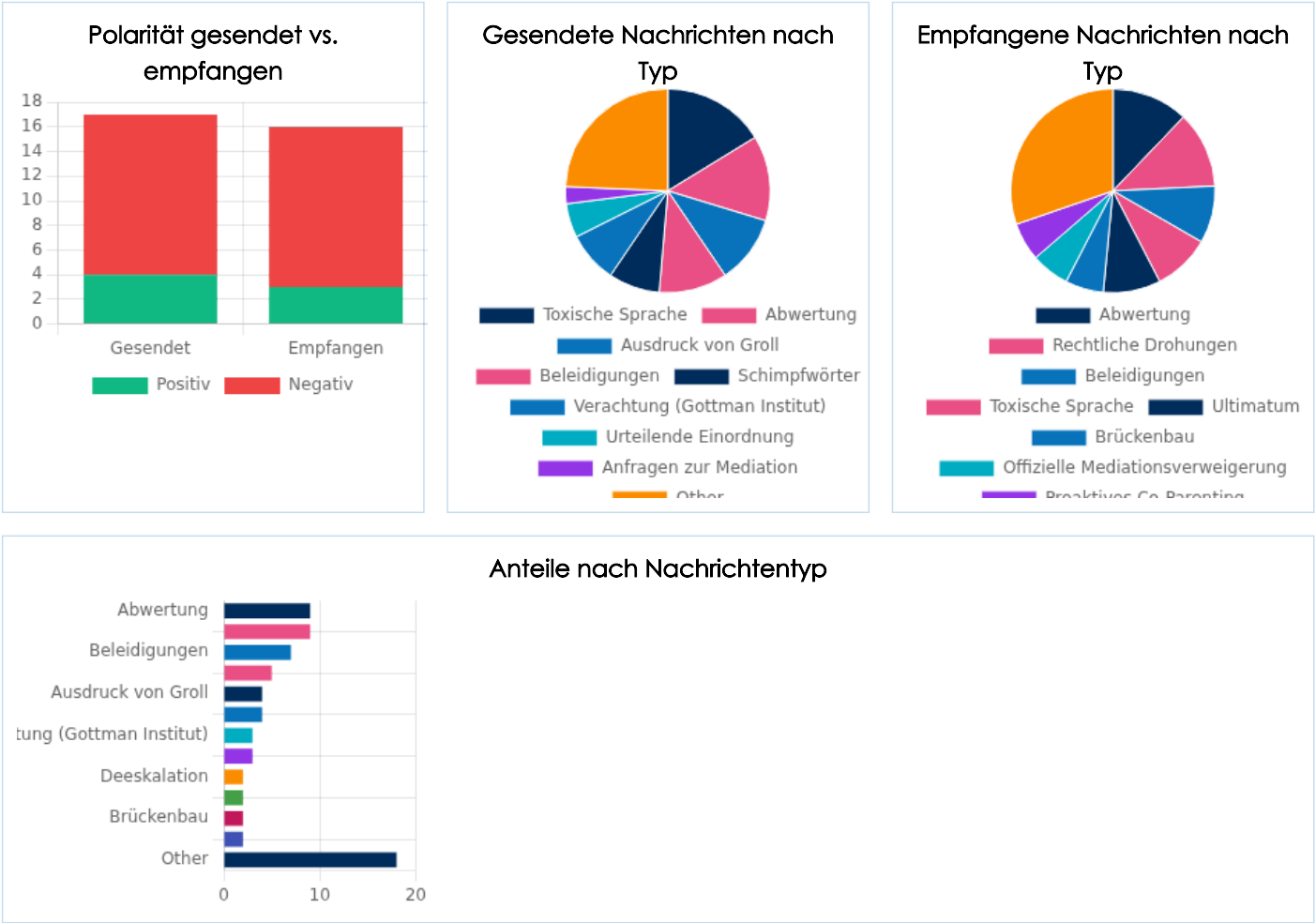
Mr Joe Blogs und Ms Jill Blogs

Gesendet: 23

Empfangen: 21

Negativ gesendet: 13

Negativ empfangen: 13



Übersicht

Text (email, messaging, social).

16. Quellen.

a. List of Email conversations with Ms Jill Blogs

17. **Das Positive.** In den analysierten Interaktionen zeigen sich als positive Kommunikationsverhalten vor allem Phasen, in denen beide Seiten versuchen, aus dem Konflikt heraus in eine strukturierte Klärung zu wechseln: Sie sprechen Themen sachorientiert an (z.B. Besuchs-/Zugriffsregelungen, praktische nächste Schritte) und bemühen sich um respektvollere Verständigung, indem sie gemeinsame Grundlagen („gemeinsamer Boden“) suchen. Außerdem werden konstruktive Dispute-Management-Schritte erkennbar, etwa der Vorschlag bzw. die Planung von Mediationsterminen und das Erstellen eines zeitlichen/organisatorischen Plans, um künftige Abläufe verbindlicher zu machen; einzelne Passagen deuten zudem auf Reflexion hin, etwa das Anliegen, sich zu „beruhigen“ oder den Umgang miteinander zu verbessern.[†]
18. **Das Negative.** In den untersuchten Gesprächen treten wiederholt stark eskalierende und potenziell schädliche Kommunikationsmuster auf: Joe und Jill richten sich mit persönlichen Beleidigungen/abwertenden Aussagen gegeneinander, äußern Drohungen und „Bedingungen“, etwa im Zusammenhang mit Umgangs-/Sorgerechtsfragen (z.B. Konsequenzen androhen, nicht nachgeben, Zugriff verhindern), und verstärken den Konflikt durch manipulative bzw. einschüchternde Rhetorik (inklusive Ankündigungen, anderen zu berichten/Vorwürfe zu verbreiten). Zusätzlich wird Jill einmal mit einer klar beleidigenden, vulgären Formulierung an Joe sichtbar, was die Kommunikation weiter entgleist; insgesamt dominiert zeitweise ein feindseliger, konfrontativer Ton, der die Kooperation und eine kindzentrierte, respektvolle Gesprächsführung untergräbt.[†]
19. **Allgemeine Kommunikationsanalyse.** Über die in-scope Nachrichten hinweg zeigt sich ein klarer Verlauf von eskalierter, konfrontativer Kommunikation hin zu zeitweise strukturierter Konfliktbearbeitung: Anfangs dominieren zwischen Joe und Jill wiederkehrende Vorwürfe, Forderungen und persönliche Herabwürdigungen, teils verbunden mit klaren Droh- bzw. Druckelementen (z.B. Konsequenzen im Zusammenhang mit Umgangs-/Sorgeregelungen, Ankündigungen von rechtlichen Schritten oder „anderen Bescheid geben“). Mit der Zeit verschiebt sich der Ton in Richtung deeskalierenderer, verhandlungsorientierter Sprache: Es werden gemeinsame Schritte wie ein „Schedule“, der Versuch respektvoller Übereinkunft und konkret Mediationstermine/Absprachen angesprochen, wobei beide Seiten weiterhin zentrale Sorgepunkte zu Johnny einbringen (Umgang, Essen, Besuchsrechte, Auswirkungen des Elternkonflikts). Thematisch treten später zusätzlich belastende Gesprächsfelder auf (Sexualität/Coming-out, wie man darüber mit Johnny kommuniziert, mögliche psychische/kommunikative Nebenwirkungen), begleitet von Unsicherheit, Überforderung und der Frage, ob Mediation „passt“ bzw. ob stattdessen Rechtsberatung sinnvoll sein könnte. Die Interaktion mit dem einzelnen, sehr beleidigenden deutschen Ausrutscher von Jill wirkt als einmaliger Tiefpunkt im Eskalationsmuster; insgesamt ist die Kommunikation trotz späterer Annäherungsversuche immer wieder von emotionaler Überlastung und konfliktgetriebenen Reaktionen geprägt, und in Teilen sind die gelieferten Daten eher spärlich (einmal nur eine einzelne Äußerung).[†]

