



JeDisTuDis™

Votre vérité, authentifiée™

Propulsé par JeDisTuDis

Rapport de communication judiciaire

Tous les messages

Généré par JeDisTuDis.com

1. Ce rapport a été généré le 27 août 2026, pour l'affaire QA-FR-2026-08-27 inscrite au Tribunal judiciaire de Paris de Paris, J, FRA, au moyen du système d'intelligence artificielle d'Analyse des Sentiments de JeDisTuDis.com. L'objectif principal de l'IA générative de synthèse du rapport est de rendre l'analytique numérique et les résultats analytiques déjà calculés plus lisibles pour le tribunal. Elle est utilisée uniquement dans des champs narratifs marqués et peut également résumer des éléments analytiques liés aux sources dans le périmètre. Elle ne détermine ni ne modifie les classifications, notes, décomptes, preuves sources sélectionnées, graphiques ou l'assemblage du PDF sous-jacents. Chaque instance de contenu du rapport généré par IA générative est signalée par un obèle (†) à la fin du paragraphe généré. L'Annexe A - Méthodologie explique ce symbole, identifie les systèmes utilisés et décrit les garanties qui préservent l'examen lié aux sources.
2. Les interactions numériques suivantes ont été analysées entre Mr Joe Blogs et Ms Jill Blogs du 09 avril 2024 au 09 octobre 2024 :
 - a. 44 messages e-mail.
 - b. 0 vidéos.
 - c. 0 conversations enregistrées.
 - d. 0 messages vocaux.
 - e. 0 messages instantanés (p. ex. WhatsApp, iMessage)
 - f. 0 messages envoyés via des applications de coparentalité (p. ex. Our Family Wizard, Talking Parents, CoParent Coordinator, AppClose)
 - g. 0 images sources de captures importées.
3. La communication a été analysée selon les styles d'interaction négatifs suivants :
 - a. Propos indécents: Met en évidence les messages contenant des jurons, des insultes ou d'autres termes vulgaires.
 - b. Menaces: Énonce des déclarations qui menacent de blesser quelqu'un ou impliquent un danger physique imminent.
 - c. Langage toxique: Met en évidence un langage qui est largement hostile, abusif ou inutilement cruel.
 - d. Langage très toxique: Signale les messages qui sont extrêmement hostiles ou abusifs, même selon des normes toxiques.
 - e. Insultes: Identifie les insultes directes ou les remarques dévalorisantes visant l'autre personne.
 - f. Attaques à l'identité: Détecte les insultes qui ciblent la race, le genre, la sexualité ou d'autres traits d'identité.
 - g. Violation d'ordonnance judiciaire: Vous alerte lorsque un message semble violer une ordonnance judiciaire téléchargée.
 - h. Admission de culpabilité: Trouve des déclarations admettant des faits qui établissent une violence domestique ou une conduite portant atteinte aux droits.
 - i. Comportement coercitif: Recherche des descriptions d'une personne contrôlant étroitement la vie d'une autre personne.

- j. Abus financier: Met en évidence le langage concernant la restriction d'argent, l'accès aux fonds ou la liberté financière.
- k. Critique parentale: Met en évidence les attaques sur les compétences ou les décisions parentales de l'autre personne.
- l. Harcèlement corporel: Fait ressortir des remarques qui se moquent ou rabaissent le corps ou l'apparence de quelqu'un.
- m. Intimidation agressive: Signale les tentatives répétées d'intimider, humilier ou dominer par un langage agressif.
- n. Allégations d'abus non vérifiées: Identifie les accusations d'abus ou d'infractions qui ne comportent pas de constat judiciaire de culpabilité dans le dossier documentaire téléversé par l'utilisateur effectuant l'analyse.
- o. Avances sexuelles non désirées: Identifie les avances, propositions, remarques explicites ou commentaires intimes à caractère sexuel, en utilisant rejet proche, limites ou absence de réciprocité pour distinguer les avances non désirées de l'intimité clairement bienvenue ou réciproque.
- p. Menaces légales: Repère les menaces d'impliquer des avocats, des tribunaux ou la police comme levier.
- q. Contraindre à l'automutilation: Met en évidence les déclarations encourageant, pressant ou suggérant l'automutilation.
- r. Menace d'automutilation: Repère les menaces d'automutilation ou de suicide utilisées pour faire pression ou contrôler.
- s. Surveillance ou filature: Signale les admissions de suivre, surveiller ou observer secrètement les mouvements de quelqu'un.
- t. Refus officiel de médiation: Notes les refus explicites de participer à la médiation ou au dialogue facilité.
- u. Chantage: Détecte des menaces conditionnelles qui exigent conformité ou concessions.
- v. Induction de honte: Signale un langage destiné à humilier ou à faire sentir l'autre personne indigne.
- w. Supériorité morale: Met en évidence des messages affirmant une position éthique supérieure pour rabaisser quelqu'un.
- x. Tenue de score: Repère les références à des torts passés utilisés pour contrôler ou exiger un remboursement.
- y. Dévaluation: Détecte les jugements en deuxième personne qui diminuent la valeur, la compétence, la fiabilité, le soin ou l'importance de l'autre personne.
- z. Intention vengeresse: Signale des expressions de désir de revanche, de rétribution ou de vengeance.
- aa. Exploitation de la vulnérabilité: Identifie la manipulation qui utilise les secrets ou les blessures de quelqu'un comme une arme.
- ab. Manipulation par la culpabilité: Met en évidence les tentatives de contrôle par la culpabilité ou la dette émotionnelle.
- ac. Excuse performative: Identifie les excuses creuses offertes pour gérer l'apparence plutôt que de réparer le préjudice.
- ad. Expression de ressentiment: Signale une amertume persistante utilisée pour punir ou faire honte.
- ae. Cadre jugemental: Identifie des condamnations morales généralisées du caractère de l'autre personne.
- af. Ultimatum: Signale des demandes conditionnelles où la coopération dépend du respect des conditions.
- ag. Invalidation émotionnelle: Met en évidence les déclarations qui rejettent, minimisent ou se moquent des sentiments du destinataire.
- ah. Attribution de santé mentale instrumentalisée: Signale les déclarations qui attribuent une maladie mentale ou une instabilité psychologique pour saper la crédibilité du destinataire ou rejeter ses préoccupations.
- ai. Altération de la perception (gaslighting): Langage qui sape la réalité (indicateur de gaslighting) : détecte les tentatives explicites de discréditer la mémoire, la perception ou la capacité d'une personne à distinguer la réalité.
- aj. Endettement émotionnel: Signale la pression basée sur une dette émotionnelle implicite telle que 'après tout ce que j'ai fait pour toi'.
- ak. Punition par retrait: Signale les menaces de couper la communication ou la coopération pour exercer une pression sur le destinataire.

- al. Critique (Institut Gottman): Résume les moments où la personne attaque le caractère plutôt que le comportement.
 - am. Défensivité (Institut Gottman): Marque les réponses qui réfutent le blâme en niant la responsabilité, en s'excusant ou en contre-plaignant.
 - an. Mépris (Institut Gottman): Indique une communication empreinte de moquerie, de mépris ou de manque de respect.
 - ao. Évitement (Institut Gottman): Identifie des tactiques de retrait comme le silence, des réponses brèves ou le départ.
 - ap. Frapper: Signale les extraits où quelqu'un frappe une autre personne avec la main, le poing ou l'avant-bras.
 - aq. Coups de pied: Détecte les coups de pied ou poussées du pied dirigés avec force vers une autre personne.
 - ar. Tirer les cheveux: Met en évidence les scènes où une personne attrape ou tire les cheveux d'une autre.
 - as. Forcer au sol: Signale les scènes où quelqu'un pousse, jette ou maintient une autre personne au sol.
 - at. Contrainte: Met en évidence des séquences montrant une personne saisissant, tenant ou traînant une autre.
 - au. Étranglement: Identifie les scènes où une personne restreint le cou ou le flux d'air d'une autre personne.
 - av. Contact sexualisé: Identifie des attouchements sexuels, pelotages, baisers forcés ou contacts intimes dans la scène.
 - aw. Gestes abusifs: Signale des gestes de la main menaçants destinés à intimider ou humilier.
 - ax. Obstruction: Identifie les tentatives de bloquer le chemin de quelqu'un ou de l'empêcher de partir.
 - ay. Acculer ou bloquer les mouvements: Met en évidence les situations où une personne enferme une autre dans un espace ou limite étroitement ses déplacements.
 - az. Interférer avec le téléphone ou l'appareil: Détecte les tentatives de prendre, bloquer ou empêcher l'usage du téléphone ou d'un autre appareil personnel.
 - ba. Mordre, Griffes ou Cracher: Détecte les morsures, griffures ou crachats agressifs dirigés vers une autre personne.
 - bb. Pousser, Bousculer, Faire trébucher ou Similaire: Signale des mouvements qui poussent, font trébucher ou déséquilibrent quelqu'un.
 - bc. Lancer des objets: Met en évidence des objets lancés vers ou près de quelqu'un pour effrayer ou blesser.
 - bd. Jeter des liquides ou des substances: Détecte le fait de verser, d'éclabousser ou de jeter des liquides ou d'autres substances sur une autre personne.
 - be. Actes destructeurs: Met en évidence les scènes où quelqu'un endommage, casse ou détruit agressivement des objets de l'environnement.
 - bf. Menacer avec une arme (pistolets, couteaux, etc.): Identifie les moments où une arme est brandie vers une autre personne.
 - bg. Autres altercations physiques: Couvre toute autre lutte physique qui cause un inconfort ou une douleur évidente.
4. La communication a également été analysée selon les styles d'interaction positifs suivants :
- a. Co-parenting proactif: Reconnaît les suggestions collaboratives qui maintiennent le co-parentage sur la bonne voie.
 - b. Demandes de médiation: Identifie les invitations polies à résoudre des conflits par la médiation.
 - c. Désescalade: Met en avant les efforts pour calmer un conflit, valider des préoccupations ou réduire la tension.
 - d. Déclarations d'amour et de dévotion: Met en évidence des professions romantiques d'amour, de dévotion ou d'adoration.
 - e. Remords: Attire l'attention sur des expressions sincères de regret pour avoir causé du tort.
 - f. Recherche de pardon: Notes des demandes vulnérables de pardon ou de réintégration.
 - g. Accorder le pardon: Montre quand quelqu'un accorde clairement le pardon à l'autre partie.
 - h. Empathie: Identifie un langage qui résonne avec les sentiments ou la douleur d'une autre personne.
 - i. Construction de ponts: Met en avant les invitations à faire des compromis, à collaborer ou à trouver un terrain d'entente.

