

Pseudonymisation, singularité situationnelle et risque de réidentification dans les rapports sociaux

Un cadre d'analyse technique, juridique et éthique

Patrick Danto

ConfidensIA SAS

`patrick.danto@confidensia.fr`

Juin 2026

Statut du document

Ce document est une version de travail destinée à la discussion scientifique. Les résultats d'annotation sont préliminaires : l'arbitrage inter-annotateur est en cours et les scores définitifs seront consolidés après décision du troisième annotateur. Le modèle de risque présenté est en cours de formalisation et fera l'objet d'une évaluation empirique via un protocole d'annotation dédié.

Résumé

Ce travail présente deux contributions complémentaires. La première est **LaPlume**, un pipeline de pseudonymisation NER hybride (CamemBERT + règles + gazetteers) développé pour le secteur social et médico-social français (ESSMS), couvrant 101 catégories d’entités nommées incluant les entités directement et indirectement identifiantes. Les premiers résultats d’annotation humaine sur 10 rapports complexes donnent un F1 pondéré par criticité de 94,15% à 97,51% selon les annotateurs. La seconde contribution est un **modèle de risque de réidentification** $R(d, \ell)$ à quatre composantes — résidu de pseudonymisation $S(d)$, unicité situationnelle $U(d)$, accessibilité des sources auxiliaires $A(d)$, et connaissance a priori du lecteur $K(\ell)$ — ancré dans la jurisprudence récente (CJUE OC c/ Commission 2024, CE Cegedim 2026). Un protocole d’annotation original à trois profils de lecteurs est proposé pour estimer ce modèle empiriquement. Ce cadre vise à opérationnaliser le standard de preuve que la jurisprudence fait peser sur le responsable de traitement.

Contents

1	Introduction et positionnement	2
1.1	Contexte : le pipeline LaPlume	2
2	Premiers résultats d'évaluation	3
2.1	Corpus et protocole	3
2.2	Résultats par annotateur	4
2.3	Accord inter-annotateur	4
3	Limite structurelle : la singularité situationnelle	5
3.1	Ce que la pseudonymisation par entités ne peut pas traiter	5
3.2	Le coût de la généralisation	5
3.3	Position défendue	6
4	Cadre juridique	7
4.1	Hierarchie des textes de référence	7
4.2	Ce qui distingue notre cas de Cegedim	7
5	Modèle de risque de réidentification $R(d, \ell)$	8
5.1	Cadre général	8
5.2	Composante 1 : Résidu de pseudonymisation $S(d)$	8
5.3	Composante 2 : Unicité intrinsèque $U(d)$	8
5.4	Composante 3 : Accessibilité des sources auxiliaires $A(d)$	9
5.5	Composante 4 : Connaissance a priori du lecteur $K(\ell)$	10
5.6	Modèle combiné	10
6	Protocole d'évaluation empirique	11
6.1	Principe	11
6.2	Les trois profils opérationnels	11
6.3	Variables à collecter	12
7	Perspectives : vers un deuxième article	12
	Références	13

1. Introduction et positionnement

Le traitement de rapports sociaux par des systèmes d'intelligence artificielle soulève une question que le débat public tend à mal poser. La question n'est pas de savoir si un tel traitement peut s'effectuer sans risque — la réponse est non, et elle le sera toujours pour tout texte narratif riche. La question pertinente est de savoir si le risque résiduel après pseudonymisation est suffisamment réduit pour être juridiquement et éthiquement acceptable dans un contexte professionnel.

Ce document présente deux contributions imbriquées. La première est **LaPlume**, un pipeline de pseudonymisation développé spécifiquement pour les établissements et services sociaux et médico-sociaux (ESSMS) français, dont nous documentons ici les premiers résultats d'évaluation. La seconde est un **cadre d'analyse du risque de réidentification résiduelle** sur texte narratif social, qui constitue la contribution principale d'un preprint en cours de finalisation.

1.1. Contexte : le pipeline LaPlume

LaPlume est un système de pseudonymisation NER hybride combinant un modèle de langue (CamemBERT fine-tuné), des règles métier, et des gazetteers sectoriels. Il est conçu pour opérer sur des rapports sociaux — documents narratifs produits par des travailleurs sociaux dans des contextes de prise en charge variés : CHRS, CADA, ESAT, SESSAD, protection de l'enfance, CPOM, etc.