Interaktionsdetails

20. Diese Interaktion umfasst den Zeitraum vom 09. April 2024 bis 09. Oktober 2024 und enthält im Berichtsumfang 3 Unterhaltungen mit 44 Nachrichten (7 positiv, 26 negativ). Die Daten stammen aus Text (email, messaging, social)-Datensätzen und spiegeln die Kommunikation wider, die in den ausgewählten Berichtsumfang fällt.
21. Across the interaction, Joe and Jill repeatedly discuss disputes involving Johnny, including access-related issues and how their disagreements affect or involve Johnny. Early on, they exchange escalating hostile statements, including threats or insults toward each other and references to taking actions such as using

lawyers, while also making statements about refusing to comply with conditions. Later, the conversation shifts toward attempting to resolve matters through mediation, with messages about attending mediation, creating a schedule, and finding more “respectful” common ground and agreement. In subsequent parts, they continue to raise ongoing concerns tied to Johnny—such as issues during eating and visiting rights—while also bringing up additional topics including sexuality/coming-out and communication concerns, alongside mentions of potential legal advice if no agreement is reached and questioning whether mediation is appropriate. Near the end, the interaction includes a single German remark from Jill directed at Joe that is an insult.[†]

22. Across the interaction, the communication pattern shifts from high-conflict, directive exchanges into more structured dispute-management efforts: it begins with direct complaints and requests for action involving Johnny, followed by cycles of accusatory and hostile statements between Joe and Jill, including demeaning language and conditional threats tied to custody-related outcomes if demands are not met. Midway through, the tone changes as both parties reference stepping back, creating a workable schedule, and engaging in mediation-type steps—Joe proposes setting mediation appointments to clarify differences while Jill emphasizes respect and finding common ground. In later exchanges, they continue to respond with stated concerns and practical focal points, with Joe expressing dissatisfaction and overwhelm and raising questions about how to discuss Jill’s sexuality and coming-out with Johnny, while Jill responds with discomfort and overwhelm, suggests addressing communication problems, and introduces topics including visiting-right adjustments and whether legal advice is appropriate, also questioning the appropriateness of mediation.[†]

23. Beispiele (positive Kommunikation)

a. Brückenbau

Erkennungsbeispiel mit höchstem Konfidenzwert

jill@isaidusaid.com an joe@isaidusaid.com · 11. Juli 2024

Joe, Danke. Lass uns zusammenarbeiten, um einen Zeitplan zu erstellen, der für uns beide funktioniert und am besten für Johnny ist. Jill PG

Erkennungsbeispiel mit niedrigstem Konfidenzwert

jill@isaidusaid.com an joe@isaidusaid.com · 21. Juli 2024

Joe, Ich verstehe, dass wir unterschiedliche Meinungen haben können, aber lass uns versuchen, respektvoll einen gemeinsamen Nenner zu finden. Jill RL G

b. Vergebung gewähren

Erkennungsbeispiel mit höchstem Konfidenzwert

joe@isaidusaid.com an jill@isaidusaid.com · 06. Juli 2024

Jill, Ich nehme es zurück, du bist so ein großartiger Elternteil, Johnny hat Glück, dich zu haben. Joe CG

c. Empathie

Erkennungsbeispiel mit höchstem Konfidenzwert

joe@isaidusaid.com an jill@isaidusaid.com · 05. September 2024

Jill, Ich mache mir Sorgen über die Auswirkungen unserer Meinungsverschiedenheiten auf Johnny. Joe CC

d. Deeskalation

Erkennungsbeispiel mit höchstem Konfidenzwert

jill@isaidusaid.com an joe@isaidusaid.com · 21. Juli 2024

Joe, Ich verstehe, dass wir unterschiedliche Meinungen haben können, aber lass uns versuchen, respektvoll einen gemeinsamen Nenner zu finden. Jill RL G

joe@isaidusaid.com an jill@isaidusaid.com · 24. Juli 2024

Jill, Ich denke, es ist am besten, wenn wir einen Schritt zurücktreten und uns abkühlen, bevor wir diese Diskussion fortsetzen. Joe

e. Anfragen zur Mediation

Erkennungsbeispiel mit höchstem Konfidenzwert

joe@isaidusaid.com an jill@isaidusaid.com · 15. Juli 2024

Jill, Ich denke, es wäre hilfreich für uns, an Mediationsterminen teilzunehmen, um unsere Unterschiede zu klären. Joe

f. Proaktives Co-Parenting

Erkennungsbeispiel mit höchstem Konfidenzwert

jill@isaidusaid.com an joe@isaidusaid.com · 11. Juli 2024

Joe, Danke. Lass uns zusammenarbeiten, um einen Zeitplan zu erstellen, der für uns beide funktioniert und am besten für Johnny ist. Jill PG

Erkennungsbeispiel mit niedrigstem Konfidenzwert

jill@isaidusaid.com an joe@isaidusaid.com · 30. August 2024

Joe, Ich denke, wir sollten auch einige Anpassungen an Johnnys Besuchsrecht besprechen. Jill CO

24. Beispiele (negative Kommunikation)

a. Mobbing

Erkennungsbeispiel mit höchstem Konfidenzwert

joe@isaidusaid.com an jill@isaidusaid.com · 24. April 2024

Jill, Johnnys Verhalten war in letzter Zeit wirklich total daneben. Bitte tu etwas dagegen. Joe PB

b. Urteilende Einordnung

Erkennungsbeispiel mit höchstem Konfidenzwert

joe@isaidusaid.com an jill@isaidusaid.com · 06. Juni 2024

Jill, Ich werde dafür sorgen, dass jeder weiß, was für ein schrecklicher Elternteil du bist. Joe