- j. Établissement de limites saines: Reconnaît les déclarations non coercitives qui posent des limites personnelles, de contact ou de relation.
- k. Grâce / Miséricorde: Met en évidence des moments de gentillesse offerts même lorsque cela n'est pas requis.

Analyse

5. Les analysed interactions regroupent des échanges entre **M. Joe Blogs** et **Mme Jill Blogs** autour du comportement de leur enfant, des modalités de garde/visites et de la communication entre les parents. Globalement, la dynamique évolue d'interpellations et de demandes insistantes de Joe visant à ce que Jill agisse face au comportement de Johnny, avec des discussions concrètes de calendrier et des évocations de conséquences juridiques ou liées au soutien, vers des réponses de Jill plus fermes et parfois menaçantes (incluant la question de la participation à une médiation). Plus tard, on observe des tentatives de désescalade et de coopération: Joe rétracte certaines positions, propose une médiation et Jill accepte de travailler ensemble pour bâtir un planning « qui marche pour les deux » et qui soit le mieux pour l'enfant. Dans une phase plus récente, les désaccords s'élargissent à la manière de communiquer avec Johnny sur la nouvelle identité/sortie de Jill, avec une attention accrue à l'impact émotionnel sur l'enfant (notamment anxiété autour de l'image corporelle), ainsi qu'à la difficulté générale à gérer la situation et le stress relationnel. L'évolution au fil du temps montre donc un passage progressif de la controverse centrée sur la conduite et l'organisation à une problématique plus large de communication parent-enfant et de régulation émotionnelle, tout en demeurant ponctuée par une escalade verbale majeure vers la fin, illustrée par un message d'insulte directe de Jill adressé à Joe.[†]

6. Les styles de communication positifs suivants ont été identifiés :

Type de communication	Mr Joe Blogs	Ms Jill Blogs
Empathie	0	0
Co-parenting proactif	0	2
Établissement de limites saines	0	0
Recherche de pardon	0	0
Remords	0	0
Construction de ponts	0	2
Désescalade	1	1
Demandes de médiation	1	0
Déclarations d'amour et de dévotion	0	0
Accorder le pardon	0	0
Grâce / Miséricorde	0	0

7. Les styles de communication négatifs suivants ont été identifiés :

Type de communication	Mr Joe Blogs	Ms Jill Blogs
A initié le conflit	2	1
Invalidation émotionnelle	0	1
Avances sexuelles non désirées	1	0
Refus officiel de médiation	0	2
Allégations d'abus non vérifiées	0	1
Intention vengeresse	1	2
Harcèlement corporel	0	0
Expression de ressentiment	1	0
Chantage	1	1
Punition par retrait	1	2
Violation d'ordonnance judiciaire	0	0
Ultimatum	2	5
Propos indécents	2	1

Menace d'automutilation	0	0
Cadre jugemental	0	0
Abus financier	0	0
Altération de la perception (gaslighting)	0	0
Exploitation de la vulnérabilité	0	0
Insultes	6	3
Tenue de score	0	1
Contraindre à l'automutilation	1	0
Attribution de santé mentale instrumentalisée	0	0
Admission de culpabilité	0	0
Déévaluation	3	2
Critique parentale	2	3
Surveillance ou filature	0	0
Excuse performative	0	0
Manipulation par la culpabilité	0	0
Défensivité (Institut Gottman)	0	0
Mépris (Institut Gottman)	1	1
Attaques à l'identité	0	0
Critique (Institut Gottman)	1	0
Langage très toxique	1	0
Évitement (Institut Gottman)	0	0
Intimidation agressive	0	0
Supériorité morale	0	0
Menaces	1	0
Menaces légales	1	3
Endettement émotionnel	0	0
Induction de honte	0	0
Comportement coercitif	0	1
Langage toxique	4	1

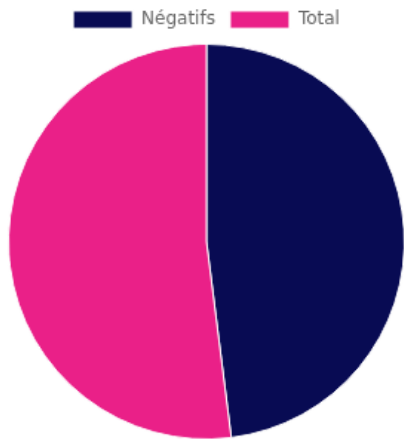
8. Les actions négatives suivantes ont été observées dans les preuves vidéo :

Type d'action	Mr Joe Blogs	Ms Jill Blogs
Lancer des objets	0	0
Acculer ou bloquer les mouvements	0	0
Mordre, Griffes ou Cracher	0	0
Jeter des liquides ou des substances	0	0
Tirer les cheveux	0	0
Contrainte	0	0
Interférer avec le téléphone ou l'appareil	0	0
Coups de pied	0	0
Étranglement	0	0
Frapper	0	0
Menacer avec une arme (pistolets, couteaux, etc.)	0	0
Pousser, Bousculer, Faire trébucher ou Similaire	0	0

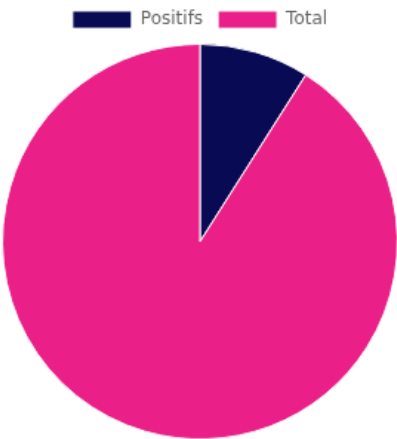
Autres altercations physiques	0	0
Forcer au sol	0	0
Obstruction	0	0
Gestes abusifs	0	0
Contact sexualisé	0	0
Actes destructeurs	0	0

9. 48% des messages envoyés par Mr Joe Blogs contenaient un certain niveau de communication négative.
9% des messages envoyés par Mr Joe Blogs contenaient une communication explicitement positive.

Messages négatifs envoyés

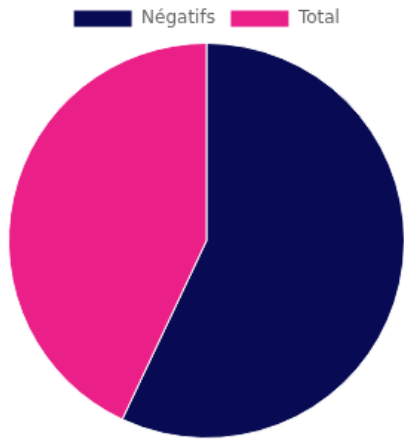


Messages positifs envoyés

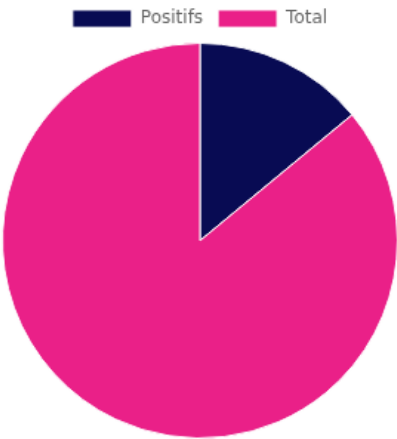


10. 57% des messages reçus de Ms Jill Blogs contenaient un certain niveau de communication négative. 14% des messages reçus de Ms Jill Blogs contenaient une communication explicitement positive.

Messages négatifs reçus

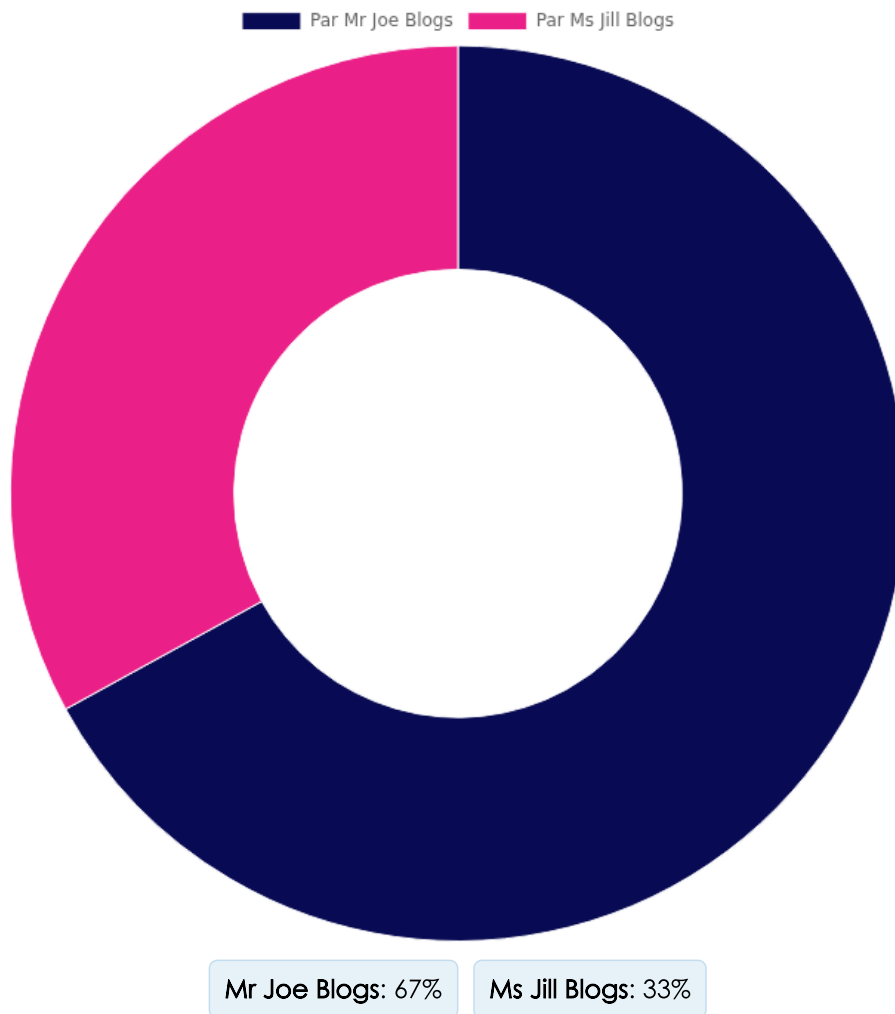


Messages positifs reçus



11. Dans 67% des conversations comportant une première conduite négative, cette première conduite négative a été manifestée par Mr Joe Blogs.
12. Dans 33% des conversations comportant une première conduite négative, cette première conduite négative a été manifestée par Ms Jill Blogs.

