



Thomas LEQUEUX Consulting x Opti'

L'éthique de l'IA dans l'humanitaire

De la théorie à la pratique

Guide récapitulatif

Le règlement européen sur l'IA, le cadre SAFE AI et deux démonstrations : les données clés à garder sous la main

Webinaire du samedi 26 septembre 2026

Thomas LEQUEUX · Consultant MEAL · Concepteur d'Opti'

Comment lire ce guide

Il reprend l'essentiel du webinaire en une dizaine de pages : les chiffres, les deux grilles de lecture, les deux cas pratiques et les liens pour aller plus loin.

À jour au 21 septembre 2026. Ce guide n'est pas un avis juridique. Les dates du règlement européen ont changé deux fois en un an : vérifiez-les avant de vous y engager.

1. L'essentiel en cinq idées

- 1. L'écart est le vrai risque.** 75 % des humanitaires utilisent l'IA chaque jour ou chaque semaine, mais 23 % seulement travaillent dans une organisation qui a une politique IA. En Afrique subsaharienne, 19 %.
- 2. La loi fixe un plancher.** Le règlement européen (AI Act) classe les usages de l'IA par niveau de risque. Interdictions et formation depuis 2025, transparence depuis août 2026, haut risque fin 2027.
- 3. SAFE AI fixe le standard du secteur.** Le droit de savoir, les principes humanitaires, trois niveaux de risque, un parcours en quatre étapes, et le droit de dire non.
- 4. On classe un usage, pas un outil.** Le même outil peut relever du niveau le plus bas ou du plus élevé, selon ce que l'on fait de son résultat.
- 5. Ce qui peut être vérifié par une machine doit l'être.** Ce qui ne peut pas l'être reste une décision humaine, et s'écrit dans une fiche de transparence.

2. L'état des lieux en 2026

Enquête de la Humanitarian Leadership Academy et de Data Friendly Space, janvier 2026 : 1 729 humanitaires dans plus de 120 pays, dont près de la moitié en Afrique subsaharienne. Entre parenthèses, l'évolution par rapport à mai et juin 2025.

Indicateur	Valeur	Ce que ça dit
Ont déjà utilisé l'IA pour leur travail	95,3 % (+2,2)	L'IA est entrée partout.
L'utilisent chaque jour ou chaque semaine	75,0 % (+4,6)	L'usage devient quotidien.
Utilisent déjà des agents IA	29,8 %	Des outils qui agissent à la place de l'utilisateur.
Ont reçu une formation de leur organisation	34,0 % (+2,4)	La plupart se forment seuls.
Organisation dotée d'une politique IA formelle	23,3 % (+1,5)	Le cadre ne suit pas l'usage.
IA intégrée à l'échelle de l'organisation	9,0 % (+1,2)	L'adoption reste individuelle.

Qui a une politique IA ?	Part	
Organisations locales	13 %	Les plus fortes utilisatrices quotidiennes, les moins encadrées.
Secteur privé	22 %	
ONG internationales	31 %	
Agences des Nations unies	39 %	
Afrique subsaharienne (toutes organisations)	19 %	Avec 50 % d'usage quotidien.

La phrase à retenir

L'adoption individuelle devance largement la préparation des organisations. Et l'écart ne se referme pas : la part de politiques IA n'a gagné que 1,5 point en huit mois.

Ce qui a bougé en 2026

- **Mai** : publication de SAFE AI, premier cadre opérationnel partagé pour l'IA humanitaire.
- **1er et 6-7 juillet** : rapport préliminaire du Panel scientifique international de l'ONU sur l'IA, puis premier Dialogue mondial sur la gouvernance de l'IA, à Genève.
- **20 juillet** : lignes directrices de la Commission européenne sur l'obligation de signaler l'IA (article 50).
- **27 juillet** : entrée en vigueur de l'omnibus IA, règlement (UE) 2026/1744, qui modifie le règlement européen.
- **2 août** : application générale du règlement européen.
- **Tendance** : 362 incidents liés à l'IA documentés en 2025, contre 233 en 2024 (Stanford AI Index 2026, tous secteurs).

3. Le règlement européen sur l'IA (AI Act)

Règlement (UE) 2024/1689, en vigueur depuis le 1er août 2024. Il régit des usages, pas une technologie : plus un usage peut nuire aux personnes, plus les obligations sont lourdes. Il s'ajoute au RGPD sans le remplacer, et ne prévoit **aucune exception humanitaire**.