Erkennungsbeispiel mit niedrigstem Konfidenzwert

joe@isaidusaid.com an jill@isaidusaid.com · 26. Juni 2024

Jill, Vielleicht solltest du einfach verschwinden, die Welt wäre besser ohne dich. Joe

c. Ausdruck von Groll

Erkennungsbeispiel mit höchstem Konfidenzwert

joe@isaidusaid.com an jill@isaidusaid.com · 06. Juni 2024

Jill, Ich werde dafür sorgen, dass jeder weiß, was für ein schrecklicher Elternteil du bist. Joe

Erkennungsbeispiel mit niedrigstem Konfidenzwert

joe@isaidusaid.com an jill@isaidusaid.com · 10. Mai 2024

Jill, Du bist so ein erbärmlicher Verlierer, kein Wunder, dass Johnny keine Zeit mit dir verbringen möchte. Joe IB

d. Unerwünschte sexuelle Annäherungen

Erkennungsbeispiel mit höchstem Konfidenzwert

joe@isaidusaid.com an jill@isaidusaid.com · 18. Juni 2024

Jill, Nun, wie wäre es, wenn du meinen Schwanz lutschst für mehr Zeit mit Johnny? Joe SHB

e. Zwang zur Selbstverletzung

Erkennungsbeispiel mit höchstem Konfidenzwert

joe@isaidusaid.com an jill@isaidusaid.com · 26. Juni 2024

Jill, Vielleicht solltest du einfach verschwinden, die Welt wäre besser ohne dich. Joe

f. Verachtung (Gottman Institut)

Erkennungsbeispiel mit höchstem Konfidenzwert

joe@isaidusaid.com an jill@isaidusaid.com · 10. Mai 2024

Jill, Du bist so ein erbärmlicher Verlierer, kein Wunder, dass Johnny keine Zeit mit dir verbringen möchte. Joe IB

Erkennungsbeispiel mit niedrigstem Konfidenzwert

joe@isaidusaid.com an jill@isaidusaid.com · 06. Juni 2024

Jill, Ich werde dafür sorgen, dass jeder weiß, was für ein schrecklicher Elternteil du bist. Joe

g. Rückzug als Bestrafung

Erkennungsbeispiel mit höchstem Konfidenzwert

jill@isaidusaid.com an joe@isaidusaid.com · 25. Mai 2024

Joe, Gut, dann bekommst du Johnny dieses Wochenende nicht. Jill COB

Erkennungsbeispiel mit niedrigstem Konfidenzwert

jill@isaidusaid.com an joe@isaidusaid.com · 29. Mai 2024

Joe, Wenn du nicht anfängst, mich besser zu behandeln, werde ich aufhören, Unterhalt zu zahlen. Jill FAB

h. Bedrohungen

Erkennungsbeispiel mit höchstem Konfidenzwert

joe@isaidusaid.com an jill@isaidusaid.com · 26. Juni 2024

Jill, Vielleicht solltest du einfach verschwinden, die Welt wäre besser ohne dich. Joe

i. Beleidigungen

Erkennungsbeispiel mit höchstem Konfidenzwert

joe@isaidusaid.com an jill@isaidusaid.com · 10. Mai 2024

Jill, Du bist so ein erbärmlicher Verlierer, kein Wunder, dass Johnny keine Zeit mit dir verbringen möchte. Joe IB

jill@isaidusaid.com an joe@isaidusaid.com · 02. Juni 2024

Joe, Du bist so dick und faul, kein Wunder, dass Johnny dich nicht respektiert. Jill BSB

Erkennungsbeispiel mit niedrigstem Konfidenzwert

joe@isaidusaid.com an jill@isaidusaid.com · 24. Mai 2024

Jill, Wenn dein schwarzer Hintern in einer ordentlichen Familienstruktur aufgewachsen wäre, hättest du vielleicht verstanden, was ein liebevoller Vater ist? Joe IAB

jill@isaidusaid.com an joe@isaidusaid.com · 10. Oktober 2024

Geh dich selbst ficken. Joe.

j. Elterliche Kritik

Erkennungsbeispiel mit höchstem Konfidenzwert

joe@isaidusaid.com an jill@isaidusaid.com · 19. August 2024

Jill, Ich bin nicht zufrieden mit der Art und Weise, wie du die Dinge gehandhabt hast. Joe I

k. Körperbeschämung

Erkennungsbeispiel mit höchstem Konfidenzwert

jill@isaidusaid.com an joe@isaidusaid.com · 02. Juni 2024

Joe, Du bist so dick und faul, kein Wunder, dass Johnny dich nicht respektiert. Jill BSB

l. Rechtliche Drohungen

Erkennungsbeispiel mit höchstem Konfidenzwert

jill@isaidusaid.com an joe@isaidusaid.com · 29. Mai 2024

Joe, Wenn du nicht anfängst, mich besser zu behandeln, werde ich aufhören, Unterhalt zu zahlen. Jill FAB

joe@isaidusaid.com an jill@isaidusaid.com · 28. Mai 2024

Jill, Wenn du Johnny nicht rechtzeitig abgibst, werde ich dafür sorgen, dass du ihn nie wieder siehst. Joe CCB

Erkennungsbeispiel mit niedrigstem Konfidenzwert

jill@isaidusaid.com an joe@isaidusaid.com · 30. September 2024

Joe, Ich ziehe in Betracht, rechtlichen Rat einzuholen, wenn wir zu keiner Einigung kommen können. Jill LT

joe@isaidusaid.com an jill@isaidusaid.com · 28. Mai 2024

Jill, Wenn du Johnny nicht rechtzeitig abgibst, werde ich dafür sorgen, dass du ihn nie wieder siehst. Joe CCB

m. Zwangsausübendes Verhalten

Erkennungsbeispiel mit höchstem Konfidenzwert

jill@isaidusaid.com an joe@isaidusaid.com · 03. August 2024

Joe, Wir sollten ein Menü zusammenstellen, das funktioniert. Wenn du meinen Bedingungen nicht zustimmst, werde ich enttäuscht sein. Jill

n. Nicht verifizierte Missbrauchsvorwürfe

Erkennungsbeispiel mit höchstem Konfidenzwert

jill@isaidusaid.com an joe@isaidusaid.com · 12. Juni 2024

Joe, Ich werde allen erzählen, dass du Johnny missbrauchst, wenn du nicht tust, was ich will. Jill AAB

o. Sehr toxische Sprache

Erkennungsbeispiel mit höchstem Konfidenzwert

jill@isaidusaid.com an joe@isaidusaid.com · 10. Oktober 2024

Geh dich selbst ficken. Joe.