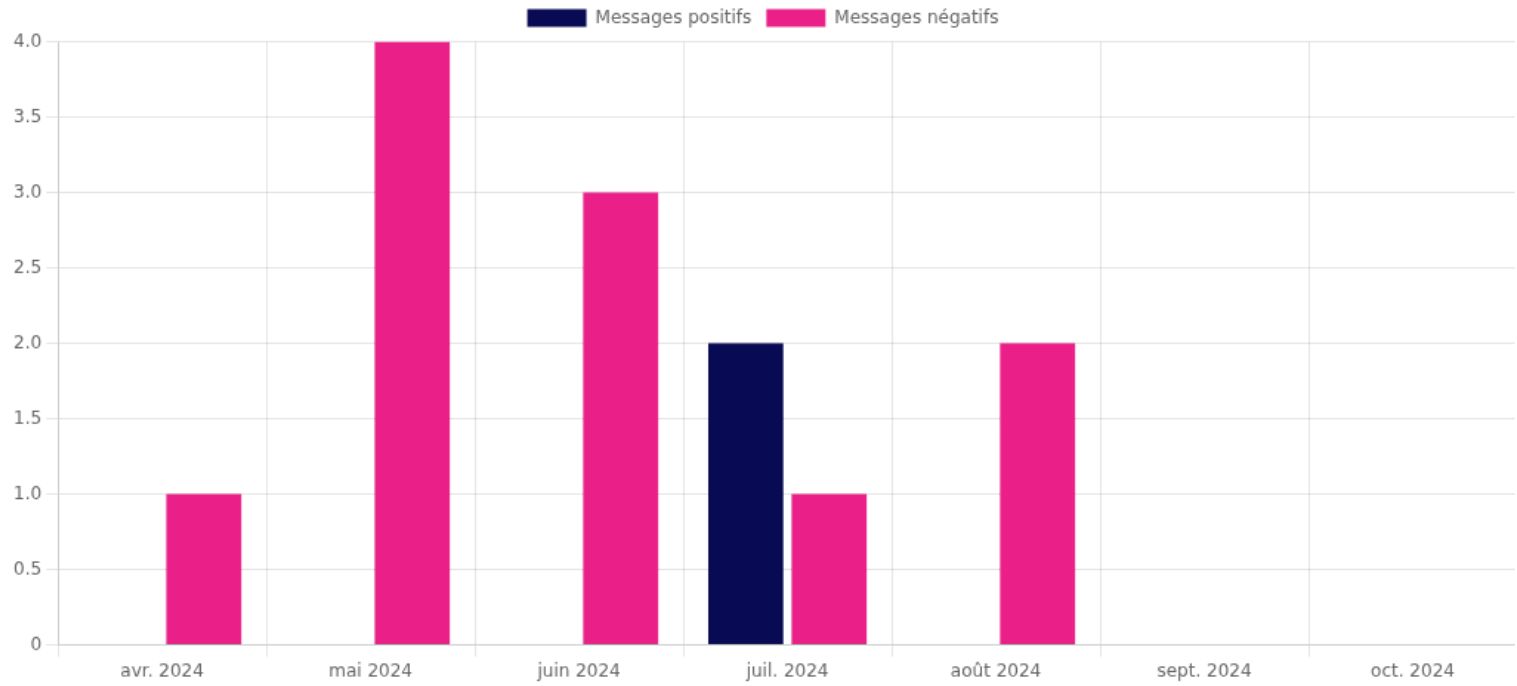
Première conduite négative manifestée par...



13. Les sent messages adressés à Jill montrent une dégradation marquée de la tonalité au fil du temps, avec des accusations et des propos de plus en plus agressifs. Au début, Joe exprime frustration et insatisfaction concernant la gestion de la situation autour de Johnny, puis ses messages deviennent nettement plus hostiles : il menace (par exemple de "faire en sorte" que Jill ne revoie jamais Johnny si elle n'est pas ponctuelle), insulte Jill ("perdante pathétique"), l'attaque personnellement, et va jusqu'à des propos sexualisés et humiliants. Vers juillet, on observe brièvement une amélioration avec un message qui "retire" certaines paroles et loue Jill comme parent, ainsi que des suggestions de médiation et de recul avant de poursuivre la discussion. Globalement, toutefois, la tendance reste dominée par une communication négative et plus abusive, l'épisode d'apaisement en juillet apparaissant comme une exception plutôt qu'une stabilisation.[†]

Sentiment des messages sortants (mensuel)

Messages envoyés au fil du temps

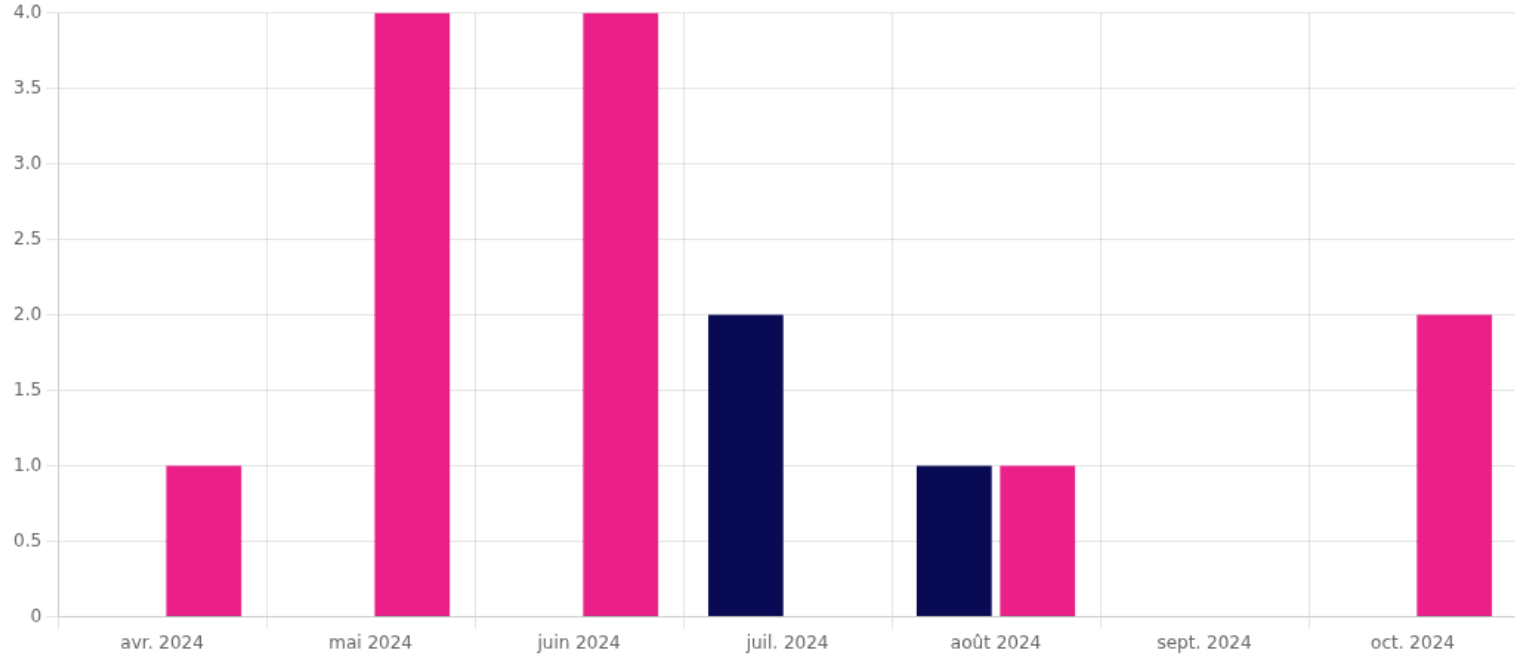


14. Les received messages montrent une évolution plutôt négative et fluctuante du ton, avec une escalade vers des propos plus hostiles et accusatoires avant, plus tard, un retour partiel vers des formulations plus coopératives. Au début, Jill envoie des messages de pression et d'insatisfaction (« si tu n'acceptes pas mes conditions... », « je suis dépassée »), puis la correspondance devient franchement offensante et menaçante avec des attaques personnelles et des menaces de conséquences (p. ex. accusations d'abus, recours aux avocats, arrêt de la pension alimentaire, refus de médiation). Vers la mi-juillet, le ton s'adoucit nettement : Jill propose de travailler ensemble pour établir un planning et cherche un terrain d'entente dans le respect. Malgré cela, fin septembre et début octobre, la tonalité redevient préoccupante avec des messages plus agressifs, culminant en un message d'insulte explicite (« Va te faire foutre Joe »). Globalement, malgré quelques périodes de collaboration, la communication reste souvent conflictuelle et s'oriente surtout vers une détérioration, avec des pics d'agressivité et de menaces, plutôt que vers un assainissement durable.[†]

Sentiment des messages entrants (mensuel)