Ses caractéristiques techniques principales sont les suivantes :

- Taxonomie à **101 catégories** d'entités nommées, versionnée et documentée ;
- Architecture hybride : modèle distillé *titibongbong v2* (11L→9L, 172 Mo FP16, F1 86,8%) + règles métier + gazetteers sectoriels ;
- Couverture des entités **directement identifiantes** (personnes, dates, lieux, identifiants administratifs) *et* **indirectement identifiantes** (institutions, services, organisations, pathologies) ;
- Résolution d'adresses avec mode **equilibre**, normalisation des labels via la taxonomie ;
- **Aucun stockage** : le pipeline opère en mémoire, les données pseudonymisées sont restituées au système appelant sans persistance côté ConfidensIA.

Ce qui distingue LaPlume des pipelines existants. Les solutions comparables — pipelines NER classiques et outils comme OpenAI Privacy Filter — couvrent principalement les entités directement identifiantes. LaPlume couvre en outre les entités indirectement identifiantes, notamment les institutions, services et organisations nommés, qui permettent de situer une personne dans un réseau local même en l'absence de tout nom propre.

Finalité du traitement. Le pipeline est utilisé à des fins de formalisation de l’écrit professionnel et d’alimentation de systèmes RAG (*retrieval-augmented generation*) pour les retours d’expérience dans le secteur ESSMS. Il ne s’agit pas d’un traitement à des fins statistiques ou commerciales — distinction qui a une portée juridique directe au regard de la jurisprudence récente (voir section 4).

2. Premiers résultats d’évaluation

2.1. Corpus et protocole

L’évaluation a été conduite sur deux régimes distincts, dont nous présentons ici les résultats préliminaires.

Régime 1 — Annotation humaine indépendante

Le corpus manuel comprend **10 vrais rapports ESSMS** sélectionnés pour leur charge de difficulté : pathologies, adresses imbriquées, entités multiples dans une même phrase, contextes administratifs denses (CADA, CHRS, CJM, ESAT, SESSAD). Sur ce corpus, le pipeline détecte **2018 entités**, dont la répartition par criticité est donnée au tableau 1.

Table 1: Entités détectées par le pipeline sur les 10 rapports d’annotation humaine

Criticité	Entités
CRIT	150
ELEV	860
MOY	576
FAIB	267
INCONNUE	165
Total	2018

Le score de référence est le **F1 pondéré par criticité** :

$$F1_{\text{pondéré}} = \frac{4 \times F1_{\text{CRIT}} + 2 \times F1_{\text{ELEV}} + 1 \times F1_{\text{MOY}} + 0,5 \times F1_{\text{FAIB}}}{7,5} \quad (1)$$

Ce schéma de pondération reflète les enjeux métier : les entités CRIT (identifiants directs) et ELEV (entités indirectement identifiantes à fort risque) pèsent respectivement quatre et deux fois plus que les entités de criticité moyenne.

Régime 2 — Audit LLM sur corpus de distillation

Une campagne d’audit automatique a été conduite sur 592 rapports issus du corpus de distillation, avec un LLM auditeur indépendant du LLM générateur. Le pipeline y détecte **25 921 entités**. Les scores obtenus ($F1_{\text{pondéré}} > 99\%$) constituent un **indicateur de cohérence interne** — ils vérifient que le pipeline ne se dégrade pas sur des textes

peu ambigus. Ils ne constituent pas un indicateur de généralisation : le LLM générateur avait accès à la taxonomie, ce qui introduit une homogénéité conceptuelle favorable. L’indicateur de généralisation métier reste le F1 pondéré humain sur les 10 rapports difficiles.

2.2. Résultats par annotateur

Deux annotateurs indépendants ont audité les entités du pipeline sur les 10 rapports. Le tableau 2 présente les scores F1 globaux et pondérés, avec et sans les entités de criticité FAIB.