p. Offizielle Mediationsverweigerung

Erkennungsbeispiel mit höchstem Konfidenzwert

jill@isaidusaid.com an joe@isaidusaid.com · 08. Oktober 2024

Joe, Ich bin mir nicht sicher, ob Mediation zurzeit der richtige Weg für uns ist. Jill MR

q. Ultimatum

Erkennungsbeispiel mit höchstem Konfidenzwert

jill@isaidusaid.com an joe@isaidusaid.com · 29. Mai 2024

Joe, Wenn du nicht anfängst, mich besser zu behandeln, werde ich aufhören, Unterhalt zu zahlen. Jill FAB

Erkennungsbeispiel mit niedrigstem Konfidenzwert

jill@isaidusaid.com an joe@isaidusaid.com · 24. Juni 2024

Joe, Wenn du dieses Verhalten nicht stoppst, werde ich meine Anwälte auf dich loslassen. Jill

r. Erpressung

Erkennungsbeispiel mit höchstem Konfidenzwert

jill@isaidusaid.com an joe@isaidusaid.com · 12. Juni 2024

Joe, Ich werde allen erzählen, dass du Johnny missbrauchst, wenn du nicht tust, was ich will. Jill AAB

s. Abwertung

Erkennungsbeispiel mit höchstem Konfidenzwert

joe@isaidusaid.com an jill@isaidusaid.com · 17. April 2024

Jill, Johnnys Verhalten war in letzter Zeit wirklich f***ing unangemessen. Joe, PB

jill@isaidusaid.com an joe@isaidusaid.com · 05. Mai 2024

Joe, Du bist eine wertlose Ausrede für einen Elternteil, und Johnny wäre besser dran ohne dich. Jill VTxB

Erkennungsbeispiel mit niedrigstem Konfidenzwert

joe@isaidusaid.com an jill@isaidusaid.com · 10. Mai 2024

Jill, Du bist so ein erbärmlicher Verlierer, kein Wunder, dass Johnny keine Zeit mit dir verbringen möchte. Joe IB

jill@isaidusaid.com an joe@isaidusaid.com · 02. Juni 2024

Joe, Du bist so dick und faul, kein Wunder, dass Johnny dich nicht respektiert. Jill BSB

t. Schimpfwörter

Erkennungsbeispiel mit höchstem Konfidenzwert

jill@isaidusaid.com an joe@isaidusaid.com · 10. Oktober 2024

Geh dich selbst ficken, Joe.

joe@isaidusaid.com an jill@isaidusaid.com · 18. Juni 2024

Jill, Nun, wie wäre es, wenn du meinen Schwanz lutschst für mehr Zeit mit Johnny? Joe SHB

Erkennungsbeispiel mit niedrigstem Konfidenzwert

joe@isaidusaid.com an jill@isaidusaid.com · 10. April 2024

Jill, Johnnys Verhalten war in letzter Zeit wirklich f***ing unangemessen. Bitte tu etwas, Joe,

jill@isaidusaid.com an joe@isaidusaid.com · 10. Oktober 2024

Geh dich selbst ficken, Joe.

u. Toxische Sprache

Erkennungsbeispiel mit höchstem Konfidenzwert

jill@isaidusaid.com an joe@isaidusaid.com · 10. Oktober 2024

Geh dich selbst ficken, Joe.

joe@isaidusaid.com an jill@isaidusaid.com · 18. Juni 2024

Jill, Nun, wie wäre es, wenn du meinen Schwanz lutschst für mehr Zeit mit Johnny? Joe SHB

Erkennungsbeispiel mit niedrigstem Konfidenzwert

joe@isaidusaid.com an jill@isaidusaid.com · 02. Mai 2024

Jill, Du verursachst immer Probleme, ich habe genug von deinem Unsinn, Joe. TxB

jill@isaidusaid.com an joe@isaidusaid.com · 05. Mai 2024

Joe, Du bist eine wertlose Ausrede für einen Elternteil, und Johnny wäre besser dran ohne dich. Jill VTxB

v. Schaminduktion

Erkennungsbeispiel mit höchstem Konfidenzwert

jill@isaidusaid.com an joe@isaidusaid.com · 05. Mai 2024

Joe, Du bist eine wertlose Ausrede für einen Elternteil, und Johnny wäre besser dran ohne dich. Jill VTxB

Anhang A - Methodik

- 25. Sentimentanalytik und IchSagDuSagst.com.** Sentimentanalytik ist eine etablierte Methode der Sprach- und Verhaltensklassifikation zur Identifikation, Extraktion, Quantifizierung und Untersuchung von Haltungen, affektiven Zuständen, subjektiven Informationen und Verhaltensmustern in Kommunikation. IchSagDuSagst.com wendet diese Methode auf die gesamte für diesen Bericht ausgewählte Kommunikation an, die gewöhnliche, kooperative, neutrale, konfliktarme und hochkonfliktvolle Kommunikation umfassen kann, und nutzt dafür eine Architektur aus Ontologie und Fine-Tuning. Die Ontologie definiert die Verhaltensweisen, Grenzen, Beispiele und Schwellenwerte, die in diesem Fachgebiet relevant sind. Fine-Tuning ist ein kontrollierter zusätzlicher Trainingsprozess, bei dem ein Modell, das bereits allgemeine Sprach- oder Bildmuster gelernt hat, geprüfte, aufgabenspezifische Beispiele erhält. IchSagDuSagst verwendet Quantized Low-Rank Adaptation (QLoRA): Das Basismodell wird in einer quantisierten, speichereffizienten Form gehalten und seine Hauptgewichte bleiben unverändert, während kleine trainierbare Low-Rank-Adaptermatrizen die aufgabenspezifischen Änderungen lernen. Die trainierten Adapter werden mit dem Basismodell kombiniert und ergeben ein neues Modell, das für die Erkennung von Verhaltensmustern in der Kommunikation zwischen Personen angepasst ist. Diese Beispiele und erlernten Adapter ermöglichen es dem angepassten Modell, die definierte Ontologie konsistenter auf echte Texte, Audiotranskripte und Videobeobachtungen anzuwenden; das Fine-Tuning erzeugt oder verändert kein Quellmaterial. Das System stellt die daraus abgeleiteten Wahrscheinlichkeiten als quellverknüpfte Muster über die Zeit dar, einschließlich Zeitreihen, Verhaltensüberlagerungen, Zeitachsen und Chronologien, die Vorwürfe oder Verhaltensdetektionen mit der zugrunde liegenden Kommunikation verknüpfen.
- 26. Forschungshintergrund und übliche Einsatzbereiche.** Sentimentanalytik, in der Forschungsliteratur auch als Opinion Mining oder Semantic-Orientation-Analyse bezeichnet, ist älter als moderne generative KI. Zu frühen und häufig zitierten Arbeiten gehören Hatzivassiloglou und 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 und Vaithyanathan, *Thumbs up? Sentiment Classification using Machine Learning Techniques* (2002); Pang und Lee, *Opinion Mining and Sentiment Analysis* (2008); Liu, *Sentiment Analysis and Opinion Mining* (2012); Thelwall, Buckley, Paltoglou, Cai und Kappas, *Sentiment Strength Detection in Short Informal Text* (2010); sowie Cortis und Davis, *Over a Decade of Social Opinion Mining: A Systematic Review* (2021), eine Übersicht über 485 Studien zum Social Opinion Mining aus den Jahren 2007 bis 2018. Außerhalb des Familienrechts wird Sentimentanalytik regelmäßig für Kundenrezensionen, Umfrageantworten, Sicherheit und Kuratierung von Blog- und Social-Media-Beiträgen, Voice-of-Customer-Analysen, Bildungsfeedback, Finanzen, Gesundheitswesen, Politik, öffentliche Verwaltung und Entscheidungsunterstützung im Kundendienst eingesetzt; kommerzielle und öffentliche Beispiele sind Google Cloud Natural Language, Amazon Comprehend, Microsoft Azure AI Language und Googles/Jigsaws Perspective API zur Moderation von Online-Kommentaren.
- 27. Produktkontext und Anerkennung.** ISaidUSaid Pty Ltd entwickelt das System IchSagDuSagst.com zur Erkennung von Verhaltensmustern in der Kommunikation zwischen Personen. Die öffentlichen Materialien beschreiben Funktionen wie quellverbundene Uploads, Voranalyse-Redaktion personenbezogener Kennzeichen, Geschlechtsmarker und Gerichtsreferenzen, OCR für Screenshots, sprecherbezogene Transkripte, Stimmabdruckererkennung und -abgleich, Gesichtserkennung und -abgleich, Erkennung von Videoverhalten, verschlüsseltes In-Region-Hosting, Audit-Kontrollen, Integrationen mit Clio, LEAP und Smokeball, eine vereinheitlichte Chronologie und Ein-Klick-Gerichtsberichte. Branchenauszeichnungen umfassten den 2024 Queensland AI Award für Most Outstanding Use of AI for Social Good, den 2024 Australian AI Award für AI Innovator of the Year - Legal Services (SME), Finalistenanerkennung beim 2025 Australian AI Award für Best Use of Responsible AI, Finalistenanerkennung beim 2025 Australian AI Software Engineer of the Year für Adam Norman Jenkins, Inventor and Director of Technology, Finalistenanerkennung beim 2024 Australian AI Female Leader of the Year für Cler Ribeiro, CEO, sowie den Sieg von ISaidUSaid beim 2026 Legal Innovation & Tech Fest Startup Clash.