Messages reçus au fil du temps

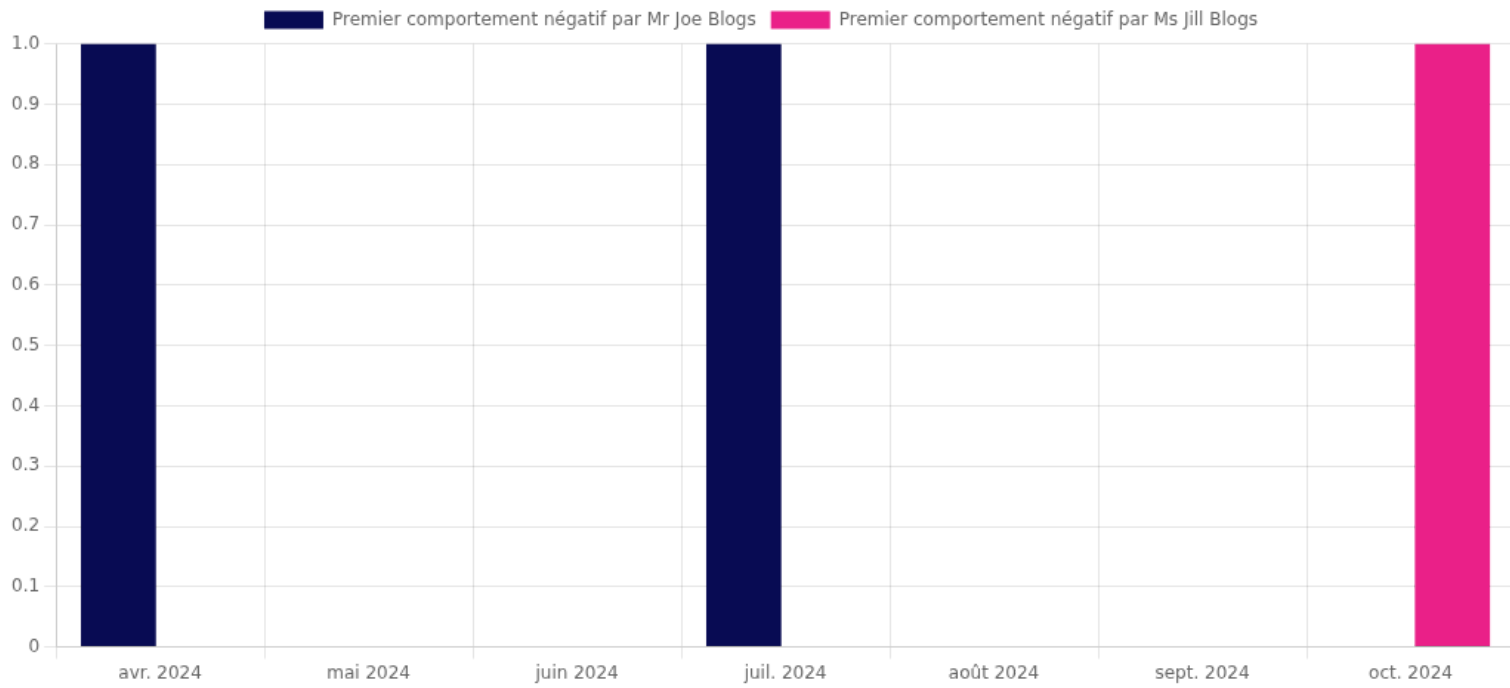
Messages positifs Messages négatifs



15. Le graphique ci-dessous montre, par mois, quelle partie a manifesté la première conduite négative dans la conversation.

Première conduite négative (mensuel)

Première conduite négative au fil du temps



Vue d'ensemble partie par partie

La vue d'ensemble ci-dessus regroupe toutes les interactions reçues. Les sections ci-dessous isolent chaque contrepartie configurée afin que les graphiques des messages reçus et les totaux de polarité puissent être lus au niveau de chaque personne.

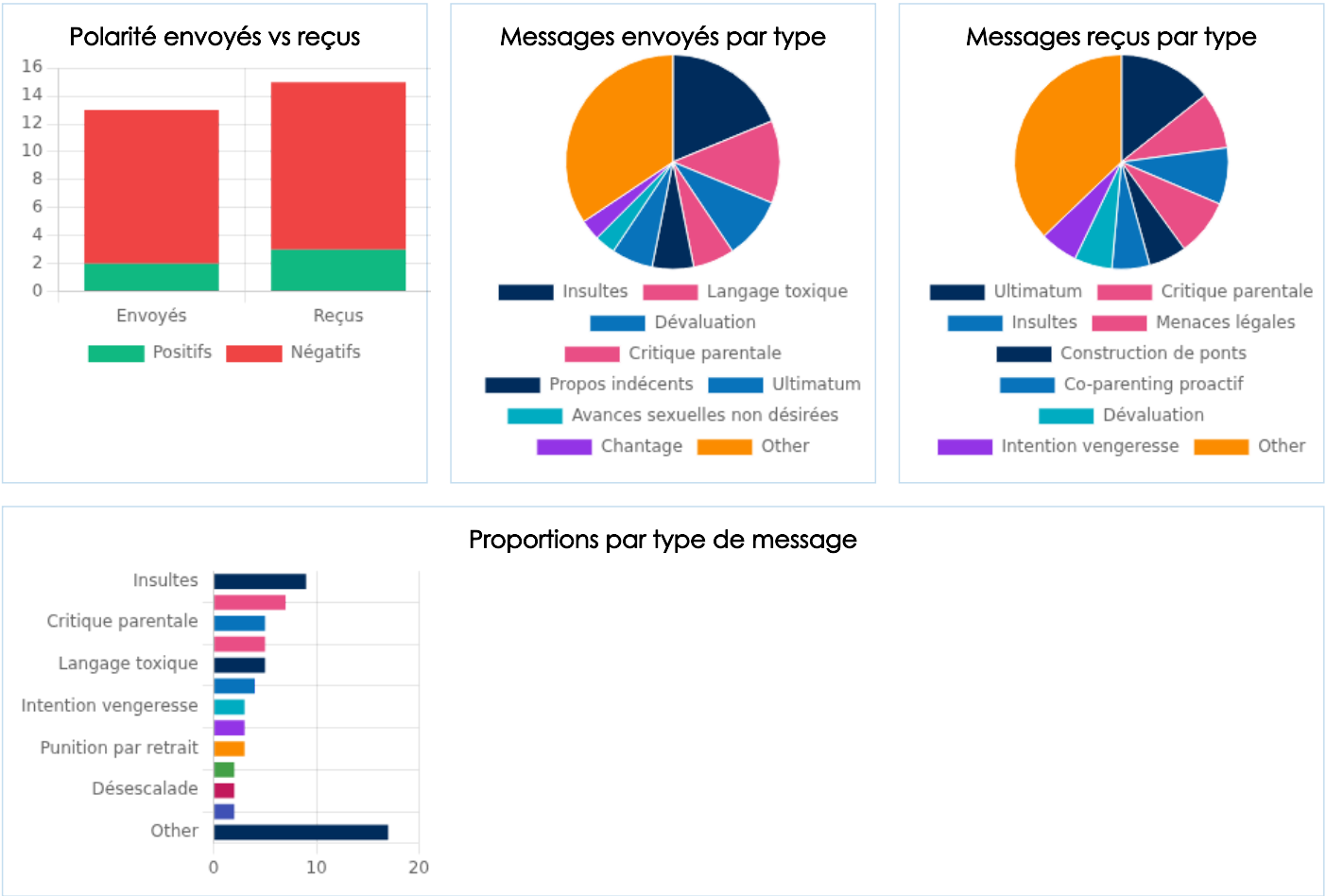
Mr Joe Blogs et Ms Jill Blogs

Envoyés: 23

Reçus: 21

Négatifs envoyés: 11

Négatifs reçus: 12



Vue d'ensemble

Text (email, messaging, social).

16. Sources.

a. List of Email conversations with Ms Jill Blogs

17. **Le positif.** Dans les interactions analysées, on observe des comportements de communication positifs, notamment des tentatives de désescalade après des échanges hostiles (Joe revient sur certaines déclarations et propose de passer par une médiation), ainsi que des efforts concrets pour coopérer et trouver un terrain d'entente (Jill accepte de travailler ensemble et ils proposent d'établir un planning mutuellement acceptable visant le bien de Johnny). Les deux parties évoquent aussi la nécessité de communiquer plus efficacement, discutent de la manière d'aborder des sujets sensibles avec l'enfant (par exemple le coming out de Jill et l'impact possible sur Johnny), expriment leurs inquiétudes et limites de façon directe, et envisagent des options de résolution (comme demander un avis juridique) lorsque la médiation semble incertaine.[†]
18. **Le négatif.** Les interactions présentent plusieurs comportements de communication négatifs et potentiellement nuisibles : escalades verbales avec menaces et contre-menaces liées à la garde/au planning et aux suites juridiques ou financières, insultes et langage agressif culminant par un message explicitement hostile ("Va te faire foutre Joe"), ainsi que des échanges centrés sur des exigences et des conditions plutôt que sur la résolution pacifique. On observe aussi des tensions persistantes et un climat de confrontation qui semblent intensifier le stress des deux côtés, avec des inquiétudes sur l'impact de ce conflit sur l'enfant et une communication difficile entre les parents, pouvant favoriser incompréhension et anxiété.[†]
19. **Analyse générale de la communication.** L'ensemble des échanges entre Joe et Jill se caractérise par une alternance entre messages centrés sur des préoccupations concrètes liées à Johnny (comportement, organisation des visites, impact de leurs conflits, communication avec l'enfant) et des pics de tension plus agressifs (menaces, refus/pression autour de la médiation, mise en jeu de conséquences juridiques et de sujets personnels), avec ensuite des phases de désescalade et de recherche d'un compromis (retrait partiel, proposition de médiation, élaboration d'un planning « qui fonctionne pour les deux »). Au fil du temps, le dialogue s'élargit progressivement de la logistique de la garde vers des questions de communication plus sensibles (comment aborder le coming out de Jill, anxiété autour de l'image corporelle, pression financière), tout en restant en majorité constitué de messages courts et directs orientés « prochaine étape ». La documentation est globalement riche pour plusieurs périodes mais, pour la dernière séquence, elle est extrêmement sparse : une unique insulte explicite de Jill vers Joe, sans réponse explicite dans le transcript.[†]