Table 2: Scores F1 par annotateur (scores d’audit sur référence humaine)

Annotateur	F1 global <i>avec FAIB</i>	F1 pondéré <i>avec FAIB</i>	F1 global <i>sans FAIB</i>	F1 pondéré <i>sans FAIB</i>
A	91,26%	94,15%	94,64%	95,79%
B	92,66%	95,82%	96,09%	97,51%

L’analyse par niveau de criticité (tableau 3) montre que les deux annotateurs sont très solides sur CRIT et ELEV, ce qui est l’essentiel pour la protection effective des personnes. La catégorie FAIB est structurellement plus volatile : elle regroupe des entités de contexte faible, moins discriminantes pour juger la robustesse métier du marquage.

Table 3: F1 par criticité

Criticité	Annotateur A	Annotateur B
CRIT	96,67%	98,67%
ELEV	95,58%	96,63%
MOY	92,71%	94,62%
FAIB	71,16%	72,28%

2.3. Accord inter-annotateur

Le tableau 4 présente les métriques d’accord entre les deux annotateurs, calculées par le script `compare_audit_annotations.py`.

L’accord ajusté de 0,852 (0,865 sans missing suspects) est défendable pour une tâche de cette complexité. La faible couverture d’appariement (0,312) reflète des différences de segmentation et de périmètre entre les deux annotateurs, que l’arbitrage du troisième annotateur permettra de trancher.

Table 4: Métriques d'accord inter-annotateur

Indicateur	Valeur
Paires appariées	113
Left only (A)	144
Right only (B)	105
Couverture d'appariement	0,312
Accord extraction	0,814
Accord complet	0,434
Accord inter-annotateur ajusté	0,852
Accord ajusté sans missing suspects	0,865

État du chantier

Les scores présentés sont des **scores d'audit relatifs**. L'arbitrage inter-annotateur est en cours ; les classeurs ont été générés mais les décisions du troisième annotateur ne sont pas encore consolidées. La référence *gold* sera établie après cet arbitrage et les scores définitifs pourront s'en écarter.

3. Limite structurelle : la singularité situationnelle

3.1. Ce que la pseudonymisation par entités ne peut pas traiter

LaPlume couvre de façon robuste les entités directement et indirectement identifiantes. Toutefois, un risque résiduel subsiste pour les **descriptions narratives longues et détaillées de situations individuelles**. La combinaison de détails contextuels — configuration familiale, séquence de prise en charge, pathologies co-occurentes — peut conférer une **singularité situationnelle** suffisante pour permettre une réidentification, même en l'absence de tout span identifiant détecté.

Ce risque ne relève pas de la détection d'entités mais de la **généralisation du récit**, et constitue une limite structurelle du paradigme NER appliqué aux documents narratifs sociaux.

3.2. Le coût de la généralisation

Le tableau 5 illustre les trois niveaux de traitement possibles et leur impact sur la valeur informationnelle du document.

La version intermédiaire (après LaPlume) est la seule combinaison viable : elle préserve la structure sémantique, les catégories d'entités et les relations entre éléments, tout en supprimant les identifiants. La généralisation du récit n'est pas une option viable comme standard général.

Table 5: Illustration de l’impact du niveau de traitement sur la valeur métier

Niveau	Exemple	Valeur RAG / écrit professionnel
Texte brut	“M. <i>Prénom Nom</i> , né le <i>12/03/1987</i> , suivi au <i>CHRS L’Escale</i> depuis janvier 2023 porte un trouble du spectre autistique, avec déficience intellectuelle. Il est suivi à l’Hôpital de Creil et père de deux enfants placés à l’ASE de l’Oise.”	Maximale — données personnelles en clair
Après LaPlume	“M. $\langle PER_1 \rangle$, né le $\langle DATE_2 \rangle$, suivi au $\langle ETAB_CHRS_1 \rangle$ $\langle DATE_1 \rangle$ porte un trouble du $\langle PATHOLOGY_2 \rangle$, avec $\langle PATHOLOGY_1 \rangle$. Il est suivi à l’ $\langle ORG_CH_CHU_1 \rangle$ et père de deux enfants placés à l’ASE $\langle LOC_1 \rangle$.”	Maximale — données protégées, sens intact
Après généralisation	“Une personne adulte est suivie dans une structure depuis un certain temps. Elle a des difficultés et une situation familiale complexe.”	Nulle — inutilisable

3.3. Position défendue

Trois conditions permettent de défendre que le traitement pseudonymisé est acceptable :

1. Le pipeline couvre les entités directement *et* indirectement identifiantes ;
2. Le risque de réidentification pratique est négligeable au sens du RGPD (considérant 26) ;
3. La comparaison est honnête : non avec un idéal inaccessible, mais avec les pratiques effectives du secteur, où le traitement de données non filtrées par des outils grand public est déjà courant.