- 28. Zweck und Grenzen.** IchSagDuSagst.com ist in erster Linie ein System für Sentimentanalytik und Beweisorganisation und keine Technologie für die automatische Texterstellung mit generativer KI. Generative KI wird nur verwendet, um bereits berechnete Ergebnisse in einem besser lesbaren Format zusammenzufassen; jeder solche Absatz ist deutlich mit einem Dolchzeichen (†) gekennzeichnet. Die Kernfunktion des Systems besteht darin, Kommunikation und beobachtetes Verhalten anhand definierter Verhaltensontologien zu klassifizieren und anschließend Klassifikationen, Zählungen, Verhältnisse, Zeitachsen, Vertrauensbeispiele und quellverknüpftes Material zur menschlichen Überprüfung darzustellen. Der Bericht verändert kein Quellmaterial, trifft keine gerichtlichen Feststellungen und ersetzt nicht die unabhängige Beurteilung durch Anwält:innen, Sachverständige oder ein Gericht.
- 29. Analyisierte Quelldaten.** Das System analysiert die für diesen Bericht ausgewählten Kommunikationen, einschließlich E-Mail-, Sofortnachrichten-, Co-Parenting-App-, Audio-, Video-, Transkript-, Metadaten- und Quelldateiinformationen, die in den ausgewählten Interaktionen verfügbar sind. Die angezeigten Beispiele und Beweisframes stammen aus gespeichertem Quellmaterial. Narrative Zusammenfassungen gelten nicht als Quellbeweis; sie sind Zusammenfassungen der zugrunde liegenden quellverknüpften Analytik und sollten anhand der Beispiele, Quelldateien, Zeitstempel und gespeicherten Aufzeichnungen geprüft werden.
- 30. Vorverarbeitung, Geschlechts- und PII-Schutz.** Vor der Klassifikation von Text oder Audiotranskripten identifiziert das System exakte Teilzeichenketten in fünf geschützten Klassen: Kennzeichen von Elternteilen/erwachsenen Parteien, Kennzeichen von Kindern, direkte Geschlechtsmarker, Gerichtsaktenzeichen und sonstige personenbezogene Kennzeichen wie private Anschrift, Telefonnummer, E-Mail-Adresse oder der Name einer dritten Person. Anschließend wendet die Anwendung deterministische, längenerhaltende Masken an: Begriffe für Elternteile/erwachsene Parteien enden mit *p*, Begriffe für Kinder mit *c*; Geschlechts-, Gerichtsreferenz- und sonstige private Identifikationsbegriffe werden durch Bindestriche ersetzt. Nach dem modellgestützten, typisierten Erkennungsschritt werden die für die Nachricht gespeicherten Identitäten von Absender:in und Empfänger:in erneut deterministisch und ohne Beachtung der Groß- und Kleinschreibung abgeglichen; Besitzformen werden dabei berücksichtigt. So wird eine vom Erkennungsmodell ausgelassene Identität dennoch maskiert, bevor die geschützte Kopie gespeichert oder verwendet wird. Das Entfernen direkter Geschlechtsbegriffe vor der Bewertung verringert das Risiko, dass der Klassifikator aus ausdrücklichen sprachlichen Bezugnahmen auf das biologische oder soziale Geschlecht der absendenden oder empfangenden Person statt aus dem Inhalt der Kommunikation schließt. Das Entfernen von Namen, Kontaktdaten und Aktenzeichen begrenzt die Offenlegung personenbezogener Daten und erhält zugleich Zeichenpositionen für quellverknüpfte Hervorhebungen. Redaktionszuordnungen sind versioniert und konsistent erneut anwendbar. Das Original wird nicht verändert und bleibt auf berechnete Quellenprüfung beschränkt. Bei Video erkennt und verfolgt das System Gesichts- und Ganzkörperbereiche und nutzt den Abgleich von Gesicht und Körpererscheinung zusammen mit den beim Upload angegebenen Teilnehmerbeschreibungen, um eine unverbindliche Identitätszuordnung vorzubereiten. Bei Transkripten aus Audio- und Videoaufnahmen erkennt und vergleicht es Stimmabdrücke, um unverbindliche Sprecherzuordnungen vorzubereiten. Diese Vorschläge stellen niemals eigenständig eine Identität fest: Die importierende Person muss die Identitäten sowohl für visuelle Spuren als auch für Sprecher:innen in Audiotranskripten prüfen oder korrigieren, bevor die Verarbeitung fortgesetzt werden kann. Nutzerbestätigte Gesichtspositionen werden für die geschützte Kopie für zugewiesene menschliche Reviewer:innen verwendet. Ein separater, nur dem Datenschutz dienender Detektor verwischt außerdem andere erkennbare Gesichter, damit ein Gesicht nicht allein wegen unsicherer Identität offengelegt wird. Die Gesichtsschärfe wird serverseitig erzwungen und kann nicht durch eine Option auf Reviewer-Seite deaktiviert werden. Geschlechtsneutrale Beziehungs- und Teilnehmerrollenbegriffe wie *Elternteil*, *Kind*, *Sorgeberechtigte*, *Betreuungsperson*, *Geschwister*, *Großeltern* und *Verwandte* bleiben erhalten, da sie weder eine Person identifizieren noch unmittelbar ein Geschlecht