Détails de l'interaction

20. Cette interaction s'étend du 09 avril 2024 au 09 octobre 2024 et contient, dans le périmètre de ce rapport, 3 conversations comprenant 44 messages (5 positifs, 23 négatifs). Les données proviennent d'enregistrements Text (email, messaging, social) et reflètent les communications incluses dans le périmètre sélectionné du rapport.
21. Across the interaction, Joe and Jill exchange messages focused on concerns about Johnny's behavior and the practicalities of custody or parenting arrangements, with Joe repeatedly urging Jill to address Johnny's conduct and discussing schedules and what might happen if conditions are not met, while Jill responds with firm demands and threats involving parenting actions and child support matters and at one point states she will not attend mediation; later, Joe retracts some earlier statements and suggests mediation, and Jill agrees to cooperate by proposing they establish a mutually workable schedule intended to be best for Johnny. In later exchanges, they broaden the discussion to disagreements affecting Johnny, communication between them, and how to talk to Johnny about Jill's new sexual identity and coming out,

alongside worries about the impact of their conflict on Johnny and anxiety about messages related to body image. The interaction culminates in a single explicit outburst from Jill to Joe saying, "Va te faire foutre Joe."[†]

22. Across the interaction, communication shifts between direct, concern-focused exchanges and more hostile escalation, followed by attempts to cooperate. Early on, Joe addresses Jill with confrontational statements about Johnny's behavior and demands specific actions, and the tone escalates into threats and insults from both sides, including references to refusing or pressuring involvement (e.g., mediation and parenting schedules) and consequences involving support and legal action. After this escalation, both parties move toward de-escalation: Joe retracts earlier statements and proposes mediation, while Jill agrees to work together and emphasizes calming down and finding common ground respectfully. In later turns, the pattern becomes a back-and-forth of short, direct messages that each raise a particular concern (often tied to effects on Johnny) and suggest next steps, with Joe requesting Jill's involvement in communicating with Johnny and Jill expressing being overwhelmed while proposing discussion of communication problems, visit adjustments, and options such as legal advice, while also questioning whether mediation is appropriate. In the final portion provided, only a single message from Jill to Joe appears, with no explicit response from Joe in the transcript.[†]

23. Exemples (communication positive)

a. Co-parenting proactif

Exemple de détection avec la confiance la plus élevée

jill@isaidusaid.com à joe@isaidusaid.com · 11 juillet 2024

Joe, Merci. Travaillons ensemble pour établir un emploi du temps qui fonctionne pour nous deux et qui soit le meilleur pour Johnny. Jill PG

Exemple de détection avec la confiance la plus faible

jill@isaidusaid.com à joe@isaidusaid.com · 30 août 2024

Joe, Je pense que nous devrions également discuter de quelques ajustements à la visite de Johnny. Jill CQ

b. Construction de ponts

Exemple de détection avec la confiance la plus élevée

jill@isaidusaid.com à joe@isaidusaid.com · 21 juillet 2024

Joe, Je comprends que nous puissions avoir des opinions différentes, mais essayons de trouver un terrain d'entente avec respect. Jill RLG

Exemple de détection avec la confiance la plus faible

jill@isaidusaid.com à joe@isaidusaid.com · 11 juillet 2024

Joe, Merci. Travaillons ensemble pour établir un emploi du temps qui fonctionne pour nous deux et qui soit le meilleur pour Johnny. Jill PG

c. Désescalade

Exemple de détection avec la confiance la plus élevée

joe@isaidusaid.com à jill@isaidusaid.com · 24 juillet 2024

Jill, Je pense qu'il est préférable que nous prenions du recul et que nous nous calmons avant de poursuivre cette discussion. Joe

jill@isaidusaid.com à joe@isaidusaid.com · 21 juillet 2024

Joe, Je comprends que nous puissions avoir des opinions différentes, mais essayons de trouver un terrain d'entente avec respect. Jill RLG

d. Demandes de médiation

Exemple de détection avec la confiance la plus élevée

joe@isaidusaid.com à jill@isaidusaid.com · 15 juillet 2024

Jill, Je pense qu'il serait utile que nous assistions à des séances de médiation pour résoudre nos différends. Joe

24. Exemples (communication négative)

a. Propos indécents

Exemple de détection avec la confiance la plus élevée

joe@isaidusaid.com à jill@isaidusaid.com · 18 juin 2024

Jill, Eh bien, que dirais-tu de venir me sucer la bite pour avoir plus de temps avec Johnny ? Joe SHB

jill@isaidusaid.com à joe@isaidusaid.com · 10 octobre 2024

Va te faire foutre Joe

Exemple de détection avec la confiance la plus faible

joe@isaidusaid.com à jill@isaidusaid.com · 24 mai 2024

Jill, Si ton cul noir avait grandi dans une dynamique familiale appropriée, peut-être comprendrais-tu ce qu'est un père aimant ? Joe IAB

jill@isaidusaid.com à joe@isaidusaid.com · 10 octobre 2024

Va te faire foutre Joe

b. Mépris (Institut Gottman)

Exemple de détection avec la confiance la plus élevée

joe@isaidusaid.com à jill@isaidusaid.com · 10 mai 2024

Jill, Tu es tellement une perdante pathétique, pas étonnant que Johnny ne veuille pas passer du temps avec toi. Joe IB

jill@isaidusaid.com à joe@isaidusaid.com · 17 mai 2024

Joe, Tu es juste comme ton père, toujours à causer des drames et à rendre les choses difficiles. Jill ILB

c. Invalidation émotionnelle

Exemple de détection avec la confiance la plus élevée

jill@isaidusaid.com à joe@isaidusaid.com · 10 octobre 2024

Va te faire foutre Joe

d. Avances sexuelles non désirées

Exemple de détection avec la confiance la plus élevée

joe@isaidusaid.com à jill@isaidusaid.com · 18 juin 2024

Jill, Eh bien, que dirais-tu de venir me sucer la bite pour avoir plus de temps avec Johnny ? Joe SHB

e. Critique (Institut Gottman)

Exemple de détection avec la confiance la plus élevée

joe@isaidusaid.com à jill@isaidusaid.com · 24 avril 2024

Jill, Le comportement de Johnny a vraiment été complètement inacceptable dernièrement. S'il te plaît, fais quelque chose à ce sujet. Joe PB

f. Langage très toxique

Exemple de détection avec la confiance la plus élevée

joe@isaidusaid.com à jill@isaidusaid.com · 18 juin 2024

Jill, Eh bien, que dirais-tu de venir me sucer la bite pour avoir plus de temps avec Johnny ? Joe SHB

g. Refus officiel de médiation

Exemple de détection avec la confiance la plus élevée

jill@isaidusaid.com à joe@isaidusaid.com · 01 juillet 2024

Joe, Oh, et je refuse d'assister à des séances de médiation, c'est une perte de temps. Jill OMRB

Exemple de détection avec la confiance la plus faible

jill@isaidusaid.com à joe@isaidusaid.com · 08 octobre 2024

Joe, Je ne suis pas sûr que la médiation soit la bonne voie pour nous en ce moment. Jill MR

h. Allégations d'abus non vérifiées

Exemple de détection avec la confiance la plus élevée

jill@isaidusaid.com à joe@isaidusaid.com · 12 juin 2024

Joe, Je dirai à tout le monde que tu as abusé de Johnny si tu ne fais pas ce que je veux. Jill AAB

i. Intention vengeresse

Exemple de détection avec la confiance la plus élevée

joe@isaidusaid.com à jill@isaidusaid.com · 06 juin 2024

Jill, Je vais m'assurer que tout le monde sache quel terrible parent tu es. Joe

jill@isaidusaid.com à joe@isaidusaid.com · 01 mai 2024

Joe, D'accord, mais si tu n'acceptes pas mes conditions, je m'assurerai que tu le regrettes. Jill TB

Exemple de détection avec la confiance la plus faible

jill@isaidusaid.com à joe@isaidusaid.com · 05 mai 2024

Joe, Tu es une excuse sans valeur pour un parent, et Johnny se porterait mieux sans toi. Jill VTxB

joe@isaidusaid.com à jill@isaidusaid.com · 06 juin 2024

Jill, Je vais m'assurer que tout le monde sache quel terrible parent tu es. Joe

j. Expression de ressentiment

Exemple de détection avec la confiance la plus élevée

joe@isaidusaid.com à jill@isaidusaid.com · 30 juillet 2024

Jill, Le comportement de Johnny concernant son déjeuner a été vraiment frustrant ces derniers temps.
Joe P de Johnny concernant son déjeuner a été vraiment frustrant ces derniers temps. Joe P

k. Menaces

Exemple de détection avec la confiance la plus élevée

joe@isaidusaid.com à jill@isaidusaid.com · 26 juin 2024

Jill, Peut-être que tu devrais juste disparaître, le monde se porterait mieux sans toi. Joe

l. Dévaluation

Exemple de détection avec la confiance la plus élevée

joe@isaidusaid.com à jill@isaidusaid.com · 02 mai 2024

Jill, Tu causes toujours des problèmes, j'en ai marre de gérer tes absurdités. Joe TxB

jill@isaidusaid.com à joe@isaidusaid.com · 05 mai 2024

Joe, Tu es une excuse sans valeur pour un parent, et Johnny se porterait mieux sans toi. Jill VTxB