L’alternative réelle pour les petites et moyennes structures n’est pas l’absence de traitement LLM, mais des modèles locaux aux capacités inférieures, avec un taux d’hallucination plus élevé et une accessibilité restreinte.

4. Cadre juridique

4.1. Hiérarchie des textes de référence

Le cadre applicable est consolidé par une séquence de décisions convergentes.

RGPD, considérant 26 (2018). Fondement : pour déterminer si une personne est identifiable, il convient de prendre en considération l'ensemble des moyens raisonnablement susceptibles d'être utilisés, en tenant compte de facteurs objectifs : coût, temps, technologies disponibles.

CJUE, Breyer c/ Bundesrepublik Deutschland, C-582/14, 19 octobre 2016. Arrêt fondateur sur la notion de “moyens raisonnables”. La CJUE pose le *test de relativité* : une même donnée peut être personnelle pour l'un et non personnelle pour l'autre selon leurs capacités respectives d'identification. La réidentification n'est pas raisonnablement possible lorsqu'elle impliquerait un effort démesuré en termes de temps, de coût et de main-d'œuvre.

CJUE, OC c/ Commission, C-479/22 P, 7 mars 2024. Précision décisive : un moyen n'est pas susceptible d'être raisonnablement mis en œuvre lorsque le risque d'identification paraît *en réalité insignifiant*, en ce que l'identification est interdite par la loi ou *irréalisable en pratique*, notamment parce qu'elle impliquerait un effort démesuré. C'est cet arrêt que le Conseil d'État cite dans l'affaire Cegedim.

CJUE, CEPD c/ CRU, C-413/23 P, 4 septembre 2025. Confirmation du test par destinataire : les données pseudonymisées ne constituent pas en toute hypothèse et pour toute personne des données à caractère personnel. L'évaluation dépend du point de vue du destinataire et de ses capacités raisonnables d'identification.

CE, GERS / Cegedim Santé, n° 498628, 13 février 2026. Décision pivot pour le droit administratif français. Le Conseil d'État confirme les sanctions CNIL et ancre le principe selon lequel la pseudonymisation ne fait pas sortir les données du champ du RGPD dès lors que le risque de réidentification n'est pas insignifiant. Points saillants pour notre propos : la qualification dépend de l'appréciation objective du risque effectif ; la charge de la démonstration pèse sur le responsable de traitement ; l'affaire portait sur des traitements à des *finalités statistiques et commerciales* sur des volumes massifs — contexte structurellement différent du nôtre.

4.2. Ce qui distingue notre cas de Cegedim

Cegedim illustre ce qu'il ne faut pas faire. Il ne condamne pas tout traitement LLM de données pseudonymisées — il fixe le standard de preuve que nous cherchons précisément à satisfaire avec le modèle présenté en section 5.

Table 6: Comparaison des contextes Cegedim et LaPlume/ESSMS

Dimension	Cegedim	LaPlume / ESSMS
Finalité	Statistique et commerciale	Formalisation professionnelle, RAG interne
Stockage	Bases de données massives persistantes	Aucun — traitement en mémoire
Volume	Millions de dossiers patients	Corpus limité par organisation
Destinataire	Tiers commerciaux	Professionnel ayant accès au dossier
Pipeline de protection	Pseudonymisation par codes	101 catégories, entités directes et indirectes

5. Modèle de risque de réidentification $R(d, \ell)$

5.1. Cadre général

Le risque de réidentification est modélisé comme une **propriété relationnelle** — une fonction du document, de son environnement informationnel externe, et du lecteur :

$$R(d, \ell) = f(S(d), U(d), A(d), K(\ell)) \quad (2)$$

Les quatre composantes sont de nature différente : S et U sont des propriétés du document ; A est une propriété de l’environnement informationnel autour de la situation décrite ; K est une propriété du lecteur. Le risque émerge de leur interaction.