Votre rôle : fournisseur ou déployeur

Rôle	Définition
Fournisseur	Développe un système d'IA, ou le fait développer, et le met sur le marché sous son nom (Mistral, OpenAI, l'éditeur d'un outil).
Déployeur	Utilise un système d'IA sous son autorité, dans un cadre professionnel. C'est le rôle de votre organisation dès que vos équipes utilisent l'IA pour le travail.

Votre organisation est-elle concernée ?

- **Oui**, si elle est établie dans l'Union européenne (par exemple une ONG internationale dont le siège est en Europe), ou si le résultat de l'IA est utilisé dans l'Union.
- **Pas directement**, pour une organisation nationale sans lien avec l'Union. Mais ses bailleurs et partenaires européens y sont soumis et le répercutent dans leurs contrats. Et les lois nationales de protection des données s'appliquent déjà.

Quatre niveaux de risque

Niveau	Régime	Exemples et obligations
Risque inacceptable	Interdit	Exploiter les vulnérabilités liées à l'âge, au handicap ou à la situation sociale ou économique ; notation sociale ; manipulation ; reconnaissance des émotions au travail et à l'école ; catégorisation biométrique ; moissonnage d'images faciales ; prédiction d'infractions par profilage ; identification biométrique en temps réel à des fins répressives. Depuis le 2 février 2025.
Haut risque	Encadré	Éligibilité aux aides et services publics essentiels (par ou pour une autorité publique), triage d'urgence, migration et asile, recrutement. Gestion des risques, documentation, supervision humaine. Reporté au 2 décembre 2027.
Transparence	Signalé	Dire qu'on parle à une IA, marquer les contenus générés, signaler les hypertrucages et les textes publiés. Depuis le 2 août 2026.
Risque minimal	Libre	La grande majorité des usages courants : aucune obligation spécifique.

Deux obligations qui s'appliquent déjà à vous

La transparence (article 50)

Un texte généré par l'IA et publié pour informer le public sur une question d'intérêt général doit être signalé comme tel, sauf s'il a été relu par une personne qui en assume la responsabilité éditoriale. Les hypertrucages (image, son, vidéo imitant une personne ou un événement réel) doivent toujours être signalés.

La littératie IA (article 4), reformulée le 27 juillet 2026

Avant : garantir, dans toute la mesure du possible, un niveau suffisant de maîtrise de l'IA du personnel.

Depuis : prendre des mesures pour soutenir le développement de cette maîtrise, sans devoir garantir un niveau donné. C'est une obligation de moyens, qui pèse sur l'organisation et non sur la personne.

Concrètement : pouvoir montrer des consignes écrites, une formation adaptée aux rôles et une voie de signalement.

Le calendrier

Date	Ce qui s'applique
1er août 2024	Entrée en vigueur
2 février 2025	Pratiques interdites et littératie IA
2 août 2025	Modèles d'IA à usage général, gouvernance, sanctions
27 juillet 2026	Omnibus IA en vigueur, règlement (UE) 2026/1744
2 août 2026	Application générale, dont la transparence
2 décembre 2026	Nouvelles interdictions, fin du délai de marquage des contenus pour les systèmes existants
2 décembre 2027	Haut risque, annexe III
2 août 2028	Haut risque intégré à des produits réglementés

Sanctions maximales : 35 millions d'euros ou 7 % du chiffre d'affaires mondial pour une pratique interdite, 15 millions ou 3 % pour les autres obligations.

4. Le cadre SAFE AI

Standards and Assurance Framework for Ethical Artificial Intelligence, version 1.1, mai 2026. Publié par le CDAC Network, l'Alan Turing Institute et Humanitarian AI Advisory, avec un financement du FCDO britannique. C'est un cadre volontaire : il ne crée aucune obligation juridique et ne délivre aucune certification.

Le principe qui organise tout : le droit de savoir

Les personnes touchées par des décisions influencées par l'IA ont le droit de savoir que l'automatisation est à l'œuvre, de comprendre comment elle façonne ces décisions, et de les contester. Collectivement aussi : plusieurs organisations dans une même zone doivent être tenues au même standard.