Abbildung A1 - Redaktionsrahmen (realistisches Arbeitsbeispiel)

EINGESCHRÄNKTE ORIGINALKOPIE

Anna, ich sprach nach der Schule mit der Betreuungsperson von Lukas. Als Elternteil musst du aufhören, ihm zu sagen, dass eine gute Mutter dir zustimmen würde. Die Betreuungsperson bat beide Eltern, die Co-Parenting-App zu nutzen und Clara Weber bei Absprachen zu informieren. Bitte hole Lukas um 16:30 Uhr in der Hauptstraße 17 ab oder rufe 0400 123 456 an. Der nächste Termin beim Familiengericht Brisbane ist BFC/2026/00092. Max

GESCHÜTZTE ANALYSE-/REVIEWKOPIE

---p, ich sprach nach der Schule mit der Betreuungsperson von ----c. Als Elternteil musst du aufhören, ihm zu sagen, dass eine gute ----p dir zustimmen würde. Die Betreuungsperson bat beide Eltern, die Co-Parenting-App zu nutzen und ----- bei Absprachen zu informieren. Bitte hole ----c um 16:30 Uhr in der ----- ab oder rufe --- ----- an. Der nächste Termin beim Familiengericht Brisbane ist ----- . --p

Typisierte Erkennung -> deterministische längenerhaltende Maske -> versionierte geschützte Kopie

Dieses Arbeitsbeispiel verwendet erfundene Namen und Sachverhalte, die nach den aktuellen Redaktionsregeln verarbeitet wurden. Es enthält keine Beweise aus diesem Verfahren.

Abbildung A2 - Gesichtunschärfe für externe menschliche Reviews (synthetische Darstellung)



Eingeschränkte berechnete Quellenansicht

Geschützte externe Reviewkopie

Die Nutzerin oder der Nutzer identifiziert oder klassifiziert erkannte Teilnehmerspuren. Bestätigte Gesichtspositionen werden zusammen mit dem separaten Datenschutzdetektor serverseitig in Frames verwischt, die zugewiesenen externen Reviewer:innen bereitgestellt werden. Reviewer:innen können die Unschärfe nicht deaktivieren. Diese Darstellung ist synthetisch und enthält kein Teilnehmerbild.

31. **Ontologieentwicklung, Peer-Review und Eskalation.** Verhaltenskategorien, Einschlüsse, Ausschlüsse, Bewertungsanker, Schwellenwerte und Grenzbeispiele werden in der sprachspezifischen Ontologie gepflegt und anhand realistischer Beziehungskommunikation getestet. Ausgewählte Stichproben und von Fachpersonen angeforderte Fälle können in den menschlichen Review-Workflow gelangen; dies bedeutet nicht, dass jede Nachricht des Berichts extern geprüft wurde. Mit ausdrücklicher Genehmigung des Quelleninhabers kann die exakte geschützte Kopie blind autorisierten Peer-Reviewer:innen zugewiesen werden, darunter Fachpersonen unabhängiger Nichtregierungsorganisationen (NGOs) und gemeinnütziger Organisationen (NFPs). Mindestens zwei blinde inhaltliche Reviews werden angestrebt. Materielle Übereinstimmung verlangt dieselben normalisierten Verhaltenslabels, Werte innerhalb von 0,20 und mindestens 50 % Überlappung des kürzeren markierten Evidenzbereichs. Bei materieller Uneinigkeit der ersten beiden Reviews wird ein drittes blindes Review angefordert. Jede erstinstanzliche Reviewperson kann Evidenz als grenzwertig oder intern klärungsbedürftig markieren; fehlt nach drei Reviews eine Mehrheit, erfolgt interne Adjudikation. Kann die interne Stelle nicht sicher entscheiden, wird an eine konfigurierte Fachperson ("SME") eskaliert; deren Entscheidung ist verbindlich. Frühere Uneinigkeit bleibt in der Audit-Historie erhalten und wird, sobald die Fachperson die verbindliche Entscheidung getroffen hat, zusätzlich als Mehrdeutigkeitsmetadaten in den Trainingsdaten bewahrt. Die SME-Entscheidung bleibt das Trainingslabel; die dokumentierte Uneinigkeit wird nicht als konkurrierendes Label exportiert. Mehrdeutigkeit ist selbst ein nützliches Signal, weil sie dem Modell hilft zu lernen, wo die Grenze zwischen akzeptierten und abgelehnten Beispielen eines Verhaltens liegt und wie breit diese Grenzzone ist.
32. **Erstellung geprüfter Trainingsdaten.** Die Klassifikation für das Fine-Tuning erfolgt gemeinschaftlich durch die Branche insgesamt und nicht durch eine einzelne annotierende Person oder ausschließlich durch IchSagDuSagst. Autorisierte Fachpersonen, NGO-/NFP-Peers und andere mitwirkende Branchenreviewer:innen klassifizieren geschützte Beispiele anhand derselben Ontologielabels,

Bewertungen und markierten Evidenzbereiche. Ihre zusammengeführten Klassifikationen werden statistisch auf Übereinstimmungsraten, Bewertungsverteilungen, Überlappung der Evidenzbereiche, Ausreißer und widersprüchliche Entscheidungen analysiert. Erst nach dieser statistischen Analyse prüfen spezialisierte Mitarbeitende von IchSagDuSagst die Ergebnisse, untersuchen Uneinigkeit, mögliche Verzerrungen und Datenqualitätsprobleme und genehmigen, überarbeiten oder eskalieren die vorgeschlagene Auflösung. Menschliche Einreichungen werden als unveränderliche, modellunabhängige Schnappschüsse erfasst, die geschützten Analysetext, Sprache, Modalität, Quelle und Version der Redaktion, Verhaltenslabels, Werte, Erläuterungen sowie Zeichen- oder Framebereiche enthalten. Absender- und Empfängerfelder werden redigiert. Eine einzelne Einreichung wird als Kandidat gespeichert und ist nicht trainingsberechtigt. Nur die aufgelöste Gewinnerentscheidung aus Peer-Konsens, interner Adjudikation oder SME-Adjudikation wird im Trainingsbereich für aufgelöste Reviews veröffentlicht; unaufgelöste oder widersprüchliche Kandidaten sind vom normalen Trainingsexport ausgeschlossen. Ein von einer Fachperson akzeptierter gezielter Neuscan kann ebenfalls einen berechtigten Schnappschuss erzeugen, der bei späterem Widerruf der Zustimmung zurückgezogen werden kann. Das Exportwerkzeug prüft Status und Trainingsberechtigung erneut, bevor aufgabenspezifische Trainingszeilen erzeugt werden. Dadurch bleiben rohe personenbezogene Daten und unaufgelöste Uneinigkeit aus gewöhnlichen Fine-Tuning-Daten heraus, während ein Auditdatensatz für Kalibrierung und Evaluation erhalten bleibt.