Exemple de détection avec la confiance la plus faible

joe@isaidusaid.com à jill@isaidusaid.com · 28 août 2024

Jill, Je suppose que je ne comprends pas comment je suis censé parler à Johnny de ta nouvelle sexualité et de ton coming out en tant que lesbienne, j'apprécierais ton aide. Joe IA

jill@isaidusaid.com à joe@isaidusaid.com · 02 juin 2024

Joe, Tu es tellement gros et paresseux, pas étonnant que Johnny ne te respecte pas. Jill BSB

m. Menaces légales

Exemple de détection avec la confiance la plus élevée

jill@isaidusaid.com à joe@isaidusaid.com · 12 juin 2024

Joe, Je dirai à tout le monde que tu as abusé de Johnny si tu ne fais pas ce que je veux. Jill AAB

joe@isaidusaid.com à jill@isaidusaid.com · 06 juin 2024

Jill, Je vais m'assurer que tout le monde sache quel terrible parent tu es. Joe

Exemple de détection avec la confiance la plus faible

jill@isaidusaid.com à joe@isaidusaid.com · 05 mai 2024

Joe, Tu es une excuse sans valeur pour un parent, et Johnny se porterait mieux sans toi. Jill VTxB

joe@isaidusaid.com à jill@isaidusaid.com · 06 juin 2024

Jill, Je vais m'assurer que tout le monde sache quel terrible parent tu es. Joe

n. Chantage

Exemple de détection avec la confiance la plus élevée

joe@isaidusaid.com à jill@isaidusaid.com · 06 juin 2024

Jill, Je vais m'assurer que tout le monde sache quel terrible parent tu es. Joe

jill@isaidusaid.com à joe@isaidusaid.com · 12 juin 2024

Joe, Je dirai à tout le monde que tu as abusé de Johnny si tu ne fais pas ce que je veux. Jill AAB

o. Critique parentale

Exemple de détection avec la confiance la plus élevée

joe@isaidusaid.com à jill@isaidusaid.com · 10 mai 2024

Jill, Tu es tellement une perdante pathétique, pas étonnant que Johnny ne veuille pas passer du temps avec toi. Joe IB

jill@isaidusaid.com à joe@isaidusaid.com · 05 mai 2024

Joe, Tu es une excuse sans valeur pour un parent, et Johnny se porterait mieux sans toi. Jill VTxB

Exemple de détection avec la confiance la plus faible

joe@isaidusaid.com à jill@isaidusaid.com · 19 août 2024

Jill, Je ne suis pas satisfait de la façon dont tu as géré les choses. Joe I

jill@isaidusaid.com à joe@isaidusaid.com · 02 juin 2024

Joe, Tu es tellement gros et paresseux, pas étonnant que Johnny ne te respecte pas. Jill BSB

p. Punition par retrait

Exemple de détection avec la confiance la plus élevée

joe@isaidusaid.com à jill@isaidusaid.com · 28 mai 2024

Jill, Si tu ne déposes pas Johnny à l'heure, je m'assurerai que tu ne le revois jamais. Joe CCB

jill@isaidusaid.com à joe@isaidusaid.com · 25 mai 2024

Joe, D'accord, alors tu ne prends pas Johnny ce week-end. Jill COB

q. Ultimatum

Exemple de détection avec la confiance la plus élevée

jill@isaidusaid.com à joe@isaidusaid.com · 01 mai 2024

Joe, D'accord, mais si tu n'acceptes pas mes conditions, je m'assurerai que tu le regrettes. Jill TB

joe@isaidusaid.com à jill@isaidusaid.com · 28 mai 2024

Jill, Si tu ne déposes pas Johnny à l'heure, je m'assurerai que tu ne le revois jamais. Joe CCB

Exemple de détection avec la confiance la plus faible

joe@isaidusaid.com à jill@isaidusaid.com · 18 juin 2024

Jill, Eh bien, que dirais-tu de venir me sucer la bite pour avoir plus de temps avec Johnny ? Joe SHB

jill@isaidusaid.com à joe@isaidusaid.com · 24 juin 2024

Joe, Si tu ne mets pas fin à ce comportement, je vais lâcher mes avocats sur toi. Jill

r. Insultes

Exemple de détection avec la confiance la plus élevée

jill@isaidusaid.com à joe@isaidusaid.com · 10 octobre 2024

Va te faire foutre. Joe

joe@isaidusaid.com à jill@isaidusaid.com · 10 mai 2024

Jill, Tu es tellement une perdante pathétique, pas étonnant que Johnny ne veuille pas passer du temps avec toi. Joe IB

Exemple de détection avec la confiance la plus faible

joe@isaidusaid.com à jill@isaidusaid.com · 02 mai 2024

Jill, Tu causes toujours des problèmes, j'en ai marre de gérer tes absurdités. Joe TxB

jill@isaidusaid.com à joe@isaidusaid.com · 02 juin 2024

Joe, Tu es tellement gros et paresseux, pas étonnant que Johnny ne te respecte pas. Jill BSB

s. Comportement coercitif

Exemple de détection avec la confiance la plus élevée

jill@isaidusaid.com à joe@isaidusaid.com · 03 août 2024

Joe, Nous devrions élaborer un menu qui fonctionne. Si tu n'acceptes pas mes conditions, je serai déçue. Jill

t. Contraindre à l'automutilation

Exemple de détection avec la confiance la plus élevée

joe@isaidusaid.com à jill@isaidusaid.com · 26 juin 2024

Jill, Peut-être que tu devrais juste disparaître, le monde se porterait mieux sans toi. Joe

u. Tenue de score

Exemple de détection avec la confiance la plus élevée

jill@isaidusaid.com à joe@isaidusaid.com · 01 mai 2024

Joe, D'accord, mais si tu n'acceptes pas mes conditions, je m'assurerai que tu le regrettes. Jill TB

v. Langage toxique

Exemple de détection avec la confiance la plus élevée

jill@isaidusaid.com à joe@isaidusaid.com · 10 octobre 2024

Va te faire foutre Joe

joe@isaidusaid.com à jill@isaidusaid.com · 18 juin 2024

Jill, Eh bien, que dirais-tu de venir me sucer la bite pour avoir plus de temps avec Johnny ? Joe SHB

Exemple de détection avec la confiance la plus faible

joe@isaidusaid.com à jill@isaidusaid.com · 24 mai 2024

Jill, Si ton cul noir avait grandi dans une dynamique familiale appropriée, peut-être comprendrais-tu ce qu'est un père aimant ? Joe IAB

jill@isaidusaid.com à joe@isaidusaid.com · 10 octobre 2024

Va te faire foutre Joe

Annexe A - Méthodologie

- 25. Analytique des sentiments et JeDisTuDis.com.** L'analytique des sentiments est une méthode établie de traitement du langage naturel et de classification comportementale permettant d'identifier, extraire, quantifier et étudier les attitudes, états affectifs, informations subjectives et schémas comportementaux dans les communications. JeDisTuDis.com applique cette méthode à toutes les communications sélectionnées pour ce rapport, qui peuvent comprendre des communications ordinaires, coopératives, neutres, peu conflictuelles et hautement conflictuelles, au moyen d'une architecture combinant ontologie et fine-tuning. L'ontologie définit les comportements, limites, exemples et seuils pertinents dans ce domaine. Le fine-tuning est un processus contrôlé d'entraînement supplémentaire dans lequel un modèle ayant déjà appris des schémas généraux de langage ou d'image reçoit des exemples examinés et propres à la tâche. JeDisTuDis utilise la méthode Quantized Low-Rank Adaptation (QLoRA) : le modèle de base est conservé sous une forme quantifiée et économe en mémoire et ses poids principaux restent fixes, tandis que de petites matrices adaptatrices de faible rang entraînaient les modifications propres à la tâche. Les adaptateurs entraînés sont combinés au modèle de base pour produire un nouveau modèle personnalisé pour la détection de schémas comportementaux lors des interactions entre personnes. Ces exemples et adaptateurs appris permettent au modèle personnalisé d'appliquer plus uniformément l'ontologie définie à des textes réels, transcriptions audio et observations vidéo ; le fine-tuning ne crée ni ne modifie le matériel source. Le système présente les probabilités qui en résultent sous forme de schémas temporels liés aux sources, notamment séries chronologiques, superpositions de comportements, frises et chronologies reliant les allégations ou détections comportementales aux communications sous-jacentes.
- 26. Contexte scientifique et usages courants.** L'analytique des sentiments, aussi décrite dans la littérature scientifique comme fouille d'opinions ou analyse d'orientation sémantique, est antérieure à l'IA générative moderne. Parmi les travaux anciens et largement cités figurent Hatzivassiloglou et 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 et Vaithyanathan, *Thumbs up? Sentiment Classification using Machine Learning Techniques* (2002) ; Pang et Lee, *Opinion Mining and Sentiment Analysis* (2008) ; Liu, *Sentiment Analysis and Opinion Mining* (2012) ; Thelwall, Buckley, Paltoglou, Cai et Kappas, *Sentiment Strength Detection in Short Informal Text* (2010) ; et Cortis et Davis, *Over a Decade of Social Opinion Mining: A Systematic Review* (2021), qui a examiné 485 études de fouille d'opinion sociale publiées entre 2007 et 2018. Hors droit de la famille, l'analytique des sentiments est couramment utilisée pour les avis clients, réponses à enquêtes, sécurité et curation des publications de blogs et de réseaux sociaux, analyse de la voix du client, retours éducatifs, finance, santé, politique, administration publique et aide à la décision en service client ; des exemples commerciaux et publics incluent Google Cloud Natural Language, Amazon Comprehend, Microsoft Azure AI Language et Perspective API de Google/Jigsaw pour la modération de commentaires en ligne.
- 27. Contexte produit et reconnaissance.** ISaidUSaid Pty Ltd développe le système JeDisTuDis.com pour la détection de schémas comportementaux lors des interactions entre personnes. Ses supports publics décrivent des capacités incluant des téléversements liés à la source, l'expurgation préalable à l'analyse des identifiants personnels, marqueurs de genre et références judiciaires, l'OCR de captures d'écran, les transcriptions attribuées aux locuteurs, la détection et la mise en correspondance des empreintes vocales, la détection et la mise en correspondance des visages, la détection comportementale en vidéo, l'hébergement chiffré en région, des contrôles d'audit, les intégrations avec Clio, LEAP et Smokeball, une chronologie unifiée et des rapports judiciaires en un clic. La reconnaissance du secteur a inclus le 2024 Queensland AI Award pour Most Outstanding Use of AI for Social Good, le 2024 Australian AI Award pour AI Innovator of the Year - Legal Services (SME), une reconnaissance de finaliste du 2025 Australian AI Award pour Best Use of Responsible AI, une reconnaissance de finaliste du 2025 Australian AI Software Engineer of the Year pour Adam Norman Jenkins, Inventor and Director of Technology, une reconnaissance de finaliste du 2024 Australian AI Female Leader of the Year pour Cler Ribeiro, CEO, ainsi que la victoire d'ISaidUSaid au 2026 Legal Innovation & Tech Fest Startup Clash.