Les principes humanitaires, appliqués à l'IA

Principe	Le risque avec l'IA	Ce que SAFE AI demande
Humanité	Un système peut nuire tout en améliorant l'efficacité.	Regarder le coût de l'erreur et la supervision humaine minimale.
Impartialité	Exclusion par la langue, les trous dans les données, la connectivité.	Vérifier que l'effet peut être détecté et corrigé.
Neutralité	Alignement, réel ou perçu, avec une partie au conflit.	Reconnaître, justifier, atténuer.
Indépendance	Dépendance à un fournisseur.	Garder l'autorité de suspendre, modifier ou retirer.

Les communautés dans la boucle : concevoir avec les personnes concernées, pas seulement pour elles. Le cadre désigne la participation de façade comme le mode d'échec dominant.

Les trois niveaux de risque

Niveau	Type d'usage	Ce qu'il faut produire
1 • Base	Consultatif, réversible, surtout interne, sans données de bénéficiaires. Une à deux journées.	Évaluation d'impact (1re partie), fiche de transparence allégée.
2 • Renforcé	Orienté la priorisation, le ciblage ou des décisions opérationnelles. Utilise des données des communautés. Une à deux semaines.	Évaluation complète, fiche complète, assurance technique.
3 • Haut risque	Affecte directement l'accès à l'aide, à la protection ou à l'information. Coût de l'erreur élevé.	Tous les outils, co-conception obligatoire.

Deux règles : un usage interne ne garantit pas le niveau 1 ; en cas de doute, on applique le niveau supérieur.

La méthodologie en quatre étapes

Étape	L'essentiel
1 • Définir le problème	L'énoncer sans parler d'IA. Se demander si l'IA est vraiment la solution. Vérifier la préparation de l'organisation et les données.
2 • Concevoir	Risques et coût de l'erreur, communautés, choix délibéré de l'architecture.

Étape	L'essentiel
3 · Développer	Acheter avec les bonnes clauses (droit d'audit, sortie, notification des changements), tester sur le terrain et dans les langues parlées, valider avec un responsable nommé.
4 · Déployer et suivre	Publier la fiche de transparence avant la mise en service, retours dès le premier jour, réévaluation, incidents, retrait si nécessaire.

Chaque étape se ferme par un point de décision : continuer, reconcevoir, suspendre ou arrêter.

Le refus responsable

Ne pas utiliser l'IA est une issue valide. On suspend, on reconçoit ou on arrête si l'une de ces conditions ne peut pas être atténuée :

- aucune voie de recours réelle pour les personnes concernées ;
- personne n'est clairement responsable, ni habilité à arrêter le système ;
- les données sont réutilisées au-delà de ce que les personnes pouvaient attendre ;
- des personnes qui ne peuvent pas refuser sont touchées de façon disproportionnée ;
- l'IA remplace un jugement humain essentiel ;
- l'adoption se fait sous la pression de l'urgence, du coût ou d'un bailleur.

Les outils du cadre

Liste de préparation de l'organisation (5 sections) · liste de préparation du cas d'usage · évaluation d'impact en deux temps · stratégies d'atténuation des risques · guide de co-conception · guide d'architecture · guide d'achat · assurance technique · fiche de transparence. Tout est gratuit, en anglais, sur cdacnetwork.org/safe-ai.

La phrase à retenir

« Si ce n'est pas consigné dans la fiche de transparence, ce n'est pas considéré comme gouverné par SAFE AI. »

5. Les deux cas pratiques

Deux cas réels, sur les mêmes données de suivi post-distribution, classés avec les deux grilles. La leçon : le même outil, deux usages, deux classements.

	Créer un questionnaire et le déployer sur Kobo	Rédiger le rapport narratif
Règlement européen	Risque minimal	Transparence : signalement si publié, sauf relecture humaine assumée
SAFE AI	Niveau 1 · fiche allégée	Niveau 2 · fiche complète, assurance technique
Données de bénéficiaires	Aucune pour générer	Oui, les réponses collectées
Ce qui ferait monter d'un niveau	Un score qui décide de l'aide	Un classement de ménages sans relecture

Qui vérifie quoi : trois acteurs

Acteur	Son rôle
Humain	Fixe le besoin, tranche aux points de décision, relit, valide, et porte la responsabilité de ce qui sort.
IA	Propose, structure, rédige, dans un cadre fixé à l'avance. Elle ne décide de rien.
Machine	Vérifie de façon déterministe : même entrée, même résultat, sans modèle de langage.