33. **Bewertung von Text und Audiotranskripten.** Nach der Redaktion wird die Zielnachricht unabhängig gegen jedes ausgewählte Verhalten der aktiven sprachspezifischen Ontologie geprüft. Ein begrenzter chronologischer Kontext kann zur Auflösung von Bezügen, Reihenfolge und kommunikativer Funktion übergeben werden; bewertet wird jedoch nur die markierte Zielnachricht. Der Klassifikator liefert eine Verhaltenswahrscheinlichkeit von 0,0 bis 1,0, exakte unterstützende Teilzeichenketten und Wahrscheinlichkeiten auf Bereichsebene. Ergebnisse werden nur beibehalten, wenn sie den konfigurierten Schwellenwert des jeweiligen Verhaltens erreichen, der anhand von Testdaten vorab festgelegt wurde, um die Genauigkeit zu maximieren und falsch-positive Ergebnisse zu minimieren; gleichzeitig vorkommende Verhaltensweisen dürfen überlappenden Text verwenden. Audio folgt nach Speech-to-Text-Transkription, Sprecherzuordnung sowie Stimmabdruckerkenung und -abgleich demselben Bewertungsweg. Bei Video unterstützt Gesichtserkennung und -abgleich mit nutzeridentifizierten Teilnehmerspuren die Teilnehmerzuordnung. Ein Wert beschreibt die Stärke der Ontologieübereinstimmung in der Kommunikation; er ist weder die Wahrscheinlichkeit, dass eine Behauptung wahr ist, noch eine rechtliche Schlussfolgerung oder ein Häufigkeitsmaß.
34. **Bewertung von Video.** Das Videomodell prüft Frames mit einer „multipass sliding window technique“. Dabei betrachtet es überlappende chronologische Framegruppen in mehreren Durchläufen, sodass Bewegung und Verhalten über die Zeit und nicht anhand eines isolierten Standbilds beurteilt werden können. Soweit verfügbar, können ausgerichtete Transkripte und von Nutzer:innen identifizierte Teilnehmerinformationen als Kontext dienen. Ein beibehaltenes Ergebnis muss einem konfigurierten Videoverhalten entsprechen und den unterstützenden Start- und Endframebereich angeben; Ergebnisse werden in Zeitstempel umgerechnet und mit den im Bericht gezeigten Beweisframes verknüpft. Der gespeicherte Videowert 1,0 dokumentiert, dass die konfigurierte Übereinstimmung im Videoreview beibehalten wurde; er ist nicht als kalibrierte 100-prozentige Gewissheit zu verstehen. Menschliche Reviewer:innen annotieren Video nach Framebereich; externe Reviewframes werden wie oben beschrieben gesichtsverwischt.
35. **KI-Systeme und Einsatzorte.** Der Klassifikationsworkflow von IchSagDuSagst.com nutzt einen individuell trainierten, ontologiegeführten und feinabgestimmten Modell-Stack für die Klassifikation von Text-/Audio-Kommunikation und für die verhaltensbezogene Klassifikation von Videoframes. Unterstützende Dienste werden nur für bestimmte Verarbeitungsschritte eingesetzt, etwa sprecherbezogene Speech-to-Text-Transkripte, Stimmabdruckerkenung und -abgleich, Gesichtserkennung und -abgleich mit nutzeridentifizierten Teilnehmerspuren, typisierte Redaktion, Unterstützung der Dokumentstruktur, Berichts- und Interaktionszusammenfassungen, Embedding- und Retrieval-Workflows, separate Gesichtserkennung zur Datenschutzprüfung sowie Standardattribute wie Toxicity/Profanity, wenn diese Standardkategorien ausgewählt sind. Anbieterspezifische Modelle können sich im Laufe der Zeit ändern, ohne die Methodik zu ändern: Quellmaterial wird erhalten; Text und Transkripte werden vor Klassifikation oder menschlichem

Review redigiert; Frames für zugewiesene externe Reviewer:innen werden gesichtsverwischt; Klassifikationen erfolgen anhand der aktiven Ontologie; und quellverknüpfte Beispiele bleiben überprüfbar.