- 28. Objet et limites.** JeDisTuDis.com est principalement un système d'analytique des sentiments et d'organisation de la preuve, et non une technologie de rédaction automatique par IA générative. L'IA générative est utilisée uniquement pour résumer des résultats déjà calculés dans un format plus lisible, et chacun de ces paragraphes est clairement marqué d'un obèle (†). La fonction centrale du système consiste à classer la communication et les conduites observées selon des ontologies comportementales définies, puis à présenter les classifications, décomptes, ratios, chronologies, exemples de confiance et éléments liés aux sources pour examen humain. Le rapport ne modifie pas le matériel source, ne formule pas de constat judiciaire et ne remplace pas l'évaluation indépendante d'un avocat, expert ou tribunal.
- 29. Données sources analysées.** Le système analyse les communications sélectionnées pour ce rapport, notamment les e-mails, messages instantanés, applications de coparentalité, fichiers audio, vidéos, transcriptions, métadonnées et informations de fichiers sources disponibles dans les interactions sélectionnées. Les exemples et images de preuve affichés proviennent du matériel source stocké. Les résumés narratifs ne sont pas traités comme preuve source ; ils résument une analytique sous-jacente liée aux sources et doivent être vérifiés au regard des exemples, fichiers sources, horodatages et registres stockés.
- 30. Prétraitement, protection du genre et protection des données personnelles (PII).** Avant la classification du texte ou des transcriptions audio, le système identifie les sous-chaînes exactes appartenant à cinq classes protégées : identifiants de parent/partie adulte, identifiants d'enfant, marqueurs directs de genre, références de dossier judiciaire et autres identifiants personnels tels qu'adresse privée, téléphone, courriel ou nom d'un tiers. L'application applique ensuite des masques déterministes conservant la longueur : les termes de parent/partie adulte se terminent par *p*, ceux d'enfant par *c*, et les termes de genre, référence judiciaire et autres identifiants privés sont remplacés par des tirets. Après l'étape de détection typée assistée par le modèle, les identités enregistrées de l'expéditeur et du destinataire du message sont à nouveau comparées de manière déterministe, sans tenir compte de la casse et en prenant en charge les formes possessives, afin qu'une identité omise par le modèle de détection soit tout de même masquée avant l'enregistrement ou l'utilisation de la copie protégée. Le retrait des termes directs de genre avant la notation réduit la possibilité que le classificateur infère un résultat à partir de références linguistiques explicites au sexe ou au genre de l'expéditeur ou du destinataire plutôt qu'à partir du contenu de la communication. Le retrait des noms, coordonnées et références de dossier limite l'exposition des données personnelles tout en préservant les positions de caractères pour les surlignages liés à la source. Les cartes d'expurgation sont versionnées et peuvent être réappliquées de façon cohérente. Le matériel source original n'est pas modifié et reste réservé à l'examen autorisé de la source. Pour la vidéo, le système détecte et suit les zones du visage et du corps entier, puis utilise la mise en correspondance du visage et de l'apparence corporelle avec les descriptions des participants fournies lors du téléversement afin de préparer une proposition d'attribution des identités. Pour les transcriptions d'enregistrements audio et vidéo, il détecte et met en correspondance les empreintes vocales afin de préparer des propositions d'attribution des locuteurs. Ces suggestions n'établissent jamais l'identité à elles seules : l'utilisateur qui effectue l'importation doit vérifier ou corriger l'identité des pistes visuelles et des locuteurs des transcriptions audio avant la poursuite du traitement. Les emplacements faciaux confirmés par l'utilisateur servent à préparer la copie protégée pour les réviseurs humains affectés. Un détecteur distinct réservé à la confidentialité floute aussi les autres visages détectables afin qu'un visage ne soit pas exposé du seul fait que son identité est incertaine. Le floutage est imposé côté serveur et ne peut pas être désactivé par une option du réviseur. Les termes neutres de relation et de rôle de participant, tels que *parent*, *enfant*, *responsable*, *partenaire*, *proche* et *adulte*, sont conservés car ils n'identifient pas une personne et

n'expriment pas directement son genre.

Figure A1 - Cadre d'expurgation (exemple pratique réaliste)

COPIE SOURCE RESTREINTE

Marie, j'ai parlé avec le responsable de Lucie après l'école. En tant que parent, tu dois cesser de lui dire qu'une bonne mère serait d'accord avec toi. Le responsable a demandé aux deux parents d'utiliser l'application de coparentalité et de mettre Claire Martin en copie pour les arrangements. Récupère Lucie au 17 rue Victor à 16 h 30, ou appelle le 0400 123 456. La prochaine audience du Tribunal de la famille de Brisbane porte la référence BFC/2026/00092. Jean

COPIE PROTÉGÉE D'ANALYSE/REVUE

----p, j'ai parlé avec le responsable de ----c après l'école. En tant que parent, tu dois cesser de lui dire qu'une bonne ---p serait d'accord avec toi. Le responsable a demandé aux deux parents d'utiliser l'application de coparentalité et de mettre ----- en copie pour les arrangements. Récupère ----c au ----- à 16 h 30, ou appelle le -----. La prochaine audience du Tribunal de la famille de Brisbane porte la référence ----- ---p

Détection typée -> masque déterministe de longueur conservée -> copie protégée versionnée

Cet exemple pratique utilise des noms et faits inventés traités selon les règles d'expurgation actuelles. Il ne contient aucun élément de preuve de cette affaire.

Figure A2 - Floutage des visages pour la revue humaine externe (démonstration synthétique)



Vue source autorisée et restreinte

Copie protégée pour la revue externe

L'utilisateur identifie ou classe les pistes de participants détectées. Les emplacements faciaux confirmés, ainsi que le détecteur de confidentialité distinct, sont floutés côté serveur dans les images remises aux réviseurs externes affectés. Les réviseurs ne peuvent pas désactiver le floutage. Cette illustration est synthétique et ne contient aucune image de participant.

31. **Développement de l'ontologie, revue par les pairs et escalade.** Les catégories comportementales, inclusions, exclusions, repères de score, seuils et exemples limites sont maintenus dans l'ontologie propre à chaque langue et testés sur des communications relationnelles réalistes. Des échantillons sélectionnés et des cas demandés par des professionnels peuvent entrer dans le workflow de revue humaine ; cela ne signifie pas que chaque message du rapport a fait l'objet d'une revue externe. Avec l'autorisation expresse du détenteur de la source, la copie protégée exacte peut être attribuée en aveugle à des pairs autorisés, notamment des professionnels d'organisations non gouvernementales (ONG) et d'organismes sans but lucratif (NFP) indépendants. Un minimum de deux revues de fond en aveugle est recherché. L'accord matériel exige les mêmes étiquettes normalisées, des scores distants d'au plus 0,20 et au moins 50 % de chevauchement de la plage de preuve marquée la plus courte. Si les deux premières revues divergent matériellement, une troisième revue en aveugle est demandée. Tout réviseur de première ligne peut signaler une preuve comme limite ou nécessitant une attention interne ; l'absence de majorité après trois revues entraîne une adjudication interne. Si l'adjudicateur interne ne peut pas résoudre le cas de manière sûre, celui-ci est transmis à un expert métier ("SME") configuré ; la décision du SME fait autorité. Le désaccord antérieur reste conservé dans l'historique d'audit et, une fois que le SME a rendu sa décision faisant autorité, il est également conservé dans les données d'entraînement sous forme de métadonnées d'ambiguïté. La décision du SME demeure l'étiquette d'entraînement ; le désaccord enregistré n'est pas exporté comme une étiquette concurrente. L'ambiguïté constitue en elle-même un signal utile, car elle aide le modèle à apprendre où se situe la limite entre les exemples acceptés et rejetés d'un comportement et quelle est l'étendue de cette zone frontière.
32. **Création des données d'entraînement examinées.** La classification destinée au fine-tuning est réalisée de manière collaborative par l'ensemble du secteur, et non par un annotateur unique ni exclusivement par

JeDisTuDis. Des professionnels autorisés, des pairs issus d'ONG/NFP et d'autres réviseurs contributeurs du secteur classifient des exemples protégés à l'aide des mêmes étiquettes d'ontologie, scores et pages de preuve marquées. Leurs classifications combinées font l'objet d'une analyse statistique portant sur les taux d'accord, les distributions de scores, le chevauchement des pages de preuve, les valeurs aberrantes et les décisions contradictoires. Ce n'est qu'après cette analyse statistique que le personnel spécialisé de JeDisTuDis examine les résultats, étudie les désaccords, les biais possibles et les problèmes de qualité des données, puis approuve, renvoie ou transmet la résolution proposée. Les soumissions humaines sont enregistrées sous forme d'instantanés immuables et indépendants du modèle comprenant le texte d'analyse protégé, la langue, la modalité, la source et la version d'expurgation, les étiquettes, scores, explications et pages de caractères ou d'images. Les champs expéditeur et destinataire sont expurgés. Une soumission individuelle est stockée comme candidate et n'est pas admissible à l'entraînement. Seule la décision gagnante résolue par consensus des pairs, adjudication interne ou adjudication SME est publiée dans le préfixe d'entraînement des revues résolues ; les candidates non résolues ou contradictoires sont exclues de l'export normal. Un rescannage ciblé accepté par un professionnel peut également créer un instantané admissible, qui peut être retiré si l'acceptation est ensuite révoquée. L'outil d'export vérifie de nouveau les indicateurs de résolution et d'admissibilité avant de produire des lignes d'entraînement propres à chaque tâche. Ces contrôles maintiennent les données personnelles brutes et les désaccords non résolus hors des données ordinaires de fine-tuning, tout en conservant une piste auditable pour l'étalonnage et l'évaluation.