Le questionnaire : les contrôles automatiques

- chaque ajout de l'IA est validé, et chaque référence d'un saut ou d'un calcul est confrontée à la liste des variables ;
- les scores et indices sont vérifiés arithmétiquement ;
- avant l'envoi : modalités en double, références orphelines, groupes vides, listes manquantes, noms de variables invalides ;
- validation par pyxform et ODK Validate, le même contrôle que KoboToolbox au déploiement.

Côté humain : valider le plan, relire et corriger chaque section, tester le formulaire, puis le piloter sur le terrain.

Le rapport : le chiffre ne vient jamais du modèle

- les tableaux croisés sont calculés à partir des données, sans IA ;
- chaque valeur reçoit un repère unique dans un registre ;
- avant de rédiger, l'IA soumet ses lectures : on valide, corrige ou rejette, et on ajoute le contexte du terrain ;
- chaque chiffre écrit est confronté au registre, et c'est toujours la valeur du registre qui est publiée.

6. Classer un usage en cinq questions

Une grille à appliquer à n'importe quel usage de l'IA, un outil acheté comme un compte personnel.

	Question	Si la réponse est oui
1	L'usage fait-il partie des pratiques interdites de l'AI Act ?	On ne le fait pas.
2	Décide-t-il de l'accès de personnes à l'aide, à la protection ou à l'information ?	SAFE AI niveau 3, possible haut risque au sens de la loi.
3	Orienté-t-il un ciblage, une priorisation ou une décision opérationnelle ?	SAFE AI niveau 2 au minimum.
4	Utilise-t-il des données de bénéficiaires ?	Pas de niveau 1 ; vérifier aussi le consentement et votre AIPD.
5	Le résultat est-il publié ou montré à des personnes ?	Transparence : signaler l'IA, ou relire et assumer la responsabilité éditoriale.

Et toujours : en cas de doute, le niveau supérieur.

7. Trois gestes pour lundi

1. **Faire l'inventaire.** Quels outils d'IA votre équipe utilise-t-elle vraiment, comptes personnels compris ? Sans jugement : le but est de rendre l'usage visible.
2. **Classer trois usages.** Les trois plus fréquents, avec la grille ci-dessus.
3. **Remplir la liste de préparation SAFE AI.** Cinq sections, une réunion de direction. Chaque « non » devient une ligne de votre plan d'action.

8. Pour aller plus loin

- La formation SAFE AI, offerte : myoptibot.com/app/?formation=safe_ai

Neuf étapes, une heure environ, dans Opti' Academy. Gratuite sur tous les forfaits, y compris le compte gratuit.

- L'AI Act expliqué aux organisations humanitaires : myoptibot.com/ia-act
- Le cadre SAFE AI, expliqué : myoptibot.com/safe-ai
- Les textes et outils officiels SAFE AI : cdacnetwork.org/safe-ai
- L'enquête HLA et Data Friendly Space : humanitarianleadershipacademy.org
- Opti' a un an : 14 jours d'accès complet, à activer avant le 8 octobre : myoptibot.com/anniversaire

Contact

Thomas LEQUEUX · Consultant MEAL · consulting@lequeux-thomas.com · linkedin.com/in/thomas-lequeux · lequeux-thomas.com

Sources

- Humanitarian Leadership Academy et Data Friendly Space, How are humanitarians using artificial intelligence? January 2026 pulse survey insights, note de synthèse, mars 2026.
- Règlement (UE) 2024/1689 sur l'intelligence artificielle ; règlement (UE) 2026/1744 (omnibus IA), JOUE du 24 juillet 2026 ; Commission européenne, lignes directrices sur l'article 50, 20 juillet 2026.
- H. McElhinney, A. Mazumder, M. Tjalve, S. Madigan, S. Spencer, SAFE AI: A Governance Framework for Humanitarians using AI, et SAFE AI Tools and Guidance, version 1.1, mai 2026, CDAC Network.
- ONU, Panel scientifique international indépendant sur l'IA, rapport préliminaire, 1er juillet 2026 ; Stanford HAI, AI Index 2026, cité par la Humanitarian Leadership Academy, 14 avril 2026.

Ce guide est la propriété de Thomas LEQUEUX Consulting. Utilisation, reproduction ou diffusion autorisées pour un usage non commercial, à condition de citer la source.