5.2. Composante 1 : Résidu de pseudonymisation $S(d)$

$$S(d) = \sum_{c \in \mathcal{C}} w_c \cdot n_c(d) \quad (3)$$

où $n_c(d)$ est le nombre d’entités résiduelles de catégorie c après pseudonymisation, et w_c le poids de criticité défini par le schéma de l’équation (1). Dans le cas idéal, $S(d) \approx 0$ sur CRIT et ELEV. Ce terme capture les défauts résiduels du pipeline, distinct de l’unicité situationnelle.

5.3. Composante 2 : Unicité intrinsèque $U(d)$

Unicité situationnelle U_{sit}

Soit $\mathbf{x}(d) = (x_1, \dots, x_p)$ un vecteur de traits binaires ou ordinaux extraits du document pseudonymisé. L’unicité situationnelle est modélisée comme l’information propre de la combinaison observée :

$$U_{\text{sit}}(d) = -\log \hat{P}(\mathbf{x}(d)) \quad (4)$$

où $\hat{P}(\mathbf{x}(d))$ est estimée empiriquement sur le corpus annoté (voir note 5.3). Les traits candidats incluent notamment : pathologie rare (binaire), comorbidité multiple (ordinal), mesure de protection judiciaire (binaire), cumul de mesures (ordinal), configuration familiale atypique (binaire), parcours migratoire spécifique (binaire), séquence de structures atypique (binaire).

Unicité temporelle U_{temp}

Les séquences temporelles résiduelles constituent une dimension non encore couverte par le pipeline :

$$U_{\text{temp}}(d) = g(n_{\text{dates}}(d), n_{\text{dures}}(d), n_{\text{séquences}}(d)) \quad (5)$$

Unicité globale

$$U(d) = \alpha \cdot U_{\text{sit}}(d) + (1 - \alpha) \cdot U_{\text{temp}}(d) \quad (6)$$

avec α à estimer empiriquement sur le corpus annoté.

Note d'estimation ($\hat{P}$)

En l'absence de données populationnelles ESSMS, la distribution $\hat{P}(\mathbf{x}(d))$ est estimée empiriquement sur le corpus d'annotation. Une approche non-paramétrique (distance aux plus proches voisins dans l'espace des traits) sera comparée à l'estimateur empirique pour s'affranchir de l'hypothèse distributionnelle.

5.4. Composante 3 : Accessibilité des sources auxiliaires $A(d)$

Motivation. Deux situations peuvent présenter le même résidu $S(d)$ et la même unicité $U(d)$ et pourtant avoir des risques de réidentification très différents — parce que l'une a fait l'objet d'un article de presse, d'un jugement publié, d'un communiqué institutionnel ou d'une page Facebook accessible, et l'autre non. Cette dimension est explicitement couverte par la jurisprudence Breyer : les “moyens raisonnablement susceptibles d'être mis en œuvre” incluent les sources *légalement et pratiquement accessibles* à un tiers.

$$A(d) = \sum_{s \in \mathcal{S}} \lambda_s \cdot a_s(d) \quad (7)$$

où $a_s(d) \in \{0, 1\}$ indique la présence d'une source auxiliaire de type s , et λ_s son poids de recoupabilité. Le tableau 7 donne les types de sources envisagés.

Propriété importante. $A(d)$ est indépendant du pipeline. Un rapport parfaitement pseudonymisé ($S(d) = 0$) concernant une situation médiatisée reste exposé via $A(d)$.

Table 7: Types de sources auxiliaires et poids indicatifs

Source s	λ_s	Exemples
Jugement publié (Légifrance, presse judiciaire)	Élevé	Décision de placement, mesure pénale
Article de presse identifié	Élevé	Fait divers, signalement public
Communiqué municipal / institutionnel	Moyen	Annonce de prise en charge
Réseaux sociaux publics	Moyen	Page Facebook, post identifiable
Annuaire / registre public	Faible	Annuaire associatif

Inversement, une situation à forte unicité ($U(d)$ élevé) sans aucune trace publique a un $A(d) \approx 0$ qui réduit significativement le risque pratique.