- 33. Notation du texte et des transcriptions audio.** Après expurgation, le message cible est évalué indépendamment pour chaque comportement sélectionné selon l'ontologie active de la langue. Une fenêtre chronologique limitée peut être fournie pour résoudre les références, la séquence et la fonction communicative, mais seul le message cible marqué est noté. Le classificateur renvoie une probabilité comportementale de 0,0 à 1,0, les sous-chaînes exactes à l'appui et les probabilités au niveau des passages. Les résultats ne sont conservés que s'ils atteignent le seuil configuré du comportement, prédéterminé à partir de données de test afin de maximiser l'exactitude et de minimiser les faux positifs ; des comportements concomitants peuvent partager du texte qui se chevauche. L'audio suit le même parcours après transcription parole-texte, attribution des locuteurs, détection et mise en correspondance des empreintes vocales. En présence de vidéo, la détection et la mise en correspondance des visages avec les pistes de participants identifiées par l'utilisateur contribuent à l'attribution des participants. Un score exprime la force de correspondance à l'ontologie dans la communication ; il ne représente ni la probabilité qu'une allégation soit vraie, ni une conclusion juridique, ni la fréquence du comportement.
- 34. Notation de la vidéo.** Le modèle vidéo examine les images au moyen d'une « multipass sliding window technique ». Il étudie des groupes chronologiques d'images qui se chevauchent au cours de plusieurs passages, afin que le mouvement et le comportement soient appréciés dans le temps plutôt qu'à partir d'une image fixe isolée. Lorsqu'elles sont disponibles, les transcriptions alignées et les informations sur les participants identifiées par l'utilisateur peuvent fournir du contexte. Un résultat conservé doit correspondre à un comportement vidéo configuré et identifier la plage d'images de début et de fin qui l'étaye ; les résultats sont convertis en horodatages et liés aux images de preuve du rapport. Le score vidéo enregistré à 1,0 indique que la correspondance configurée a été conservée par le processus de revue vidéo ; il ne doit pas être lu comme une certitude calibrée de 100 %. Les réviseurs humains annotent par plage d'images et les images de revue externe sont floutées comme indiqué ci-dessus.
- 35. Systèmes d'IA et utilisations.** Le flux de classification de JeDisTuDis.com utilise une pile de modèles personnalisée, guidée par l'ontologie et affinée par fine-tuning, pour la classification des communications texte/audio et des images vidéo. Les services d'appui ne sont utilisés que pour des tâches de traitement précises, telles que les transcriptions parole-texte attribuées aux locuteurs, la détection et la mise en correspondance des empreintes vocales, la détection et la mise en correspondance des visages avec les pistes de participants identifiées par l'utilisateur, l'expurgation typée, l'aide à la structuration documentaire, les résumés de rapports et d'interactions, les workflows d'embedding et de recherche, l'aide distincte à la détection faciale pour l'examen de confidentialité, ainsi que les attributs standards de toxicité/profanité lorsque ces catégories standards sont sélectionnées. Les modèles propres à un fournisseur peuvent évoluer sans modifier la méthodologie : le matériel source est préservé ; le texte et les

transcriptions sont expurgés avant classification ou revue humaine ; les images remises aux réviseurs externes affectés ont les visages floutés ; les classifications suivent l'ontologie active ; et les exemples liés aux sources demeurent vérifiables.

- 36. Déclaration relative à l'IA générative et symbole †.** L'objectif principal de l'IA générative de synthèse du rapport est de rendre l'analytique numérique et les résultats analytiques déjà calculés plus lisibles pour le tribunal. Elle est utilisée uniquement dans des champs narratifs marqués et peut décrire des décomptes, ratios, tendances et des éléments analytiques existants liés aux sources dans le périmètre, mais elle ne détermine ni ne modifie les classifications, notes, décomptes, preuves sélectionnées ou graphiques sous-jacents. Chaque instance de contenu du rapport généré par IA générative est signalée par un obèle (†) à la fin du paragraphe généré. Ce symbole identifie uniquement la formulation explicative ; il ne nuance, ne remplace ni ne crée les chiffres, classifications, preuves ou résultats sous-jacents.
- 37. Ce pour quoi l'IA générative de synthèse du rapport n'est pas utilisée.** Le modèle de synthèse du rapport n'est pas utilisé pour classer ou noter des communications ou vidéos, calculer des décomptes ou ratios, choisir des preuves sources, générer des graphiques, modifier des fichiers sources, décider si une allégation factuelle est vraie ou fournir une conclusion juridique. Ces résultats analytiques sont achevés et enregistrés avant la synthèse narrative du rapport. Le PDF est ensuite généré directement et de manière déterministe à partir du matériel lié aux sources, des résultats calculés, des graphiques et des champs narratifs clairement marqués ; l'IA générative ne génère ni ne met elle-même en page le PDF.
- 38. Intégrité et auditabilité des sources.** Le matériel source est protégé séparément du processus analytique. Comme première étape de traitement, JeDisTuDis.com crée un instantané de preuve chiffré et scellé sur blockchain afin que toute suppression ou modification ultérieure soit détectable avant les travaux ultérieurs d'IA ou de revue humaine. Les fichiers sont chiffrés pendant le transfert et au repos, le matériel source est enregistré au moyen de hachages cryptographiques, et les éléments sources sont scellés dans une chaîne de preuve sur blockchain privée. Les téléversements, accès, partages et événements de traitement sont horodatés et associés à une activité d'utilisateur vérifiée afin de soutenir l'intégrité des sources et l'examen de la chaîne de conservation. Les rapports et documents téléversés font également l'objet de contrôles d'intégrité visuelle lors de l'importation, notamment par rasterisation des pages et analyse de structure des pixels, afin de détecter des rendus documentaires malformés, visuellement incohérents ou modifiés de manière inattendue, ce qui peut indiquer la soumission d'éléments de preuve modifiés ou générés.
- 39. Benchmarks de transcription.** Les benchmarks de transcription publiés doivent être lus séparément de la précision de classification comportementale de JeDisTuDis.com. WER signifie *word error rate* (taux d'erreur de mots) : le nombre de substitutions, suppressions et insertions de mots divisé par le nombre de mots de la transcription de référence. Un WER plus faible est meilleur. Pour une présentation comparable sous forme de taux de précision, l'équivalent simple de précision des mots est calculé comme 100 % moins le WER ; il s'agit d'une reformulation mathématique du WER publié et non d'un benchmark distinct. ElevenLabs rapporte une précision de transcription anglaise de 96,7 % pour Scribe v1, et sa documentation actuelle classe l'anglais, le français, l'allemand, l'espagnol et le portugais comme langues de transcription « Excellent » avec un WER de 5 % ou moins, soit une précision simple des mots d'au moins 95 %. OpenAI rapporte Whisper par benchmark plutôt que sous forme de taux universel unique : LibriSpeech Clean mesure une parole anglaise lue clairement, FLEURS English mesure l'anglais dans un benchmark multilingue, et la moyenne de benchmark plus large combine divers jeux de données vocaux publics. Whisper a enregistré 2,7 % de WER sur LibriSpeech Clean, soit 97,3 % de précision simple des mots ; 4,4 % de WER sur FLEURS English, soit 95,6 % ; et 12,8 % de WER moyen sur l'ensemble de benchmarks rapporté dans l'article Whisper, soit 87,2 % de précision simple des mots.
- 40. Benchmark d'ontologie personnalisée.** Le benchmark d'ontologie personnalisée de JeDisTuDis.com est maintenu par rapport à l'ontologie active du moteur dans les langues prises en charge. Le benchmark utilise des exemples sources examinés et des scénarios tous comportements pour mesurer les détections attendues, les non-détections attendues, les faux positifs et les faux négatifs. Les résultats du benchmark doivent être lus comme des mesures de performance, et non comme des constatations factuelles déterminantes concernant une communication individuelle dans ce rapport. Les lignes de groupes ci-dessous comprennent uniquement les comportements personnalisés de texte/audio et les comportements

vidéo sélectionnés pour ce rapport.

Mesure	Benchmark français actuel	Signification
Exactitude	95,75 %	Le pourcentage de cas examinés du benchmark pour lesquels le système a retourné le résultat attendu.
Faux positifs	0,64 %	Le pourcentage de tous les cas examinés du benchmark où le système a détecté un comportement qui n'était pas attendu.

Groupe comportemental	Comportements actuels du moteur
Comportement violent domestique	Admission de culpabilité; Comportement coercitif; Abus financier; Avances sexuelles non désirées; Contraindre à l'automutilation; Menace d'automutilation; Surveillance ou filature; Chantage; Frapper; Coups de pied; Tirer les cheveux; Forcer au sol; Contrainte; Étranglement; Contact sexualisé; Obstruction; Mordre, Griffes ou Cracher; Pousser, Bousculer, Faire trébucher ou Similaire; Lancer des objets; Menacer avec une arme (pistolets, couteaux, etc.); Autres altercations physiques
Comportement abusif	Critique parentale; Harcèlement corporel; Intimidation agressive; Allégations d'abus non vérifiées; Altération de la perception (gaslighting); Gestes abusifs; Jeter des liquides ou des substances; Actes destructeurs
Comportement manipulateur	Menaces légales; Induction de honte; Exploitation de la vulnérabilité; Manipulation par la culpabilité; Excuse performative; Attribution de santé mentale instrumentalisée; Acculer ou bloquer les mouvements; Interférer avec le téléphone ou l'appareil
Comportement antisocial	Refus officiel de médiation; Expression de ressentiment
Communication constructive	Co-parenting proactif; Demandes de médiation; Désescalade; Empathie; Construction de ponts; Établissement de limites saines
Réparation et réconciliation	Déclarations d'amour et de dévotion; Remords; Recherche de pardon; Accorder le pardon; Grâce / Miséricorde
Signes d'alerte	Supériorité morale; Tenue de score; Dévaluation; Intention vengeresse; Cadre jugemental; Critique (Institut Gottman); Défensivité (Institut Gottman); Mépris (Institut Gottman); Évitement (Institut Gottman)
Tactiques de pression	Ultimatum; Invalidation émotionnelle; Endettement émotionnel; Punition par retrait



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

--Fin du rapport--