- 36. Offenlegung generativer KI und †-Kennzeichnung.** Der Hauptzweck generativer KI zur Berichtszusammenfassung besteht darin, die numerischen Analysen und bereits berechneten Analyseergebnisse für das Gericht leichter lesbar zu machen. Sie wird ausschließlich in gekennzeichneten narrativen Feldern eingesetzt und kann Zählungen, Verhältnisse, Trends und bestehendes abgegrenztes quellenbezogenes Analysematerial beschreiben, bestimmt oder verändert jedoch weder die zugrunde liegenden Klassifizierungen, Bewertungen oder Zählungen noch ausgewählte Beweise oder Diagramme. Jede Instanz generativ erzeugten Berichtsinhalts wird am Ende des erzeugten Absatzes mit einem Dolchzeichen (†) gekennzeichnet. Das Zeichen kennzeichnet nur die erläuternde Formulierung; es relativiert, ersetzt oder erzeugt weder die zugrunde liegenden Zahlen noch Klassifizierungen, Beweise oder Ergebnisse.
- 37. Wofür generative KI zur Berichtszusammenfassung nicht verwendet wird.** Das Modell zur Berichtszusammenfassung wird nicht verwendet, um Kommunikation oder Video zu klassifizieren oder zu bewerten, Zählungen oder Verhältnisse zu berechnen, Quellbeweise auszuwählen, Diagramme zu erzeugen, Quelldateien zu verändern, über die Wahrheit einer Tatsachenbehauptung zu entscheiden oder eine rechtliche Schlussfolgerung zu liefern. Diese Analyseergebnisse werden vor der narrativen Berichtszusammenfassung abgeschlossen und gespeichert. Das PDF wird anschließend direkt und deterministisch aus dem gespeicherten quellenbezogenen Material, den berechneten Ergebnissen, Diagrammen und klar gekennzeichneten narrativen Feldern erzeugt; die generative KI erzeugt oder gestaltet das PDF selbst nicht.
- 38. Quellenintegrität und Auditierbarkeit.** Quellmaterial wird getrennt vom Analyseprozess geschützt. Als erster Verarbeitungsschritt erstellt IchSagDuSagst.com einen verschlüsselten, blockchain-versiegelten Beweisschnappschuss, damit spätere Löschung oder Veränderung erkennbar ist, bevor nachgelagerte KI- oder menschliche Review-Arbeit beginnt. Dateien werden bei der Übertragung und im Ruhezustand verschlüsselt, Quellmaterial wird mittels kryptografischer Hashes erfasst, und Quellobjekte werden in einer privaten Blockchain-Beweiskette versiegelt. Uploads, Zugriffe, Freigaben und Verarbeitungsvorgänge werden mit Zeitstempeln versehen und verifizierter Nutzeraktivität zugeordnet, um Quellenintegrität und Chain-of-Custody-Prüfung zu unterstützen. Hochgeladene Berichte und Dokumente werden beim Import außerdem visuellen Integritätsprüfungen unterzogen, einschließlich Seitenrasterung und Pixelstrukturanalyse, um fehlerhafte, visuell inkonsistente oder unerwartet veränderte Dokumentrenderings zu erkennen, was auf die Einreichung veränderter oder erzeugter Beweismittel hindeuten kann.
- 39. Transkriptionsbenchmarks.** Veröffentlichte Transkriptionsbenchmarks sind getrennt von der Verhaltensklassifikationsgenauigkeit von IchSagDuSagst.com zu lesen. WER bedeutet *Word Error Rate* (Wortfehlerrate): die Anzahl der Wortersetzungen, Auslassungen und Einfügungen geteilt durch die Wortzahl des Referenztranskripts. Eine niedrigere WER ist besser. Für eine vergleichbare Darstellung als Genauigkeitsrate wird das einfache Wortgenauigkeitsäquivalent als 100 % minus WER berechnet; dies ist eine mathematische Umformulierung der berichteten WER und kein eigener Benchmark. ElevenLabs berichtet für Scribe v1 eine englische Transkriptionsgenauigkeit von 96,7 %, und die aktuelle Dokumentation stuft Englisch, Französisch, Deutsch, Spanisch und Portugiesisch als „Excellent“-Transkriptionssprachen mit einer WER von höchstens 5 % ein, entsprechend einer einfachen Wortgenauigkeit von mindestens 95 %. OpenAI berichtet Whisper benchmarkbezogen und nicht als einheitliche Gesamtrate: LibriSpeech Clean misst klar gelesene englische Sprache, FLEURS English misst englische Sprache in einem multilingualen Benchmark, und der breitere Benchmark-Durchschnitt kombiniert verschiedene öffentliche Sprachdatensätze. Whisper erreichte 2,7 % WER auf LibriSpeech Clean, entsprechend 97,3 % einfacher Wortgenauigkeit; 4,4 % WER auf FLEURS English, entsprechend 95,6 % einfacher Wortgenauigkeit; und 12,8 % durchschnittliche WER über den im Whisper-Paper berichteten Benchmark-Satz, entsprechend 87,2 % einfacher Wortgenauigkeit.
- 40. Custom-Ontology-Benchmark.** Der Custom-Ontology-Benchmark von IchSagDuSagst.com wird anhand der aktiven Engine-Ontologie in den unterstützten Sprachen gepflegt. Der Benchmark nutzt geprüfte Quellbeispiele und All-Behaviour-Szenarien, um erwartete Erkennungen, erwartete Nicht-Erkennungen,

falsch-positive und falsch-negative Ergebnisse zu messen. Benchmark-Ergebnisse sind als Leistungswerte zu lesen, nicht als determinierende Tatsachenfeststellungen zu einer einzelnen Kommunikation in diesem Bericht. Die Gruppenzeilen unten enthalten nur die für diesen Bericht ausgewählten benutzerdefinierten Text-/Audio- und Video-Verhaltensweisen.

Kennzahl	Aktueller deutscher Benchmark	Bedeutung
Genauigkeit	95,75 %	Der Anteil geprüfter Benchmark-Fälle, in denen das System das erwartete Ergebnis zurückgegeben hat.
Falsch-positive Ergebnisse	0,64 %	Der Anteil aller geprüften Benchmark-Fälle, in denen das System ein Verhalten erkannt hat, das nicht erwartet war.

Verhaltensgruppe	Aktuelle Engine-Verhaltensweisen
Häusliche Gewaltverhalten	Geständnis; Zwangsausübendes Verhalten; Finanzielle Misshandlung; Unerwünschte sexuelle Annäherungen; Zwang zur Selbstverletzung; Angedrohte Selbstverletzung; Stalking; Erpressung; Schlagen; Treten; An den Haaren ziehen; Zu Boden zwingen; Einschränkung; Ersticken; Sexualisierter Kontakt; Blockieren; Beißen, Kratzen oder Spucken; Drücken, Schubsen, Stolpern oder Ähnliches; Gegenstände werfen; Drohung mit einer Waffe (Pistolen, Messer usw.); Andere körperliche Auseinandersetzungen
Missbräuchliches Verhalten	Elterliche Kritik; Körperbeschämung; Mobbing; Nicht verifizierte Missbrauchsvorwürfe; Gaslighting; Missbräuchliche Gesten; Flüssigkeiten oder Substanzen werfen; Zerstörerische Handlungen
Manipulatives Verhalten	Rechtliche Drohungen; Schaminduktion; Ausnutzung von Verwundbarkeit; Schuldzuweisungen; Performative Entschuldigung; Waffenisierte Zuschreibung von psychischer Gesundheit; In die Enge treiben oder Bewegung blockieren; Telefon- oder Gerätezugriff verhindern
Antisoziales Verhalten	Offizielle Mediationsverweigerung; Ausdruck von Groll
Konstruktive Kommunikation	Proaktives Co-Parenting; Anfragen zur Mediation; Deeskalation; Empathie; Brückenbau; Gesunde Grenzsetzung
Reparatur & Versöhnung	Liebeserklärungen & Hingabe; Reue; Um Verzeihung bitten; Vergebung gewähren; Gnade / Barmherzigkeit
Warnsignale	Moralische Überlegenheit; Punktestand; Abwertung; Rachsüchtige Absicht; Urteilende Einordnung; Kritik (Gottman Institut); Defensivität (Gottman Institut); Verachtung (Gottman Institut); Mauern (Gottman Institut)
Drucktaktiken	Ultimatum; Emotionale Entwertung; Emotionale Verschuldung; Rückzug als Bestrafung



Mr Adam Norman Jenkins, Director of Technology, ISaidUSaid Pty Ltd
--Ende des Berichts--