5.5. Composante 4 : Connaissance a priori du lecteur $K(\ell)$

$$K(\ell) \in [0, 1] \quad (8)$$

Table 8: Profils de lecteurs, valeurs de $K(\ell)$ et ancrage juridique

Profil ℓ	$K(\ell)$	Ancrage juridique	Interprétation
Tiers quelconque	$K_0 \approx 0$	Considérant 26 RGD — standard de référence	Aucune connaissance du secteur
Tiers motivé	$K_1 \approx 0,3$	Lignes directrices CEPD — données sensibles	Connaissance du secteur, accès aux sources publiques
Professionnel du territoire	$K_2 \approx 0,6$	—	Connaissance du réseau local
Professionnel du service	$K_3 \approx 0,85$	—	Connaissance des situations en cours
Insider / rédacteur	$K_4 = 1$	Hors périmètre (CE Cegedim 2026)	Relève du contrôle d'accès

5.6. Modèle combiné

$$R(d, \ell) = \sigma(\beta_0 + \beta_1 S(d) + \beta_2 U(d) + \beta_3 A(d) + \beta_4 K(\ell) + \beta_5 U(d) \cdot K(\ell) + \beta_6 A(d) \cdot K(\ell)) \quad (9)$$

où σ est la fonction logistique, ce qui contraint $R \in [0, 1]$.

Lecture des termes d'interaction. Le terme $\beta_5 U(d) \cdot K(\ell)$ capture le fait que l'unicité situationnelle n'est exploitable que si le lecteur a une connaissance suffisante pour en tirer parti. Le terme $\beta_6 A(d) \cdot K(\ell)$ capture le fait que les sources auxiliaires ne deviennent

dangereuses que pour un lecteur capable de les identifier et de les croiser. Un jugement publié sur Légifrance ne présente pas de risque pratique pour un quidam sans connaissance sectorielle.

Lecture juridique. Le seuil $R(d, \ell_0) < \varepsilon$ pour le tiers quelconque ℓ_0 opérationnalise le critère “insignifiant” de la jurisprudence OC c/ Commission (2024) repris dans Cegedim (2026). Le seuil $R(d, \ell_1) < \varepsilon$ pour le tiers motivé ℓ_1 répond au standard plus exigeant des lignes directrices CEPD sur les données sensibles. Ce modèle est conçu comme un outil de démonstration du respect du standard de preuve que la jurisprudence fait peser sur le responsable de traitement.

6. Protocole d’évaluation empirique

6.1. Principe

Pour estimer les paramètres β du modèle, nous proposons un protocole d’annotation portant non sur des entités mais sur des **jugements de réidentifiabilité** par profil de lecteur, sur des paragraphes pseudonymisés issus du corpus réel.

La tâche est formulée de façon identique pour tous les profils :

“À partir de ce paragraphe pseudonymisé, pensez-vous pouvoir identifier la personne ou reconnaître la situation ?”

Échelle de réponse (variable cible $y(d, \ell)$) :

- 0 Aucune réidentification possible
- 1 Doute — des éléments suggèrent une situation connue
- 2 Reconnaissance probable
- 3 Certitude

Trois colonnes $y(d, \ell_0)$, $y(d, \ell_1)$, $y(d, \ell_3)$ permettent d’estimer les β par régression logistique.

6.2. Les trois profils opérationnels

Profil 1 — Quidam (tiers quelconque ℓ_0). Aucune connaissance du secteur social ni médico-social. Opérationnalise le standard de référence du considérant 26 RGPD.

Profil 2 — Expert du secteur (tiers motivé ℓ_1). Connaissance des pratiques, des structures, du cadre légal ESSMS ; pas de connaissance du territoire ni des situations individuelles. Opérationnalise le tiers motivé des lignes directrices CEPD. Tâche complémentaire : renseigner $A(d)$ par recherche active de sources publiques (Légifrance, presse, réseaux sociaux, annuaires). Question spécifique : “Avez-vous trouvé une source publique permettant de recouper cette situation ? Si oui, laquelle et de quel type ?”

Profil 3 — Annotateur avec corpus partiel (insider partiel ℓ_3). Dispose de 10 situations non pseudonymisées comme base de reconnaissance. Tâche : reconnaître si le paragraphe pseudonymisé correspond à l’une des situations connues. Ce profil opérationnalise la **distinguabilité intra-corpus** — le risque pratique le plus réaliste dans le secteur médico-social (un professionnel qui croise un dossier qu’il a déjà traité). Ce profil est, à notre connaissance, original dans la littérature sur l’évaluation du risque de réidentification.

6.3. Variables à collecter

Deux séries de variables sont collectées en complément du score $y(d, \ell)$:

Traits $\mathbf{x}(d)$ pour alimenter U_{sit} (tous profils) : présence de pathologie rare ou combinaison de pathologies, cumul de mesures de protection, configuration familiale atypique, parcours migratoire spécifique, séquence de structures atypique, repères temporels résiduels.

Sources auxiliaires $a_s(d)$ pour alimenter $A(d)$ (profil 2 uniquement) : jugement publié, article de presse, communiqué municipal ou institutionnel, page ou post de réseau social public, présence dans un annuaire ou registre public.

Considérations éthiques

Le profil 3 implique un accès temporaire à des situations réelles non pseudonymisées. Le protocole prévoit : (1) consentement éclairé des annotateurs ; (2) destruction sécurisée des données après annotation ; (3) justification explicite dans l’article : ce protocole *mesure* un risque existant dans les pratiques sectorielles, il ne le crée pas.

7. Perspectives : vers un deuxième article

Les résultats de l’évaluation empirique du modèle $R(d, \ell)$ feront l’objet d’un preprint indépendant dont les grandes lignes sont les suivantes.

Contribution principale. Premier travail à formaliser le risque de réidentification résiduelle sur texte narratif social français, posant le risque comme propriété relationnelle document/environnement/lecteur. La composante $A(d)$ est originale : elle opérationnalise dans un cadre quantitatif les “sources légalement accessibles” de la jurisprudence Breyer. Le protocole à trois profils de lecteurs fournit une estimation empirique du seuil “insignifiant” au sens d’OC c/ Commission (2024) et de Cegedim (2026) pour les deux profils juridiquement pertinents (ℓ_0 et ℓ_1).

Plan envisagé : (1) Introduction — pourquoi le débat est mal posé ; (2) État de l’art — pseudonymisation NER, limites, cadre RGPD ; (3) Pipeline LaPlume — architecture, résultats, comparaison alternatives ; (4) Singularité situationnelle — formalisation, illustration Avant/Après ; (5) Cadre juridique — hiérarchie des textes, calibration des profils ; (6) Modèle $R(d, \ell)$ — quatre composantes, termes d’interaction, ancrage juridique ; (7)

Protocole d’annotation ; (8) Résultats ; (9) Arbitrage pratique — position normative, comparaison avec les pratiques sectorielles ; (10) Conclusion.

Lien avec le présent article. Le preprint peut citer les résultats LaPlume sans en dépendre. Les deux contributions sont complémentaires : l’une documente le pipeline, l’autre formalise la question du risque résiduel que tout pipeline — aussi complet soit-il — ne peut entièrement éliminer.

Références

References

- [1] Règlement (UE) 2016/679 du Parlement européen et du Conseil (RGPD), 27 avril 2016. Considérant 26.
- [2] Cour de justice de l’Union européenne, *Patrick Breyer c/ Bundesrepublik Deutschland*, affaire C-582/14, 19 octobre 2016. ECLI:EU:C:2016:779.
- [3] Cour de justice de l’Union européenne, *OC c/ Commission européenne*, affaire C-479/22 P, 7 mars 2024. ECLI:EU:C:2024:215.
- [4] Cour de justice de l’Union européenne, *Contrôleur européen de la protection des données c/ Conseil de résolution unique*, affaire C-413/23 P, 4 septembre 2025.
- [5] Conseil d’État, *GERS / Cegedim Santé c/ CNIL*, n° 498628, 13 février 2026.
- [6] Commission nationale de l’informatique et des libertés (CNIL), *L’anonymisation de données personnelles*, 2020. <https://www.cnil.fr/fr/lanonymisation-de-donnees-personnelles>
- [7] Danto, P., *LaPlume v2 : modèle de pseudonymisation NER pour le secteur médico-social français*, Zenodo, 2026. DOI: 10.5281/zenodo.17689395

Disponibilité. Le pipeline LaPlume et le modèle distillé sont disponibles sur Hugging Face : <https://huggingface.co/jmdanto>. Le DOI de l’article de distillation est 10.5281/zenodo.17689395.